

ÉCOLE DE TECHNOLOGIE SUPÉRIEURE
UNIVERSITÉ DU QUÉBEC

THÈSE PRÉSENTÉE À
L'ÉCOLE DE TECHNOLOGIE SUPÉRIEURE

COMME EXIGENCE PARTIELLE
À L'OBTENTION DU
DOCTORAT EN GÉNIE
Ph. D.

PAR
Chafika TATA

CONCEPTION D'UN MODÈLE NOVATEUR AMÉLIORANT LA PERFORMANCE
DANS LES RÉSEAUX DE LA SÉCURITÉ PUBLIQUE SUR LTE HÉTÉROGÈNES

MONTREAL, LE 18 SEPTEMBRE 2014



Chafika Tata, 2014



Cette licence Creative Commons signifie qu'il est permis de diffuser, d'imprimer ou de sauvegarder sur un autre support une partie ou la totalité de cette œuvre à condition de mentionner l'auteur, que ces utilisations soient faites à des fins non commerciales et que le contenu de l'œuvre n'ait pas été modifié.

PRÉSENTATION DU JURY

CETTE THÈSE A ÉTÉ ÉVALUÉE

PAR UN JURY COMPOSÉ DE :

M. Michel Kadoch, directeur de thèse
Département de génie électrique à l'École de technologie supérieure

M. Alain April, président du jury
Département de génie logiciel et des TI à l'École de technologie supérieure

M. François Gagnon, membre du jury
Département de génie électrique à l'École de technologie supérieure

M. Roch Glitho, examinateur externe
Professeur associé, Institut d'Ingénierie des Systèmes d'Information de l'Université
Concordia

ELLE A FAIT L'OBJET D'UNE SOUTENANCE DEVANT JURY ET PUBLIC

LE 19 AOÛT 2014

À L'ÉCOLE DE TECHNOLOGIE SUPÉRIEURE

REMERCIEMENTS

Mes premiers remerciements iront à M. Alain April, à M. François Gagnon et à M. Roch Glitho, pour m'avoir fait l'honneur de participer à ce jury de thèse.

Je souhaite exprimer ma plus profonde gratitude à mon directeur de recherche, M. Michel Kadoch, pour m'avoir conseillée et orientée tout au long de la thèse avec patience et disponibilité, et pour la confiance qu'il m'a accordée.

Mes remerciements vont particulièrement à mon mari Hamid qui m'a constamment encouragée et soutenue pendant toute la durée de ma thèse et plus particulièrement durant les derniers mois de rédaction qui n'ont pas toujours été des plus agréables. C'est également le moment d'exprimer ma gratitude à mes enfants Younes et Zakaria pour la part certaine qu'ils ont jouée dans la réalisation de cette thèse.

Je tiens à remercier vivement Djedjiga Benzid pour les échanges fructueux qu'on a eu pendant toutes ces années. Je n'oublierai pas Nada Hakam et Nassima Fellag pour leur encouragement et soutien moral. À toutes les trois, je vous souhaite beaucoup de réussite.

Je voudrais également adresser mes remerciements à toute ma famille pour m'avoir toujours encouragée et s'être inquiétée du bon déroulement de mes études.

Que toutes ces personnes et celles qui ne sont pas nommément citées mais qui ont également de près ou de loin pris part à ce travail, reçoivent ici l'assurance de ma gratitude.

CONCEPTION D'UN MODÈLE NOVATEUR AMÉLIORANT LA PERFORMANCE DANS LES RÉSEAUX DE LA SÉCURITÉ PUBLIQUE SUR LTE HÉTÉROGÈNES

Chafika TATA

RÉSUMÉ

Durant les situations d'urgences, la disponibilité des moyens de télécommunications est cruciale et indispensable pour les usagers des réseaux de la Sécurité Publique (PSN). En revanche, durant de tels moments, le besoin en échange d'information croît d'une façon spectaculaire. Par conséquent, l'accès au médium radio devient congestionné très rapidement. Malheureusement, durant ces moments, les ressources dédiées aux réseaux (PSN) ne semblent pas être suffisantes pour satisfaire toutes les requêtes d'établissement des nouveaux bearers. Les réseaux LTE viennent donc contribuer à la résolution de cette problématique, en offrant l'accès à la Radio Commerciale Partagée pour le réseau PSN, avec une certaine priorisation, afin d'améliorer les communications PSN lors des situations d'urgences. Néanmoins, cet accès ne doit pas accaparer toutes les ressources du réseau commercial.

De plus, la technologie LTE permet l'utilisation des communications Device-to-Device qui consiste à échanger l'information directement entre les équipements sans avoir à passer par l'eNodeB. Les communications D2D doivent donc être exploitées pour contourner les problèmes de congestion, surtout lors des désastres.

Par ailleurs, l'amélioration de la performance des réseaux PSN ne se limite pas dans la gestion efficace des ressources radio. L'allocation des ressources de bande passante au niveau du réseau Backhaul et du réseau cœur LTE doit aussi être améliorée.

Dans cette thèse un nouveau modèle novateur a été conçu pour l'amélioration de la performance dans les réseaux de la sécurité publique sur les réseaux LTE hétérogènes. Ce modèle qui compte dix solutions, intervient sur les trois réseaux composant le réseau LTE, à savoir le réseau d'accès, le réseau Backhaul et le réseau cœur LTE.

Nos différentes solutions ont toutes été validées par simulations, et ont toutes apporté une amélioration par rapport à aux approches classiques ou par rapport à d'autres approches existantes dans la littérature.

Mots clés : LTE HetNet, WMN, Sécurité Publique, D2D, Codage Réseau, Multipath, Prémption, QoS, Ressources Radio, Bande Passante, Sécurité, Réseau Papillon.

INNOVATIVE CONCEPTION OF A MODEL ENHANCING PERFORMANCE IN PUBLIC SAFETY OVER LTE HETEROGENEOUS

Chafika TATA

ABSTRACT

During emergencies, the availability of telecommunications is crucial and essential for network users of Public Safety Networks (PSN). Therefore, during such times, the need of information exchange increases dramatically. Therefore, access to the radio medium becomes congested quickly. Unfortunately, during these times, the resources dedicated for PSN does not seem to be sufficient to meet all new bearers' requests. LTE networks are therefore contributed to address this issue by providing access to Shared Commercial Radio for PSN. This access is done with some prioritization to improve PSN communications during emergencies. However, this access should not monopolize all the commercial network resources.

In addition, the LTE technologie allows the use of communications Device-to-Device which is the exchange of information directly between devices without having to go through the eNodeB. Hence, the D2D communications should be used to address the congestion issue, especially during disasters.

Moreover, improving the performance of PSN networks is not limited in the reliable management of radio resources. The allocation of bandwidth resources at the Backhaul network and the core network of LTE should also be improved.

In this thesis a new innovative model has been designed to improve the performance in Public Safety networks over LTE heterogeneous networks. This model has ten solutions. It operates in the three parts of the LTE network, namely the access network, the backhaul network and the core network of LTE.

Our various solutions have all been validated by simulations, and all have made an improvement compared to classical approaches or compared to other existing approaches in the literature.

Keywords: LTE HetNet, WMN, Public Safety, D2D, Network Coding, Multipath, Preemption, QoS, Radio Resources, Bandwidth, Security, Butterfly Networks.

TABLE DES MATIÈRES

	Page
INTRODUCTION	1
CHAPITRE 1 Généralité sur les réseaux LTE HetNets	11
1.1 Introduction.....	11
1.2 Architecture du réseau LTE	11
1.2.1 Le réseau d'accès	12
1.2.2 Le réseau Backhaul	12
1.2.3 Le réseau cœur : EPC.....	13
1.2.3.1 MME (Mobility Managment Entity)	14
1.2.3.2 S-GW (Serving Gateway)	14
1.2.3.3 P-GW (Packet Data Network Gateway)	14
1.2.3.4 PCRF (Policy and Charging Rules Function).....	15
1.2.3.5 HSS (Home Subscriber Server)	15
1.2.4 Les interfaces de connexion.....	16
1.3 Modèle en couches du réseau LTE	17
1.3.1 Couche Physique.....	17
1.3.2 Couche liaison de données.....	18
1.3.2.1 Sous couche PDCP	18
1.3.2.2 Couche RLC.....	19
1.3.2.3 Sous couche MAC	19
1.3.3 Couche RRC	19
1.4 Plan de contrôle et plan usager	20
1.5 Transmission des données dans LTE	20
1.5.1 Service data Flow.....	20
1.5.2 Les Bearers EPS.....	21
1.5.3 Traffic Flow Template	22
1.6 La QoS dans LTE.....	23
1.6.1 QoS Class Identifier	23
1.6.2 Allocation and Retention Priority	24
1.6.3 Les paramètres du débit dans LTE.....	25
1.7 Conclusion	275
CHAPITRE 2 Généralités sur le codage réseau	27
2.1 Introduction.....	27
2.2 Principe de fonctionnement du Network Coding.....	28
2.3 Types de codage réseau	30
2.3.1 Codage réseau linéaire	30
2.3.2 Codage réseau linéaire aléatoire	32
2.3.3 Codage réseau partiel.....	32
2.4 Conclusion	32

CHAPITRE 3	Revue de la littérature	35
3.1	Introduction.....	35
3.2	Les réseaux LTE pour la sécurité publique.....	36
3.3	Gestion des ressources radio partagées dans les réseaux d'accès LTE	37
3.4	Gestion des ressources de bande passante dans les réseaux Backhaul LTE.....	42
3.5	Communications Device to Device dans LTE Hétérogène	43
3.5.1	Intégration du D2D dans l'architecture LTE	44
3.5.2	D2D pour le processus d'offloading.....	44
3.5.3	D2D pour l'amélioration de la QoS et la consommation d'énergie.....	45
3.6	Le codage réseau dans les réseaux sans fil	49
3.7	Sécurité dans les réseaux locaux sans fil	53
3.8	Conclusion	54
CHAPITRE 4	Gestion des ressources dans le réseau cœur et dans le Backhaul LTE	57
4.1	Introduction.....	57
4.2	Modèles d'allocation de bande passante avec contraintes.....	58
4.2.1	Maximum Allocation Model.....	59
4.2.2	Russian Dolls bandwidth constraints Model.....	60
4.3	CAM : Courteous bandwidth Allocation constraints Model	63
4.3.1	Le modèle mathématique du CAM.....	63
4.3.1.1	Condition d'application du CAM	63
4.3.1.2	Modèle de gestion de files d'attente du modèle CAM	64
4.3.2	Description de l'algorithme CAM	66
4.3.3	Application du CAM sur LTE	67
4.3.4	Simulations et résultats	69
4.3.4.1	Simulation d'un trafic FTP plus dense que le trafic de la voix..	69
4.3.4.2	Simulation d'un trafic FTP moins dense que le trafic de la voix	76
4.4	Conclusion	76
CHAPITRE 5	Gestion des ressources dans le réseau cœur et dans le Backhaul LTE	79
5.1	Introduction.....	79
5.2	Accès à la RAN commerciale partagée pour les premiers répondants LTE.....	80
5.3	Classification des bearers et calcul de priorité.....	83
5.4	Modèles d'allocation de bandes de fréquences avec contraintes.....	87
5.4.1	Modèle de base: L'approche classique	87
5.4.2	Courteous Constraints Allocation Model for Frequencies.....	91
5.4.3	G-CAMF	92
5.4.4	Radio Usage Situation based Courteous Constraints Allocation Model for Frequencies	95
5.5	Gestion efficace des ressources radio dans le réseau d'accès LTE HetNet.....	98
5.5.1	CPA: Accès avec priorité et courtoisie à la radio fréquence commerciale	99
5.5.1.1	Processus d'accès à la Radio de fréquence	99
5.5.1.2	Description de l'algorithme CPA.....	100
5.5.1.3	Simulations et résultats	105

5.5.2	CPAwO : Algorithme d'allocation de bandes de fréquences partagées avec offloading.....	110
5.5.2.1	Description.....	110
5.5.2.2	Data Offloading to Small Cells Algorithm.....	117
5.5.2.3	Simulations et résultats	119
5.6	Conclusion	130
CHAPITRE 6	Routage efficace et sécurité pour améliorer les transmissions D2D dans le réseau de sécurité publique sur LTE HetNets.....	133
6.1	Introduction.....	133
6.2	Routage efficace et sécurité pour les communications D2D	134
6.3	Algorithmes de signalisation.....	137
6.3.1	RBC.....	138
6.3.1.1	Description de RBC	138
6.3.1.2	Simulations et résultats	144
6.3.2	LBS-AOMDV.....	150
6.3.2.1	Best Child.....	151
6.3.2.2	LBS-AOMDV Description	154
6.3.2.3	Simulations et résultats	160
6.4	Sécurisation des communications D2D pour les réseaux PSN.....	168
6.4.1	L'algorithme G-SNCDS	173
6.4.1.1	Le mécanisme Data Splitting et l'opération de codage de G-SNCDS.....	175
6.4.1.2	Exemple de codage réseau basé sur le mécanisme DS de G-SNCDS.....	179
6.4.1.3	Le processus de décodage et le rassemblement des données de G-SNCDS.....	183
6.4.1.4	Exemple pour le décodage de l'information avec G-SNCDS ..	186
6.4.2	G-SNCDS pour éviter les attaque de confidentialité	186
6.4.3	G-SNCDS for data integrity attack avoidance.....	188
6.4.4	G-SNCDS for data availability attack avoidance	189
6.4.5	Simulations et résultats	192
6.5	Conclusion	200
CONCLUSION		203
RECOMMANDATIONS		209
PUBLICATIONS.....		211
RÉFÉRENCES		213

LISTE DES TABLEAUX

	Page
Tableau 4.1 Paramètres de la simulation	70
Tableau 4.2 Résultats numériques	71
Tableau 5.1 Définition des variables utilisées pour les modèles de contraintes.....	88

LISTE DES FIGURES

	Page
Figure 1.1	Les trois réseaux de la technologie LTE.....12
Figure 1.2	Réseau Backhaul LTE.....13
Figure 1.3	Le réseau EPS15
Figure 1.4	Les interfaces du réseau LTE.....17
Figure 1.5	Les couches du réseau LTE18
Figure 1.6	Le plan de contrôle et le plan usager20
Figure 1.7	Le Single Data Flot21
Figure 1.8	Les bearers EPS22
Figure 1.9	Filtrage des trafics dans LTE23
Figure 1.10	QoS et différenciation de services dans LTE24
Figure 2.1	Bénéfice du codage réseau : Multicast.....28
Figure 2.2	Bénéfice du codage réseau : Réduire le nombre de transmissions29
Figure 2.3	Exemple d'un codage réseau linéaire31
Figure 4.1	Bandwidth allocation in MAM59
Figure 4.2	Bandwidth allocation in RDM61
Figure 4.3	Exemple d'application pour RDM.....62
Figure 4.4	Allocation de la bande passante avec CAM63
Figure 4.5	Application du modèle CAM pour LTE68
Figure 4.6	Modèle de file d'attente adaptée par CAM.....65
Figure 4.7	Délai de la voix71
Figure 4.8	Délais de FTP.....72
Figure 4.9	Longueur moyenne de la file d'attente voix73

Figure 4.10	Longueur moyenne de la file d'attente FTP	74
Figure 4.11	Perte de parquets de type FTP	75
Figure 5.1	Nouvelle solution pour une meilleure gestion des ressources radio LTE..	82
Figure 5.2	Courteous Allocation Model for frequencies.....	90
Figure 5.3	Generalized CAMF	93
Figure 5.4	Radio Usage Situation based CAMF	96
Figure 5.5	Resources Radio Management.....	100
Figure 5.6	Nombre des bearers bloqués	107
Figure 5.7	Nombre des bearers CN actifs	108
Figure 5.8	Nombre des bearers PS actifs dans le réseau	109
Figure 5.9	Preemption of Commercial bearers	109
Figure 5.10	Modèle du réseau LTE HetNet	120
Figure 5.11	Nombre des bearers PS emergency actifs dans le réseau (CPAwO)	121
Figure 5.12	Nombre des bearers PS Non-Emergency actifs (CPAwO).....	122
Figure 5.13	Nombre des bearers CN Emergency actifs dans le réseau (CPAwO)	123
Figure 5.14	Nombre des bearers CN Non-Emergency actifs (CPAwO).....	124
Figure 5.15	Nombre des bearers CN Non-Emergency bloqués (CPAwO).....	125
Figure 5.16	Nombre des bearers PS Emergency bloqués (CPAwO)	126
Figure 5.17	Nombre des bearers PS Non-Emergency bloqués (CPAwO)	127
Figure 5.18	Nombre de bearers CN Non-Emergency interrompus (CPAwO).....	128
Figure 6.1	Amélioration du routage et de la sécurité pour les réseaux D2D	134
Figure 6.2	Effet papillon dans un WMN sans fil	139
Figure 6.3	Étapes de construction d'un réseau papillon dans un WMN	140
Figure 6.4	Butterfly effects performed with RBC algorithm: Exemple 1	145

Figure 6.5	Butterfly effects performed with RBC algorithm: Exemple 2.....	145
Figure 6.6	Primary and Backup Butterfly effects constructed with RBC algorithm	146
Figure 6.7	Butterfly based- load balancing in the WMN	148
Figure 6.8	Backup après Butterfly failure	149
Figure 6.9	Best Child selection process	153
Figure 6.10	Bande passante disponible pour le plus court chemin	161
Figure 6.11	Bande passante disponible pour le plus court chemin	163
Figure 6.12	Nombre de sauts dans le plus court chemin.....	164
Figure 6.13	Nombre de sauts dans le plus court chemin.....	165
Figure 6.14	Nombre des requêtes RREQ transmises dans le réseau.....	166
Figure 6.15	Nombre des requêtes RREQ transmises dans le réseau.....	167
Figure 6.16	Codage réseau avec attaque de sécurité	170
Figure 6.17	Le mécanisme Data Splitting.....	176
Figure 6.18	Transmission de paquets avec codage réseau et Data Splitting.....	177
Figure 6.19	Exemple de Data Splitting	180
Figure 6.20	Exemple de l'utilisation du DS pour le codage réseau	181
Figure 6.21	Attaques interne et externe de confidentialité dans le WMN	187
Figure 6.22	Attaque interne d'intégrité dans un réseau WMN	188
Figure 6.23	Attaque de disponibilité interne dans un réseau WMN	190
Figure 6.24	Construction d'un effet papillon par le système RBC: Attaque interne ..	192
Figure 6.25	Nombre de bits interceptés: Attaque interne de confidentialité.....	194
Figure 6.26	Construction d'un effet papillon par le système RBC: Attaque externe..	195
Figure 6.27	Nombre de bits interceptés: Attaque externe de confidentialité	195
Figure 6.28	Attaque d'intégrité dans un réseau WMN	196

Figure 6.29	Attaque de disponibilité dans un réseau WMN	198
-------------	---	-----

LISTE DES ALGORITHMES

	Page
Algorithme 5.1	CPAwO.....112
Algorithme 5.2	ResourcesAllocation.....113
Algorithme 5.3	Fonction Courtoisie.....114
Algorithme 5.4	Fonction PreEmption.....116
Algorithme 5.5	Algorithme HOw.....118
Algorithme 6.1	Algorithme LBS-AOMDV.....155
Algorithme 6.2	Fonction Multipath.....157
Algorithme 6.3	Fonction BestChild.....159

LISTE DES ABRÉVIATIONS, SIGLES ET ACRONYMES

2G	Deuxième Génération
3G	Troisième Génération
3GPP	3rd Generation Partnership Project
AAA	Authentication, Authorization, Accounting/Auditing
AES	Advanced Encryption Standard
AP	Access Point
ARP	Allocation and Retention Priority
ARQ	Automatic Repeat reQuest
BC	Bandwidth Constraints
BE	Best Effort
CN	Commercial User
CNR	Carrier-to-Noise Ratio
CoA	Care-Of-Address
CT	Class of Traffic
D2D	Device to Device
DHCP	Dynamic Host Configuration Protocol
DiffServ	Differentiated Services
DS-TE	DiffSev Traffic Engineering
eNodeB	evolved Node B
ePDG	Evolved Packet Data Gateway
E-UTRAN	Evolved Universal Terrestrial Radio Access Network
EPC	Evolved Packet Core
EPS	Evolved Packet System
EXP	experimental bits
FTP	File Transfer Protocol
GBR	Garanteed Bit Rate

GSM	Global System for Mobile Communications
HetNet	Heterogeneous Network
HSS	Home Subscriber Server
HTTP	HyperText Transfer Protocol
IMS	IP Multimedia Subsystem
IP	Internet Protocol
LAN	Local Area Network
LSP	Label Switched Path
LTE	Long Term Evolution
LTE-A	Long Term Evolution Advanced
MAM	Maximum Allocation Model
Mbps	megabits per second
MBR	Maximum Bit Rate
MD5	Message Digest 5
MME	Mobility Management Entity
MN	Mobile Node
MPLS	Multiprotocol Label Switching
MPLS-TP	MPLS Transport Profile
NAS	Non-Access Stratum
NC	Network Coding
OFDM	Orthogonal Frequency Division Multiplexing
OFDMA	Orthogonal Frequency Division Multiple Access
PDCCP	Packet Data Convergence. Protocol
PDU	Packet Data Unit
P-GW	Packet Data Network Gateway
PQ	Priority Queueing

PS	Public Safety
PSN	Public Safety Network
RB	Resources Bloc
RDM	Russian Dolls Model
RLC	Radio Link Control
RRC	Radio Resource Control
RR	Round Robin
QCI	QoS Class Identifier
QoS	Quality of Service
RAN	Radio Access Network
RCRF	Policy and Charging Rules Function
RSA	Ronald Rivest, Adi Shamir Leonard Adleman
RSVP	Resource Reservation Protocol
SDF	Service Data Flow
SDU	Service Data Unit
SGSN	Serving GPRS Support Node
S-GW	Serving Gateway
TB	Transmission Bloc
TCP/IP	Transmission Control Protocol /the Internet Protocol
TE	Traffic Engineering
TFT	Traffic Flow Template
UE	User Equipement
UMTS	Universal Mobile Telecommunications System
VaNet	Vehicular Ad-hoc NETwork
VPN	Virtual Private Network
WiMAX	Worldwide Interoperability for Microwave Access

WLAN	Wireless Local Area Network
WMN	Wireless Mesh Network
WSN	wireless sensor network
xPON	X version of Passive Optical Network

INTRODUCTION

La gestion des désastres représente certainement une importante problématique autant pour les premiers répondants que pour les chercheurs en technologie d'information et en télécommunication. Durant les moments de crises, toute information est inhérente pour sauver des vies humaines. C'est pour cette raison que les réseaux de la sécurité publique (PSN) sont déployés. En Amérique du nord, les bandes de fréquences 700 Mhz sont dédiées aux réseaux PSN (Wang, Sicker et al. 2013).

Durant les situations d'urgences, telles que les désastres naturels, les accidents de la route ou les émeutes sportives (Hallahan and Peha 2010), la disponibilité des moyens de télécommunication est cruciale et indispensable pour les usagers des réseaux PSN. En revanche, durant de tels moments, le besoin en échange d'information croît d'une façon spectaculaire, que ce soit pour les premiers intervenants (service de police, pompiers, ambulanciers... etc.) ou pour les simples citoyens. Par conséquent, l'accès au médium radio devient congestionné très rapidement. Plusieurs appels seront mis en attente et d'autres bloqués. Toutefois, les appels d'urgence doivent être acheminés, car ils sont parfois indispensables pour la gestion de la crise, certains peuvent même sauver des vies. Malheureusement, durant les catastrophes, les ressources dédiées au PSN ne semblent pas être suffisantes pour satisfaire toutes les requêtes d'établissement de bearers de type Sécurité Publique (PS).

Les réseaux Long Term Evolution (LTE) (Astely, Dahlman et al. 2013) viennent contribuer à la résolution de cette problématique, en offrant l'accès à la Radio Commerciale Partagée (RCP) pour le réseau PSN, avec une certaine priorisation (Blom, de Bruin et al. 2008, Hallahan and Peha 2010, Borkar, Roberson et al. 2011, Chadchan and Akki 2011, Hallahan and Peha 2013), afin d'améliorer les communications PSN lors des situations d'urgence.

Problématique de recherche

L'accès à la radio commerciale partagée peut être retardé durant les moments de congestion dans les réseaux LTE, autant pour les usagers du Réseau Commercial (CN) que pour les premiers répondants. La gestion d'allocation des ressources radio dédiées et celles partagées entre le réseau PSN et le réseau CN est assurée par le mécanisme Allocation and Retention Priority (ARP) (Simic 2012) via un mécanisme de contrôle d'admission de bearers (voir la définition du bearer dans la section 1.5.2 du chapitre 1). Quand les ressources radios sont disponibles, le protocole ARP permet l'établissement d'un nouveau bearer comme il assure la modification d'un bearer déjà actif dans le réseau afin de répondre aux exigences de la Qualité de Service (QoS) des flux du trafic transportés via ce bearer. Dans le cas où les ressources sont limitées, et en appliquant le mécanisme ARP, un ou plusieurs bearers actifs dans le réseau ayant une faible priorité et étant vulnérables à la préemption seront interrompus et leurs ressources seront allouées au nouveau bearer de plus haute priorité. Autrement, la requête d'établissement ou de modification du bearer sera rejetée. On parlera alors d'un bearer bloqué. Par conséquent, suivant une telle politique de gestion de ressources radio, le nombre de bearers bloqués de basses priorités va augmenter quand les ressources sont limitées, et les trafics de faibles priorités vont voir leur niveau de la QoS se détériorer. Il est à noter que l'ARP assure un accès privilégié aux usagers PS au sein de la bande de fréquence commerciale partagée. En effet, quand les ressources sont limitées, les clients commerciaux sont ceux qui vont souffrir le plus des problèmes d'accès au réseau. L'article (Borkar, Roberson et al. 2011) propose une solution pour un accès PSN prioritaire aux ressources radio commerciales partagées. Cependant, la problématique de la détérioration de la QoS du réseau CN s'expose toujours.

Par conséquent, en tenant compte de l'importance des communications PSN lors des situations de crises, et en considérant l'exigence des usagers commerciaux en terme de QoS, la solution à concevoir doit assurer un accès prioritaire aux premiers répondants tout en garantissant un bon niveau de QoS aux clients CN.

Toutefois, les ressources radio au niveau d'une même macro cellule LTE étant restreintes, le risque de congestion ne pourra être éliminé. Une situation de blocage total pourra survenir à tout moment, d'autant plus lors des désastres. On entend par un blocage total, le rejet de toute requête formulée pour établir un nouveau bearer ou modifier un bearer actif dans le réseau LTE, que ce soit pour le profit du PSN ou bien du CN. Une telle situation peut avoir lieu dès que toutes les ressources radio de la macro cellule LTE seront allouées.

Certainement, sauver des vies requiert une disponibilité continue des moyens de communication, alors que malheureusement, une pénurie en ressources radio peut bloquer le processus de gestion du désastre et amplifier le bilan de la catastrophe. Sans oublier que durant une crise, les exigences en communication relative au réseau CN augmentent à leur tour. Cet état de choses est relié au fort besoin des citoyens de s'informer de l'avancement de la résolution de la crise et d'avoir des nouvelles de leurs familles et amis impliqués dans le désastre. De ce fait, le nouveau défi à relever sera donc de trouver des ressources radio additionnelles quand celles fournies par la station de base LTE (eNodeB) sont toutes allouées.

Un autre point aussi important à soulever est le fait que dans le cas où les ressources additionnelles sont fournies par un réseau sans fil, suite à un offloading des bearers des macros cellules vers les petites cellules LTE, il devient improbable de fournir le même niveau de sécurité que celui offert par le réseau LTE. Une solution de sécurité doit donc être développée pour sécuriser les échanges de données Device-to-Device, entre les différents usagers du PSN et CN au niveau des réseaux locaux sans fil (WLAN) opérant dans les petites cellules. La contrainte reliée à la conception d'un tel modèle est de garder un bon niveau de QoS pour les classes de trafics PS et CN. Autrement dit, l'approche ne doit pas ajouter un trafic de contrôle supplémentaire dans le réseau.

D'autre part, étant donné que les deux réseaux PSN et CN partagent une même portion du spectre radio, il devient difficile de gérer un accès équitable à cette partie de ressources communes. L'un ou l'autre des deux réseaux pourra accaparer toutes les ressources pendant

de courtes ou de longues durées. Il serait alors pertinent de concevoir un modèle d'allocation de ressources radio avec contraintes. Ces contraintes consistent à considérer les besoins de chaque réseau et à développer un modèle d'allocation de ressources radios aux différents bearers des deux réseaux PSN et CN, suivant des conditions prédéfinies et relativement aux besoins et à la nature de chaque réseau.

Par ailleurs, l'amélioration de la performance des réseaux PSN ne se limite pas dans la gestion efficace des ressources radio. L'allocation des ressources de bande passante au niveau du réseau Backhaul et du réseau cœur LTE doit aussi être améliorée. Les deux réseaux appliquent le principe d'ordonnancement pour la gestion de la transmission de différents types de trafics. Pour la même raison que pour les ressources radios, la gestion des ressources de la bande passantes doit respecter les contraintes relatives à chaque type de trafic afin de garantir un bon niveau de QoS pour chacun de ces trafics.

Ce travail propose donc la conception d'un modèle novateur pour l'amélioration de la performance des communications PSN sur les réseaux LTE Hétérogènes (HetNets). Ce modèle se compose de dix solutions qui peuvent être appliquées conjointement ou séparément. En effet, chacune d'entre elle apporte sa valeur ajoutée pour la performance de notre modèle.

Objectifs

L'objectif principal de notre travail est de développer un nouveau modèle pour l'amélioration de la performance des réseaux de la Sécurité Publique opérant sur les réseaux LTE HetNets, lors des situations de désastres. Cela se traduit par l'amélioration de la gestion des ressources radio et de bande passante au niveau du réseau d'accès, du réseau Backhaul et du réseau cœur LTE. Les objectifs secondaires de notre travail sont les suivants :

- gérer efficacement l'allocation des ressources radio commerciales partagées dans les macros cellules du réseau LTE pour assurer l'accès avec priorité aux usagers PS, tout en améliorant le niveau de la QoS aux usagers CN ;
- développer un modèle d'allocation de ressources radio avec contraintes afin de garantir un niveau minimal de QoS pour l'ensemble des clients LTE. Ce modèle doit empêcher les classes de trafics prioritaires d'accaparer les ressources radio, spécialement lors des moments de crise ;
- développer un modèle d'allocation de bande passante avec contraintes afin d'améliorer l'ordonnancement des trafics dans le réseau Backhaul et dans le réseau cœur LTE HetNet. Tout comme les ressources radio, les ressources de bande passante peuvent être accaparées par certaines classes de trafics prioritaires, ce qui pénalise les flux de trafics moins prioritaires ;
- concevoir un système de basculement (Offloading) des nouveaux bearers arrivant à la macro cellule vers les réseaux locaux sans fil, quand les ressources radio de l'eNodeB sont épuisées ;
- améliorer les communications D2D pour garantir une transmission efficace des flux basculés de la macro cellule, ainsi que pour assurer une extension de la couverture des cellules LTE à travers :
 - l'implémentation d'une solution pour le routage dans les réseaux locaux sans fil ;
 - développement d'un nouveau système pour la sécurisation des communications D2D au niveau des réseaux locaux sans fil sans affecter le niveau de la QoS.

Méthodologie de recherche

Les solutions proposées dans la littérature pour l'allocation des ressources radio focalisent sur la priorisation des usagers PS lors des situations d'urgences. Toutefois, ces solutions affectent la QoS des trafics CN. Notre modèle propose le même avantage pour le réseau PSN tout en améliorant la QoS du réseau CN. Pour offrir un accès privilégié au réseau d'accès LTE pour PSN, le présent travail se base initialement sur l'article (Borkar, Roberson et al. 2011) que nous nommons méthode classique dans cette thèse. La principale idée de la

méthode classique consiste à assurer un accès à la radio avec priorité pour les clients PSN. De ce fait, une portion de la bande de fréquence commerciale partagée est affectée aux premiers répondants, d'une façon privilégiée, autrement dit, tout bearer CN utilisant cette portion sera interrompu dès qu'une requête d'allocation est formulée par un client PS.

Ainsi, une amélioration de la solution proposée dans l'étude citée ci-dessus a été effectuée dans cette thèse pour le volet priorisation des usagers PS pour l'accès à la bande de fréquence partagée, ainsi que pour l'amélioration de la QoS des usagers CN.

En premier lieu, trois modèles d'allocation de ressources radio avec contraintes ont été implémentés, à savoir Courteous Allocation constraints Model for Frequencies (CAMF), Radio Usage Situation based Courteous constraints Allocation Model for Frequencies (RUS-CAMF) et Generalized Courteous constraints Allocation Model for Frequencies (G-CAMF). Chaque modèle consiste à développer un système mathématique pour l'attribution des ressources radio aux clients PSN et aux clients CN tout en respectant les contraintes d'allocations, la priorité et le type du chaque trafic. Chaque modèle attribue les ressources radio à deux ou plusieurs groupes d'utilisateurs. Chaque groupe d'utilisateurs détient un ensemble de caractéristiques communes, dont la priorité et le type de données à échanger. En outre, chaque groupe plus prioritaire peut céder certaines de ses ressources allouées au groupe moins prioritaire adjacent. L'opération de cessation n'est pas obligatoire, elle s'effectue plutôt par courtoisie et elle ne doit en aucun cas affecter la QoS du groupe courtois. L'implantation de ces trois mécanismes se base sur le système Courteous bandwidth Allocation constraints Model (CAM) que nous avons conçu initialement pour l'allocation des ressources de bande passante dans le réseau MPLS (DiffServ-aware MPLS Traffic Engineering) et que nous utilisons dans cette thèse pour la gestion de la bande passante avec contraintes dans le réseau Backhaul LTE reliant les eNodeB et le EPC (Evolved Packet Core), ainsi qu'au niveau du réseau cœur LTE pour l'allocation de la bande passante et l'ordonnancement des trafics qui y transitent en direction du réseau externe. Le mécanisme CAM est implémenté pour améliorer la gestion des ressources de bande passante proposée par les deux mécanismes MAM (Le Faucheur 2005) et RDM (Le Faucheur 2005).

Par la suite, les systèmes modélisés dans cette étude pour la gestion de bande de fréquence par contraintes, ont été utilisés pour le développement de l'algorithme Courteous Priority Access (CPA) et de l'algorithme CPA with Offload (CPAwO). CPA assure l'allocation des ressources radios commerciales partagées avec priorisation pour les usagers PSN dans les macros cellules LTE. CPAwO étend le fonctionnement de CPA en basculant des bearers de la macro cellule vers les petites cellules LTE. CPAwO allouera donc des ressources radio additionnelles, aux bearers basculés, au sein des réseaux locaux sans fil opérant dans les petites cellules. Cette solution est pertinente dans le cas où les ressources radio sont limitées dans la macro cellule.

Les bearers basculés vers les réseaux locaux sans fil vont utiliser les ressources radio publiques. De telles communications sont connues sous le nom de transmissions Device-to-Device (D2D) (Lin, Andrews et al. 2013). Afin d'améliorer la gestion des ressources dans les réseaux D2D, deux solutions ont été proposées dans cette thèse pour le routage des paquets, à savoir Reliable Butterfly effect Construction algorithm (RBC) et Load Balancing based Selective AOMDV algorithm (LBS-AOMDV). RBC est un algorithme de signalisation dont le but est de construire des effets papillon dans un réseau maillé sans fil (WMN). Le but de construire les réseaux papillon dans le réseau WMN est de permettre l'application du codage réseau (Network Coding ou NC) pour la transmission de données au sein du WMN, du fait que la revue de littérature a prouvé que l'utilisation de la méthode NC avait apporté son amélioration de la QoS pour les réseaux filaires et sans fils (Ho, Koetter et al. 2003, Li and Li 2004, Gkantsidis and Rodriguez 2005, Fragouli, Widmer et al. 2006, Katti, Rahul et al. 2006, Matsuda, Noguchi et al. 2011). Rappelons que l'application du NC à travers un réseau papillon élimine le risque d'un éventuel échec de décodage. En effet, certaines topologies empêchent le décodage des paquets codés par le mécanisme NC, étant donné que le processus de décodage exige la réception, par la destination, d'un certain nombre de paquets via des routes différentes lui permettant de réaliser des opérations mathématiques afin d'extraire les paquets natifs. Sachant qu'un réseau papillon détient une topologie spécifique, ce ne sont pas tous les réseaux WMN qui englobent des effets papillon. Ceci dit, une solution alternative a été développée dans ce travail afin d'assurer le routage des données en l'absence

des réseaux papillon dans le WMN. L'approche proposée tient à construire un multipath entre la source de données et sa destination lors de la transmission d'informations. Le multipath ayant des routes disjointes permet d'équilibrer la charge (Load Balancing) dans le réseau WMN. Cette solution a été nommée LBS-AOMDV. Basée sur AOMDV, elle permet de définir un ensemble de routes disjointes entre une source et une destination données en respectant les contraintes de la QoS requise par la source et en générant moins de trafic de contrôle que le protocole AOMDV.

Finalement, le problème de la sécurité dans les réseaux locaux sans fil a été abordé pour sécuriser les transmissions D2D. Un algorithme nommé Generalised Secure Network Coding based Data Splitting (G-SNCDS) a été conçu. Il s'agit d'une solution qui se base sur l'utilisation d'une technique de division et de mixage des données à transmettre via les effets papillon lors de l'application du codage réseau. Trois types d'attaques sont considérés dans ce travail, à savoir, l'attaque de confidentialité, l'attaque d'intégrité et l'attaque de disponibilité. Une solution partielle intégrée dans G-SNCDS, nommée SNCDS, a été développée initialement pour contrer les attaques de confidentialité. Elle est présentée dans notre article (Tata and Kadoch 2014). Elle traite la première partie relative à G-SNCDS présenté dans cette thèse.

Contributions

La principale contribution de ce travail est humanitaire. Elle se traduit par l'augmentation du nombre des vies sauvées lors d'un désastre. Ceci se réalise grâce à l'accroissement du nombre d'informations échangées via les premiers répondants lors des situations de crises, par une meilleure gestion des ressources du réseau. Ainsi que par l'utilisation de la solution D2D qui permet de fournir des ressources additionnelles au réseau LTE et assure une extension de la couverture de celui-ci pour atteindre des zones isolées non couvertes les macros cellules LTE. Ceci peut contribuer fortement à sauver des vies, alors que cela est presque impossible quand les ressources de l'eNodeB sont épuisées ou quand son rayon n'atteint pas la zone de la catastrophe.

L'autre contribution de ce travail est écologique, du fait que tous les modèles de gestion des ressources avec contrainte conçus dans cette thèse et qui sont basés sur la courtoisie fonctionnent de telle sorte à minimiser le gaspillage des ressources, qui peut être interprété par une sous-utilisation des ressources allouée par des usagers PSN ou CN. L'idée est que chaque classe de trafic prioritaire peut céder certaines ressources allouées à une autre classe moins prioritaire si cette cessation n'affecte pas sa QoS.

Cette solution a aussi des retombées économiques grâce à l'utilisation des fréquences publiques gratuites lors des communications D2D, une approche tout à fait transparente aux clients, capable de réduire le coût global de certaines transmissions, et grâce à la sécurisation des communications D2D sans ajouter un trafic de contrôle additionnel.

Sans oublier la contribution technologique qui se résume, entre autres, en un ensemble de dix algorithmes et modèles mathématiques permettant d'améliorer les communications PSN sur les réseaux LTE HetNets.

Structure de la dissertation

Cette thèse est organisée comme suit. Le chapitre 1 relate quelques généralités sur le réseau LTE. Le chapitre 2 présente des généralités sur la technologie codage réseau. Le chapitre 3 résume la revue de littérature. Le chapitre 4 présente notre solution pour l'allocation avec contrainte des ressources de bande passante pour le réseau Backhaul et le réseau cœur LTE. Le chapitre 5 détaille notre approche pour la gestion des ressources radio au niveau du réseau d'accès radio LTE. Le chapitre 6, en plus de concevoir une solution de transmission sécurisée pour la transmission des données D2D, expose nos approches modélisées pour la signalisation dans les réseaux locaux sans fils. Une conclusion et des recommandations concluent la thèse.

CHAPITRE 1

Généralité sur les réseaux LTE HetNets

1.1 Introduction

Le réseau Long Term Evolution (LTE) est une norme 3GPP. Elle offre un débit allant jusqu'à 50 Mbps pour le lien montant et 100 Mbps pour le lien descendant. Le réseau LTE se caractérise par la transmission des informations, quelles qu'elles soient par IP (Internet Protocol), y compris les applications multimédias, d'où son appellation réseau All-IP ou Full-IP. Il est divisé en trois parties principales, à savoir le réseau d'accès, ayant une interface radio, le réseau cœur fonctionnant avec la technologie IP et le réseau Backhaul. Le réseau d'accès et le réseau cœur sont reliés via le réseau Backhaul dont le rôle est d'agréger le trafic transmis par les différentes stations de base LTE (eNodeB) vers le réseau cœur, appelé aussi Evolved Packet Core (EPC).

Le présent chapitre relate des généralités de la technologie LTE qui sont pertinentes pour la compréhension des différentes solutions développées dans cette thèse. Le contenu de ce chapitre est extrait de (Monfreid 2009), (Bouguen, Hardouin et al. 2012), (Yahiya 2011), et de (Ergen 2009).

1.2 Architecture du réseau LTE

L'architecture LTE comprend trois domaines (figure 1.1), à savoir le réseau d'accès (Radio Access Network, RAN), le réseau Backhaul et le réseau cœur EPC (Evolved Packet Core). Les différents composants du réseau LTE sont reliés via des interfaces spécifiques (voir figure 1.3). Cette section introduit les différents domaines d'un réseau LTE, ainsi que ses interfaces.

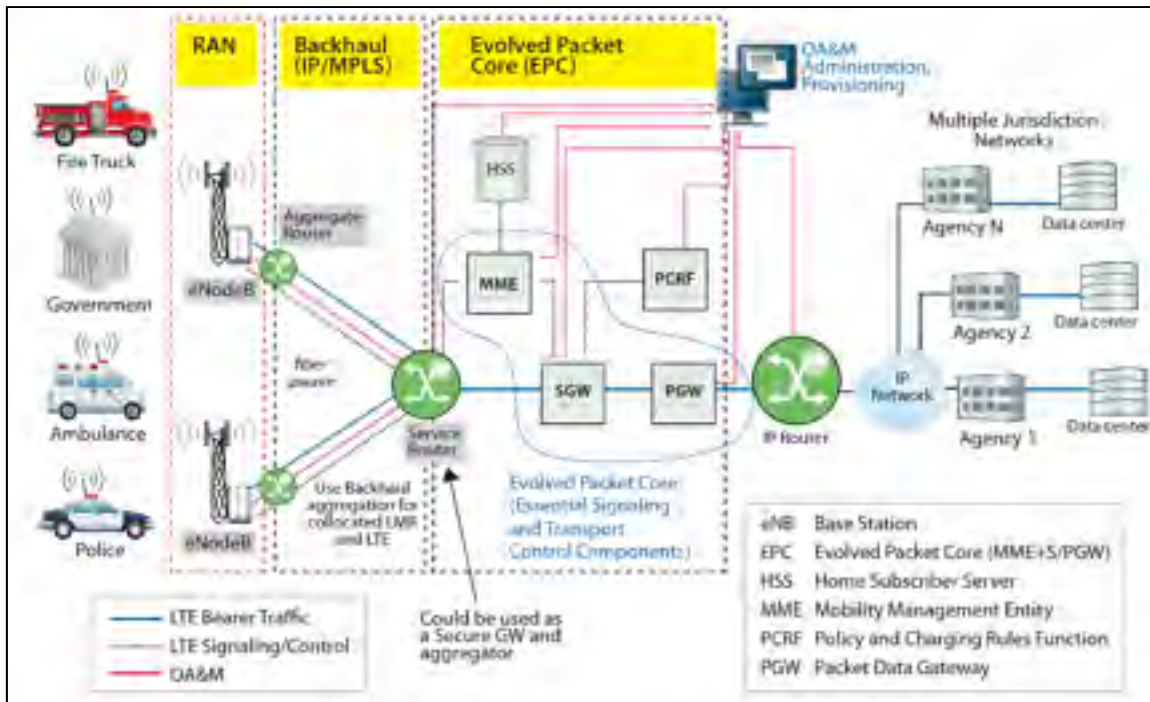


Figure 1.1 Les trois réseaux de la technologie LTE
Tirée de Lucent (2010)

1.2.1 Le réseau d'accès

Le réseau d'accès représente l'interface radio du réseau LTE (Radio Access Network, RAN). Il a pour rôle d'assurer l'accès aux ressources radio à l'UE (User Equipment). Il est composé d'une station de base, l'eNodeB, qui relie l'UE à l'EPC via le plan usager et le plan de contrôle (voir ci-dessous les définitions correspondantes). Notons que le contrôleur de stations de base introduit dans les 2G et 3G a été supprimé dans l'architecture LTE. Ses fonctionnalités ont été intégrées en partie à l'eNodeB et en autre partie à EPC.

1.2.2 Le réseau Backhaul

Le réseau Backhaul LTE est le réseau d'agrégation responsable du regroupement du trafic généré par plusieurs stations eNodeBs reliées entre elles. Il représente le lien de connexion entre les nœuds eNodeBs et le réseau cœur EPC. La figure 1.2 illustre un exemple de solution IP/MPLS proposée par la compagnie Alcatel-Lucent. Dans cette figure, on peut constater les

différentes technologies qui permettent le déploiement d'un réseau Backhaul LTE. Entre autres, on distingue la fibre optique et le réseau hertzien. La technologie MPLS étant la plus rentable est devenue la meilleure solution pour l'agrégation des débits en provenance du réseau d'accès et allant vers les réseaux cœur LTE. L'intégration de la technologie MPLS dans le réseau Backhaul LTE est prometteuse. Cette technologie permet l'amélioration de la QoS de bout en bout dans les réseaux LTE, comme elle assure de nouveaux services mobiles.

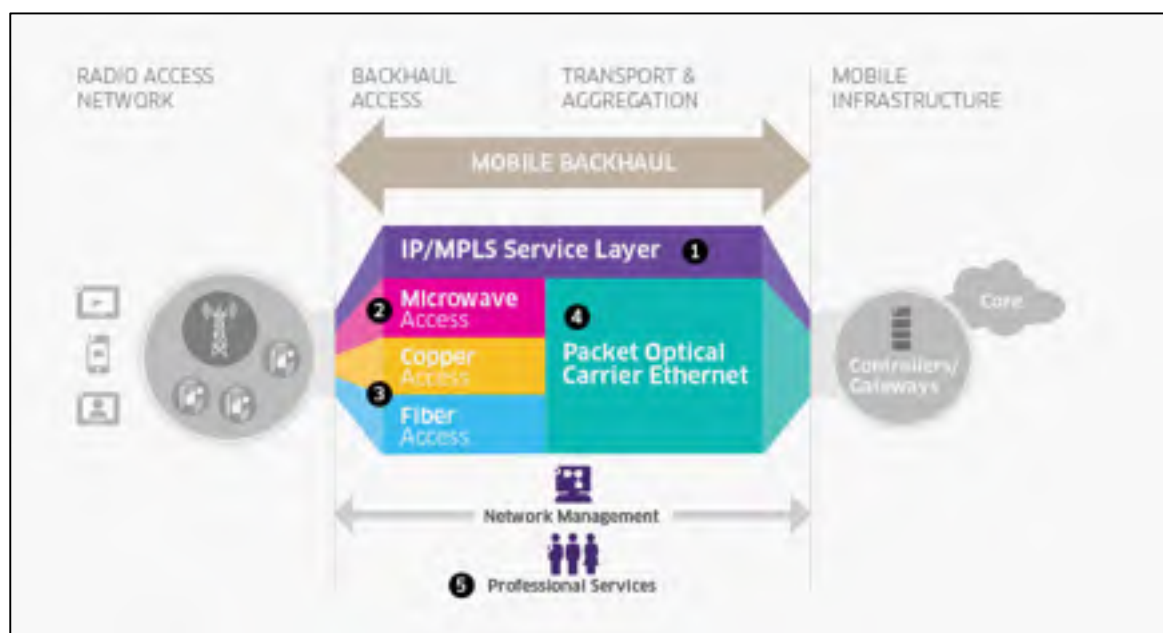


Figure 1.2 Réseau Backhaul LTE
Tirée de Lucent (2014)

1.2.3 Le réseau cœur : EPC

Le réseau cœur regroupe des équipements reliés par des fils terrestres assurant certaines fonctionnalités du réseau LTE, tels que l'enregistrement des utilisateurs, la sécurité, la mobilité et le contrôle des appels.

EPC est un réseau All-IP, autrement dit, contrairement aux générations précédentes, le réseau cœur LTE ne considère plus la commutation de circuit pour l'acheminement des appels, même quand il s'agit de la voix ou de la vidéo. La tendance actuelle veut que le réseau cœur

LTE soit de type IP/MPLS. L'intégration de MPLS dans LTE permet d'améliorer la QoS dans ce réseau.

Tel que le montre la figure 1.3, EPC se compose de plusieurs éléments, à savoir le MME, le P-GW, le S-GW, le HSS, le PCRF et l'ePDG. Ces principaux composants sont définis dans ce qui suit :

1.2.3.1 MME (Mobility Management Entity)

MME est le nœud de l'EPC qui effectue la signalisation entre l'UE et le réseau cœur. Le MME s'occupe de la gestion des bearers en termes d'établissement, de maintenance et de relâche de ces bearers. Il est aussi responsable de la gestion de la connexion de signalisation et de la sécurité entre l'UE et le réseau cœur. La connexion d'un UE est maintenue tant qu'il demeure enregistré dans le réseau.

1.2.3.2 S-GW (Serving Gateway)

S-GW a pour rôle principal la gestion de flux de données du plan usager transitant entre l'eNodeB et le P-GW. Elle sert d'un point de passage de tous les paquets destinés aux usagers LTE. Elle s'occupe aussi de la gestion de certaines fonctions au niveau du réseau visité lors de l'itinérance (Roaming), telle que l'envoi des informations sur la facturation.

1.2.3.3 P-GW (Packet Data Network Gateway)

P-GW est une passerelle qui assure la connexion du réseau LTE aux autres types de réseaux, tels que WMAX et UMTS. Elle a pour rôle d'attribuer une adresse IP à l'UE, de concrétiser l'application de la QoS et d'analyser les paquets du plan usager (voir la section 1.4). Elle s'occupe aussi de l'application des règles prédéfinies relatives aux différents clients, comme elle permet d'appliquer une politique de facturation par flux de données en considérant les règles établies par le PCRF.

1.2.3.4 PCRF (Policy and Charging Rules Function)

PCRF est un nœud optionnel pour le réseau EPC. Son rôle principal est d'appliquer les règles de gestion de facturation de l'utilisateur en fonction des règles prédéfinies qui s'appliquent sur lui.

1.2.3.5 HSS (Home Subscriber Server)

HSS est le serveur d'enregistrement des usagers UE dans l'architecture LTE. Il contient toutes les informations de souscription relatives aux utilisateurs, telles que le profil de la QoS ou la restriction d'accès en Roaming. Il détient aussi l'information relative à l'identité du MME auquel est attaché un abonné.

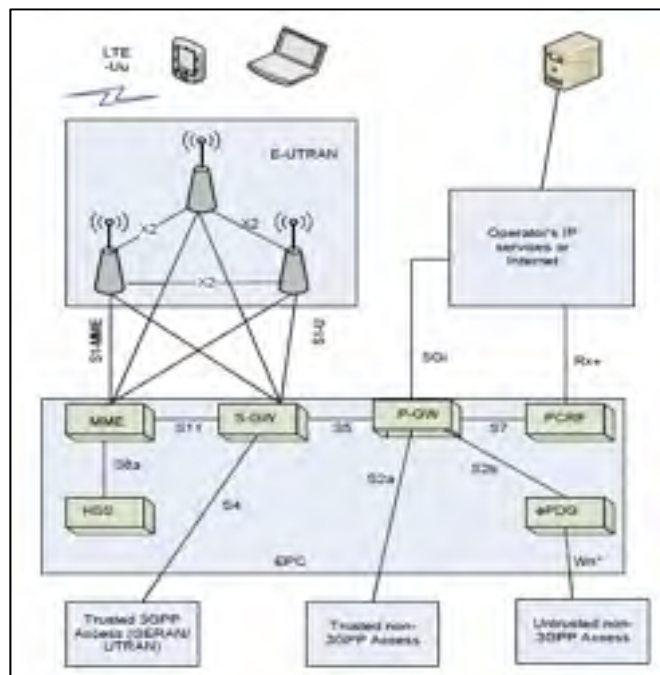


Figure 1.3 Le réseau EPS
Tirée de Wang, Sicker et al. (2013)

1.2.4 Les interfaces de connexion

La technologie LTE a défini plusieurs types d'interfaces selon les équipements qu'elles relient. La figure 1.4 illustre la connexion de certains équipements LTE par les interfaces appropriées. Entre autres, on distingue :

- deux interfaces X2, à savoir X2-CP et X2-US. La première est porteuse de trafic de contrôle entre les différents eNodeB, alors que la deuxième transporte le trafic de données entre ces nœuds;
- deux interfaces S1, à savoir S1-MME et S1-U. La S1-MME relie l'eNodeB avec le MME. Elle transporte la signalisation du plan de contrôle entre ces deux équipements. La S1-U transmet les données utilisateurs entre l'eNodeB et le S-GW;
- trois interfaces S2, notamment S2a, S2b et S2c. L'interface S2a relie les réseaux crédibles non 3GPP à P-GW, l'interface S2b interconnecte le P-GW et le ePDG et la S2c sert de jointure entre le P-GW et les réseaux non crédibles de type IP ou 3GPP;
- une interface S3. Celle-ci permet l'échange des informations relatives aux usagers et aux bearers entre SGSN (Serving GPRS Support Node) et MME;
- une interface S4 qui transporte le trafic entre les usagers 2G et le P-GW.
- une interface S5 qui sert d'interconnexion entre le S-GW et le P-GW pour l'échange du trafic data correspondant au plan usager;
- une interface S6 qui relie le MME avec le HSS. Cette connexion sert à accomplir les tâches d'enregistrement, d'authentification et d'autorisations des usagers UEs au niveau du HSS;

et des mesures radio pour évaluer la qualité du signal et le niveau de puissance de réception. La couche physique est responsable aussi de détecter la présence de cellules voisines et de s'y connecter.

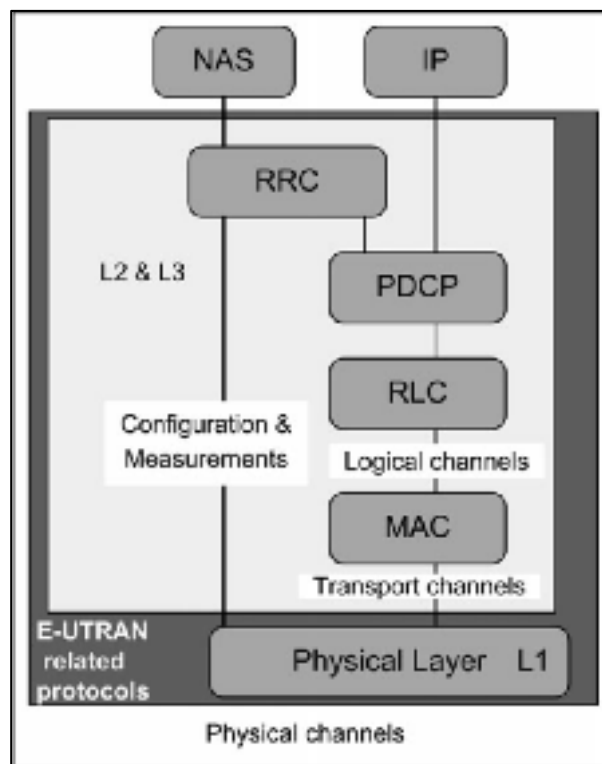


Figure 1.5 Les couches du réseau LTE
Tirée de Yahiya (2011)

1.3.2 Couche liaison de données

Elle est composée de trois sous couches, à savoir, la PDCP (Packet Data Compression Protocol), la RLC (Radio Link Control) et la MAC (Medium Link Control).

1.3.2.1 Sous couche PDCP

La sous-couche Packet Data Compression Protocol (PCDP) assure l'envoi de données et la signalisation de la couche supérieure RRC (Radio Resource Control) sous un format crypté. De plus elle assure la compression de l'en-tête Service Data Unit (SDU) PCDP grâce au

protocole RoHC (Robust Header Compression). Cette technique est bénéfique pour améliorer la transmission radio, surtout pour le trafic temps réel, tel que la voix. D'autre part, elle s'occupe de la détection et de la suppression des unités de données (PDU) PCDP doublant. Cette fonction est fort utile lors des processus de handover entre les cellules LTE.

1.3.2.2 Couche RLC

La sous-couche Radio Link Control (RLC) est responsable du transfert des PDU de la couche supérieure, de la concaténation, de la segmentation et du l'assemblage de RLC SDU. RLC est également responsable de la réorganisation des PDU de données RLC, de la détection des doublons RLC, ainsi que de la détection et de la correction d'erreur par ARQ (Automatic Repeat reQuest).

1.3.2.3 Sous couche MAC

La sous-couche Medium Link Control (MAC) est responsable du multiplexage des MAC SDU d'un ou plusieurs canaux logiques sur des blocs de transport (TB) qui doivent être livrés à la couche physique sur les liens de transport. Elle s'occupe aussi du multiplexage de MAC SDU d'un ou de plusieurs liens logiques de blocs de transport (TB) livrés à partir de la couche physique sur les liens de transport. Elle se charge aussi de la correction d'erreur, et de la gestion de priorité entre les canaux logiques d'un UE.

1.3.3 Couche RRC

La Radio Resource Control (RRC) est connue sous le nom de Layer 3. Les principaux services et fonctions de la couche RRC comprennent la diffusion des informations sur le système liées au protocole Non Access Stratum (NAS), la diffusion des informations sur le système qui sont liées à la couche d'accès, la pagination, la création, l'entretien et la libération d'une connexion RRC entre l'UE et E-UTRAN, les fonctions de sécurité, y compris la gestion des clés, la création, la configuration, la maintenance et la libération des bearers radio.

1.4 Plan de contrôle et plan usager

Le réseau LTE est conçu de telle sorte à ce que les données des utilisateurs et celles reliées à la signalisation soient séparées et acheminées sur des médias à part. Les données des utilisateurs sont transmises par le plan usager (User Plane), par contre, les données de la signalisation sont portées par le plan de contrôle (Control Plane). Chaque plan comprend des protocoles utiles pour offrir des services spécifiques. La figure 1.6 illustre la pile protocolaire de chaque plan. Tel que le montre la figure, les deux plans partagent un ensemble de protocoles communs résidants dans les couches inférieures. Toutefois, la compression des entêtes n'est pas applicable pour le plan de contrôle au niveau de ces couches. Par ailleurs, les deux plans utilisent des protocoles différents dans les couches supérieures.

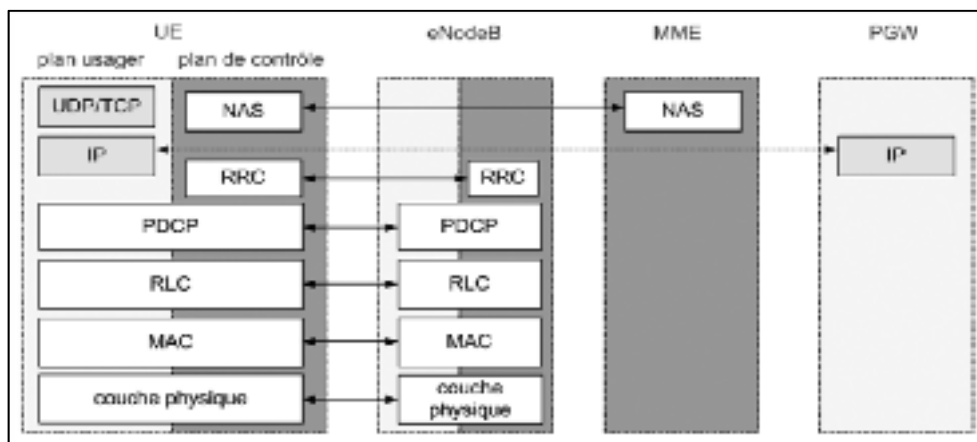


Figure 1.6 Le plan de contrôle et le plan usager
Tirée de Bouguen, Hardouin et coll. (2012)

1.5 Transmission des données dans LTE

1.5.1 Service data Flow

Le service Data Flow (SDF) est une agrégation de flux IP d'utilisateurs LTE correspondants à un même service (figure 1.7). Le mapping des flux IP dans des SDF correspond à la politique appliquée par le fournisseur du service pour l'allocation des ressources et la taxation de chaque service offert.

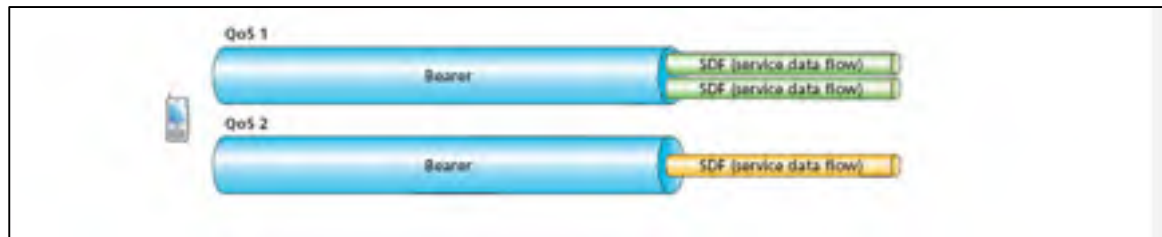


Figure 1.7 Le Single Data Flot
Tirée de Lucent (2009)

1.5.2 Les Bearers EPS

« Un bearer peut être vu comme un tuyau entre deux entités du réseau qui communiquent entre elles sur une interface, tuyau dont certaines caractéristiques sont négociées entre ces entités lors de son établissement » (Bouguen, Hardouin et al. 2012).

Un bearer EPS est un regroupement logique de plusieurs SDFs appartenant à une même classe de QoS qui interconnecte le réseau d'accès au réseau cœur EPC via les deux entités du réseau UE et P-GW. Le regroupement de plusieurs SDFs dans un bearer EPS spécifique s'effectue en utilisant un filtre de trafic appelé Traffic Flow Template (TFT) qui sera détaillé un peu plus loin dans ce chapitre.

Un bearer EPS est composé de trois types de bearers, à savoir, un bearer radio, un bearer S1 et un bearer S5/S8 (figure 1.8). Le bearer radio, appelé aussi Air bearer est établi entre l'utilisateur UE et l'eNodeB à travers l'interface radio, le bearer S1 est déployé entre l'eNodeB et le S-GW à travers l'interface S1-U et le bearer S5/S8 est établi entre les deux passerelles S-GW et P-GW.

Par ailleurs, on peut distinguer deux types de bearers EPS. Notamment, un bearer par défaut et un bearer dédié. En effet, pour pouvoir échanger des informations sur le réseau LTE, un utilisateur LTE (UE) doit d'abord s'enregistrer dans le réseau. Pour ce faire, il doit établir un bearer EPS initial. Ce bearer est appelé le bearer par défaut. Ce bearer est affecté à l'UE d'une façon permanente tant qu'il demeure attaché au réseau. L'interruption du bearer par

défaut d'un UE s'effectue donc suite à sa déconnexion. D'autre part, un bearer par défaut est un bearer qui ne garantit aucune QoS. Par conséquent, l'acheminement de flux de trafic ayant des besoins spécifiques de QoS, tel que la voix, exige l'établissement d'un autre type de bearers. Il s'agit de bearer dédié. Un bearer dédié est un bearer EPS établi entre un UE et un P-GW ayant la capacité de garantir un certain niveau de QoS pour le trafic qu'il transporte.

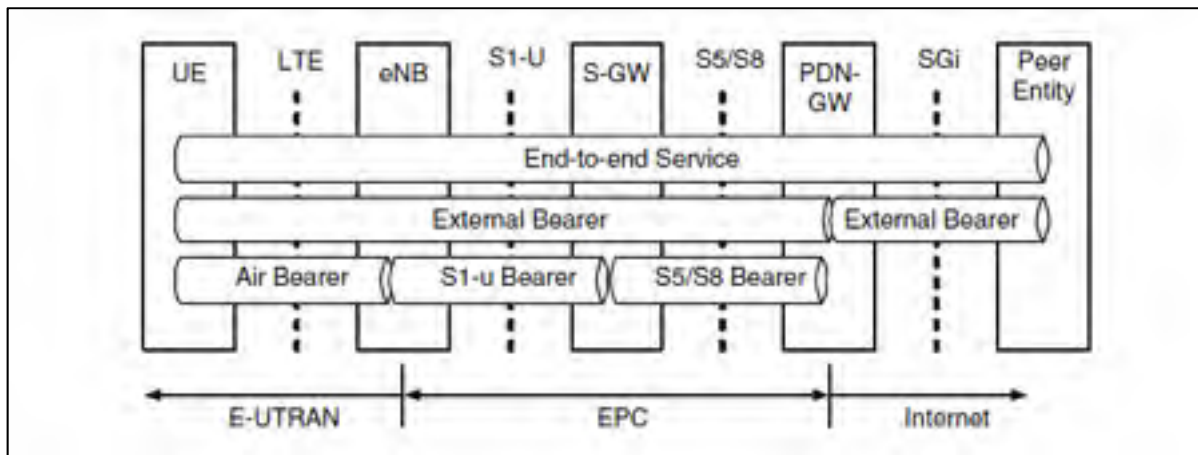


Figure 1.8 Les bearers EPS
Tirée de Ergen (2009)

1.5.3 Traffic Flow Template

Un Traffic Flow Template (TFT) est un classificateur de trafics data du réseau LTE. Son rôle est de regrouper un ou plusieurs SDF ayant les mêmes exigences de QoS dans un même bearer EPS. Il sert d'un outil d'application de politique de priorisation de trafics dans le réseau LTE qui est établie par le fournisseur de service.

La figure 1.9 illustre la classification de trois SDF dans deux bearers EPS suite au processus de filtrage de paquets. En effet, les flux appartenant aux deux SDF de gauche sont forcés à intégrer le bearer d'en haut, alors que les flux de trafic appartenant au troisième SDF doivent être acheminé via le bearer EPS d'en bas.

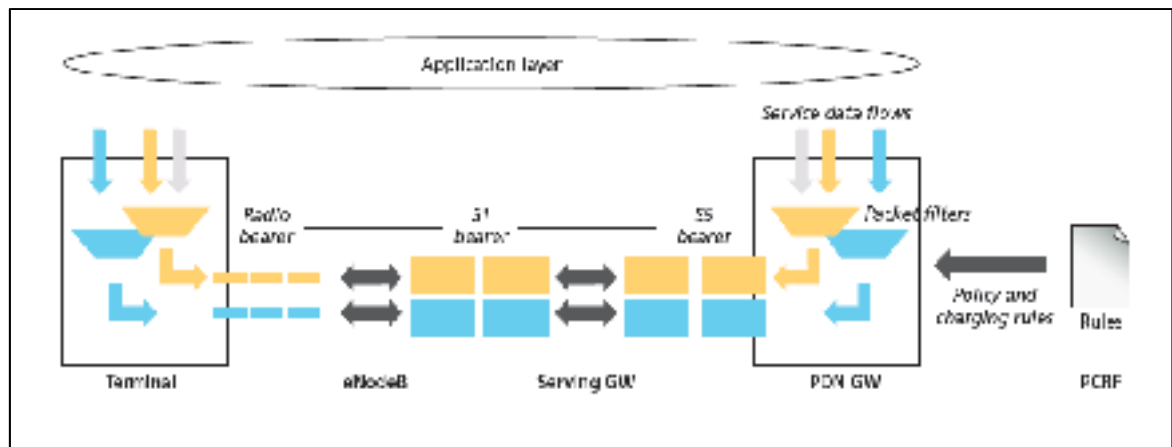


Figure 1.9 Filtrage des trafics dans LTE
Tirée de Lucent (2009)

1.6 La QoS dans LTE

La QoS d'un bearer EPS est définie par certains paramètres. Notamment, la classe de service déterminée par un identificateur de classe de QoS (QoS Classe Identifier, QCI), la priorité d'allocation et de rétention (Allocation and Retention Priority, ARP), le débit garanti (Guaranteed BitRate, GBR) et le débit maximal ou Maximal Bit Rate (Maximal Bit rate, MBR)

1.6.1 QoS Class Identifier

Le paramètre QCI (QoS Class Identifier) est introduit par la technologie LTE afin de déterminer les niveaux de QoS dans un même bearer EPS. Le standard LTE a défini neuf valeurs QCI pour permettre d'appliquer la différenciation de service entre les différents flux de trafics transitant entre l'utilisateur UE et l'entité P-GW. La différenciation de service au sein du réseau EPS permet de garantir un bon niveau de QoS pour les trafics GBR, tel que la voix, ainsi que d'assurer un bon service pour les trafics Non-GBR sans les laisser accaparer toutes les ressources du réseau EPS.

Resource Type	QCI	Priority	Packet Delay Budget ¹²	Packet Error Loss Rate ¹³	Example Services
GBR	1	2	100 ms	10^{-2}	Conversational Voice
	2	4	150 ms	10^{-3}	Conversational Video (Live Streaming)
	3	3	50 ms	10^{-3}	Real Time Gaming
	4	5	300 ms	10^{-6}	Non-Conversational Video (Buffered Streaming)
Non-GBR	5	1	100 ms	10^{-6}	IMS Signalling
	6	6	300 ms	10^{-6}	Video (Buffered Streaming)
	7	7	100 ms	10^{-3}	TCP-based (e.g., www, e-mail, chat, ftp, p2p file sharing, progressive video, etc.)
	8	8	300 ms	10^{-6}	Voice, Video (Live Streaming), Interactive Gaming
	9	9	300 ms	10^{-6}	Video (Buffered Streaming)
					TCP-based (e.g., www, e-mail, chat, ftp, p2p file sharing, progressive video, etc.)
					QCI typically used for the default bearer of a UE/PDN

Figure 1.10 QoS et différenciation de services dans LTE
Tirée de Hallahan and Peha (2010)

Tel que le montre la figure 1.10, les quatre premières valeurs QCI sont réservées pour le trafic GRB et les cinq autres sont allouées au trafic Non-GBR. Notons que plus la valeur QCI est petite, plus le trafic relatif est prioritaire. Par ailleurs, la classification QCI est appliquée lors de la transmission des paquets à travers le réseau.

1.6.2 Allocation and Retention Priority

Le mécanisme Allocation and Retention Priority (ARP) est conçu pour décider d'accepter ou de refuser l'admission d'un nouveau bearer quand les ressources sont limitées dans le réseau LTE. En effet, suite à l'impossibilité de satisfaire une requête d'établissement d'un nouveau bearer à cause du manque des ressources, ARP aura le choix entre le refus du nouveau bearer et l'interruption d'un bearer actif dans le réseau pour accepter l'établissement du nouveau bearer. Notons que le bearer à interrompre doit avoir une priorité plus basse que celle du nouveau bearer à établir, comme il doit être vulnérable à la préemption, alors que le nouveau bearer doit avoir la capacité de préemption égale à (oui binaire). D'autre part, il est important

de mentionner qu'une fois que le bearer est établi, le mécanisme ARP n'intervient point pour la gestion du trafic, son rôle se limite donc juste lors de l'établissement du nouveau bearer. Dans le cas où les ressources ne seront plus disponibles pour un bearer déjà établi, ce bearer sera interrompu.

1.6.3 Les paramètres du débit dans LTE

Dans cette section on définit les paramètres du débit LTE, à savoir, le GBR, le Non-GBR et le MBR.

- le GBR (Guaranteed Bit Rate) est le paramètre qui caractérise le débit garanti dans le réseau EPS. Il est offert par les bearers GBR pour les services exigeant un certain niveau de QoS. Notons que le débit réel peut être inférieur au débit garanti si le nombre de paquets délivrés est réduit;
- le Non-GBR (non Guaranteed Bit Rate) est le paramètre de débit qui correspond au débit offert par les bearers Non-GBR. Tel que son nom l'indique, le débit Non-GBR n'est pas garanti, donc une application Non-GBR ne peut pas avoir la garantie d'allouer les ressources requise pour l'acheminement du flux de données qui lui correspond;
- le MBR (Maximum Bit Rate) est le taux maximal de débit autorisé sur un bearer EPS. Tout débit dépassant cette valeur va subir une opération de lissage par l'équipement qui le transmet.

1.7 Conclusion

Ce chapitre ne se veut pas un document détaillant la technologie LTE. Il ne contient plutôt que des généralités que nous avons jugé capital de connaître et apprendre avant la lecture de cette thèse afin de faciliter sa compréhension. En effet, les notions de base qui sont présentes dans ce chapitre aident à faciliter la compréhension des différentes solutions développées dans ce document, du moment que ces solutions intègrent les concepts utilisés par la technologie LTE, tels que le bearer, le UE, le eNodeB, le mécanisme ARP et QCI. D'autres

notions qui sont pertinentes à la compréhension des termes cités dans cette thèse ont aussi été présentées, telles que le plan contrôle, le plan usager, ainsi que le modèle en couche LTE.

Comme il a été mentionné dans l'introduction, la matière présentée dans ce chapitre a été extraite des ouvrages dont les références sont (Monfreid 2009), (Bouguen, Hardouin et al. 2012), (Yahiya 2011) et (Ergen 2009).

CHAPITRE 2

Généralités sur le codage réseau

2.1 Introduction

La transmission de données au sein des réseaux sans fil, notamment, dans les réseaux maillés sans fils, s'effectue via des diffusions multicast. Tous les nœuds se trouvant dans le rayon de diffusion vont recevoir le signal. En considérant ce mode de transmission, le problème des nœuds cachés et celui des nœuds exposés apparaissent et rétrogradent, par conséquent, la fiabilité de ce type de réseaux. Les transmissions simultanées vers un même nœud causent donc des collisions à son niveau.

Certaines études, telle que (Wang, Zhang et al. 2008), ont été menées afin de passer outre ces problèmes, et rendre la communication sans fil plus efficace en permettant la transmission simultanée entre nœuds même s'ils sont cachés ou exposés l'un par rapport à l'autre. Ces études convergent vers l'utilisation du codage réseau dans les réseaux sans fil pour permettre d'envoyer simultanément, plusieurs symboles en provenance de plusieurs sources.

Établie par le chercheur Ahlswede (Ahlswede, Cai et al. 2000), le codage réseau (Network Coding) est une technique permettant la transmission simultanée de plusieurs flux de données parvenant d'une ou de plusieurs sources vers une ou plusieurs destinations tout en éliminant le risque d'interférences. En effet, lors du processus d'acheminement de paquets, les nœuds constituant un réseau traditionnel se contentent de copier et de diffuser l'information, une opération connue sous le nom de « Copy and forward ». Alors que les nœuds aptes à effectuer un codage réseau codent l'information reçue avant de la retransmettre, on parlera donc de l'approche « Copy, code and farward ». Certainement, le codage réseau effectué par les nœuds encodeurs apporte des améliorations remarquables sur la performance des réseaux de télécommunications, en termes de réduction des délais de transmission et des pertes de paquets, ainsi que d'augmentation du débit (Ho, Koetter et al. 2003, Li and Li 2004,

Gkantsidis and Rodriguez 2005, Fragouli, Widmer et al. 2006, Katti, Rahul et al. 2006, Matsuda, Noguchi et al. 2011).

Le chapitre présent fournit les notions fondamentales de la technologie codage réseau. Notamment, son principe de fonctionnement, ses bénéfices et ses types.

2.2 Principe de fonctionnement du Network Coding

Une meilleure façon pour comprendre le fonctionnement du processus codage réseau est d'analyser l'application du codage réseau dans un réseau papillon et autres. (voir figure 2.1).

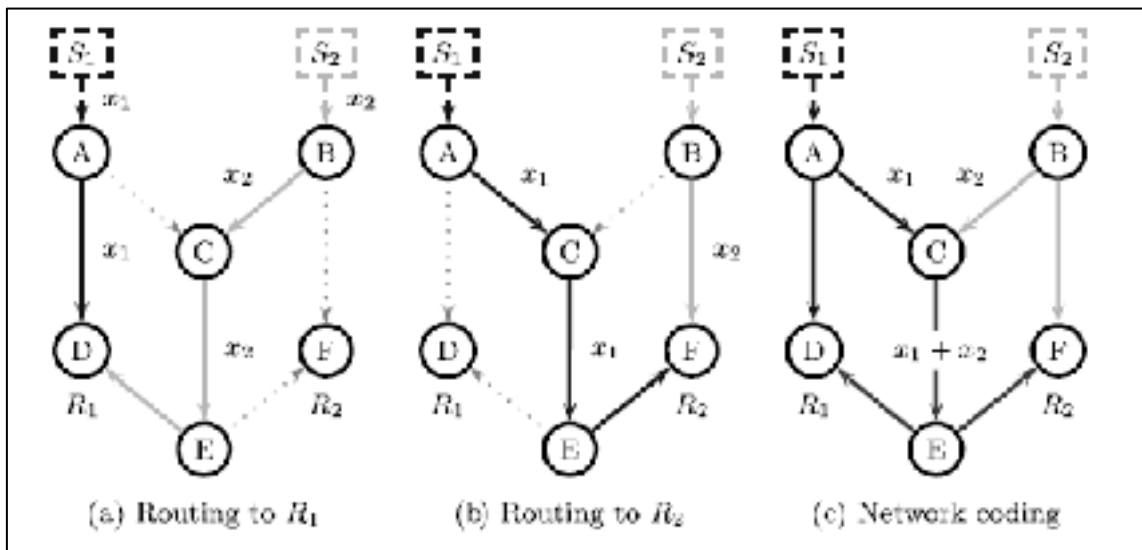


Figure 2.1 Bénéfice du codage réseau : Multicast
Tirée de Fragouli and Sojanin (2007)

La figure ci-dessus représente une comparaison entre un routage multicast classique (parties a et b) et un routage multicast avec codage réseau (partie c) dans un réseau papillon. Dans cet exemple, les deux sources, S_1 et S_2 , veulent transmettre des données aux deux destinations D et F. afin d'éviter les problèmes de collusions dans le réseau. Le routage classique effectue le

multicast en un total de quatre transmissions (partie a et b). Cependant, le routage avec codage réseau effectue la même opération en seulement trois transmissions.

L'application du codage du réseau permet aux deux nœuds A et B de transmettre leurs paquets simultanément vers le nœud C. Celui-ci au lieu de retransmettre l'information reçue telle qu'elle est, il la code d'abord. Le codage utilisé dans cet exemple est représenté par un simple XOR des informations reçues de la part des nœuds émetteurs. Notons qu'une telle transmission permet aux deux nœuds récepteurs de décoder l'information reçue. Comme R_1 (le nœud D) reçoit x_1 et $x_1 \oplus x_2$, alors il sera en mesure de déduire x_2 . De la même manière, R_2 (le nœud F) déduira x_1 . À travers cet exemple, le codage réseau se montre comme une solution pertinente pour réduire les collusions dans les réseaux de télécommunications tout en améliorant leurs performances.

Un autre bénéfice du codage réseau consiste dans la réduction de la consommation d'énergie, qui représente un défi important pour les réseaux sans fil, et donc une meilleure exploitation des ressources du réseau. Cette idée est illustrée dans la figure 2.2

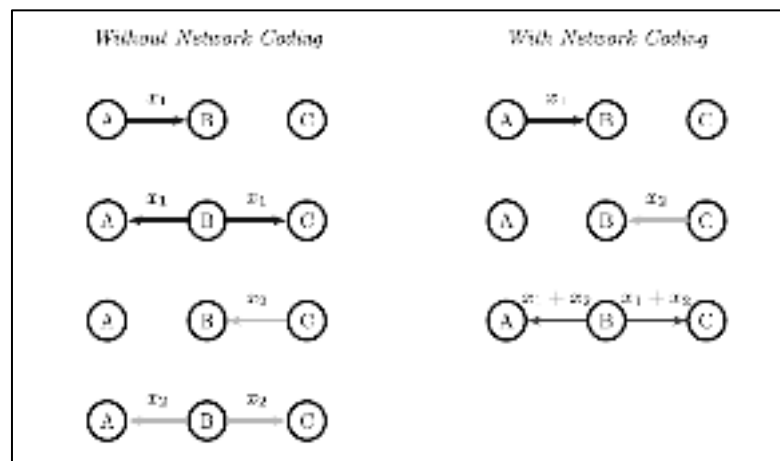


Figure 2.2 Bénéfice du codage réseau : Réduire le nombre de transmissions
Tirée de Fragouli and Sojanin (2007)

La figure 2.2 décrit la transmission de paquets entre trois nœuds sans collusions. Cette transmission concerne l'envoi de paquet de A vers C et de C vers A en passant par B, avec et sans codage réseau. Le codage réseau permet de réduire le nombre d'opérations de transmission effectuées de quatre à trois, cette réduction se représente par 25% du nombre de transmissions totales. Par conséquent, l'énergie consommée sera diminuée et les ressources préservées.

2.3 Types de codage réseau

2.3.1 Codage réseau linéaire

Le codage réseau linéaire (Linear Network Coding) est une variante du codage réseau classique, basée sur l'attribution de codes linéaires aux différentes informations reçues par un nœud, avant de procéder à leur transmission aux nœuds voisins. Les valeurs des codes établis sont choisies dans un intervalle numérique fini.

La figure 2.3 illustre un exemple concret d'un codage réseau linéaire tiré de (Perillo 2007). Dans cet exemple, la source S transmet quatre messages M_1, M_2, M_3 et M_4 à la destination D. Chaque nœud récepteur procède au codage de l'information reçue avant sa retransmission de nouveau à la destination. La destination va recevoir $j, j \in \mathbb{N}$ paquets qui sont des combinaisons de M_i , tel que i est le nombre de messages originaux. Pour cet exemple $i=4$ et $j=3$.

Chaque X_j détient un vecteur de codage noté $(g_{j1}, \dots, g_{jn})$. Alors, X_j se calcule comme suit :

$$X_j = \sum_{i=1}^n g_{ji} M_i \quad (2.1)$$

Tirée de (Perillo 2007)

Dans l'exemple, illustré dans la figure 2.3, l'auteur a considéré les valeurs suivantes :

$m = 4, n = 3, i = (1,2,3,4)$ et $j = (1,2,3)$.

Soit la matrice G , telle que

$$G = \begin{pmatrix} g_{11} & \cdots & g_{1m} \\ \vdots & \ddots & \vdots \\ g_{n1} & \cdots & g_{nm} \end{pmatrix}, g_{ji} \in G, i = 1 \text{ à } m \text{ et } j = 1 \text{ à } n, n \text{ et } m \in \mathbb{N} \quad (2.2)$$

La matrice G est la matrice de codage utilisée pour le codage des symboles. Pour l'exemple cité ci-dessus on prend $m = 4$ et $n = 3$. Chaque élément $g_{ji} \in G, i = 1 \text{ à } 4 \text{ et } j = 1 \text{ à } 3$ représente un code pouvant être utilisé par un nœud codeur pour accomplir la tâche de codage ou de décodage. Chaque nœud codeur utilise un et un seul code pour chaque symbole. Autrement dit, le codage de deux symboles nécessite l'utilisation de deux codes.

La figure 2.3 illustre le codage final de X_1 et de X_2 en utilisant la matrice de codage G et en respectant la formule 2.1

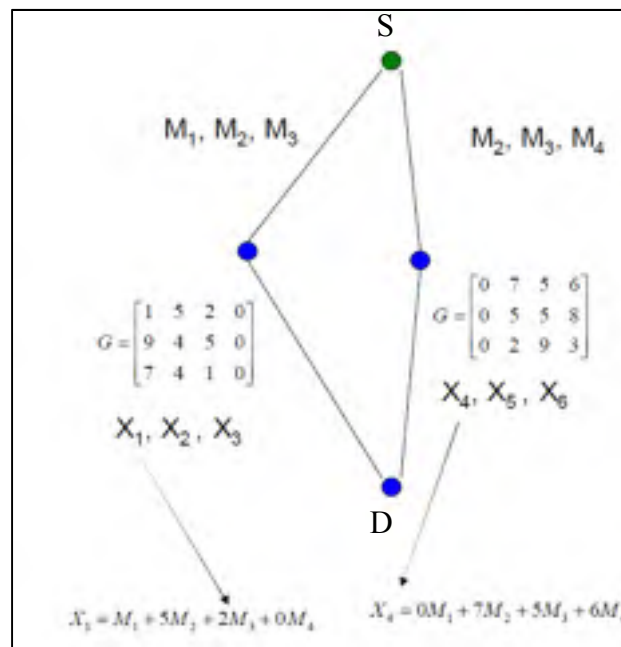


Figure 2.3 Exemple d'un codage réseau linéaire
Tirée de Perillo (2007)

2.3.2 Codage réseau linéaire aléatoire

L'inconvénient du codage linéaire consiste dans le format fixe de la matrice de codage. Autrement dit, l'application de la formule 2.1 pour l'attribution des codes nécessite une connaissance préalable de la topologie du réseau. Par conséquent, cette solution n'est pas applicable quand l'architecture du réseau n'est pas identifiée auparavant. Le codage réseau linéaire aléatoire (Random Linear Network Coding) est conçu pour résoudre cette problématique. Un codage réseau linéaire aléatoire est semblable au codage réseau linéaire. La différence entre les deux consiste dans la façon de choisir les codes. En effet, le codage réseau linéaire aléatoire n'applique pas la formule 2.1 pour le choix de ses codes, mais il procède plutôt à un choix aléatoire des codes parmi ceux qui sont inclus dans la matrice de codage G .

2.3.3 Codage réseau partiel

Il consiste à coder seulement une partie du symbole original. L'article (Wang, Zhang et al. 2008) propose une telle solution pour les réseaux sans fil capteurs (Wireless Sensor Networks, WSN). Les WSN sont connus par leurs ressources restreintes, et l'application d'un codage réseau dans ces réseaux cause des difficultés dans la gestion de l'information, notamment une perte de données considérable du fait de l'impossibilité de conserver toutes les informations pertinentes au décodage des messages dans les buffers des capteurs. Un codage partiel va donc réduire la taille de ces informations, ainsi que le nombre d'opérations effectuées pour réaliser un décodage.

2.4 Conclusion

Le codage réseau est une nouvelle approche conçue pour la transmission de données d'une façon simultanée tout en évitant les collisions. Cette technique est applicable autant pour les réseaux filaires que pour les réseaux sans fils. Il s'agit d'un mécanisme qui consiste à coder l'information en provenance de plusieurs sources afin de générer un seul paquet prêt à l'envoi vers la destination. Le codage réseau est utilisé dans le but d'améliorer la

performance dans les réseaux de télécommunications en termes de délais, débit et pertes de paquets, ainsi qu'en termes de conservation de l'énergie. Cette technologie se montre donc prometteuse pour l'amélioration de la QoS dans les réseaux de la sécurité publique opérant sur les réseaux LTE Hétérogènes.

L'application du codage réseau est effectuée dans cette thèse au niveau des petites cellules LTE pour garantir une meilleure transmission dans les réseaux locaux sans fil. Le chapitre 6 englobe les solutions utilisant le codage réseau pour la transmission des données dans les petites cellules LTE.

CHAPITRE 3

Revue de la littérature

3.1 Introduction

L'utilisation de la technologie LTE est encourageante pour améliorer la performance du réseau de la sécurité publique. Le partage des ressources radio commerciales du réseau LTE avec le réseau de la sécurité publique peut être fructueux pour contourner le problème de manque de ressources radio menant à des situations de congestions, surtout en moments de gestion de crises, là où les ressources sont le plus sollicitées et le plus utilisées. Cependant, les ressources radio de la macro cellule LTE sont limitées. Par conséquent, même un partage idéal des ressources radio entre les usagers commerciaux et les premiers répondants de la sécurité publique ne réussira point à éviter la pénurie des ressources. Une solution pertinente à cette problématique consiste dans l'utilisation de la technologie LTE Hétérogène où les usagers LTE sont autorisés d'employer les fréquences à licences publiques des petites cellules, pour effectuer leurs transmissions en cas de pénurie de ressources dans la macro cellule LTE. Toutefois, l'utilisation des fréquences publiques des petites cellules LTE ne sera pas en mesure d'apporter le même niveau de QoS et de sécurité que les fréquences privées assurent.

Par ailleurs, l'amélioration de la qualité des communications de la sécurité publique ne se limite pas dans l'obtention de ressources radio additionnelles. Le développement de modèles novateurs pour la gestion des ressources au niveau du réseau Backhaul et du réseau cœur LTE n'est pas moins pertinent que rendre meilleure la gestion des ressources radio du réseau d'accès LTE.

Ce chapitre se veut une exploration de la revue de littérature relativement à l'utilisation de LTE pour la sécurité publique, au partage du spectre radio commerciale LTE entre le réseau commercial et le réseau de la sécurité publique, au mécanisme d'offloading des macros cellules LTE vers les petites cellules, à l'efficacité de la gestion des ressources dans les

réseaux Backhaul, ainsi qu'à l'utilisation du codage réseau pour l'amélioration de la performance des réseaux locaux sans fil et aussi relativement à la sécurité dans ces réseaux.

3.2 Les réseaux LTE pour la sécurité publique

Les fréquences radio dédiées à la gestion des crises ne suffisent plus à cause de l'augmentation du nombre d'accès à ces fréquences par les premiers répondants. Cette situation a mené vers la recherche d'autres ressources et même d'autres technologies afin de répondre aux exigences des clients de ces réseaux. De ce fait, le réseau LTE a été utilisé pour résoudre cette problématique. Plusieurs études ont été effectuées afin de montrer l'utilisation de LTE pour la sécurité publique. Certaines de ces analyses sont citées ci-dessous.

Les auteurs de l'article (Doumi, Dolan et al. 2013) présente quelques concepts relatifs à la technologie de la sécurité publique (PS), à savoir les fonctionnalités et les exigences des réseaux de la sécurité publique, l'utilisation de LTE pour PS, les communications D2D, les communications de groupe sur LTE, ainsi que les bandes de fréquences PS.

Les auteurs de l'analyse (Astely, Dahlman et al. 2013) offre un résumé sur les technologies les plus importantes de LTE Advanced (LTE-A). Ils abordent le sujet des réseaux hétérogènes qui caractérise les réseaux LTE-A, ainsi que la technologie Device to Device, Machine to Machine et la technique de multiantennes.

L'article (Hallahan and Peha 2010) étudie la problématique de l'utilisation du réseau LTE pour le réseau de la sécurité publique. Les auteurs jugent que le réseau LTE détient plusieurs niveaux de priorité, ce qui se peut être convenable pour une technologie telle que la sécurité publique où les premiers répondants peuvent traiter des problèmes de différentes exigences en termes de besoins en ressources du réseau. Cela va permettre de donner plus d'importance aux trafics plus prioritaires. L'article propose aussi l'ajout d'un nouveau facteur qui puisse affecter la priorité des trafics, à savoir le facteur humain. Cela est utile pour pouvoir changer la priorité d'un trafic selon la situation de sa transmission, ainsi que pour donner un certain

pouvoir aux utilisateurs décideurs. Dans cet article plusieurs définitions relatives à LTE sont présentées et expliquées.

D'autres articles présentent des détails sur la technologie LTE et son utilisation pour les réseaux PS. Entre autres, on trouve les articles dont les références sont (Wu, Li et al. , Monfreid 2009, Lucent 2010, Ali, Taha et al. 2013, Stanze and Weber 2013). Le chapitre 1 de cette thèse offre des détails pertinents relativement aux réseaux LTE qui aident à la bonne compréhension des solutions offertes dans cette thèse.

3.3 Gestion des ressources radio partagées dans les réseaux d'accès LTE

Plusieurs études ont abordé la problématique de l'allocation des ressources radio commerciales partagées dans le réseau LTE. L'objectif est de fournir des systèmes efficaces de gestion des ressources radio dans les réseaux LTE HetNets et LTE. Les ressources partagées peuvent être utilisées conjointement par le réseau commercial et le réseau de la Sécurité Publique. Ce dernier exige une certaine priorité lors de la transmission de ses données.

L'article (Qian, Huang et al. 2009) propose une nouvelle solution pour l'allocation des ressources radio à travers mécanisme nommé « Radio Admission Control » (RAC). RAC propose de diviser les bearers en trois groupes, à savoir, le premier groupe qui comprend les bearers vulnérables à la préemption, le deuxième groupe, y compris les bearers qui ne subissent pas de préemption, mais qui peuvent interrompre ceux du premier groupe, finalement le troisième groupe est celui qui contient les bearers qui ne peuvent ni être interrompus ni interrompre les autres bearers. En outre, les bearers d'un même groupe peuvent partager la même quantité de ressources radio. Par conséquent, une fois que les ressources du deuxième groupe sont limitées, ses bearers peuvent interrompre les autres bearers du premier groupe. Cette solution permet d'organiser l'allocation des ressources radio selon les exigences des flux de trafics. Chaque trafic sera mappé à un groupe de bearers

parmi les trois déterminés par RAC, donc chaque trafic va hériter des avantages et des limites du groupe du bearer auquel il correspond.

Les auteurs de l'article (Kwan, Arnott et al. 2010) proposent un système de contrôle d'admission associé à un système de contrôle de congestion comme une solution pour réduire le blocage de bearers et la probabilité de perte d'appels dans le système. Les auteurs proposent un modèle mathématique pour le calcul de la charge dans le réseau. Ce système contribue dans le processus de réduction de la charge, comme il contribue dans le processus d'admission des nouveaux bearers dans le réseau. En effet, l'admission d'un bearers ne s'effectue que si la charge du réseau n'a pas dépassé une certaine valeur. Le principe de priorisation des bearers est considéré dans cette étude et la préemption se réalise selon les règles appliquées par le mécanisme ARP.

Dans (Chadchan and Akki 2011), les bearers prioritaires peuvent interrompre les bearers actifs moins prioritaires jusqu'à l'obtenir des ressources nécessaires pour atteindre le niveau minimum requis de qualité de service. Cette étude propose deux solutions, à savoir un algorithme de préemption partielle nommé PS Minimum QoS Preemption Algorithm (PS-MQPA) et le Total Preemption Algorithm (TPA). Pour chaque nouvel appel, les auteurs calculent les ressources radio totales qui peuvent être obtenues après la préemption de tous les bearers de priorités plus basses que celle du nouveau bearer, ainsi que la quantité des ressources minimales qui correspond à la quantité des ressources qui peut être obtenue après la reconfiguration de la QoS. La reconfiguration de la QoS est effectuée dans le réseau pour garantir seulement le niveau minimal de QoS acceptable pour chaque classe de trafic. Dans le cas où les ressources totales sont inférieures à la quantité des ressources minimales citée ci-dessus, alors celui-là sera rejeté. Dans le cas contraire, deux cas de figure se présentent. Dans le cas où les ressources requises sont de quantité inférieure à celle qui peut être fournie par la quantité des ressources minimales, définie ci-dessus, alors l'algorithme PS-MQPA sera exécuté sinon le deuxième algorithme sera appliqué. L'algorithme PS-MQPA consiste à obtenir des ressources en interrompant des bearers de plusieurs classes de trafic de telle façon à garder certains bearers actifs pour chaque classe de priorité plus basse que celle à laquelle

appartient le nouvel appel. Les auteurs proposent d'interrompre moins de bearers dans les classes plus prioritaires concernées par l'interruption. En d'autres termes, plus la priorité de la classe est basse, plus le nombre de bearers interrompus est important. Cette technique sert à garantir la QoS des classes de hautes priorités. Par ailleurs, dans le cas où la quantité requise est inférieure à la quantité minimale déterminée en haut, alors ce sera l'algorithme TPA qui sera exécuté en commençant par interrompre les bearers actifs de la plus basse priorité, suivis des bearers de la seconde classe moins prioritaire, ainsi de suite jusqu'à atteindre la quantité requise en termes de ressources radio.

D'autres études se penchent sur le problème de l'accès à la radio commerciale partagé dans les réseaux LTE et les LTE HetNets.

L'article (Borkar, Roberson et al. 2011) développe une solution pour fournir un accès prioritaire à la radio commerciale commune pour les clients de la sécurité publique (PS) et les usagers commerciaux (CN) sur le réseau LTE. Dans cet article, les utilisateurs PS peuvent utiliser seulement une partie de la radio partagée, même si aucun utilisateur commercial ne demande des ressources radio. Par contre, les utilisateurs commerciaux peuvent utiliser toutes les ressources radio, y compris celles réservées pour les clients PS. L'allocation des ressources PS pour les clients commerciaux ne se fait que durant l'absence des premiers. De plus, les clients commerciaux utilisant les ressources PS seront interrompus dès qu'elles sont demandées par les clients PS. Certaines conditions s'appliquent pour la préemption d'un bearer actif (AB) dans le réseau par un nouveau bearer (NB). En effet, un bearer moins prioritaire ne sera pas en mesure d'interrompre un autre plus prioritaire. De plus, un nouveau bearer doit avoir la valeur Capacity to pre-emption (CP) équivalente à Yes. Un bearer ayant la valeur CP équivalent à No, ne sera pas capable d'interrompre les bearers actifs dans le réseau. D'autre part, le bearer susceptible d'être interrompu doit être vulnérable à la préemption. Cette valeur, exprimée par VP, doit être égale à Yes.

L'approche proposée dans (Shajaiah, Abdel-Hadi et al. 2014) permet l'accès aux premiers répondants avec priorité à la radio commerciale partagée. Cette étude classe les trafics en

quatre catégories, selon qu'ils soient des trafics temps réel ou non-temps réel, ou selon qu'ils soient des trafics du réseau de la sécurité publique ou du réseau commercial. Les auteurs considèrent que le trafic temps réel est prioritaire que le trafic non temps réel au sein du même type de réseau. D'autre part, il considère que le trafic généré par les premiers répondants est plus prioritaire que le trafic commercial. Le système d'allocation de ressources proposé dans cette analyse n'attribue aucune ressource au réseau commercial quand celles-ci sont limitées. Par contre, quand la quantité des ressources disponibles est supérieure à la quantité minimale des ressources requises par le réseau de la sécurité publique, le système attribue au réseau de la sécurité publique la quantité de ressources minimale qu'il exige, ensuite il alloue le reste aux deux réseaux, commercial et sécurité publique, en respectant la priorité des trafics à acheminer. Cette technique bien qu'elle garantie la QoS requise pour le réseau de la sécurité publique, elle pénalise les usagers du réseau commercial quand ceux-ci transmettent beaucoup plus du trafic du type non temps réel.

Les auteurs de l'article (Blom, de Bruin et al. 2008) propose un concept de sécurité publique basé sur le principe de commutation de paquets autant pour le data que pour la voix. Le modèle conçu doit répondre aux exigences des premiers répondants relativement à la gestion des crises. Pour développer leur solution, les auteurs utilisent les deux technologies IMS (IP Multimedia Subsystem) et PoC (Push-to-talk over Cellular). Ils justifient ce choix par le fait que ces applications sont utilisées par la norme 3GPP. Pour cette solution deux réseaux radio sont considérés, notamment le réseau commercial et le réseau gouvernemental. Les deux réseaux combinés représentent le réseau de la sécurité publique. Le réseau gouvernemental étend la couverture du réseau commercial quand la couverture de ce dernier n'est pas suffisante. Les deux réseaux sont connectés au réseau « government home network » qui représente le réseau cœur basé sur IMS, et qui est géré par des agences de la sécurité publique. Par ailleurs, des passerelles sont utilisées pour la communication avec les autres types de réseaux, autres que ceux de la sécurité publique. Les résultats de l'analyse relative à l'implémentation de ce modèle montrent une amélioration dans la surface de couverture, dans la latence, ainsi que dans la capacité au niveau du réseau.

L'article (Tung, Lu et al. 2013) traite la problématique de gestion des ressources radio dans le Vehicular Network (Vanet). Au lieu d'utiliser un mécanisme de contrôle d'admission pour la gestion des priorités des usagers, les auteurs préfèrent l'implémentation d'un système de contrôle de congestion. Étant donné que le statut des véhicules ne pourra être connu qu'après leur admission, les usagers de hautes priorités peuvent donc être rejetés par erreur si le mécanisme de contrôle d'admission est utilisé. Deux métriques sont utilisées pour décider de la priorité d'un véhicule. La première consiste dans la distance entre le véhicule et l'intersection. Plus cette distance est petite, plus le véhicule est prioritaire. La deuxième est représentée par le temps restant pour un véhicule pour atteindre une intersection. Une valeur plus petite de ce temps correspondra à une priorité plus haute pour le véhicule. L'algorithme de contrôle de congestion utilise ces deux métriques pour classer les bearers correspondants aux différents véhicules. Lors de la réduction de la charge d'une cellule, ce sont les bearers moins prioritaires qui seront interrompus. L'interruption des bearers se fait tant que la charge tolérée n'est pas atteinte. Les résultats de la simulation montrent que cette méthode peut contrôler la charge dans les cellules LTE tout en assurant les services d'urgence.

Une solution basée sur l'utilisation de la radio cognitive est proposée par les auteurs de l'article (Al-Hourani and Kandeepan 2013). L'objectif de cette approche est d'améliorer la couverture dans le réseau régulier de la sécurité publique sur LTE, ainsi que de réduire le niveau d'interférences. Les petites cellules sont représentées par des clusters. Chaque cluster a un cluster head. Le rôle de cluster head consiste à acheminer le trafic allant vers et arrivant du réseau cellulaire, considéré comme le réseau parent, ainsi que de gérer les communications entre les membres du son cluster afin de réduire la signalisation. Les échanges intra-cluster s'effectuent via des communications D2D. Entre autres, un cluster est activé dans le cas où le niveau de SINR est élevé pour un ensemble d'utilisateur UEs. Ceux-ci vont former un cluster, choisir un cluster head et communiquer en utilisant la radio cognitive via des transmissions D2D. Par ailleurs, un UE dont le niveau de SINR est plus élevé que la valeur tolérée va chercher de rejoindre un cluster déjà établi. D'autre part, la désactivation d'un cluster est effectuée pour maintes causes, entre autres, après que les nœuds du cluster aient

un niveau acceptable du SINR dans le réseau parent, ou quand le cluster head perd sa capacité pour assumer son rôle.

3.4 Gestion des ressources de bande passante dans les réseaux Backhaul LTE

Tel que nous l'avons mentionné dans le chapitre 1, le réseau Backhaul LTE est le réseau d'agrégation responsable du regroupement du trafic généré par plusieurs stations eNodeB reliées entre elles. Il représente le lien de connexion entre les nœuds eNodeB et le réseau cœur EPC. Il existe différentes technologies qui permettent le déploiement d'un réseau Backhaul LTE. Entre autres, on distingue la fibre optique et le réseau hertzien, ainsi que la technologie MPLS. La technologie MPLS étant la plus rentable est devenue la meilleure solution pour l'agrégation des débits en provenance du réseau d'accès et allant vers les réseaux cœur LTE. L'intégration de la technologie MPLS dans le réseau Backhaul LTE est prometteuse. Cette technologie permet l'amélioration de la QoS de bout en bout dans les réseaux LTE, comme elle assure de nouveaux services mobiles.

Il n'existe malheureusement pas beaucoup d'études qui traitent le sujet d'intégration de MPLS dans le réseau Backhaul LTE. Alcatel-Lucent (Lucent 2014) offre une solution pour l'intégration de IP/MPLS dans le réseau d'agrégation LTE et relate les bénéfices d'une telle technologie en termes de performance du réseau Backhaul et aussi de mise à l'échelle. De sa par Cisco, via son document (Cisco 2011), confirme que la technologie MPLS a déjà été utilisée avec succès pour les réseaux cœurs et elle peut être réutilisée pour les réseaux d'agrégation. Cela peut apporter une belle amélioration dans le contrôle du trafic et la gestion des ressources, ainsi que d'assurer les opérations, l'administration et la maintenance (OAM) dans ces réseaux.

L'article (Venkatesan and Kulkarni 2008) affirme que le protocole MPLS supporte le routage P2MP (Point to Multi-Point), donc un seul tunnel P2MP déployé entre l'eNodeB et l'EPC suffira pour permettre d'effectuer des routages Unicast et Multicast à la fois.

Les auteurs de l'étude (Nurcahyaningsih, Munadi et al. 2014) proposent l'utilisation de la technologie MPLS-VPN dans les réseaux Backhaul LTE pour améliorer leur performance en terme de délai, débit, gigue et perte de paquets. La raison pour laquelle ils font ce choix est le fait que la technologie MPLS-VPN supporte plusieurs protocoles, permet l'application de l'ingénierie du trafic donc assure le Load Balancing. De plus, cette technologie permet la mise à l'échelle. Les résultats de la simulation expriment une amélioration dans la QoS du réseau Backhaul dans le cas de l'utilisation de MPLS-VPN par rapport au cas où cette technologie n'est pas appliquée.

Tenant en compte les avantages d'utiliser la technologie MPLS, les auteurs de l'article (Mo, Yuan et al. 2015) propose une solution MPLS pour les réseaux Backhaul. Cette solution n'est pas spécifique à LTE. Le réseau Backhaul considéré dans cette approche peut regrouper des trafics en provenance de plusieurs types de réseaux d'accès radio, tel que GSM et LTE. L'architecture du réseau proposé dans cette étude contient des liens de type xPON (X version of Passive Optical Network) à l'entrée et à la sortie du réseau Backhaul, et des liens de types Ethernet (couche 2) à l'intérieur. Les auteurs de cet article proposent un nouveau modèle d'intégration de la technologie MPLS dans les réseaux Backhaul garantissant la mise à l'échelle de ces réseaux en assurant d'offrir plusieurs types de services. Ceci est possible en utilisant le champ CoS (Classe of Service) de la couche 2 pour déterminer la priorité des paquets au niveau du réseau Ethernet et d'utiliser le champ EXP (experimental bits) de MPLS pour spécifier la priorité du trafic à l'entrée du réseau Backhaul. De plus, les outils OAM (Operation, Administration end Maintenance) peuvent être utilisés grâce au déploiement de la solution MPLS-TP (Transport Profil MPLS) qui supporte ce mécanisme.

3.5 Communications Device to Device dans LTE Hétérogène

La technologie Device-to-Device (D2D) vient ajouter de la performance au réseau LTE en permettant aux usagers LTE (UEs) de communiquer entre eux d'une façon directe et sans avoir à passer par l'eNodeB. Les communications D2D sont utilisées pour diverses raisons, à savoir l'extension de la couverture d'une macro cellule LTE, l'offloading des cellules

congestionnées, la réduction de la consommation d'énergie et l'amélioration de la QoS dans le réseau en termes de réduction des délais et des pertes de paquets, ainsi qu'en termes d'augmentation de throughput et du Packet Delivery Ratio. Dans cette section, nous présentons quelques études relatives à la technologie D2D et qui traitent les causes citées ci-dessus.

3.5.1 Intégration du D2D dans l'architecture LTE

L'article (Raghothaman, Deng et al. 2013) propose une solution pour l'intégration de la fonctionnalité D2D dans l'architecture existante du réseau cœur LTE. Les auteurs proposent l'ajout d'un nouvel équipement, nommé D2D Server dont le rôle est de fournir, maintenir, et sauvegarder les identifiants des usagers D2D. Une demande de communication D2D pour l'obtention d'un service donné, qui est formulée par un usager UE_i , est transmise au D2D Server via l'équipement MME du réseau cœur EPC. Suite à la réception de la requête, D2D Server va solliciter le PCEF afin de vérifier le droit de l'utilisateur UE_i d'utiliser le service demandé. Dans le cas où il a le droit de l'utiliser, alors D2D Server va lui fournir un identifiant, autrement, sa requête sera rejetée. Cet article montre les différents raccordements du serveur D2D server avec les autres équipements déjà implémentés dans EPC. Ces différentes connexions permettent à D2D Server de coopérer avec les autres éléments de EPC afin d'accomplir certaines tâches, telles que la découvertes des usagers D2D, la gestion de la mobilité de ces usagers ainsi que l'établissement des appels D2D.

3.5.2 D2D pour le processus d'offloading

Le partage du spectre radio commercial LTE pour y permettre l'accès aux clients PS rend la gestion des ressources radio efficace et rentable. Toutefois, une fois que toutes les ressources sont allouées, et qu'aucun bearer actif dans le réseau ne puisse être interrompu, le niveau de qualité de service des différents trafics sera rétrogradé. Pour résoudre ce problème, Le mécanisme d'offloading des macros cellules vers les petites cellules ou d'autres macros cellules LTE est appliqué. Notons que les communications Device-to-Device (D2D) naissent suite à ce phénomène. Les communications D2D se caractérisent par l'échange des données

directement entre les équipements sans passer par l'eNodeB. Dans ce qui suit, certaines études récentes sont explorées relativement à la technologie D2D.

L'étude (Pyattaev, Johnsson et al. 2013) traite le basculement des bearers d'une macro cellule LTE vers les cellules du réseau WiFi. Cette approche considère le cas où le processus de découverte de la localisation des petites cellules est appliqué par le réseau cellulaire, afin d'optimiser le nombre de connexions Device-to-Device. Les auteurs montrent que leur modèle peut fournir une amélioration significative de la capacité et de la consommation d'énergie.

Les auteurs de (Liu, Kawamoto et al. 2014) proposent un algorithme pour l'offloading des bearers arrivant vers une cellule LTE eNodeB congestionnée. La particularité de cette étude est l'utilisation de la technique D2D pour le basculement des utilisateurs UEs de la macro cellule vers d'autres macros cellules ou vers des picos ou femtos cellules. Cette technique permet à un usager de communiquer à la fois sur le réseau cellulaire et le réseau D2D. De plus, un usager appartenant à une macro cellule donnée peut communiquer avec un usager D2D appartenant à une autre cellule adjacente si celui-ci est localisé dans son rayon de couverture, cet usager voisin deviendra donc un nœud relai pour le premier utilisateur. Par conséquent, l'utilisateur appartenant à la cellule d'origine pourra communiquer avec l'eNodeB de la cellule adjacente via le relai D2D. Par ailleurs, L'offloading proposé dans ce travail offre la possibilité à une cellule cible d'être déchargée d'un ou de plusieurs UEs pour pouvoir intégrer un usager basculé de la cellule d'origine. En outre, le basculement des usagers peut correspondre à plusieurs cellules à la fois, ce qui va assurer un partage de la charge dans le réseau (Load Balancing). Chaque usager basculé communiquera avec la station de base de la cellule cible via les relais D2D.

3.5.3 D2D pour l'amélioration de la QoS et la consommation d'énergie

En plus d'assurer le processus d'offloading du réseau cellulaire, les communications D2D peuvent apporter une amélioration dans la qualité de service des réseaux cellulaires. Elles

peuvent aussi réduire la consommation de l'énergie par les nœuds du réseau. Dans ce qui suit, quelques approches proposées dans la littérature sont présentées.

La proposition faite par les auteurs de l'article (Le, Keller et al. 2014) consiste en un modèle où plusieurs usagers peuvent coopérer pour le téléchargement d'une même vidéo. Habituellement, le téléchargement d'un même fichier par plusieurs utilisateurs s'effectue un nombre de fois égale au nombre des usagers qui le font. La contribution de cet article consiste dans la réduction du nombre de téléchargements du fichier en question à un seul téléchargement. Cela est concrétisé via l'utilisation des communications D2D. Pour ce faire, les usagers adjacents voulant télécharger un même fichier vont s'organiser pour former un réseau WiFi. Un seul nœud parmi eux va télécharger le fichier voulu en utilisant le lien radio cellulaire. Ensuite le fichier sera diffusé dans le réseau WiFi et tous les utilisateurs pourront en bénéficier. Les résultats de la simulation montrent que cette approche améliore la performance dans le réseau cellulaire sans affecter pour autant l'utilisation des batteries des terminaux des utilisateurs.

D'autre part, Les auteurs de l'article (Feng, Lu et al. 2013) traitent la problématique des interférences qui peuvent être générées avec les communications D2D utilisant les ressources radio du réseau cellulaire. Ils proposent entre autres un nouveau modèle de contrôle d'admission des paires D2D au sein d'un réseau cellulaire. Les auteurs ont conçu leur modèle de telle sorte que les usagers D2D partagent les ressources d'un Cellular User (CU). En effet, l'admission d'une paire d'utilisateurs D2D n'est possible que si les conditions suivantes sont vérifiées :

- la valeur du Signal to Interference plus Noise Ratio (SINR) des usagers D2D, ainsi que celle du CU correspondant, sont supérieures ou égales au SINR minimal requis pour chacun d'eux;
- la puissance de transmission de chacun des nœuds cités ci dessus doit être inférieure ou égale à la puissance de transmission permise pour chacun d'eux.

Cette méthode permet le partage des ressources entre les usagers D2D et les usagers CU tout en assurant un niveau acceptable d'interférence.

L'approche (Yaacoub and Kubbar 2012) propose l'utilisation des communications D2D pour améliorer les communications du réseau de la sécurité publique sur le réseau LTE. L'objectif est de minimiser la consommation d'énergie au niveau des terminaux. Pour ce faire, les nœuds forment une coalition. Chaque coalition est représentée sous forme d'un cluster. Un cluster head est choisi dans chaque cluster pour communiquer avec le réseau cellulaire. Par la suite, les utilisateurs de certains clusters seront rassemblés ensemble afin de minimiser la consommation d'énergie. Pour ce faire, les auteurs proposent de sélectionner un premier cluster ayant un niveau de consommation d'énergie élevé, soit C_k ce cluster, la prochaine étape consiste à trouver un autre cluster C_i , telle que la somme de consommation d'énergie des deux clusters ensemble soit inférieure à la somme de la consommation de l'énergie de chaque cluster individuellement. Autrement dit,

$$E_{C_k \cup C_i} < E_{C_k} + E_{C_i} \quad (3.1)$$

Avec E_{C_j} est l'énergie consommée par le cluster j , $j = i$ ou k . Dans le cas où C_i existe alors, les membres de C_i seront fusionnés au cluster C_k et le cluster C_i sera supprimé. Dans le cas contraire, la fusion ne s'effectuera pas et les éléments de C_i ne vont plus participer au processus de sélection de nœuds à fusionner avec les nœuds du cluster C_k . Les résultats de la simulation montrent que le système de coalition proposé réduit la consommation de l'énergie totale comparant avec la méthode conventionnelle n'appliquant pas le principe de coalition. Les résultats démontrent aussi une réduction du délai de transmission et une amélioration de la qualité de la vidéo.

Tout comme la solution donnée par (Yaacoub and Kubbar 2012), l'article (Asadi and Mancuso 2013) propose d'utiliser la technique de clustering pour acheminer les communications D2D. Les nœuds D2D sont regroupés dans des clusters WiFi en utilisant leur deuxième interface réseau, de type WiFi. La formation des clusters ne s'effectue que si la communication au niveau cellulaire est jugée être de bonne qualité. Un Cluster Head (CH)

est déterminé dans chaque cluster de façon à ce qu'il contienne la meilleure qualité du lien parmi les membres de son cluster. Le CH est responsable de la communication avec l'eNodeB pour lui transmettre le trafic agrégé des membres de son cluster. Tous les clusters formés doivent s'enregistrer au niveau du réseau LTE via deux procédures, à savoir la procédure Cluster Notification et la procédure Cluster Verification. La première procédure consiste à informer l'eNodeB de la formation du cluster et la deuxième permet à l'eNodeB de vérifier que tous les clients d'un cluster font vraiment partie de ce cluster, autrement dit ils sont des membres de ce cluster. La vérification effectuée par l'eNodeB se réalise via une requête de sécurité envoyée à tous les clients du cluster en question. Dans le cas où le CH ne pourra plus jouer son rôle de Cluster Head à cause de la détérioration de sa qualité du lien, alors un autre membre du cluster sera désigné pour prendre sa place. Ce mécanisme améliore le throughput dans le réseau et réduit la consommation d'énergie, ainsi que le délai. Toutefois, cette étude ne prend pas en considération la différenciation de services, étant donné que les communications D2D peuvent acheminer du trafic temps réel comme elles peuvent transmettre des trafics non temps réel.

La solution proposée dans l'article (Wang, Wei et al. 2013) a pour objectif d'améliorer les communications cellulaires en utilisant la technologie D2D. L'algorithme développé par les auteurs se base sur la coopération des nœuds libres du réseau cellulaire (UEs libres) avec les nœuds UEs actifs afin d'acheminer les données de ces derniers. Les UEs libres se joignent aux UEs actifs afin de former un réseau multi-sauts pour pouvoir effectuer des communications D2D entre eux en utilisant les fréquences de licences privées. Les utilisateurs UEs libres jouent le rôle de relais pour les UEs actifs. L'allocation des ressources s'effectue de façon équitable entre les UEs opérant en D2D et ceux opérant dans le réseau cellulaire. L'algorithme Round Robin (RR) est utilisé pour la gestion de ces ressources afin d'assurer l'équité de leur allocation. Les résultats de la simulation démontrent que l'utilisation de la solution proposée dans cet article améliore la performance des communications cellulaires en termes de valeur de SINR. Bien que les résultats obtenus sont encourageants, mais les auteurs ne considèrent pas la priorisation des trafics dans leur étude.

Rappelons que les exigences en termes de QoS diffèrent d'une classe de trafic à une autre, donc l'intégration de la priorisation des trafics dans une telle approche est nécessaire.

L'approche des auteurs de l'article (Chen, Zhao et al. 2013) consiste à partager les ressources radio entre des utilisateurs du réseau cellulaire (CU) et ceux du réseau D2D en réduisant les interférences. La démarche de cette solution tient à assigner pour chaque usager D2D un CU. Cette approche propose le partage des ressources avec les CUs en suivant l'une des méthodes ci-dessous :

1. Les usagers D2D doivent partager les ressources d'un CU assez loin afin de garantir un niveau acceptable du Signal-to-Interference-plus-Noise Ratio (SINR);
2. Un utilisateur D2D peut utiliser les ressources de plusieurs CUs tant qu'il n'obtient pas les ressources requises pour atteindre le niveau de QoS demandé;
3. Les ressources d'un CU sont divisées en blocs (Blocs de ressources ou RB). Un usager D2D peut utiliser plusieurs RBs d'un même CU. Les RBs utilisés par D2D doivent être consécutifs. Les RBs sont choisis de telle sorte à ce qu'ils contribuent à offrir une valeur SINR inférieure au seuil SINR qui est toléré;
4. Utiliser le mécanisme Orthogonal Frequency-Division Multiple Access (OFDMA) afin de permettre aux utilisateurs D2D de partager des RBs non consécutifs au sein des ressources d'un même CU.

Les résultats de la simulation montrent que la quatrième méthode, notamment celle qui utilise OFDMA, offre un meilleur niveau de QoS en termes de valeur de SINR.

3.6 Le codage réseau dans les réseaux sans fil

Dans cette section, quelques solutions pertinentes du network coding seront présentées. Elles permettent de mieux comprendre son fonctionnement et d'illustrer son bénéfice.

L'article (Katti, Rahul et al. 2008) présente une implantation du codage réseau dans un réseau mesh sans fil. L'étude représente un ensemble de technique à appliquer et la règle pertinente pour la constitution d'une combinaison de symboles codés. Cette solution, nommée COPE, utilise trois techniques pour l'établissement d'un codage réseau, à savoir Opportunistic Listening, Opportunistic Coding et Learning Neighbor State. Le principe de la technique « Opportunistic Listening » consiste à ce que tous les nœuds du réseau écoutent toutes les communications et stockent les paquets pendant une période T ($T = 0.5$ ms). Chaque nœud diffuse un rapport de réception afin d'informer ses voisins des paquets qu'il a reçus et qu'ils stockent à son niveau. Par ailleurs, un nœud ne possédant pas de paquets à transmettre, diffuse son rapport de réception dans des paquets de contrôle spéciaux. La technique « Opportunistic Coding » permet à chaque nœud voisin appelé NextHop de décoder le paquet reçu pour extraire le paquet natif. Pour ce faire, deux files d'attente sont utilisées, à savoir le Packet Pool et la file Output Queue. Le Packet Pool est un buffer servant au stockage des paquets entendus pendant la période T , cité ci-dessus. Output Queue est la file d'attente des paquets à diffuser.

Une règle spécifique doit être vérifiée afin de transmettre n paquets $P_1 \dots P_n$ à n nœuds voisins NextHop $R_1 \dots R_n$. Notamment, un nœud K_i réalise une opération de XOR de n paquets ensembles et met le message résultant dans la file Output Queue, seulement si chaque NextHop dispose dans son Packet Pool, de tous les $(n-1)$ paquets P_j , tel que $j \neq i$. Chaque nœud choisit une valeur de n la plus large possible afin de tirer un maximum de profit du codage réseau, tout en respectant la contrainte citée ci-dessus. La troisième technique utilisée pour cette solution est nommée Learning Neighbor Stat. En effet, le rapport de réception sert à informer les nœuds voisins d'un nœud donné de sa liste des paquets stockés dans la file Output Queue. Or, ces rapports peuvent être perdus à cause d'une congestion ou arriver aussi tard que le nœud voisin ait déjà pris sa décision de codage. À cet effet, chaque nœud doit pouvoir deviner les paquets résidant dans la file Output Queue de ses voisins. Une solution est de considérer la probabilité de délivrance de paquets entre deux nœuds x et y . La métrique ETX (Expected Transmission Count) peut être utilisée pour cette fin. Autrement dit, cette approche estime la probabilité qu'un nœud dispose d'un paquet à

son niveau comme la probabilité de délivrance du lien reliant le nœud et son voisin possédant le paquet. Dans le cas où l'estimation est erronée, le paquet ne sera pas décodé. Par conséquent, le paquet natif sera codé avec une nouvelle combinaison de paquets et retransmis par la suite dans un nouveau paquet codé. Les résultats de la simulation montrent que l'algorithme COPE améliore la performance du réseau mesh WMN en termes de Throughput.

La proposition de l'article (Fang-Chun, Kun et al. 2009) présente, entre autres, une solution pour améliorer la gestion des pertes de paquets dans les réseaux sans fil. Le principe est d'utiliser le codage réseau lors de la retransmission des paquets perdus. En effet, un nœud source va procéder à la transmission des données selon la méthode classique, donc sans codage réseau. Dans le cas où plusieurs nœuds ne reçoivent pas leurs paquets, la source va procéder à une retransmission des paquets perdus. Seulement, au lieu d'envoyer chaque paquet à part, elle transmet une combinaison codée de tous les symboles perdus. De cette manière le temps de retransmission va se réduire d'une façon remarquable.

L'article (Peng, Song et al. 2014) est une nouvelle approche de routage basée sur le codage réseau linéaire et le mécanisme « Fault Tolerance Routing ». Son objectif est l'amélioration du « Packet Delivery Ratio », du délai de bout en bout, ainsi que du throughput dans les réseaux WMN. Les auteurs proposent un algorithme dont les étapes sont les suivantes :

1. construction d'un multipath entre la source S et la destination D. Trois chemins, au minimum doivent être définis;
2. soit N le nombre de paquet envoyé et $i = N/k$, avec $k \leq N$. On note par k le nombre de paquets à envoyer simultanément;
3. les k paquets sont codés linéairement suivant la technique de codage réseau linéaire. Les codes utilisés pour le codage réseau sont inclus dans l'entête du paquet codé;
4. le paquet généré par codage réseau est transmis via toutes les routes du multipath;
5. à la réception des codes, la destination utilise les codes inclus dans l'entête du paquet reçu pour générer la matrice de codage. Par la suite, elle détermine le rang de la matrice

de codage générée. Dans le cas où ce rang est inférieur à k , elle constatera qu'un ou plusieurs paquets sont perdus. Elle transmettra donc une requête de retransmission de paquets. Dans le cas contraire, la destination procèdera au décodage des paquets reçus afin d'en extraire les natifs. Le décodage des paquets s'effectue en groupe de k paquets à la fois jusqu'à atteindre le nombre de N paquets.

La solution de codage réseau proposée dans l'article (Yang, Ling et al. 2012) est conçue pour les réseaux Manet. Basé sur l'algorithme AOMDV pour la construction des chemins du multipath, son objectif est d'appliquer le mécanisme du Load Balancing lors de la transmission des données. Cette solution (NC-AOMDV), comparée avec AOMDV, montre une amélioration de la QoS dans Manet à travers la réduction du trafic de contrôle et l'augmentation du taux de paquets délivrés (Packet Delivery Ratio). L'approche adoptée par NC-AOMDV consiste d'abord à construire un multipath entre le nœud source et le nœud destination en utilisant l'algorithme AOMDV. Dans le cas où le multipath n'existe pas, alors le NC-AOMDV ne sera pas mis en œuvre. Autrement, le processus de codage réseau s'initie à travers le codage des symboles natifs afin de générer le paquet codé à envoyer. Les auteurs utilisent le codage réseau linéaire comme type de codage réseau et ils incluent les codes dans le paquet envoyé. Finalement, la destination utilise les codes inclus dans chaque paquet reçu pour générer la matrice de codage inverse. Cette matrice est utilisée dans le processus de décodage.

Dans l'article (Júnior, Vieira et al. 2014) on propose un nouveau mécanisme pour la transmission de données, nommé CodeDrip. CodeDrip est un protocole de diffusion de données dans les réseaux capteurs sans fil (WSN). Il utilise le codage réseau lors de la transmission des paquets afin d'améliorer la consommation d'énergie, ainsi que la fiabilité et la vitesse de diffusion. La perte de paquets est réduite car les paquets perdus peuvent être récupérés par le décodage des autres paquets codés incluant ces paquets perdus. En évitant la retransmission, CodeDrip arrive à minimiser le nombre de diffusion.

3.7 Sécurité dans les réseaux locaux sans fil

De nombreuses études sont réalisées pour améliorer la sécurité dans les réseaux sans fil et à contourner les attaques de la confidentialité, de l'intégrité et de la disponibilité. L'une des solutions adaptées pour assurer une communication sécurisée dans les réseaux sans fil est d'utiliser le mécanisme de chiffrement (cryptage), plusieurs études proposent des solutions de cryptage afin de sécuriser les transmissions (Shah, Rashmi et al. 2013, Tang 2013, Zhao, Kent et al. 2013, Zhou 2013). Cependant, il n'est pas possible d'éviter le trafic de contrôle supplémentaire généré par les différentes opérations mathématiques complexes effectuées par les algorithmes de cryptage. Par conséquent, le niveau de la qualité de service (QoS) va certainement se détériorer en raison des usages complémentaires des ressources radio et de bande passante. Or, les réseaux WMN sont connus pour utiliser des transmissions de redondance. En fait, les ressources de réseau deviennent vite limitées.

D'autre part, le mécanisme d'authentification se montre comme une solution inhérente à empêcher les nœuds non autorisés à accéder au réseau. De nombreuses solutions sont proposées pour assurer le processus d'authentification. Transport Layer Security (TLS) (Fischer and Rödig 2000), Extensible Authentication Protocol (EAP) (Aboba, Blunk et al. 2004), l'authentification, l'autorisation et de comptabilité (AAA) (Metz 1999), Message Digest 5 (MD5) (Rivest 1992) sont tous des mécanismes proposés dans la littérature pour éviter l'accès non autorisé au réseau pour certains utilisateurs. Notons qu'en plus de l'ajout de trafic supplémentaire qui est généré par l'application de ces algorithmes et par les informations échangées entre les différents nœuds dans le réseau, ces solutions ne sont pas en mesure de surmonter les attaques internes de disponibilités.

D'autres propositions sont faites par (Kent and Liebrock 2011, Sen, Koilakonda et al. 2011, Shah and Valiveti 2012, Prasad and Giri 2014) pour améliorer le niveau de sécurité dans les réseaux sans fil, mais chacune d'entre elles, génère du trafic supplémentaire dans le réseau. Rappelons que lors des situations de crises les ressources du réseau s'épuisent rapidement, à cause de l'augmentation spectaculaire du nombre d'appels établi. De ce fait, une solution de

sécurité qui engendre un trafic supplémentaire ne sera pas adéquate dans de tels cas. Il serait pertinent de développer une nouvelle approche qui offre un bon niveau de sécurité lors de la transmission des données et qui surtout garantit un niveau de QoS acceptable dans le réseau.

3.8 Conclusion

L'exploration de la revue de la littérature a ressorti l'importance de l'utilisation de la technologie LTE pour améliorer la performance dans les réseaux de la sécurité publique. D'une part, le partage de la radio commerciale avec le réseau PS et d'autre part l'utilisation du mécanisme d'offloading, ne font que rendre plus efficace la gestion des ressources en évitant les situations de congestion dans le réseau. Cependant, toutes les solutions proposées dans la littérature et présentées dans ce chapitre mettent l'emphasis sur la performance des communications prioritaires. Or, même les flux de basses priorités peuvent être importants et peuvent avoir besoin d'être acheminés efficacement et à temps. À titre d'exemple, des usagers commerciaux ont besoin de communiquer avec leurs proches et leurs amis lors des situations de crises afin d'avoir de leurs nouvelles. Dans un tel cas de figure, étant donné que le trafic commercial est moins prioritaire que le trafic de la sécurité publique, plusieurs utilisateurs commerciaux se voient refuser l'accès aux ressources radio à cause de la congestion. Rappelons que durant les désastres, le nombre d'appels augmente d'une façon fulgurante. De plus, une partie non négligeable de ces appels est établie par les premiers répondants. Ces derniers étant généralement déployés dans de petites aires, représentant la surface de l'emplacement de la crise, peuvent donc accaparer toutes les ressources radio. Par conséquent, toutes les solutions proposées ci-dessus ne vont pas être bénéfiques pour les usagers commerciaux.

D'autre part, les approches proposées pour effectuer des communications D2D en partageant les ressources d'un usager LTE garantissent un bon niveau de QoS pour les usagers D2D, ainsi que pour le CU, quand les ressources sont disponibles. De plus, elles permettent de réduire le niveau d'interférences entre les paires D2D et les CU correspondants. En revanche, comme elles utilisent les fréquences du réseau LTE, donc elles ne seront pas en mesure de fournir les ressources radio manquantes lors des moments de congestions, où les ressources

sont limitées. L'utilisation des fréquences publiques se montre donc pertinente lorsque les ressources LTE sont limitées ou épuisées. Bien que les réseaux locaux sans fil n'offrent pas le même niveau de performance que les réseaux cellulaires, mais ils peuvent quand même fournir des ressources radio supplémentaires, qui peuvent garantir un certain niveau de QoS s'ils sont dotés de bons algorithmes de routage et de sécurité.

Le codage réseau est l'une des nouvelles approches de transmission de données, qui peuvent améliorer la performance dans les réseaux. Les solutions de codage réseau illustrées ci-dessus démontrent chacune une amélioration remarquable dans la performance des réseaux dans lesquelles elles sont appliquées. De plus, ce sont des solutions qui permettent de réussir le décodage de l'information reçue. Toutefois, il est clair que les codes utilisés dans le processus de codage-décodage de l'information sont accessibles par tout nœud recevant le paquet codé à cause de l'inclusion des codes dans le paquet transmis. Ceci dit, aucune de ces solutions ne s'avère pertinente pour application lors de l'échange des données D2D dans les petites cellules LTE. En effet, les transmissions dans les réseaux sans fil sont basées sur le principe du Broadcast, ainsi la présence d'un nœud malveillant dans le rayon de transmission des nœuds légitimes du réseau sans fil augmente le risque d'attaques de sécurité. Cependant, les communications D2D sont nombreuses à appartenir au réseau de la sécurité publique, donc elles doivent être qualifiées d'importantes et de confidentielles. Par conséquent, des attaques de sécurité qui atteignent la confidentialité, l'intégrité et la disponibilité de l'information peuvent nuire d'une façon remarquable à la performance des communications D2D dans les petites cellules. De ce fait, un nouveau modèle de codage réseau doit être développé afin de sécuriser les communications D2D des réseaux LTE Hétérogènes.

Quant aux solutions de sécurité explorées dans cette thèse, elles offrent toutes des solutions efficaces pour sécuriser les réseaux sans fils des différents types d'attaques de sécurité. En revanche, il est important de se rappeler que lors des moments de crises, les ressources s'épuisent très rapidement, donc toute épargne de ressources est utile. Or, les approches de sécurité présentées dans ce chapitre utilisent toutes des ressources additionnelles et

augmentent toutes la taille de trafic de contrôle. Cela va à l'encontre du principe d'économiser l'utilisation des ressources radio lors des moments de crises.

N'oublions pas que l'amélioration de la performance dans les réseaux LTE Hétérogène ne se limite pas dans ses deux parties représentées par le réseau d'accès et le réseau cœur. Le réseau Backhaul LTE doit aussi être perfectionné afin d'assurer une bonne QoS de bout en bout pour tous les trafics LTE. Les études citées dans cette thèse relatent l'importance et le bénéfice d'intégration de MPLS dans le réseau Backhaul LTE.

Un nouveau modèle novateur doit être conçu pour améliorer la performance dans les réseaux de la sécurité publique sur les réseaux LTE hétérogènes. Ce modèle ne doit toutefois pas négliger la performance des réseaux commerciaux. Les chapitres 4, 5 et 6 présentent notre nouvelle approche pour construire ce modèle. Le chapitre 4 développe une nouvelle approche pour l'allocation des ressources de bande passante avec contraintes dans le réseau Backhaul et le réseau cœur LTE. Le chapitre 5 présente un nouveau modèle pour la gestion efficace des ressources radio dans LTE HetNet et le chapitre 6 modélise une nouvelle solution pour l'amélioration du routage et de la sécurité dans les petites cellules LTE HetNets.

CHAPITRE 4

Gestion des ressources dans le réseau cœur et dans le Backhaul LTE

4.1 Introduction

L'essor des applications multimédia et mobile, ainsi que la variation des services offerts en télécommunications, tels que la voix sur IP et la vidéo-conférence, ont augmenté la demande en ressources de bande passante d'une façon remarquable. Une telle situation est fortement présente dans les réseaux cœurs et les réseaux Backhaul des réseaux cellulaires, notamment dans les réseaux LTE. Un phénomène qui ne cesse de compliquer la gestion des ressources de bande passante, encore plus durant les moments de congestion. À cet effet, de nouveaux systèmes de gestion de bande passante doivent être développés. Ces systèmes auront pour rôle d'empêcher l'accaparement des ressources par les classes de trafics de haute priorité. Cela va garantir un certain niveau de QoS pour les trafics de basses priorités.

La nouvelle tendance pour les réseaux Backhaul LTE tend vers l'utilisation de la technologie Multi-Protocol Label Switching (MPLS). Cette pratique permet de rendre plus efficace l'agrégation des trafics en provenance de différents eNodeB, et allant vers le réseau cœur LTE.

Par ailleurs, La différence qui existe entre les exigences de chaque classe de trafic transmis dans un réseau cœur LTE, nécessite d'instaurer des politiques de gestion de ressources basées sur des contraintes d'allocation de ces ressources. Les contraintes d'allocation des ressources permettent un contrôle continu des quantités de bande passante attribuées à chaque flux de données. Aucune classe ne pourra donc maintenir plus de ressources que la quantité qui lui est permise et réservée initialement par le système de gestion des ressources. L'objectif est de garantir un bon niveau de QoS autant pour les classes de trafics prioritaires que pour les classes de trafic moins prioritaires.

En outre, l'amélioration de la performance du réseau cœur et du réseau Backhaul va permettre l'augmentation du nombre de bearers EPS admis dans le réseau LTE. Ces bearers peuvent être de type commercial, comme ils peuvent être de type PS. La croissance du nombre des bearers PS va jouer un rôle important dans l'amélioration de la qualité des communications établies par les premiers répondants des réseaux de la Sécurité Publique. Surtout lors des moments de gestion de crises, où toute information est nécessaire pour sauver des vies

Dans ce chapitre, une nouvelle approche, Courteous Allocation bandwidth constraints Model (CAM), a été développée, pour l'allocation des ressources de bande passante au sein du réseau Backhaul et du réseau cœur LTE. L'algorithme CAM représente un modèle d'allocation de bande passante avec contraintes dans un réseau MPLS-Traffic Engineering-DiffServ (DS-TE). CAM améliore l'allocation des ressources pour les classes de trafics moins prioritaires, tels que le trafic FTP. Sans affecter pour autant la QoS des trafics de haute priorité. CAM apporte une amélioration de la QoS pour l'ensemble des trafics moins prioritaires par rapport aux deux approches MAM (Le Faucheur 2005) et RDM (Le Faucheur 2005). La solution développée dans ce chapitre est tirée principalement de notre article (Tata and Kadoch 2013).

4.2 Modèles d'allocation de bande passante avec contraintes

MPLS Traffic Engineering (MPLS-TE) assure la création des routes de bout en bout au sein d'un réseau filaire, avec une réservation de bande passante (Adami, Callegari et al.). Le MPLS-TE combiné à DiffServ (DS-TS), peut fournir un LSP (Label Switched Path) pour chaque classe de service à travers la génération d'une relation (DiffServ, LSP), en d'autres termes, chaque traffic DiffServ sera affecté à un LSP.

DS-TE peut traiter les trafics selon leur classe d'appartenance (Din, Hakimie et al. 2007). En effet, la contribution principale du DS-TE est la modélisation d'un système d'allocation de bande passante avec contraintes (BC) qui décrit comment est allouée la bande passante pour

les différentes classes de trafic (CT) (Xu, Liu et al. 2006). Notons que le nombre maximal des CT_s défini par DS-TE, est de l'ordre de huit classes de trafics allant de CT₀ à CT₇, en affectant le trafic Best Effort (BE) à la classe CT₀.

DS-TE introduit les concepts suivants pour le modèle BC.

Class-Type (CT) est un groupe de classes de trafic basé sur leur exigence en termes de QoS. En fait, chaque CT partage une même quantité de bande passante réservée. De plus, chaque CT peut regrouper une ou plusieurs classes de service (Molnar and Vlcek 2009).

Bandwidth Constraint (BC) est la limite en pourcentage de la bande passante d'un lien qu'un CT peut utiliser (Molnar and Vlcek 2009). Plusieurs modèles BC sont proposés dans la littérature pour les réseaux DS-TE. Les deux modèles de base sont Maximum Allocation Model (MAM) and the Russia Dolls allocation Model (RDM).

Dans ce qui suit, les modèles MAM et RDM sont présentés, ainsi que notre solution CAM.

4.2.1 Maximum Allocation Model

Maximum Allocation Model (MAM) est le premier modèle conçu pour l'allocation de bande passante avec contraintes où au plus huit CTs peuvent être actives dans le réseau. MAM fonctionne de telle sorte à assigner une bande passante maximale pour chaque CT.

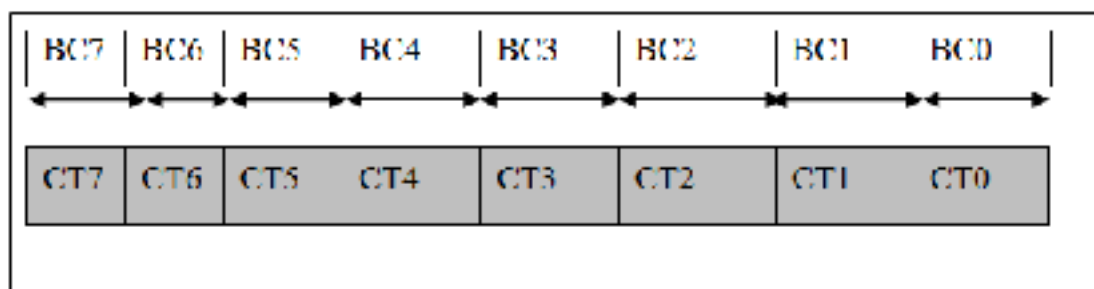


Figure 4.1 Bandwidth allocation in MAM

La figure 4.1 représente l'architecture MAM. Pour ce système, la bande passante réservée pour un CT_i , où $i = 0$ à 7 , ne pourra être utilisée par un autre CT_j , $j \neq i$, d'où l'inconvénient de cette approche. Bien que MAM garantisse l'allocation des ressources pour chaque CT, il ne gère pas efficacement les portions de bande passante non utilisées, autrement dit, certaines ressources seront sous-utilisées sans aucune possibilité qu'elles soient redistribuées.

Le modèle analytique de contraintes représentant la solution MAM, est donné comme suit :

1. Pour chaque $i \in [0, K-1]$

$$N_i \leq BC_i \leq M \quad (4.1)$$

Où K est le nombre des CT_s actives dans le réseau, BC_i est la contrainte en bande passante relative à CT_i et M est la quantité maximale de bande passante réservable par l'ensemble des CT_i actives.

2. Avec les contraintes

$$\sum_{i=0}^{K-1} N_i \leq M \quad (4.2)$$

Où N_i est la quantité de bande passante allouée par le groupe de classe CT_i .

3. Finalement

$$\sum_{i=0}^{K-1} BC_i \geq M \quad (4.3)$$

4.2.2 Russian Dolls bandwidth constraints Model

Le second modèle BC est le RDM (figure 4.2). Tout comme MAM, huit CT_s au maximum peuvent être actives dans le réseau. Contrairement au MAM, la contribution principale de ce modèle est de permettre le partage de la bande passante entre plusieurs CTs. Le système RDM peut être exprimé comme suit :

- tous les LSPs relatifs à CT_7 ne peuvent utiliser plus de BC_7 en tant que bande passante allouée;
- tous les LSPs relatifs aux CT_6 et CT_7 ne peuvent utiliser plus que BC_6 ;
- tous les LSPs relatifs aux CT_5 , CT_6 et CT_7 ne peuvent utiliser plus que BC_5 ;
- ...etc.
- tous les LSPs relatifs aux CT_0 , $CT_1, \dots, CT_7$ ne peuvent utiliser plus que $BC_0 = M$.

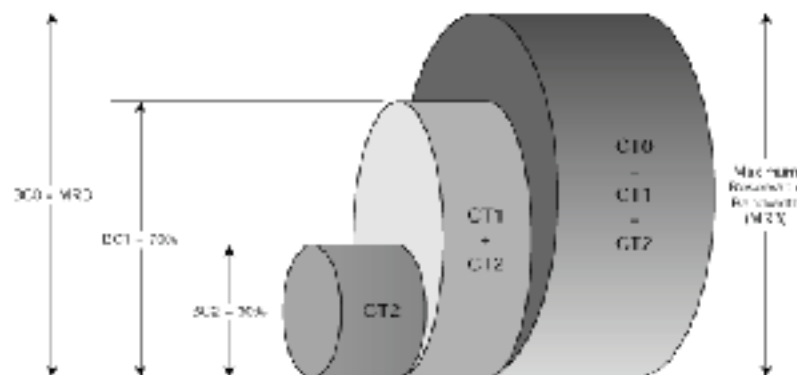


Figure 4.2 Bandwidth allocation in RDM¹

Le modèle analytique de contraintes représentant la solution RDM est donné tel qu'il est présenté par les auteurs dans l'article (Adami, Callegari et al. 2008). Pour ce modèle, le nombre de CT_s maximal actifs dans le réseau est donné par K . N_i est la quantité de bande passante allouée par la classe CT_i . BC_i est la contrainte en bande passante relative à CT_i . M est la quantité maximale de bande passante qui peut être réservée pour l'ensemble des CT_s

1. Pour tout $i \in [0, K-1]$

$$\sum_{j=i}^{K-1} N_j \leq BC_i \leq M \quad (4.4)$$

¹ <http://www.ciscopress.com/articles/article.asp?p=520184&seqNum=3>

2. Avec les contraintes

$$BC_0 = M \quad (4.5)$$

$$\sum_{i=0}^{K-1} N_i \leq M \quad (4.6)$$

Exemple d'application pour RDM

Soit l'exemple représenté par la figure 4.3 illustrant un modèle RDM avec un nombre de classe de trafic $CT_s = C=3$. En appliquant le principe de RDM on aura :

- a) tous les LSPs correspondants à CT_2 ne peuvent utiliser plus de BC_2 ;
- b) tous les LSPs correspondants CT_1 et CT_2 ne peuvent utiliser plus de BC_1 ;
- c) tous les LSPs correspondant à CT_0 , CT_1 et CT_2 ne peuvent utiliser plus de BC_0

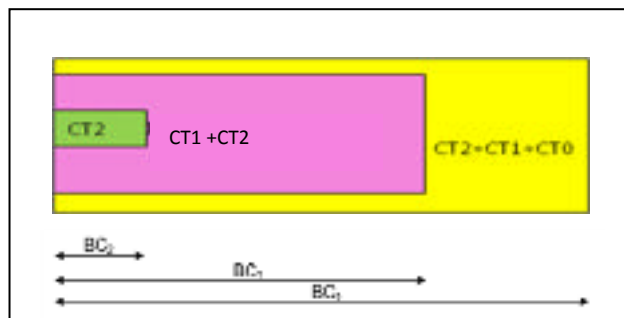


Figure 4.3 Exemple d'application pour RDM

Pour cet exemple, on utilise CT_2 pour la voix et CT_1 pour FTP et CT_0 pour le http.

1. BC_0 = Bande passante Maximale Réserveable = 3.5 Mbps : La contrainte de réservation en bande passante pour les trafics Voix + Data+Http est limitée à 3.5 Mbps;
2. $BC_1=3$ Mbps : La contrainte de réservation en bande passante pour les trafics Voix + Data est limitée à 3 Mbps;

3. $BC_2 = 1$ Mbps : La contrainte de réservation en bande passante pour les trafics Voix est limitée à 1 Mbps.

4.3 CAM : Courteous bandwidth Allocation constraints Model

4.3.1 Le modèle mathématique du CAM

4.3.1.1 Condition d'application du CAM

Afin de simplifier notre étude, nous allons considérer deux classes CT_1 et CT_0 (Figure 4.4). Soit CT_0 la classe de trafic FTP et CT_1 la classe de trafic Voix. Considérons $Pr(CT_0)$ la priorité des applications FTP et $Pr(CT_1)$ celle des applications Voix.

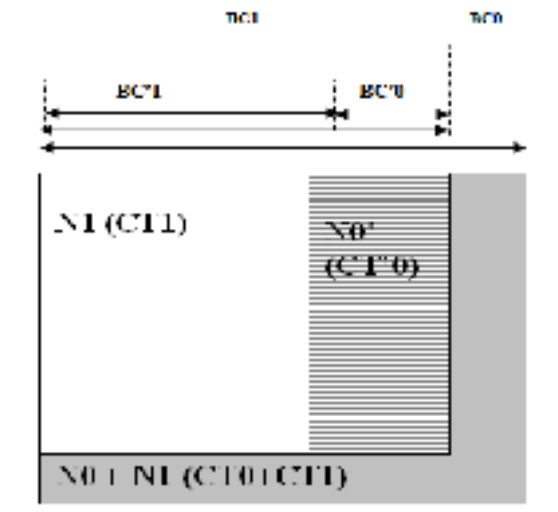


Figure 4.4 Allocation de la bande passante avec CAM

Les conditions d'application de la courtoisie se résument dans ce qui suit :

Condition 1

La priorité de la classe courtoise doit être plus haute que celle de la classe qui bénéficie des ressources additionnelles par courtoisie. Autrement dit :

$$Pr(CT_0) < Pr(CT_1) \quad (4.7)$$

Condition 2

L'application de la courtoisie ne doit pas affecter, en aucun moment, la QoS de la classe courtoise. Donc, le taux de perte de paquets, τ_v de la classe prioritaire ne doit pas dépasser le seuil τ_1 , qui représente le seuil de tolérance de perte de paquets pour la classe CT courtoise.

$$\tau_v < \tau_1 \quad (4.8)$$

Condition 3

La QoS de la classe de trafic bénéficiaire de la courtoisie doit être détériorée. Ce qui s'exprime par un taux de perte de paquet, τ_f qui dépasse le seuil de perte de paquets toléré pour cette classe, noté τ_1

$$\tau_f > \tau_2 \quad (4.9)$$

Condition 4

La portion des ressources de bande passante supplémentaires allouée par CT₀, $BW_{courteous_0}$ ne doit pas excéder la quantité de bande passante cédée par courtoisie par la classe CT₁, nommée $BW_{courteous_1}$. Donc :

$$BW_{courteous_0} \leq BW_{courteous_1} \quad (4.10)$$

Condition 5

Les prévisions en matière du trafic de la classe prioritaire ne doit pas prévoir une augmentation qui dépasse le seuil $\theta(v)$. Le seuil $\theta(v)$ est exprimé en pourcentage. Cette condition va assurer que les ressources offertes à la classe moins prioritaire ne seront pas utiles pour la classe prioritaire pendant l'exécution de la courtoisie. Notons que dans cette thèse on ne présente pas d'étude relative à ces prévisions.

4.3.1.2 Modèle de gestion de files d'attente du modèle CAM

Le modèle de contraintes d'allocation de bande passante CAM se représente par un modèle de gestion de files d'attente incluant quatre files d'attente simulant un comportement de « Bandwidth Constraints Model ». Les deux premières files Q1 et Q2 permettent de simuler

la première condition du modèle CAM (équation 4.11). Les deux autres, Q1Q et Q2Q, servent à simuler les autres contraintes de ce modèle (équations 4.12 à 4.14). La figure 4.5 illustre le système de file d'attente adapté pour CAM. Les deux serveurs, Server 1 et Server 2, appliquent la politique de gestion de file d'attente FIFO, car ils traitent chacun des paquets de mêmes types. Le dernier, Server 3, la politique de gestion de file d'attente Courteous Priority Queuing (CPQ) (Tata 2009). Les deux types de trafics considérés dans cette étude mathématique sont la voix et FTP. La voix est dotée d'une plus haute priorité et passe par les deux files Q1 et Q1Q. FTP a une basse priorité et passe par les files Q2 et Q2Q.

Les arrivées vers les files d'attente Q1 et Q2 suivent la loi de poisson alors que les départ suivent la lois exponentielle.

Notons que CPQ offre la possibilité de servir des paquets FTP au lieu des paquets Voix en appliquant la courtoisie. Pour ce faire, les conditions de la courtoisie, citées plus haut dans ce chapitre, doivent être vérifiées.

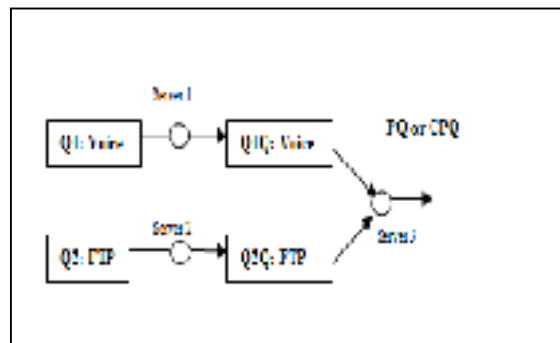


Figure 4.5 Modèle de file d'attente adaptée par CAM

Rappelons que Th_Q1Q et Th_Q2Q sont les seuils de remplissage des files Q1Q et Q2Q, respectivement. Une taille de la file Q1Q inférieure à Th_Q1Q indique que le taux de perte de paquet de types voix est acceptable. Dans ce cas, le système pourra appliquer l'algorithme de courtoisie pour la gestion des files d'attente si la taille de Q2Q dépasse Th_Q2Q . Ce qui s'interprète par un taux élevé de perte de paquets FTP. Voir (Tata 2009) pour plus de détails

sur le modèle CPQ Queuing Model et les seuils Th_Q1Q et Th_Q2Q . Rappelons que CAM est le modèle qui définit les contraintes d'allocations des ressources que CPQ utilise pour l'attribution des ressources disponibles dans le réseau pour chaque classe de trafic.

4.3.2 Description de l'algorithme CAM

Tout comme RDM, le modèle CAM (Courteous bandwidth Allocation constraints Model) permet le partage de la bande passante entre les différentes classes de service (CTs). Or, contrairement au RDM, une partie de la bande passante réservable au groupe CT_i peut être cédée au groupe CT_{i-1} , $i \in [0, K-1]$ et K est le nombre des CTs actifs dans le réseau. Cette cessation n'est pas obligatoire ni automatique, elle se fait plutôt par courtoisie quand les ressources sont limitées pour le groupe CT_{i-1} .

Le Modèle de contraintes d'allocation de bande passante CAM est détaillé dans ce qui suit :

Comme le RDM et MAM, le nombre de CT_s maximal actives dans le réseau est donné par K , N_i est la quantité de bande passante allouée par la classe CT_i . BC_i est la contrainte en bande passante relative à CT_i . M est la quantité maximale de bande passante qui peut être réservée pour l'ensemble des CTs

1. Pour chaque $i \in [0, K-1]$

$$\sum_{j=i}^{K-1} N_j \leq BC_i \leq M \quad (4.11)$$

2. Avec la contrainte

$$\sum_{i=0}^{K-1} N_i \leq M \quad (4.12)$$

3. Pour chaque $i= 1$ à K

$$N'_{i-1} = N'_{i-1} + N_{i-1} \quad (4.13)$$

Où N_{i-1} est la bande passante réservable pour la classe CT_{i-1} , N'_{i-1} est la quantité de la bande passante qui peut être cédée par courtoisie par CT_i pour le profit de CT_{i-1} . N''_{i-1} est la quantité de bande passante totale qui peut être réservée par la classe CT_{i-1} .

4. Finalement

$$\sum_{i=0}^{K-1} BC_i \geq M \quad (4.14)$$

4.3.3 Application du CAM sur LTE

Le modèle CAM, n'étant pas une solution exclusive au réseau WAN, il peut certainement être appliqué aux autres types de réseaux filaires ou sans fil, notamment sur LTE, tant que la différenciation de service est requise et est applicable et tant que le LTE supporte le protocole IP/MPLS. La figure 4.6 illustre l'application du CAM sur la technologie LTE. En effet, les ressources de bande passante seront allouées selon la politique du CAM au niveau du Backhaul LTE, ainsi qu'au niveau du réseau cœur LTE.

Alcatel affirme que « The transition to LTE and small cells will help address the need to deliver more capacity and coverage into the mobile network. But both will have a major impact on the mobile backhaul network. »². Par conséquent, en tirant profit des avantages de la solution CAM, la gestion des ressources sera améliorée au niveau du réseau Backhaul LTE. Ce progrès va certainement réduire l'impact de l'utilisation de LTE et des petites cellules sur la QoS au niveau des réseaux Backhaul. Par ailleurs, une meilleure gestion des ressources de bande passante du réseau Backhaul évitera la congestion à son niveau. Cela va améliorer la performance du réseau LTE de bout en bout. En outre, la résolution du problème de congestion dans le réseau Backhaul assure l'acheminement d'un nombre plus important de paquets pendant des périodes plus courtes. En conséquence, des ressources de bande passante seront donc libérées au niveau du Backhaul. De plus, plusieurs bearers radio mis en attente seront donc admis dans le réseau d'accès.

² <http://www.alcatel-lucent.com/solutions/mobile-backhaul>

Finalement, l'application du CAM sur le réseau cœur LTE, permettra une gestion équitable des ressources de bande. Ceci est possible en répondant efficacement aux exigences de chaque classe de trafic et particulièrement, en offrant un certain privilège aux classes de trafics de basse priorité, habituellement négligés à cause de la présence des trafics prioritaires tels que la voix ou la vidéo.

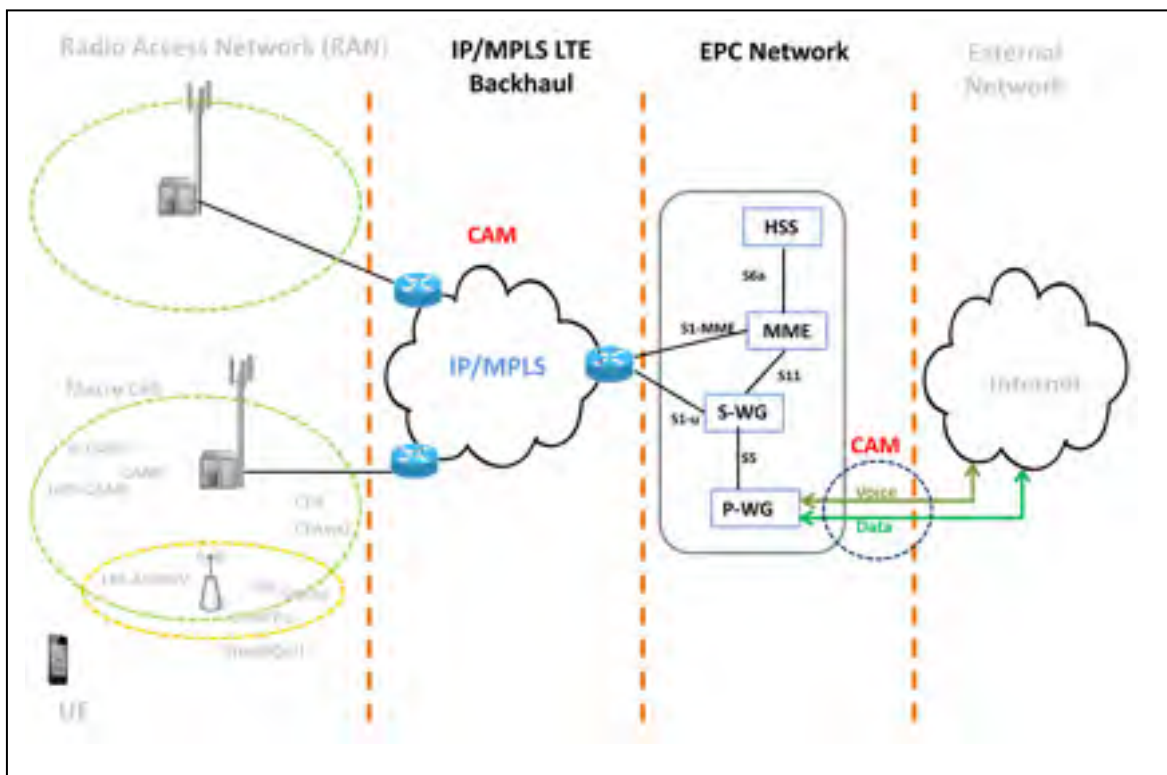


Figure 4.6 Application du modèle CAM pour LTE

La figure 4.6 illustre l'utilisation du CAM dans les deux réseaux Backhaul et cœur LTE. L'application du CAM sur LTE s'effectue de la même façon au niveau des deux réseaux spécifiés par la figure 4.6, à savoir le Backhaul LTE et le lien downlink arrivant vers le réseau cœur. Le réseau LTE permet la transmission de plusieurs types de trafics simultanément. En effet, LTE assure l'ordonnancement des flux de trafic en les regroupant, selon leur valeur QCI (QoS Class Identifier) dans des SDF (Service data Flow). Chaque SDF contient des flux de trafics ayant les mêmes exigences en termes de QoS. Les différents SDF

seront par la suite transmis via des bearers EPS. D'autre part, notons que le niveau de QoS est interprété par LTE par une valeur scalaire QCI. Ceci dit, l'application de CAM pour LTE s'effectuera telle qu'il a été décrit et conçu plus haut dans ce chapitre. La seule adaptation que nous proposons pour l'utilisation du CAM par LTE est de regrouper un ou plusieurs SDF dans une même classe CT (Classe Traffic).

4.3.4 Simulations et résultats

La validation de notre solution a été effectuée par la simulation, avec Matlab, d'un système de file d'attente des trois modèles MAM, RDM et CAM. Deux types de trafic ont été considérés, à savoir le trafic voix, avec une priorité plus haute, et le trafic FTP avec une basse priorité. Le modèle de gestion de files d'attente considéré pour les trois modèles est identique à celui représenté par la figure 4.5, à la différence que pour les deux modèles MAM et RDM, le serveur 3 utilise la politique Priority Queuing comme politique de gestion de files d'attente. Dans le but d'examiner le comportement des trois modèles nous avons effectué plusieurs simulations. Nous avons variés les paramètres de nos expériences pour conclure l'efficacité de notre modèle CAM et connaître ses limites. Pour chaque scénario plusieurs expériences ont été simulées. Les résultats représentés dans ce travail représentent les résultats moyens obtenus suite à toutes ces simulations.

4.3.4.1 Simulation d'un trafic FTP plus dense que le trafic de la voix

Ce scénario consiste à comparer les trois modèles MAM, RDM et CAM en considérant le cas d'un réseau transmettant plus de trafic FTP que de trafic Voix. FTP représente 60% du trafic global alors que la voix en représente 40%.

Les paramètres de la simulation sont présentés dans le tableau 4.1.

Tableau 4.1 Paramètres de la simulation

Paramètres de la simulation	CAM	RDM	MAM
Inter-arrival Q1 (s)	0.0012	0.0012	0.0012
Inter-arrival Q2 (s)	0.0008	0.0008	0.0008
Q1Q & Q2Q Service “ μ ” (s)	0.002	0.002	0.002
Q2 service (s)	0.002	0.002	0.0015
Q1 service (s)	0.001	0.001	0.001
Max Waiting Time allowed (Voice) (s)	0.2	0.2	0.2
Taille limite de Q1 : Q1_limite (paquets)	120	120	120
Taille limite de Q2 : Q2_limite (paquets)	150	150	150
Taille limite de Q1Q et Q2Q (paquets)	150	150	150
Th_Q1Q (paquets)	50		
Th_Q2Q (paquets)	130		
Paquets FTP (%)	60	60	60
Paquets Voix (%)	40	40	40

Le tableau 4.2 résume les résultats numériques les plus importants, relatifs à la simulation des trois modèles MAM, RDM et CAM. L’interprétation des résultats numériques montre que le modèle MAM offre le meilleur délai d’attente moyen pour la voix. En revanche, le modèle CAM améliore le délai moyen pour le trafic data ainsi que son taux de perte de paquets. En outre, bien que CAM augmente le taux de perte de paquets Voix, mais ce taux reste acceptable vu qu’il ne dépasse pas 10 % qui représente le taux de perte de paquets tolérable pour la voix. Rappelons qu’aucun mécanisme de correction d’erreur n’a été implémenté dans cette solution. L’ajout d’un système de correction d’erreur permet de réduire l’impact de la perte de paquet sur la QoS de la voix, ainsi que de FTP dans le système. Notons que plus de 61% des paquets FTP ont bénéficié de la courtoisie.

Tableau 4.2 Résultats numériques

Résultats	MAM	RDM	CAM
Temps de la simulation (s)	23.9	23.9	23.9
Délai moyen de la voix (s)	0.000801	0.000829	0.017928
Délai moyen de FTP (s)	0.11075	0.1053	0.05998
Paquets arrivés (Voix)	20032	20032	20032
Paquets perdus (Voix) (%)	0	0	7%
Paquets arrivés (FTP)	29962	29962	29962
Paquets perdus (FTP) (%)	7.13%	6.2%	1%
Paquets servis par courtoisie			18433
% des paquets servis par courtoisie			61.5213%

Les figures 4.7 à 4.11 représentent les résultats graphiques de notre simulation pour l'allocation des ressources de la bande passante en utilisant les trois algorithmes CAM, RDM et MAM.

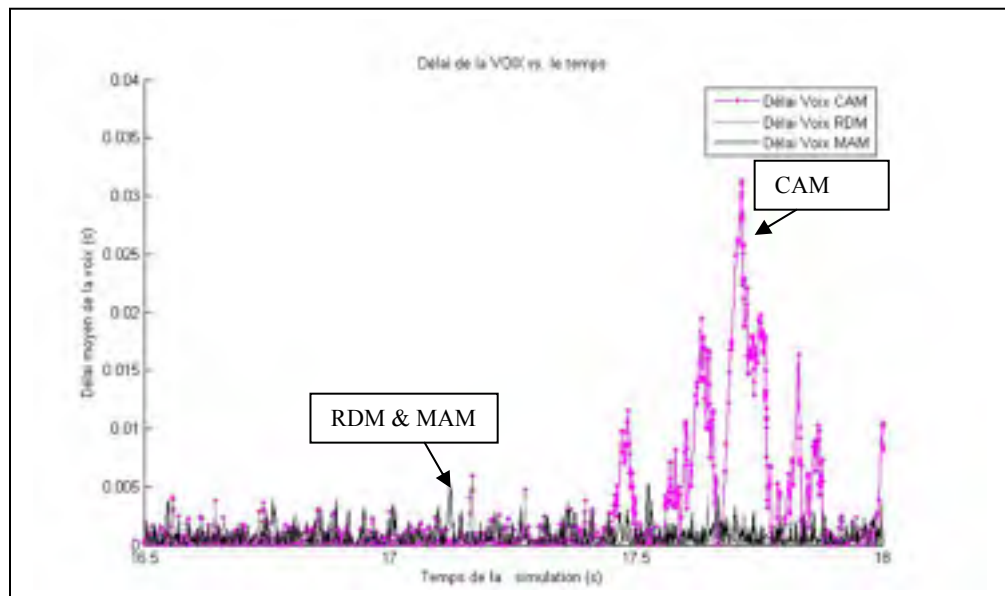


Figure 4.5 Délai de la voix

La figure 4.7 illustre les résultats de délais de la voix. Il est clair que CAM obtient les délais les plus longs pour la voix par rapport aux deux autres solutions. Notons que ce résultat reste acceptable tant que le délai de la voix reste inférieur à 150 ms. En effet, le délai supplémentaire de la voix lors de l'application de CAM est justifié par l'utilisation de la courtoisie. Notons que la garantie de la QoS de la voix est l'une des conditions de la l'application de la courtoisie. Rappelons que le principe de la courtoisie consiste à libérer certaines ressources réservées pour les classes prioritaires pour qu'elles soient utilisées par les trafics moins favorisés. Cela n'est possible que si la QoS des classes courtoises n'est pas affectée. En effet, avant de céder des ressources par la voix pour FTP, un temps, dis de tolérance, est calculé par le système CPQ pour déterminer le temps d'attente supplémentaire des paquets voix dans le système de files d'attente, sans que la QoS de la voix ne soit affectée pour autant.

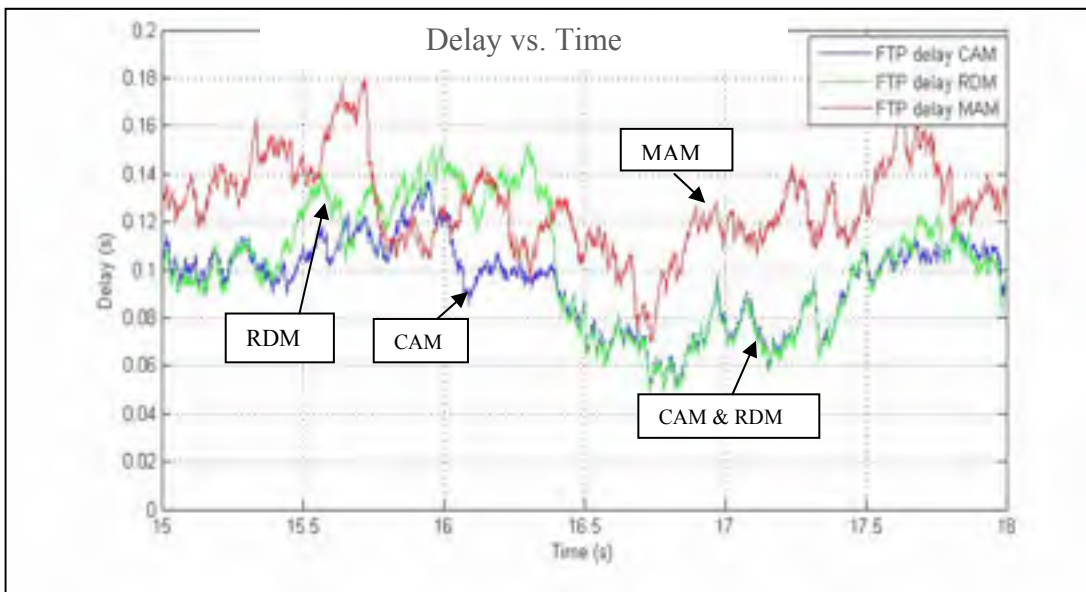


Figure 4.6 Délais de FTP

Par ailleurs, nos résultats de simulations montrent que l'algorithme CAM offre le meilleur délai pour FTP par rapport aux deux approches RDM et MAM (figure 4.8). En tout temps, le délai de FTP est plus petit en appliquant CAM, comparativement aux deux autres approches.

D'autre part, MAM affiche les plus longs délais pour FTP à cause du principe de l'allocation fixe de la bande passante pour chaque type de trafic. L'inconvénient de MAM est le fait que la bande passante réservée pour chaque classe de trafic soit déterminée initialement par le système et demeurera inchangé, même si les besoins en ressources réseaux changent dans le temps pour l'une ou l'autre des classes trafics. De ce fait, certaines ressources seront parfois accaparées par une classe de service, la voix dans cette simulation, sans les utiliser. On parlera donc d'une sous utilisation de la bande passante. Ceci arrive malheureusement, alors que la classe de basse priorité, FTP pour cette expérience, souffre d'un délai d'attente et d'un taux de perte de paquets élevé. RDM et CAM sont des systèmes qui sont conçus pour résoudre cette problématique. Ils permettent tous les deux le partage de la bande passante entre la voix et FTP. Cette technique améliore l'utilisation des ressources du réseau et réduit le délai moyen des trafics moins prioritaires. Toutefois, l'application de la courtoisie fait du CAM un meilleur système d'allocation de ressources pour les trafics moins prioritaires. Tant qu'il arrive à fournir des ressources supplémentaires pour les classes moins prioritaires. Par conséquent, il parvient, mieux que RDM, à réduire les délais d'attente de ces classes.

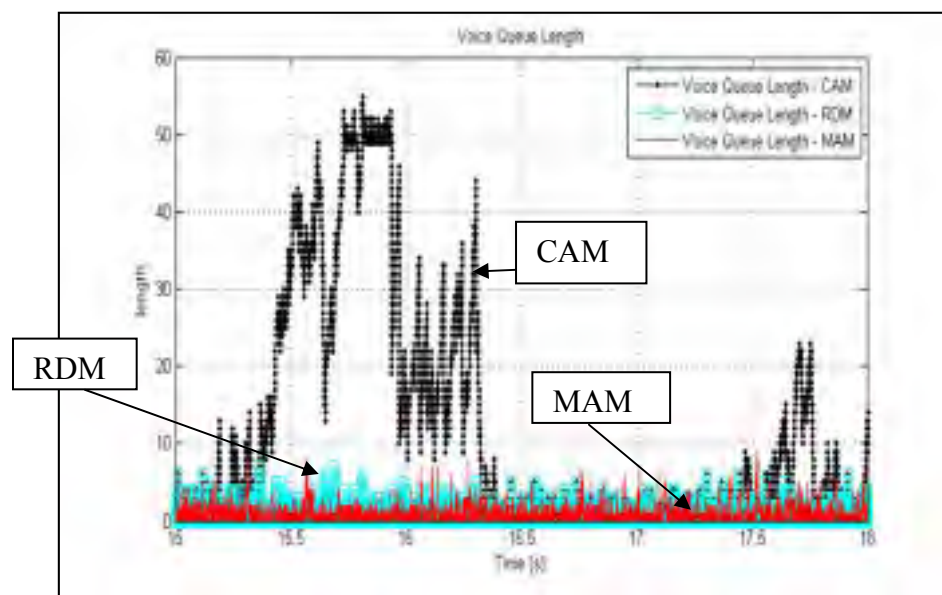


Figure 4.7 Longueur moyenne de la file d'attente voix

D'un autre côté, les résultats de nos simulations indiquent que la longueur moyenne des files d'attente adapte le même comportement que le délai. En effet, la taille des files d'attente est plus longue dans le cas de l'application de MAM que pour RDM et CAM. Le CAM donne par contre le meilleur résultat (figure 4.9). Encore une fois, le mécanisme de courtoisie a été avantageux comme système d'allocation de bande passante avec contraintes. CAM a contribué à la réduction de la taille de la file d'attente FTP. D'autre part, on remarque que le niveau de la QoS de la voix reste acceptable (figures 4.9 et tableau 4.2), la taille de la file de la voix n'affiche pas une augmentation important et la perte de paquets de la voix reste acceptable (figure 4.9). En outre, les résultats numériques correspondant à la perte de paquets de types voix affichent une perte maximale de 7% de paquets (tableau 4.2). Finalement, tel que nous l'avons cité ci-dessus, le délai maximal des paquets de la voix n'atteint pas 200 ms, donc demeure acceptable (figure 4.7).

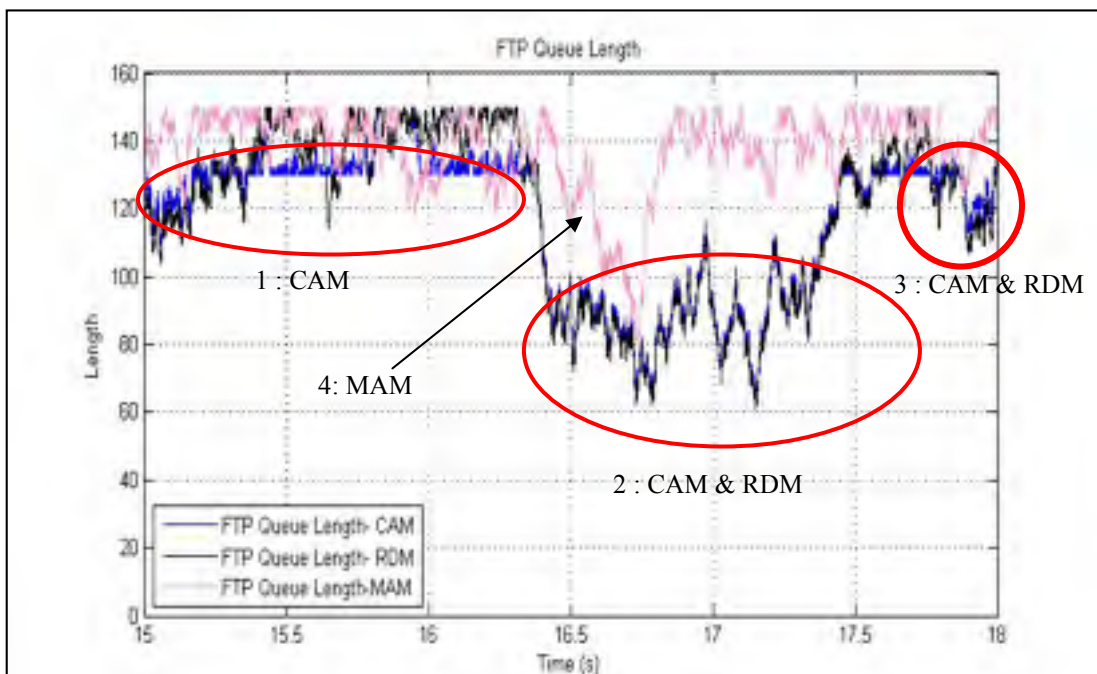


Figure 4.8 Longueur moyenne de la file d'attente FTP

La figure 4.10 illustre la variation de la taille de la file d'attente FTP, relative à l'application de chaque modèle de gestion de file d'attente avec contraintes adaptées dans cette simulation, à savoir MAM, RDM et CAM. Les principaux résultats de CAM sont entourés dans la figure

4.10. La plupart du temps CAM se montre comme le meilleur système de gestion des ressources de bande passante avec contraintes pour réduire la taille de la file d'attente FTP. Toutefois, entre 16,5 s et 17,5 s de la simulation, CAM et RDM donne les mêmes résultats. Si on revient aux résultats relatifs aux délais de la voix et de FTP, on remarque qu'entre 16,5 s et 17,5 s, le délai moyen de la voix, ainsi que celui de FTP est considéré comme le plus bas délai durant toute la simulation. Cela signifie que moins de trafic arrive pendant ce temps pour la voix et FTP. Ceci dit, la courtoisie n'aura pas raison d'être appliquée. Par conséquent, CAM aura le même comportement que RDM.

La figure 4.11 illustre une autre contribution de notre solution CAM. Cette figure montre que CAM a réduit d'une façon remarquable le taux de perte de paquets FTP, par rapport à MAM et RDM.

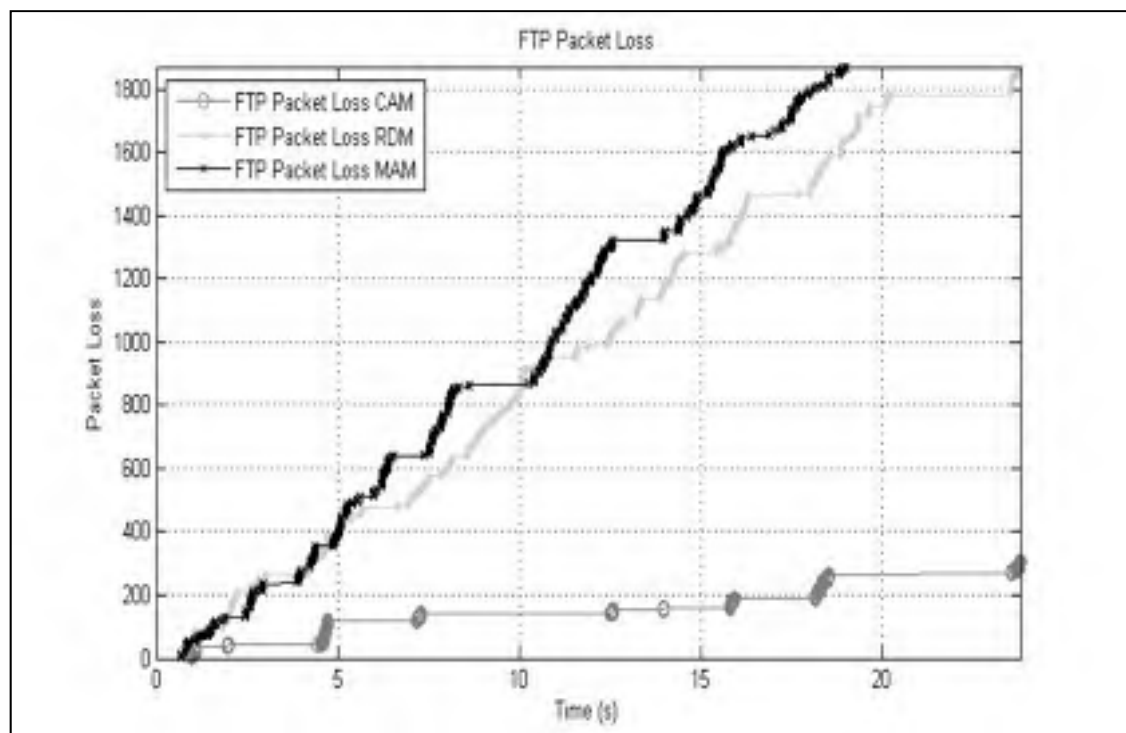


Figure 4.9 Perte de paquets de type FTP

4.3.4.2 Simulation d'un trafic FTP moins dense que le trafic de la voix

Ce scénario consiste à comparer les trois modèles MAM, RDM et CAM en considérant le cas d'un réseau transmettant moins de trafic FTP que de trafic Voix. FTP représente 30% du trafic global alors que la voix en représente 70%.

4.4 Conclusion

L'amélioration de la gestion de la bande passante dans le réseau Backhaul et le réseau cœur LTE a sans doute un impact positif sur la QoS des bearers EPS des réseaux LTE, cela s'applique sur le réseau commercial, comme sur le réseau de la sécurité publique. Cette partie filaire du réseau LTE pourra souffrir, tout comme le réseau d'accès, des problèmes de congestion suite au manque de ressources. Cela pourrait survenir principalement dans les moments où le nombre d'appels établis dans les macros et petites cellules LTE augmente. Une situation omniprésente lors des désastres et des moments de crises. Néanmoins, c'est en ces moments que les moyens de télécommunications sont les plus importants et cruciaux pour la gestion de ces situations difficiles, d'autant plus pour les premiers répondants du réseau de la sécurité publique dont toute information est importante pour sauver des vies. De plus, c'est durant ces moments que la qualité de l'information émise et reçue joue un rôle inhérent dans l'efficacité de la gestion de la crise. Par conséquent, l'amélioration de l'utilisation des ressources de bande passante dans le réseau Backhaul et le réseau cœur LTE n'est pas moins importante que celle reliée au réseau d'accès radio (RAN).

Une meilleure façon de faire pour garantir une gestion efficace et équitable des ressources de bande passante dans le réseau Backhaul et le réseau cœur LTE est l'adaptation d'un système de gestion de ressources basé sur les contraintes d'allocation de ces ressources. Pour cela, une nouvelle approche pour la gestion des ressources de bande passante avec contraintes a été développée dans ce chapitre. Ce système, appelé CAM, applique le principe de la discrimination positive lors de l'allocation des ressources pour répondre aux exigences de chaque classe de trafic, et donc garantir la QoS de chaque classe. On entend par une discrimination positive, la prise en considération de la priorité du trafic lors de l'allocation

des ressources. Toutefois, cette discrimination n'est pas absolue, des contraintes d'allocation viennent encadrer la gestion de la bande passante et empêchent tout accaparement des ressources dans le réseau Backhaul, comme dans le réseau cœur LTE. La principale contribution du CAM par rapport aux deux systèmes de base MAM et RDM est l'introduction de la courtoisie dans l'allocation des ressources. En d'autres termes, CAM montre une certaine flexibilité en termes de contraintes d'allocation des ressources. Cette flexibilité se traduit par la permission offerte aux classes moins prioritaires souffrant d'un niveau de QoS détérioré de faire exception à la règle d'allocation des ressources de bande passante et de s'offrir une portion déjà réservée aux classes de haute priorité. Cette pratique n'est ni automatique ni obligatoire, d'où son nom courtoisie. La cessation de ressources par les classes prioritaires ne se fait que dans le cas où la QoS de ces classes est bonne et demeure ainsi même après l'application de la courtoisie. Les résultats de nos simulations montrent que CAM améliore la QoS pour les trafics moins prioritaires par rapport à MAM et RDM. Cette amélioration n'affecte pas pour autant la QoS des classes prioritaires.

Rappelons que cette technique est applicable sur LTE du moment que l'architecture du réseau Backhaul et du réseau cœur supporte la technologie MPLS.

Dans le prochain chapitre, nous allons traiter la problématique de gestion des ressources radio dans LTE pour l'amélioration de la QoS du réseau de la sécurité publique.

CHAPITRE 5

Gestion efficace des ressources radio dans le réseau d'accès LTE HetNet.

5.1 Introduction

Les réseaux de la sécurité publique (PSN) ont été déployés afin de garantir la gestion efficace des crises et des catastrophes naturelles. Les moyens de communication offrent une belle opportunité aux usagers de ces réseaux pour assumer pleinement leur rôle de premiers répondants, en ayant une meilleure évaluation des faits, à travers l'accès efficace et continu à l'information. En revanche, les bandes de fréquences réservées aux réseaux PSN n'offrent pas toutes les ressources nécessaires pour une gestion optimale des désastres. Des ressources additionnelles ont pu finalement être fournies par les réseaux LTE. En effet, les ressources radio commerciales LTE ont été partagées pour permettre un accès commun aux usagers commerciaux (CN) et aux usagers de la sécurité publique (PS). À partir de là, l'accès des clients PS au réseau d'accès radio (RAN) commercial devient un défi à soulever pour le réseau LTE.

Dans ce chapitre, d'une part, l'algorithme Courteous Priority Access (CPA) a été développé pour l'accès des utilisateurs PS à la RAN LTE commerciale partagée. Notre approche permet aux clients PS un accès à la RAN commerciale, à la fois avec priorité et courtoisie. La priorité permet aux premiers répondants du réseau de la sécurité publique d'accomplir leurs tâches avec le même niveau de QoS que celui offert dans les bandes de fréquences PSN dédiées. La courtoisie, quant à elle, intervient pour l'amélioration du niveau de la QoS des classes moins prioritaires, en leur offrant des ressources radio supplémentaires. Notons que même si la courtoisie appliquée pour la gestion de la bande passante et celle utilisée pour les fréquences radio dérivent d'une même idée, qui est de céder sa place généreusement quand cela est possible, la structure de chaque algorithme est différente pour répondre spécifiquement à la politique de gestion des ressources, soit radio ou bien bande passante. Ceci dit, l'algorithme de courtoisie que nous utilisons dans ce chapitre, c'est celui qui s'applique à la gestion de fréquences radio.

D'autre part, une nouvelle approche, à savoir, CPA with Offloading (CPAwO), a été conçue pour assurer la décharge de la macro cellule LTE lors des moments de congestion. CPAwO est une solution généralisée de CPA. De plus, elle fournit des ressources radio additionnelles offertes par les petites cellules via le processus d'offloading.

Par ailleurs, tout comme l'allocation de la bande passante aux trafics à différentes exigences de QoS, la gestion des ressources radio nécessite de suivre un modèle d'allocation de ressources basé sur des contraintes, afin de garantir un partage juste de ces ressources entre les différentes classes de services. Ces contraintes permettent d'empêcher certaines classes d'accaparer toutes les ressources disponibles. Dans ce travail, nous avons adapté le CAM, proposé dans le chapitre précédent, pour qu'il soit utilisable pour les fréquences. Trois modèles ont été conçus, à savoir le modèle de base CAMF (CAM for frequencies), le modèle généralisé de CAMF, à savoir G-CAMF (Generalized CAMF) et le modèle RUS-CAMF (Radio Usage Situation based CAMF) basé sur la situation d'utilisation des ressources, à savoir situation d'urgence ou de non-urgence. Ces modèles seront décrits dans le présent chapitre.

5.2 Accès à la RAN commerciale partagée pour les premiers répondants LTE

Dans ce chapitre une nouvelle approche est proposée pour le partage du spectre radio commercial entre les usagers commerciaux et les usagers PS. L'objectif principal de cette solution est d'assurer l'accès avec priorisation et courtoisie pour les utilisateurs PS aux ressources radio partagées du réseau commercial au sein d'un réseau LTE HetNet. Deux algorithmes ont été développés, à savoir Courteous Priority Access (CPA) et CPA with Offloading (CPAwO). Notre système, en plus d'offrir des ressources commerciales supplémentaires au réseau PS, assure une certaine priorité pour les usagers commerciaux en leur attribuant des quantités de ressources radio supplémentaires par le biais du processus de gestion des ressources radio avec courtoisie. Le processus de courtoisie pour les fréquences est détaillé plus loin dans ce chapitre. Notons que la courtoisie n'est applicable que si la QoS du trafic plus prioritaire PS est acceptable. Cette approche permet de retarder la préemption

et le blocage de bearers commerciaux quand les ressources radio sont limitées. Toutefois, la courtoisie ne peut pas être appliquée quand les ressources sont épuisées dans l'eNodeB du réseau LTE. Pour cette raison, notre solution a été étendue afin de fournir des ressources additionnelles permettant de retarder ou même d'éviter la préemption des bearers dans le réseau LTE. Ce mécanisme consiste à utiliser les réseaux locaux sans fil tels que WiFi, Ad hoc ou WMN afin de basculer les bearers non admis dans la macro cellule LTE vers les petites cellules. Cette approche est faisable dans les réseaux LTE HetNets qui comportent et gèrent, entre autres, des macros cellules et des petites cellules conjointement. Ce basculement, ou offloading, doit être complètement transparent aux utilisateurs LTE (UEs). Par conséquent, les UEs vont utiliser les fréquences publiques pour échanger les données.

Les deux algorithmes CPA et CPAwO conçus spécialement pour l'allocation des ressources radio ne sauront accomplir leur tâche sans qu'ils respectent certaines contraintes lors de l'allocation des ressources. Ces contraintes, tel qu'on l'a déjà mentionné ci-haut, sont liées aux exigences de chaque type de trafic. En d'autres termes, l'allocation des quantités de ressources pour chaque classe de service doit s'effectuer de façon à lui garantir un bon niveau de QoS, sans toutefois accaparer la totalité des ressources disponibles. Néanmoins, il est clair que certaines classes, dites prioritaires, seront privilégiées par rapport à d'autres qui se trouvent moins exigeantes en termes de QoS. Ce privilège se traduit par des contraintes d'allocation moins serrées pour les classes de trafic plus prioritaires, sans que ces contraintes soient complètement annulées. Notons que l'annulation de contraintes d'allocation pour les trafics prioritaires risque de détériorer la QoS des trafics moins prioritaires.

Par conséquent, des modèles d'allocations de ressources radio avec contraintes ont été implantés dans ce chapitre. Ces modèles garantissent un bon niveau de QoS pour les classes de plus haute priorité tout en assurant un bon service pour les moins prioritaires.

Par ailleurs, l'allocation des ressources radio dans LTE s'interprète par l'admission et l'établissement des bearers radio dans le réseau d'accès LTE. LTE utilise les valeurs ARP (Allocation and retention priority) des bearers comme critère de contrôle d'admission de ces

bearers au niveau du réseau d'accès. De plus, chaque bearers exige l'acquisition d'une quantité de ressources radio. En fait, l'admission d'un bearer n'est possible que si la quantité de ressources qu'il requiert est disponible. Toutefois, quand les ressources sont limitées, certains nouveaux bearers auront la possibilité d'interrompre d'autres bearers déjà actifs dans le réseau, et s'accaparer leurs ressources. Il est à noter que la préemption de certains bearers par d'autres s'effectue selon un processus de préemption et suivant des conditions établies par LTE. Entre autres, un bearer n'est apte à interrompre un autre seulement et seulement s'il est plus prioritaire. Les deux algorithmes CPA et CPAwO prennent cette condition en considération pour la préemption des bearers actifs dans le réseau d'accès LTE. En effet, un nouveau mécanisme de classification des bearers et de calcul de priorité a été implémenté dans ce chapitre et présenté ci-après. Ce modèle de classification de bearers est celui utilisé par CPA et CPAwO.

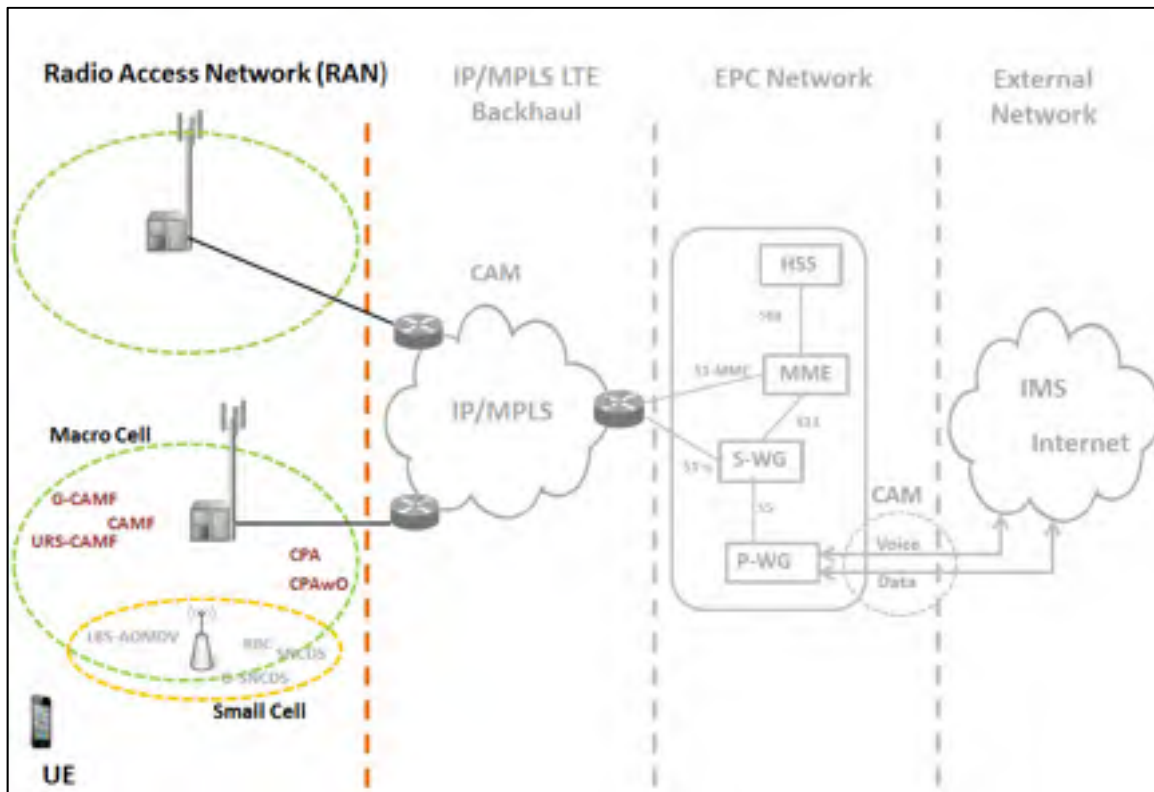


Figure 5.1 Nouvelle solution pour une meilleure gestion des ressources radio LTE

La figure 5.1 illustre l'ensemble des solutions développées dans ce chapitre, leur domaine d'application et leur position par rapport aux autres solutions implémentées dans cette thèse afin d'améliorer la qualité des communications du réseau de la sécurité publique dans le réseau LTE hétérogène. La figure 5.1 présente nos modèles d'allocation des ressources avec contraintes, à savoir, CAMF, G-CAMF, et RUS-CAMF ainsi que les algorithmes les utilisant, CPA et CPAwO. Ces solutions sont implémentées au niveau du réseau d'accès radio LTE. L'utilisation du CPAwO engendre le basculement de certains bearers vers les petites cellules. Par conséquent, afin de garantir une meilleure QoS pour les bearers admis dans ces petites cellules, certaines solutions ont été développées. Celles-ci seront détaillées dans le prochain chapitre.

5.3 Classification des bearers et calcul de priorité

Dans cette section nous discutons la classification et le calcul de la priorité des différents bearers dans le réseau LTE HetNets. En effet, dans cette étude, nous tenons à classer des bearers de type commercial et des bearers de type PS tous ensemble. Par conséquent, il faudra déterminer la priorité des bearers au sein d'un même réseau PS ou CN, ainsi que par rapport aux bearers de chaque réseau indépendamment de l'autre.

La classification des bearers, selon leur priorité, est très pertinente pour le mécanisme ARP lors du processus d'admission des nouveaux bearers, modification des bearers actifs dans le réseau LTE HetNets et aussi lors de la préemption de certains bearers pour libérer les ressources radio au profit des bearers plus prioritaires.

Le calcul de la priorité est effectué lors de la demande d'établissement d'un nouveau bearer. Trois éléments sont susceptibles d'affecter le calcul de la priorité d'un bearer de Sécurité Publique (PS) ou commercial (CN), à savoir l'état de l'utilisation (*Usage Status*), La valeur ARP du bearer et finalement sa valeur QCI (Quality of Service Class Identifier). Rappelons que l'état d'utilisation, noté α , spécifie si l'établissement du bearer est relatif à la gestion d'une situation urgente ou non urgente. Tout au long de ce chapitre, nous considérons les

notations suivantes : Un bearer établis pour la gestion d'une situation urgente est nommé emergency bearer. De même, un bearer utilisé pour traiter une situation non urgente est appelé non-emergency bearer

Les trois facteurs cités ci-dessus sont combinés ensembles afin de calculer le coefficient de priorité d'un bearer. De ce fait, nous définissons la formule suivante pour le calcul du coefficient de priorité d'un bearer

$$\Phi(Br) = \alpha * \vartheta(Br) * (\beta + QCI(Br)) \quad (5.1)$$

Où

- $\Phi(Br)$: Coefficient de priorité d'un bearer Br;
- α : Usage status = $\left\{ \begin{array}{ll} 0.1 & \text{si l'appel est urgent} \\ 3 & \text{sinon} \end{array} \right\}$;
- $\vartheta(Br)$: La valeur ARP du bearer Br;
- $QCI(Br)$: QoS Class Identifier;
- $\beta = \begin{cases} \max(\vartheta(Br)) = 8 & \text{Si le type de bearer = CN.} \\ 0 & \text{Sinon} \end{cases}$

Notons que plus la valeur de $\Phi(Br)$ est petite, plus la priorité du bearer est haute. En outre, dans notre solution nous considérons que tous les bearers urgents de type PS sont plus prioritaires que les bearers urgents de type CN pour l'accès à la radio commerciale partagée. Pareillement, les bearers non urgents de type PS sont plus prioritaires que les bearers non urgents de type CN pour l'accès à la même bande de fréquence partagée. La priorisation des bearers PS par rapport aux bearers CN lors d'une même situation d'utilisation est accomplie dans notre approche grâce à l'utilisation du facteur β . Le coefficient β garanti que la priorité la plus basse des bearers de type PS aura une valeur plus petite pour Φ que celle de la priorité la plus haute des bearers de type CN. De plus, les appels urgents, PS ou CN, doivent avoir une priorité plus haute que les appels non urgents, quel que soit leur type. Pour répondre à cette problématique, le facteur Usage Status, α , a été considéré dans le calcul de Φ . En effet,

les deux valeurs de α permettent de ranger les coefficients de tous les bearers urgents dans l'intervalle $[0.01, 2.55]$ (équations 5.6 et 5.7) et les coefficients de tous les bearers non urgents dans l'intervalle $[3, 765]$ (équations 5.11 et 5.12).

Rappelons que les valeurs de QCI varient entre 1 et 9 inclusivement et que celles de l'ARP varient de 1 à 15 inclusivement. Une modélisation numérique du calcul des priorités des bearers est donnée ci-dessous. L'objectif est de démontrer que notre approche permet de privilégier tous les trafics prioritaires par rapport aux moins prioritaires, ainsi que le trafic PS par rapport au trafic CN lors d'une même situation d'utilisation.

Les formules (5.2 à 5.12) montrent le calcul des coefficients de priorité de trafic relativement à deux situations d'usages différents, notamment une situation d'urgence et une situation non urgente.

$$\text{Min}(\Phi_{\text{status}}(\text{Br})) \leq \Phi_{\text{status}}(\text{Br}) \leq \text{Max}(\Phi_{\text{status}}(\text{Br})) \quad (5.2)$$

Si Status = urgence, alors $\alpha = 0.01$ et :

$$\begin{aligned} \text{Min}(\Phi_{\text{emergency}}(\text{Br})) &= \text{Min } \alpha * \text{Min}(\vartheta(\text{Br}_{\text{PS}})) * \text{Min}(QCI(\text{Br}_{\text{PS}})) \\ &= 0.01 * 1 * 1 = 0.01 \end{aligned} \quad (5.3)$$

$$\begin{aligned} \text{Max}(\Phi_{\text{emergency}}(\text{Br})) &= \text{Max}(\Phi_{\text{emergency}}(\text{Br}_{\text{CN}})) \\ &= \alpha * \text{Max}(\vartheta(\text{Br}_{\text{CN}})) * (\beta + \text{Max}(QCI(\text{Br}_{\text{CN}}))) = 0.01 * 15 * (8 + 9) = 2.55 \end{aligned} \quad (5.4)$$

Donc, on pourra écrire :

$$0.01 \leq \Phi_{\text{emergency}}(\text{Br}) \leq 2.55 \quad (5.5)$$

De même, en appliquant la formule (5.1) on pourra écrire

$$\begin{aligned} 0.01 &\leq \Phi_{\text{emergency}}(\text{Br}_{\text{PS}}) \leq 0.72 \\ 0.81 &\leq \Phi_{\text{emergency}}(\text{Br}_{\text{CN}}) \leq 2.55 \end{aligned}$$

Si Status = non-urgence , alors $\alpha = 3$ et

$$\begin{aligned} \text{Min} \left(\Phi_{\text{non-emergency}}(Br) \right) &= \text{Min} \left(\Phi_{\text{non-emergency}}(Br_{ps}) \right) \\ &= \alpha * \text{Min}(\vartheta(Br_{ps})) * \text{Min} \left(QCI(Br_{ps}) \right) = 3 * 1 * 1 = 3 \end{aligned} \quad (5.3)$$

$$\begin{aligned} \text{Max} \left(\Phi_{\text{non-emergency}}(Br) \right) &= \text{Max} \left(\Phi_{\text{non-emergency}}(Br_{CN}) \right) = \alpha * \text{Max}(\vartheta(Br_{CN})) * \\ &\left(\beta + \text{Max}(QCI(Br_{CN})) \right) = 3 * 15 * (8 + 9) = 765 \end{aligned} \quad (5.4)$$

Donc, on pourra écrire:

$$3 \leq \Phi_{\text{non-emergency}}(Br_i) \leq 765, i = \text{PS or CN} \quad (5.5)$$

Aussi

$$3 \leq \Phi_{\text{non-emergency}}(Br_{PS}) \leq 216 \quad (5.6)$$

$$243 \leq \Phi_{\text{non-emergency}}(Br_{CN}) \leq 765 \quad (5.7)$$

Le modèle numérique proposé ci-dessus pour le calcul de coefficients de priorité, assure que lors de la transmission de données effectuée pour un même type d'utilisation, notamment celui relié à une situation d'urgence ou non-urgence, le trafic PS sera plus prioritaire que le trafic CN, formules (5.6) et (5.7) et formules (5.11) et (5.12). De plus, il garantit que le trafic CN relatif aux appels urgents sera doté d'une priorité plus haute que le trafic PS non urgent, formules (5.7) et (5.11). Ce modèle est basé sur les valeurs QCI et ARP définies actuellement par le standard LTE. Un changement éventuel des valeurs de QCI et ARP n'affectera pas pour autant notre mécanisme de calcul de priorité et de classification des bearers. En effet, un choix convenable des valeurs de α et β doit se faire de façon à ce que les formules (5.1) et (5.2) soient toujours valides.

5.4 Modèles d'allocation de bandes de fréquences avec contraintes

Dans cette section trois modèles d'allocation de ressources avec contraintes ont été développés, à savoir CAMF, RUS-CAMF et G-CAMF, en plus du modèle de base généré par l'adaptation de l'étude effectuée dans l'article (Borkar, Roberson et al. 2011). Ce dernier modèle est conçu pour déterminer un modèle de référence pour les trois premiers modèles. Les différentes variables utilisées dans les quarts modèles analytiques d'allocations des ressources radio avec contraintes sont définies dans le tableau 5.1.

5.4.1 Modèle de base: L'approche classique

L'article (Borkar, Roberson et al. 2011) décrit l'accès à la bande de fréquence commerciale partagée. Cette approche que nous nommons approche classique illustre le partage des ressources radio entre les bearers de type PS et les bearers de type CN. Cette solution permet à tous les bearers CN d'utiliser les ressources radio de la bande partagée tant qu'aucun bearer PS ne réclame l'accès à ces ressources radio. Par contre les bearers de type PS ne peuvent utiliser qu'une portion de la bande partagée même si toutes les ressources sont libres. D'autre part, une préemption doit être appliquée sur les bearers commerciaux opérant sur les fréquences PS dès l'arrivée d'un nouveau bearer de type PS.

Certaines conditions s'appliquent pour la préemption d'un bearer actif (AB) dans le réseau par un nouveau bearer (NB). En effet, un bearer moins prioritaire ne sera pas en mesure d'interrompre un autre plus prioritaire. De plus, un nouveau bearer doit avoir la valeur Capacity to pre-emption (CP) équivalente à Yes, que nous représentons dans ce travail par la valeur positive (+1). Un bearer ayant la valeur CP équivalente à No, qui est égale à (-1) pour notre étude, ne sera pas capable d'interrompre les bearers actifs dans le réseau. D'autre part, le bearer susceptible d'être interrompu doit être vulnérable à la préemption. Cette valeur, exprimée par VP, doit être égale à Yes. Dans notre solution nous attribuons à VP la valeur positive +1 si le bearer correspondant est vulnérable à la préemption, -1 sinon.

Tableau 5.1 Définition des variables utilisées pour les modèles de contraintes

Variable	Definition
$\varphi_s(\text{PS})$	PS shared radio resources constraint
$\varphi_s(\text{CN})$	CN shared radio resources constraint
$\varphi_{s,e}(\text{PS})$	PS shared radio resources constraint for emergency situations (ES)
$\varphi_{s,e}(\text{CN})$	CN shared radio resources constraint for ES
$\varphi_{s,ne}(\text{PS})$	PS shared radio resources constraint for non-emergency situations (NES)
$\varphi_{s,ne}(\text{CN})$	CN shared radio resources constraint for NES
$\varphi'_{s,e}(\text{CN})$	CN shared radio resources constraint for ES, got by courteous scheme
$\varphi'_{s,e}(\text{PS})$	PS shared radio resources constraint for ES, given by courteous to CN
$\varphi'_{s,ne}(\text{CN})$	CN shared radio resources constraint for NES, got by courteous scheme
$\varphi'_{s,ne}(\text{PS})$	PS shared radio resources constraint for non NES, given by courteous to CN
$RR_s(\text{PS})$	PS allocated radio resources
$RR_s(\text{CN})$	CN allocated radio resources
$RR_{s,e}(\text{PS})$	PS allocated radio for ES
$RR_{s,e}(\text{CN})$	CN allocated radio for ES
$RR_{s,ne}(\text{PS})$	PS allocated radio for NES
$RR_{s,ne}(\text{CN})$	CN allocated radio for NES
$RR'_{s,e}(\text{CN})$	CN allocated radio for ES, among the PS shared frequencies, after the courteous process
$RR'_{s,ne}(\text{PS})$	PS allocated radio for NES, among the PS and CN shared frequencies for ES, after courteous process
$RR'_{s,ne}(\text{CN})$	CN allocated radio for NES, among the others shared frequencies after the courteous process
M	Maximum reservable radio resources. $M=\varphi_s(\text{CN})$
φ_{Max}	Maximum amount of radio resources
$\varphi_d(\text{PS})$	PS dedicated radio resources constraint
$\varphi_d(\text{CN})$	CN dedicated radio resources
$\varphi_s(\text{CN})$	CN shared radio resources constraint
$\varphi'_s(\text{CN})$	CN shared radio resources constraint got by courteous scheme
$\varphi'_s(\text{PS})$	PS shared radio resources constraint given by courteous to CN
$RR_d(\text{PS})$	PS dedicated radio resource
$RR_d(\text{CN})$	PS allocated radio among the PS dedicated frequencies
$RR_s(\text{PS})$	PS allocated radio among the PS shared frequencies
$RR_s(\text{CN})$	CN allocated radio among the PS dedicated frequencies
$RR'_s(\text{PS})$	PS allocated radio among the PS shared frequencies after courteous process
$RR'_s(\text{CN})$	PS allocated radio among the PS shared frequencies after the courteous process
$R_{th,s}(\text{PS})$	Courteous threshold

En se basant sur le processus de préemption décrit dans l'article (Borkar, Roberson et al. 2011), nous avons établi les formules suivantes.

$$Condition_1 = \left\{ \begin{array}{l} Pri(NB) > Pri(BA_{vp}) \\ \text{and} \\ CP(NB) = +1 \\ \text{and} \\ VP(BA_{vp}) = +1 \end{array} \right\} \quad (5.13)$$

$$Condition_2 = \left\{ \begin{array}{l} Pri(NB) \leq Pri(BA_{vp}) \\ \text{or} \\ CP(NB) = -1 \\ \text{or} \\ VP(BA_{vp}) = -1 \end{array} \right\} \quad (5.14)$$

$$Preemption(NB, BA_{vp}) = \begin{cases} True & \text{if } Condition_1 \\ False & \text{if } Condition_2 \end{cases} \quad (5.15)$$

La figure (5.2.a) représente le modèle classique d'allocation des ressources radio avec contrainte que nous avons développé en nous basant sur les concepts décrits dans l'article (Borkar, Roberson et al. 2011). En se basant sur ce même l'article, nous avons défini un modèle analytique pour l'allocation des ressources radio avec contraintes relativement à l'approche classique. Ce modèle est donné comme suit.

$$1. \text{ Pour } i = CN, \text{ et } j = PS \quad (5.16)$$

$$RR_s(i) + RR_s(j) \leq \varphi_s(i)$$

2. Avec la contrainte

$$\sum_{x=\{CN,PS\}} RR_d(x) + RR_s(x) \leq \varphi_{Max}$$

3. Finalement

$$\sum_{x=\{CN,PS\}} \varphi_d(x) + \varphi_s(x) \geq \varphi_{Max}$$

L'approche classique est modélisée pour la gestion des ressources radio allouées aux deux réseaux PSN et CN lors des situations de crises. Notons que chaque réseau détient une partie

de ressources radio qui lui est dédiée et une partie partagée avec l'autre réseau. La gestion du spectre partagé s'effectue comme suit :

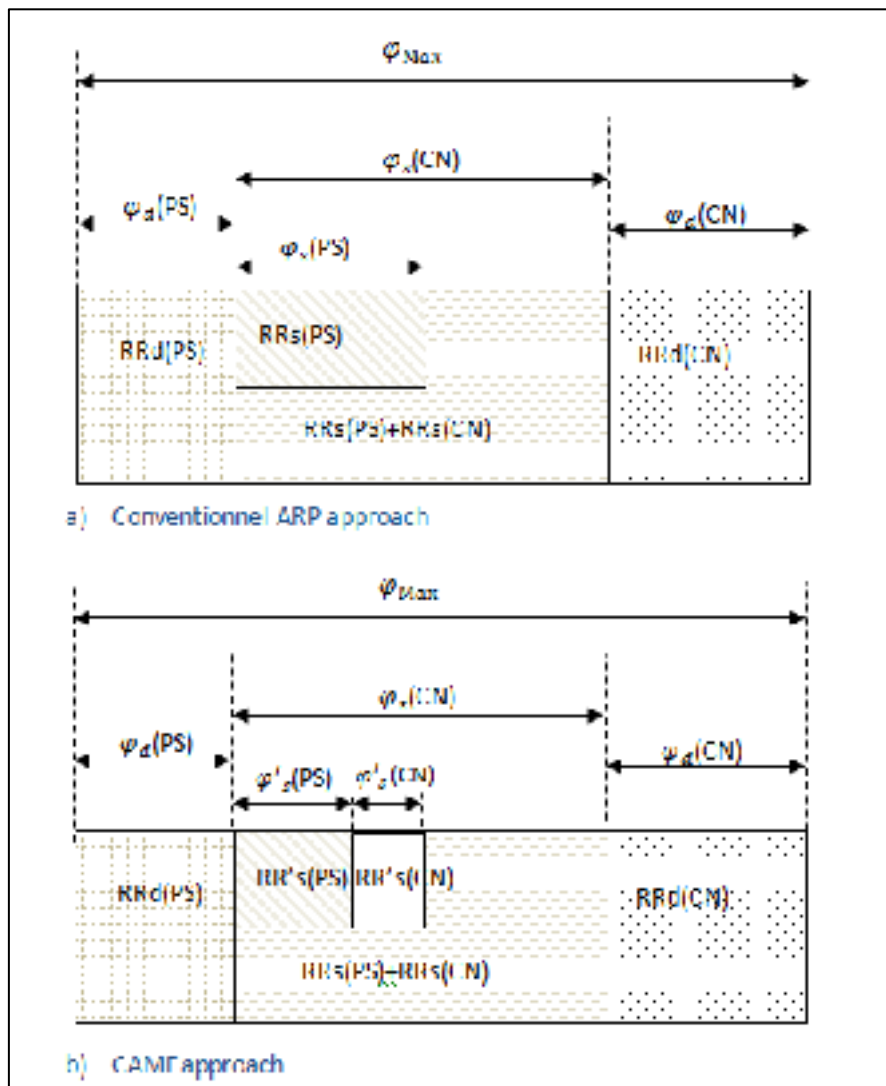


Figure 5.2 Courteous Allocation Model for frequencies

- la radio LTE partagée appartient à l'origine au réseau commercial ;
- les deux réseaux PSN et CN peuvent tous les deux allouer des ressources radio parmi celles qui sont partagées ;

- le réseau PSN ne peut allouer plus de $\varphi_s(\text{PS})$ de ressources radio au sein de la bande commerciale partagée. Cette valeur représente le seuil maximal d'allocation de la radio partagée aux bearers de type PS ;
- les bearers CN peuvent allouer la totalité de la bande de fréquence commerciale partagée, mais doivent céder la portion réservée au PSN dès qu'une demande d'établissement de bearer soit formulée par PSN.

5.4.2 Courteous Constraints Allocation Model for Frequencies

Le modèle CAMF (figure 5.2.b) est essentiellement basé sur le modèle CAM, décrit dans le chapitre 4. CAMF représente une application du modèle CAM aux fréquences radio des réseaux sans fil.

Le modèle analytique CAMF se donne comme suit

1. Pour $i = \text{CN}$ et $j = \text{PS}$

$$RR_s(i) + RR_s(j) \leq \varphi_s(i) \quad (5.17)$$

2. Avec la contrainte

$$\sum_{x=\{\text{CN}, \text{PS}\}} RR_d(x) + RR_s(x) \leq \varphi_{\text{Max}}$$

3. Et

Pour $i = \text{CN}$ et $j = \text{PS}$

$$RR_s(j) = RR'_s(j) + RR'_s(i)$$

$$RR'_s(j) \leq R_{th,s}(j)$$

4. Finalement

$$\sum_{x=\{\text{CN}, \text{PS}\}} \varphi_d(x) + \varphi_s(x) \geq \varphi_{\text{Max}}$$

Le CAMF diffère de l'approche classique par la façon dont la bande de fréquence partagée est allouée aux deux réseaux PSN et CN. En effet, le CAMF évite l'interruption des bearers CN occupant des portions radio PSN tant que le niveau de la QoS du trafic PSN demeure

acceptable. Ce niveau de QoS se traduit par le taux de blocage des bearers PS dans le réseau LTE. De ce fait, le réseau PSN applique le principe de la courtoisie pour céder une partie de ses ressources au profit du réseau CN. Ceci n'est applicable que si la QoS de service du réseau CN est détériorée alors que celle du PSN est acceptable.

5.4.3 G-CAMF

La figure 5.3 présente la généralisation du modèle d'allocation de ressources radio CAMF sur n fréquences $n \in \mathbb{N}$, en considérant les contraintes d'allocation relatives à chaque groupe de classes de fréquence, nommé Frequency Class Group (FCG_k), $k \in \{1, C\}$, C est le nombre de FCGs actifs dans le réseau.

Notons que tous les bearers appartenant à un même FCG détiennent une valeur du coefficient de priorité comprise entre cp_1 et cp_2 , tel que cp_1 est la valeur du coefficient de priorité du bearer le plus prioritaire et cp_2 est la valeur du coefficient de priorité du bearer le moins prioritaire de ce FCG. L'objectif principal de la modélisation d'un système d'allocation de fréquences avec contraintes et courtoisie est de garantir une meilleure QoS pour les bearers de hautes priorités, ainsi que d'améliorer la performance des bearers moins prioritaires.

Dans ce modèle, tous les bearers actifs dans le réseau partagent les ressources radio selon leur appartenance aux différents groupes FCG en respectant les conditions suivantes :

- tous les bearers actifs inclus dans FCG_n ne peuvent utiliser plus de $\varphi_n(FCG_n)$ de ressources radio ;
- tous les bearers actifs inclus dans FCG_n et FCG_{n-1} ne peuvent utiliser plus que $\varphi_{n-1}(FCG_{n-1}) + \varphi'_{n-1}(FGC_{n-1})$;
- etc... ;
- tous les bearers actifs inclus dans $FCG_n + FCG_{n-1} + \dots + FCG_1$ ne peuvent utiliser plus que $\varphi_{n-1}(FCG_{n-1}) + \varphi'_{n-1}(FGC_{n-1}) + \varphi_{n-2}(FCG_{n-2}) + \varphi'_{n-2}(FGC_{n-2}) + \dots + \varphi_1(FCG_1) + \varphi'_1(FGC_1)$;

- la somme de toutes les ressources de la bande de fréquence allouées doit être inférieure à la quantité maximale de ressources radio qui peuvent être allouée, notée par M.

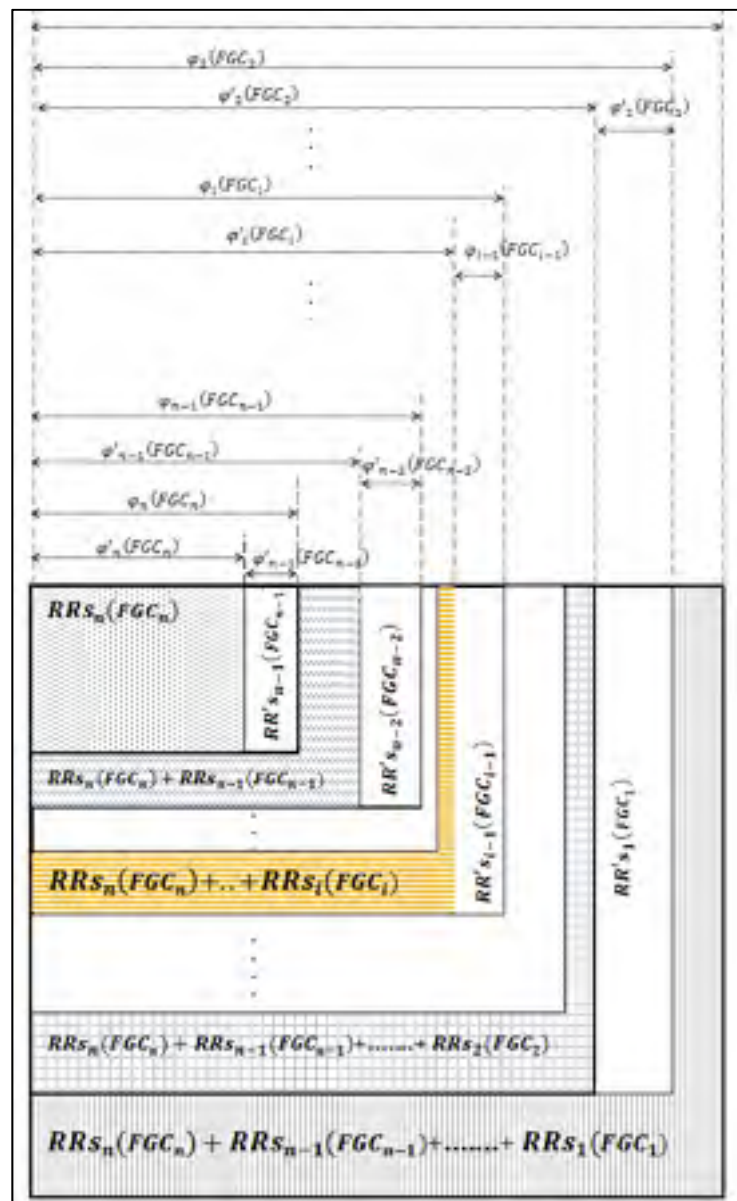


Figure 5.3 Generalized CAMF

- chaque bearer appartenant à un $FCG_k, k \in \{1, C-1\}$, C est le nombre de FCG actifs dans le réseau peut allouer RRs_k parmi $\varphi_k(FCG_k)$ de ressources radio et $RR's_{k+1}$ parmi $\varphi_{k+1}(FCG_{k+1})$ de radios fréquences ;
- les bearers appartenant au groupe FCG_n ayant la plus haute priorité ne peuvent pas allouer des ressources radio supplémentaires par courtoisie ;
- cependant, la somme des $\varphi_k(FCG_k)$ pour $k \in \{0, C-1\}$ peut aller au-delà de M .

Considérant les conditions citées ci-dessus, la modélisation mathématique du modèle G-CAMF est donnée par l'ensemble de formules (5.18).

Soit C le nombre de FCGs actifs dans le réseau. Pour chaque bearer dans $FCG_k, k \in \{1, C\}$, C est le nombre des groupes FCG actifs et $C = n$

$$1. \sum_{i=k}^c RRs_i \leq \sum_{i=k}^c \varphi_i \leq M$$

2. Avec la contrainte

$$RRs_C + \sum_{i=0}^{C-1} (RRs_i + RR's_i) \leq M \quad (5.18)$$

3. Pour chaque bearer dans $FCG_i, i \in \{1, C-1\}, C = n$,
 $RR's_i = RRs_i + RR's_{i+1}$

4. Pour $C = n$,
 $RR's_n = RRs_n$

5. Finalement, pour $C = n$
 $\sum_{i=0}^C \varphi_{s,i} \geq M$

5.4.4 Radio Usage Situation based Courteous Constraints Allocation Model for Frequencies

CAMF tel qu'il a été développé dans cette thèse, gère l'accès des clients CN et PSN à la bande de fréquence commerciale partagée sans prendre en considération le facteur d'urgence du trafic. Bien que CAMF considère que le trafic PS détient une priorité plus haute que celle de CN, mais les bearers au sein d'un même réseau, PS ou CN sont considérés ayant un même niveau de priorité. Autrement dit, les appels urgents seront traités de la même manière que les appels moins urgents au niveau de chaque réseau. Dans le bus de gérer efficacement les situations de crises, l'information échangée lors des moments de catastrophes doit être acheminée d'une façon prioritaire. À cet effet, le modèle CAMF a été adapté dans cette étude afin d'offrir une haute priorité au trafic relié à la gestion des crises.

Le but de cette nouvelle approche est tel que cité ci-dessus, de donner une priorité plus haute au trafic urgent par rapport au trafic quotidien, et aussi, de prioriser le trafic urgent CN par rapport au trafic non urgent de PS. D'autre part, le processus de courtoisie est repensé de telle sorte que les ressources dédiées au trafic urgent de type PS ne seront offertes, par courtoisie, que pour le trafic urgent de type CN. Cela va réduire le taux de perte de paquets PS échangés lors de la gestion de la crise. Or, le trafic non urgent de type CN ne va pas être privé de toutes les ressources radio supplémentaires parmi celles cédées par le réseau PS. Cependant, il ne bénéficiera que de ressources radio dédiées au trafic PS non urgent, offertes par courtoisie.

La figure 5.4 illustre l'adaptation du modèle CAMF pour le partage de la bande de fréquences commerciales par les trafics urgents et non urgents des deux réseaux PSN et CN

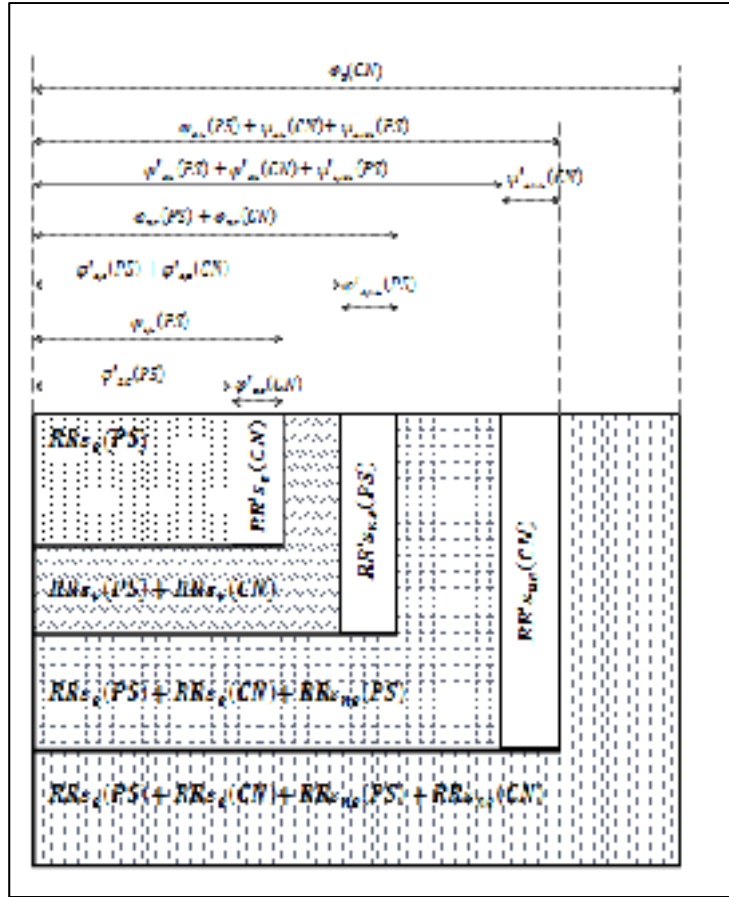


Figure 5.4 Radio Usage Situation based CAMF

Le modèle RUS-CAMF est conçu à partir du modèle G-CAMF. RUS-CAMF est un cas particulier du G-CAMF. Il est obtenu en considérant la situation d'utilisation de la radio par les deux réseaux PSN et CN. La situation d'utilisation de la radio peut avoir deux cas de figure, soit une utilisation pour une situation d'urgence (emergency) ou bien une utilisation pour une situation non urgente (non-emergency).

RUS-CAMF est défini par les clauses suivantes

- Frequency Class Groups (FCG) est un ensemble de bearers appartenant à un même réseau PSN ou CN, et ayant un coefficient de priority appartenant au même intervalle

$[cp_1, cp_2]$, tel que cp_1 est la valeur du coefficient de priorité du bearer le plus prioritaire et cp_2 est la valeur du coefficient de priorité du bearer le moins prioritaire de ce FCG ;

- les bearers sont affectés aux FCG comme suit :
 - a) les bearers CN non-emergency sont affectés au FCG_1 ;
 - b) les bearers PS non-emergency sont affectés au FCG_2 ;
 - c) les bearers CN emergency sont affectés au FCG_3 ;
 - d) les bearers PS emergency sont affectés au FCG_4 .

- l'allocation des ressources dépend du type de réseau PSN ou CN et de la situation d'utilisation de la radio, urgente (emergency) ou non urgente (non-Emergency). Le regroupement des appels dans des ensembles relativement à la situation de leur établissement et de leur type, donne les combinaisons suivantes :
 - a) CN non-emergency est représenté par la valeur 1;
 - b) PS non-emergency est représenté par la valeur 2;
 - c) CN emergency est représenté par la valeur 3;
 - d) PS emergency est représenté par la valeur 4.

Finalement, le modèle analytique représentant le mécanisme RUS-CAMF se donne comme suit :

Soit C le nombre des FCG groups. $C = 4$

1. Pour chaque bearer dans $FCG_k, k \in \{1, C\}, C$

$$\sum_{i=k}^C RRs_i \leq \sum_{i=k}^C \varphi_i \leq \varphi_s(CN) \quad (5.19)$$

2. Avec la contrainte

$$RRs_C + \sum_{i=0}^{C-1} (RRs_i + RR's_i) \leq \varphi_s(CN)$$

3. Pour chaque bearer dans $FCG_k, i \in \{1, C-1\}, C = 4$,

$$RR's_i = RRs_i + RR's_{i+1}$$

4. Pour $C = 4$,

$$RR's_n = RRs_n$$

5. Finalement, $C = 4$

$$\sum_{i=0}^C \varphi_{s,i} \geq \varphi_s(CN)$$

5.5 Algorithmes d'allocation de ressources radio dans LTE HetNets

Cette section détaille notre solution pour le partage du spectre radio commercial avec le réseau PS. Nous proposons une solution pour l'accès avec priorisation pour les usagers PSN aux ressources radio commerciales partagées. Notre approche en plus d'offrir des ressources commerciales supplémentaires au réseau PSN, elle assure une certaine priorité pour les usagers commerciaux en leur attribuant des quantités de ressources radio supplémentaires par le biais du processus de courtoisie. Ceci n'est possible que si la QoS du trafic PS est acceptable. Cette approche permet de retarder la préemption et le blocage de bearers quand les ressources radio sont limitées. L'autre volet de l'approche CPA est d'appliquer le principe d'offloading afin de réduire l'impact de la congestion des macros cellule. Cette technique consiste à basculer les nouveaux bearers arrivant vers les macros cellules LTE vers des petites cellules. Une solution totalement transparente aux usagers et efficace par sa réduction du nombre de bearers bloqués et interrompus. Par conséquent, l'allocation des ressources radio via le mécanisme de courtoisie et l'utilisation des fréquences sans licences des petites cellules, tels que WiFi, WMN ou Ad hoc, permettent de réduire le taux de blocage des nouveaux bearers et de retarder la préemption des bearers actifs dans le réseau LTE HetNets, ainsi que de réduire le cout des communications de bout en bout vue l'utilisation des fréquences publiques gratuites.

5.5.1 CPA: Accès avec priorité et courtoisie à la radio fréquence commerciale

Dans cette section nous détaillons notre algorithme CPA pour le partage du spectre radio commercial avec le réseau PS. Nous proposons une solution pour l'accès avec priorisation des utilisateurs PS aux ressources radio du réseau commerciales. Notre solution en plus d'offrir des ressources radio supplémentaires au réseau PS, elle assure une certaine priorité pour les usagers commerciaux en leur attribuant des quantités de ressources radio additionnelles par le biais du processus de courtoisie. Ceci n'est possible que si la QoS du trafic PS est acceptable. Cette approche permet de retarder la préemption et le blocage de bearers quand les ressources radio sont limitées.

5.5.1.1 Processus d'accès à la Radio de fréquence

La figure 5.5 illustre les différents cas de figure reliés à l'utilisation de la radio fréquence dédiée et celle partagée aux et entre les deux réseaux CN et PS. Ce modèle est inspiré de celui élaboré en (Borkar, Roberson et al. 2011). Dans le premier cas, la fréquence radio partagée est libre. Seules les portions dédiées sont exploitées. Les bearers de type commercial ou PS sont donc tous acceptés par ARP. Dans le deuxième cas, le nombre de demandes d'établissement de bearers commerciaux et PS accroit, ce qui invoque un accès à la radio fréquence commerciale partagée. Les nouveaux bearers demeurent acceptés par le système ARP tant que les ressources sont disponibles. Les contraintes d'allocation de ressources sont spécifiées par le mécanisme CAMF, alors que le processus d'allocation est exécuté par CPA. Cette approche permet de retarder la préemption des usagers commerciaux sans détériorer la QoS des usagers PS. La description de l'algorithme CPA est détaillée plus loin dans ce chapitre.

Le troisième cas montre le scénario de l'épuisement de toutes les ressources radio partagées par les deux réseaux CN et PS. En effet, quoique l'application du CPA ait réussi à retarder la congestion, cette solution n'a pas empêché son apparition. La préemption des appels moins prioritaires devient donc inévitable. Pour empêcher une telle situation, une autre approche alternative a été proposée. Il s'agit d'appliquer un offloading des nouveaux appels arrivés au

lieu d'interrompre les appels en cours. Cet offloading consiste à basculer les bearers arrivant vers l'eNodeB de la macro cellule vers les réseaux locaux sans fil opérant dans les petites cellules. Cela permet de libérer certaines ressources radio au niveau de la macro cellule LTE. La congestion sera donc éliminée et le retour vers le cas 1 se montre très probable.

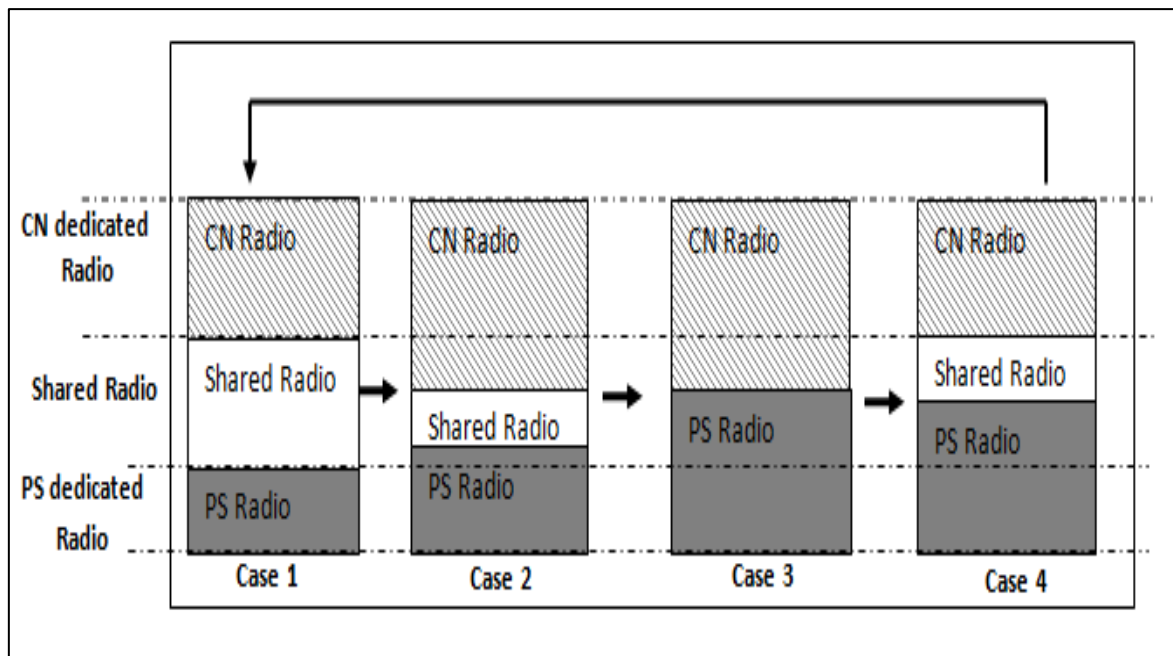


Figure 5.5 Resources Radio Management
Adaptée de Borkar, Roberson et al. (2011)

5.5.1.2 Description de l'algorithme CPA

CPA est un algorithme conçu pour l'allocation des ressources radio commerciales partagées aux bearers PS et CN lors des situations d'urgence où les ressources deviennent de plus en plus limitées à cause de l'augmentation de la demande. L'objectif de CPA est de retarder la préemption des bearers CN actifs dans le réseau tout en offrant un accès priorisé à la radio commerciale partagée pour les usagers PSN. L'idée est de concevoir un mécanisme d'allocation dynamique de ressources basé sur deux principes, à savoir la priorisation des bearers PS et la courtoisie de ces mêmes bearers qui se traduit par la cessation d'une portion

de leurs ressources radio au profit des bearers CN nécessaires, quand les conditions de la courtoisie s'appliquent.

Notons que CPA est appliqué seulement pour la gestion des ressources radio partagée entre les deux réseaux CN et PSN. En fait, CPA ne s'intéresse pas à la gestion de la bande de fréquence dédiée au PSN ou celle dédiée à CN. Rappelons aussi que l'algorithme CPA se base sur le modèle d'allocation de bande de fréquence avec contraintes CAMF pour la gestion des ressources radio.

L'accès des bearers PS à la radio commerciale partagée permet l'amélioration de la performance du réseau PSN en lui fournissant des ressources radio supplémentaires. Toutefois, les bearers CN verront une partie de leurs ressources accaparée par le PSN. Ce qui peut engendrer une détérioration de leur QoS, d'autant plus, durant les moments de congestion. Par le biais de la courtoisie, notre approche CPA permet aux bearers CN d'allouer des ressources radio PS tant que le taux de blocage des bearers PS demeure acceptable. Autrement dit, le PSN peut céder ses ressources malgré que certains de ces bearers soient bloqués, pourvu que cette situation n'affecte pas la QoS de l'ensemble du réseau PSN. Cette technique contribue à l'amélioration du niveau de la QoS du réseau CN sans affecter pour autant celui du réseau PSN. L'allocation des ressources par courtoisie n'est pas un processus irréversible. Le point fort de la courtoisie, est le fait de permettre au réseau PSN la récupération une partie ou l'ensemble de ses ressources dès qu'il en aura besoin. Par conséquent, une préemption des bearers allouant ces ressources aura lieu afin de garantir un niveau acceptable de QoS au réseau PSN. De ce fait, la courtoisie permettra une utilisation optimale des ressources radio partagées entre les deux réseaux PSN et CN.

La description de l'algorithme CPA est donnée comme suit :

ÉTAPE 1

CPA est initialisé par la réception d'une requête d'établissement ou de modification de bearer au niveau de la macro cellule LTE. Les informations relatives à la quantité de ressources

requis, le type de bearer et sa priorité sont fournis au système de contrôle d'admission afin de mesurer la possibilité d'accepter ou de rejeter la demande d'établissement du bearer. L'algorithme CPA exécute l'instruction suivante pour avoir les trois informations citées ci-dessus

$$\text{Get}(RR_{req}, Type, Priority) \quad (5.20)$$

Suivant le type du bearer, CPA invoque le processus d'allocation des ressources correspondant au PSN ou celui correspondant à CN.

ÉTAPE 2

Pour un bearer CN, le système tente initialement d'allouer des ressources parmi celles disponibles et réservées pour le réseau CN. Ceci est interprété par l'instruction suivante

$$\begin{aligned} &\text{If } RR_s(CN) + RR_{req} \leq \varphi_s(CN): \text{ Then} \\ &\quad RR_s(CN) = RR_s(CN) + RR_{req} \\ &\text{End} \end{aligned} \quad (5.21)$$

Avec

- $RR_s(CN)$ sont les ressources radio allouées à CN au sein de la radio commerciale partagée ;
- RR_{req} sont les ressources requises pour répondre à la demande du bearer en question ;
- $\varphi_s(CN)$ est la portion maximale de ressources radio commerciales partagées réservables par CN

Dans le cas où les ressources sont indisponibles pour CN, autrement dit

$$RR_s(CN) = \varphi_s(CN) \quad (5.22)$$

CPA applique l'algorithme de courtoisie pour obtenir des ressources supplémentaires. Rappelons que la courtoisie consiste à céder certaines ressources PS au profit de CN quand le

réseau CN affiche un niveau de QoS détérioré alors que le PSN n'atteint pas le seuil de blocage toléré de bearers CN. En effet, le processus de la courtoisie est invoqué quand

$$\tau_{CN} > \gamma_{CN} \quad (5.23)$$

Où τ_{CN} est le taux de blocage des bearers CN et γ_{CN} est le seuil de blocage toléré de ces mêmes bearer.

Notons que l'application de la courtoisie n'est possible que si le taux de blocage des bearers PS, τ_{PS} , est inférieur au seuil de blocage toléré de bearers PS, nommé γ_{PS} . Donc

$$\tau_{PS} < \gamma_{PS} \quad (5.24)$$

Si les conditions de la courtoisie, énumérées ci-dessus, sont vérifiées, CPA va calculer la quantité des ressources radio qui peuvent être cédées par PS. Pour déterminer si cette quantité, nommée $RR_{courteous}$, est suffisante pour acquiescer à la demande du bearer. Cette opération est effectuée sans procéder à la cessation des ressources par PSN.

Dans le cas où

$$RR_{courteous} \geq RR_{req} \quad (5.25)$$

Alors, les ressources requises RR_{req} vont être fournies parmi celles réservées au PSN. Autrement, la courtoisie ne va pas être appliquée et le bearer CN sera bloqué

Notons que PSN offre seulement la quantité radio RR_{req} par courtoisie, au lieu de $RR_{courteous}$. En outre, aucune ressource ne sera cédée par courtoisie si

$$RR'_s(CN) + RR_{req} \geq \varphi'_s(CN) \quad (5.26)$$

Où

- $\varphi'_s(CN)$ est la portion radio maximale qui peut être allouée à CN par courtoisie. Sachant que cette portion fait partie de la radio PSN réservée dans la radio partagée.

L'allocation des ressources CN par courtoisie est donnée comme suit.

$$\begin{aligned} &\text{If } RR'_s(CN) + RR_{req} < \varphi'_s(CN) \text{ Than} \\ &\quad RR'_s(CN) = RR'_s(CN) + RR_{req} \\ &\text{End} \end{aligned} \tag{5.27}$$

Où

- $RR'_s(CN)$ est la quantité radio allouée à CN par courtoisie parmi les ressources PSN;
- RR_{req} Ressources radios requises par le bearer;
- $\varphi'_s(CN)$ est la portion radio maximale qui peut être allouée à CN par courtoisie. Sachant que cette portion fait partie de la radio PSN réservée dans la radio partagée.

ÉTAPE 3

Si les ressources nécessaires pour l'établissement d'un bearer PS sont disponibles parmi celles réservées au PSN, alors CPA va satisfaire la demande du bearer. Sinon, si les ressources PS disponibles sont insuffisantes ou épuisées et que certaines sont allouées à CN alors deux cas de figure peuvent se présenter, à savoir le cas où le taux de blocage des bearers PS est inférieur au seuil de blocage toléré pour le réseau PSN et le cas où ce taux est supérieur au même seuil. Dans la première situation la requête d'établissement ou de modification du bearer sera rejetée et les ressources demeurent la propriété du CN. Dans l'autre cas, des bearers CN seront interrompus afin de récupérer les ressources PS requises pour satisfaire la demande du bearer PS. Cette politique de partage et d'allocation de ressources se traduit par les instructions suivantes

If $RR_s(PS) + RR'_s(CN) + RR_{req} \leq \varphi_s(PS)$ Than (5.28)

$$RR_s(PS) = RR_s(PS) + RR_{req}$$

Else if $RR_s(PS) \geq Th_{PS}$ and $RR'_s(CN) \geq RR_{req}$ Than

Préemption des bearers CN bearers qui utilisent les ressources $RR'_s(CN)$

Else

Bloquer le bearer

End

Où

- $RR_s(PS)$ sont les ressources radio allouées à PSN au sein de la radio commerciale partagée ;
- RR_{req} Ressources radio requises par le bearer ;
- $\varphi_s(PS)$ Portion maximale de ressources qui peuvent être allouées à PSN ;
- $RR'_s(CN)$ Radio PS allouée à CN par courtoisie ;
- Th_{PS} Le seuil d'allocation radio du réseau PSN. Avec $Th_{PS} = \varphi'_s(PS)$.

ÉTAPE 4

Attendre une nouvelle requête

5.5.1.3 Simulations et résultats

Cette section présente les résultats d'implantation de l'algorithme CPA. L'objectif principal de notre simulation, effectuée sous Matlab, est d'illustrer l'avantage de l'utilisation du processus de courtoisie pour la gestion de la bande de fréquence commerciale partagée dans le réseau LTE. Notons que les temps d'arrivée et de terminaison des appels sont générés par l'application « RandomWeight Function » de Matlab. Deux types de trafics sont considérés, à savoir le trafic PS et le trafic CN. L'acheminement de chaque flux nécessite l'établissement d'un ou plusieurs de bearers de mêmes types. Dans ces simulations, tous les bearers de type PS sont supposés avoir la même priorité et les bearers de type CN sont à leur tour considérés avoir la même priorité. Par ailleurs, les bearers de type PS ont une priorité plus haute que les bearers CN.

D'autre part, le nombre total des arrivées est de l'ordre de 6000 appels pour chaque type de trafic. De plus, 150 bearers de type PS et 50 bearers de type CN peuvent être établis simultanément dans le réseau LTE. Notons que le nombre maximal de bearers qui puissent être admis dans le réseau suite à l'application du mécanisme de la courtoisie est fixé à 50.

Les graphes des figures illustrées ci-dessous, représentent une comparaison de l'allocation des ressources radios pour l'établissement des bearers PS et CN, avec l'application de la courtoisie via l'exécution de l'algorithme CPA, versus l'application de l'approche classique où chaque bearer ne puisse être admis que dans les fréquences qui lui sont dédiées initialement.

La figure 5.6 montre le nombre des bearers bloqués dans le réseau LTE. Les résultats de la simulation montrent que l'application de l'algorithme CPA réduit le taux de blocage des bearers CN comparativement à l'approche classique. Le taux maximal de blocage des bearers CN est de l'ordre de 2,6% et de 1,9 %, avec l'application de la méthode classique et de la courtoisie, respectivement. En revanche, le taux de blocage des bearers PS augmente, mais il reste acceptable tant que ce taux reste inférieur à 1,6% de la totalité des bearers PS qui doivent être établis. Notons que ce ne sont pas tous les bearers PS ou CN qui sont reliés à la gestion de crises. Certains d'entre eux sont relatifs à des messages quotidiens non urgents, de type PS ou de type CN. En outre, l'impact du taux de perte de paquets peut être réduit par l'application d'un mécanisme de correction d'erreur, tel FEC. Une autre contribution de l'algorithme CPA, qui se montre clairement dans la figure 5.6, est le retardement du blocage du premier bearer de type CN d'environ 300 secondes par rapport à l'approche classique.

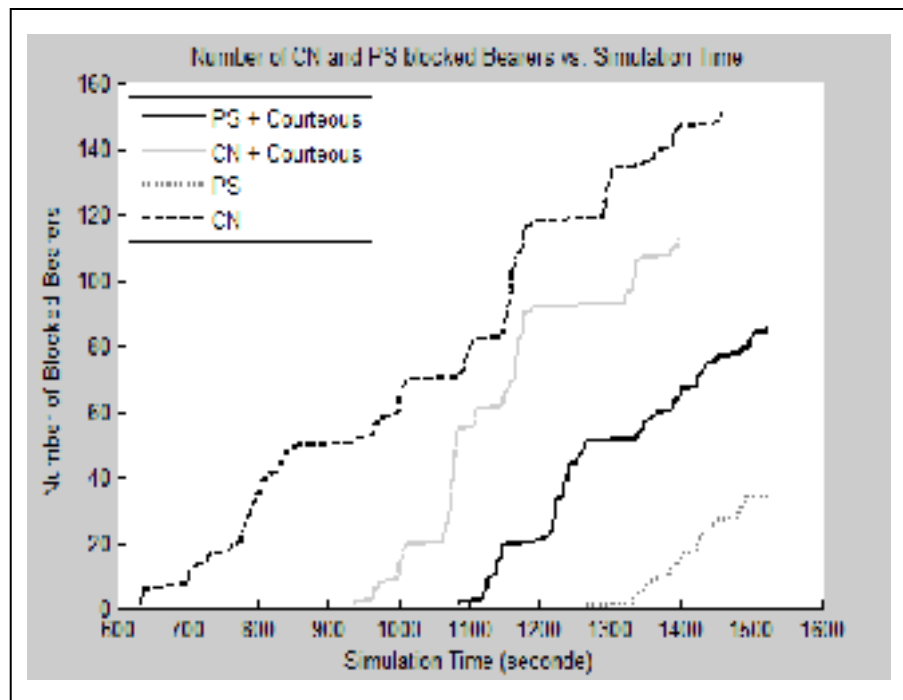


Figure 5.6 Nombre des bearers bloqués

Les résultats d'allocation des ressources pour les bearers de différents types sont illustrés par les figures 5.7 et 5.8.

Les graphes de la figure 5.7 montrent que le processus de courtoisie est invoqué à partir de 600 secondes après le début de la simulation et s'applique pendant 800 secondes. Durant cette période, plus de 100 bearers de type CN peuvent être admis simultanément dans le réseau LTE, où ils seront actifs. Rappelons que la bande de fréquence allouée pour les bearers CN ne peut permettre d'admettre plus de 50 bearers CN simultanément. De ce fait, il est clair que le surplus en nombre de bearers CN admis dans le réseau résulte de l'application de l'algorithme de courtoisie, lors de l'allocation des ressources radio.

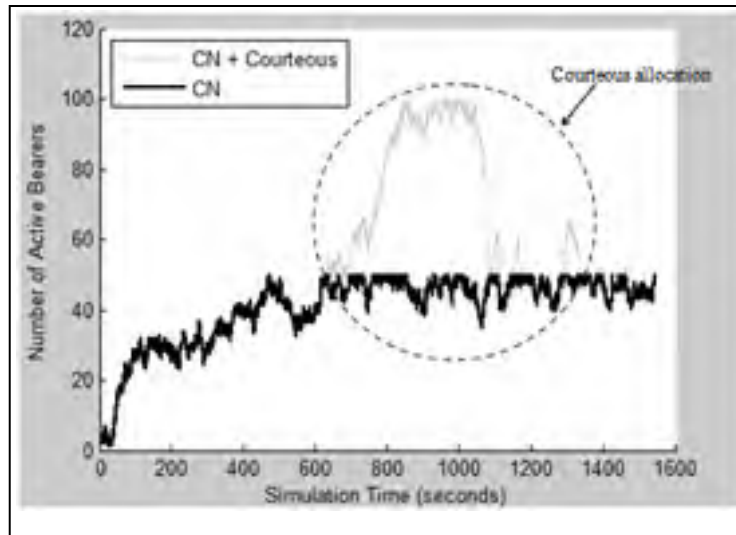


Figure 5.7 Nombre des bearers CN actifs

En outre, la figure 5.8 confirme cette constatation en montrant qu'à partir de 1100 secondes du début de la simulation le nombre de bearers PS admis ne dépasse pas 100, alors que la bande de fréquence dédiée à PS est apte à en admettre 150. Rappelons que le processus de courtoisie peut permettre l'admission des bearers de type CN en leur offrant des ressources radio de type PS, jusqu'à l'occurrence de 50 bearers CN. Remarquons que durant la période appartenant à l'intervalle de simulation [600s, 1100s] l'application de la courtoisie n'a pas eu d'impact sur le nombre de bearers PS actif dans le réseau. Cela est dû au fait que la somme des bearers PS actifs dans le réseau combiné avec le nombre des bearers CN bénéficiant de la courtoisie n'atteignent pas le nombre maximal de bearers qui peuvent être admis dans les fréquences dédiées au réseau PS, à savoir 150 bearers.

Rappelons qu'il a été constaté à travers la figure 5.6 que le blocage des bearers CN a été retardé de 300 secondes. Or, quand les ressources sont limitées, le réseau PSN ne sera pas capable de céder ses ressources au réseau CN. De ce fait une préemption des bearers CN admis par courtoisie sera effectuée (figure 5.9). Cet état de choses arrive après environ 1100 secondes à partir du début de la simulation.

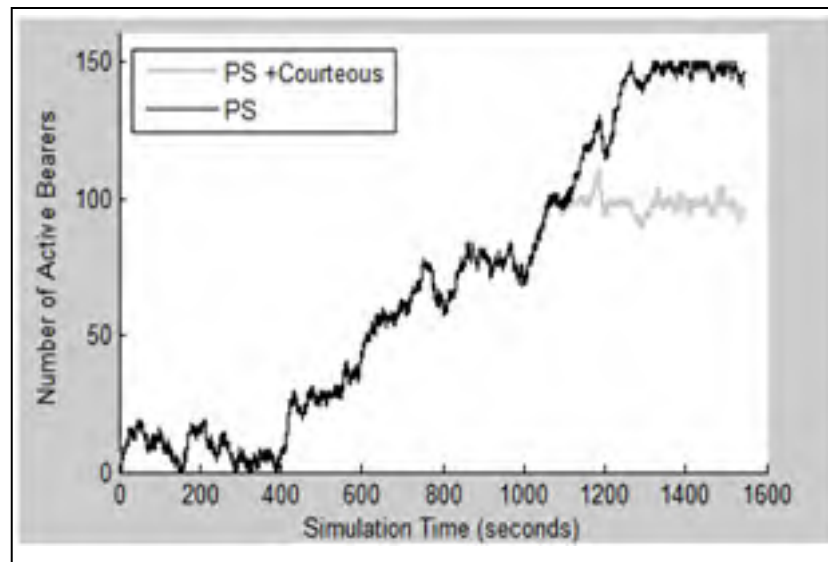


Figure 5.8 Nombre des bearers PS actifs dans le réseau

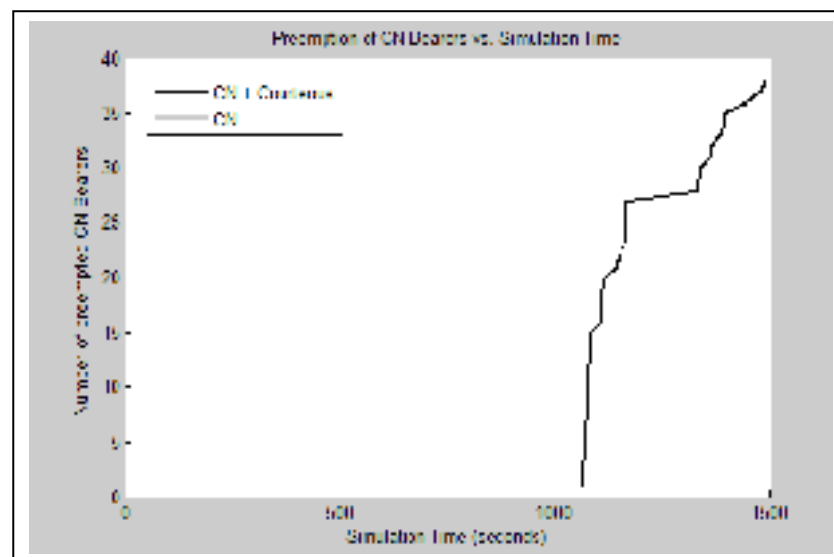


Figure 5.9 Preemption of Commercial bearers

Certainement, l'application du processus de courtoisie ne peut s'appliquer d'une façon continue et à tout moment durant l'expérimentation. Aux moments où les ressources radio sont limitées pour le réseau PS, ce dernier ne pourra point céder aucune des ressources qui lui sont dédiées. De plus, tant que la pénurie en ressources est présente, le réseau PS procédera à

la récupération de ces ressources allouées à CN par courtoisie. Des bearers CN seront donc interrompus via un mécanisme de préemption. Tel qu'il est illustré dans la (figure 5.9) et relativement aux résultats des figures 5.7 et 5.8, après 1100 s du début de la simulation, les ressources radio deviennent limitées pour le réseau PS. Par conséquent, la préemption des bearers CN sera effectuée. Les bearers CN interrompus seront choisis parmi ceux qui sont actifs dans les ressources radio réservées à PS. Notons qu'aucune interruption n'avait lieu quand l'approche classique est appliquée. En effet, aucun bearer CN n'alloue les ressources radio dédiées à PS lors de l'utilisation de l'approche classique. Par conséquent, aucun bearers de ce type n'est disponible pour interruption.

5.5.2 CPAwO : Algorithme d'allocation de bandes de fréquences partagées avec offloading

5.5.2.1 Description

L'autre volet de notre solution pour l'allocation des ressources radio est d'appliquer le principe d'offloading afin de réduire l'impact de la congestion des macros cellules. Cette technique, nommée CPAwO (CPA with Offloading) consiste à basculer les nouveaux bearers arrivant vers les macros cellules LTE, vers des petites cellules, telles que des cellules des réseaux WiFi, WMN et Ad hoc. Une solution totalement transparente aux usagers et efficace par sa réduction du nombre de bearers bloqués et interrompus. De plus, elle se montre moins coûteuse, car elle utilise des fréquences publiques pour décharger la macro cellule.

CPAwO est une généralisation de la solution CPA décrite précédemment dans ce chapitre. La première contribution de l'algorithme CPAwO par rapport à CPA est la mise à l'échelle de ce dernier à n groupes de classes de fréquences (FCG). Rappelons que le système CPA opère seulement sur deux groupes, notamment le groupe des bearers de type sécurité publique (PS) et le groupe des bearers de type commercial (CN). De plus, la technique CPAwO est une solution flexible, car elle permet un libre choix quant aux regroupements d'ensembles de bearers ayant des caractéristiques semblables au sein d'un même groupe FCG. Ceci est possible du moment que ces bearers ont des valeurs de coefficient de priorité

Φ inclus dans un même intervalle $[a,b]$, a et $b \in \mathbb{R}$. Par conséquent, cette solution devient applicable sur une infinité de nombre de groupes FCG.

L'algorithme CPAwO est un mécanisme d'allocation de ressources radio avec contraintes sur les quantités de ressources réservées à chaque groupe FCG. La quantité de la bande de fréquence allouée dépend de la valeur des priorités des bearers inclus dans le groupe FCG en question. Ce principe est appliqué afin de garantir la disponibilité d'une portion de ressources aux bearers de hautes priorités. Ainsi que pour assurer le service aux bearers moins prioritaires grâce à la contrainte relative à la limitation de la quantité maximale de ressources réservables par les bearers de classes prioritaires. Toutefois, tout comme CAP, CPAwO ne se limite pas à la gestion de l'allocation des ressources radio, mais il s'intéresse aussi à la redistribution des ressources sous-utilisées et cédées par les classes de hautes priorités au profit des classes défavorisées. Cela est applicable sous certaines conditions, appelées conditions de courtoisie. À savoir,

- les ressources cédées par courtoisie doivent appartenir au groupe de bearers de hautes priorités adjacent ;
- les bénéficiaires des ressources supplémentaires par courtoisie doivent appartenir à des groupes de basses priorités ;
- le taux de blocage de bearer de la classe prioritaire doit être acceptable est inférieur au seuil de blocage tolérable ;
- le taux de blocage de bearers de la classe moins prioritaire doit être critique, donc supérieur au seuil tolérable.

Certainement, la réallocation des ressources sous-utilisées permet une amélioration de la qualité de service des bearers appartenant à des groupes de classes défavorisés, comme elle permet de retarder la préemption des bearers actifs dans le réseau et réussir même à retarder le blocage de bearers, car il y aura moins de ressources inutilisées dans l'ensemble du réseau. Cela va permettre l'augmentation du nombre de requêtes acceptées pour l'établissement de nouveaux bearers.

Les algorithmes 5.1 à 5.4 détaillent la solution CPAwO, ainsi que les fonctions qui lui sont reliées.

Algorithme 5.1 Algorithme CPAwO

Algorithme : CPAoW	
1. Set	
	C = number of active bearers group
2. FOR EACH $i = 1$ to $C-1$	
3. Set $RR(FCG_i) = 0$	
4. Set $RR'(FCG_i) = 0$	
5. END	
6. Set $RR(FCG_C) = 0$	
7. WHILE Arrival(NB)	
8. Get ($RR_{req}, \alpha, Piority$)	
9. $\Phi = \alpha * \vartheta(NB) * (\beta + QCI(NB))$	
10. ResoucesAllocation (RR_{req}, FCG_{NB})	
11. END WHILE	

Dans l'état initial, toutes les ressources sont disponibles et le taux de réservation de la radio est nul. Ces ressources sont allouées au fur et à mesure que de nouvelles requêtes pour l'établissement ou la modification de bearers se formulent. Pour chaque requête un traitement sera effectué afin de déterminer la quantité de ressources requises pour l'établissement du bearer, ainsi que son groupe FCG d'appartenance et la situation de son utilisation, à savoir, l'utilisation des ressources pour la gestion de crises ou pour une simple communication. Ces informations sont nécessaires pour le calcul du coefficient de priorité Φ (formule 5.1).

Le classement des bearers sous forme de groupe FCG s'effectue relativement à la valeur de Φ . Rappelons que les bearers ayant des valeurs de coefficient de priorité Φ inclus dans un même intervalle $[a,b]$ a et $b \in \mathbb{R}$ seront regroupés dans le même FCG. La détermination du groupe d'appartenance FCG est pertinente pour la définition des contraintes d'allocation relatives au type du bearer qui seront appliquées par le modèle G-CAMF illustré par la figure 5.3.

Algorithme 5.2 Fonction ResourcesAllocation

Fonction : ResourcesAllocation(RR_{req}, FCG_i)	
1.	IF $RR_{req} < \varphi(FCG_i) - (RR(FCG_i) + RR'(FCG_{i-1}))$
2.	New Bearer Accepted
3.	$RR(FCG_i) = RR(FCG_i) + RR_{req}$
4.	ELSE IF $RR_{req} < \varphi(FCG_i) - \varphi'(FCG_{i+1})$
5.	CourteousAllocation(RR_{req}, FCG_i)
6.	ELSE IF IsOffloading
7.	HOW
8.	ELSE
9.	Pre-emption(RR_{req}, FCG_i)
10.	ENDIF

La prochaine étape consiste à l'enclenchement du processus de l'allocation de la radio avec contraintes (Algorithme 5.2). Le système d'allocation de fréquences tente initialement la réservation des ressources requises pour le bearer de groupe FCG_i au sein de la bande de fréquences $\varphi(FCG_i)$, avec i est le numéro de FCG incluant le bearers demandant la réservation des ressources.

Algorithme 5.3 Fonction Courtoisie

Fonction: Courteous Allocation(RR_{req}, FCG_i)	
1.	Set ω_i the number of blocked bearer belonging to FCG_i classe
2.	Set ψ_i the FCG_i classe threshold tolerated blocked bearer
3.	Set $RR_{courteous}(FCG_{i+1})$ the maximum Radio quantity can be given to FCG_i classe by FCG_{i+1} classe by courteous scheme
4.	IF $\omega_i \geq \psi_i \&\& RR'(FCG_{i+1}) < RR_{courteous}(FCG_{i+1})$
5.	$\tau = RR_{courteous}(FCG_{i+1}) - RR'(FCG_{i+1})$
6.	IF $\tau \geq RR_{req}$
7.	New Bearer Accepted
8.	$RR'(FCG_{i+1}) = RR(FCG_{i+1}) + RR_{req}$
9.	ENDIF
10.	ENDIF

Si les ressources sont limitées alors l'algorithme sollicitera le groupe adjacent supérieur, notamment le FCG_{i+1} . Par conséquent, l'algorithme de courtoisie s'exécutera. Dans le cas où les conditions de la courtoisie, citées un peu plus haut dans ce chapitre, sont vérifiées et que les ressources inutilisées ou aptes à être cédées, peuvent satisfaire la demande du nouveau bearer, alors certaines des ressources destinées au groupe FCG_{i+1} seront allouées au groupe FCG_i . Temporairement, tant que la qualité de service du groupe donateur est acceptable et que le groupe receveur des ressources souffre d'une mauvaise QoS. Dans le cas contraire, la requête d'allocation de ressources supplémentaires par courtoisie sera rejetée (Algorithme 5.3).

Certainement, le rejet de la requête d'allocation de bande de fréquence par courtoisie entrainera une détérioration dramatique de la QoS du FCG en question. D'autant plus que la demande de courtoisie ne s'effectue par les bearers nécessaires qu'en moment de congestion,

où quand leur taux de perte de paquets est non acceptable. Pour pallier cette situation, le CPAwO présente une solution pertinente pour améliorer la performance du trafic défavorisé. Une contribution qui viendra s'ajouter à celle déjà apportée par notre mécanisme de gestion de ressources radio. En effet, dans les moments de congestion et de saturation des stations eNodeB du réseau LTE, CPAwO applique l'algorithme HandOff to the WLAN (HOW) détaillé un peu plus loin dans ce chapitre. HOW est une solution de « Offloading » de la Macro cellule vers les petites cellules dans le réseau LTE Hétérogène. Au lieu de rejeter les nouveaux bearers ou d'interrompre les bearers de basses priorités actifs dans le réseau, l'algorithme HOW propose le basculement des nouveaux bearers au réseau WLAN, tels que WMN, Ad hoc ou WiFi. Notons que HOW opte pour la redirection d'un nouveau bearer de haute priorité vers les petites cellules au lieu d'effectuer la préemption d'un bearers moins prioritaire et de le basculer vers le WLAN. Cette approche est adoptée afin d'assurer un plus haut niveau de stabilité dans le réseau LTE HetNets.

La solution proposée pour le basculement des bearers vers les petites cellules se base sur deux principes pertinents pour l'amélioration de la performance de la QoS dans le réseau LTE HetNets. Notamment, le retardement de la préemption des bearers actifs dans le réseau en redirigeant les nouvelles requêtes d'établissement de bearers vers les petites cellules, ainsi que l'augmentation du nombre des bearers actifs dans le système incluant les macros et petites cellules, formant ainsi le réseau LTE hétérogène. Notons que le nombre global de bearers actif dans tout le réseau est égal à la somme des bearers actifs dans la macro cellule et les petites cellules. De ce fait, si on considère n comme le nombre de petites cellules contribuant au Data Offload et si on prend γ_j comme le nombre de bearers LTE actifs dans une petite cellule j alors on notera par Ω le gain en nombre de bearers supplémentaire admis dans l'ensemble du réseau LTE HetNets. La formule suivante illustre le calcul de Ω

$$\Omega = \sum_{j=1}^n \gamma_j \quad (5.29)$$

Algorithme 5.4 Fonction PreEmption

Fonction : Pre-emption ($\mathbf{RR}_{req}$, $\mathbf{FCG}_j$)
1. $i=1$ 2. LP = lowest priority in $\mathbf{FCG}_j$ 3. $BS(i)$= ActifBearers (LP) 4. Γ = RadioQ($BS(i)$) 5. WHILE $\Gamma < \mathbf{RR}_{req} \&\& LP < Priority_{NB}$ 6. IF $\exists$ ActifBearers (LP) $\&\& VP(ActifBearers)=+1$ 7. Γ = RadioQ($BS(i)$) 8. $i=i+1$ 9. $BS(i)$= ActifBearers(LP) 10. LP=Next lowest priority in $\mathbf{FCG}_j$ 11. END 12. END WHILE 13. IF $\Gamma \geq \mathbf{RR}_{req}$ 14. $RR_{gained}=0$ 15. FOR $i < \text{Length}(BS) \&\& RR_{gained} < \mathbf{RR}_{req}$ 16. Preempt $BS(i)$ 17. $RR_{gained} = RR_{gained} + RR_{BS(i)}$ 18. ENDFOR 19. New Bearer is Accepted 20. ELSE 21. New Bearer Blocked 22. ENDIF

D'autre part, telle la macro cellule, les petites cellules sont sujettes à la saturation. Suite à l'épuisement des ressources disponible, aucune petite cellule n'acceptera la réception des bearers LTE. On parlera donc d'une pénurie de ressources radio. Dans un tel cas, que nous

décrivons par le cas extrême de congestion, le CPAwO procède à la préemption des bearers actifs dans le réseau (Algorithme 5.4). La préemption s'effectue selon les règles établies par le standard LTE en se basant sur les conditions illustrées par les formules (5.13 à 5.15).

Notons que la préemption des bearers ne se fait que lorsque la quantité des ressources qui pourrait être suite à leur interruptions soit suffisante pour accepter la demande d'établissement des nouveaux bearers. De ce fait, la fonction de préemption commence par le calcul de gain de préemption comme le montre l'algorithme 5.4. La première étape de la fonction de préemption est de déterminer les bearers appartenant au groupe FCG ayant la plus basse priorité. Elle calcule par la suite le gain qui pourrait être obtenu suite à l'interruption d'un ou plusieurs bearers actifs de ce groupe, qui sont vulnérables à la préemption. Le processus de calcul de gain se poursuit tant que la quantité des ressources requise pour l'établissement des nouveaux bearers prioritaires n'est pas atteinte et tant qu'il existe des bearers actifs moins prioritaires et vulnérable à la préemption. Ces bearers peuvent appartenir à un ou plusieurs groupes FCG. Finalement, si le gain en ressources radio est supérieur ou égal à la quantité radio requise alors la préemption des bearers se concrétisera et le ou les nouveaux bearers seront admis dans le réseau LTE. Dans le cas contraire, la demande d'établissement de bearers sera rejetée. Rappelons que seules les bearers dont les ressources sont utilisées pour le calcul du gain en ressources radio Γ seront interrompues.

5.5.2.2 Data Offloading to Small Cells Algorithm

Le principe de l'algorithme HOW est de basculer les nouveaux bearers arrivant vers les macros cellules LTE HetNets, lors des moments de congestion, vers les petites cellules WLAN. Un nouveau bearers pourrait être redirigé vers les petites cellules même s'il est doté d'une plus haute priorité. D'autre part, les bearers de basses priorités demeureront actifs dans les macros cellules LTE, malgré l'arrivée des bearers plus prioritaires. L'objectif est de minimiser le nombre de handoff des macros cellules vers les petites cellules. Toutefois, cette stratégie ne doit pas influencer sur la QoS des bearers prioritaires.

Algorithme 5.5 Algorithme HOW

Algorithme HOW
1. $\Omega = 0$ 2. $j=1$ 3. Set n = number of small Cells 4. While $j \leq n \ \&\& \ \Omega < RR_{req}$ 5. $\Omega = 0$ 6. IF $RR_{available}(j) > 0$ 7. IF $TimeLive(RR_{available}(j)) > TimeLive(NB)$ 8. $\Omega = RR_{available}(j)$ 9. End 10. End 11. $j=j+1$ 12. End

Notons que le basculement des nouveaux bearers vers les réseaux publics ne s'effectue que si les conditions suivantes s'appliquent :

- les ressources radio doivent être épuisées au sein de la macro cellule. Ceci s'explique par la formule 5.30. Cette formule interprète le cas où pour chaque groupe FCG_i , i est le nombre de groupe de classes de fréquences actives dans le réseau, la quantité radio réservée doit atteindre la quantité maximale de ressources radio réservable ;

$$\forall i \in C, RR(FCG_i) = \varphi(FCG_i) \quad (5.30)$$

- des ressources radio doivent être disponibles dans une parmi les petites cellules et leurs durées de vie doivent être assez suffisantes pour satisfaire la requête d'établissement des nouveaux bearers (ligne 1 à 4 de l'algorithme 5.5).

- la durée de vie des ressources disponibles dans la petite cellule doit être supérieure ou égale à la durée de vie du bearer à établir (ligne 7 de l'algorithme 5.5).

5.5.2.3 Simulations et résultats

Dans le but de valider le modèle CPAwO, une simulation de cet algorithme a été effectuée dans Matlab. Le modèle de réseau considéré dans cette expérience est un réseau LTE HetNet représenté par une macro cellule LTE avec quatre petites cellules incluant chacune des fréquences radio publiques (figure 5.10). Certains usagers LTE peuvent appartenir à la fois à la macro cellule et à une petite cellule. Les usagers appartenant à une même petite cellule peuvent former un réseau WMN et sont donc aptes à s'échanger des données entre eux via des communications Device to Device (D2D). Dans cette simulation nous supposons que la localisation des différentes cellules, ainsi que de l'ensemble des usagers sont connus par le eNodeB.

Deux types de trafics sont considérés, notamment le trafic de sécurité publique (PS) et le trafic Commercial (CN). Chaque appel arrivant vers le réseau LTE est affecté à un groupe dépendamment de la situation d'utilisation et du type de l'appel en question. Deux situations d'utilisations sont prises en considération dans ces simulations, à savoir, une utilisation de la communication pour la gestion de crise, ce qui donne l'aspect Urgent au flux de données correspondant et l'utilisation du trafic pour une communication quotidienne, ce qui donne un aspect non urgent au flux de données. En effet, les quatre groupes de trafics simulés dans cette section sont le trafic Public Safety Urgent (PS emergency), le trafic Public Safety Non Urgent, (PS non-emergency), le trafic Commercial Urgent (CN emergency) et le trafic Commercial Non Urgent (CN non-emergency). Les modèles d'allocation de ressources radio avec contraintes adaptés pour cette simulation est le RUS-CAMF pour CPAwO et CAMF pour CPA. 2500 requêtes d'établissement de bearers de type PS emergency ont été simulées, contre 3000 requêtes correspondantes aux bearers CN non emergency, 8000 pour les bearers CN emergency et 2500 pour les bearers PS non-emergency.

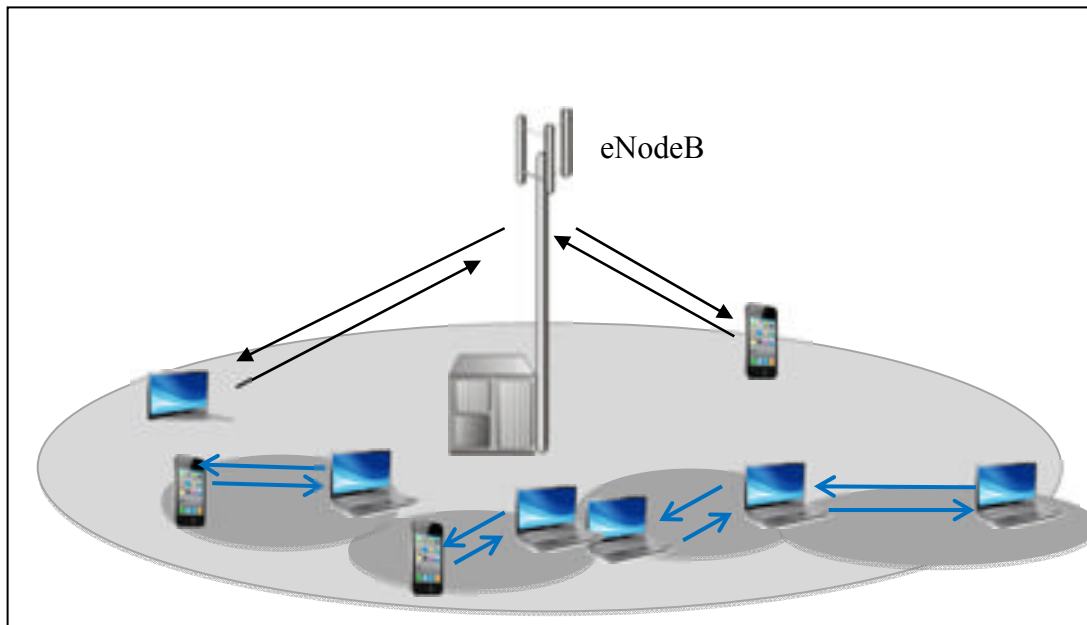


Figure 5.10 Modèle du réseau LTE HetNet

Les temps d'arrivées et de départs des bearers vers la macro cellules sont générés aléatoirement par la fonction « RandomWeight Function ». D'autre part, le trafic PS arrive 120 secondes après l'arrivée du premier appel CN. La durée de la simulation est de 600 secondes

Afin de valider notre approche CPAwO, trois modèles d'allocation de ressources radio ont été simulés. En plus du CPAwO, représenté par le modèle « Courteous+Offload », on a considéré le modèle classique implanté dans (Borkar, Roberson et al. 2011), et l'algorithme CPA, détaillé plus haut dans ce chapitre. CPA est représenté par le modèle « Courteous » dans cette simulation. Rappelons que CPA est le nouveau modèle de gestion de l'allocation des ressources radio basé sur la méthode classique à laquelle est ajoutée la technique de la courtoisie. En outre, CPQwO est la combinaison de CPA avec le mécanisme Offloading.

Par ailleurs, rappelons que l'approche CPAwO alloue des ressources au niveau de la macro cellule LTE, ainsi qu'au niveau des petites cellules WLAN. Dans ce travail, une seule macro cellule est utilisée conjointement avec quatre petites cellules. Pour le cas de CPA et l'approche classique, seule les ressources de la macro cellules sont allouées aux bearers. La

quantité des ressources radio disponibles dans les petites cellules est initialisée aléatoirement au début de la simulation. L'arrivée des bearers vers les petites cellules est relative au processus d'offload de la macro cellule, appliqué via l'algorithme HOW, défini dans ce chapitre. Autrement dit, pour chaque opération d'offload de la macro cellule est associée une arrivée de bearer vers la petite cellule. Or, le temps de départ des bearers établis dans les petites cellules est généré aléatoirement dans Matlab tout comme ceux établis dans la macro cellule.

Les figures 5.11 à 5.15 illustrent le nombre des bearers actifs dans le réseau LTE HetNet pour l'ensemble des classes de trafic simulées dans cette étude. Les bearers actifs représentent les bearers qui sont admis et qui demeurent établis dans le réseau. Rappelons que l'admission de chacun des bearer requiert la disponibilité d'une quantité de ressources radio supérieure ou égale à la quantité requise par le bearer en question. Par conséquent, le nombre de bearers actifs dans le réseau relatif à un type de trafic, reflète la quantité dans ressources allouées pour cette classe de trafic. En effet, plus le nombre de bearers actifs appartenant à une classe de trafic, est important, plus la quantité de ressources, allouées pour cette classe de trafic, est importante.

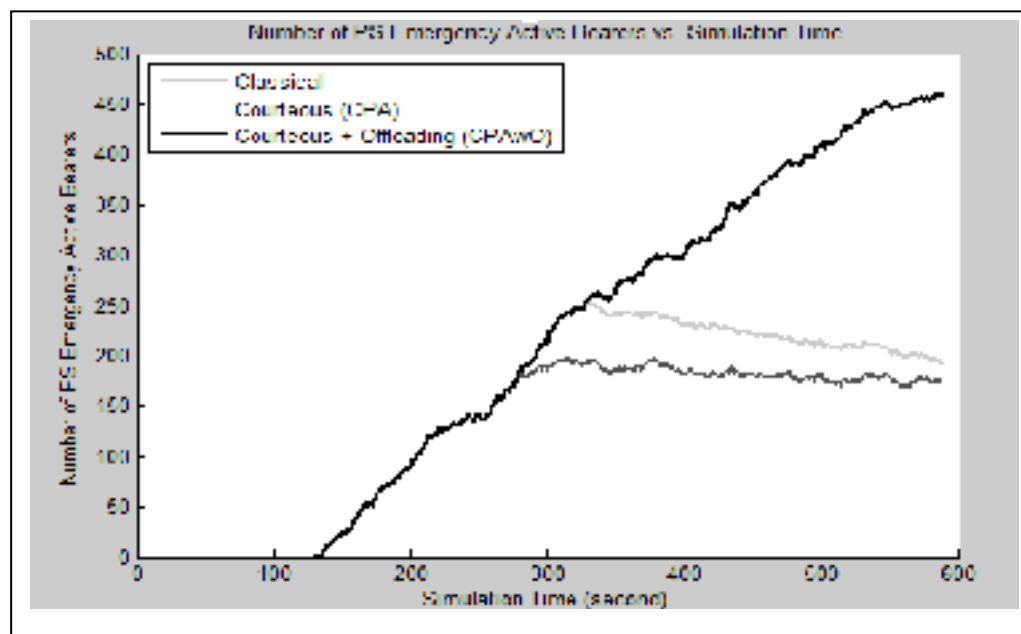


Figure 5.11 Nombre des bearers PS emergency actifs dans le réseau (CPAwO)

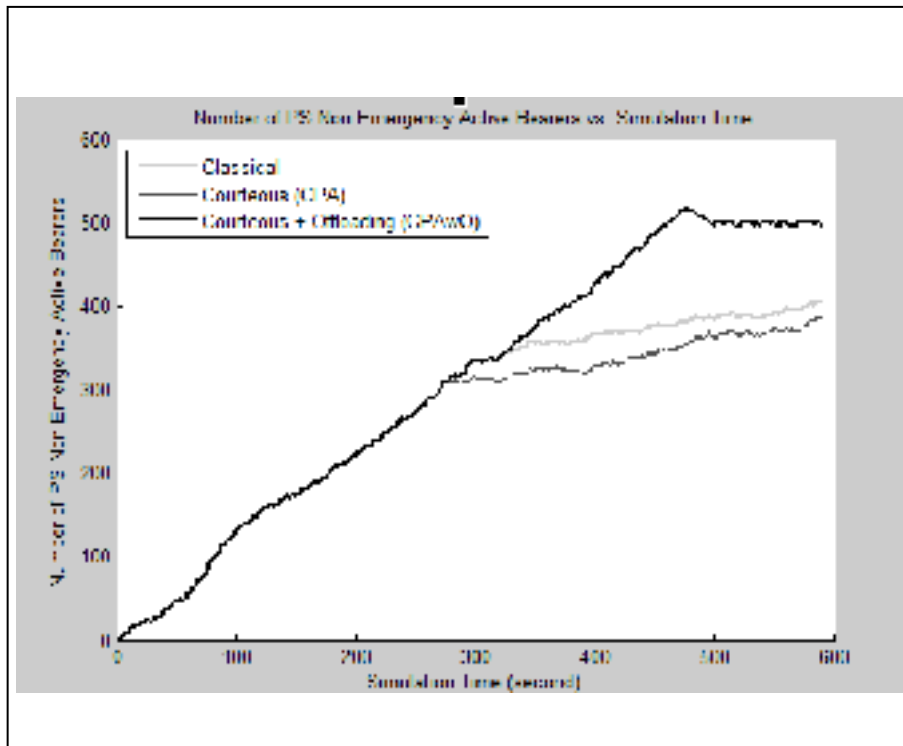


Figure 5.12 Nombre des bearers PS Non-Emergency actifs dans le réseau (CPAwO)

Les figures 5.11 et 5.12 montrent que le nombre de bearers PS actifs dans le réseau diminue avec l'application de la solution CPA, représentée par le graphe courteous, comparant à l'approche classique. Ceci est le résultat de l'application du processus de la courtoisie, qui offre certaines ressources PS au réseau CN, tant que la QoS du réseau PS ne sera pas affectée, alors que les ressources dédiées à CN sont épuisées. En effet, cela explique l'augmentation du nombre de bearers CN actifs dans le réseau quand l'approche CPA est utilisée, comparant au cas où la méthode classique est adoptée (figures 5.13 et 5.14).

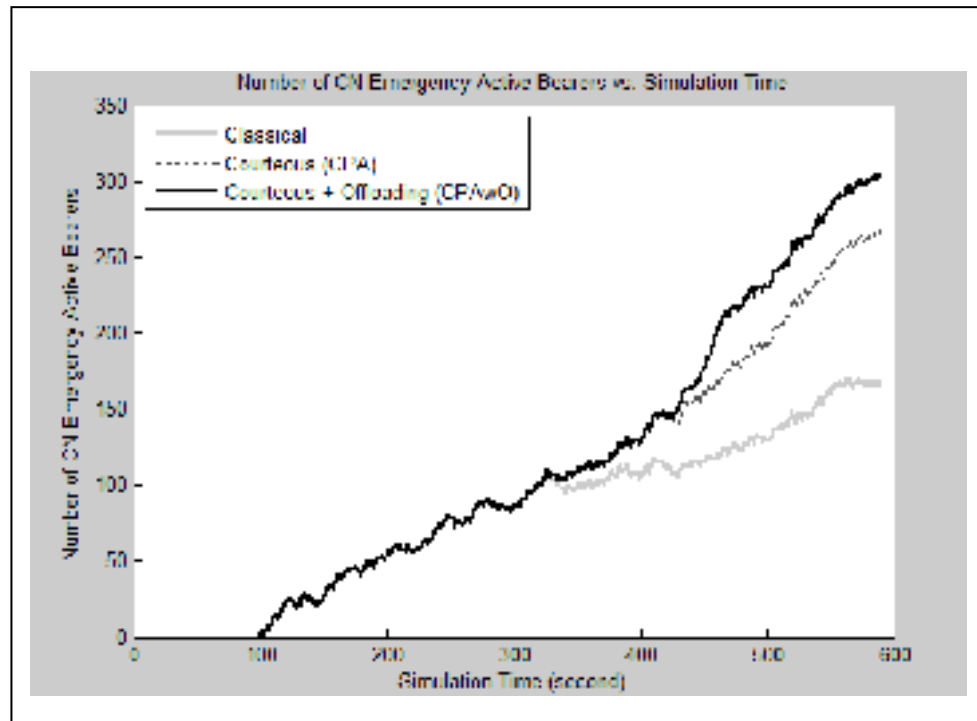


Figure 5.13 Nombre des bearers CN Emergency actifs dans le réseau (CPAwO)

D'autre part, les figures 5.11 à 5.14 montrent que l'approche CPAwO représentée par le graphique « Courteous+Offload » accroît le nombre des bearers CN actifs dans le réseau, tel que CPA, tout en augmentant le nombre des bearers PS actifs. En effet, le CPAwO est basé sur deux processus principaux, à savoir l'algorithme de courtoisie et l'algorithme d'Offloading « HOW ». Le mécanisme de courtoisie assure l'allocation de ressources radio supplémentaires pour les usagers CN durant les moments de pénurie de ressources CN. Ceci est réalisé de la même manière que cela est effectué avec CPA. Or, le mécanisme Offloading lui fournit aussi des ressources supplémentaires non seulement pour le réseau CN mais aussi pour le réseau PS. Ces ressources additionnelles sont réservées au niveau des petites cellules WLAN. Par conséquent, l'utilisation de la courtoisie combinée à l'algorithme HOW offre un meilleur modèle de gestion des ressources radio pour l'ensemble des classes de trafics et

durant les deux situations d'utilisation, à savoir les moments de crises et les appels quotidiens.

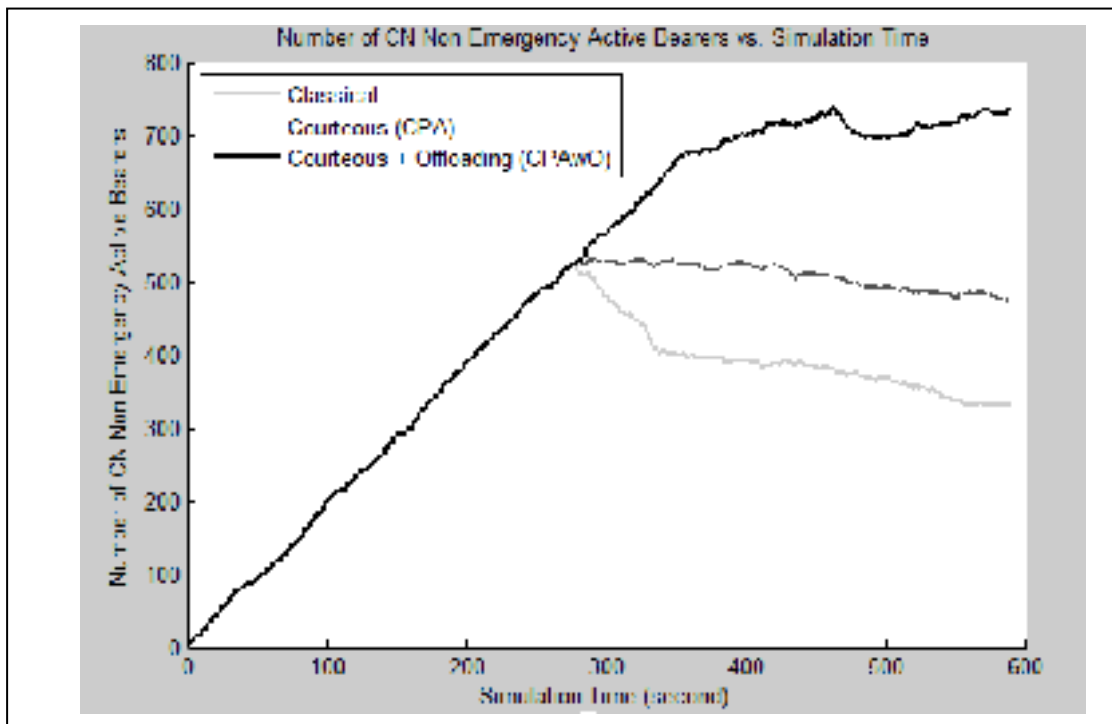


Figure 5.14 Nombre des bearers CN Non-Emergency actifs dans le réseau (CPAwO)

Notons que les trois approches ont le même comportement durant les premières 300 secondes de la simulation (figures 5.11 à 5.14). Cela est dû au fait que pendant ce temps les ressources radio sont disponibles dans le réseau pour l'ensemble des bearers arrivés. Autrement dit, quand les ressources sont disponibles, toutes les requêtes d'établissement ou de modification de bearer seront satisfaites. Or, une fois que les ressources commencent à s'épuiser, le processus d'admission de bearers allouera de moins en moins de ressources pour les bearers de faibles priorités, jusqu'à leur blocage.

D'autre part, les résultats de cette simulation illustrés à la figure 5.15 montrent que l'algorithme CPA, représenté par l'approche « courteous », améliore la qualité de service des trafics moins prioritaires. Cette amélioration apparaît dans la réduction du nombre de bearers bloqués de type CN non-emergency. En outre, l'approche CPAwO obtient de meilleurs

résultats, en éliminant le nombre des bearers CN bloqués (figure 5.15). Cette amélioration est réalisée grâce à l'utilisation de la radio fréquence des petites cellules. Rappelons que le mécanisme « HOW » permet de basculer les bearers non admis dans la macro cellule à cause du manque de ressources, vers les petites cellules Wlan. Cette technique augmente le nombre des bearers admis et arrive même à accepter toutes les requêtes d'établissement de bearers quand les ressources radio sont disponibles dans les petites cellules. Notons que le basculement des bearers de la macro cellule s'effectue d'une façon toute à fait transparente à l'utilisateur, ainsi, le degré de satisfaction du client quant à la qualité de service reçu par le fournisseur de service ne sera pas affecté.

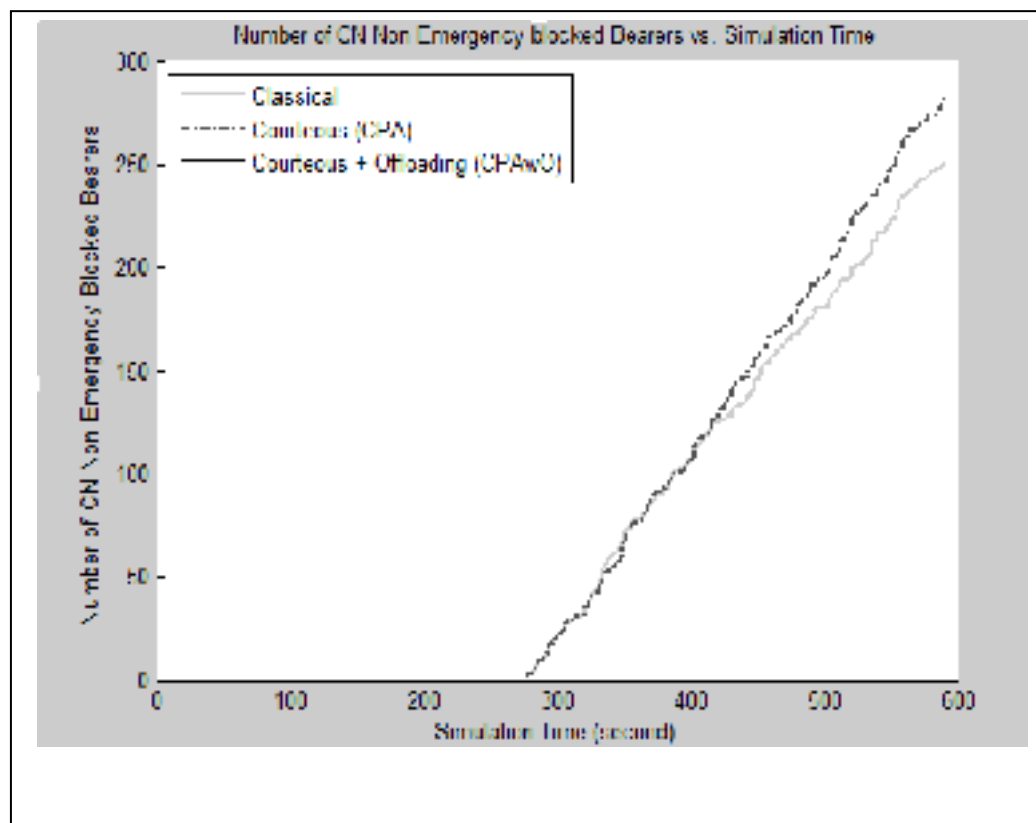


Figure 5.15 Nombre des bearers CN Non-Emergency bloqués (CPAwO)

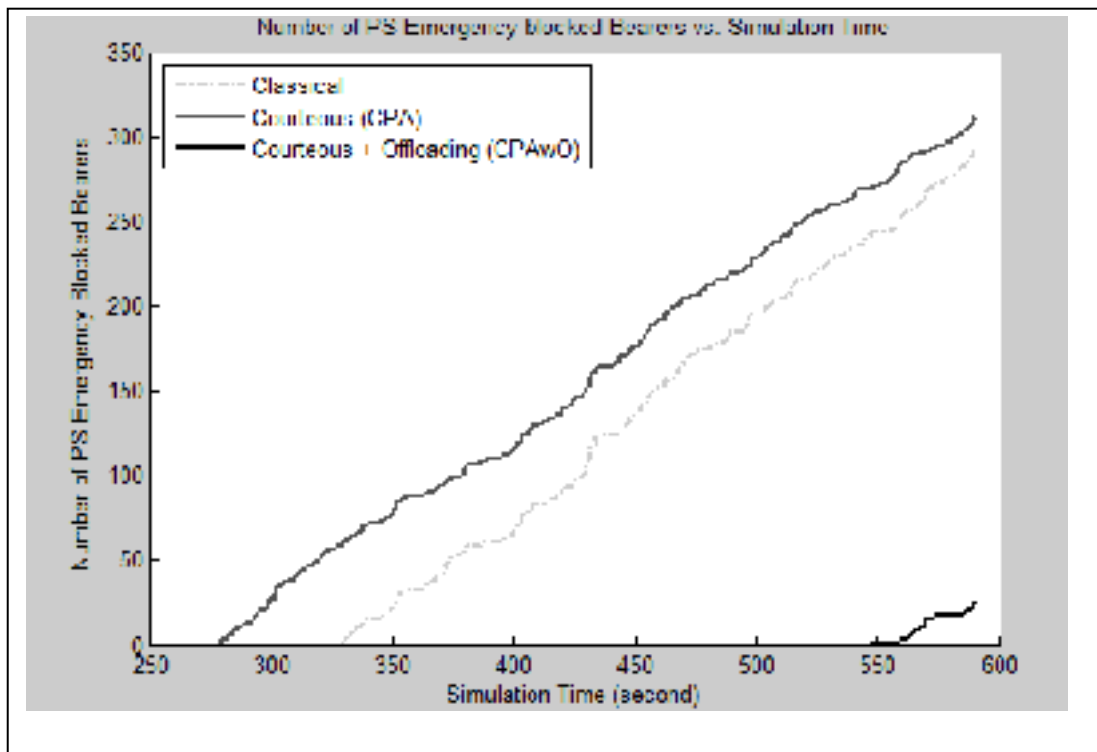


Figure 5.16 Nombre des bearers PS Emergency bloqués (CPAwO)

De la même façon, les deux figures 5.16 et 5.17, nous permettent aussi de voir que le nombre des bearers PS bloqués augmente lors de l'application de la technique CPA. Cette augmentation par rapport à l'approche classique est interprétée par l'effet de la courtoisie qui offre des ressources radio PS aux bearers commerciaux. Cette courtoisie s'applique tant que le taux de blocage des bearers PS n'atteint pas le seuil de blocage toléré pour le réseau PSN. Par ailleurs, la solution CPAwO vient renforcer le système CPA en offrant les avantages de celui-ci tout en évitant ses inconvénients. Ceci dit, CPAwO ne se contente pas d'éviter l'accroissement du nombre de bearers PS bloqués dans le système, mais il ira jusqu'à réduire considérablement ce nombre par rapport à la méthode classique (figures 5.16 et 5.17) grâce à l'algorithme HOW qui assure des ressources radio supplémentaires en utilisant les

fréquences publiques des réseaux locaux sans fil opérant dans les petites cellules (figure 5.10). Cette contribution est fort pertinente pour l'amélioration de la qualité des transmissions PS et aussi CN dans le réseau LTE HetNet.

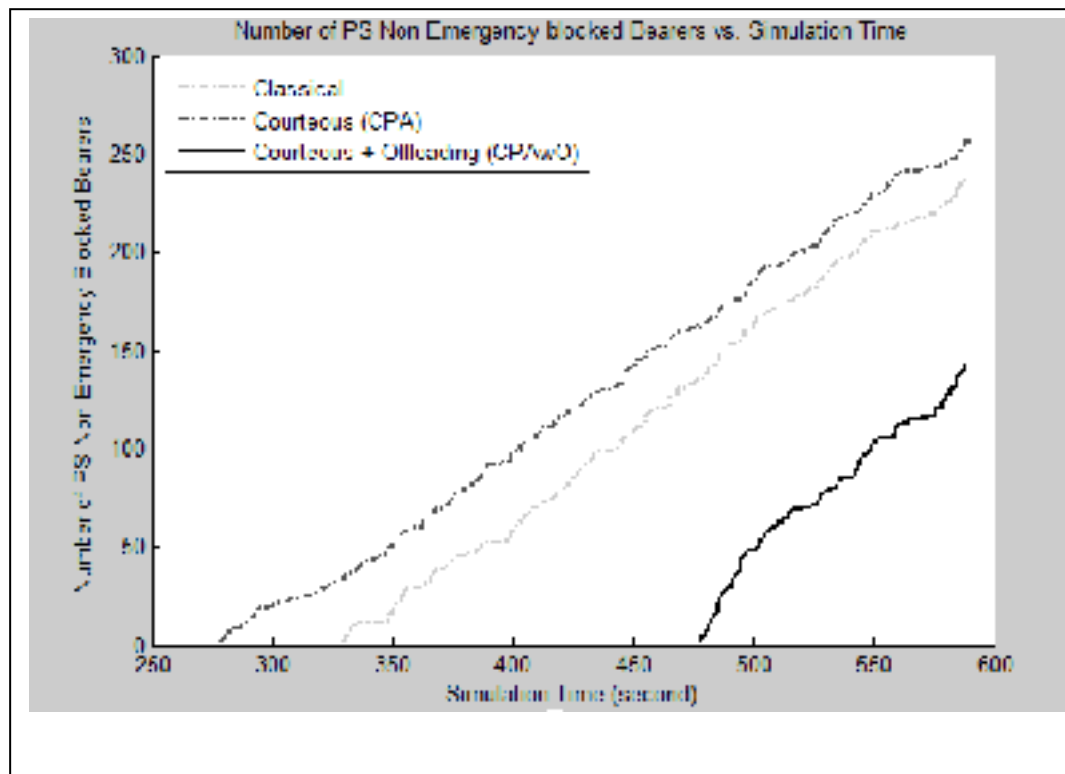


Figure 5.17 Nombre des bearers PS Non-Emergency bloqués (CPAwO)

La figure 5.18 illustre le nombre des bearers interrompus dans le réseau lors de l'application des trois approches, notamment l'approche classique, l'approche CPA (Courteous) et l'approche CPAwO (Courteous + Offloading). Rappelons que l'algorithme CPAwO utilise des ressources radio appartenant à la macro cellules, ainsi que des ressources incluses dans les petites cellules. Par conséquent, la préemption des bearers pourrait avoir lieu dans les fréquences publiques et les fréquences privées, selon la disponibilité des ressources et la priorité des bearers arrivés vers le réseau LTE HetNet. Pour cette simulation nous considérons que seuls les bearers CN Non-emergency sont des bearers vulnérables à la

préemption. Rappelons que certaines conditions s'appliquent pour la préemption d'un bearer actif dans le réseau par un nouveau qui arrivant vers la macro cellule ou même vers la petite cellule. En effet, seuls les bearers de haute priorité peuvent interrompre ceux ayant une basse priorité. De plus, les nouveaux bearers, aptes à interrompre les autres bearers, doivent avoir la valeur Capacity to Pre-emption (CP) équivalente à (+1). D'autre part, les bearers susceptibles d'être interrompus doivent être vulnérables à la préemption. Cette valeur, exprimée par Vulnerability to Pre-emption (VP), doit être égale à (+1).

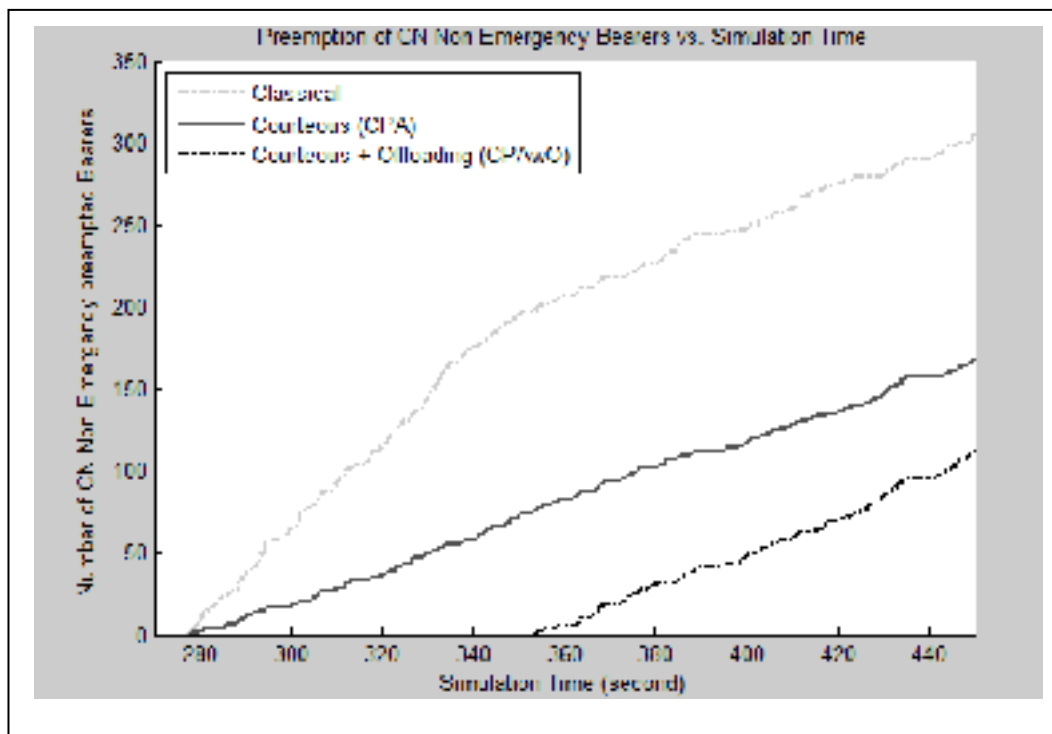


Figure 5.18 Nombre de bearers CN Non-Emergency interrompus (CPAwO)

Par ailleurs, les résultats de la simulation qui sont représentés par les graphes de la figure 5.18 montrent que non seulement l'application de la courtoisie permet de retarder la préemption des bearers CN non-emergency, mais aussi elle diminue le nombre des bearers interrompus dans l'ensemble du réseau LTE pour ce même type de trafic. Ce résultat est obtenu grâce à la réallocation, par courtoisie, de certaines ressources détenues par les usagers plus prioritaires pour les clients moins favorisés à savoir les usagers transmettant des trafics

de type CN non-emergency. De ce fait, la courtoisie se montre comme une approche pertinente pour améliorer la QoS des trafics de basse priorité. En outre, cette technique assure d'éviter toute sous utilisation des ressources radio dans le réseau en procédant à la redistribution des ressources non employées et en évitant toute opération d'accaparement de ressources par certaines classes de trafics plus prioritaires. La courtoisie est donc un moyen de servir les différents usagers du réseau LTE avec équité et discrimination positive. L'équité se montre dans le privilège qu'obtiennent les trafics moins prioritaires en bénéficiant de ressources supplémentaires, en plus de celles qui leur sont réservées préalablement. La discrimination positive est faite d'adopter le principe de priorisation des trafics afin d'assurer l'uniformité des communications exigeantes en termes de QoS. Ces dernières sont considérées dans ce travail comme des communications urgentes.

En effet, la meilleure solution développée dans cette étude relativement à la préemption des bearers actifs dans le réseau est offerte par CPAwO (figure 5.18). Cette solution, grâce à la combinaison des deux techniques ; courtoisie et Offloading, assure un meilleur résultat pour la préemption des bearers actifs dans l'ensemble du réseau LTE HetNet. Rappelons que lorsqu'un nouveau bearer de hautes priorités arrive vers la macro cellule, et que celle-ci ne sera pas en mesure de lui offrir les ressources demandées, CPAwO exécute le mécanisme HOW et bascule ce bearer vers une petite cellule non congestionnée. CPAwO privilégie le basculement du nouveau bearer à l'interruption d'un bearer moins prioritaire déjà actif dans le réseau. De ce fait, la préemption n'aura lieu qu'en moment de pénurie totale des ressources radio dans l'ensemble du réseau LTE HetNet, composé du réseau LTE et des réseaux locaux sans fil. Cette technique assure une certaine stabilité pour les communications courantes en évitant leur interruption. Sans oublier que cette valeur ajoutée pour les deux réseaux PS et CN se réalise sans causer des frais élevés aux fournisseurs de service, étant donné que les ressources supplémentaires fournies par LTE HetNet appartiennent à des fréquences publiques sans licences disponibles dans les petites cellules.

5.6 Conclusion

Dans ce chapitre deux solutions ont été proposées, pour l'amélioration de la gestion des ressources radio commerciales partagée entre les usagers commerciaux et ceux de la sécurité publique. À savoir l'approche CPA et l'approche CPAwO. La dernière technique est une amélioration de la première, sauf que son application nécessite l'intégration des petites cellules Wlan dans le réseau LTE. Autrement dit, CPAwO s'applique au sein du LTE HetNet seulement. L'objectif principal des deux approches est de permettre un accès avec priorité aux premiers répondants sans pénaliser les clients commerciaux. Ceci est applicable grâce au processus de la courtoisie, qui réalloue certaines ressources réservées aux réseaux PSN aux utilisateurs commerciaux. Rappelons que la courtoisie n'est applicable que si les trafics relatifs à la gestion des crises ont un niveau de QoS acceptable alors que les trafics commerciaux ont une QoS détériorée. Les résultats de la simulation ont montré que CPA réduit le nombre des bearers CN bloqués, retarde le début du blocage, ainsi il fait de même pour la préemption des bearers. De plus, il augmente le nombre des bearers commerciaux actifs dans le réseau LTE. Tout cela est réalisable en garantissant un niveau de QoS acceptable pour les premiers répondants au sein du réseau commercial partagé.

Par ailleurs, il est à noter que quand les ressources deviennent limitées pour les clients de la sécurité publique, le processus de courtoisie cesse de s'exécuter et aucune ressource additionnelle ne sera allouée aux bearers commerciaux. Par conséquent, la préemption de ces derniers devient, malheureusement, inévitable. Pour résoudre cette problématique, l'algorithme CPAwO a été conçu. Tel qu'il a été déjà mentionné ci-dessus, CPAwO est une amélioration du CPA afin de renforcer sa contribution en matière de gestion efficace des ressources radio. La valeur ajoutée de CPAwO par rapport à CPA est l'utilisation du mécanisme d'Offloading how, nommé HOW, afin de décharger les macro-cellules LTE. Ce déchargement se traduit par le transfert des nouveaux bearers arrivant à une macro cellule congestionnée vers une petite cellule contenant des ressources radio disponibles. La contribution principale de CPAwO par rapport à CPA est le fait de fournir des ressources radio supplémentaires et gratuites quand les macro-cellules deviennent congestionnées et

quand la courtoisie cesse de fonctionner. Les résultats de la simulation ont montré que les ressources supplémentaires offertes par les petites cellules ont contribué à l'amélioration de la QoS de l'ensemble des trafics dans le réseau LTE HetNet y compris les trafics moins prioritaires. Les résultats illustrent que la gestion des ressources radio est nettement efficace par rapport aux deux approches CPA et la méthode classique.

Un autre volet de notre solution présentée dans ce chapitre pour l'amélioration de la gestion des ressources radio commerciales partagées dans le réseau LTE est le développement de trois modèles de gestion de l'allocation des fréquences avec contraintes. Ces modèles nommés CAMF, GCAMF et RUS-CAMF sont conçus respectivement pour la gestion d'allocation avec contraintes pour deux classes de trafics, n classes de trafics, $n \in \mathbb{R}$, et quatre classes de trafic. Le CAMF définit les contraintes d'allocation des ressources de fréquences pour les deux réseaux PSN et CN. Le G-CAMF est une généralisation de CAMF pour définir les contraintes d'allocation de ressources à n classes de trafics $n \in \mathbb{R}$. Finalement, le RUS-CAMF est une application du modèle G-CAMF à $n = 4$. L'objectif du RUS-CAMF est de prendre en considération la situation de l'utilisation des ressources comme contrainte, lors de l'allocation des ressources radio. Autrement dit, un appel de type PS établi lors d'un désastre ne doit pas avoir la même priorité qu'un appel de même type établi par un premier répondant afin d'effectuer un test de routine d'un système de sécurité résidentiel. Ces trois systèmes d'allocation de ressources avec contraintes sont utilisés par les algorithmes CPA et CPAwO lors de l'allocation des ressources. Leur contribution consiste dans la garantie d'une portion de fréquence radio pour tous les types de trafics en partageant les ressources disponibles d'une façon à répondre aux exigences de chacun d'eux en terme de QoS.

D'autre part, rappelons que la solution CPAwO proposée dans ce chapitre pour la gestion efficace des ressources radio commerciales partagées se repose en partie sur le basculement des bearers vers les réseaux sans fil locaux. Or, les réseaux WLAN, tels que Ad hoc, WMN et WiFi n'offrent point le même niveau de sécurité que celui assuré par LTE. De ce fait, sécuriser les communications D2D, effectuées dans les petites cellules suite au processus

d'offloading devient un défi difficile à relever, surtout que certains bearers transférés vers les fréquences publiques sont critiques, car ils dépendent de la gestion des désastres et ont un caractère hautement confidentiel. Un autre point aussi important à traiter est le routage au sein des petites cellules. La problématique à résoudre consiste en l'amélioration de la signalisation dans les réseaux locaux sans fil afin de rendre efficace la gestion de la bande passante disponible.

Le chapitre suivant propose une solution novatrice pour la sécurité des communications D2D sans ajouter de trafic de contrôle supplémentaire. Deux nouvelles approches pour le routage dans les réseaux locaux sans fil ont été aussi développées.

CHAPITRE 6

Routage efficace et sécurité pour améliorer les transmissions D2D dans le réseau de sécurité publique sur LTE HetNets

6.1 Introduction

La technologie Device-to-Device (D2D) des réseaux LTE Hétérogènes (LTE HetNets), apparaît comme une solution efficace pour les communications de sécurité publique dans des zones de crises se trouvant en dehors de la couverture du réseau LTE, ainsi que quand les macro-cellules LTE sont congestionnées et quand leurs ressources radio sont limitées. Les communications D2D utilisent les fréquences publiques, gratuites, des réseaux locaux sans fil comme médium de transmission. Toutefois, il est connu que les réseaux locaux sans fil ne sont pas dotés d'un même niveau de sécurité dont jouissent les réseaux cellulaires. Rappelons que plusieurs communications échangées dans les petites cellules LTE HetNets sont transférées à partir des macro-cellules. De plus, elles peuvent être des communications extrêmement confidentielles établies par les premiers répondants du réseau PS. Par conséquent, le développement d'une solution de sécurité pour les communications D2D devient primordial afin d'assurer la confidentialité, l'intégrité et la disponibilité de l'information. D'autre part, le routage au sein des réseaux WLAN présente toujours un défi à soulever pour une meilleure gestion des ressources du réseau, qui garantit un bon niveau de QoS pour l'ensemble des classes de trafics transférés dans les petites cellules.

Dans ce chapitre, une nouvelle approche a été conçue pour assurer un routage efficace, robuste et sécurisé pour les flux de données D2D dans les petites cellules LTE HetNets. Cette solution est baptisée Generalized Secure Network Coding based data Splitting (G-SNCDS). Elle est conçue principalement pour les réseaux WMN. L'algorithme G-SNCDS repose sur deux approches, à savoir, le codage réseau pour l'amélioration de la QoS dans le réseau WMN, ainsi que le principe du découpage des paquets de données (Data Splitting) afin d'améliorer le niveau de sécurité dans ce même réseau. Dans le but de garantir le processus

de codage-décodage des paquets, un nouvel algorithme a été développé pour la construction des effets papillon dans le WMN, nommé Reliable Butterfly Construction (RBC). Une deuxième approche, nommée Load Balancing based Selective On-Demand Multipath Distance Vector (LBS-AOMDV), a été implémentée comme solution alternative à RBC dans le cas où les réseaux papillons n'existent pas dans le WMN. L'objectif de LBS-AOMDV est de déterminer un multipath dans le WMN reliant la source de données à sa destination et assurant la continuité des communications D2D dans le WMN.

6.2 Routage efficace et sécurité pour les communications D2D

Certainement, les communications D2D opérant dans des réseaux locaux sans fil ne bénéficieront point du même niveau de sécurité ni de QoS que les communications LTE. Ceci est causé par les caractéristiques de la technologie des réseaux sans fil WLAN. Dans le but d'améliorer le niveau de sécurité, ainsi que le niveau de la QoS pour les communications D2D, trois solutions sont proposées, à savoir RBC et LBS-AOMDV comme solutions de routage dans les réseaux WMN, ainsi que G-SNCDS, incluant SNCDS (Secure Network Coding based data Splitting), comme solution de sécurité.

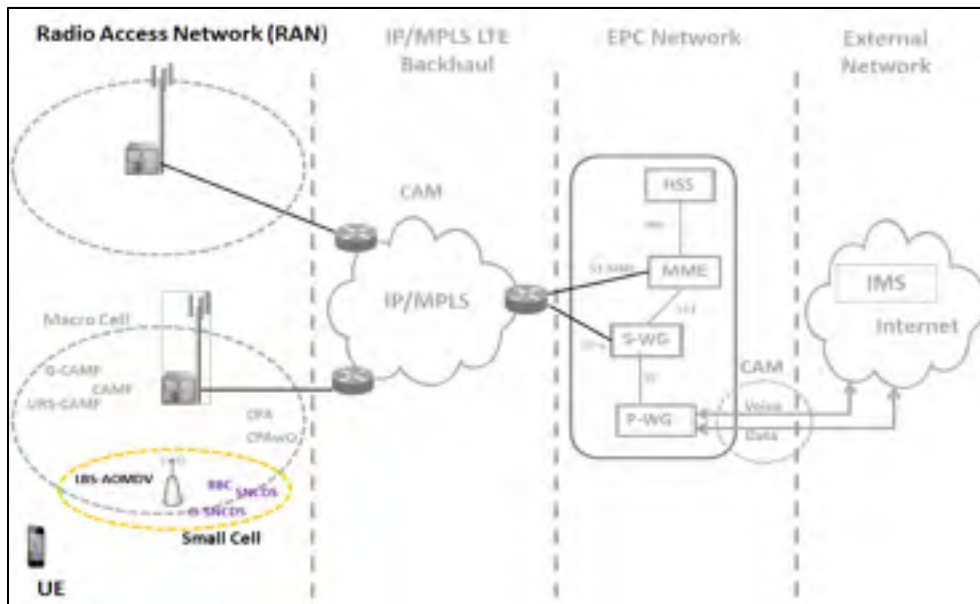


Figure 6.1 Amélioration du routage et de la sécurité pour les réseaux D2D

La figure 6.1 montre l'ensemble des algorithmes développés pour améliorer le routage et la sécurité pour les communications D2D, ainsi que la position de ces solutions par rapport à l'approche globale implémentée dans cette thèse pour l'amélioration de la performance des réseaux de la sécurité publique dans les réseaux LTE HetNets.

La première partie de notre approche consiste en une solution de routage dans le réseau local sans fil au sein d'une petite cellule LTE HetNet. Le réseau sans fil considéré pour notre système de signalisation est un réseau maillé sans fil (WMN). D'autres réseaux sans fil peuvent être utilisés, tels qu'un WiFi ou un Ad hoc. Le but est de trouver un ensemble de routes reliant une source de données, nommé S à une destination D. Pour ce faire, deux algorithmes de signalisation ont été proposés. Le premier, nommé Reliable Butterfly effect Construction (RBC) (Tata and Kadoch 2013) et le deuxième est appelé Load Balancing Based Selective Ad hoc On-Demand Multipath Distance Vector (LBS-AOMDV) (Tata and Kadoch 2014).

L'algorithme RBC permet de construire un ou plusieurs effets papillon (Butterfly effects) au sein d'un réseau sans fil WMN. Le but de cette opération est de réaliser un routage de données de S vers D, basé sur le codage réseau. Les études menées dans les articles (Ho, Koetter et al. 2003, Li and Li 2004, Gkantsidis and Rodriguez 2005, Fragouli, Widmer et al. 2006, Katti, Rahul et al. 2006, Matsuda, Noguchi et al. 2011) ont démontré la contribution du codage réseau pour l'amélioration des performances des réseaux de télécommunications filaires et sans fil. Ces performances s'illustrent, entre autres, dans l'augmentation du débit, la réduction des délais de bout en bout, ainsi que dans la diminution des taux de pertes de paquets. Toutefois, à défaut de recevoir tous les paquets nécessaires pour le décodage de l'information, le décodeur se trouve incapable de mener à bien sa mission de décodage. En effet, ce ne sont pas toutes les topologies réseau qui permettent de fournir toute l'Information requise pour réussir le décodage. En fait, une parmi les topologies qui garantissent de réussir le processus de codage et décodage réseau est l'effet papillon (Ahlsvede, Cai et al. 2000). La première contribution de RBC est de construire un ou plusieurs effets papillon dans un réseau

WMN. D'autre part, l'utilisation de RBC permet non seulement de trouver plus d'un chemin reliant la source à la destination, offrant ainsi une solution de load Balancing et aussi de redondance pour la transmission de données, mais il assure la construction de plusieurs effets papillons, garantissant donc l'implantation d'un processus de Back-up pour le routage dans le réseau WMN. Cette dernière caractéristique du RBC lui donne la qualification de robustesse.

Par ailleurs, il est à noter que la construction d'un effet papillon à l'intérieur d'un réseau WMN n'est pas toujours possible. La structure d'un réseau papillon exige la connexion des différents nœuds voisins de telle sorte à former une topologie bien particulière. La figure 2.1 du chapitre 2 illustre un modèle d'un réseau papillon. Ceci dit, dans certaines situations, l'algorithme RBC se trouve dans l'impossibilité de construire des effets papillon. Par conséquent, le routage de l'information qui doit s'effectuer entre les deux nœuds du réseau, Set D ne se réalisera point. Afin de pallier ce problème, une autre alternative a été proposée dans cette recherche afin d'éviter l'interruption des communications D2D. Il s'agit de l'algorithme LBS-AOMDV qui permet de construire un multipath au sein du réseau WMN. Tout comme RBC, LBS-AOMDV permet d'effectuer un load balancing, ainsi qu'une redondance lors de la transmission de données en utilisant deux ou plusieurs chemins du Multipath. L'existence de plusieurs chemins reliant S à D garanti le processus du Back-Up en cas de rupture d'un ou des chemins de transmission.

L'autre partie de notre solution pour l'amélioration des communications D2D, consiste en la sécurité de l'information transmise entre les différents nœuds du réseau WMN. Il est à noter que la solution de sécurité proposée dans ce chapitre est adaptable seulement pour le routage RBC. Nos travaux futurs vont être consacrés en partie pour compléter notre solution de sécurité des communications D2D afin de l'adapter aussi au routage LBS-AOMDV. Deux algorithmes ont été développés dans cette recherche pour résoudre la problématique de sécurisation des flux LTE transmis à travers les WMN. La première solution est appelé Secure Network Coding based Data splitting algorithm (SNCDS) (Tata and Kadoch 2014). La contribution principale du SNCDS est de construire un mécanisme de codage réseau basé sur la fraction et le mélange des données à la source de transmission. L'objectif de cette

technique est d'empêcher les attaques de confidentialité interne comme externe au niveau des réseaux papillon. La deuxième solution est nommée Generalized Secure Network Coding based Data splitting algorithm (G-SNCDS) est une extension de l'algorithme SNCDS afin d'inclure la résolution des attaques d'intégrité et de disponibilité au sein d'un réseau papillon. Il est à noter que l'algorithme SNCDS est une partie intégrante du G-SNCDS, il s'agit du module qui traite les attaques de la confidentialité. Par conséquent, dans ce chapitre on ne développera pas une section à part pour SNCDS. On détaillera tout simplement l'approche G-SNCDS. Par contre, dans la partie simulation les résultats des deux approches, notamment SNCDS et G-SNCDS seront implémentées afin de démontrer l'intérêt de généraliser la solution SNCDS pour avoir G-SNCDS.

Les différentes approches citées ci-dessus sont détaillées dans les sections qui suivent.

6.3 Algorithmes de signalisation

Cette section détaille nos deux solutions de signalisation dans le réseau WMN pour améliorer les communications D2D au sein des petites cellules LTE HetNet. La contribution principale des deux solutions, à savoir RBC et LBS-AOMDV consiste à trouver plusieurs routes reliant une source de données S à sa destination D. Ces multiples chemins permettent d'effectuer un routage avec Load Balancing, comme ils assurent la redondance de données et la restauration de la topologie en cas de perte d'un ou plusieurs liens ou nœuds. Le Load Balancing est un processus qui contribue fortement à l'augmentation du débit dans le réseau en transmettant plus d'un paquet simultanément à travers plusieurs chemins. Par conséquent, les délais de bout en bout seront réduits et les taux de perte de paquets seront minimisés. D'autre part, la redondance offre un moyen de vérification de l'intégrité de l'information via des mécanismes de correction et de détection d'erreurs, tels que Forward Error Correction (FEC). Pour sa part, le processus de Back-Up assure la disponibilité de l'information à travers la restauration des chemins de transmission.

6.3.1 RBC

6.3.1.1 Description de RBC

Le codage réseau est une solution efficace pour améliorer le niveau de la QoS au sein des réseaux de télécommunications, que ce soit des réseaux filaires ou sans fil. Le codage réseau est une technique qui permet d'augmenter le débit dans l'ensemble du réseau à travers sa capacité de transmettre plusieurs paquets simultanément alors que le routage conventionnel les transmet chacun son tour. Cette transmission simultanée des paquets réduit le nombre de time-slots utilisés pour la transmission. Par conséquent, les délais de bout en bout, ainsi que les taux de perte de paquets seront réduits. Les avantages d'utiliser le codage réseau pour la transmission de données dans les réseaux de télécommunications ont été largement étudiés par les chercheurs. Certains de ces avantages sont donnés par les articles (Ho, Koetter et al. 2003, Li and Li 2004, Gkantsidis and Rodriguez 2005, Fragouli, Widmer et al. 2006, Katti, Rahul et al. 2006, Matsuda, Noguchi et al. 2011).

Telles que toutes les solutions de routage et de transmission de données existantes, le codage réseau souffre de quelques problématiques. L'un des défis à relever, relativement à cette technologie, est d'établir des routes dans le réseau, qui garantissent la réussite du processus du codage-décodage de l'information transmise. Parfois, même si l'opération de codage s'effectue à merveille, ce ne sera pas le cas, malheureusement, pour le processus de décodage. L'impossibilité de réaliser un décodage est très souvent reliée à la topologie du réseau. Une architecture de réseau qui ne détient pas le nombre de routes requises par le décodeur pour recevoir ses paquets, ne sera pas en mesure de réussir la mission d'extraction des paquets natifs de ceux qui sont reçus codés. Par conséquent, elle ne sera pas utilisable pour l'établissement du codage réseau.

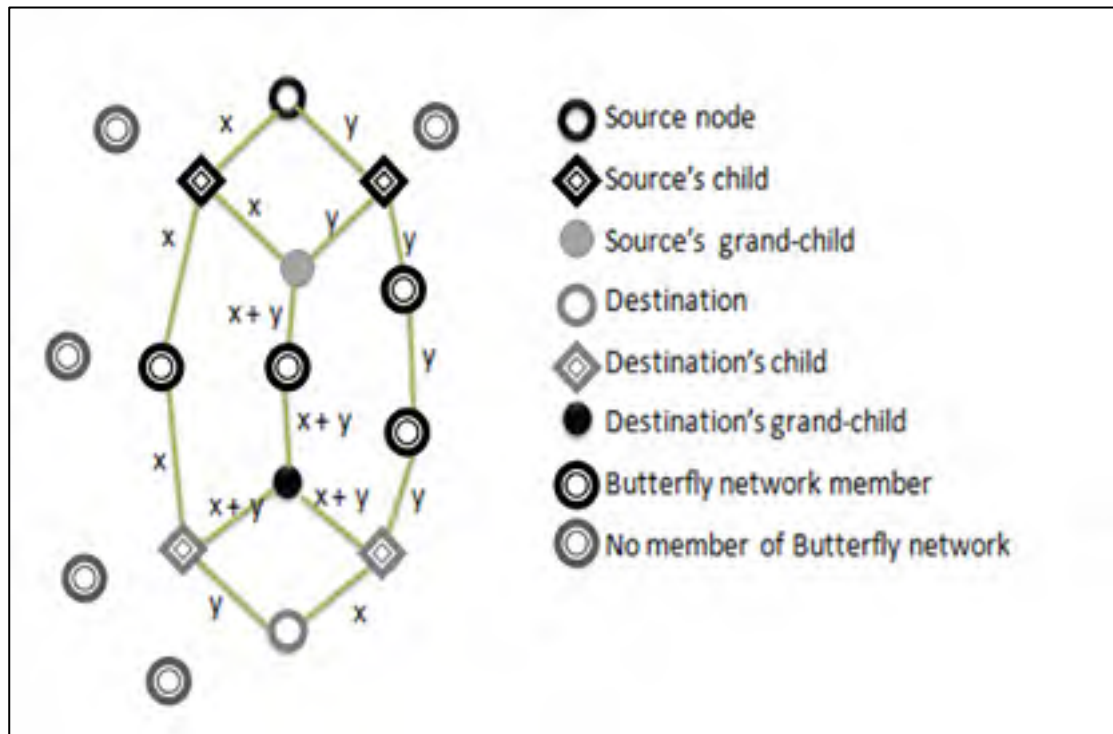


Figure 6.2 Effet papillon dans un WMN sans fil

Une topologie, connue par sa conformité pour la réussite de l'opération du codage réseau est le réseau papillon, appelé aussi effet papillon (figure 6.2). La figure 6.2 montre clairement que les enfants de la destination, qui jouent le rôle de décodeurs dans cette architecture de réseau, reçoivent les paquets entrant via deux chemins distincts. Un premier chemin leur fournit le paquet codé $x \oplus y$ et le second leur transmet un paquet natif x parmi les deux qui sont inclus dans le paquet codé. Une opération XOR (ou logique exclusive) entre les deux paquets reçus $x \oplus y$ et x permet, d'extraire le paquet y .

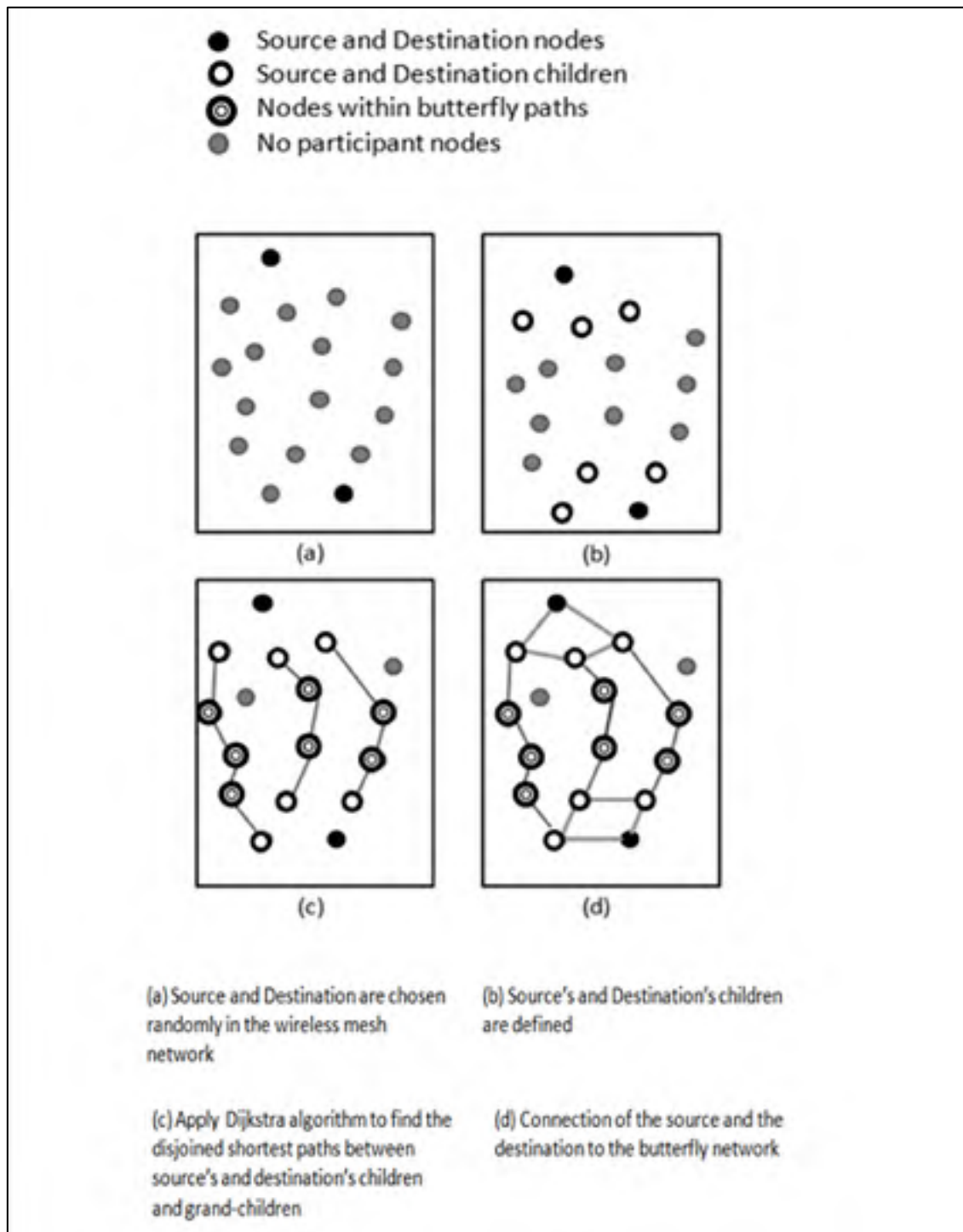


Figure 6.3 Étapes de construction d'un réseau papillon dans un WMN

L'algorithme RBC est notre solution pour construire des effets papillon dans un réseau WMN, afin de garantir la réussite du processus de codage réseau. Le modèle réseau considéré pour expliquer l'approche RBC est représenté par la figure 6.2. Dans cette topologie, les voisins directs de la source et de la destination sont appelés « source's children » et « destination's children », respectivement. De plus, les voisins à deux sauts de la source et de la destination sont appelés « source's grand-children » et « destination's grand-children », respectivement. Notons que certains nœuds appartiennent au WMN ne peuvent toutefois être inclus dans le réseau papillon. Le nœud codeur, respectivement décodeur, est un nœud parmi les nœuds grand-children de la source, respectivement de la destination.

La figure 6.3 illustre les différentes étapes adaptées par l'algorithme RBC pour construire un réseau papillon (Butterfly effect) dans un réseau WMN sans fil.

RBC repose sur le principe de déterminer deux chemins reliant deux nœuds source's children à deux nœuds destination's children, dont au moins un nœud enfant commun aux deux nœuds source's children choisis, peut être relié à au moins un nœud enfant commun aux deux nœuds destination's children choisis. Les trois chemins principaux du butterfly seront donc formés. Reste à relier chaque parent à son enfant afin de former le réseau papillon.

Les étapes de l'algorithme RBC sont données comme suit:

Soit $G = (V, E)$ le graphe représentant le réseau WMN. V étant l'ensemble des nœuds du WMN et E l'ensemble des liens entre chaque deux nœuds voisins.

Soit R le rayon de transmission des différents nœuds du réseau WMN. Un nœud u est donc considéré comme un voisin d'un autre nœud v si et seulement si la distance entre u et v est inférieure ou égale à R . Si on exprime cette distance par la fonction $Distance(v, u)$ on pourra écrire

$$v \in V, u \in V: Distance(v, u) \leq R \rightarrow e = (v, u) \in E \quad (6.1)$$

ÉTAPE 1

Déterminer le nœud source S et un nœud destination D parmi les nœuds de l'ensemble V .

ÉTAPE 2

Trouver tous les nœuds enfants de S

$$C(S) = \{v \in V, \text{Distance}(v, S) \leq R\} \quad (6.2)$$

Notons que $\text{Distance}(v, S)$ est la distance entre v et S

ÉTAPE 3

Pour tout $u, v \in C(S)$ trouver l'ensemble des enfants communs de u et v , nommé $Cc(u, v)$.

$$Cc_s(u, v) = C(u) \cap C(v) \quad (6.3)$$

Chaque nœud $w \in Cc_s(u, v)$ a deux parents communs, formant ainsi l'ensemble $Pr_s(w)$.

$$Pr_s(w) = \{u, v \in C(S)\} \quad (6.4)$$

L'ensemble des petits enfants de S (grand-children of S), ayant chacun deux parents est nommé $GCc(s)$.

$$GCc(s) = \bigcup_{u, v \in C(S)} Cc_s(u, v) \quad (6.5)$$

ÉTAPE 4

Répetons les étapes de 1 à 3 pour la destination D afin d'établir les formules et les ensembles suivants :

$$C(D) = \{v \in V, \text{Distance}(v, D) \leq R\} \quad (6.6)$$

$$Cc_d(u, v) = C(u) \cap C(v) \quad (6.7)$$

$$GCc(D) = \bigcup_{u,v \in C(D)} Cc_D(u, v) \quad (6.8)$$

$$Pr_d(w) = \{u, v \in C(D)\} \quad (6.9)$$

ÉTAPE 5

Soit $\alpha = 0$, où α est le nombre de réseaux papillon dans le réseau WMN, initialement mis à zéro.

ÉTAPE 6

Pour chaque nœud $\in GCc(s)$

1. Trouver le chemin le plus court qui relie $w_1 \in GCc(S)$ à $w_2 \in GCc(D)$, noté e_1 ;
2. Si e_1 exists, alors trouver, e_2 , le lien le plus cour reliant $u_1 \in Pr_s(w_1)$ to $u_2 \in Pr_d(w_2)$;
3. Si e_1 et e_2 existent, alors trouver e_3 , le plus court chemin reliant $v_1 \in Pr_s(w_1)$ à $v_2 \in Pr_d(w_2)$ Où $v_1 \neq u_1$ and $v_2 \neq u_2$;
4. Si $i = \{1,2,3\}$ existent et sont disjointe, alors soit $\alpha = \alpha + 1$
5. Construire le $\alpha^{ième}$ réseau papillon, en reliant
 - a. w_1 à u_1 et v_1 , u_1 and $v_1 \in Pr_s(w_1)$
 - b. w_2 à u_2 et v_2 , u_2 and $v_2 \in Pr_d(w_2)$
 - c. S à u_1 et v_1
 - d. D à u_2 et v_2
6. Représenterle $\alpha^{ième}$ réseau papillon comme suit

$$G_{Bfly}^\alpha = (V_{Bfly}, E_{Bfly}) \quad (6.10)$$

$$V_{Bfly} = \{S, D, w_1, w_2\} \cup Pr_s(w_1) \cup Pr_d(w_2) \cup \{V_1, V_2, V_3\} \quad (6.11)$$

$$V_i = \{v, (u, v) \in e_i, i = 1,2,3\} \quad (6.12)$$

Tel qu'il a déjà été mentionné, la principale contribution de l'algorithme RBC est de trouver un ou plusieurs effets papillon, reliant la source de donnée à sa destination, à l'intérieur du réseau WMN sans fil. En effet, la disponibilité de plusieurs réseaux papillon simultanément, assure la restauration des chemins de transmission en cas de brise lien ou d'un Butterfly Failure. On entend par Butterfly Failure la perte de certains liens qui font perdre à la topologie du réseau son allure de papillon. D'autre part, l'utilisation d'un seul effet papillon pour la transmission avec codage réseau permet l'échange d'information avec redondance, en envoyant le même paquet sur les deux chemins partant de S, comme il peut assurer le Load Balancing, qui être réalisé en utilisant deux réseaux papillons pour la transmission des paquets, de façon à équilibrer la charge dans le réseau global. Par ailleurs, RBC apporte une valeur ajoutée à notre recherche menée pour l'amélioration des communications D2D. En effet, d'une part, RBC assure un routage avec codage réseau dans le WMN, ce qui permet de tirer profit des avantages de ce mécanisme pour l'amélioration de la performance dans le WMN. D'autre part, il permet l'application de notre solution de sécurité G-SNCDS, qui inclut le SNCDS, pour faire face aux attaques de confidentialité, d'intégrité et de disponibilité.

6.3.1.2 Simulations et résultats

Des simulations ont été effectuées dans cette thèse afin de montrer l'efficacité du RBC dans la construction des effets papillon dans un réseau WMN sans fil. Ces simulations ont été réalisées avec Matlab. Pour les différents scénarios, 80 Nœuds ont été générés dans une surface de 800 m x 800 m. L'emplacement des nœuds dans l'aire de travail a été généré aléatoirement pour chaque scénario. Le rayon de transmission est de l'ordre de 250 m. Par ailleurs, pour chaque scénario, la source et la destination ont été choisies aléatoirement parmi les nœuds du WMN.

Afin d'obtenir des résultats crédibles, on a répété nos expériences dix fois. Pour simplifier la l'apparence des résultats, on à effectuer des simulations pour construire soit un, soit deux réseaux papillons.

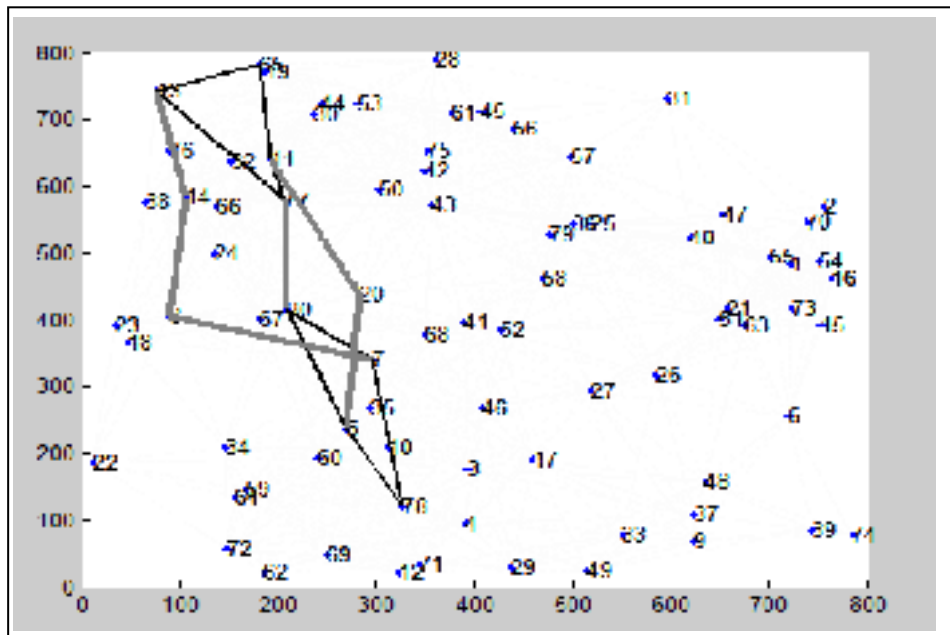


Figure 6.4 Réseau papillon construit avec RBC: Exemple 1

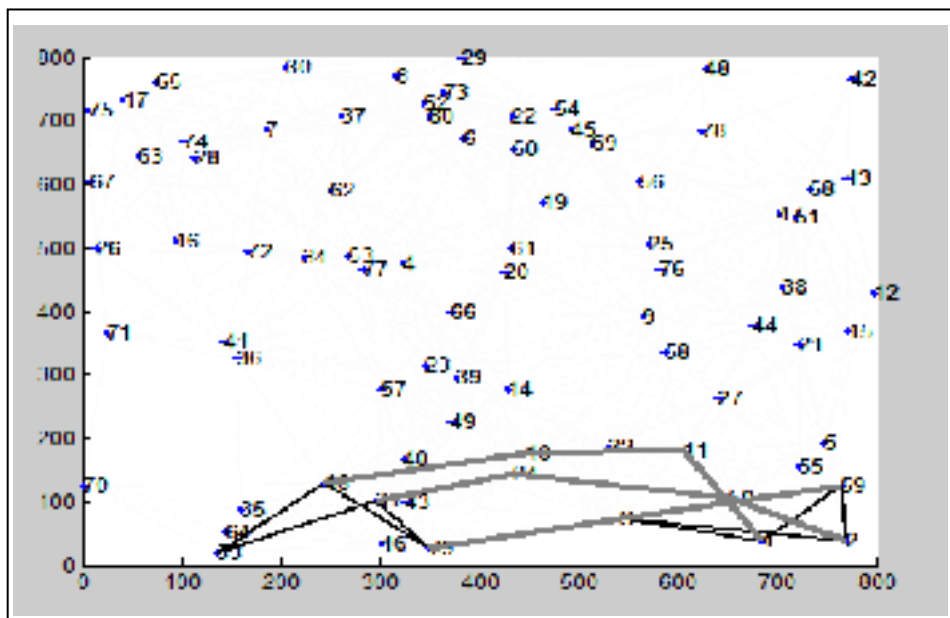


Figure 6.5 Réseau papillon construit avec RBC: Exemple 2

La figure 6.4 et 6.5 montrent la construction d'un réseau papillon dans deux topologies WMN différentes et avec des paires (source, destination) différentes. Cette différence est illustrée par

l'emplacement des nœuds et la distance entre chaque couple de nœuds. Encore une fois, pour la clarté des figures, on n'a pas représenté les liens entre les nœuds WMN. Seules les arêtes du réseau papillon y apparaissent.

Il est important de mentionner que l'existence d'un effet papillon dans un réseau sans fil WMN dépend catégoriquement de la disponibilité de trois routes distinctes, reliant les nœuds fils de la source S aux nœuds fils de la destination D et le nœud petit fils (grand-child) de s au petit fils de D. Ces trois chemins forment ce qu'on appelle le réseau papillon cœur. Tel qu'il est représenté dans les deux figures 6.4 et 6.5, l'algorithme RBC établit trois chemins distincts pour former chaque réseau papillon cœur, l'extension de ces routes vers S et D, a formé par la suite l'effet papillon. Il est à noter que RBC utilise l'algorithme Dijkstra afin de trouver les chemins du réseau papillon cœur.

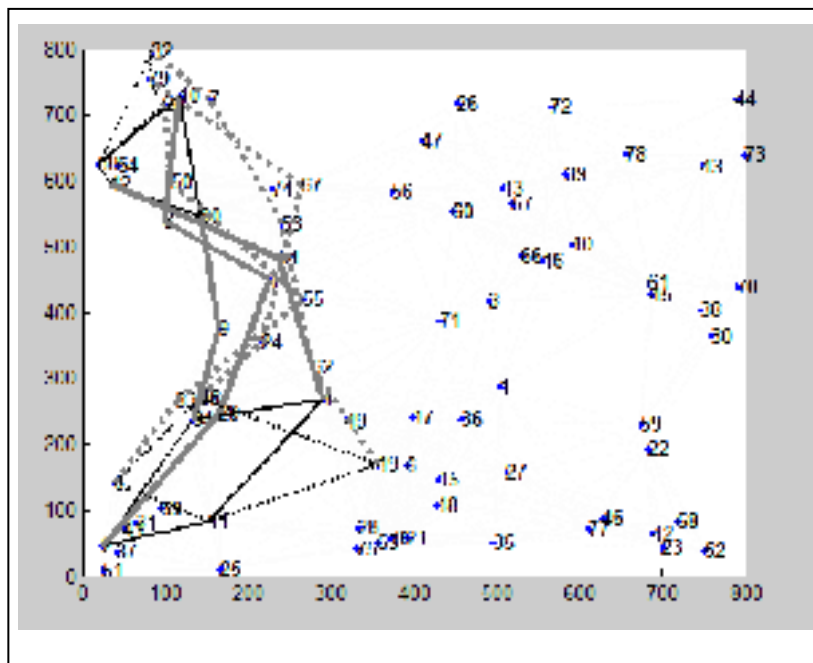


Figure 6.6 Primary and Backup Butterfly effects constructed with RBC algorithm

En outre, les résultats de nos simulations démontrent non seulement que RBC permet de construire un effet papillon dans une topologie de réseau WMN, mais aussi de former plusieurs réseaux papillon tout à fait disjoints, au sein d'un même réseau WMN et reliant une même source à une destination (figure 6.6). L'objectif derrière ce mécanisme est de fournir une solution qui assure l'application du Load Balancing lors de la transmission de données avec codage réseau. Rappelons que la réalisation du Load Balancing nécessite l'utilisation de deux ou plusieurs effets papillon.

La construction de plus d'un réseau papillon permet aussi la restauration de la topologie en cas de sa perte. Les deux effets papillon illustrés dans la figure 6.6 représentent un modèle parfait pour la restauration de la topologie du réseau de transmission. À titre d'exemple, le réseau ayant des liens continus peut être considéré comme le réseau principal, par contre, celui ayant des liens discontinus sera considéré comme le réseau Back-Up. Ces deux réseaux étant établis simultanément peuvent garantir la restauration de la topologie sans avoir à recalculer de nouvelles routes. À partir de là un nouvel avantage peut être ajouté à la liste des bénéfices du RBC. IL s'agit ici de fournir une communication uniforme aux utilisateurs D2D. Autrement dit, en cas de butterfly failure, les communications D2D bascule automatiquement vers le réseau BackUp sans délai significatif, étant donné que le réseau existe et que RBC n'est pas appelé à recalculer les chemins de S vers D formant un effet papillon. En fait, c'est cette caractéristique qui rend l'algorithme RBC fiable d'où son nom « Reliable ».

Un résultat important de l'implémentation de RBC, dont on a déjà cité est sa capacité d'appliquer le mécanisme Load Balancing dans le WMN afin d'augmenter le débit. En effet, il est évident que l'utilisation de plusieurs effets papillonne. Dans ce cas la charge du réseau WMN sera partagée entre les réseaux papillon participants au routage de l'information. Une autre façon de faire est possible pour réaliser un Load Balancing si l'application du RBC n'offre qu'un seul réseau papillon. La figure 6.7 confirme cette proposition. Cette figure montre une topologie d'un réseau papillon construite au sein d'un réseau WMN. Lors du codage réseau, la source de données va transmettre des paquets vers la destination. Afin de

réaliser un load Balancing sur un seul réseau papillon il faut envoyer une partie des paquets sur vers un premier fils de S et l'autre partie du trafic vers le second fils de S. Dans la figure 6.7 la première partie du flux est représentée par les flèches pleines, l'autre partie est illustrée par des flèches à tirer et le trafic codé par des flèches en pointillés. L'application du Load Balancing augmente donc le débit dans le réseau du moment que plus d'un chemin est utilisé pour la transmission de données.

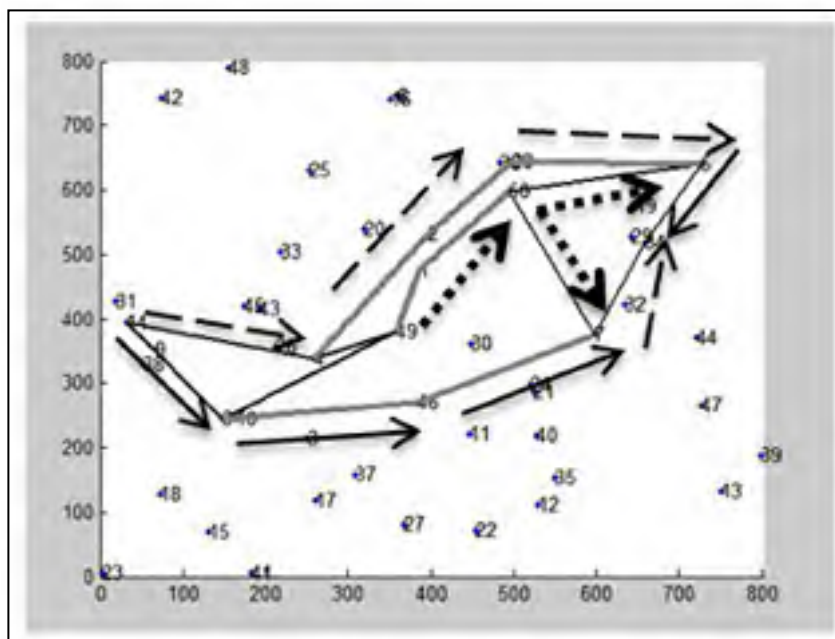


Figure 6.7 Butterfly based- load balancing in the WMN

La figure 6.8 décrit le processus de restauration de la transmission de paquets après un butterfly failure. Le réseau WMN représenté par la figure 6.8 contient deux réseaux papillon. Un réseau papillon principal illustré en ligne pleine et un réseau de restauration représenté par des lignes en tiret. Deux types de transmissions sont présents dans le schéma de la figure 6.8. Une transmission principale de paquets via le réseau principal. Celle-ci est illustrée avec des flèches en tiret. Une deuxième transmission de paquets, dite de restauration (Back-Up Transmission en anglais). Celle-là est représentée par des flèches pleines. Initialement, la transmission de paquets s'effectue en utilisant les routes du réseau principal. Dès qu'une rupture de lien survient, un Butterfly failure aura lieu. La rupture de lien est représentée par

une croix sur le schéma de la figure 6.8. Cette situation empêche la continuité du processus du codage réseau dans le réseau principal. Donc, la restauration de la topologie devient nécessaire. À ce moment-là, RBC intervient afin de rendre disponible une architecture réseau qui garantit l'uniformité de la communication. Par conséquent, la transmission bascule du premier réseau vers le réseau de Back-Up. Un point fort de la solution offerte par RBC est le fait que le passage du réseau principal vers le réseau Back-Up s'effectue sans causer de temps additionnel. Cela est possible étant donné qu'à chaque invocation de RBC, plusieurs réseaux papillon seront déterminés. Certains seront utilisés comme des réseaux principaux de transmission, les autres seront maintenus pour une éventuelle restauration.

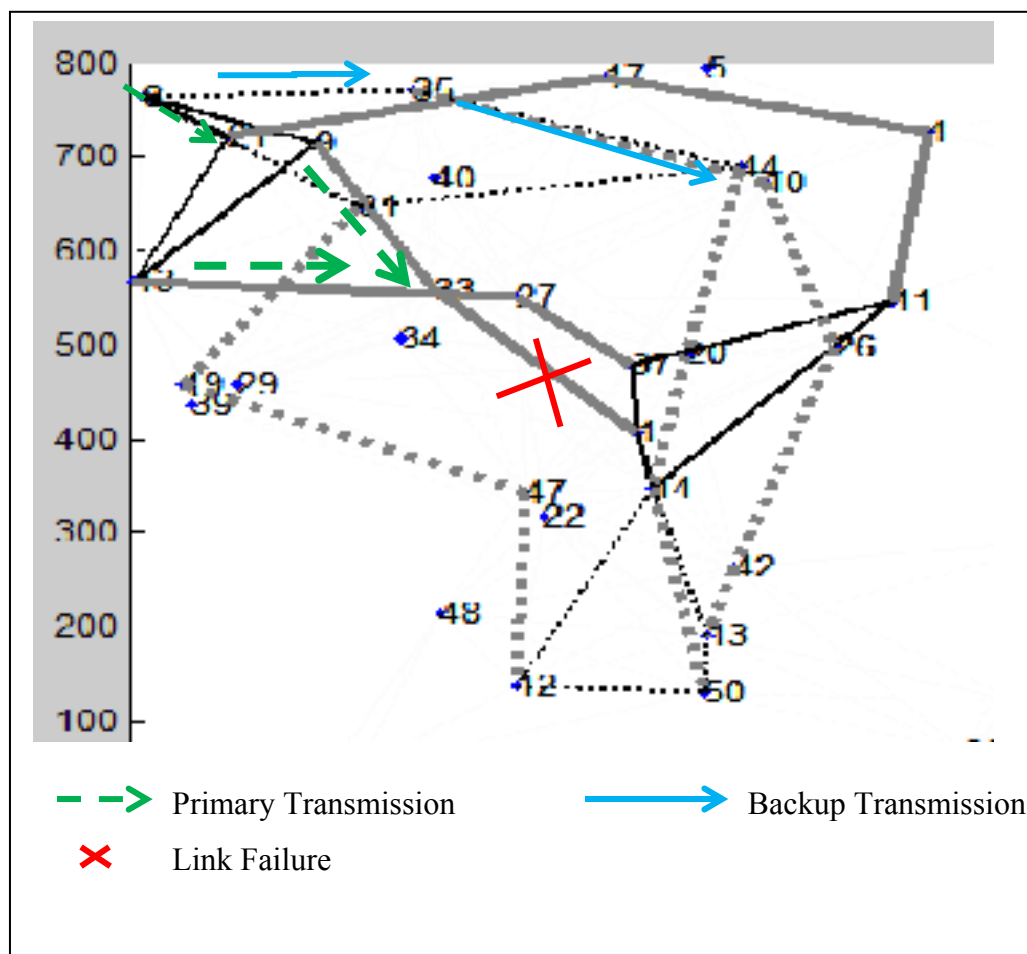


Figure 6.8 Backup après Butterfly failure

Les résultats illustrés ci-dessus ressortent les avantages de RBC quant à l'application du load Balancing, la redondance et aussi quant à la restauration de la communication après un Butterfly Failure. Or, dans certaines situations la construction d'un effet papillon dans le réseau WMN se trouve impossible à cause de l'emplacement des différents nœuds du réseau qui empêche de trouver des routes formant le réseau papillon. Cette situation empêche la transmission de données, à cause de l'absence d'un protocole de routage dans le réseau.

Dans cette thèse, on propose une nouvelle solution pour le routage dans les réseaux WMN afin d'assurer les communications D2D. Il s'agit de notre algorithme LBS-AOMDV publié dans (Tata and Kadoch 2014). Cette solution offre la possibilité de construire un multipath dans le réseau WMN. Ce multipath reliant la source S à sa destination D, est node-disjoint, autrement dit, tous les nœuds du multipath sont disjoints. Cette architecture permet d'assurer les deux fonctions le Load Balancing et la redondance. En plus de la possibilité de la restauration des chemins, tant que LBS-AOMDV peut calculer plusieurs chemins possibles reliant S à D. Le détail de cette approche est donné par la section suivante.

6.3.2 LBS-AOMDV

Les solutions de routage traditionnelles ne procèdent au calcul de nouvelles routes qu'après la rupture des routes utilisées ou lors des congestions des liens. Cela va sûrement causer des délais supplémentaires qui affecteront sans doute la QoS des trafics transmis dans le réseau. La résolution de cette problématique est possible avec l'algorithme RBC. Or, son application nécessite l'existence d'un réseau papillon à l'intérieur du réseau WMN. Notre nouvelle solution, nommée LBS-AOMDV, est une approche pour la construction d'un multipath dans un réseau WMN, afin d'assurer le routage des données D2D. L'objectif principal de cet algorithme est de permettre l'application du Load Balancing lors des transmissions de paquets. L'autre objectif de cette approche consiste dans la restauration de la topologie de pertes de chemins de routage.

Load Balancing Based Selective Ad hoc On-Demand Multipath Distance Vector (LBS-AOMDV) est un protocole de routage connectant une source de donnée S à une destination D via plusieurs routes distinctes. Il est appliqué pour la communication D2D au sein des petites cellules LTE HetNet. LBS-AOMDV est une amélioration de l'algorithme AOMDV (Ad hoc On-Demand Multipath Distance Vector). Une particularité de cet algorithme, par rapport à AOMDV, est sa capacité d'offrir l'information relative à la quantité de bande passante disponible pour chaque route du multipath. De plus, notre approche réduit la quantité du trafic du contrôle du moment qu'il diffuse moins de requêtes RREQ que l'algorithme AOMDV. Ceci est faisable, du moment qu'au lieu d'inonder le réseau avec les requêtes RREQ, LBS-AOMDV sélectionne les récepteurs de ces requêtes parmi les nœuds aptes à les recevoir. Ce comportement lors de la transmission des requêtes RREQ lui donne le nom de Selectif AOMDV. D'autre part, la construction de plusieurs routes disjointe permet l'application d'un load balancing pour la transmission des données, d'où l'introduction de ces deux termes dans l'appellation de notre solution. Par ailleurs, tel qu'il a été mentionné ci-dessus, la transmission des RREQ s'effectue à travers un processus de sélection des nœuds récepteurs. Avant chaque retransmission de RREQ, un nœud enfants récepteurs est sélectionné par le nœud parent transmetteur. Cet enfant est appelé le Best Child. La définition détaillée du nœud Best Child sera donnée dans ce qui suit.

6.3.2.1 Best Child

Le nœud Best Child, c_b , du nœud parent p_r est celui qui, parmi tous ses frères, détient la meilleure quantité de bande passante disponible, notée $BestBW$. Dans ce qui suit, on explique la façon de choisir le nœud Best Child.

Soit p_r un nœud parent d'un ensemble de nœuds membres du réseau WMN, soit $Ch(pr)$ l'ensemble de ses enfants. Avec

$$Ch(pr) = \{ci / i \in N \text{ and } distance(pr, ci) \leq R\} \quad (6.13)$$

On considère que $\text{distance}(pr, ci)$ est la distance entre le nœud parent pr , et son enfant ci . Où, R est le rayon de transmission des nœuds dans le réseau sans fil WMN. Une quantité minimale de bande passante disponible BW_{min} est requise pour qu'un nœud soit candidat à devenir un membre du multipath. Chaque nœud c appartenant au WMN, détient une bande passante disponible notée $BW(c) \geq 0$.

On définit $BestBW$ comme la meilleure bande passante qui puisse être fournie par un nœud enfant de pr . Elle est calculée comme suit :

$$BestBW = \{BW(c_j) / \forall c_i \in Ch(pr), BW_{min} \leq BW(c_j) \leq BW(c_i)\} \quad (6.14)$$

$BestBW$ représente la quantité de bande passante minimale parmi celles qui sont fournies par les enfants du nœud pr et qui sont supérieures ou égales à la bande passante minimale requise BW_{min} . Notons que BW_{min} est la bande passante minimale requise pour un nœud enfant pour qu'il soit candidat à rejoindre le multipath. Le nœud Best Child du nœud parent pr est déterminé par l'ensemble des formules suivantes

If $\exists ci \in Ch(pr), ci \neq cj : BW_{av} \geq BW_{min}$

$$\begin{aligned} BestChild(pr) &= \{cj / cj \in Ch(pr) \cap \forall ci \in Ch(pr), ci \neq cj : \\ BW(ci) > BW_{min} \Rightarrow BW(cj) \leq BW(ci)\} \end{aligned} \quad (6.15)$$

Else If $\forall ci \in Ch(pr), ci \neq cj : BW_{av} < BW_{min}$ Then

$$BestBW = \left\{ \begin{array}{l} \sum_{cj \in Ch(pr)} BW(cj) / cj \in Ch(pr), j \leq i, \quad i = Card(Ch(pr)) : \\ BW_{min} \leq \sum_{cj \in Ch(pr)} BW(cj) \leq \sum_{ci \in Ch(pr)} BW(ci) \end{array} \right\} \quad (6.16)$$

$$\begin{aligned} BestChild(pr) &= \{\cup_{cj \in Ch(pr)} cj / j \leq i, \\ i &= Card(Ch(pr)) \cap BW(\cup_{cj \in Ch(pr)} cj) = BestBW \} \end{aligned} \quad (6.17)$$

End

Les formules ci-dessus sont déterminées la façon de choisir le nœud Best Child parmi les nœuds enfants d'un nœud parent.

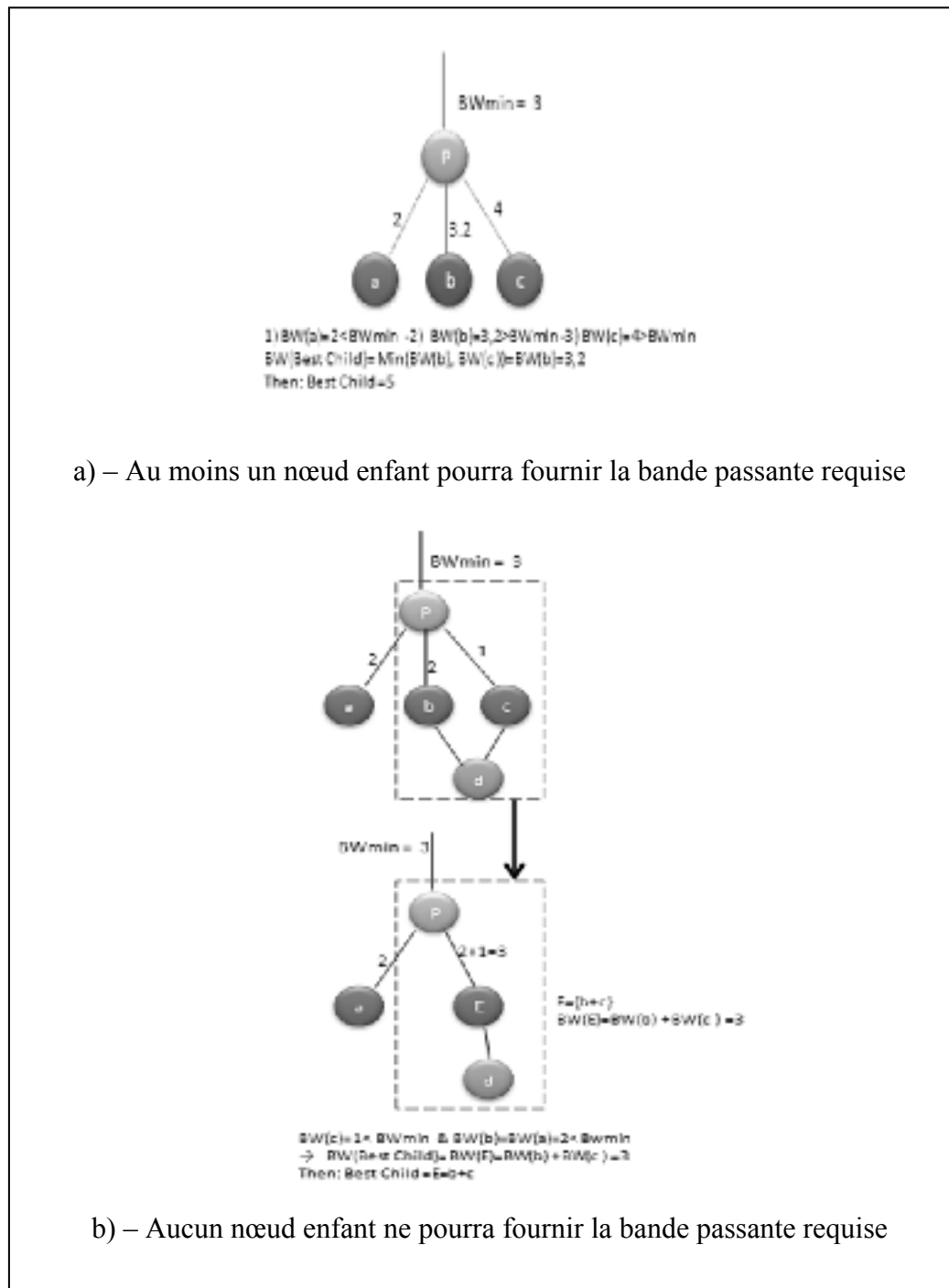


Figure 6.9 Best Child selection process

Le processus de sélection d'un nœud Best Child repose sur le choix du nœud ayant la plus petite quantité de bande passante parmi ceux qui détiennent une bande passante disponible supérieure ou égale à la bande passante BW_{min} (figure 6.9.a). Si aucun nœud n'est en mesure de fournir une telle quantité de bande passante, le parent choisira donc plusieurs nœuds, dont la somme de bande passante disponible est supérieure ou égale à BW_{min} . Ces nœuds sont vus par le nœud parent comme un seul nœud logique. Dans ce cas-là, lors de la transmission des paquets, l'information sera scindée en plusieurs parties. Le nombre de divisions est égal au nombre d'enfants formant le nœud logique. La division de l'information ne se fera pas forcément en quantité égale, chaque quantité sera proportionnelle à la quantité de bande passante qui peut être fournie par le nœud correspondant (figure 6.9.b).

6.3.2.2 LBS-AOMDV Description

Tel qu'il a été mentionné plus haut dans cette thèse, L'algorithme RBC ne réussit pas à toutes les reprises à construire un réseau papillon dans le réseau WMN. Rappelons que le réseau papillon est construit dans le but d'effectuer un routage avec codage réseau entre une source de donnée et sa destination. Dans un tel cas, l'algorithme LBS-AOMDV sera exécuté pour la détermination d'un multipath dans le WMN assurant une communication D2D avec application du Load Balancing (figure 6.10) et assurant la restauration de la topologie en cas de rupture de liens, tout en minimisant le temps de restauration. LBS-AMDV s'initie quand la source S exprime son intention de transmettre des données à la destination D. S commencera par établir une connexion avec D en utilisant l'algorithme Dijkstra. Une fois que le plus court chemin reliant S et D soit établi, S transmet sa requête SREQ (Send Request) à D afin de l'informer de sa volonté de lui envoyer des paquets. La requête SREQ, destinée à D, contient une demande de la liste des voisins directs de D ainsi que de la bande passante disponible pour chacun de ces voisins. Une fois que SREQ est reçue par D. Cette dernière inonde le réseau avec un message Hello (HelloMsg) avec un Time To Live (TTL) égale à un. Cette opération a pour but de déterminer l'ensemble des voisins directs de D, ainsi que de connaître la bande passante disponible qui peut être fournie par chaque nœud recevant HelloMsg.

Algorithme 6.1 Algorithme LBS-AOMDV

Algorithm LBS-AOMDV
DATA : G = (E, V) such that G is a graphe RESULT : $G_m = (N, P)$ Where N is multipath routes and P is Multipath nodes s = source , d = destination e = Dijkstra(s, d), $BW_{min} = \alpha$ mbps children(s) = $\emptyset$, children(d) = $\emptyset$, /* Unicast SREQ to d through “e” S_{req}(d, struct(Child(d), ChildBW)) /*d broadcast Hello message with TTL=1 Hello_{msg}(d, TTL) FOREACH n $\in$ G receiving Hello_{msg}(d, TTL) do /* n replays to d with Hello response Hello_{Rep}(n, d, BW) children(d) = children(d) $\cup$ {n} ENDFOR /*d unicast SRep to s (as a response to Rreq) S_{rep}(d, s, struct(Child(d), ChildBW)) /* S defines its children set Hello_{msg}(s, TTL) FOREACH n $\in$ G receiving Hello_{msg}(s, TTL) do /* n replays to d with Hello response Hello_{rep}(n, s, BW) children(s) = children(s) $\cup$ {n} ENDFOR /* s constructs the sets of children belonging to the multipath ch(s) = {$n_k \in \text{children}(s) \mid BW(n_k) \geq BW_{min}$} ch(d) = {$n_l \in \text{children}(d) \mid BW(n_l) \geq BW_{min}$} /* s sends RREQ to d RREQ(Rreq_{id}, ch(s), ch(d), BW_{min}) RreqTx(s) = off RreqRx(s) = off FOREACH n_k $\in$ ch(s)do RreqRx(s) = off ENDFOR Multipath()

Tout nœud recevant le paquet HelloMsg va répondre avec un message Hello Replay (HelloRep) informant D de son existence et de la bande passante disponible à son niveau. Une fois que la liste des enfants (voisins directs) de D est obtenue par celle-ci, la destination va procéder à la transmission de cette liste à la source. L'envoi de la liste des voisins, avec l'information sur leur bande passante, s'effectue via la réponse Send Replay (SREP). À la réception de SREP, la source S procédera à la détermination de la liste de ses voisins directs. Il s'agit de ces nœuds enfants. Pour ce faire, elle procédera de la même manière que D. Elle va donc inonder le réseau avec un message HelloMsg et s'attendra à recevoir un message HelloRep de la part de ses voisins directs l'informant de leur existence et de leur bande passante disponible.

Les messages HelloMsg et HelloRep sont échangés périodiquement entre tous les nœuds du multipath. Cet échange permet de détecter tout changement de la topologie suite à un bris lien.

Dans le but d'éviter la formation des boucles lors de la construction du multipath, deux entités ont été définies pour chaque nœud du réseau WMN. Il s'agit du RReqRx et RReqTx, toutes les deux peuvent avoir deux valeurs possibles, à savoir « Off » et « On ». Notons qu'elles sont initialisées à « On ». Par la suite, d'une part, chaque Nœud qui envoie une requête RREQ va mettre la valeur de RReqTx qui lui correspond à « off ». D'autre part, chaque nœud qui reçoit la requête RREQ va mettre la valeur de RReqRx qui lui est reliée à « off ».

Ces deux opérations s'interprètent comme suit : chaque nœud qui reçoit une requête RREQ va refuser la réception de toute requête RREQ du même identifiant. De la même façon, chaque nœud ayant envoyé une requête RREQ va s'abstenir d'en envoyer d'autres ayant le même identifiant. De cette manière, une requête RREQ, déterminée par un numéro d'identification unique ne va pas transiter par un même chemin plus d'une fois. Cela empêchera la construction de boucle dans le multipath.

Algorithme 6.2 Fonction Multipath

Multipath Function
<pre> Parent = ch(s) P = 0 path(p) = $\emptyset$ FOREACH $n_j \in \text{Parent}$ do p = p + 1 HelloMsg(n_j; TTL = 1) ENDFOR FOREACH $n_i \in G$ receiving HelloMsg do /* n_i replays to n_j with HelloRsp packet */ HelloRsp(n_i, n_j, BW); children(n_j) = children(n_j) $\cup \{n_i\}$ ENDFOR /* n_j builds ch(n_j) ch(n_j) = {$n \in \text{children}(n_j) \mid \text{RReqRx}(n) = \text{on}$} struct($n_c, \text{BW}(n_c)$) = BestChild(ch(n_j)) RReqTx(n_j) = off SendRReq(RReq_{id}, n_j, n_c RReqRx(n_c) = off path(p) = path(p) $\cup (n_i, n_c)$ IF $n_c \in \text{ch}(d)$ THEN path(p) = {path(p) $\cup (n_c, d)$} /* d sends RRep to s throw reverse path p RRep($d, s, \text{reverse}(\text{p})$) IF number(RRep) == Card(ch(d)) THEN RReqRx = off ENDIF ELSE nj = nc ENDIF </pre>

Deux exceptions s'appliquent. La première concerne la source, qui ne met son sa valeur RReqTx à « off » qu'une fois elle transmet la requête RREQ à tous ses enfants. La deuxième exception est relative à la destination. Celle-ci ne met sa valeur RReqRx à « off » qu'une fois qu'elle reçoit toutes les requêtes RREQ de tous ces enfants. Ces deux cas particuliers sont nécessaires pour former le multipath. Autrement, une seule route sera définie entre S et D. Notons que les valeurs RReqTx et RReqRx sont mises à jour au fur et à mesure que les nœuds du réseau reçoivent ou transmettent des requêtes RREQ. Tel que le montre l'algorithme 6.1, une fois que les listes des voisins de S et ceux de D, sont connues. Le nœud S initialise la construction du multipath en envoyant la requête RREQ vers tous ses enfants. Pour ce faire, la fonction Multipath est appelée par l'algorithme LBS-AOMDV (Algorithme 6.2).

L'algorithme 6.2 représente la fonction multipath. Cette fonction a pour objectif de construire un multipath entre la source et sa destination. La construction de ce multipath se réalise en reliant les enfants de S aux enfants de D via des routes disjointes. La découverte des routes s'effectue via la transmission des requêtes RREQ de père en fils. Néanmoins, contrairement aux autres algorithmes de découverte de chemins, tel que AOMDV, LBS-AOMDV n'inonde pas le réseau par les requêtes RREQ, a plutôt chaque parent doit sélectionner un meilleur enfant parmi ces enfants pour lui transmettre le paquet RREQ. Le choix du meilleur enfant se fait par la fonction BestChild.

L'algorithme 6.3, illustre le processus de choix du meilleur enfant (BestChild). À travers les instructions de cette fonction, qui ont déjà été expliqués dans la section précédente liée à la sélection du meilleur nœud enfant. Il est évident que l'approche basée sur l'inondation du réseau avec des requêtes RREQ n'est pas adoptée comme solution de découvertes de routes du multipath. À la place, une nouvelle méthode est appliquée. Cette méthode, considérée comme l'une des contributions principales du LBS-AMDV, repose sur le principe de sélection des récepteurs du RREQ. Un seul ou certains nœuds seulement vont recevoir cette requête. La réception de la requête RREQ par plusieurs nœuds n'est qu'une exception pour LBS-AOMDV. Cette pratique contourne le problème où route ayant une bonne QoS ne soit

pas choisie, car elle possède un nœud de faible QoS. Dans un tel cas de figure, deux ou quelques nœuds seront choisis comme un seul nœud logique pour jouer le rôle d'un nœud enfant (voir la figure 6.9). Ceci dit, l'algorithme LBS-AOMDV contribue fortement pour réduire le nombre de requêtes RREQ qui circulent dans le réseau, ce qui diminue la quantité du trafic du contrôle.

Algorithme 6.3 Fonction BestChild

Fonction BestChild

IF $\forall c_i \in \text{ch}(\text{pr}) \Rightarrow \text{BW}(c_i) < \text{BW}_{\min}$ **THEN**

BestBW

$$= \left\{ \begin{array}{l} \sum_{c_j \in \text{Ch}(\text{pr})} \text{BW}(c_j) / c_j \in \text{Ch}(\text{pr}), j \leq i, \quad i = \text{Card}(\text{Ch}(\text{pr})): \\ \text{BW}_{\min} \leq \sum_{c_j \in \text{Ch}(\text{pr})} \text{BW}(c_j) \leq \sum_{c_i \in \text{Ch}(\text{pr})} \text{BW}(c_i) \end{array} \right\}$$

BestChild(pr) = $\{ \cup_{c_j \in \text{Ch}(\text{pr})} c_j \mid j \leq i,$

$i = \text{Card}(\text{Ch}(\text{pr})) \cap \text{BW}(\cup_{c_j \in \text{Ch}(\text{pr})} c_j) = \text{BestBW} \}$

ELSE

BestBW = $\{ \text{BW}(c_j) \mid \forall c_i \in \text{Ch}(\text{pr}), \text{BW}_{\min} \leq \text{BW}(c_j) \leq \text{BW}(c_i) \}$

BestChild(pr) = $\{ c_j / c_j \in \text{Ch}(\text{pr}) \cap \forall c_i \in \text{Ch}(\text{pr}), c_i \neq c_j : \text{BW}(c_i) > \text{BW}_{\min} \Rightarrow \text{BW}(c_j) \leq \text{BW}(c_i) \}$

ENDIF

6.3.2.3 Simulations et résultats

Dans cette section, des simulations avec Matlab ont été effectuées pour la construction d'un multipath en utilisant l'algorithme LBS-AOMDV et l'algorithme AOMDV. Les résultats des deux approches ont été comparés afin de valider la solution LBS-AOMDV. Rappelons que l'utilisation du LBS-AOMDV pour la construction du multipath dans un réseau WMN a pour but de déterminer plusieurs routes disjointes possibles connectant une source S à sa destination D . Chaque route calculée est apte à fournir la bande passante minimale requise par la source S pour la transmission des paquets. Certaines routes seront utilisées pour la transmission et d'autres seront marquées comme chemins alternatifs (Back-Up). Une telle solution permet l'application du Load Balancing et de la redondance, ainsi que la restauration de la communication en cas de rupture des routes de transmission. L'application du LBS-AOMDV s'effectue dans un réseau WMN hébergé dans une petite cellule LTE HetNet. De ce fait, les données échangées dans ce réseau concerneront des communications D2D.

Pour réaliser nos expériences, plusieurs réseaux WMN ont été générés dans Matlab, avec un nombre variable de nœuds. Le nombre d'hôtes de chaque réseau varie entre 30 et 80 nœuds. Chaque topologie est incluse dans une aire de 800 m x 800 m. La localisation des nœuds a été générée aléatoirement avec un rayon de transmission de l'ordre de 250 m. Chaque lien du réseau est doté d'une bande passante disponible de 1 Mbps. Notons que la source et la destination sont choisies aléatoirement parmi les nœuds du WMN. Pour LBS-AOMDV deux scénarios ont été considérés relativement à la bande passante minimale BW_{min} requise par S . Notamment, $BW_{min} = 1\text{ Mbps}$ et $BW_{min} = 2\text{ Mbps}$. Rappelons que chaque nœud doit posséder une bande passante disponible supérieure ou égale à BW_{min} pour qu'il soit candidat à joindre le multipath. Cette condition ne concerne pas l'algorithme AOMDV, qui ne prend pas en considération la valeur de la bande passante disponible dans un nœud lors de la construction du multipath.

À Travers les différents scénarios, il est sujet de démontrer que l'approche LBS-AOMDV est capable de réduire la quantité des paquets RREQ transmis dans le réseau comparativement à

AOMDV, ce qui réduira donc le trafic du contrôle. En outre, il est à illustrer que notre solution offre à la source l'information sur la QoS offerte par chaque route calculée. Cette information est interprétée par la quantité de la bande passante disponible dans cette route. Notons que la bande passante disponible pour une route est égale à la valeur minimale des quantités de bandes passantes disponibles dans tous les liens de cette route. Finalement, ces simulations ont pour finalité de montrer que LBS-AOMDV réduit le nombre d'appels bloqués dans le réseau WMN par rapport à AOMDV.

Les figures 6.10 et 6.11 représentent la quantité de la bande passante disponible pour le plus court chemin versus le nombre de nœuds du réseau WMN pour les deux approches LBS-AOMDV et AOMDV.

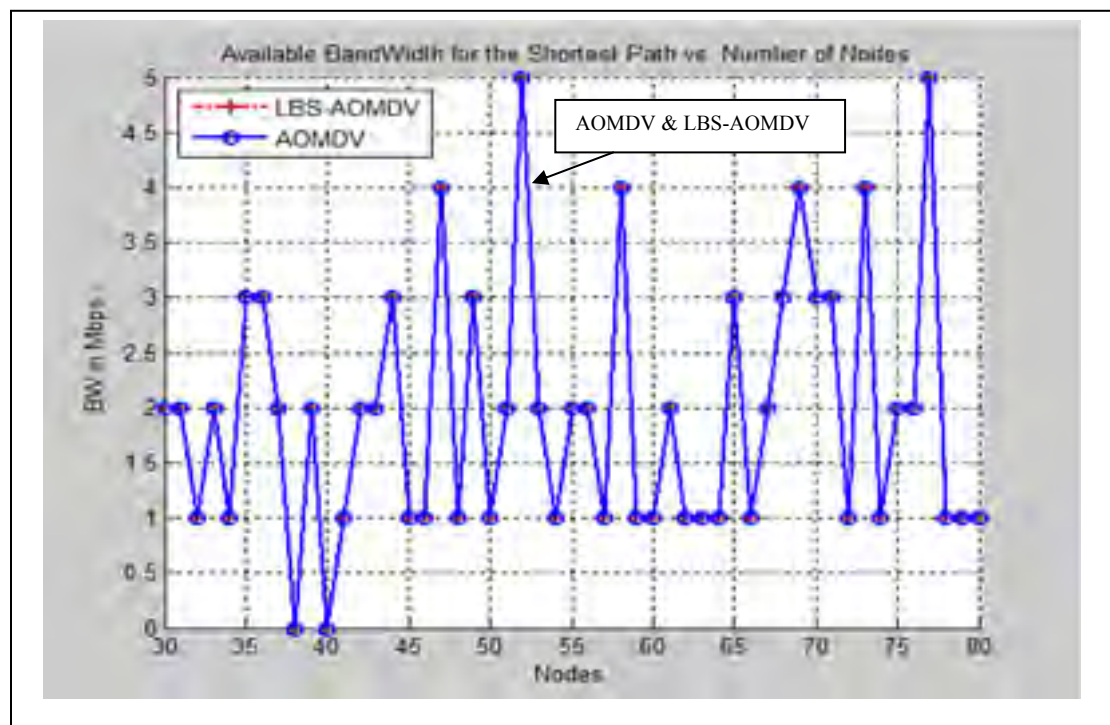


Figure 6.10 Bande passante disponible pour le plus court chemin
(BW required = 1Mbps)

La figure 6.10 montre que les deux algorithmes obtiennent des résultats identiques. Cette situation est produite, car pour ce scénario la bande passante minimale requise pour accepter

un nœud dans le multipath est égale à la bande passante disponible dans chaque nœud. Autrement dit, tous les nœuds du réseau WMN sont aptes à joindre le multipath. De ce fait, LBS-AOMDV choisira les mêmes chemins que ceux pris par AOMDV pour chaque topologie WMN considérée. Ces chemins sont ceux qui contiennent un nombre minimal de sauts entre la source et la destination.

Par contre, la figure 6.11 donne des résultats différents quant à la quantité de la bande passante offerte par le plus court chemin. La figure 6.11 montre que dans certains cas, AOMDV détermine un chemin reliant la source à sa destination, alors que LBS-AOMDV n'en calcule aucun. C'est le cas pour le réseau WMN à 50 nœuds par exemple. Il ne s'agit pas d'une faiblesse de LBS-AOMDV. En effet, dans tous les cas où LBS-AOMDV ne détermine aucun chemin alors qu'AOMDV en offre un, il s'avère qu'il s'agit d'un lien non conforme aux exigences de la source S relativement à la quantité de la bande passante minimale requise pour le choix de la route. Rappelons qu'un chemin n'offrant pas la quantité minimale de la bande passante requise va causer une détérioration de la QoS des trafics y transitant. En revanche, l'objectif de notre recherche est de développer une solution de routage apte à améliorer la qualité des communications D2D dans les petites cellules. Par conséquent, les résultats obtenus par LBS-AOMDV dans la simulation représentée par la figure 6.11 sont considérés comme meilleurs que ceux obtenus par AOMDV pour le même scénario, dans le sens où tous les chemins calculés par LBS-AOMDV sont des chemins conformes aux exigences de la source S en termes de QoS. En plus, ceux sont des chemins qui sont capables de garantir un niveau de QoS acceptable pour les communications D2D. D'autre part, il est évident que les routes calculées par AOMDV ne peuvent pas garantir, en tout temps, un acheminement de paquets avec une bonne QoS.

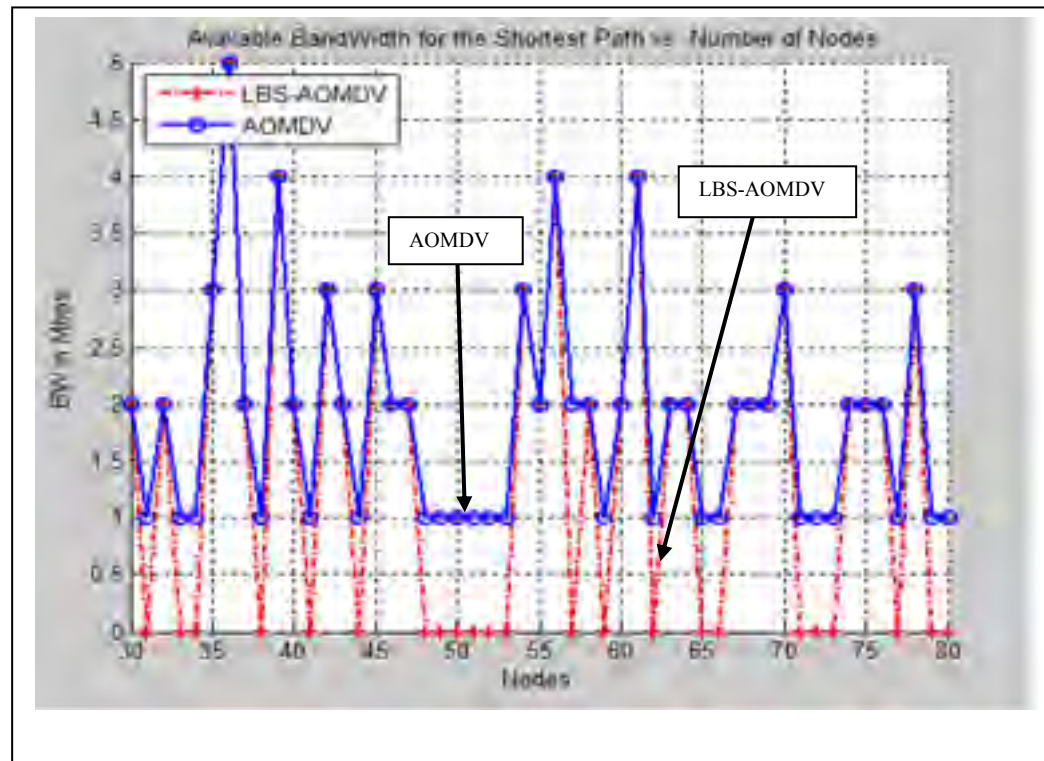


Figure 6.11 Bande passante disponible pour le plus court chemin
(BW required = 2Mbps)

Les résultats de la figure 6.11 peuvent être vus sous un autre angle. Étant donné que plusieurs chemins calculés par AOMDV ne répondent pas aux conditions établies par la source pour la sélection de routes, tous les paquets qui y transitent vont souffrir de manque de ressources de bande passante. Le manque de ressources de bande passante sur ces chemins va créer une congestion lors du routage des données. Par conséquent, beaucoup de paquets seront rejetés comparativement au cas d'utilisation des chemins calculés par LBS-AOMDV, où plus de ressources sont disponibles.

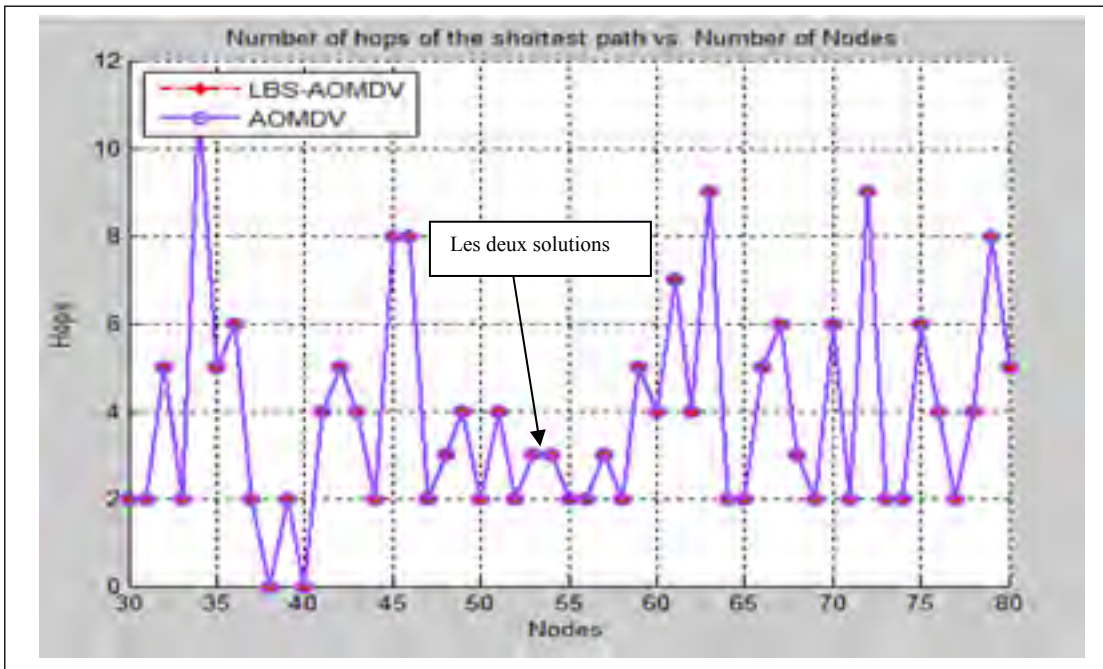


Figure 6.12 Nombre de sauts dans le plus court chemin
(BW required = 1Mbps)

La figure 6.12 représente le nombre de sauts du plus court chemin du multipath construit par chacune des solutions AOMDV et LBS-AOMDV. Il s'agit du cas où la bande passante minimale requise par la source pour la sélection du nœud est égale à 1 Mbps. Rappelons que pour toutes les topologies WMN simulées dans ce scénario, tous les nœuds disposent d'une bande passante disponible de 1Mbps. Cela à mener à avoir les mêmes résultats pour les deux approches. Donc comme tous les nœuds du réseau sont acceptables pour joindre le multipath, alors le LBS-AOMDV aura le même comportement qu'AOMDV quant au choix du plus court chemin. Il s'agit ici du chemin contenant le moins nombre de sauts entre S et D.

Contrairement aux résultats illustrés dans la figure 6.12, la figure 6.13 montre que LBS-AOMDV calcule des chemins plus longs que ceux calculés par AOMDV pour la majorité des topologies WMN simulées. Ceci est possible tant que certains nœuds inclus dans les chemins calculés par AOMDV n'offrent pas la bande passante minimale requise. De ce fait, LBS-AOMDV ne les sélectionne pas lors de la construction de son multipath. D'autres nœuds appartenant à des chemins ayant plus de sauts peuvent par contre offrir une bonne quantité de

bande passante disponible. Par conséquent, LBS-AOMDV va préférer ces nœuds par rapport aux nœuds choisis par AOMDV. Encore une fois, ces résultats ne seront pas pris comme les plus mauvais résultats, du moment que les chemins offerts par LBS-AOMDV, garantissent un bon niveau de QoS, alors que ce n'est pas le cas pour les chemins calculés par AOMDV, malgré qu'ils soient plus courts. Rappelons que le but de LBS-AOMDV n'est pas d'optimiser la routage entre la source S et sa destination D, mais plutôt de déterminer des chemins aptes à offrir la QoS exigée par S.

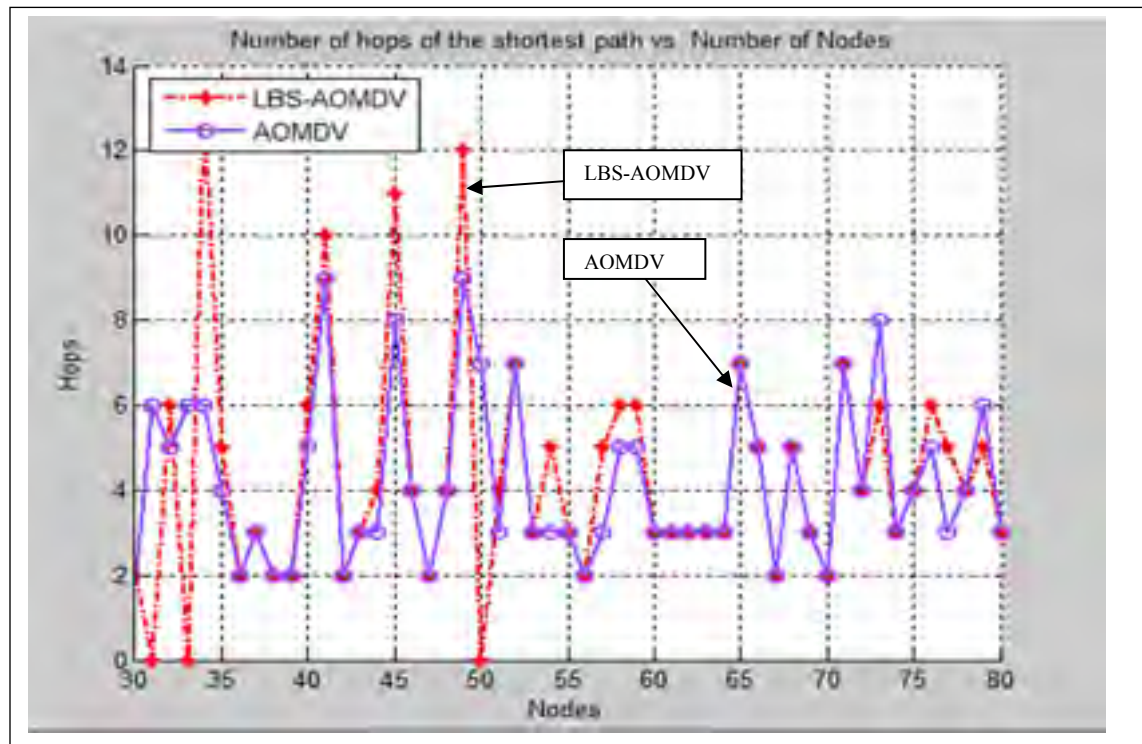


Figure 6.13 Nombre de sauts dans le plus court chemin
(BW required = 2Mbps)

Les deux figures 6.14 et 6.15 illustrent le nombre de requêtes RREQ diffusées dans le réseau pour les différentes topologies WMN simulées dans ce chapitre. Deux scénarios sont considérés, à savoir le cas où BWmin = 1 mbps (figure. 6.14) et le cas où BWmin = 2 Mbps (figure 6.15). Les résultats obtenus pour les deux cas de figure montrent que le trafic généré par la transmission des paquets RREQ est nettement moins important dans le cas de

l'application de LBS-AOMDV pour la construction du multipath. Notons que même si les routes calculées par les deux approches sont parfois identiques (voir figure 6.12), le nombre de paquets RREQ générés pour la construction de ces routes est supérieur dans le cas de l'utilisation d'AOMDV (figure 6.14). Ce résultat est attendu, du fait que la découverte des routes avec AOMDV se base sur l'inondation du réseau avec les paquets RREQ. Alors que LBS-AOMDV n'envoie la requête RREQ qu'au BestChild. Rappelons que le Best Child est le meilleur candidat parmi les nœuds enfants du transmetteur, pour recevoir la requête RREQ. C'est celui qui détient la valeur de la bande passante minimale parmi les nœuds qui peuvent offrir le niveau de la QoS exigé par la source S.

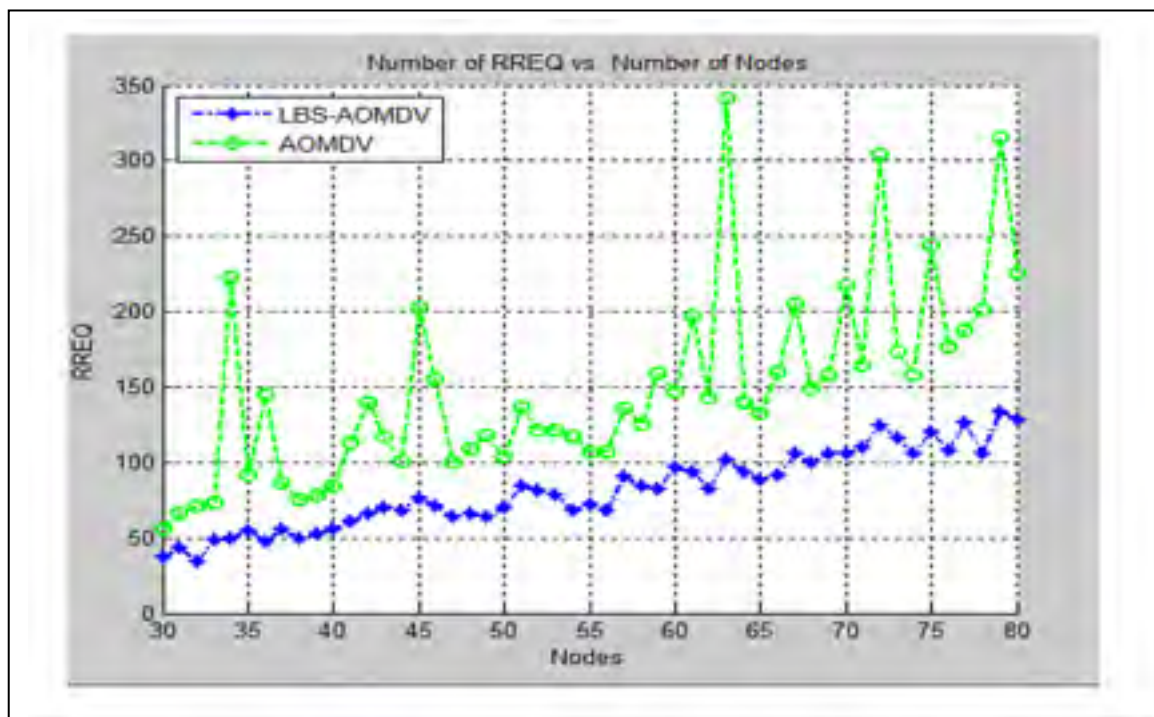


Figure 6.14 Nombre des requêtes RREQ transmises dans le réseau
(BW required = 1Mbps)

D'autre part, dans le cas où $BW_{min} = 2$ Mbps, la figure 6.15 donne des résultats semblables à ceux donnés par la figure 6.14. Autrement dit, même pour le cas où les chemins calculés sont souvent plus longs, LBS-AOMDV réduit le nombre de paquets RREQ transmis dans le réseau pour la construction du multipath. De plus, si on retourne vers les résultats de la figure

6.13, on constate que LBS-AOMDV offre des chemins plus longs que ceux calculés par AOMDV. Cependant, le nombre des requêtes RREQ diffusées dans le réseau demeure inférieur au cas de l'application d'AOMDV. La raison est identique à celle citée pour le cas précédant. Il s'agit évidemment du principe de l'inondation du réseau par les paquets RREQ utilisé par AOMDV pour la construction des routes du multipath.

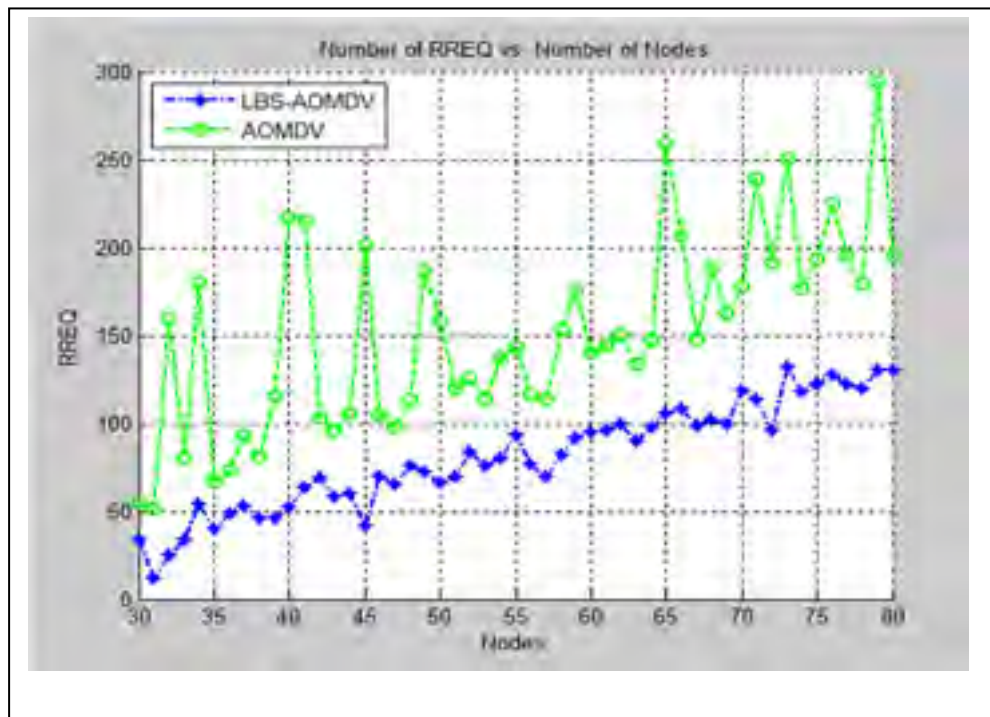


Figure 6.15 Nombre des requêtes RREQ transmises dans le réseau
(BW required = 2Mbps)

Notons que la construction du multipath s'effectue à l'intérieur d'un réseau maillé (WMN) où le nombre de liens entre un nœud et ses voisins est sensiblement plus important que dans les autres réseaux locaux sans fil. De ce fait, le nombre de requêtes RREQ transmises par AOMDV va être important, du moment que le principe d'inondation est la technique de transmission adaptée par ce protocole.

Finalement, on constate que l'application du LBS-AOMDV pour la construction d'un multipath dans le réseau WMN apporte des avantages pour l'amélioration des

communications D2D au sein des petites cellules LTE HetNet. Avantages se traduisent par la réduction de la quantité du trafic du contrôle qui transite dans le réseau, ainsi que par la construction de routes garantissant un bon niveau de QoS pour les trafics qui y transitent. Rappelons que l'objectif de LBS-AOMDV est loin de chercher des routes optimales en termes de nombre de sauts, mais plutôt de calculer des chemins entre S et D offrant un bon niveau de QoS.

6.4 Sécurisation des communications D2D pour les réseaux PSN

La gestion des désastres a toujours attiré l'attention des chercheurs, ainsi que celle des premiers répondants des réseaux de la Sécurité Publique (PS). Durant les moments de crises, toute information est importante pour sauver des vies. C'est pour cette raison que les réseaux de la Sécurité Publique ont été déployés et des fréquences radio ont leur été réservé. Aujourd'hui, les premiers répondants peuvent même accéder à radio fréquence commerciale grâce à la technologie LTE. Cet accès contribue fortement à l'amélioration des performances des communications du réseau PS en rendant disponibles des ressources radio supplémentaires capables de minimiser le blocage des appels des premiers répondants. Ces appels sont en grand nombre reliés à des processus de gestion de désastres et de sauvetage de vies humaines. Toutefois, l'utilisation des ressources commerciales LTE semble parfois impossible. Plusieurs raisons sont en cause, entre autres, on cite la localisation de la zone de la crise se trouve en dehors de la couverture des cellules LTE ou la pénurie des ressources dans ces cellules. De telles causes bloquent l'accès aux ressources commerciales, pour les premiers répondants. Cet état de choses pourra affecter profondément la bonne résolution de la crise. Par conséquent, le besoin de trouver de nouvelles ressources a donné naissance à la technologie Device-to-Device (D2D) (Lin, Andrews et al. 2013, Raghothaman, Deng et al. 2013). Pour communiquer, les nœuds utilisant la technologie D2D peuvent utiliser une partie des ressources radio LTE, comme ils peuvent s'en passer pour opérer dans les fréquences publiques, via des réseaux locaux sans fil. Ces réseaux locaux, tels que WMN ou WiFi, assurent la transmission de données à l'intérieur de petites cellules, de rayon de couverture plus petit que celui de la macro cellule LTE. Afin de tirer un meilleur profit de l'intégration

de ces petites cellules au réseau LTE, leurs ressources radio seront utilisées pour décharger la macro cellule LTE lors des moments de congestions. En effet, certains bearers arrivant vers la macro cellule seront basculés d'une manière transparente aux clients, vers les petites cellules. Cette solution est très efficace pour résoudre le problème de pénurie des ressources radio, d'autant plus que les fréquences publiques sont des ressources radio gratuites. D'ailleurs, plusieurs recherches ont montré l'efficacité de l'utilisation des fréquences des réseaux locaux sans fil dans la gestion des ressources radio des réseaux LTE HetNets (Hagos 2012, Lei, Zhang et al. 2013, Nagpal, Choudhury et al. 2013, Pyattaev, Johnsson et al. 2013, Andreev, Pyattaev et al. 2014, Fodor, Sorrentino et al. 2014, Mumtaz, Lundqvist et al. 2014). En revanche, ces réseaux WLAN sont connus par leur vulnérabilité aux attaques de sécurité, alors que les trafics LTE doivent bénéficier d'un bon niveau de sécurité, tel est le cas pour les données qui utilisent les fréquences privées LTE.

Les attaques de sécurité peuvent être des attaques internes ou externes. Les attaques internes sont celles qui sont effectuées par des nœuds malveillants, membres du WLAN. Par contre les attaques externes sont effectuées par des pirates externes capables d'intercepter l'information qui circule. Les différentes attaques peuvent affecter d'une façon remarquable la confidentialité, l'intégrité et la disponibilité des données échangées dans le réseau. La confidentialité des données est définie par le fait que l'information ne doit être accessible et connue que par des usagers autorisés. Autrement dit, aucune information ne doit être divulguée à un utilisateur non autorisé. L'intégrité des données se traduit par le concept de préservation de l'information sous une forme complète et inchangée durant toute sa durée de vie. Quant à la disponibilité des data, cela signifie que l'information doit être disponible pour utilisation et consultation d'une façon permanente.

Rappelons que les réseaux sans fil utilisent les fréquences radio pour l'échange d'informations. Cet échange est basé sur la diffusion de l'information par l'émetteur et par l'interception de celle-ci par le récepteur en utilisant des antennes radio aptes à capturer le signal. Ceci dit, un pirate peut se placer dans la zone de diffusion de l'information et sera

donc capable d'intercepter le signal tout comme l'utilisateur autorisé. L'information doit être donc sécurisée afin de pallier les problèmes de sécurité.

Une solution qui apparaît pertinente pour réduire la vulnérabilité des réseaux locaux sans fil est l'utilisation du codage réseau. Cette technique se base sur le principe de codage des paquets avant leur envoi. Bien que les paquets codés s'avèrent plus difficiles à lire vu qu'ils ne sont plus identiques aux paquets natifs envoyés par la source, mais il reste que certains paquets sont transmis sans codages. Ces derniers sont nécessaires pour le décodage de l'information. Par conséquent, l'utilisation du codage réseau tel qu'il est défini initialement par (Ahlsweide, Cai et al. 2000) ne garantis pas un bon niveau de sécurité de l'information.

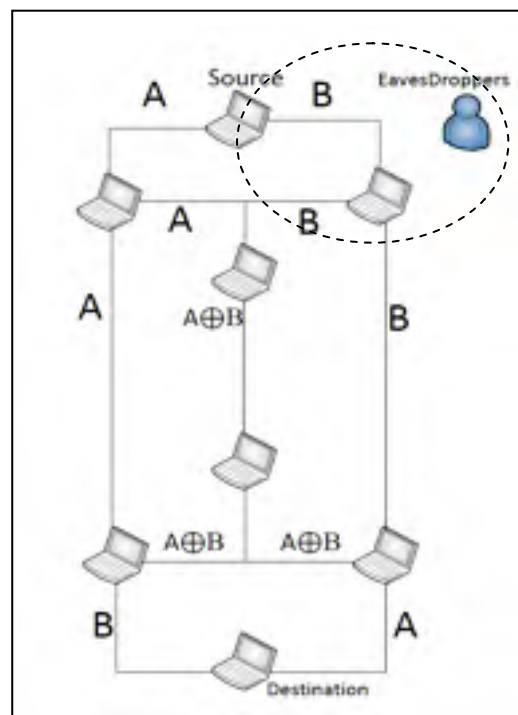


Figure 6.16 Codage réseau avec attaque de sécurité

La figure 6.16 illustre un cas de routage des données basé sur le codage réseau avec la présence d'un eavesDropper. Un eavesDropper est un nœud malveillant dont le but est d'intercepter les informations qui transitent à l'intérieur de son rayon de transmission, pour des fins d'espionnage. Il s'agit donc d'une attaque de confidentialité. Ainsi, le pirate

représenté dans la figure 6.16 peut intercepter tous les paquets qui passent par sa zone de couverture, représentée par le cercle en pointillés. Par conséquent, tous les paquets qui sont envoyés par S vers D via le chemin latéral droit du réseau papillon seront lus par l'eavesDropper.

Maintes solutions ont été développées pour améliorer le niveau de sécurité dans les réseaux sans fils afin d'éviter les attaques de confidentialité, d'intégrité et de disponibilité. L'une des approches proposées et utilisées d'une façon ré pondue consiste en la cryptographie (Shah, Rashmi et al. 2013, Tang 2013, Zhao, Kent et al. 2013, Zhou 2013). Néanmoins, il n'est pas possible d'éviter la croissance du trafic du contrôle quand un mécanisme de cryptage est adapté comme solution de sécurité. Surtout que ces mécanismes reposent sur de complexes opérations mathématiques. Par conséquent, quand les ressources réseau sont limitées, la QoS de l'ensemble du réseau se détériore.

Par ailleurs l'authentification des différents nœuds du réseau se montre comme une solution encourageante pour assurer la sécurité de l'information transmise. Plusieurs mécanismes ont été développés pour cette fin. Le protocole Transport Layer Security (Fischer and Röd ig 2000), le mécanisme Extensible Authentication Protocol (EAP) (Aboba, Blunk et al. 2004) , l'algorithme authentication, authorization and accounting (AAA) (Metz 1999) et le système Message Digest 5 (MD5) (Rivest 1992) sont tous des solutions d'authentification proposées dans la littérature pour empêcher l'accès au réseau aux les nœuds non autorisés. Remarquons qu'en plus du trafic du contrôle additionnel qui est généré suite à l'application de ces algorithmes, l'attaque de disponibilité ne sera pas résolue. Le déni de service (DoS) est l'une de ces attaques de disponibilité. DoS peut avoir lieu en submergeant le réseau du trafic inutile. Une simple requête Ping peut causer un DoS si elle est utilisée par un hacker.

D'autres propositions sont faites dans (Kent and Liebrock 2011, Sen, Koilakonda et al. 2011, Shah and Valiveti 2012, Prasad and Giri 2014) pour améliorer le niveau de sécurité dans le réseau. Toutefois, toutes ces approches augmentent le trafic de contrôle dans le réseau. Par conséquent, leur impact sur la QoS apparait certainement quand les ressources sont limitées.

Dans ce chapitre, une nouvelle approche, appelée Generalized Secure Network Coding based Data Splitting algorithm (G-SNCDS), a été développée pour sécuriser les réseaux locaux sans fil. Cette solution est appliquée sur les WMN, mais cela n'empêche pas son utilisation pour les autres réseaux WLAN.

L'algorithme G-SNCDS est une de sécurité pour la transmission de donnée D2D sur des réseaux LTE HetNets. G-SNCDS utilise le codage réseau pour la transmission de données. Il est basé aussi sur la technique de Data Splitting (DS) que nous développons dans cette thèse dans le but de sécuriser les informations contre les attaques internes et externes.

D'autre part, le codage réseau est établi via un réseau papillon. Ce dernier est construit en utilisant l'algorithme RBC, détaillé au début de ce chapitre. Rappelons que RBC est une solution de construction de réseaux papillon dans les réseaux locaux sans fil, spécialement dans les réseaux WMN. L'utilisation des réseaux papillon garantit le succès du processus du codage-décodage des données. Dans cette étude, nous montrons comment l'utilisation du codage réseau combiné à la technique Data Splitting arrive à éviter les attaques de confidentialité, d'intégrité et de disponibilité, sans augmenter le trafic de contrôle dans le réseau. Ainsi, la contribution principale du G-SNCDS est d'augmenter le niveau de sécurité des communications D2D au sein des petites cellules LTE HetNets sans affecter le niveau de la QoS dans le réseau. Cette solution est pertinente quand les ressources réseau sont limitées.

Il est à noter qu'une deuxième solution appelée SNCDS a été développée et publiée dans l'article [37]. Comme cette solution est une partie intégrante de G-SNCDS, on a jugé inutile de lui consacrer une section à part. En effet, SNCDS qui a été développé avant G-SNCDS, représente la partie qui traite les attaques de confidentialité.

Dans ce qui suit, la solution G-SNCDS sera présentée en détail. Elle sera succédée par des simulations de G-SNCDS et de SNCDS.

6.4.1 L'algorithme G-SNCDS

Dans ce travail, une nouvelle approche est développée pour la transmission sécurisée des données D2D sur les réseaux LTE HetNets. L'objectif de cette approche est de garantir la confidentialité des données transmises sur le réseau, leur intégrité, ainsi que leur disponibilité. Notre solution consiste à appliquer le mécanisme de la division des données (Data Splitting) lors de la transmission des symboles sur un réseau papillon. En d'autres termes, au lieu d'envoyer des paquets entiers à travers les chemins du réseau papillon, chaque paquet sera divisé en fragments de six bits. Par la suite, l'ordre des bits de chaque fragment sera mélangé selon la séquence de positions aléatoires (Random Sequence Position ou RSP) générée par la source. RSP est obtenue en effectuant un mélange aléatoire des positions des bits d'un fragment. Par ailleurs, les différents bits d'un même fragment seront transmis à la destination par l'intermédiaire de deux routes distinctes. La source émet le premier bit du fragment via le premier chemin latéral du réseau papillon. Le second bit sera envoyé via l'autre chemin latéral, et ainsi de suite. L'opération est ensuite réinitialisée avec le fragment suivant, et le processus est répété tant que des bits attendent d'être acheminés. Pour cette approche on suppose que la séquence RSP générée par la source sera cryptée et envoyée à la destination au début de la transmission. D'autre part, on suppose que le format de la matrice de codage est connu par la destination, mais pas ses valeurs. La destination utilise les éléments de RSP pour construire la matrice de codage.

Soit C la matrice de codage qui est utilisée pour la transmission des paquets avec codage réseau. Cette matrice est donnée comme suit :

$$C = \begin{pmatrix} c1 & c2 \\ c3 + c4c1 & c4*c2 \\ c5*c1 & c5c2 + c6 \end{pmatrix} \quad (6.18)$$

Avec C_i , $i=1$ à 6 , sont les codes utilisés par les codeurs du réseau papillon afin de réaliser un codage réseau des informations transmises. Tel que le montre la figure 6.18, le processus de codage réseau nécessite l'utilisation de deux codes par chaque nœud codeur. Rappelons que

quand le codage réseau est adapté comme solution de routage, deux symboles sont envoyés simultanément par un nœud transmetteur, contre un seul symbole envoyé par l'approche classique de transmission. La transmission simultanée de deux symboles exige par contre de les mixer ensemble avant leur transmission. Une façon de faire pour réaliser ce mixage est d'utiliser les codes de codage.

Par ailleurs soit $P = (p_1 p_2 p_3 p_4 p_5 p_6)$ la séquence RSP générée par la source S , avec $p_i \in [1,6]$. La destination construit la matrice C en substituant chaque code C_i par la valeur de P_i . Notons que les $p_i \in [1,6]$ représentent les positions des bits dans un même fragment d'un paquet donné. Tel qu'il a été mentionné ci-dessus, la séquence $RSP = P$ est cryptée par la source S avant sa transmission à la destination D . Il s'agit de la seule information cryptée par notre solution. En outre, la source peut envoyer une nouvelle séquence RSP autant de fois qu'elle souhaite changer le mécanisme de mixage de données qu'elle envoie. Cela peut avoir lieu après qu'un pirate arrive à réussir une attaque ou périodiquement afin de réduire la probabilité de la découverte de la séquence RSP par des nœuds malhonnêtes. Par conséquent, cette pratique ne peut que rendre plus robuste notre solution de sécurité.

Cette approche empêche les pirates d'obtenir des informations utiles en interceptant des données confidentielles. En effet, l'attaquant ne peut pas obtenir tous les bits du paquet envoyé, parce que certains bits transitent via un autre chemin de transmission qui se trouve hors de la zone de couverture de l'attaquant. De plus, le pirate ignore l'existence de la séquence RSP utilisée par la source pour mélanger les données envoyées. Par conséquent, ce sera compliqué pour lui de reconstruire le paquet natif. Ainsi, le G-SNCDS peut donc éviter les attaques de confidentialité via le mécanisme DS.

En outre, le système G-SNCDS utilise l'algorithme RBC pour obtenir un ensemble d'effets papillon. La transmission de données par plus d'un effet papillon permet d'éviter les attaques d'intégrité et de disponibilité. Cela sera détaillé plus loin dans ce travail. Notons que le mécanisme Data Splitting (DS) est appliqué pour chaque réseau papillon. De plus, chaque réseau papillon détient sa propre séquence RSP. Ces séquences peuvent être différentes. Cela

augmente le degré de difficulté relatif à la résolution des données transmises. Le système DS sera détaillé dans ce qui suit.

6.4.1.1 Le mécanisme Data Splitting et l'opération de codage de G-SNCDS

La figure 6.17 présente le paquet natif à envoyer par la source à la destination. Ce paquet est scindé en fragments de six bits chacun via le mécanisme DS du G-SNCDS (figure 6.17.b). Le choix de la valeur six comme nombre de bits de chaque fragment n'est pas le fruit du hasard. Notre objectif pour le choix de cette valeur est de faire une relation entre le nombre de positions de bits dans le fragment et le nombre de codes utilisés pour coder les symboles transitant via un réseau papillon. En fait, une telle topologie exige d'effectuer trois codages, en d'autres termes, trois nœuds codeurs vont accomplir la tâche de codage dans ce réseau dont chacune nécessite l'utilisation de deux codes. Par conséquent, nous proposons d'utiliser les six de RSP comme codes pour le processus de codage réseau. Pour ce faire, le premier bit à envoyer est celui qui possède la position, du fragment initial, définie par le premier élément de RSP. Le deuxième est celui qui en possède la deuxième valeur de RSP et ainsi de suite.

Rappelons que RSP donne la position initiale de chaque bit reçu par la destination. De la sorte, comme la destination connaît la séquence RSP et connaît le format de la matrice de codage, alors elle pourra à la fois générer la matrice de codage et aussi déterminer la position initiale dans son fragment, de chaque bit reçu.

Par ailleurs, la figure 6.17 montre les différentes étapes du mécanisme DS. Cette opération commence par scinder un paquet en plusieurs fragments de six bits (figure 6.17.b). Par la suite, la source utilise la séquence RSP pour générer les nouveaux fragments à envoyer à la destination (figure 6.17.c). Chaque nouveau fragment verra la position de ses bits modifiée en respectant les valeurs de RSP, les bits du nouveau fragment sont représentés par les variables b_i , avec $i = 1$ à 6. Finalement, le processus d'envoi de fragment basé sur le codage réseau sera initié (figure 6.17.d).

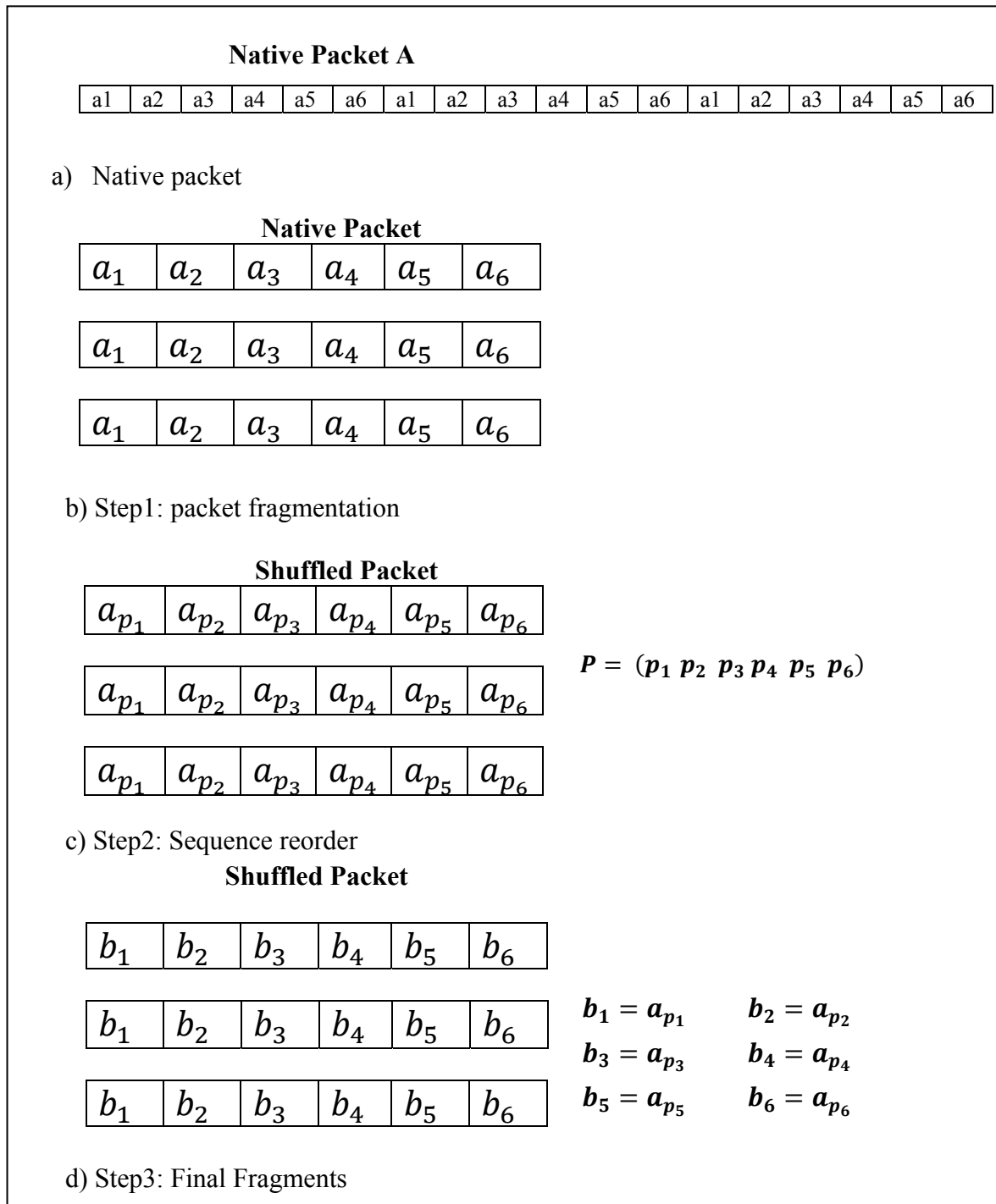


Figure 6.17 Le mécanisme Data Splitting

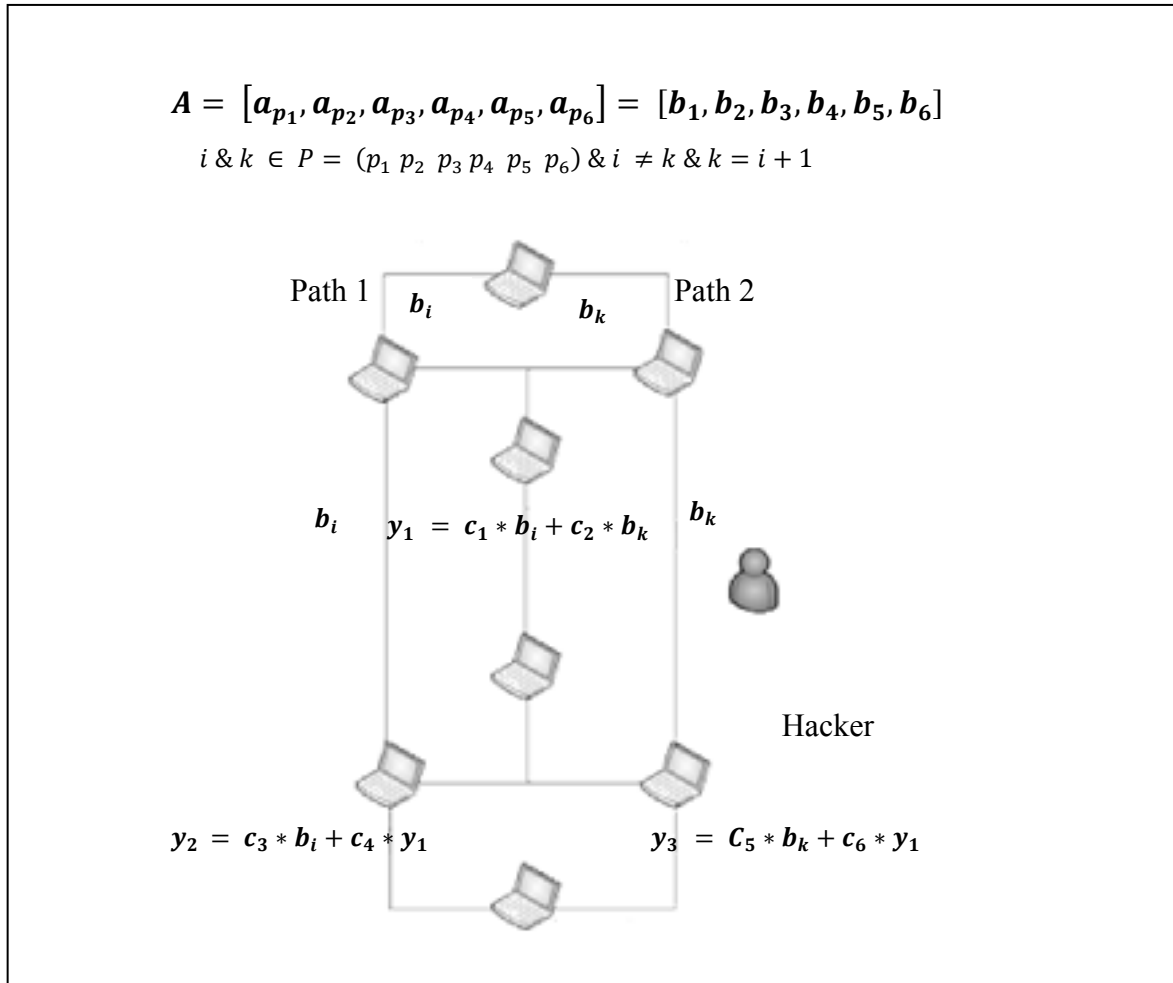


Figure 6.18 Transmission de paquets avec codage réseau et Data Splitting

La figure 6.18 montre un cas général de transmission de données en utilisant le codage réseau combine au mécanisme DS. Notons que tous les bits b_i , avec $i = 1$ à 6, de positions impaires d'un même fragment seront envoyés via le premier chemin. D'autre part, les autres bits seront transmis via l'autre chemin. Ce mécanisme permet d'éviter les attaques de confidentialité et réduire le taux de paquets affectés par une attaque d'intégrité. En effet, le système DS est capable d'éviter une attaque de confidentialité du moment qu'elle empêche le pirate d'intercepter la totalité de l'information transmise. Ceci est possible tant qu'une partie des données sont envoyées à travers un chemin qui se trouve en dehors de la zone de couverture du pirate. De plus, l'impact de l'attaque d'intégrité est réduit en utilisant DS pour

la même raison citée ci-dessus. En fait, 50% des paquets empruntent un chemin différent de celui qui est accessible par l'attaquant. Par conséquent, 50% des bits transmis seront à l'abri d'une éventuelle attaque d'intégrité. Il est à missionner qu'une solution complémentaire est proposée par G-SNCDS pour éviter les attaques d'intégrité d'une façon intégrale. Il s'agit d'utiliser plusieurs effets papillon pour la transmission de données. Cette approche est détaillée un peu plus loin dans ce chapitre.

Dans ce qui suit, la technique de transmission de données avec codage réseau et Data Splitting est détaillée. Soit t b_i, b_k les deux bits à transmettre par codage réseau, de S vers D, via le réseau papillon, avec $i \& k \in P = (p_1 p_2 p_3 p_4 p_5 p_6) \& i \neq k$. Considérons que C est la matrice de codage réseau. La matrice C est donnée par la formule 6.18 plus haut. Finalement, soit y_j les symboles codés, avec $j \in (1, 2, 3)$. La formulation mathématique relative au codage des bits citées ci-dessus est données comme suit:

$$\begin{aligned} y_1 &= c_1 * b_i + c_2 * b_k \\ y_2 &= c_3 * b_i + c_4 * y_1 \\ y_3 &= c_5 * b_k + c_6 * y_1 \end{aligned} \tag{6.19}$$

Par ailleurs, comme seuls Y_2 et Y_2 sont reçus par la destination et que Y_1 peut s'écrire en fonction de Y_2 et Y_2 , on ne va considérer que Y_2 et Y_2 pour modéliser notre analyse mathématique. En utilisant le système d'équations donné en 6.19 on pourra déduire le système d'équations 6.20.

$$\begin{aligned} y_2 &= c_3 * b_i + c_4 * y_1 = (c_3 + c_4 * c_1) b_i + c_4 * c_2 * b_k \\ y_3 &= c_5 * b_k + c_6 * y_1 = c_1 * c_6 * b_i + (c_5 + c_6 * c_2) b_k \end{aligned} \tag{6.20}$$

Finalement:

$$\begin{aligned} y_2 &= (c_3 + c_4 * c_1) b_i + c_4 * c_2 * b_k \\ y_3 &= c_1 * c_6 * b_i + (c_5 + c_6 * c_2) b_k \end{aligned} \tag{6.21}$$

Donc, le système de codage sera représenté comme suit

$$\begin{pmatrix} y_2 \\ y_3 \end{pmatrix} = \begin{pmatrix} c_3 + c_4 c_1 & c_4 * c_2 \\ c_5 * c_1 & c_5 c_2 + c_6 \end{pmatrix} * \begin{pmatrix} b_i \\ b_k \end{pmatrix} \quad (6.22)$$

Avec $i \& k \in P = (p_1 p_2 p_3 p_4 p_5 p_6)$ & $i \neq k$. Rappelons que $c_j = p_j$ pour $j \in [1,6], p_j \in P$.

Par ailleurs, à chaque fois que la destination reçoit les deux symboles codés y_i , avec $i = 1$ à 2 , elle exécutera le processus de décodage afin d'extraire les symboles natifs.

6.4.1.2 Exemple de codage réseau basé sur le mécanisme DS de G-SNCDS

Un exemple d'application du système DS au routage avec codage réseau est illustré dans la figure 6.19. Soit A le fragment initial de données à transmettre. Avec

$$A = (a_1, a_2, a_3, a_4, a_5, a_6) \quad (6.23)$$

La séquence RSP adaptée pour cet exemple est donnée par le vecteur P , avec:

$$P = (p_1 p_2 p_3 p_4 p_5 p_6) = (4, 2, 5, 6, 1, 3) \quad (6.24)$$

Par conséquent, la transmission de bits s'effectue selon la séquence B , tel que:

$$B = (b_1, b_2, b_3, b_4, b_5, b_6) = (a_4, a_2, a_5 a_6, a_1, a_3) \quad (6.25)$$

Avec b_i , $i=1$ à 6 , sont les bits envoyés par la source en respectant la séquence RSP. La source envoie les bits via deux chemins d'une façon alternative.

Donc, b_1 , b_3 et b_5 seront transmis via le premier chemin et b_2 , b_4 et b_6 via le deuxième.

La figure 6.20 illustre une transmission de bits basée sur le codage réseau et le mécanisme DS. Comme il apparaît sur la figure 6.20, un attaquant (EavesDropper) est présent et est capable d'intercepter une partie de l'information qui circule dans le réseau.

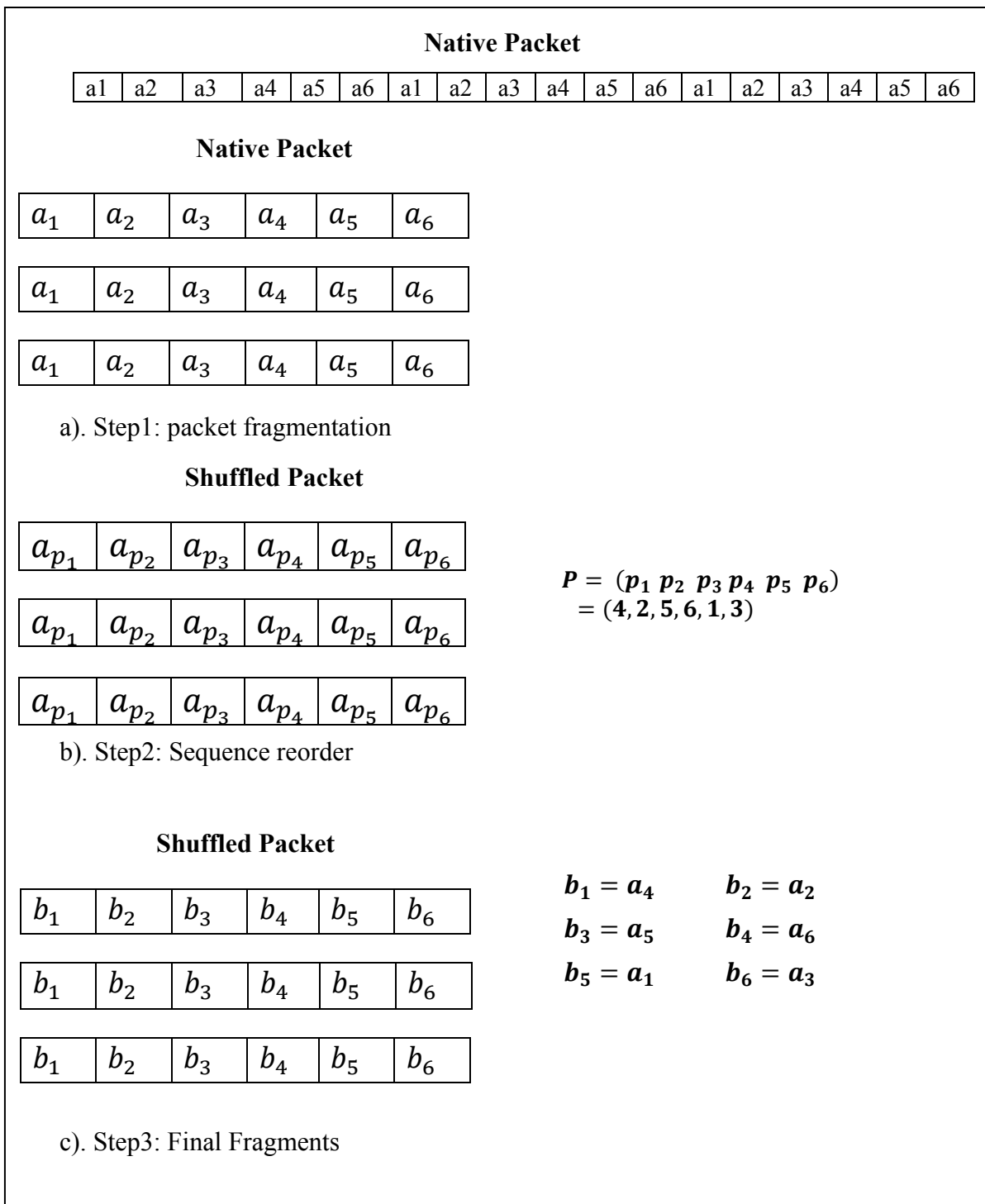


Figure 6.19 Exemple de Data Splitting

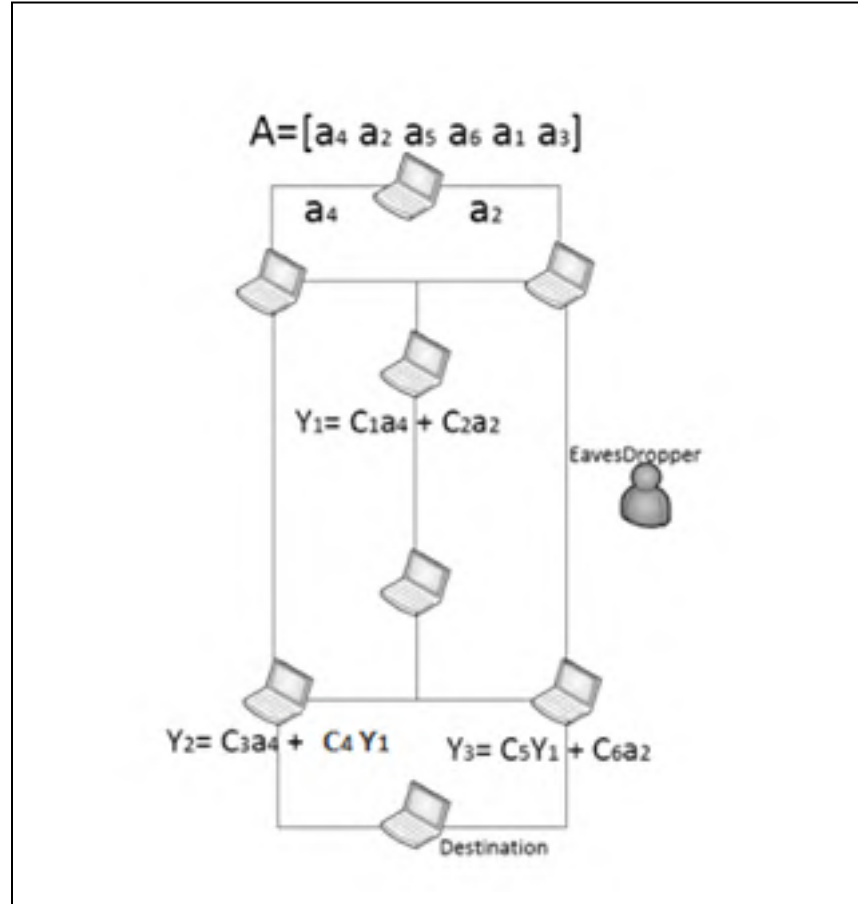


Figure 6.20 Exemple de l'utilisation du DS pour le codage réseau

La formulation mathématique de cette transmission de données avec codage réseau et Data Splitting se donne comme suit:

$$\begin{aligned}
 y_1 &= c_1 * a_4 + c_2 * a_2 \\
 y_2 &= c_3 * a_4 + c_4 * Y_1 = (c_3 + c_4 c_1) a_4 + c_4 c_2 a_2 \\
 y_3 &= c_5 * Y_1 + c_6 * a_2 = c_5 c_1 a_4 + (c_5 c_2 + c_6) a_2
 \end{aligned}
 \tag{6.26}$$

Alors le système de codage relatif à y_2 et y_3 sera donné comme suit:

$$\begin{pmatrix} y_2 \\ y_3 \end{pmatrix} = \begin{pmatrix} c_3 + c_4 c_1 & c_4 c_2 \\ c_5 c_1 & c_5 c_2 + c_6 \end{pmatrix} * \begin{pmatrix} a_4 \\ a_2 \end{pmatrix}
 \tag{6.27}$$

De la même façon, quand les symboles natifs $a5$ et $a6$ sont transmis, on aura les symboles $y1$, $y2$ et $y3$ qui sont relatifs aux deux symboles natifs cités ci-dessus. Les équations linéaires générées par le processus de codage des symboles natifs donnent ce qui suit:

$$\begin{aligned} y1 &= c1 * a5 + c2 * a6 \\ y2 &= c3 * a5 + c4 * Y1 = (c3 + c4c1)a5 + c4c2 a6 \\ y3 &= c5 * Y1 + c6 * a6 = c5c1 a5 + (c5c2 + c6)a6 \end{aligned} \quad (6.28)$$

Donc, le système de codage donne le système linéaire suivant:

$$\begin{pmatrix} y2 \\ y3 \end{pmatrix} = \begin{pmatrix} c3 + c4c1 & c4c2 \\ c5c1 & c5c2 + c6 \end{pmatrix} * \begin{pmatrix} a5 \\ a6 \end{pmatrix} \quad (6.29)$$

Un raisonnement identique, lors de la transmission de $a1$ et $a3$, permet d'écrire ce qui suit:

$$\begin{aligned} y1 &= c1 * a1 + c2 * a3 \\ y2 &= c3 * a1 + c4 * Y1 = (c3 + c4c1)a1 + c4c2 a3 \\ y3 &= c5 * Y1 + c6 * a3 = c5c1 a1 + (c5c2 + c6)a3 \end{aligned} \quad (6.30)$$

Donc, on pourra considérer l'écriture matricielle suivante pour le codage réseaux des deux symboles $a1$ et $a3$:

$$\begin{pmatrix} y2 \\ y3 \end{pmatrix} = \begin{pmatrix} c3 + c4c1 & c4c2 \\ c5c1 & c5c2 + c6 \end{pmatrix} * \begin{pmatrix} a1 \\ a3 \end{pmatrix} \quad (6.31)$$

Rappelons que $c_j = p_j$ pour $j \in [1,6]$, $p_j \in P = (p_1 p_2 p_3 p_4 p_5 p_6)$.

La section suivante définit le processus de décodage de l'information, ainsi que le rassemblement des données, effectuées par la destination afin de récupérer le paquet natif.

6.4.1.3 Le processus de décodage et le rassemblement des données de G-SNCDS

Le processus de décodage est effectué par la destination. Cela augmente le niveau de sécurité dans le réseau en n'autorisant la découverte de l'information à d'autres noeuds que la source et la destination. Le raisonnement mathématique suivant illustre le processus de décodage de l'information en effectuant des opérations d'algèbre linéaire sur le système matriciel relatif aux symboles codés représenté par la formule 6.22.

Dans le but de décoder les symboles reçus, nous considérons le système 6.22 et les trois formules suivantes:

$$Y = \begin{pmatrix} y_2 \\ y_3 \end{pmatrix} \quad (6.32)$$

$$C = \begin{pmatrix} c_3 + c_4c_1 & c_4c_2 \\ c_5c_1 & c_5c_2 + c_6 \end{pmatrix} \quad (6.33)$$

Et

$$A = \begin{pmatrix} a_1 \\ a_3 \end{pmatrix} \quad (6.34)$$

Donc le système linéaire 6.22 pourra se reformuler comme suit:

$$Y = C * A \quad (6.35)$$

L'application des propriétés de l'algèbre linéaire, on peut écrire ce qui suit:

$$\left\{ \begin{array}{l} C^{-1} * Y = C^{-1} * C * A \\ C^{-1} * Y = I_d * A \text{ avec } I_d \text{ est la matrice identité} \\ C^{-1} * Y = \hat{A} \\ \text{Alors} \\ A = C^{-1} * Y \end{array} \right. \quad (6.36)$$

De cette façon, on peut déduire l'écriture matricielle 6.37, obtenue par la substitution de Y, A et C par leurs équivalences données par les formules 6.32, 6.33 et 6.34.

$$\begin{pmatrix} a_i \\ a_j \end{pmatrix} = \begin{pmatrix} c_3 + c_4 c_1 & c_4 c_2 \\ c_5 c_1 & c_5 c_2 + c_6 \end{pmatrix}^1 * \begin{pmatrix} y_2 \\ y_3 \end{pmatrix} \quad (6.37)$$

Par ailleurs, la matrice C^{-1} doit être définie afin d'avoir les valeurs finales de a_i et a_j .

Soit M une matrice carrée. M est donnée comme suit:

$$M = \begin{pmatrix} a & b \\ c & d \end{pmatrix} \quad (6.38)$$

En algèbre linéaire, l'inverse de la matrice M, notée M^{-1} est calculée comme suit:

$$M^{-1} = \frac{1}{\text{Det}(M)} (\text{CoFactors}(M))^t \quad (6.39)$$

Avec $\text{Det}(M)$ est le déterminant de la matrice M et $(\text{CoFactors}(M))^t$ est la matrice transposée de la matrice des cofacteurs de M.

Le déterminant de la matrice M est donné par la formule 6.40

$$\text{Det}(M) = a * d - c * b \quad (6.40)$$

D'autre part, la matrice des cofacteurs de la matrice M est donnée comme suit:

$$\text{CoFactors}(M) = \begin{pmatrix} d & -c \\ -b & a \end{pmatrix} \quad (6.41)$$

Donc, on pourra écrire

$$(CoFactors(M))^t = \begin{pmatrix} d & -b \\ -c & a \end{pmatrix} \quad (6.42)$$

Notons qu'une matrice M est inversible si et seulement si le déterminant de M est non nul. Donc:

$$Det(M) \neq 0 \quad (6.43)$$

Ainsi, en considérant les formules 6.39, 6.40 et 6.42, on sera en mesure de définir les formules 6.45, 6.46 et 6.47 comme suit:

$$C^{-1} = \frac{1}{Det(C)} (CoFactors(C))^t \quad (6.45)$$

$$Det(C) = (c_3 + c_4 * c_1) * (c_5 * c_2 + c_6) - (c_5 * c_1) * (c_4 * c_2) \quad (6.46)$$

$$(CoFactors(C))^t = \begin{pmatrix} c_5 c_2 + c_6 & -c_4 c_2 \\ -c_5 c_1 & c_3 + c_4 c_1 \end{pmatrix} \quad (6.47)$$

Rappelons que $c_j = p_j$ pour $j \in [1,6]$, $p_j \in P = (p_1 p_2 p_3 p_4 p_5 p_6)$

Une fois que la destination D accomplit la tâche de décodage des symboles reçus relatifs aux six bits natifs, elle remet les bits extraits dans leur ordre initial afin de régénérer le fragment initial. L'opération de remise en ordre repose principalement sur l'utilisation de la séquence $RSP = P = (p_1 p_2 p_3 p_4 p_5 p_6)$. Telle qu'il a déjà été mentionné plus haut dans ce chapitre, la destination D obtient la séquence RSP , cryptée, de la source au début de l'échange des données. Pour mettre les bits en ordre initial, la destination considère que la position du premier bit reçu est p_1 , celle du deuxième bit reçu est p_2 , et ainsi de suite jusqu'à la détermination de la position initiale de chacun des six bits reçus. Cette opération est effectuée pour chaque fragment. Finalement, chaque fragment sera concaténé à son précédent afin de régénérer le paquet natif.

6.4.1.4 Exemple pour le décodage de l'information avec G-SNCDS

On considère l'exemple illustré par les deux figures 6.19 et 6.20. Dans cette section on n'explique que la partie relative au décodage de deux bits natifs a_4 and a_2 Appartenant à un fragment donné du paquet natif illustré par les deux figures citées ci-dessus. L'exemple relatif au décodage de ces bits a été détaillé un peu plus haut dans ce chapitre.

L'extraction des symboles a_4 and a_2 est obtenue via la formule 6.37. Donc on peut déduire la formule 6.48 suivante:

$$\begin{pmatrix} a_4 \\ a_2 \end{pmatrix} = \begin{pmatrix} c_3 + c_4 c_1 & c_4 c_2 \\ c_5 c_1 & c_5 c_2 + c_6 \end{pmatrix}^{-1} \begin{pmatrix} y_2 \\ y_3 \end{pmatrix} \quad (6.48)$$

Avec, $(c_1 \ c_2 \ c_3 \ c_4 \ c_5 \ c_6) = (p_1 \ p_2 \ p_3 \ p_4 \ p_5 \ p_6) = (4, 2, 5, 6, 1, 3)$ (voir figure 6.19)

Donc

$$\begin{pmatrix} a_4 \\ a_2 \end{pmatrix} = \frac{1}{97} * \begin{pmatrix} 5 & -12 \\ -4 & 29 \end{pmatrix} * \begin{pmatrix} y_2 \\ y_3 \end{pmatrix} \quad (6.49)$$

6.4.2 G-SNCDS pour éviter les attaque de confidentialité

La solution G-SNCDS a été conçue pour éviter les attaques de confidentialité dans un réseau WMN. Dans ce cas, elle procède comme l'algorithme SNCDS (Tata and Kadoch 2014). L'objectif de G-SNCDS lorsqu'il est appliqué pour éviter une attaque de confidentialité, est d'empêcher les attaquants internes, comme externes, d'obtenir une information compréhensible parmi celles qui sont transmises dans le réseau WMN.

Afin d'expliquer l'efficacité du G-SNCDS pour contrer les attaques de confidentialité, nous considérons l'attaque illustrée dans la figure 6.21. Cette figure représente deux types d'attaques dans un réseau papillon. Une attaque interne effectuée par l'eavesDropper 1 et une attaque externe effectuée par l'eavesDropper 2.

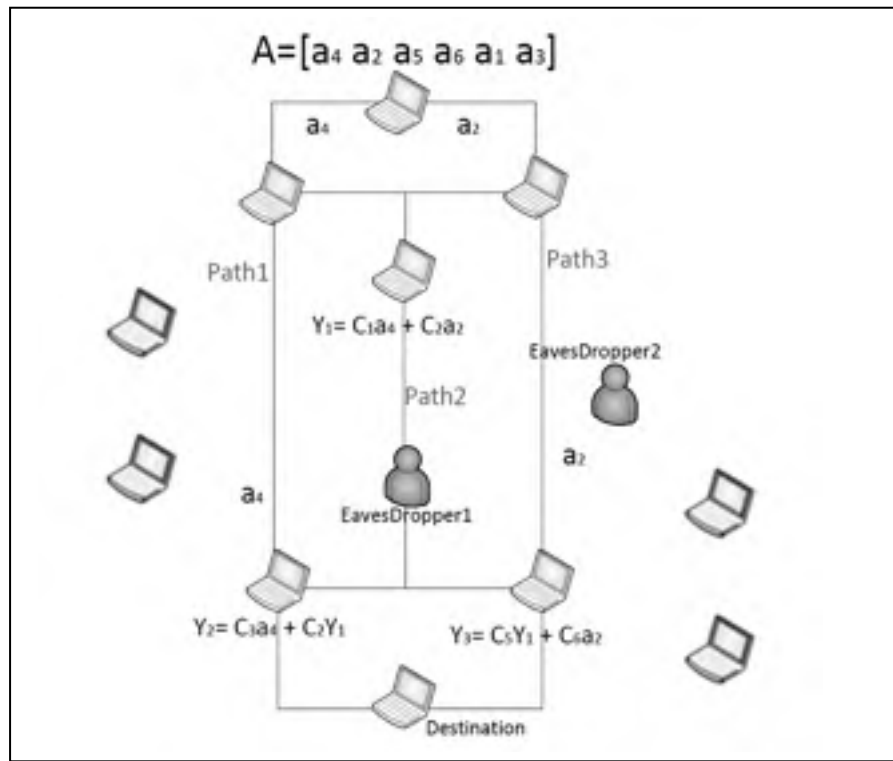


Figure 6.21 Attaques interne et externe de confidentialité dans le WMN

Dans les deux cas, les pirates ne peuvent pas obtenir l'information complète, qui est envoyée de la source vers la destination. Le eavesdropper 1 peut intercepter le symbole codé Y_1 . En revanche, il ne pourra pas résoudre ce symbole, car il n'a pas les codes de la matrice de codage. Dans un tel cas, nous supposons que les attaquants n'ont pas les capacités pour résoudre les clés de chiffrement utilisées par la source et la destination pour coder la séquence RSP. Rappelons que la matrice de codage est générée par la destination en utilisant la séquence RSP cryptée et transmise par la source. D'autre part, le deuxième attaquant ne sera pas capable d'obtenir les bits envoyés par la source via l'autre chemin qui ne lui est pas accessible. Tous les bits d'ordre impair seront transmis via la route inaccessible à cet attaquant. En outre, l'ordre des bits capturés n'est pas identique à l'ordre des bits dans le paquet natif. Par conséquent, il devient compliqué pour l'attaquant de reconstruire le paquet original.

6.4.3 G-SNCDS for data integrity attack avoidance

La figure 6.22 illustre une attaque interne d'intégrité de données dans un réseau WMN. Deux nœuds légitimes sont censés être des nœuds malveillants dans cette architecture. L'attaque d'intégrité considérée dans ce cas est le rejet de paquets par les pirates. En d'autres termes, les pirates suppriment, partiellement ou totalement, les données qu'ils reçoivent. Cette situation provoque l'altération des données transmises de la source à la destination. De ce fait, l'intégrité des données transmises sera affectée.

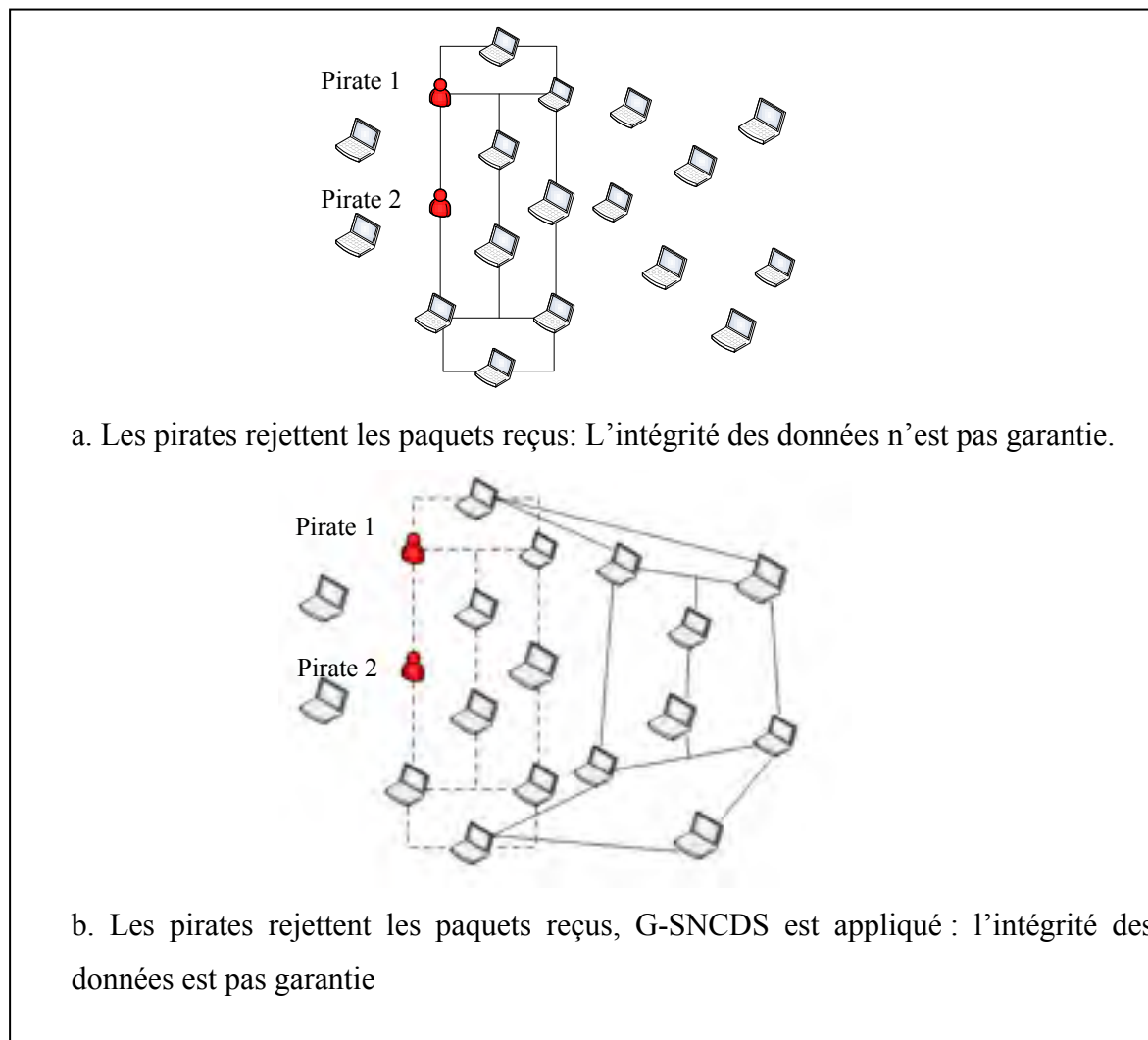


Figure 6.22 Attaque interne d'intégrité dans un réseau WMN

Quand on utilise un seul réseau papillon pour effectuer un codage réseau dans un réseau WMN (figure 6.22.a), l'intégrité de l'information ne sera pas garantie. Cela est dû au fait que l'utilisation du mécanisme DS seul ne suffit pas pour éviter toutes les attaques d'intégrité. DS permet d'envoyer une partie des données sur un chemin, et l'autre partie sur l'autre chemin. En effet, le nombre maximal de bits d'un même paquet qui peuvent être rejetés est réduit de 50%, mais ça reste qu'un autre 50% de bits peut être perdu à cause de l'existence des attaquers. Pour palier ce problème G- SNCDS utilise deux effets papillons pour effectuer le routage de l'information à travers une double opération de codage réseau (figure 6.22.b).

Dans une telle situation, les données sont transmises en utilisant un mécanisme de redondance. Chaque paquet transmis par l'intermédiaire du premier réseau papillon est dupliqué et envoyé par l'intermédiaire du second réseau papillon, de sorte que le G- SNCDS utilise la même séquence RSP pour envoyer le paquet. Si une attaque d'intégrité se produit dans le réseau, la destination la détectera après avoir constaté que certains paquets sont rejetés. Par conséquent, le réseau utilisera les données redondantes pour corriger l'information altérée.

Notons que dans le cas d'une attaque d'intégrité externe, le même processus que celui appliqué à l'attaque de l'intégrité interne sera utilisé pour surmonter le problème de paquets corrompus. De plus, la localisation des pirates n'influe point sur la façon d'appliquer G- SNCDS pour contourner les attaques d'intégrité. Autrement dit, l'utilisation de la redondance à travers deux réseaux papillon est en mesure d'éviter toutes les attaques d'intégrité, quelle que soit la position de l'attaqueur.

6.4.4 G-SNCDS for data availability attack avoidance

Figure 6.24 considère une attaque de Déni de Service (DoS) comme exemple d'attaque de disponibilité. Deux attaquers internes sont représentés dans cette figure.

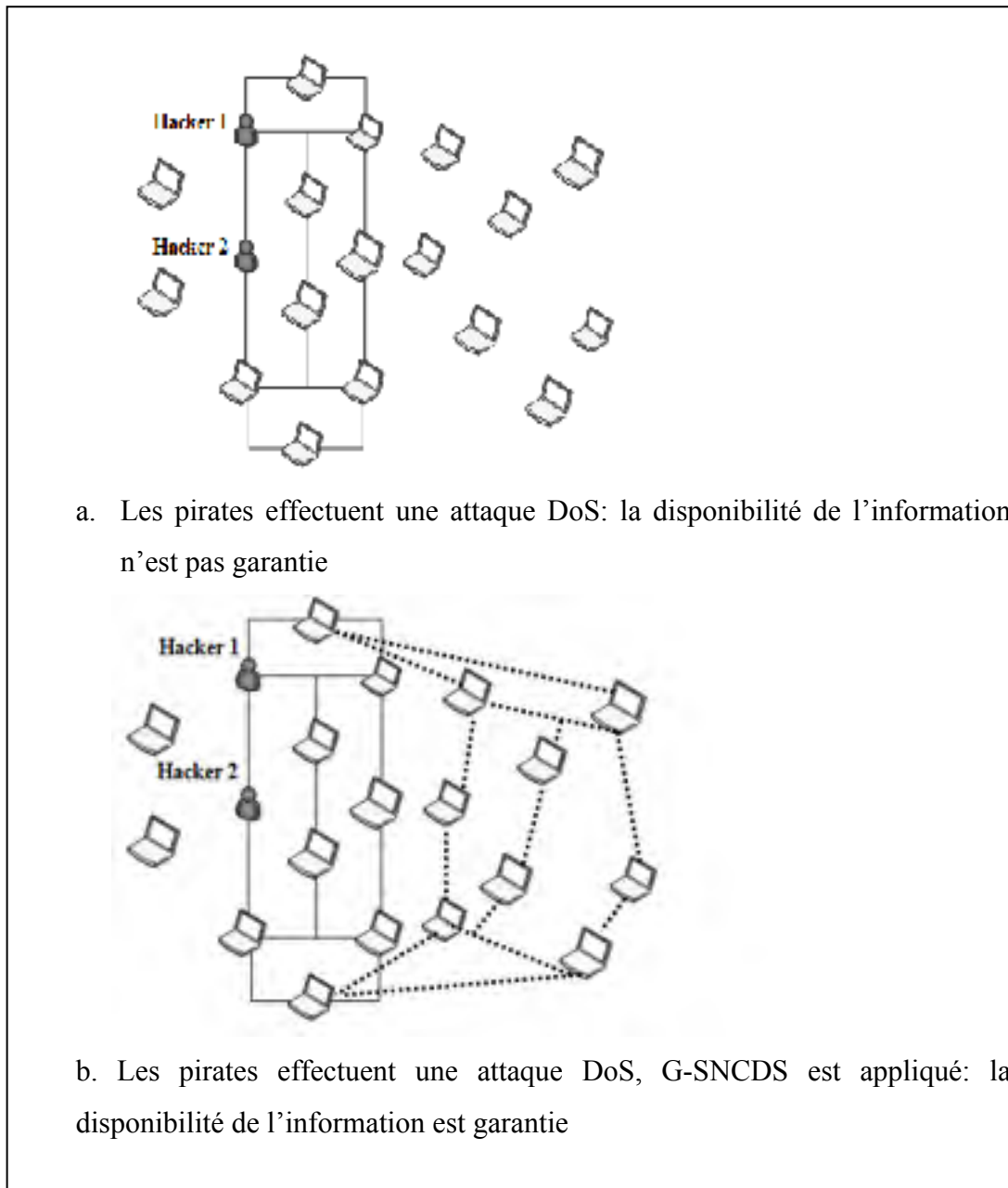


Figure 6.23 Attaque de disponibilité interne dans un réseau WMN

Les deux attaquants sont censés être des nœuds appartenant au réseau papillon utilisé pour la transmission des données entre la source et la destination. De plus, ils sont autorisés à recevoir et transmettre les paquets. La première situation (figure 6.23.a) illustre un cas d'utilisation d'un effet papillon pour transmettre des données d'une source à une destination, passant par deux nœuds malveillants au sein de ce réseau. L'autre situation (figure 6.23.b)

représente un cas d'utilisation de deux effets papillon pour la transmission de données. Ce dernier cas correspond à l'application de G-SNCDS pour pallier les problèmes des attaques de disponibilité. Dans une telle situation, le second réseau papillon est considéré comme un réseau de restauration de la communication (Back-Up). En fait, il ne sera pas utilisé jusqu'à ce que le premier devienne indisponible. Par conséquent, lorsque les pirates effectuent des attaques DoS, la solution illustrée dans la figure 6.23.a ne pourra pas garantir la disponibilité des données. En fait, la transmission sera bloquée en raison de l'attaque DoS, qui rend les ressources du réseau indisponibles. En outre, l'utilisation de G-SNCDS permet de surmonter l'attaque de disponibilité. Dès que le premier réseau papillon est submergé par des requêtes inutiles, via une attaque DoS, G-SNCDS activera le réseau papillon de restauration. Par conséquent, les données ne seront plus acheminées via le premier réseau papillon. Elles emprunteront, plutôt le réseau papillon de restauration qui prend la relève de routage des données avec codage réseau.

Par ailleurs, la résolution des attaques de disponibilité externes s'effectue de la même manière que les attaques internes. Autrement dit, pour toute attaque affectant la disponibilité de l'information dans le réseau WMN, G-SNCDS applique la politique de restauration de réseau papillon en désactivant le réseau papillon en cours et en activant un autre. Il est à noter que l'application de l'algorithme de RBC, par G-SNCDS, a permis la construction d'un ensemble d'effets papillon pour chaque couple de noeuds sources et destination, voulant échanger des données. Cela facilitera le basculement entre différents réseaux papillon si une attaque de disponibilité survient. Cela est bénéfique aussi pour les attaques d'intégrité qui nécessitent l'utilisation de plus d'un effet papillon pour la transmission des données.

Finalement, l'application de l'algorithme G-SNCDS a contribué efficacement à la résolution du problème des attaques de disponibilités dans les réseaux WMN. Cette valeur ajoutée va surement améliorer le niveau de sécurité pour les communications D2D établies dans les petites cellules des réseaux LTE HetNets.

6.4.5 Simulations et résultats

Dans cette section, nous allons procéder à la validation de notre solution G-SNCDS, qui inclut l'algorithme SNCDS. Cette validation se fait via des simulations d'attaque de sécurité dans un réseau WMN en utilisant Matlab. L'objectif de la simulation est de démontrer que notre solution G-SNCDS améliore le niveau de sécurité dans les réseaux WMN sans ajouter de trafic de contrôle supplémentaire dans le réseau. Nos résultats ont été comparés avec le cas où le codage réseau classique est utilisé comme technique de transmission.

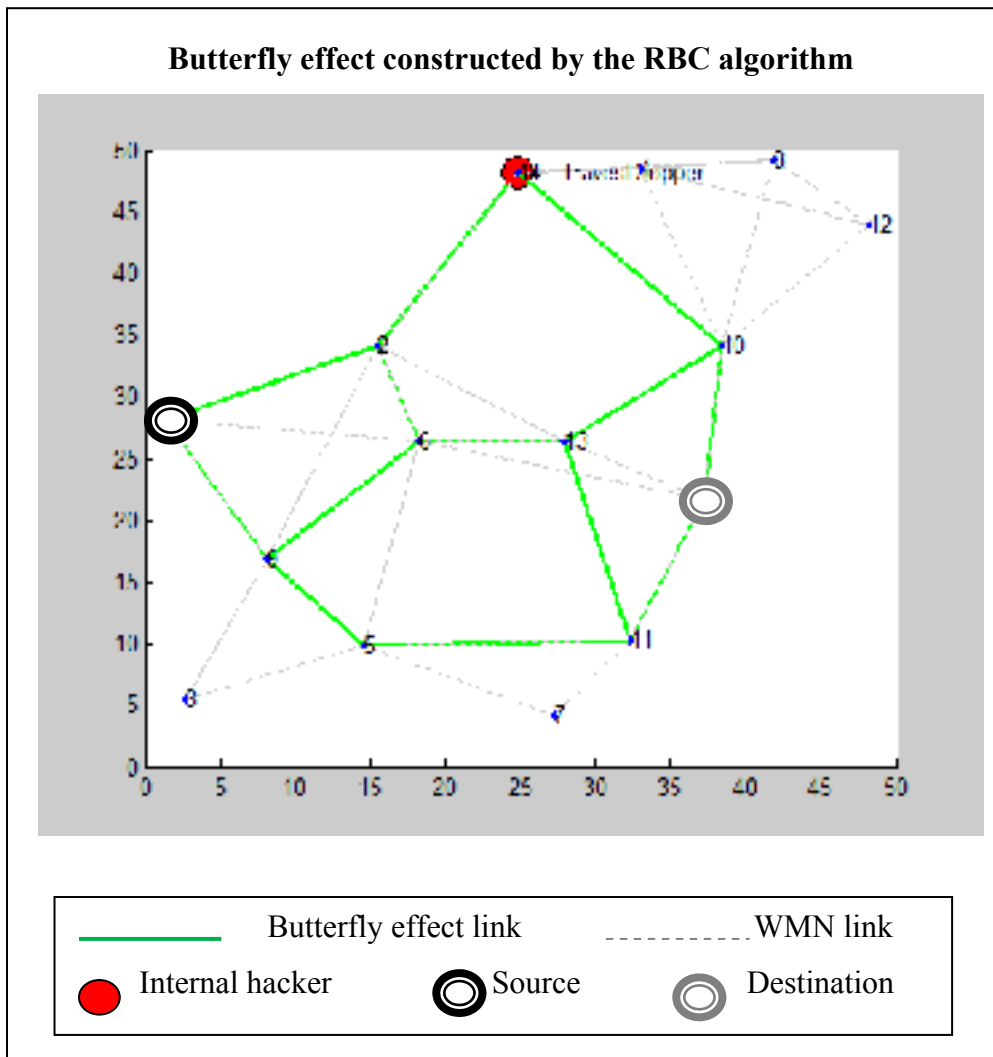


Figure 6.24 Construction d'un effet papillon par le système RBC: Attaque interne

Trois types d'attaque de sécurité ont été considérés dans le présent travail, à savoir l'attaque de confidentialité, l'attaque d'intégrité et l'attaque de disponibilité. Toutes ces attaques sont contournées par G-SNCDS. Notre algorithme utilise le codage réseau combine à la technique de Data Splitting pour éviter les attaques de confidentialité. Par ailleurs, l'utilisation de plusieurs effets papillon permet d'éviter les attaques d'intégrité et de disponibilité. En effet, une transmission de données avec redondance via deux effets papillons permet de détecter les attaques d'intégrité et de procéder la correction d'erreurs. D'autre part, l'utilisation d'un deuxième réseau papillon, comme réseau de restauration, permet d'éviter les attaques de disponibilité.

Par ailleurs, pour chaque scénario le choix de la source, de la destination et du nœud malveillant s'effectue d'une façon aléatoire. De plus, l'algorithme RBC est utilisé dans chaque expérience pour la construction des réseaux papillon à l'intérieur du réseau WMN (figure 6.24 et 6.26). Les résultats des simulations illustrés dans cette étude sont obtenus en simulant des attaques de sécurité au sein des de chaque réseau construit par RBC et représenté par les deux figures 6.24 et 6.26.

Le premier type d'attaque simulée dans cette étude est l'attaque de confidentialité. Deux scénarios sont considérés, notamment une attaque de confidentialité interne via le réseau de la figure 6.25, et une attaque de confidentialité externe via le réseau de la figure 6.27. Notons que pour pallier une attaque de confidentialité, G-SNCDS se comporte comme SNCDS. Il se contente d'utiliser un seul réseau papillon pour la transmission de données de la source vers la destination.

La figure 6.24 illustre le réseau papillon construit par RBC pour examiner l'impact d'une éventuelle attaque interne. Dans ce cas, le hacker est un nœud légitime du réseau WMN, donc du réseau papillon. Dans ce scénario, le pirate pourra recevoir une partie des données transmises de la source vers la destination. Par contre, l'application de G-SNCDS réduit le nombre de bits appartenant à un même paquet et qui sont interceptés par le nœud malveillant. Ceci est illustré dans les deux figures 6.25 et 6.27. Sur ces deux figures on remarque que le

nombre de bits interceptés par le pirate est plus grand dans le cas de l'utilisation de l'approche codage réseau classique (Network Coding ou NC) pour la transmission des données, comparativement à l'utilisation du codage réseau combiné avec G-SNCDS (NC+G-SNCDS). Cette amélioration est obtenue grâce à l'utilisation du mécanisme Data Splitting (DS) qui consiste à scinder le paquet en plusieurs fragments et à mélanger les bits d'un même fragment avant de les transmettre via le réseau papillon. De plus, DS envoie une partie du fragment via un chemin du réseau papillon et l'autre via l'autre chemin. Cet avantage n'est pas offert par le codage réseau classique, où chaque paquet est transmis en entier via un seul chemin. Cela facilite l'interception, par un espion, de l'intégralité d'un paquet transmis s'il n'est pas crypté. Finalement, les techniques adaptées par G-SNCDS ne font donc que réduire le nombre des bits d'un même paquet, interceptés par un pirate.

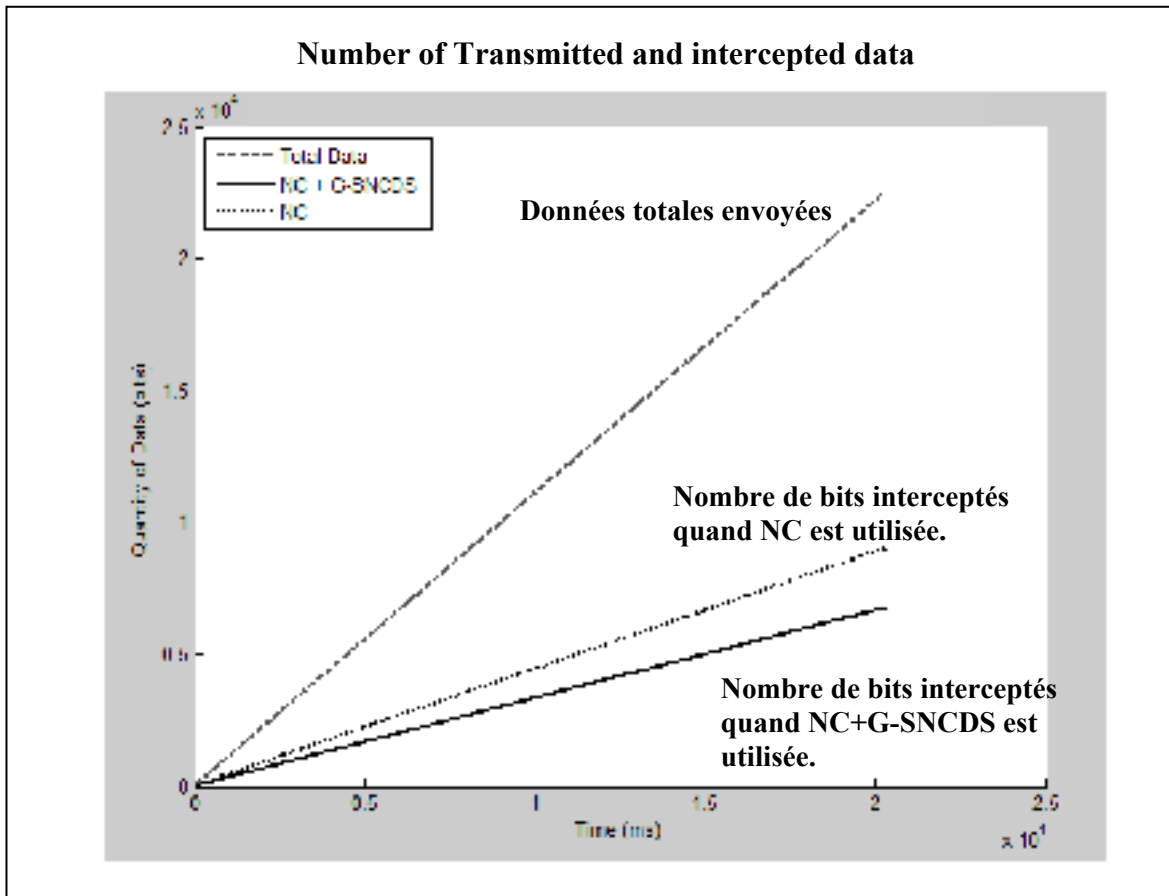


Figure 6.25 Nombre de bits interceptés: Attaque interne de confidentialité

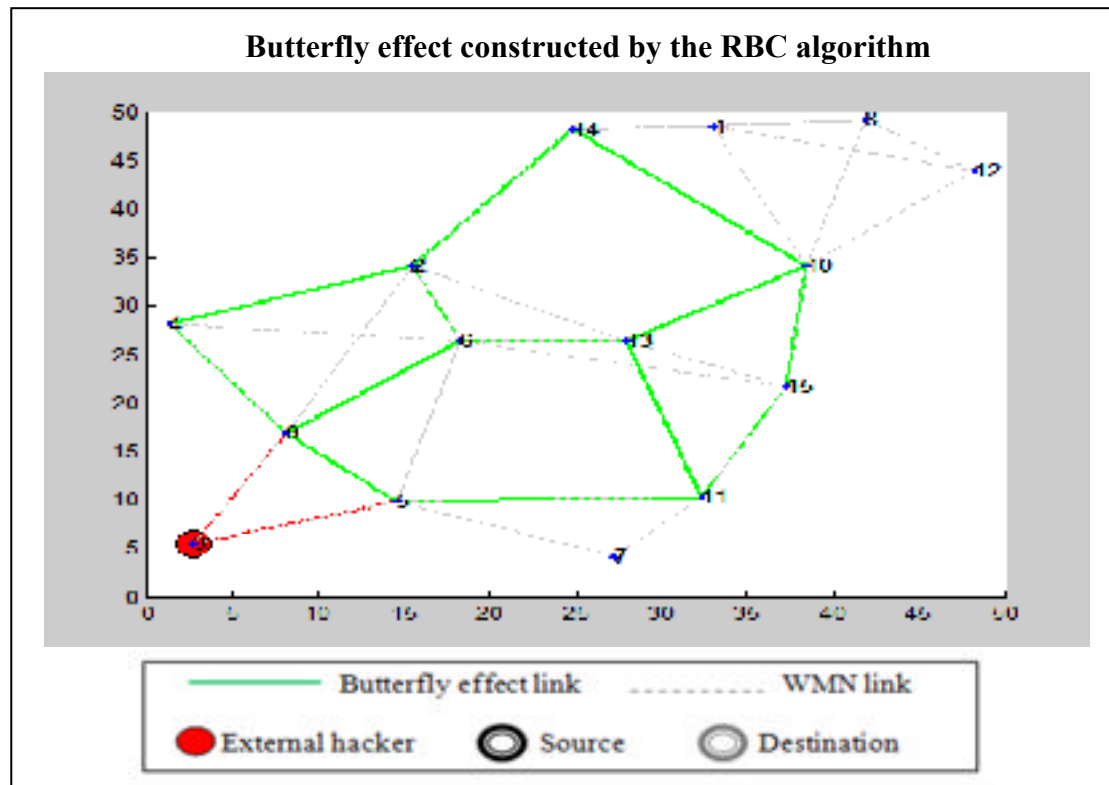


Figure 6.26 Construction d'un effet papillon par le système RBC: Attaque externe

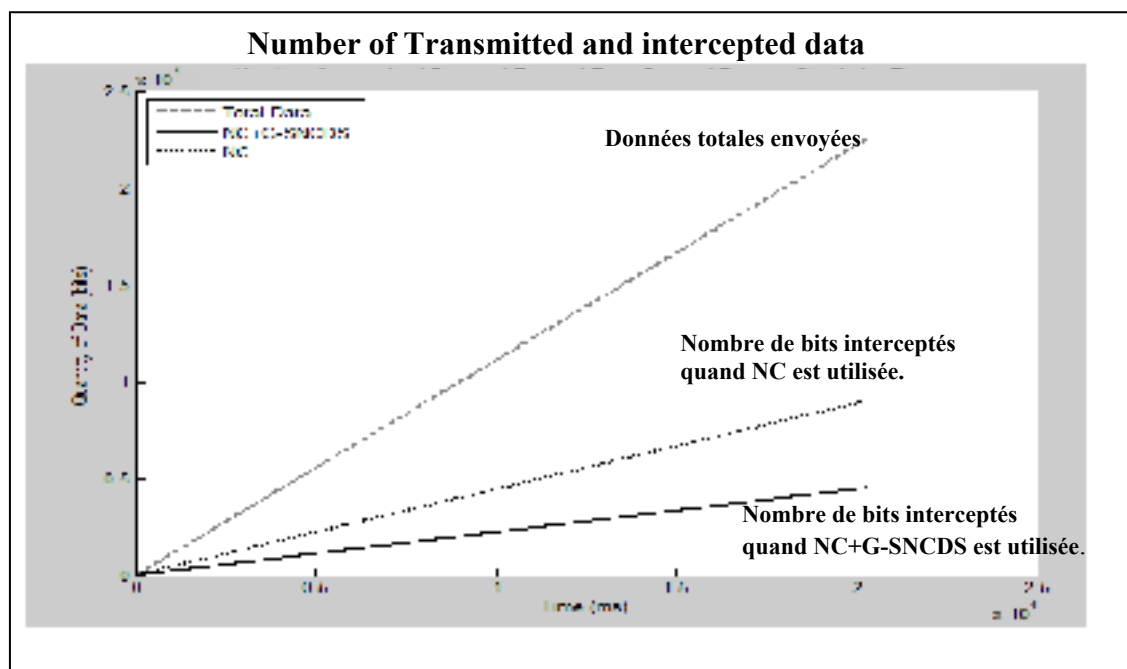


Figure 6.27 Nombre de bits interceptés: Attaque externe de confidentialité

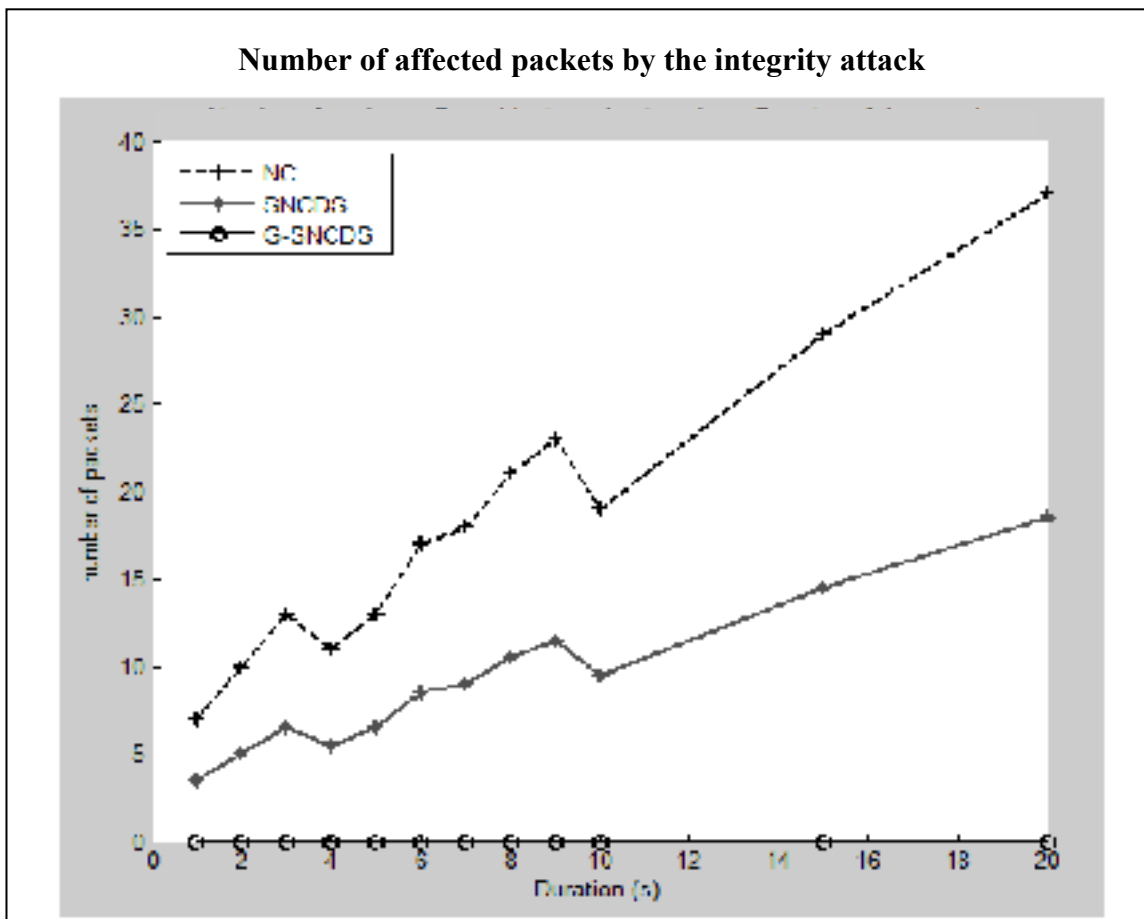


Figure 6.28 Attaque d'intégrité dans un réseau WMN

Le deuxième type d'attaques simulées dans cette étude est l'attaque d'intégrité. L'attaque d'intégrité simulée dans ce scénario est définie par une attaque de rejet de données reçues par les attaquants. Chaque attaquant, interne ou externe, procède à supprimer tous ou une partie des bits qu'il reçoit.

Dans le but de contourner l'attaque de l'intégrité, l'algorithme G-SNCDG effectue une transmission de données basée sur le codage réseau et le data Splitting à travers deux réseaux papillon. Le premier est un réseau papillon principal et le second est utilisé pour réaliser une transmission redondante de la source à la destination. Lors d'une tentative de piratage du réseau, les données interceptées par le pirate seront altérées immédiatement. Par conséquent les informations reçues par la destination via les deux réseaux papillon ne seront pas

semblables. Ainsi, la destination procédera à l'utilisation d'un mécanisme de correction pour obtenir les données natives. Elle peut aussi demander une retransmission de données. La source qui recevra la demande de retransmission de données à cause d'une altération de données va utiliser deux nouveaux autres réseaux papillon afin d'empêcher le pirate de recevoir l'information de nouveau.

Dans cette simulation, nous avons effectué trois expériences. La première se fait en réalisant une simple transmission de codage de réseau par l'intermédiaire d'un seul réseau papillon. Ceci est représenté par le graphe NC dans la figure 6.28. La deuxième consiste à appliquer la solution SNCDS interprétée en effectuant un codage réseau combinée à la technique DS sur un seul réseau papillon. Cette expérience est représentée par le graphique SNCDS de la figure 6.28. Le dernier scénario applique l'algorithme G-SNCDS qui est représenté par le graphe G-SNCDS dans la figure 6.28. Cette solution est réalisée par la transmission de données à travers deux réseaux papillon et en utilisant un codage de réseau et système DS. Les résultats de simulation montrent, d'une part, que l'application de l'algorithme SNCDS améliore le niveau de sécurité par rapport au cas où on utilise uniquement le mécanisme de codage de réseau pour transmettre des paquets. Ceci est possible grâce à l'utilisation de l'approche DS. Lorsque DS est utilisée, seule la moitié des données est transmise par la voie passant par le nœud malveillant. On observe que l'application de G-SNCDS réussit à contourner totalement l'attaque d'intégrité. L'utilisation de deux réseaux papillon pour une transmission redondante de données permet la détection des paquets corrompus. En effet, la destination procède à la comparaison de données transmises via les deux effets papillon à chaque réception. Elle peut exécuter par la suite un algorithme de correction d'erreurs pour obtenir les données d'origine. Cela explique la raison pour laquelle G-SNCDS permet de réduire à zéro le nombre de paquets altérés par les attaques d'intégrité.

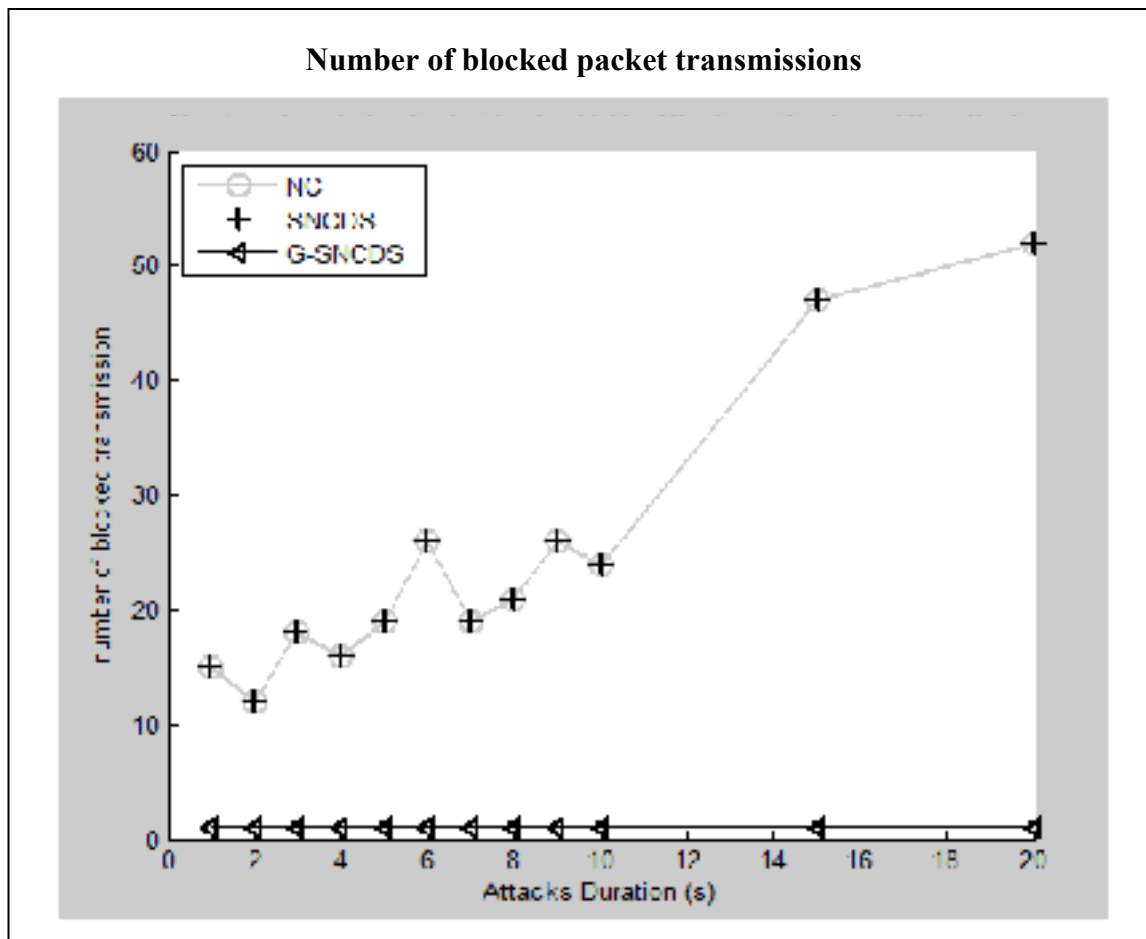


Figure 6.29 Attaque de disponibilité dans un réseau WMN

Le dernier type d'attaques est l'attaque de la disponibilité. Dans cette étude, nous considérons l'attaque de déni de service (DoS) comme une attaque de disponibilité. Les résultats de notre simulation sont illustrés dans la figure 6.29. Comme il a été effectué pour l'attaque de l'intégrité, trois scénarios sont mis en œuvre. Une transmission de données avec codage réseau (NC), une solution NC combiné avec le mécanisme DS qui représente l'algorithme SNCDS et l'approche qui représente l'algorithme G-SNCDS. Les deux premières solutions sont appliquées tel qu'il a été fait dans le scénario de l'attaque d'intégrité. La dernière solution à savoir G-SNCDS est appliquée en utilisant un seul réseau papillon tant qu'aucune attaque ne se manifeste. Dès qu'une attaque de disponibilité se produit, le second réseau de papillon sera activé pour remplacer le premier.

La figure 6.29 montre que les deux méthodes NC et SNCDS donnent les mêmes résultats pour cette expérience. En fait, les deux approches utilisent un seul réseau papillon. Par conséquent, quand une attaque DoS arrive l'ensemble du réseau sera touché et toutes les transmissions seront bloquées. Par ailleurs, le G-SNCDS peut éviter les attaques de disponibilité, telle que l'attaque DoS grâce à la restauration (Back-Up) du réseau de papillon affecté par l'attaque. Ainsi, dès que l'attaque soit détectée dans le réseau papillon principal, un réseau de Back-Up sera activé et utilisé pour relancer le processus de transmission de données. Par ailleurs, le réseau affecté par l'attaque sera désactivé et les paquets bloqués seront retransmis via le réseau Back-Up.

Par ailleurs, bien que d'autres méthodes existent pour sécuriser les transmissions des paquets dans le réseau mesh, il reste que notre approche est la plus pertinente, quand les ressources sont limitées dans le réseau. En effet, bien que le chiffrement soit une solution intéressante pour sécuriser les données acheminées dans le réseau, mais son utilisation exige d'ajouter des bits additionnels au paquet natif afin de le crypter. Les auteurs de l'article (Ganesan, Venugopalan et al. 2003) évaluent l'overhead généré par le chiffrement pour de nombreux algorithmes de cryptages. En fait, les bits ajoutés par le mécanisme de cryptage vont augmenter l'overhead dans le réseau et des ressources supplémentaires seront nécessaires pour garantir un bon niveau de QoS dans le réseau. Toutefois, fournir des ressources additionnelles lors des moments de congestion n'est pas une tâche facile. Une contribution importante de l'algorithme G-SNCDS sa capacité d'acheminer les paquets de la source vers la destination, en toute sécurité et sans ajouter de trafic de contrôle. De plus, G-SNCDS garantit un bon niveau de sécurité pour les transmissions D2D sans ajouter de trafic de contrôle autre que celui utilisé par le système de codage réseau. G-SNCDS utilise plutôt un mécanisme de fractionnement des données, combiné à un processus de mélange de bits pour éviter les attaques confidentialité, d'intégrité et de disponibilité des réseaux WMN. Rappelons que SNCDS, représentant la partie de G-SNCDS qui traite les attaques de confidentialité offre les mêmes avantages que G-SNCDS quand il est appliqué pour contrer les attaques de confidentialité.

En outre, l'approche G-SNCDS combinée à la solution RBC est une solution évolutive pour le routage sécurisé au sein du réseau WMN. En fait, la construction d'effets papillon dépend principalement de la disponibilité de chemins reliant les enfants et les petits-enfants de la source et de la destination. Il est évident que plus le nombre de nœuds du réseau WMN est important, plus le nombre d'effets papillon construits par les RBC est important. En plus, quel que soit le nombre de nœuds dans le réseau WMN, la solution G-SNCDS est suffisante pour assurer un bon niveau de sécurité, tant que la construction de réseaux de papillon est possible.

6.5 Conclusion

Les ressources radio étant très limitées dans les macros cellules, une solution a été développée dans le chapitre précédant afin d'améliorer l'allocation des ressources radio en offrant l'accès avec priorité aux premiers répondant du réseau « Sécurité Publique ». Un autre aspect de notre solution détaillée dans le chapitre précédant est d'assurer une opération d'offloading des bearers arrivant à la macro cellule LTE vers les petites cellules. Une fois que les bearers basculés vers les fréquences publiques ils devront être établis selon le mode de fonctionnement des réseaux locaux sans fil opérant dans ses fréquences. Afin de préserver un niveau de QoS acceptable pour les trafics transitant via les fréquences publiques, on a développé deux solutions de routage de paquets. La première solution est nommée RBC. Elle consiste à créer des réseaux papillon au sein des réseaux WMN afin de permettre l'application du codage réseau lors des communications D2D. Le codage réseau est une technique de transmission de paquets qui contrairement au routage conventionnel procède au mixage des paquets en provenance de plusieurs sources pour les transmettre simultanément. Cette technique apporte beaucoup d'avantages pour les transmissions de données dans les réseaux filaires et sans-fil. Plusieurs études ont montré l'amélioration de la performance des réseaux en l'utilisant (Ho, Koetter et al. 2003, Li and Li 2004, Gkantsidis and Rodriguez 2005, Fragouli, Widmer et al. 2006, Katti, Rahul et al. 2006, Matsuda, Noguchi et al. 2011). Cette performance apparait dans l'augmentation du débit (Throughput), la réduction des délais de bout en bout et en la minimisation des taux de perte de paquets. Nous avons aussi

démontré avec l'utilisation du RBC, que le codage réseau peut être utilisé pour assurer le Load Balancing et la transmission de données avec redondance. Celles-ci sont deux techniques pertinentes pour rendre le routage fiable et robuste. D'autre part, ce sont deux techniques que nous proposons pour améliorer la sécurité dans les petites cellules.

La deuxième solution de routage que nous proposons pour la transmission des données D2D est LBS-AOMDV. Cette solution vient compléter la solution RBC et consolider notre approche globale de routage dans les petites cellules. En effet, il ne faut pas oublier que la construction des réseaux papillon par RBC nécessite d'avoir une topologie WMN spécifique. Cette topologie bien qu'elle ne soit pas spéciale ou unique, mais elle doit assurer l'existence de certains chemins reliant les nœuds enfants de la source de données au nœud enfants de sa destination. Ceci dit, un troisième chemin doit être faisable. Il s'agit du lien qui relie le codeur au décodeur d'informations. Rappelons que la nature de l'architecture du WMN augmente les chances de trouver non seulement un réseau papillon, mais souvent plusieurs. Toutefois, RBC se montre dans certains cas incapable de construire des réseaux papillon, car la topologie du réseau ne le permet pas. Par conséquent, et dans le but d'apporter une amélioration permanente, au routage des données effectué dans les petites cellules, LBS-AOMDV a été développée comme solution alternative à RBC. Les résultats des simulations correspondantes à LBS-AOMDV ont démontré la pertinence de son application pour l'amélioration du routage des données dans le réseau WMN en les comparants à une solution répondue, à savoir l'algorithme AOMDV.

Par ailleurs, une solution de sécurité des communications D2D a été développée dans ce chapitre. La motivation derrière ce travail est la différence des niveaux de sécurité des réseaux LTE et des réseaux locaux sans fil. Il est clair que les communications qui transitent via les fréquences publiques sont plus vulnérables aux attaques internes et externes, que celles qui sont transmises via les réseaux utilisant des fréquences privées. Or, les bearers basculés vers les petites cellules doivent bénéficier du même niveau de sécurité offert dans les macros cellules. Ceci dit, nos deux solutions SNCDS et G-SNCDS viennent résoudre cette problématique afin d'éviter les attaques de confidentialité, d'intégrité et de

disponibilité. Les résultats obtenus ont démontré l'efficacité de nos solutions et leur pertinence pour l'amélioration de la sécurité au niveau des petites cellules. Notons que le travail mené dans cette recherche ne sécurise pas les transmissions D2D via le multipath. Il s'agit d'une étude qui rentre dans l'ensemble de nos travaux futurs.

Finalement, ce chapitre a présenté des solutions de routage et de sécurité pour l'amélioration de la qualité des communications D2D au niveau des petites cellules LTE HetNet. Ces cellules ont fortement contribué à l'amélioration de la QoS des différents trafics transmis dans les réseaux LTE HetNets. Ces contributions se montrent dans l'offre de nouvelles ressources radio lors des moments de congestion des macros cellules, ainsi que dans la couverture des zones éloignées non atteintes par les macros cellules LTE, surtout dans le cas de la gestion des catastrophes et des désastres.

CONCLUSION

Bien que l'application de la gestion anticipative des catastrophes, par les premiers répondants, puisse apporter ses fruits pour réduire l'impact matériel et surtout humain de ces événements inattendus, sa réalisation ne peut être qualifiée que de très restreinte en l'absence des moyens de communication. Une solution qui a fortement contribué à l'amélioration de la gestion des crises et à la concrétisation de leur gestion anticipative est le déploiement du réseau de la sécurité publique (PSN) pour assurer la communication et l'échange de l'information entre les premiers répondants. Ce réseau est implanté parallèlement avec le réseau commercial dans les réseaux LTE. Il utilise les bandes de fréquences 700 MHz en Amérique du Nord, et 800 MHz en Europe. Cependant, l'accroissement de la quantité de l'information échangée à travers le réseau de la sécurité publique rend les ressources radio de ce réseau insuffisantes pour assurer une meilleure gestion des désastres. Pour cette raison, le réseau LTE a permis le partage de la bande radio commerciale entre les clients commerciaux et les premiers répondants. Il s'agit donc d'une invention qui se montre prometteuse pour l'amélioration de la QoS dans le réseau de la sécurité publique. En revanche, les ressources radio sont connues par être très restreintes et coûteuses, il devient donc impératif d'optimiser leur allocation en minimisant leur sous-utilisation. D'autre part, les clients commerciaux bien qu'ils transmettent des données moins importantes que celles échangées par les premiers répondants, ils ont toujours le droit à avoir un service satisfaisant. Le défi à relever est donc de développer une solution qui permet de fournir des ressources supplémentaires pour les clients du réseau de la sécurité publique, au sein de la bande de fréquence commerciale, tout en assurant une bonne QoS pour les clients commerciaux. Cette solution doit aussi garantir un service avec priorité pour les premiers répondants.

Par ailleurs, il ne faut pas oublier que les ressources radio des réseaux cellulaires sont restreintes. Il sera donc difficile d'empêcher la congestion dans ces réseaux surtout avec l'ascension continue du nombre des abonnés mobiles. De ce fait, la réussite de la mission critique des premiers répondants va dépendre fortement de la résolution du problème de la pénurie des ressources.

D'autre part, le basculement des bearers de la macro cellule vers la petite cellule semble une solution efficace pour rendre disponible des ressources radio supplémentaires. Or, dans le cas où les réseaux des petites cellules opèrent dans des fréquences publiques, le problème de la sécurité et celui du routage efficace surgissent.

En outre, les trafics générés par les stations de base eNodeB sont transportés par des réseaux d'agrégation nommés les réseaux Backhaul. Ces réseaux sont pour la plupart de type filaire, ils assurent le lien entre le réseau d'accès et le réseau cœur LTE. Étant donné que les Backhaul peuvent transporter des trafics de différentes priorités en provenance de plusieurs cellules LTE, la pénurie en ressources de bande passante sera très probable. La problématique de gestion des ressources de bande passante apparaît aussi au niveau du réseau cœur LTE, à cause de l'augmentation du nombre des flux y transitant et de la différence entre les exigences en termes de QoS de ces flux. Par conséquent, l'implémentation d'un modèle de gestion des ressources de bande passante au niveau du réseau cœur et au niveau du réseau Backhaul LTE s'avère donc nécessaire.

Dans cette thèse, un nouveau modèle novateur a été conçu pour améliorer la performance dans les réseaux de la sécurité publique sur les réseaux LTE hétérogènes (HetNets). Ce modèle ne néglige toutefois pas la performance du réseau commercial. Les chapitres 4, 5 et 6 présentent notre nouvelle approche pour construire ce modèle. Le chapitre 4 développe une nouvelle approche pour l'allocation des ressources de bande passante avec contraintes dans le réseau Backhaul et le réseau cœur LTE. Le chapitre 5 présente un nouveau modèle pour la gestion efficace des ressources radio dans LTE HetNets et le chapitre 6 modélise une nouvelle solution pour l'amélioration du routage et de la sécurité dans les petites cellules LTE HetNets.

Le premier aspect de notre proposition consiste dans le développement des solutions de gestion de ressources radio et de ressources de bande passante dans le réseau d'accès, le réseau Backhaul et le réseau cœur LTE pour améliorer l'utilisation de ces ressources. Les modèles CPA et CPAwO ont été développés pour une meilleure gestion des ressources radio.

Ces algorithmes se basent tous les deux sur des mécanismes de contraintes d'allocations que nous avons proposés dans ce travail, notamment le CAMF, G-CAMF et RUS-CAMF. Il s'agit de modèles mathématiques qui garantissent le respect des conditions initiales d'allocation des ressources radio pour chaque classe de trafic. Donc, chaque classe de trafic n'allouera que la quantité de ressources qui lui est permise. Les résultats de la simulation montrent que le modèle CPA apporte une amélioration dans la gestion des ressources radio en assurant l'accès avec priorité aux premiers répondants et en retardant l'interruption des usagers commerciaux. L'autre approche, CPAwO se montre plus performante en offrant les avantages de CPA tout en améliorant la QoS des trafics PS. CPAwO réduit à zéro le nombre des bearers interrompus de tout type. Il fait de même pour la préemption. En contrepartie, il augmente le nombre de bearers actifs dans le réseau pour tous les types de trafics.

Pour les mêmes raisons, la gestion des ressources de bande passante doit respecter certaines contraintes d'allocation suivant les priorités et exigences de chaque type de trafic. Pour cela, le modèle CAM a été développé pour qu'il soit appliqué au niveau du réseau Backhaul et du réseau cœur LTE. Rappelons que les modèles CAMF, G-CAMF et RUS-CAMF sont basés sur le modèle CAM, qui est conçu pour les réseaux MPLS. La simulation démontre la performance de notre solution par rapport aux deux approches RDM (Russian dolls bandwidth allocation model) et MAM (Maximum Allocation Model).

Rappelons que le principe de la courtoisie a été adapté pour la gestion des ressources radio et celles de la bande passante afin de donner une certaine priorité au trafic défavorisé. Notons que même si la courtoisie appliquée pour la gestion de la bande passante et celle utilisée pour les fréquences radio dérivent d'une même idée, qui est de céder sa place généreusement quand cela est possible, la structure de chaque algorithme est différente pour répondre spécifiquement à la politique de gestion des ressources, soit radio ou bien bande passante.

Un autre aspect de notre solution développée dans cette thèse consiste à assurer une opération d'offloading des bearers arrivant à la macro cellule LTE vers les petites cellules, générant ainsi des communications Device-to-Device (D2D). Une fois que les bearers sont basculés

vers les fréquences publiques ils devront être établis selon le mode de fonctionnement des réseaux locaux sans fil opérant dans ces fréquences. Cette opération est tout à fait transparente aux clients LTE. Par conséquent, des niveaux de sécurité et de QoS acceptables doivent être fournis par les petites cellules servant les clients LTE.

Les communications D2D apparaissent comme une solution pertinente pour le réseau de sécurité publique (PSN) lorsque la couverture des macros cellules LTE ne pourra pas atteindre la zone de crise, ainsi que quand les ressources eNodeB deviennent limitées. La technologie D2D utilise les réseaux opérants dans les bandes de fréquences sans licence comme le WiFi, ad hoc et les réseaux WMN. Par conséquent, cette solution se montre moins coûteuse et plus facile à déployer par rapport à beaucoup d'autres solutions utilisant des réseaux cellulaires. Cependant, le développement de cette technologie doit faire face au défi de garantir la sécurité. Par conséquent, il est important de pouvoir concevoir une solution D2D sécurisée pour éviter les attaques de confidentialité, d'intégrité et de disponibilité qui pourront survenir dans les petites cellules LTE offrant les fréquences publiques. Dans cette thèse un nouveau système, appelé Generalized Secure Network Coding based Data Slitting algorithm (G-SNCDS), est développé pour assurer une transmission de données sécurisée lors des échanges D2D sur les réseaux LTE hétérogènes (HetNets). Notre approche consiste à effectuer du codage réseau (Network Coding, NC) basé sur le fractionnement de données (Data Spritting, DS) pour la transmission des paquets. Cette solution a été développée dans un premier temps pour éviter les attaques de confidentialité, sous le nom de Secure Network Coding based Data Splitting algorithm (SNCDS). SNCDS est une partie intégrante de G-SNCDS.

Par ailleurs, les transmissions NC sont effectuées en utilisant des effets papillons. Les effets papillons sont construits en utilisant l'algorithme Reliable Butterfly Construction algorithm (RBC). Rappelons que plusieurs échecs de décodage de paquets sont causés par la topologie du réseau. En fait, pour être en mesure de décoder le message, le décodeur doit obtenir des paquets à partir de deux routes différentes. C'est pourquoi ce ne sont pas toutes les topologies qui sont aptes à décoder les paquets codés. Ainsi, nous proposons une solution où la

transmission de données D2D dans le réseau soit effectuée via le réseau papillon afin d'assurer que les paquets codés seront décodés avec succès. Rappelons que certaines solutions proposées dans la littérature (voir chapitre 3) proposent d'intégrer les codes dans le paquet codé. Cette solution bien que simple et assure le décodage des informations, elle n'assure pas la sécurité dans le réseau. Tout nœud malveillant se trouvant dans le rayon de couverture du nœud transmetteur ou du nœud récepteur pourra intercepter le paquet codé émis et pourra extraire les paquets natifs en utilisant les codes intégrés dans le paquet codé.

Dans le cas où les effets papillons n'existent pas dans le réseau WMN, alors l'algorithme Load Balancing based Selective-AOMDV (LBS-AOMDV) est appliqué afin de construire un multipath entre la source de données et sa destination. L'algorithme LBS-AOMDV est un algorithme de routage multipath réactif. LBS-AOMDV assure la QoS à travers l'application du Load Balancing dans un réseau Ad hoc sans fil. Il est basé sur AOMDV, en terme d'échanges de paquets RREQ et RREP pour la construction de routes. LBS-AOMDV apporte une amélioration par rapport à AOMDV en termes de réduction de la quantité du trafic de contrôle servant au calcul des routes du multipath. Notre algorithme est recommandé pour le cas d'un réseau ad hoc afin d'augmenter la probabilité pour qu'un nœud parent ait un enfant participant à la connexion de la source à la destination.

Finalement, les différentes solutions proposées dans cette thèse ont chacune apporté une amélioration pour la performance du réseau PSN sur LTE HetNets. L'intervention sur les trois réseaux de LTE, à savoir le réseau d'accès, le réseau Backhaul et le réseau cœur, a permis d'offrir une solution complète pour garantir un bon niveau de QoS de bout en bout pour les réseaux de la sécurité publique sur LTE. En outre, l'utilisation de la courtoisie a permis d'offrir un certain privilège au réseau commercial quand la QoS du PSN le tolère.

RECOMMANDATIONS

Parmi les solutions présentées dans ce travail, nous avons proposé un modèle d'allocation de ressources radio opérant dans une seule macro cellule conjointement avec zéro, une ou plusieurs petites cellules. La technologie LTE permet l'échange du trafic du plan usager et le trafic du plan de contrôle entre les stations de base eNodeBs. Cet échange s'effectue via les interfaces X2. En effet, il sera intéressant d'étendre notre solution en développant un nouveau modèle pour le basculement des bearers entre les différentes macros cellules en se basant sur un modèle mathématique de contraintes. Les contraintes peuvent être résumées dans les conditions de basculement de trafic d'une cellule à une autre, comme elles peuvent être représentées par les types des trafics à basculer ou le nombre de cellules qui vont coopérer dans cette tâche. L'ensemble de nos algorithmes et modèles d'allocations de ressources radio développés dans cette thèse sera donc adapté pour fonctionner sur plusieurs macros et petites cellules radio. Cette technique va assurer un Load Balancing dans le réseau LTE HetNet en réduisant la charge dans les cellules congestionnées et en évitant la sous-utilisation des ressources des cellules moins congestionnées.

Une autre solution qui peut apporter ses fruits pour réduire l'impact de la congestion dans le réseau d'accès LTE est d'utiliser les communications D2D conjointement avec la radio cognitive. Un usager d'une cellule moins congestionnée (cellule cible) peut jouer le rôle d'un nœud relai pour un ou plusieurs usagers dépendants d'une cellule adjacente (cellule d'origine) souffrant d'un manque de ressources radio. Ces utilisateurs vont former un cluster pour pouvoir effectuer des communications D2D et le nœud relai va être désigné comme Cluster Head (CH). D'autres usagers de la cellule cible peuvent être invités à joindre le cluster. À travers le nœud relai, les utilisateurs de la cellule d'origine peuvent communiquer avec les usagers de la cellule cible qui utilisent les bandes de fréquences libres de la cellule cible. Cela est faisable grâce à la technologie Radio Cognitive. Notons que nœud CH utilisera deux interfaces de transmission, à savoir l'interface du réseau local pour communiquer avec les membres de son cluster qui sont des usagers de la cellule d'origine et l'interface radio cellulaire pour échanger les informations avec les usagers de la cellule cible qui sont aussi

membres de son cluster. Tous les nœuds de la cellule cible peuvent jouer le rôle de relai entre les usagers de la cellule d'origine et le eNodeB de la cellule cible. Cette solution permettrait l'équilibrage de la charge entre les deux macros cellules impliquées, ainsi que l'utilisation partielle des fréquences privées. Notons que les communications D2D effectuées dans les fréquences privées, grâce à la radio cognitive, seront de meilleure qualité que celles effectuées dans les fréquences publiques.

D'autre part, le modèle de transmissions D2D proposé dans cette thèse ne garantit la sécurité qu'en cas d'utilisation du codage réseau. Dans le futur, il serait intéressant de penser au développement d'une solution pour sécuriser les données transmises via le multipath construit par l'algorithme LBS-AOMDV.

Par ailleurs, une piste intéressante qui peut être explorée est l'adaptation du mécanisme Data Splitting pour l'appliquer sur le réseau Backhaul LTE. En effet, la technologie MPLS est fortement recommandée pour qu'elle soit utilisée dans ce réseau. De plus, MPLS applique souvent l'ingénierie du trafic (Traffic Engineering, TE) pour la transmission de données. Comme la technologie TE se base sur le calcul d'un nouveau chemin au moment de congestion, il serait donc intéressant d'appliquer le mécanisme DS lors du calcul d'un nouveau chemin. Son application sera identique à son utilisation dans le multipath.

Finalement, une proposition qui pourrait améliorer l'étude effectuée dans cette thèse est d'étendre les simulations de la solution d'allocation des ressources radio CPAwO pour prendre en considération la différenciation du trafic au sein du même réseau, PS ou CN et lors de la même situation d'utilisation, Emergency ou Non-Emergency. En d'autres termes, nos simulations vont traiter les trafics GRB différemment des trafics Non GBR appartenant au même réseau et utilisés pour une même situation. Pour ce faire, les 15 valeurs de ARP vont être prises en considération et le modèle G-CAMF sera appliqué. Notons que dans cette thèse CPAwO a été simulé conjointement avec le modèle RUS-CAMF.

ANNEXE I

Publications

Journals

- C. Tata and M. Kadoch. (2014) *Efficient Priority Access to the Shared Commercial Radio with Offloading for Public Safety in LTE Heterogeneous Networks*, Computer Networks and Communications Journal (In Press).
- C. Tata and M. Kadoch. (2014) *Generalized Secure Network Coding based Data Splitting for D2D Transmissions for Public Safety over LTE Heterogeneous Networks*, submitted to the IEEE Transactions on Mobile Computing.

Conférences internationales

- C. Tata and M. Kadoch. (2014) *Courteous Priority Access to the Shared Commercial Radio for Public Safety in LTE Heterogeneous Networks*. The 2nd International Conference on Future Internet of Things and Cloud (FiCloud-2014) 27-29 August 2014, Barcelona, Spain (Accepetd);
- C. Tata and M. Kadoch.(2014) *Multipath Routing Algorithm for Device-to-Device Communications for Public Safety over LTE Heterogeneous Networks*. The 1st IEEE International Conference on Information and Communication Technologies for Disaster Management (ICT-DM 2014) 24-25 march, Algiers, Algeria;
- Tata, C and Kadoch, M. (2014). *Secure Network Coding based Data Splitting for Public Safety D2D Communications over LTE Heterogeneous Networks*. 8th International Conference on Communications and Information Technology (CIT '14). 10-12 January. Spain;
- Tata, C and Kadoch, M. (2013). *RBC: Reliable Butterfly Network Construction Algorithm for Network Coding in Wireless Mesh Network*. Proceedings of the 13th International Conference on Applied Informatics and Communications (AIC '13) Valencia, Spain, World Scientific and Engineering Academy and Society (WSEAS): p.p. 291-296.
- C. Tata and M. Kadoch. (2013) *CAM: Courteous Bandwidth Constraints Allocation Model*, in *ICT 2013, The 20th International Conference on Telecommunications* Casablanca, Morocco, 2013.

BIBLIOGRAPHIE

- Aboba, B., L. Blunk, J. Vollbrecht, J. Carlson and H. Levkowitz (2004). Extensible authentication protocol (EAP), RFC 3748, June.
- Adami, D., C. Callegari, S. Giordano and M. Pagano (2008). A New NS2 Simulation Module for Bandwidth Constraints Models in DS-TE Networks, IEEE.
- Adami, D., C. Callegari, S. Giordano, M. Pagano and M. Toninelli G-RDM: A New Bandwidth Constraints Model for DS-TE Networks, IEEE.
- Ahlsweide, R., N. Cai, S.-Y. Li and R. W. Yeung (2000). "Network information flow." Information Theory, IEEE Transactions on 46(4): p.p.1204-1216.
- Al-Hourani, A. and S. Kandeepan (2013). Temporary Cognitive Femtocell Network For Public Safety LTE. Computer Aided Modeling and Design of Communication Links and Networks (CAMAD), 2013 IEEE 18th International Workshop on, IEEE.
- Ali, N. A., A.-E. Taha and H. Hassanein (2013). "Quality of service in 3GPP R12 LTE-advanced." Communications Magazine, IEEE 51(8).
- Andreev, S., A. Pyattaev, K. Johnsson, O. Galinina and Y. Koucheryavy (2014). "Cellular Traffic Offloading onto network-assisted device-to-device connections." Communications Magazine, IEEE 52(4): p.p.20-31.
- Asadi, A. and V. Mancuso (2013). WiFi Direct and LTE D2D in action. Wireless Days (WD), 2013 IFIP, IEEE.
- Astely, D., E. Dahlman, G. Fodor, S. Parkvall and J. Sachs (2013). "LTE release 12 and beyond [Accepted From Open Call]." Communications Magazine, IEEE 51(7).
- Blom, R., P. de Bruin, J. Eman, M. Folke, H. Hannu, M. Naslund, M. Stlnacke and P. Synnergren (2008). Public safety communication using commercial cellular technology. Next Generation Mobile Applications, Services and Technologies, 2008. NGMAST'08. The Second International Conference on, IEEE.
- Borkar, S., D. Roberson and K. Zdunek (2011). Priority Access for public safety on shared commercial LTE networks. Telecom World (ITU WT), 2011 Technical Symposium at ITU, IEEE.
- Bouguen, Y., E. Hardouin and F.-X. Wolff (2012). LTE et les réseaux 4G, Editions Eyrolles.
- Chadchan, S. and C. Akki (2011). "Priority-Scaled Preemption of Radio Resources for 3GPP LTE Networks." International Journal of Computer Theory and Engineering 3(6).

- Chen, Q., S. Zhao and S. Shao (2013). QoS-based resource allocation scheme for Device-to-Device (D2D) communication underlying cellular network in uplink. Signal Processing, Communication and Computing (ICSPCC), 2013 IEEE International Conference on, IEEE.
- Cisco (2011) "Extending MPLS Across the End-to-End Network: Cisco Unified MPLS."
- Din, N. M., H. Hakimie and N. Fisal (2007). "Bandwidth sharing scheme in DiffServ-aware MPLS networks."
- Doumi, T., M. F. Dolan, S. Tatesh, A. Casati, G. Tsirtsis, K. Anchan and D. Flore (2013). "LTE for public safety networks." Communications Magazine, IEEE 51(2): p.p.106-112.
- Ergen, M. (2009). Mobile broadband- Including WiMAX and LTE - Springer Verlag.
- Fang-Chun, K., T. Kun, L. XiangYang, Z. Jiansong and F. Xiaoming (2009). XOR rescue: exploiting network coding in lossy wireless networks. 2009 6th Annual IEEE Communications Society Conference on Sensor, Mesh and Ad Hoc Communications and Networks, 22-26 June 2009, Piscataway, NJ, USA, IEEE.
- FENG, D., LU, L., YUAN-WU, Y. I., YE LI, G., FENG, G., & LI, S. (2013). Device-to-Device Communications Underlying Cellular Networks. *IEEE transactions on communications*, 61(8), p.p. 3541-3551.
- Fischer, S. and D.-I. U. Rödiger (2000). Transport Layer Security. Open Internet Security, Springer: p.p.215-248.
- Fodor, G., S. Sorrentino and S. Sultana (2014). Network Assisted Device-to-Device Communications: Use Cases, Design Approaches, and Performance Aspects. Smart Device to Smart Device Communication, Springer: p.p.135-163.
- Fragouli, C. and E. Soljanin (2007). Network coding fundamentals, Now Publishers Inc.
- Fragouli, C., J. Widmer and J. L. Boudec (2006). A network coding approach to energy efficient broadcasting: from theory to practice. IEEE Infocom.
- Ganesan, P., R. Venugopalan, P. Peddabachagari, A. Dean, F. Mueller and M. Sichitiu (2003). Analyzing and modeling encryption overhead for sensor network nodes. Proceedings of the 2nd ACM international conference on Wireless sensor networks and applications, ACM.
- Gkantsidis, C. and P. R. Rodriguez (2005). Network coding for large scale content distribution. INFOCOM 2005. 24th Annual Joint Conference of the IEEE Computer and Communications Societies. Proceedings IEEE, IEEE.

- Hagos, D. H. (2012). "The Performance of WiFi Offload in LTE Networks." Master's Thesis, Lulea University of Technology, Sweden.
- Hallahan, R. and J. M. Peha (2010). Policies for public safety use of commercial wireless networks. 38th Telecommunications Policy Research Conference.
- Hallahan, R. and J. M. Peha (2013). "Enabling Public Safety Priority Use of Commercial Wireless Networks."
" <http://repository.cmu.edu/cgi/viewcontent.cgi?article=1145&context=epp>"
- Ho, T., R. Koetter, M. Medard, D. R. Karger and M. Effros (2003). "The benefits of coding over routing in a randomized setting." ISIT 2003, Japan June 29- July 4, 2003
- Júnior, N. d. S. R., M. A. Vieira, L. F. Vieira and O. Gnawali (2014). CodeDrip: Data Dissemination Protocol with Network Coding for Wireless Sensor Networks. Wireless Sensor Networks, Springer: p.p.34-49.
- Katti, S., H. Rahul, W. Hu, D. Katabi, M. Médard and J. Crowcroft (2006). XORs in the air: practical wireless network coding. ACM SIGCOMM Computer Communication Review, ACM.
- Katti, S., H. Rahul, H. Wenjun, D. Katabi, M. Medard and J. Crowcroft (2008). "XORs in the air: practical wireless network coding." IEEE/ACM Transactions on Networking 16(Copyright 2008, The Institution of Engineering and Technology): p.p.497-510.
- Kent, A. D. and L. M. Liebrock (2011). Secure communication via shared knowledge and a salted hash in Ad-hoc environments. Computer Software and Applications Conference Workshops (COMPSACW), 2011 IEEE 35th Annual, IEEE.
- Kwan, R., R. Arnott, R. Trivisonno and M. Kubota (2010). On pre-emption and congestion control for LTE systems. Vehicular Technology Conference Fall (VTC 2010-Fall), 2010 IEEE 72nd, IEEE.
- Le, A., L. Keller, H. Seferoglu, B. Cici, C. Fragouli and A. Markopoulou (2014). "MicroCast: Cooperative Video Streaming using Cellular and D2D Connections."
<http://arxiv.org/pdf/1405.3622v1.pdf>
- Le Faucheur, F. (2005). "Maximum allocation bandwidth constraints model for DiffServ-aware MPLS traffic engineering.", RFC 4125
- Le Faucheur, F. (2005). "Russian dolls bandwidth constraints model for DiffServ-aware MPLS traffic engineering.", RFC 4127

- Lei, L., Y. Zhang, X. Shen, C. Lin and Z. Zhong (2013). "Performance Analysis of Device-to-Device Communications with Dynamic Interference Using Stochastic Petri Nets." *IEEE Transactions on Wireless Communications* p.p. 1-21
- Li, Z. and B. Li (2004). Network coding in undirected networks, CISS04.
<http://pages.cpsc.ucalgary.ca/~zongpeng/publications/ciss04.pdf>
- Lin, X., J. G. Andrews, A. Ghosh and R. Ratasuk (2013). "An Overview on 3GPP Device-to-Device Proximity Services." *Communications Magazine, IEEE* 52, no. 4 (2014): 40-48
- Liu, J., Y. Kawamoto, H. Nishiyama, N. Kato and N. Kadowaki (2014). "Device-to-device communications achieve efficient load balancing in LTE-advanced networks." Wireless Communications, IEEE 21(2): p.p.57-65.
- Lucent, A. (2009) "Introduction to Evolved Packet Core." www3.alcatel-lucent.com
- Lucent, A. (2010) "A HOW-TO GUIDE for LTE in Public Safety ".
https://www.cmu.edu/silicon-valley/dmi/files/howto_guide.pdf
- Lucent, A. (2014). "Mobile Backhaul." from
<http://www.alcatel-lucent.com/solutions/mobile-backhaul>.
- Matsuda, T., T. Noguchi and T. Takine (2011). "Survey of network coding and its applications." IEICE transactions on communications 94(3): p.p.698-717.
- Metz, C. (1999). "AAA protocols: authentication, authorization, and accounting for the Internet." Internet Computing, IEEE 3(6): p.p.75-79.
- Mo, L. M. L., F. Yuan and J. Yang (2015). "Mobile Backhaul Solutions." ZTECOMMUNICATIONS.
- Molnar, K. and M. Vlcek (2009). Evaluation of bandwidth constraint models for MPLS networks, Electronics. ANNUAL JOURNAL OF ELECTRONICS, 2009, ISSN 1313-1842
- Monfreid, C. D. (2009) "The LTE Network Architecture: A comprehensive tutorial." <http://www.alcatel-lucent.com/>.
- Mumtaz, S., H. Lundqvist, K. M. S. Huq, J. Rodriguez and A. Radwan (2014). "Smart Direct-LTE communication: An energy saving perspective." Ad Hoc Networks 13: p.p.296-311.
- Nagpal, V., S. Choudhury and K. Doppler (2013). OFFLOADING TRAFFIC TO DEVICE-TO-DEVICE COMMUNICATIONS, US Patent 20,130,073,671.

- Nurcahyaning Sih, S., R. Munadi, S. N. Hertiana and T. Hasan (2014). "Transport Solution Based on Layer 3 MPLS-Virtual Private Network to Support IP Connectivity in Long Term Evolution Mobile Back-Haul." Advanced Science Letters 20(2): p.p.386-390.
- Peng, Y., Q. Song, Y. Yu and F. Wang (2014). "Fault-tolerant routing mechanism based on network coding in wireless mesh networks." Journal of Network and Computer Applications 37:p.p.259-272.
- Perillo, M. (2007) "Network Coding Overview."
<http://www.ece.rochester.edu/projects/wcng/meetings/PerilloNetworkCoding.pdf>
- Prasad, D. and M. Giri (2014). "Efficient Lightweight Hybrid Cryptography Solution to Secure Mobile Ad hoc Networks." IJRCCT 3(3): p.p.325-332.
- Pyattaev, A., K. Johnsson, S. Andreev and Y. Koucheryavy (2013). 3GPP LTE Traffic Offloading onto WiFi Direct. Proc. of the IEEE WCNC.
- Qian, M., Y. Huang, J. Shi, Y. Yuan, L. Tian and E. Dutkiewicz (2009). A novel radio admission control scheme for multiclass services in LTE systems. Global Telecommunications Conference, 2009. GLOBECOM 2009. IEEE, IEEE.
- Raghothaman, B., E. Deng, R. Pragada, G. Sternberg, T. Deng and K. Vanganuru (2013). Architecture and protocols for LTE-based device to device communication. Computing, Networking and Communications (ICNC), 2013 International Conference on, IEEE.
- Rivest, R. (1992). "The MD5 Message Digest Algorithm." RFC1321
- Sen, J., S. Koilakonda and A. Ukil (2011). A mechanism for detection of cooperative black hole attack in mobile ad hoc networks. Intelligent Systems, Modelling and Simulation (ISMS), 2011 Second International Conference on, IEEE.
- Shah, N. and S. Valiveti (2012). "Intrusion Detection Systems for the Availability Attacks in Ad-Hoc Networks." International Journal of Electronics and Computer Science Engineering (IJECS), ISSN: 2277-1956) 1(03): p.p.1850-1857.
- Shah, N. B., K. Rashmi and K. Ramchandran (2013). Secure network coding for distributed secret sharing with low communication cost. Information Theory Proceedings (ISIT), 2013 IEEE International Symposium on, IEEE.
- Shajaiah, H., A. Abdel-Hadi and C. Clancy (2014). Spectrum sharing between public safety and commercial users in 4G-LTE. Computing, Networking and Communications (ICNC), 2014 International Conference on, IEEE.

- Simic, M. B. (2012). Feasibility of long term evolution (lte) as technology for public safety. Telecommunications Forum (TELFOR), 2012 20th, IEEE.
- Stanze, O. and A. Weber (2013). "Heterogeneous Networks With LTE-Advanced Technologies." Bell Labs Technical Journal 18(1):p.p. 41-58.
- Tang, Z. (2013). "On link encryption against wiretapping attack in network coding." Networking Science: p.p.1-10.
- Tata, C. (2009). Algorithme de courtoisie: optimisation de la performance dans les réseaux WIMAX fixes, École de technologie supérieure.
- Tata, C. and M. Kadoch (2013). CAM: Courteous bandwidth constraints allocation model. Telecommunications (ICT), 2013 20th International Conference on, IEEE.
- C. Tata and M. Kadoch, "RBC: Reliable Butterfly Network Construction Algorithm for Network Coding in Wireless Mesh Network," , 13th WSEAS International Conference on APPLIED INFORMATICS and COMMUNICATIONS (AIC '13), august 2013
- Tata, C. and M. Kadoch (2014). Multipath Routing Algorithm for Device-to-Device Communications for Public Safety over LTE Heterogeneous Networks. The 1st International Conference on Information and Communication Technologies for Disaster Management ICT-DM'2014. Algiers, Algeria.
- Tata, C. and M. Kadoch (2014). Secure Network Coding based Data Splitting for Public Safety D2D communications over LTE Heterogeneous Networks. CIT 14, Tenerife, Spain.
- Tung, L.-C., Y. Lu and M. Gerla (2013). Priority-Based Congestion Control Algorithm for Cross-Traffic Assistance on LTE Networks. Vehicular Technology Conference (VTC Fall), 2013 IEEE 78th, IEEE.
- Venkatesan, G. K. and K. Kulkarni (2008). Wireless backhaul for LTE-requirements, challenges and options. Advanced Networks and Telecommunication Systems, 2008. ANTS'08. 2nd International Symposium on, IEEE.
- Wang, D., Q. Zhang and J. Liu (2008). "Partial network coding: Concept, performance, and application for continuous data collection in sensor networks." ACM Transactions on Sensor Networks (TOSN) 4(3): 14.
- Wang, D., Q. Zhang and J. Liu (2008). "Partial network coding: Concept, performance, and application for continuous data collection in sensor networks." ACM Trans. Sen. Netw. 4(3): p.p.1-22.

- Wang, P., W. Wei and L. Zhuoming (2013). System performance of LTE-advanced network with D2D multi-hop communication. Consumer Electronics, Communications and Networks (CECNet), 2013 3rd International Conference on, IEEE.
- Wang, T.-E., D. Sicker and K. Baker (2013). "Liability and Public Safety Broadband Networks." Available at SSRN: <http://ssrn.com/abstract=2239373>
- Wu, G., Q. C. Li, R. Q. Hu and Y. Qian "Overview of Heterogeneous Networks." Heterogeneous Cellular Networks: p.p.1-25.
- Xu, Feilong, Xian Liu, and Changcheng Huang. "QRP02-3: SAM: A New Bandwidth Constraint Model for Diff-Serv-Aware MPLS Networks." In *Global Telecommunications Conference, 2006. GLOBECOM'06. IEEE*, pp. 1-5. IEEE, 2006.
- Yaacoub, E. and O. Kubbar (2012). Energy-efficient Device-to-Device communications in LTE public safety networks. Globecom Workshops (GC Wkshps), 2012 IEEE, IEEE.
- Yahiya, T. A. (2011). Understanding LTE and its Performance, Springer.
- Yang, F., S. Ling, H. Xu and B. Sun (2012). Network coding-based AOMDV routing in MANET. Information Science and Technology (ICIST), 2012 International Conference on, IEEE.
- Zhao, S., R. Kent and A. Aggarwal (2013). "A key management and secure routing integrated framework for Mobile Ad-hoc Networks." Ad Hoc Networks 11(3): p.p. 1046-1061.
- Zhou, J. (2013). "Efficient and secure routing protocol based on encryption and authentication for wireless sensor networks." International Journal of Distributed Sensor Networks 2013.