



Optimisation de réseaux mobiles hybrides satellite-terrestres

Michael Crosnier

► To cite this version:

Michael Crosnier. Optimisation de réseaux mobiles hybrides satellite-terrestres. Autre. Institut National Polytechnique de Toulouse - INPT, 2013. Français. NNT : 2013INPT0045 . tel-04287194

HAL Id: tel-04287194

<https://theses.hal.science/tel-04287194>

Submitted on 15 Nov 2023

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.



Université
de Toulouse

THÈSE

En vue de l'obtention du
DOCTORAT DE L'UNIVERSITÉ DE TOULOUSE

Délivré par :

Institut National Polytechnique de Toulouse (INP Toulouse)

Discipline ou spécialité :

Réseaux, Télécommunications, Systèmes et Architecture

Présentée et soutenue par :

Michael Crosnier

le : mardi 25 juin 2013

Titre :

Optimisation de réseaux mobiles hybrides satellite-terrestre

Ecole doctorale :

Mathématiques Informatique Télécommunications de Toulouse (MITT)

Unité de recherche :

Institut de Recherche en Informatique de Toulouse - IRIT

Directeur(s) de Thèse :

André-Luc Beylot
Riadh Dhaou

Rapporteurs :

Nadia Boukhatem
Thomas Noel

Membre(s) du jury :

Michel Marot
Riadh Dhaou
Fabrice Planchou

Remerciements

Tout d’abord, je voudrais remercier mon directeur de thèse, André-Luc Beylot, qui m’a apporté un soutien incommensurable tout au long de cette thèse et qui a fourni un effort supplémentaire pour l’écriture du manuscrit. Son encadrement, ses opinions techniques et son soutien moral ont été décisifs dans la réalisation et la finalisation de cette thèse, et cela, du début à la fin, malgré une rédaction l’ayant mis à rude épreuve.

Ensuite, je tiens à exprimer ma gratitude à mon co-encadrant Riadh Dhaou pour sa gentillesse et son aide précieuse tout au long de cette thèse. Excepté lors de nombreux égarements géographiques au cours d’une conférence à Budapest, il fut toujours présent pour me fournir supports et conseils qui se révélèrent très utiles.

Mes remerciements s’adresse également à Fabrice Planchou, mon encadrant industriel pour m’avoir toujours soutenu et avoir cru en mon travail. Ces remerciements en appelleront sûrement d’autres à l’avenir.

Le service ATB3 d’Astrium a bien sur droit à un grand merci et en particulier son responsable Bernard Laurent pour m’avoir accepté au sein de son département. Je décomposerais ce service en groupe : les thésards-docteurs, Vincent, Oriol, Sylvain et bientôt JT, avec lesquels nous nous sommes soutenus mutuellement pour passer les épreuves, les experts, Eric, Jérôme et Philippe, qui ont toujours un avis aimable et pertinent, la bande à MONET, Mélanie et Valentin, qui m’a sauvé de la solitude et les inclassables mais non moins importants, Jean Christophe, Ababacar, Nico, Émilie, Thierry, Xavier, Patricia, Laura, Cyrille, Bernard, Cédric, Guillaume et Laurent. Dans cette catégorie, je terminerais par un grand merci à mes différents collègues de bureau qui se sont relayés puis presque tous échappés, mais sans qui cette thèse aurait été affadie, Veronica, Coline, Pierre et Jean Thomas.

Je suis heureux d’ajouter à ces félicités l’équipe de l’IRIT, en commençant par les joyeux lurons que furent mes compagnons de bureau, Manu, André-Luc puis Guillaume, les pétillants et pleins d’humour thésards et docteurs Clément, Renaud, Patou, Fabrice, JB et Nesrine, sans oublier les maîtres de conférence, le petit Julien, Carlos et Jérôme et enfin les savoureux cochons de lait.

Nous en arrivons aux chaleureux remerciements destinés à l’association des logeurs de Mika qui s’est étoffée au cours de cette thèse :

- Bien sûr, Clément, qui m’a supporté pendant 3 ans, malgré. Ce fût un honneur et un plaisir mon cher ami.
- Béa et JM pour leur service 4 étoiles
- Audrey et Mathieu, pour la retraite paisible en ermite dans le Lot-et-Garonne qui a sauvé mon manuscrit. Bonjour au petit Ravitalamesh.
- Laurie et Basile qui se sont retrouvés avec un gros Mika sur leur canapé
- Ségo et Reno, avec les petits marchés.
- Et mes nouveaux colocs (Vick, Brice, Nat, Manon, Julien)

Un très grand et sincère merci aux Gros, dont le soutien, l’humour et la gentillesse ont éclairé ces années (Béa, Anna, Ségo, P’tite Clem, Delph, Tiffany, Cindy, Angela, Reno, Clément, Basile, Déo,

Arnaud, Mass Matt, Nicoletta, mon petit Nico, Canard, PH Ferdi, Pépité, Cyril, Alex, Got, Pierrick, Pasca) ainsi qu' à ceux éparpillés en France et dans le monde (Alexia, Julie, Lise).

Je ne saurais finir sans remercier affectueusement ma famille : mes parents, mon frère, ma petite sœur, mes oncles et tantes, cousins, cousines, ainsi que mes grands-parents ; pour leur amour et leur humour et en espérant pouvoir réintégrer la bande à Raymond. Ma dernière pensée est destinée à Albert.

Merci à tous

Résumé

Le monde des communications par satellite est dominé par les systèmes de diffusion de la télévision. Cependant, des satellites de communication offrent aussi des services de téléphonie et de données. Ils sont regroupés dans les familles des systèmes fixes et mobiles et ciblent des marchés de niche. Dans cette thèse, nous avons la volonté d'étendre les scénarios d'utilisation de ces systèmes. Notre vision nous dicte que leur développement est lié à l'utilisation de réseaux hybrides mobiles satellite-terrestre. En effet, une utilisation complémentaire des deux segments permet de s'affranchir d'une concurrence trop féroce des réseaux de télécommunications terrestres. Pour cela, nous optons pour deux applications qui nous paraissent prometteuses : un réseau mobile LTE (*Long Term Evolution*) avec des stations de base qui possèdent un backhaul satellite et un réseau MANET (*Mobile Ad-hoc Network*) qui s'interconnecte à des réseaux extérieurs grâce à des liaisons satellite.

Nous soulevons l'un des problèmes les plus contraignants du réseau mobile LTE avec des backhails satellite : la gestion de la mobilité. L'analyse du standard nous a conduits à conclure quant à la nécessité d'optimiser les procédures du *handover*. Ceux qui nécessitent des modifications surviennent entre des stations de base qui n'utilisent pas le même backhaul satellite et entre une station de base avec un backhaul satellite vers une avec un backhaul terrestre. Deux points nous ont semblé importants : la phase de préparation et le mécanisme qui permet d'éviter les pertes. Nous proposons donc une nouvelle phase de préparation qui prend en compte le retard induit par la liaison satellite ainsi qu'une phase de préparation à double décision combinée avec une préparation de multiples stations de base. Nous tentons ainsi de maximiser les chances de réaliser un *handover* avec succès. Puis, nous avons imaginé un mécanisme qui permet à la fois d'éviter les pertes lors de l'exécution du *handover* et de sauvegarder les précieuses ressources du satellite.

Les réseaux MANET associés à des liaisons satellite offrent des caractéristiques très intéressantes pour les communications d'urgence, telles que l'indépendance vis-à-vis des infrastructures terrestres susceptibles d'être endommagées par des catastrophes ainsi qu'un déploiement rapide pour une intervention sur le théâtre des opérations. Nous avons souhaité améliorer l'un des points cruciaux dans le cadre d'une hybridation : la sélection de la passerelle satellite. Nous avons donc développé un mécanisme qui prend en compte la charge sur les passerelles satellite ainsi que le phénomène d'oscillation de passerelle souvent négligé dans la littérature.

Ces optimisations ont pour but de favoriser le développement de réseaux hybrides satellite terrestres en améliorant les performances de ces réseaux. L'avenir nous semble prometteur quant à l'utilisation de la technologie LTE avec un backhaul satellite pour lequel nous avons proposé une nouvelle gestion de la mobilité qui est primordiale pour son développement.

Abstract

Satellite communications are leaded by television broadcasting. Yet, fixed and mobile satellite systems provide voice services as well as IP-based applications. In this thesis, we try to develop user scenarios in order to extend their targeted market. Our vision to reach this objective consists to use hybrid satellite and terrestrial mobile networks. This network design avoids a competition between both segments in which a satellite success is difficult to imagine. Furthermore, hybrid networks may draw benefits from both segments. Two promising scenarios have been selected. The first one consists in a mobile LTE network (Long Term Evolution) with base stations backhauled by satellite links whereas the second scenario is composed of a Mobile Ad-hoc Network (MANET) connected to external networks thanks to satellite systems.

One of the main problems in the hybrid LTE scenario is caused by mobility procedures. As a consequence of the standard analysis, we have decided to optimize the mobility management in two cases: a handover between two base stations for which the backhaul is provided by two different satellite terminals and a handover from a base station with a satellite backhaul to one with a terrestrial backhaul. Two procedures have drawn our attention: the preparation phase and the loss avoidance mechanism during the execution phase. First of all, we design a new procedure for the preparation which takes into account the delay induced by the satellite link. This new phase is based on a twofold decision preparation associated with multiple preparations. This solution leads to an increase of handover success. The second optimization aims to avoid losses during the execution phase and, at the same time, save satellite resources.

MANET and satellite hybridization leads to very interesting characteristics for public safety communications. Indeed, these networks are independent of terrestrial infrastructures that can be impaired or destroyed. Furthermore, they can be rapidly deployed in the theater of operation. Gateway selection is a crucial problem linked to hybrid MANET. Therefore, we have focused our work on this mechanism taking into account the measured load on the satellite links as well as an often-neglected phenomenon, the gateway flapping.

These optimizations tend to promote hybrid satellite and terrestrial networks improving their performance. A promising future is foreseen for the hybrid LTE technology and we have proposed a solution to a problem that may be very detrimental to its deployment.

Sommaire

Résumé	iv
Abstract	v
Sommaire	vi
Liste des Figures	xi
Liste des Tables.....	xiv
Liste des Abréviations.....	xv
Introduction.....	1
Chapitre 1 : Les réseaux mobiles hybrides satellite-terrestres.....	3
1.1 Les réseaux terrestres mobiles.....	3
1.1.1 Les réseaux cellulaires.....	3
1.1.2 Les réseaux mobiles sans infrastructure	4
1.2 Les systèmes satellite	5
1.2.1 Les systèmes satellite de diffusion.....	5
1.2.2 Les systèmes satellite fixes.....	6
1.2.3 Les systèmes satellite mobiles	7
1.2.4 Les systèmes satellite portables.....	8
1.3 La classification des réseaux hybrides.....	9
1.3.1 Description	9
1.3.2 Les réseaux intégrés	10
1.3.2.1 Les bénéfices	10
1.3.2.2 Les réseaux cellulaires intégrés.....	11
1.3.2.3 Les réseaux sans infrastructures	13
1.3.3 Les réseaux instantanés.....	14
1.3.3.1 Les bénéfices	14
1.3.3.2 Réseaux cellulaires	14
1.3.3.3 Les réseaux sans infrastructure.....	16
1.4 Axes d'optimisations	18
Chapitre 2 : Réseaux hybrides LTE-satellite : Description et analyse du standard.....	23
2.1 Scenarios étudiés.....	23
2.2 Le standard LTE/SAE.....	25
2.2.1 Architecture standard de LTE	25
2.2.2 La pile protocolaire de LTE	26

2.2.3	La gestion des flux	27
2.2.4	La gestion de la mobilité.....	28
2.2.4.1	Idle Mode	28
2.2.4.2	Active Mode	29
2.2.4.3	S1 handover avec et sans relocalisation de la MME et de la SSGW.....	29
2.3	Les solutions de réseaux LTE hybrides avec un backhaul satellite.....	35
2.3.1	Réseau LTE hybride traditionnel.....	35
2.3.2	Architecture avec un autre réseau d'accès 3GPP.....	36
2.3.3	Architecture avec un Home eNB	37
2.3.4	Architecture avec un réseau extérieur.....	38
2.3.5	Les choix d'architectures.....	38
2.4	Analyses du standard LTE.....	39
2.4.1	Ralentissement de la phase de préparation du <i>handover</i>	39
2.4.2	Le mécanisme de <i>Forward</i> et d'ordonnancement inefficace.....	41
Chapitre 3 : Réseaux hybrides LTE-satellite : Optimisation de la phase de préparation.....		44
3.1	Etat de l'art.....	44
3.1.1	Techniques d'optimisation de la gestion de la mobilité.....	45
3.1.2	Préparation multiple	46
3.1.3	<i>Forward handover</i>	47
3.2	Les optimisations.....	48
3.2.1	Optimisation de paramètres	48
3.2.1.1	Différenciation des handovers	48
3.2.1.2	Choix des paramètres de mesure.....	49
3.2.1.3	Limitations.....	50
3.2.2	Préparation anticipée à double décision.....	50
3.2.2.1	Principe.....	50
3.2.2.2	Définition des événements	52
3.2.2.3	Concaténation des multiples préparations.	52
3.3	Conclusion	53
Chapitre 4 : Réseaux hybrides LTE-satellite : optimisation du mécanisme de <i>Forward</i>		55
4.1	Différenciation des <i>handovers</i>	55
4.2	Contraintes	56
4.3	Etat de l'art.....	57
4.3.1	Les mécanismes de <i>Forward</i>	57

4.3.1.1	Les mécanismes pour éviter les pertes pendant les handovers intra-LTE	57
4.3.1.2	Les mécanismes pour éviter les pertes pendant les handovers inter-RAT	59
4.3.1.3	Les mécanismes pour éviter les pertes pendant un handover avec un eNB relais... ..	60
4.3.2	Les mécanismes pour éviter les duplications dans le 3GPP	61
4.3.3	Les mécanismes pour réaliser l'ordonnancement après le <i>Path Switch</i> dans le 3GPP. ..	62
4.3.4	Les mécanismes intéressants	64
4.4	Optimisations	64
4.4.1	Mécanisme de <i>Forward bi-cast</i> avec <i>buffer</i> pour les <i>handovers</i> inter-satellite-terrestre	65
4.4.1.1	Principe	65
4.4.1.2	La mise en place du <i>buffer</i>	65
4.4.1.3	La gestion du <i>buffer</i> avant le déclenchement du <i>Forward</i>	68
4.4.1.4	Le déclenchement du <i>Forward</i> avec <i>buffer</i>	69
4.4.1.5	La gestion du <i>buffer</i> et le <i>Path Switch</i>	71
4.4.1.6	Récapitulatif de la procédure optimisée	73
4.4.2	Optimisation du <i>handover</i> intra-satellite par multicast	75
4.4.2.1	Principe.....	75
4.4.2.2	Mise en place du mécanisme de <i>Forward</i>	76
4.4.2.3	La gestion du <i>Forward</i> pendant avant la fin de la phase d'exécution.....	78
4.4.2.4	La duplication de paquets	78
4.4.2.5	Récapitulatif de la procédure optimisée	78
4.4.3	Minimisation de l'impact des solutions sur la signalisation	80
Chapitre 5 : Réseaux hybrides LTE-satellite : Étude des optimisations des mécanismes de <i>Forward</i> .		81
5.1	Les simulateurs	82
5.2	Implantation	83
5.2.1	La simulation	83
5.2.1.1	Le plan utilisateur	83
5.2.1.2	Le plan de contrôle	84
5.2.2	Problème d'implantation du protocole TCP dans NS3.....	85
5.2.2.1	Extensions TCP.....	85
5.2.2.2	Implantation Linux.....	85
5.2.3	L'émulation	87
5.3	Émulation du mécanisme de <i>Forward</i> pour le <i>handover</i> inter-satellite-terrestre.....	88
5.4	Émulation du mécanisme de <i>Forward</i> pour le <i>handover</i> intra-satellite	91

5.5	Conclusion	95
Chapitre 6 : MONET : Présentation de MONET et du mécanisme de sélection de passerelle.		96
6.1	Le contexte et le projet MONET	96
6.1.1	Les scénarios d'utilisation.....	96
6.1.2	Architecture.....	97
6.2	La sélection de passerelle.....	99
6.2.1	L'adressage IP et la mobilité du réseau	99
6.2.2	Les méthodes pour découvrir les passerelles.	100
6.2.3	Les algorithmes de sélection de passerelle	101
6.2.4	Le partage de charge	102
6.2.4.1	Phase de mesure et les métriques	103
6.2.4.2	La décision du partage de charge	104
6.2.4.3	L'exécution du partage de charge	104
6.2.5	Les problèmes récurrents des mécanismes de partage de charge	108
6.2.5.2	Les problèmes de simulation.....	109
6.2.6	Descriptions des principaux mécanismes partage de charge.	110
6.3	Les spécificités de MONET pour la sélection de passerelle.....	114
6.4	Conclusion	115
Chapitre 7 : MONET : Optimisation du partage de charge parmi les passerelles satellite d'un MANET		116
7.1	L'algorithme d'optimisation	116
7.1.1	Principe.....	116
7.1.2	La phase d'exécution	117
7.1.3	La phase de mesure.....	119
7.1.4	La phase de décision.....	121
7.1.5	Mécanisme contre l'oscillation de passerelle	121
7.1.5.1	Limitation de l'étendue de la zone progressive	123
7.1.5.2	Blocage du mécanisme	124
7.1.6	La sous-charge	125
7.2	Évaluation par simulation.....	125
7.2.1	Modèle par défaut.....	125
7.2.1.1	Le segment MANET	125
7.2.1.2	Le segment satellite.....	126
7.2.2	Les scénarios de trafic	126

7.2.3	Première phase de simulation.....	126
7.2.4	Deuxième phase de simulation	130
7.2.1	Conclusion	134
	Conclusion	135
	Ouverture	137

Liste des Figures

Figure 1 Evolution des standards des réseaux cellulaires	4
Figure 2 Architecture des systèmes satellite.....	5
Figure 3 Les évolutions des standards GMR-1	8
Figure 4 La classification des réseaux hybrides mobiles	9
Figure 5 Architecture des réseaux intégrés cellulaires	12
Figure 6 Architecture des réseaux intégrés sans infrastructure	13
Figure 7 Architecture d'un réseau instantané cellulaire avec un backhaul satellite.....	15
Figure 8 Architecture d'un réseau instantané cellulaire possédant une passerelle satellite	16
Figure 9 Réseau MANET instantané avec une passerelle satellite.....	17
Figure 10 Réseau MANET instantané avec une interface satellite interne	18
Figure 11 Récapitulatif de la classification	20
Figure 12 Nouveaux scénarios d'architecture LTE avec un backhaul satellite.....	24
Figure 13 Architecture standard LTE	26
Figure 14 Piles protocolaires du plan utilisateur.....	26
Figure 15 Piles protocolaires du plan de contrôle.....	27
Figure 16 La phase de préparation du <i>handover</i>	31
Figure 17 La phase d'exécution du <i>handover</i>	32
Figure 18 La phase de terminaison du <i>handover</i>	33
Figure 19 Le chemin source des données	34
Figure 20 Le chemin de <i>Forward</i>	34
Figure 21 Le chemin des données après le <i>Path Switch</i>	34
Figure 22 Fin du <i>Forward</i>	34
Figure 23 Architecture hybride LTE traditionnelle	36
Figure 24 Architecture avec un autre réseau 3GPP	37
Figure 25 Architecture avec un Home eNB	38
Figure 26 <i>Too-Late For Measure</i> RLF (type 1).	40
Figure 27 <i>Too-Late For Command</i> RLF (type 2).	40
Figure 28 <i>Too-Early</i> RLF (type 3).....	40
Figure 29 Différenciation entre le tunnel source et cible.	41
Figure 30 Transfert du contexte PDCP dans le cadre d'un <i>handover</i> intra-satellite	42
Figure 31 Préparation multiple	47
Figure 32 <i>Forward Handover</i>	48
Figure 33 Préparation anticipée à double décision.....	51
Figure 34 Événement pendant une préparation anticipée sans préparation multiple	53
Figure 35 <i>Forward</i> par <i>bi-cast</i>	58
Figure 36 Réalisation conjointe du <i>handover</i> et du <i>Path Switch</i>	59
Figure 37 Backward <i>handover</i>	59
Figure 38 <i>Forward</i> par <i>buffer</i>	60
Figure 39 Architecture des eNB relais (RN)	60
Figure 40 <i>Forward</i> pour les relais.....	61
Figure 41 Mécanisme contre les duplications.....	62
Figure 42 Réordonnancement dans l'eNB cible	63

Figure 43 <i>Forward bi-cast</i> avec <i>buffer</i> avant la phase de terminaison	65
Figure 44 <i>Forward bi-cast</i> avec <i>buffer</i> après la phase de terminaison	65
Figure 45 Mise en place du <i>buffer</i> de <i>Forward</i>	66
Figure 46 Activation du <i>buffer</i> de <i>Forward</i> dans la passerelle satellite	68
Figure 47 <i>Buffer</i> de <i>Forward</i> dans la passerelle satellite après le déclenchement du <i>Forward</i>	68
Figure 48 Déclenchement du <i>Forward</i> par le message <i>Forward Activation</i>	69
Figure 49 Mécanismes contre la duplication de paquets	71
Figure 50 Le mécanisme de <i>Forward</i> avant le <i>Path Switch</i>	72
Figure 51 Le mécanisme de <i>Forward</i> après le <i>Path Switch</i>	73
Figure 52 Procédure complète de l'optimisation du <i>handover</i> inter-segment-satellite.....	74
Figure 53 Principe du <i>Forward multicast</i> pendant un <i>handover</i> intra-satellite	75
Figure 54 Mise en place simple du mécanisme de <i>Forward</i>	76
Figure 55 Implications de la passerelle et des terminaux satellite dans le mécanisme de <i>Forward</i>	77
Figure 56 Procédure complète de l'optimisation du <i>handover</i> intra-satellite	79
Figure 57 <i>Private extension</i> du protocole GTP-C.....	80
Figure 58 Implantation des couches protocolaires sur le plan utilisateur	83
Figure 59 Implantation des couches protocolaires sur le plan utilisateur	84
Figure 60 Fenêtre de congestion de TCP.....	86
Figure 61 Architecture de simulation avec la pile protocolaire TCP/IP linux.....	87
Figure 62 Principe du mécanisme de <i>Forward</i> par <i>buffer</i>	88
Figure 63 Fenêtre de congestion lors d'un <i>handover</i> sans <i>Forward</i> déclenché lors de la phase normale	89
Figure 64 Fenêtre de congestion lors d'un <i>handover</i> avec <i>Forward</i> déclenché lors de la phase normale	89
Figure 65 Fenêtre de congestion lors d'un <i>handover</i> avec <i>Forward</i> déclenché lors de la phase de récupération.....	90
Figure 66 Principe du mécanisme de <i>Forward</i> du <i>handover</i> intra-satellite	91
Figure 67 Fenêtre de congestion lors d'un <i>handover</i> intra-satellite sans mécanisme de <i>Forward</i> et déclenché après 15s	92
Figure 68 Fenêtre de congestion lors d'un <i>handover</i> intra-satellite sans mécanisme de <i>Forward</i> et déclenché après 20s	92
Figure 69 Fenêtre de congestion lors d'un <i>handover</i> intra-satellite sans mécanisme de <i>Forward</i> et déclenché après 30s	92
Figure 70 Fenêtre de congestion lors d'un <i>handover</i> intra-satellite avec mécanisme de <i>Forward</i> et des pertes et déclenché après 15s.....	93
Figure 71 Fenêtre de congestion lors d'un <i>handover</i> intra-satellite avec mécanisme de <i>Forward</i> , sans gestion des pertes et déclenché après 20s	93
Figure 72 Fenêtre de congestion lors d'un <i>handover</i> intra-satellite avec mécanisme de <i>Forward</i> , sans gestion des pertes et déclenché après 30s	94
Figure 73 Fenêtre de congestion lors d'un <i>handover</i> intra-satellite avec mécanisme de <i>Forward</i> et sans perte	94
Figure 74 Scénario d'utilisation	97
Figure 75 Architecture du réseau MONET	98
Figure 76 Problème de gestion de la mobilité dans les MANET à plusieurs passerelles	99
Figure 77 La gestion de la mobilité dans MONET.....	100

Figure 78 Incohérence de la sélection de passerelle	105
Figure 79 Zone de partage fixe à un saut	107
Figure 80 Zone de partage progressive à un saut	107
Figure 81 Exemple d'oscillation de passerelle.....	108
Figure 82 Séparation d'un MANET	115
Figure 83 Principe du partage de charge proposé	117
Figure 84 Gestion des domaines par la correction.....	118
Figure 85 Exemple de rectification.....	119
Figure 86 Récapitulatif de la phase de mesure	121
Figure 87 Mécanisme contre l'oscillation.....	121
Figure 88 Exemple de détection du phénomène d'oscillation.....	122
Figure 89 Exemple de message HNA	123
Figure 90 Différence de correction excessive	124
Figure 91 Répartition de l'amélioration des pertes de paquets pour des flux de vidéoconférence ..	127
Figure 92 Répartition de l'amélioration des pertes de paquets pour des flux de voix.....	128
Figure 93 Répartition de l'amélioration des pertes de paquets pour des flux de vidéoconférence en fonction de la mobilité	129
Figure 94 Répartition de l'amélioration des pertes de paquets pour des flux de vidéoconférence en fonction du trafic interne	129

Liste des Tables

Tableau 1 Description des standards satellite fixes	6
Tableau 2 Caractéristiques des systèmes satellite mobiles	8
Tableau 3 les systèmes satellite portables.....	8
Tableau 4 Les caractéristiques des paramètres de mesure	30
Tableau 5 Les différents événements paramétrables dans l'UE.	31
Tableau 6 Description des différents niveaux du critère d'optimisation.....	35
Tableau 7 Description des différents niveaux du critère d'impact sur la signalisation.....	35
Tableau 8 Récapitulatif des architectures potentielles.....	39
Tableau 9 Impact du délai satellite sur la signalisation de la phase de préparation	39
Tableau 10 Description des différentes <i>Radio Link Failure</i>	40
Tableau 11 Conséquences du backhaul satellite sur le mécanisme de <i>Forward</i>	41
Tableau 12 Événement et paramètres des <i>handovers</i>	49
Tableau 13 Paramètres préconisés par le 3GPP dans (3GPP TSG-RAN WG1, March 2009)	50
Tableau 14 Définition des événements de la préparation anticipée à double décision.....	52
Tableau 15 Récapitulatif des différentes procédures pour la phase de préparation	54
Tableau 16 Contraintes sur le mécanisme contre la perte des paquets durant le <i>handover</i>	57
Tableau 17 Mécanisme contre la perte de données pendant le <i>handover</i>	64
Tableau 18 Description des paramètres des équations	67
Tableau 19 Paramètres des liaisons point à point	84
Tableau 20 Paramètres de TCP cubic sur Linux.....	88
Tableau 21 Récapitulatif des algorithmes de sélection de passerelle selon la charge	114
Tableau 22 Paramètres de simulations des applications	126
Tableau 23 Paramètres d'OLSR et du partage de charge.....	127
Tableau 24 Données statistiques des flux de voix.....	129
Tableau 25 Données statistiques des flux de vidéoconférence	130
Tableau 26 Données statistiques des flux de voix en fonction du seuil de sous-charge	130
Tableau 27 Données statistiques des flux de vidéoconférence en fonction du seuil de sous-charge	131
Tableau 28 Paramètres de l'algorithme de (Pham V. , Larsen, Kure, & Engelstad, November 2010)	133
Tableau 29 Données statistiques des flux de voix avec l'influence de l'oscillation	133
Tableau 30 Données statistiques des flux de vidéoconférence avec l'influence de l'oscillation.....	133
Tableau 31 Données statistiques sur la durée des flux de TCP avec l'influence de l'oscillation.....	133

Liste des Abréviations

4G	Quatrième Génération
3GPP	Third Generation Partnership Project
AceS	Asian Cellular System
AMBR	Aggregate Maximum Bit Rate
AMPS	Advanced Mobile Phone System
AODV	Ad hoc On Demand Distance Vector
AS	Access Stratum
BGAN	Broadband Global Access Network
BSS	Broadcast Satellite System
CDMA	Code Division Multiple Access
CDMA EV-DO	CDMA Evolution-Data Optimized
CIO	Cell Individual Offset
D-AMPS	Digital Advanced Mobile Phone System
DeNB	Donor evolved Node B
DSCP	Differentiated Services Code Point
DVB-RCS	Digital Video Broadcasting – Return Channel Satellite
DVB-S	Digital Video Broadcasting – Satellite
DVB-SH	Digital Video Broadcasting –Satellite Handheld
E-UTRAN	Evolved Universal Terrestrial Radio Access Network
EDGE	Enhanced Data rates for GSM Evolution
eNB	evolved Node B
EPC	Evolved Packet Core
EPS	Evolved Packet System
ESA	European Space Agency
ETSI	European Telecommunication Standards Institute
ETT	Expected Transmission Time
ETX	Expected Transmission Count
FA	Foreign Agent
FDMA	Frequency Division Multiple Access
FMIP	Fast Mobile Internet Protocol
FP7	The 7 th Framework Program of the European Commission
FSS	Fixed Satellite System
IETF	Internet Engineering Task Force
IMT	International Mobile Telecommunications
IMS	IP Multimedia Subsystem
IP	Internet Protocol
IPoS	IP over Satellite

IS-54/136	Interim Standard -54/136 (voir D-AMPS)
IS-95	Interim Standard 95
ISDB-S	Integrated Services Digital Broadcasting-Satellite
ITU-R	International Telecommunication Union- Radiocommunication sector
ITU-T	International Telecommunication Union- Telecommunication Standardization sector
JTACS	Japanese Total Access Communication System
GBR	Guaranteed Bit-Rate
GMR	Geo Mobile Radio interface
GMSS	Geostationary Mobile Satellite System
GPRS	General Packet Radio Service
GSM	Global System for Mobile Communications (historiquement, Groupe Spécial Mobile)
GTP-U	GPRS Tunnelling Protocol for User plane
GTP-C	GPRS Tunnelling Protocol for Control plane
GW-SOL	Gateway Solicitation
GW-ADV	Gateway Advertisement
HeNB	Home – Evolved Node B
HeNB-GW	Home – Evolved Node B Gateway
HEO	Highly Elliptical orbit
HMIP	Hierarchical Mobile Internet Protocol
HNA	Host Network Association
HSPA	High Speed Packet Access
HSS	Home Subscriber Server
LENA	LTE-EPC Network Simulator
LEO	Low Earth Orbit
LTE	Long Term Evolution
MAC	Medium Access Control
MANET	Mobile Ad-hoc NETwork
MIP	Mobile Internet Protocol
MME	Mobility Management Entity
MONET	Mechanisms for Optimization of hybrid ad-hoc networks and satellite NETworks
MPEG	Moving Experts Group
MPR	Multi-Point Relay
MRO	Mobility Robustness Optimization
MSS	Mobile Satellite System
NAS	Non Access Stratum
NS3	Network Simulator 3
NMT	Nordic Mobile Telephone
OBP	On-Board Processing
OFDM	Orthogonal Frequency Division Multiplexing
OLSR	Optimized Link State Routing

OPEX	OPerating Expense
PDCP	Packet Data Convergence Protocol
PGW	Packet Data Network Gateway
PRNet	Packet Radio Network
QoS	Quality of Service
RAT	Radio Access Technology
RFC	Request For Comments
RLC	Radio Link Control
RLF	Radio Link Failure
RN	Relay Node
RSRP	Reference Signal Received Power
RSRQ	Reference Signal Received Quality
RRC	Radio Resource Control
S-DOCSIS	Satellite- Data Over Cable Service interface Specification
S-MIM	S-band Mobile Interactive Multimedia
S-UMTS	Satellite-Universal Mobile Telecommunications System
S1-AP	S1-Application Protocol
SAE	System Architecture Evolution
SATIN	Satellite-UMTS IP-based Network
SCPC	Single Channel Per Carrier
SCTP	Stream Control Transmission Protocol)
SGW	Serving Gateway
SINUS	Satellite Integration into Networks for UMTS Services
SMS	Short Message Service
SON	Self-Organized, Self-Optimized Networks
STICS	Satellite/Terrestrial Integrated mobile Communication System
SWIPE	Space Wireless sensor networks for Planetary Exploration
TACS	Total Access Communication System
TDMA	Time Division Multiple Access
TEID	Tunnel endpoint identifier
TTT	Time To Trigger
TTL	Time-To-Live
UE	User Equipment
UMTS	Universal Mobile Telecommunications System
V2V	Vehicle to Vehicle
V2I	Vehicle to infrastructure
VANET	Vehicular Ad-hoc NETwork
VSAT	Very Small Aperture Terminal
WCDMA	Wideband Code Division Multiple Access
WiFi	Wireless Fidelity

WiMAX	Worldwide Interoperability for Microwave Access
WSN	Wireless Sensor Network

Introduction

La diffusion de programmes télévisés est le premier service par satellite avec 80% de part de marché (Satellite Industry Association (SIA), Mai 2012). Les réseaux satellite fixe et mobile ne sont relayés qu'à des marchés de niche. Pour étendre leur marché, nous proposons de diversifier les applications et les scénarios d'utilisation grâce à l'hybridation des systèmes terrestres avec des systèmes satellite. Notre objectif est d'optimiser les mécanismes issus du monde terrestre pour des réseaux hybrides en les adaptant aux caractéristiques du segment satellite. Cette voie se veut la moins concurrentielle possible avec les acteurs terrestres en choisissant une hybridation des réseaux satellite et terrestre. L'intérêt de ces réseaux hybrides réside dans la fusion des bénéfices des deux systèmes. En effet, tandis que le segment terrestre conserve ses avantages comme la faible latence et les très hauts débits, le segment satellite autorise une couverture globale et la possibilité d'un déploiement rapide. Ce mode coopératif permet d'améliorer les performances des services et d'élargir les scénarios d'application dans les pays développés. Dans les pays en voie de développement, cette solution prendrait une importance plus prononcée, car elle permettrait de compenser la faible pénétration des réseaux terrestres.

Dans l'optique de trouver de nouveaux marchés porteurs, nous avons choisi de focaliser notre travail sur les réseaux mobiles terrestres. En effet, depuis l'avènement des réseaux mobiles et de la possibilité d'une itinérance globale initiée par la norme mobile GSM, les besoins en mobilité de l'utilisateur n'ont fait que croître. Les nouvelles générations de standards mobiles ont rehaussé les exigences en termes de gestion de mobilité qui doit être la plus transparente possible pour l'utilisateur, jusqu'à des vitesses extrêmement élevées (International Telecommunication Union - Radiocommunication sector (ITU-R), November 2008) (The 3rd Generation Partnership Project (3GPP), September 2011). Les réseaux cellulaires terrestres en sont le fer de lance.

Notre objectif est donc d'optimiser l'hybridation de réseaux satellite-terrestres. Pour cela, nous définissons des scénarios et des contextes qui favorisent cette hybridation et pour chacun d'entre eux, nous proposons des mécanismes permettant de la rendre harmonieuse.

La description des scénarios choisis et les problèmes induits par l'hybridation sont exposés en même temps que l'annonce du plan ci-après. Le premier chapitre consiste en une classification des réseaux hybrides satellite terrestres. Il nous a permis de mettre en exergue les spécificités des différents réseaux hybrides mobiles ainsi que les problèmes que leur pose l'hybridation. Il s'est révélé que deux catégories de réseaux hybrides ressortent de la classification pour lesquels cette hybridation est possible, mais également pour lesquels il est fortement souhaitable d'en améliorer le fonctionnement. Le premier est l'hybridation d'un réseau mobile de quatrième génération LTE. L'hybridation dans ce scénario se situe au niveau des liens de backhaul, lorsqu'ils sont fournis par un segment satellite. Le deuxième scénario qui nous intéresse repose sur une hybridation des réseaux ad-hoc mobile. Il s'avère particulièrement intéressant dans le cadre de communications d'urgence. Dans ce cas-là, le satellite permet d'interconnecter le réseau MANET (Mobile Ad-Hoc Network) déployé sur le théâtre de l'opération avec le quartier général ainsi qu'avec les réseaux extérieurs. Le réseau qui en résulte ouvre la voie à une indépendance totale de ces réseaux vis-à-vis des infrastructures terrestres fixes. Tous deux présentent des propriétés de résistance aux catastrophes et de déploiement rapide qui les destinent à des utilisations pour des communications d'urgence.

Les chapitres 2, 3, 4 et 5 concernent alors le réseau LTE avec un backhaul satellite alors que les chapitres 6 et 7 ont pour sujet les réseaux hybrides MANET pour des scénarios de sécurité civile.

Dans le chapitre 2, après avoir décrit les scénarios potentiels d'utilisation de réseaux LTE avec un backhaul satellite, nous analysons le standard LTE pour en révéler les difficultés soulevées par l'hybridation. La gestion de la mobilité est très affectée par la liaison satellite. Le principal fautif est le délai introduit par celle-ci. En outre, des procédures de *handover* défaillantes deviendraient un frein à l'utilisation de ces scénarios. Nous proposons dans ce chapitre une architecture qui permet de réaliser des optimisations des mécanismes de mobilité tout en induisant le minimum d'impact sur le segment terrestre.

Le chapitre 3 concentre les propositions d'optimisation de la phase en amont du *handover* appelée phase de préparation. Nous nous focalisons sur l'amélioration de cette phase dans le but d'éviter au maximum les échecs de *handovers* dus à une absence de considération du lien satellite.

Le chapitre 4 présente les optimisations de la gestion des paquets de données lors du *handover*. En particulier, nous nous fixons deux objectifs. Le premier est de réduire les pertes de paquets et le temps d'interruption des communications pendant le *handover* pour que la procédure soit la plus transparente possible pour l'utilisateur. Le deuxième but est d'optimiser la gestion des ressources sur le lien et en éviter la surconsommation en raison du coût de celles-ci. Le chapitre 5 contient les simulations des mécanismes proposés dans le chapitre 4.

La présentation des scénarios ainsi que la définition de l'architecture du réseau hybride MANET satellite est effectuée dans le chapitre 6. Nous y décrivons le projet de la commission européenne MONET (*Mechanisms for Optimization of hybrid ad-hoc networks and satellite NETWORKS*) (Oliveira, Sun, Boutry, Gimenez, Pietrabissa, & Juros, 2011) qui est à l'origine de ce travail. La sélection des passerelles satellite pour atteindre l'extérieur est l'un des principaux mécanismes affectés par l'hybridation du réseau MANET. En effet, la sélection de passerelle se doit de prendre en compte à la fois les caractéristiques du lien satellite et ceux du réseau MANET pour optimiser les performances des flux applicatifs. Nous analysons donc les différentes propositions de la littérature et définissons les différentes contraintes inhérentes à l'hybridation du MANET. Le mécanisme que nous proposons est décrit dans le chapitre 7. L'algorithme prend particulièrement en considération la charge sur les liens satellite. En outre, ce chapitre contient les simulations de ce mécanisme pour en vérifier l'efficacité.

Nous sommes persuadés que l'avenir des réseaux satellite fixes et mobiles passe par une hybridation avec les réseaux terrestres. Cette thèse apporte des contributions contre les problèmes induits par cette hybridation pour deux scénarios prometteurs : le backhaul de station de base LTE et un réseau hybride satellite MANET.

Chapitre 1 : Les réseaux mobiles hybrides satellite-terrestres

Les réseaux mobiles hybrides sont dans la majorité des cas utilisés en complément des réseaux terrestres pour en étendre la couverture. Par conséquent, ils reposent, la plupart du temps, sur des architectures qui ont pour origine celles des réseaux terrestres. Nous commencerons donc par retracer brièvement les caractéristiques des réseaux mobiles terrestres puis nous proposerons une classification des réseaux hybrides mobiles sur laquelle nous fonderons notre travail. Pour finir, nous recenserons les différents problèmes liés à ces solutions hybrides et définirons les axes d'optimisations qui contribueront à leur développement et à leur déploiement.

1.1 Les réseaux terrestres mobiles

Les réseaux mobiles terrestres sont divisés en deux grandes familles. La première représente les réseaux mobiles avec infrastructure. Ce sont les réseaux de la téléphonie mobile traditionnelle. Ils possèdent le plus fort poids économique. La deuxième famille est constituée des réseaux mobiles sans infrastructure.

1.1.1 Les réseaux cellulaires.

Les réseaux cellulaires ont vu le jour dans les années 1980 avec des solutions propriétaires et se sont mondialement développés grâce à une standardisation internationale dans les années 1990. La croissance de ce marché se traduit par d'impressionnantes améliorations technologiques dont les évolutions sont classifiées sous la dénomination de générations.

L'évolution des standards des réseaux cellulaires de la première génération à la quatrième génération est relatée dans de nombreux ouvrages (Stüber, 2011) (Sesia, Toufik, & Baker, 2011). Elle est résumée de manière chronologique dans la Figure 1. À partir de la troisième génération, deux groupes de standardisation se forment, le 3GPP et le 3GPP2 (*the Third Generation Partnership Project*, (The 3rd Generation Partnership Project (3GPP)), (The 3rd Generation Partnership Project 2 (3GPP2))). Ils sont, par la même occasion, désignés comme responsable des évolutions des standards de deuxième génération. Les standards du 3GPP sont les plus utilisés actuellement avec, en particulier, la norme GSM (*Global System for Mobile Communications*) qui connaît le plus vaste déploiement de la famille des réseaux mobiles cellulaires ; le nombre d'abonnements a atteint 4,6 milliards en 2012 (Ericsson, 2012). En outre, son standard de quatrième génération, LTE-Advanced précédé de LTE (*Long Term Evolution*), est en phase de devenir l'unique représentant des standards 4G. Les différentes générations ont accordé une place de plus en plus prépondérante aux applications de données qui reposent sur la pile protocolaire TCP/IP. Ainsi, la 4^{ème} génération est caractérisée par le passage de l'ensemble du réseau au protocole IP, à la suppression du raccordement au réseau téléphonique commuté et à une augmentation sensible de la capacité.

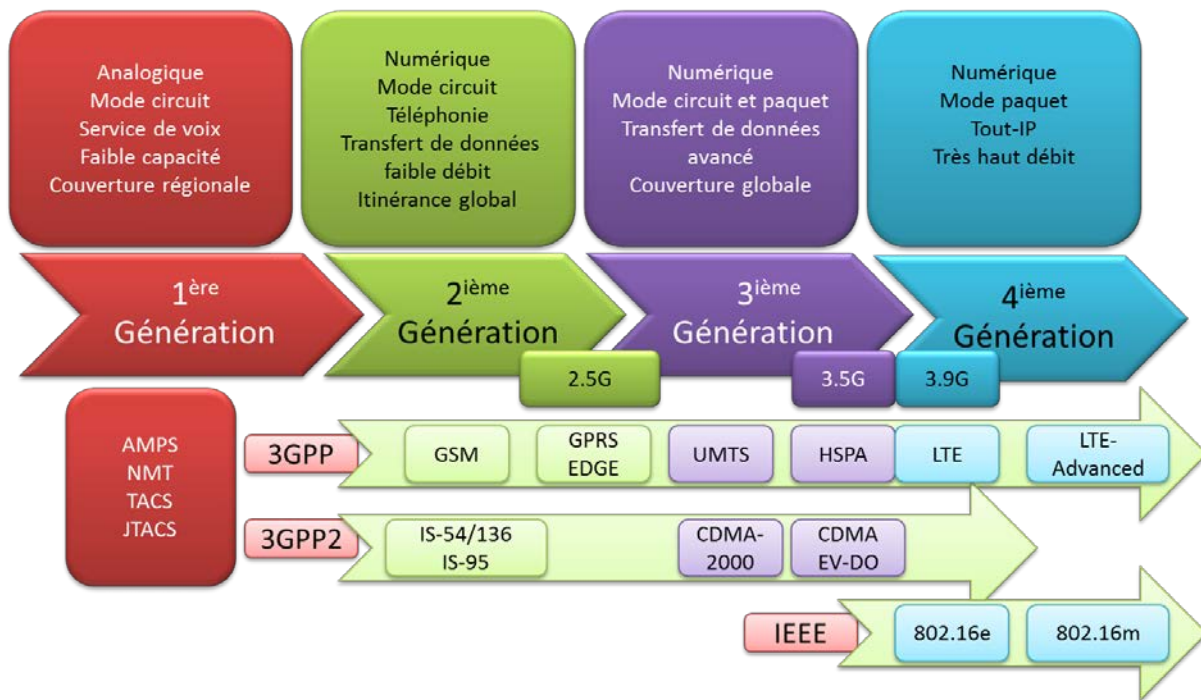


Figure 1 Evolution des standards des réseaux cellulaires

1.1.2 Les réseaux mobiles sans infrastructure

Les réseaux mobiles sans infrastructure furent inventés par les militaires américains avec le réseau PRNet (*Packet Radio Network*) en 1973 qui permet des communications ad-hoc mobiles robustes. Le groupe de travail MANET de l'IETF (*Internet Engineering Task Force* (IETF MANET Working Group)) fut formé dans le milieu des années 90. Par ailleurs, l'IEEE (*Institute of Electrical and Electronics Engineers*) a mené la standardisation de la norme 802.11 en mode maillé. Le monde de la recherche s'empara alors de ce sujet porteur. Ces réseaux se déclinent en trois familles:

- **Les MANET.** Ces réseaux mobiles fournissent des services de données et de voix. Ils sont composés de nœuds mobiles possédant une interface sans-fil, qui repose majoritairement soit sur une norme propriétaire soit sur du WiFi. Ce type de réseau a connu une effervescence au niveau de la recherche académique. Cependant, les utilisations réelles de ces réseaux se réduisent à des applications de communications tactiques.
- **Les réseaux VANET (*Vehicular Ad-hoc NETWORK*).** Cette famille n'est apparue qu'après les réseaux MANET. Les nœuds sont intégrés dans des véhicules. Cette famille englobe deux catégories: les communications entre véhicules (V2V) et les communications entre une infrastructure et les véhicules (V2I). Nous considérerons, dans cette partie, uniquement les communications V2V qui sont réellement sans infrastructure. Les applications fournies englobent un large panel qui va des services de divertissement et d'information jusqu'à des applications de sécurité.
- **Les réseaux de capteurs, WSN (*Wireless Sensor Network*).** Les contraintes de ces réseaux sont plus importantes en termes de miniaturisation, de consommation d'énergie ainsi que de coût. Des standards ont vu le jour pour répondre à ces problématiques en s'adaptant aux caractéristiques des applications.

1.2 Les systèmes satellite

Tous les systèmes satellite sont constitués de trois parties : le segment utilisateur, le segment sol et le segment spatial, représentés sur la Figure 2. Le segment utilisateur représente le terminal satellite ainsi que le réseau au-delà de ce terminal. Ce segment va permettre de définir les trois différents types de systèmes proposés dans le monde satellite en fonction de la mobilité du terminal et du service fourni. Le segment sol englobe les entités responsables du contrôle du réseau, de la gestion du satellite et de l'interconnexion avec des réseaux extérieurs. Quant au segment spatial, il est composé du satellite lui-même ou d'une constellation de satellites. La majorité des satellites de communication utilise une orbite géostationnaire et sont dits transparents. Le terme transparent signifie que le satellite n'est qu'un répéteur. Cependant, le satellite est extrêmement évolué et complexe puisqu'il amplifie le signal jusqu'à un facteur de 1000 milliards. *A contrario*, et depuis plusieurs années, des études sont menées pour faire évoluer ces satellites géostationnaires et ajouter des processeurs embarqués (OBP, *On Board Processing*) (F Vallejo, September 2005) que ce soit pour régénérer le signal ou pour déporter des fonctionnalités de routage ou d'allocation de ressources à l'intérieur du satellite. Dans le cadre de constellation de satellites, le segment spatial doit aussi gérer les communications inter-satellites (Holma & Toskala, 2011).

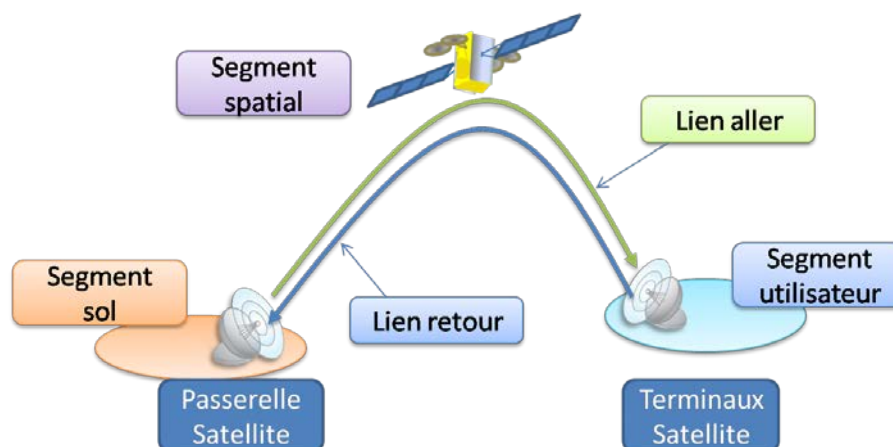


Figure 2 Architecture des systèmes satellite

Les systèmes satellite sont classés traditionnellement en trois grandes catégories. Les systèmes les plus répandus sont les satellites de diffusion (*broadcast*). En raison de leur importance économique, ils occupent une place privilégiée dans le monde des satellites de communication. Il y a ensuite les systèmes fixes et les systèmes mobiles qui fournissent des services de voix et des services de données. Nous proposons une quatrième catégorie qui prend en compte les systèmes portables. Ils utilisent des terminaux des systèmes fixes et mobiles avec comme caractéristiques d'être transportables et de se déployer rapidement.

1.2.1 Les systèmes satellite de diffusion

Les systèmes satellite de *broadcast* ou BSS (*Broadcast Satellite System*) représentent 80% du marché des télécommunications par satellite. Ces systèmes proposent des services de diffusion de la télévision et de la radio. Ils reposent sur des satellites en orbite géostationnaire. À de rares exceptions, des satellites HEO (*Highly Elliptical Orbit*) sont utilisés pour des régions à des latitudes élevées ; c'est le cas des satellites de télévision russes. Différents protocoles reposent principalement

sur du codage vidéo MPEG-2 ou MPEG-4 (*Motion Picture Experts Group*, (Motion Picture Experts Group)). La norme DVB-S (*Digital Video Broadcasting – Satellite*, (European Telecommunications Standards Institute (ETSI), August 1997)), standardisée par l'ETSI (*European Telecommunications Standards Institute*), est la plus répandue en Europe. Elle est aussi très utilisée aux États-Unis comme par exemple par Dish Networks. Quant au Japon, le standard ISDB-S (*Integrated Services Digital Broadcasting-Satellite*) fait office de référence. L'évolution de la norme de l'ETSI, DVB-S2 (European Telecommunications Standards Institute (ETSI), April 2009), connaît un succès international et tend à remplacer les anciennes normes. Elle offre des performances très supérieures en comparaison d'autres standards (de l'ordre de 30% par rapport à la précédente norme (Morello & Mignone, 2006)).

1.2.2 Les systèmes satellite fixes

Les systèmes satellite fixes FSS (*Fixed Satellite System*) fournissent des services bidirectionnels de voix et de données pour un terminal statique. De même que pour les systèmes de diffusion, les satellites suivent des orbites géostationnaires ou hautement elliptiques. Leur développement fut intimement lié à l'interconnexion de réseaux. Ils permettaient des liaisons transocéaniques en particulier pour des réseaux téléphoniques. Avec l'explosion du Web, l'utilisation de ces systèmes s'est portée vers des applications de données. Dans un premier temps, les systèmes fixes se sont fondés sur une convergence par le protocole ATM (Akyildiz & Jeong, 1997), cependant la prédominance acquise par le protocole IP en particulier dans le cadre des réseaux de nouvelles générations, a consacré l'utilisation du protocole IP comme couche de convergence.

De nombreux systèmes fixes réutilisent les protocoles DVB-S et DVB-S2 pour le lien aller. Par conséquent, la convergence IP a soulevé de nombreux problèmes, en particulier au niveau de l'encapsulation des paquets IP dans des trames MPEG-2 prévues pour la transmission de vidéo (Vasquez-Castro, Cardoso, & Rinaldo, May 2004) (Hobaya, et al., September 2010). La voie retour fut historiquement fondée sur des protocoles propriétaires tels qu'iDirect (iDirect, 2008). Avec la nécessité d'une standardisation sur la voie retour, la norme DVB-RCS (*DVB-Return Channel Satellite*) fut proposée par l'ETSI. D'autres initiatives pour la standardisation furent effectuées, Internet Protocol over Satellite (IPoS, (Telecommunications Industry Association (TIA);, 2012)) par Hugues Network Systems. Le standard S-DOCSIS (*Satellite-Data Over Cable Service interface Specification*) développés par des compagnies nord-américaines comme ViaSat, WildBlue et CableLabs a la particularité de ne pas être fondé sur le standard DVB-S et DVB-S2 pour la voie aller.

System satellite	DVB-RCS	IPoS	S-DOCSIS
Débit maximum typique du lien retour (Kb/s)	2048	2048	1203 (2406)
Débit maximum typique du lien aller (Mb/s)	DVB-S : 45/68 DVB-S2 : 150	DVB-S : 45/68 DVB-S2 : 150	50 (108)
Avantage	Standard libre, compatible avec plusieurs équipements	Terminal bon marché, un grand nombre de terminaux en service	
Inconvénients	Terminaux chers	Seulement Hugues Network commercialise cette norme	Propriétaire, beaucoup d' <i>overhead</i> sur lien retour pour des petits paquets IP

Tableau 1 Description des standards satellite fixes

Les acteurs du monde satellite sont très actifs au niveau de la recherche dans le domaine des systèmes satellite fixes afin de proposer des performances et des capacités de plus en plus importantes. Les nouveaux satellites large bande tels que (Ka-Sat, ViaSat-1) atteignent des capacités totales de 70Gb/s à 140 Gb/s. Par ailleurs, de nombreuses études proposent des améliorations qui permettront de parvenir à des capacités de l'ordre du téraoctet par seconde. Cette recherche repose sur des améliorations systèmes avec l'utilisation de bandes de fréquence plus larges, l'amélioration des techniques de réutilisation de fréquences et le développement de mécanismes de compensation d'interférences (Vidal, Verelst, Lacan, Alberty, Radzic, & Bousquet, October 2012). En outre, de nombreux mécanismes ont été implantés pour augmenter les performances des protocoles de transport comme TCP et UDP ainsi que des protocoles applicatifs comme HTTP (Caini, Firrincieli, & Lacamera, 2007) (iDirect, 2010). L'évolution de ces standards se dirige vers de nouvelles formes d'onde pour accroître leur efficacité spectrale.

1.2.3 Les systèmes satellite mobiles

Les systèmes satellite mobiles, MSS (*Mobile Satellite System*) sont traditionnellement utilisés pour des services de voix. Ils ont suivi la tendance des réseaux mobiles terrestres en offrant des services de données de plus en plus performants. Deux catégories composent les MSS : les MSS en orbite basse (LEO) et les MSS en orbite géostationnaire ou GMSS.

Les MSS LEO

L'utilisation d'une orbite basse aboutit à l'obtention d'un délai bien inférieur à celui des satellites en orbite géostationnaire ce qui est des plus bénéfiques pour des services de voix. Cependant, la couverture mondiale nécessite une constellation de plusieurs dizaines de satellites pour fournir un service continu. Les deux principaux systèmes mobiles LEO sont les constellations IRIDIUM et GLOBALSTAR (Pratt, Raines, Fossa JR., & Temple, 1999) (Globalstar).

Les GMSS

Les MSS LEO ne sont pas rentables pour des couvertures non-mondiales. Dans ce contexte, les satellites géostationnaires sont utilisés pour assurer les communications pour les terminaux mobiles. Les normes de cette catégorie sont regroupées dans la famille GMR (*GEO Mobile Radio interface*) de l'ETSI. La première norme, GMR-1, fut développée par Hugues Networks (European Telecommunications Standards Institute (ETSI), July 2009). La première version est une adaptation de l'interface air du GSM sur un lien satellite. Les versions suivantes ont eu pour but d'intégrer l'accès à l'internet, tout en s'appuyant sur les standards mobiles terrestres, GPRS et EDGE (Figure 3). La particularité des normes GMR-1 est la réutilisation de ces standards 2.5G de l'ETSI. Ainsi les couches de la Non Access Stratum sont conservées et les spécifications du terrestre sont directement applicables. Les couches de l'Access Stratum sont pour la plupart modifiées. La norme concurrente GMR-2 fut promue par ACeS (*Asia Cellular System*) pour de la téléphonie mobile en Asie et n'est maintenant utilisée que par l'ISatPhone d'Inmarsat.

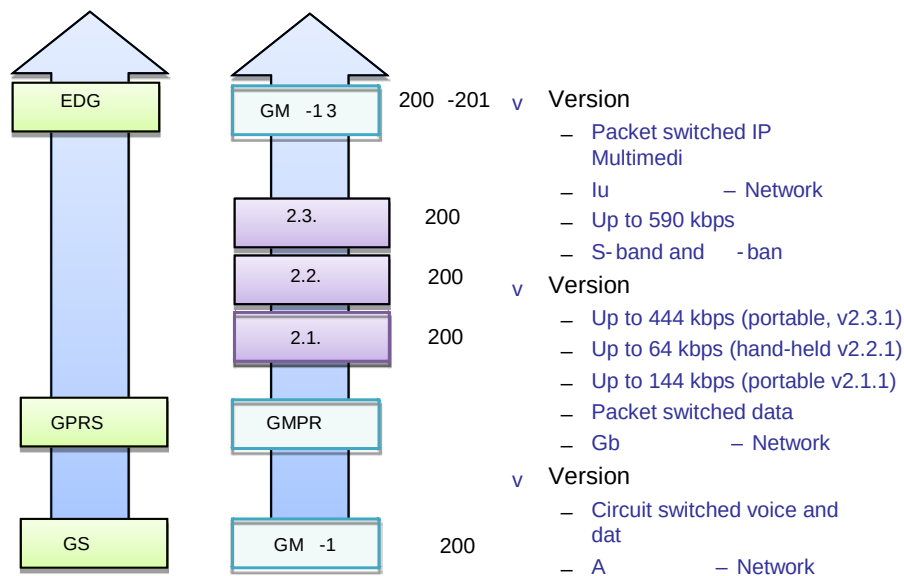


Figure 3 Les évolutions des standards GMR-1

	MSS LEO		GMSS	
Système	Iridium	Globalstar	GMR-1 (handheld)	GMR-2
Débit aller	19 kb/s	19 kb/s	160 kb/s	20 kb/s
Débit retour	2,6 kb/s	2,4 kb/s	30 kb/s	-

Tableau 2 Caractéristiques des systèmes satellite mobiles

1.2.4 Les systèmes satellite portables

Les systèmes portables ne sont traditionnellement pas une famille de systèmes à part entière. Ils correspondent aux systèmes dont les terminaux ne sont pas mobiles, mais possèdent la propriété de se déployer rapidement et sont facilement transportables. Ces terminaux spéciaux utilisent cependant des systèmes soit mobiles soit fixes. Des terminaux de systèmes mobiles peuvent subir des modifications dans le but d'atteindre de meilleures performances. En conséquence, ils ne prennent plus en charge la mobilité, mais sont opérationnels en quelques secondes ou quelques minutes comme par exemple le BGAN portable d'Inmarsat qui fournit des débits maximaux de 492 kb/s. D'autres ont pour origine les systèmes fixes. Leurs terminaux sont dotés d'une antenne de faible envergure (VSAT, *Very Small Aperture Terminal*) rendant leur installation rapide (de l'ordre d'une heure). S'ils sont équipés d'un mécanisme de pointage automatique ou d'un asservissement de l'antenne, le temps d'installation est réduit à quelques minutes pour atteindre en moyenne des débits de l'ordre de 10 Mb/s avec du DVB-S.

	MSS portable			FSS portable	
Système	BGAN	Thuraya	GMR-1 3G	VSAT (DVB-S)	VSAT (DVB-S2)
Débit aller	492 kb/s	444 kb/s	590 kb/s	~10 Mb/s	~40 Mb/s

Tableau 3 les systèmes satellite portables

1.3 La classification des réseaux hybrides

1.3.1 Description

La classification des réseaux hybrides mobiles est représentée dans la Figure 4. Elle repose sur trois critères principaux : le scénario de déploiement, l'architecture du réseau puis le positionnement du lien satellite à l'intérieur de cette architecture.

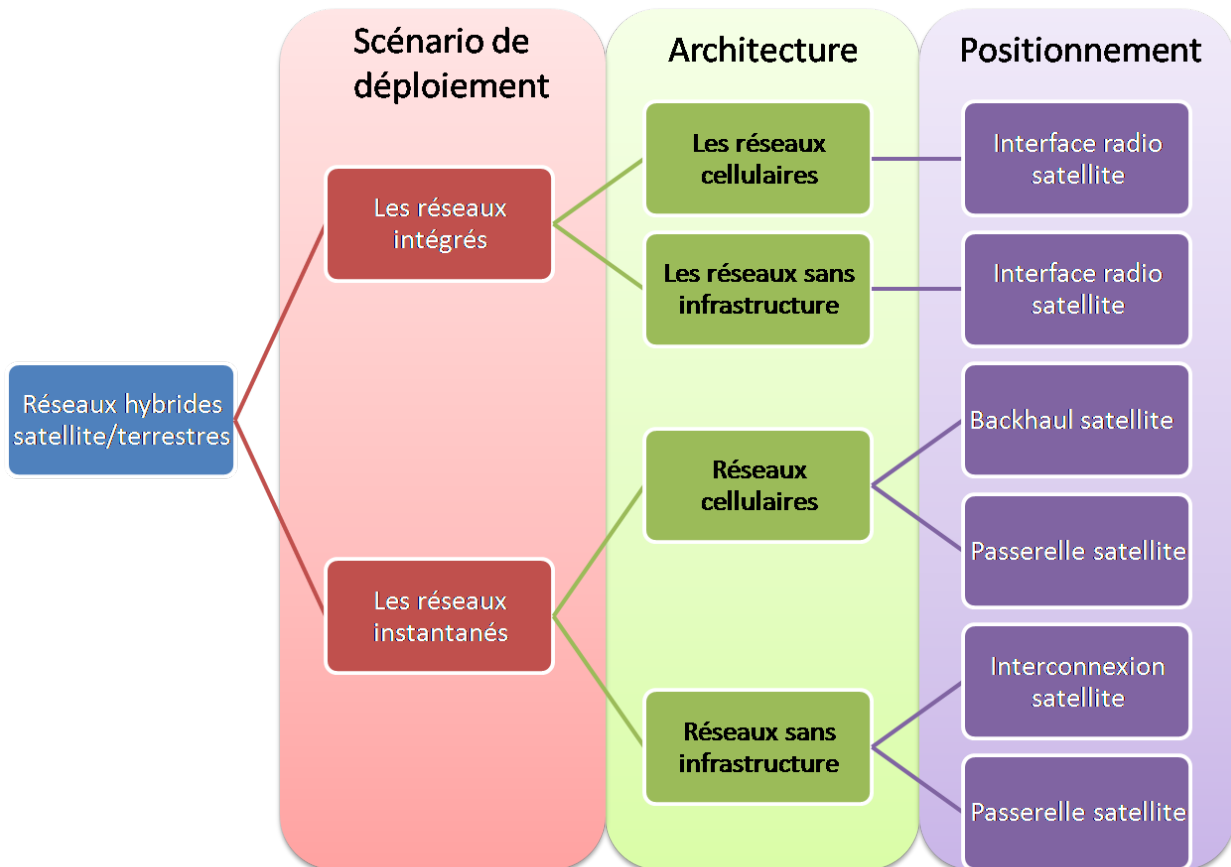


Figure 4 La classification des réseaux hybrides mobiles

Le scénario de déploiement

Le premier critère est primordial, car il définit deux manières de tirer parti des systèmes satellite dans le contexte d'une hybridation de réseaux terrestres.

- **Les réseaux intégrés** ; cette première catégorie représente des réseaux déployés de manière permanente sur la globalité de la couverture d'un satellite. Les fers de lance de cette classe de réseaux hybrides sont les réseaux cellulaires intégrés. Ils se définissent par l'utilisation d'un terminal mobile qui possède deux interfaces, une satellite et une terrestre. Cette particularité induit obligatoirement d'utiliser des systèmes satellite mobiles.

- **Les réseaux instantanés ;** cette deuxième catégorie offre la possibilité d'établir des réseaux de plus faible envergure en n'utilisant qu'une partie de la couverture offerte par le système satellite. À la différence des réseaux intégrés, le nombre de terminaux satellite est plus réduit. En effet, il est, la plupart du temps, limité à quelques-uns. Dans la majeure partie des déploiements, il est d'ailleurs réduit à un. Le lien satellite n'a, dans ce contexte, qu'un rôle d'interconnexion.

L'architecture du réseau

L'architecture du réseau hybride est le second critère. Dans la mesure où elle repose sur l'architecture du réseau mobile terrestre, nous obtenons deux nouvelles branches dans la classification avec les réseaux hybrides avec et sans infrastructure. Cette divergence est nécessaire pour aborder les problématiques spécifiques à l'hybridation, car ces deux familles de réseaux mobiles possèdent des caractéristiques totalement différentes.

Le positionnement du lien satellite

Le troisième critère est le positionnement du satellite dans l'architecture du réseau. Pour les réseaux intégrés, ce critère n'implique aucune sous-catégorie significative, car tous les terminaux possèdent deux interfaces, une satellite et une terrestre. Par conséquent, le choix du positionnement du système satellite se trouve réduit à une seule possibilité, au niveau de l'interface radio du terminal mobile. Dans le contexte de réseaux localisés instantanés, on obtient une divergence pour les deux types d'architectures :

- Le système satellite a la fonction de passerelle pour le réseau hybride vers des réseaux extérieurs.
- Le lien satellite se situe à l'intérieur et fait partie intégrante du réseau mobile. Dans le contexte des réseaux sans infrastructure, cette catégorie représente l'interconnexion de deux parties du réseau par l'intermédiaire d'un lien satellite. Dans le cadre de réseau cellulaire, cette classe comprend les stations de base dont le backhaul est effectué par un lien satellite.

1.3.2 Les réseaux intégrés

1.3.2.1 Les bénéfices

Les bénéfices liés à l'hybridation découlent essentiellement du scénario de déploiement choisi. Dans le cadre des réseaux intégrés, les bénéfices sont plus nombreux que pour les réseaux instantanés. Ce type de réseau offre à la fois des avantages aux utilisateurs ainsi qu'aux opérateurs.

Couverture globale

Cette catégorie de réseaux hybrides offre à l'utilisateur un service sur une large zone en tirant profit de la couverture du système satellite. Dans l'optique d'offrir un service continu, le segment satellite du réseau assure le service dans des endroits géographiquement inatteignables par le segment terrestre tels que les zones maritimes et les zones montagneuses. En outre, le segment terrestre fournit la connectivité dans les lieux non propices aux communications par satellite tels que les villes.

Déploiement adapté du réseau

En plus de la couverture globale, cette classe de réseaux hybrides ouvre la voie à l'optimisation du déploiement en fonction de stratégies économiques. En effet, les zones faiblement peuplées induisent un coût de déploiement bien supérieur à la rentabilité du service. La couverture globale du segment satellite permet d'éviter le surcoût de ce déploiement, tout en assurant la continuité du service. La stratégie économique s'adapte de manière cohérente aux caractéristiques techniques des deux segments. Les zones densément peuplées sont situées dans les régions urbaines où la couverture terrestre est à privilégier alors que les zones maritimes et montagneuses sont économiquement peu rentables et le système satellite y est suffisant.

Déploiement rapide

Le déploiement d'un réseau terrestre complet est très coûteux en termes d'investissements ainsi qu'en temps. L'établissement du segment satellite permet de fournir le service sur la globalité de la couverture et celui-ci ne prend qu'un temps réduit par rapport à celui du segment terrestre pour une couverture similaire.

Résistance aux catastrophes

Les terminaux utilisateurs possèdent deux interfaces radios. En conséquence, une avarie sur le système terrestre peut être compensée par le segment satellite. En particulier, lors de catastrophes d'origine naturelle ou humaine qui dégradent partiellement les infrastructures terrestres, la continuité du service est assurée par le segment satellite. Les infrastructures satellite du segment sol sont redondantes ainsi le système satellite offre une grande résistance aux catastrophes.

Optimisation de la gestion du trafic

Lorsque les segments terrestres et satellite deviennent tout deux accessibles, une gestion du trafic appropriée permet une meilleure utilisation des ressources. Les systèmes satellite étant particulièrement adaptés aux trafics *broadcast* et *multicast*, le choix du segment peut être dépendant du type de flux. De surcroît, le segment terrestre offre un service supérieur que ce soit en termes de performances ou de délai. Cependant, en cas de surcharge, le segment satellite permettra de gérer le trafic en excédant.

Optimisation de la gestion du spectre

Les segments satellite et terrestre possèdent tous les deux une bande de fréquence qui leur est attribuée de manière exclusive. En raison du coût très élevé de cette ressource, ainsi que de sa rareté grandissante, des mécanismes ont vu le jour pour autoriser le partage du spectre fréquentiel entre les deux segments.

1.3.2.2 Les réseaux cellulaires intégrés

Les réseaux intégrés cellulaires représentent le marché le plus important pour les réseaux hybrides. Ils proposent des services de téléphonie et de données mobiles. L'adaptation des couches d'accès constitue la principale problématique pour ces réseaux. Dans un premier temps, la définition des systèmes satellite est réalisée dans l'optique de réutiliser, au maximum, les entités et protocoles des systèmes terrestres, afin de minimiser les coûts de développement. *A contrario*, les couches basses

doivent être adaptées aux caractéristiques particulières du canal satellite. Les protocoles de la couche NAS (*Non Access Stratum*) ne sont que peu affectés par le satellite et ne nécessitent que des configurations adaptées pour prendre en compte le délai induit par celui-ci. Dans un deuxième temps, cette réutilisation offre la possibilité d'une intégration plus profonde entre le segment terrestre et satellite. Des travaux ont été menés pour proposer une classification de ces différents niveaux d'intégration (Deslandes, 2012) (International Telecommunication Union - Radiocommunication sector (ITU-R), May 2009).

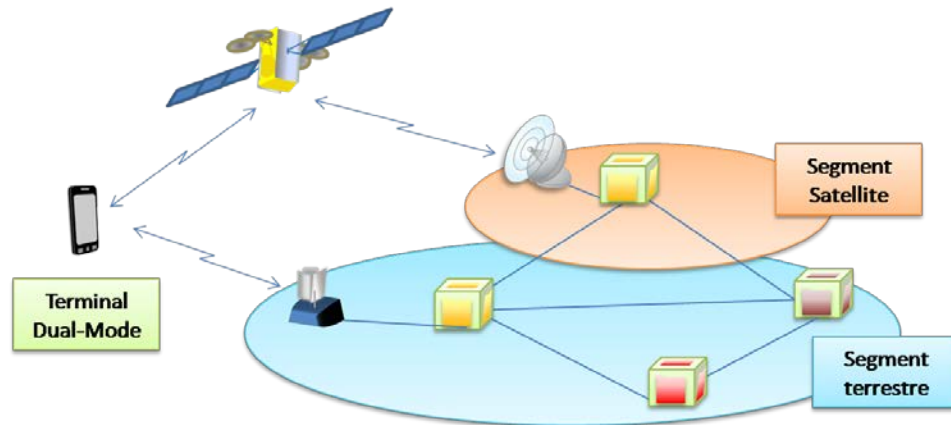


Figure 5 Architecture des réseaux intégrés cellulaires

Dans un esprit similaire à celui des acteurs terrestres, la définition du segment satellite nécessite une standardisation. C'est ainsi que les standards MSS de la famille GMR ont été développés. Ils reposent sur la standardisation des systèmes terrestres du 3GPP. Seules les couches d'accès de l'interface radio sont modifiées ; le réseau de cœur est conservé. Les deux premières versions du standard GMR-1 ne proposent qu'une gestion commune de l'itinérance avec des systèmes terrestres, alors que la dernière version, connue sous la dénomination de GMR-1 3G, ouvre la voie à une intégration plus complète en autorisant le *handover* entre un segment terrestre et satellite. Dans une approche similaire, le standard de l'ETSI, S-UMTS, fut proposé sous l'impulsion de projets de l'ESA (SATIN (Evans, et al., May 2002)). Ce standard pousse l'adaptation du WCDMA sur le satellite à un niveau supérieur à GMR-1, en proposant une intégration plus forte au niveau des couches physiques.

Les réseaux intégrés déployés sont au nombre de deux : les systèmes de Thuraya et de Terrestar (Chini, Giambene, & Kota, 2010). Thuraya propose des services dans de nombreux pays et repose sur la norme GMR-1. Cependant, il est fondé sur les deux premières versions qui ne proposent que la fonctionnalité d'itinérance (*roaming*). Dans ce contexte, le niveau d'intégration reste faible. Au contraire, Terrestar utilise le standard GMR-1 3G sur le segment satellite et grâce à des accords avec l'opérateur terrestre américain AT&T autorise le *handover* entre les segments satellite et terrestre. Ce dernier est le seul réseau intégré cellulaire opérationnel. Lighsquared a débuté le déploiement d'un réseau intégré pour l'Amérique du Nord dont le segment satellite utilise le standard GMR-1 3G. Contrairement au réseau Terrestar, le spectre est partagé entre le segment terrestre et satellite, ce qui aurait pu permettre de bénéficier des avantages du partage de spectre. En raison d'interférence avec le système de localisation GPS, le déploiement du réseau fut interrompu.

En ce qui concerne, le standard S-UMTS aucun déploiement n'a été effectué, seul un terminal de démonstration fut implanté dans un véhicule. Au Japon, le projet STICS prône l'utilisation d'un réseau

cellulaire intégré principalement en raison de sa capacité à fournir des communications après une catastrophe naturelle. Ce projet a débuté en 2008 et étudie la gestion du spectre en proposant des mécanismes de minimisation des interférences.

1.3.2.3 Les réseaux sans infrastructures

Les réseaux sans infrastructure (Figure 6) ne sont que très peu représentés et étudiés dans les réseaux intégrés. En effet, l'une des caractéristiques des réseaux sans infrastructure est le faible prix de l'équipement. L'utilisation d'un terminal dual-mode rajoute un coût qui devient rédhibitoire pour la majorité des applications.

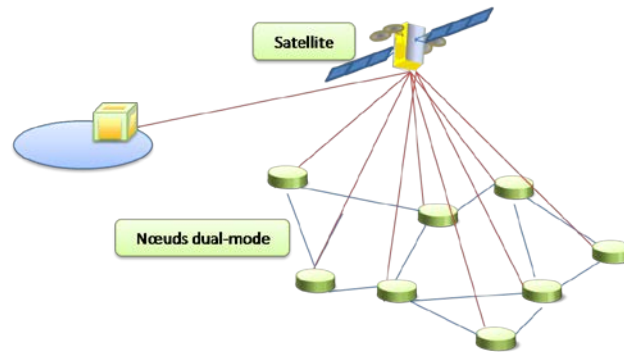


Figure 6 Architecture des réseaux intégrés sans infrastructure

L'utilisation de ce type de réseau, avec l'utilisation de terminaux dual-mode, pourrait être envisagée pour des réseaux véhiculaires hybrides avec des communications V2V et V2I. En effet, c'est le seul dans la famille des réseaux sans infrastructure, pour lequel on puisse juger acceptable le surcoût occasionné par un terminal bi-mode satellite terrestre. C'est également l'unique application qui puisse bénéficier au maximum de l'intégration et en particulier, de l'utilisation globale de la couverture mais aussi de la facilité, pour le satellite, à diffuser des flux *broadcast* et multicast. Le projet européen SafeTrip propose l'utilisation du protocole DVB-SH ainsi que du satellite W2A déjà en orbite pour fournir des terminaux VANET (SafeTrip). Les différents partenaires de ce projet ont proposé une méthode d'accès aléatoire sur la voie retour, S-MIM (Scalise, et al., June 2012), dans le but de fournir des communications bidirectionnelles. Cependant, son intégration à un réseau terrestre n'est pas encore réalisée.

Les réseaux de capteurs intégrés n'ont donné lieu à ce jour à aucune étude. En effet, pour ce type de réseau, on ne peut se permettre le surcoût induit par un terminal dual-mode. Pour cette raison, les capteurs embarquent soit une technologie terrestre, soit une technologie satellite. L'une des seules exceptions que l'on puisse rencontrer provient d'un système de positionnement de flotte maritime. Les bateaux se signalent de manière locale aux autres bateaux dans le but d'éviter les collisions. L'union européenne travaille sur un système satellite qui récupère l'information transmise par ces bateaux sans engendrer de modifications du terminal. La principale problématique dans ce cas-là provient du décodage des différents messages émis par les bateaux en utilisant des techniques de décodage récursif (Picard, Oularbi, Flandin, & Houcke, September 2012).

En ce qui concerne les réseaux MANET, les problèmes de passage à l'échelle conduisent à des performances très détériorées dès qu'il y a trop de nœuds ou de sauts. Par conséquent, la zone couverte par les solutions déployées reste restreinte. L'intérêt d'une solution intégrée est donc

faible. Par exemple, dans (Rodriguez, 2011) , l'auteur étudie un réseau intégré très spécifique où le segment satellite ne servirait qu'au flux de signalisation du protocole de routage. L'idée se résume à utiliser la capacité du système satellite pour les flux *broadcast* et multicast, pour transmettre les messages *broadcast* du protocole de routage par le système satellite et ainsi réserver le segment terrestre aux flux utilisateurs. Un des résultats importants de cette thèse a été de montrer que le gain atteint est marginal pour un surcoût non négligeable.

1.3.3 Les réseaux instantanés

1.3.3.1 Les bénéfices

Déploiement instantané

La caractéristique principale de cette catégorie provient de la rapidité de déploiement du réseau. En effet, ils sont composés de systèmes satellite et terrestres possédant cette propriété. Pour le segment satellite, cette catégorie privilégiera les systèmes satellite mobiles ou portables.

Possibilité de déploiement sur une large couverture

L'utilisation d'un système satellite permet de déployer ce type de réseau sur l'ensemble de la couverture satellite. En outre, le système satellite ne doit pas obligatoirement être dédié ; l'utilisation des systèmes satellite en fonctionnement devient primordiale pour atteindre un déploiement instantané. En fonction de la technologie satellite choisie, le réseau peut être établi par l'intermédiaire de différents satellites ce qui en augmente le champ d'action potentiel de ce réseau.

Réseaux temporaires

Leur faible envergure et leur déploiement rapide les consacrent en un réseau propice à une utilisation temporaire. En particulier, l'établissement d'un réseau d'urgence suite à une catastrophe est un des scénarios adaptés à ce type de réseaux.

Les possibilités d'optimiser la gestion du trafic sont alors restreintes. En effet, les terminaux satellite sont en tout petit nombre ; dans la majorité des cas, il n'y en a qu'un. Comme, le lien satellite est utilisé pour remplacer une technologie terrestre absente, les optimisations fondées sur le choix entre les systèmes terrestre et satellite s'avèrent inapplicables ou peu efficaces.

1.3.3.2 Réseaux cellulaires

Les réseaux cellulaires instantanés rencontrent un succès grandissant dans les pays en voie de développement. Les réseaux cellulaires terrestres n'atteignent pas le taux de couverture des pays développés, par conséquent le service mobile terrestre y est limité. Dans ce contexte, de nombreuses entreprises de services satellite proposent des équipements qui offrent la possibilité de déployer une station de base dans des régions isolées. Deux classes de réseaux hybrides sont proposées en fonction du positionnement du lien satellite dans l'architecture cellulaire.

1.3.3.2.1 Le backhaul satellite

Ce type de réseaux hybrides est très répandu pour étendre les réseaux mobiles terrestres 2G et 3G dans les zones géographiques isolées dans les pays émergents et en développement. Lorsque le lien satellite est utilisé comme backhaul, les réseaux d'accès constitués d'une seule station de base

prédominant. La simplicité du segment utilisateur assure un déploiement des plus rapides avec un minimum de mise en place d'infrastructure ainsi que de configuration. Si plusieurs stations de base sont nécessaires, l'utilisation d'un backhaul hybride est préconisée. Le backhaul hybride désigne un backhaul qui utilise une technologie satellite, avec en complément une technologie sans-fil telle que du WiMAX fixe. Dans les deux cas, les problématiques qu'entraîne l'hybridation restent très similaires. La plus importante n'est pas de nature technique mais économique, car le marché visé se situe dans les pays en voie de développement. De ce constat découle le besoin d'une optimisation drastique de la gestion des flux sur le système satellite.

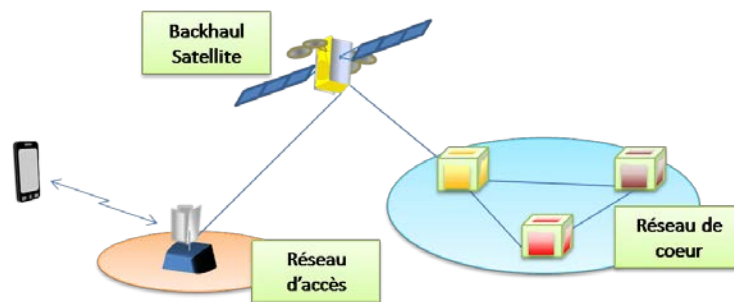


Figure 7 Architecture d'un réseau instantané cellulaire avec un backhaul satellite

La première conséquence est le passage à l'IP des flux sur le backhaul satellite en raison de terminaux satellite plus économiques. Ce changement permet aussi de réduire les coûts d'utilisation de la bande passante jusqu'à 40%. Ce type de réseaux nécessite l'utilisation de terminaux satellite VSAT qui se déploient rapidement et fournissent les meilleures performances des systèmes portables. D'un point de vue économique, le choix d'une technologie SCPC (*Single Channel Per Carrier*) est de plus en plus délaissé pour un multiplexage des flux de plusieurs terminaux satellite, que ce soit sur la voie aller ou la voie retour. En effet, l'utilisation d'une technologie SCPC réserve des ressources statiquement pour un terminal satellite ; son utilisation n'est utile que dans le cadre d'un faible nombre de stations de base avec un fort besoin en termes de débit (Anderson & Mavrakis, 2009).

De nombreuses optimisations sont mises en place pour maximiser l'utilisation de la bande passante satellite. Par exemple, des algorithmes de compression adaptés à l'interface du backhaul GSM, nommée A-bis, ainsi que de la suppression des silences sont utilisés. Comme pour les optimisations proposées dans le cadre de services satellite fixes, des mécanismes d'accélération et de cache offrent des améliorations importantes pour les trafics de données. En outre, de nombreuses offres incluent des mécanismes de communications locales pour éviter le recours au lien satellite pour des communications entre deux mobiles du même réseau hybride. Ainsi deux mobiles dans le même réseau local communiquent directement au travers de la station de base.

La grande majorité des fournisseurs de services satellite non-*broadcast* ont développé une implantation propriétaire de l'optimisation de l'interface du réseau cellulaire sur le satellite grâce à la fusion des compétences d'un acteur terrestre et d'un acteur satellite. Ainsi l'entreprise satellite Gilat s'est associé à Ericsson pour développer leur gamme de produits, de même pour ViaSat combiné avec Verso technologie (Gilat, 2010).

1.3.3.2.2 La passerelle satellite.

Le lien satellite ne sert qu'à l'interconnexion entre le réseau cellulaire et les réseaux extérieurs. Le réseau cellulaire est composé d'un seul équipement qui regroupe toutes les entités et les protocoles

du réseau de cœur ainsi que du réseau d'accès. Dans ce contexte, l'optimisation de trafic sur le lien satellite possède les mêmes caractéristiques que pour les systèmes satellite fixes. Il n'y a pas de problèmes liés à l'hybridation. La difficulté technique de ce type de réseau est située dans la simplification avancée de cet équipement, pour qu'il soit le moins coûteux et le plus miniaturisé possible. Cette simplification repose sur deux propriétés, la concaténation des entités en un seul équipement et l'unicité de la station de base. Ce déploiement ne nécessite pas de posséder un réseau de cœur déjà en fonctionnement. Cependant, le propriétaire doit détenir les droits ou posséder des accords sur des bandes de fréquences.

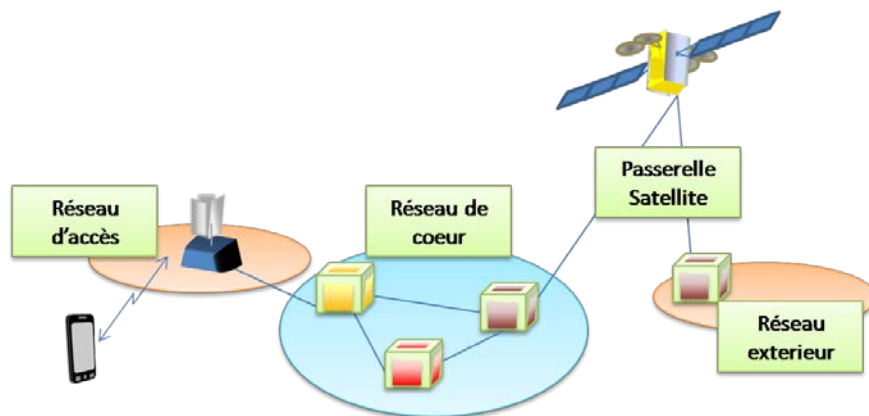


Figure 8 Architecture d'un réseau instantané cellulaire possédant une passerelle satellite

Les scénarios d'utilisation de ce réseau instantané sont principalement orientés vers des réseaux pour les communications d'urgence ou militaires. Ainsi, la compagnie Hughes Networks Systems, responsable du segment satellite, associée avec Lemko Corporation, a réalisé une démonstration en 2012 de ce réseau hybride pour des communications militaires. Le segment terrestre, développé par Lemko Corporation, adopte la technologie 4G LTE et emploie des techniques de virtualisation du réseau de cœur pour obtenir la simplification de l'équipement (Hughes Networks and Lemko Corporation, 2012).

1.3.3.3 Les réseaux sans infrastructure

Tout comme pour les réseaux cellulaires, les réseaux instantanés sans infrastructure se séparent en deux sous-catégories. La première correspond au déploiement d'un réseau, où le lien satellite est extérieur au réseau sans infrastructure et ne sert que d'interconnexion avec des réseaux externes. La seconde englobe les réseaux dont le lien satellite sert d'interface entre deux parties du réseau. Il fait alors partie intégrante du réseau sans infrastructure. Les problématiques dans ces deux cas sont différentes. En effet, dans le premier cas, on s'intéresse aux mécanismes d'interconnexion alors que dans le second, la difficulté résidera dans l'adaptation des protocoles du réseau sans infrastructure sur le lien satellite.

La spécificité de ce type de réseau n'est pas intéressante pour les VANET. En effet, un déploiement localisé n'apporterait aucun bénéfice du fait de la trop grande mobilité des nœuds. Seule l'utilisation du lien satellite comme passerelle est intéressante dans le cadre de réseaux de capteurs hybrides.

1.3.3.3.1 WSN avec une passerelle satellite

L'utilisation d'une passerelle satellite comme puits pour un réseau de capteurs est une technique intéressante, en particulier pour des zones géographiques isolées. Il conviendra alors d'optimiser l'utilisation du lien satellite grâce à un protocole adapté et un système de compression efficace. Cependant, très peu d'études ont été réalisées dans le cadre de nœuds mobiles.

Des projets européens de réseaux de capteurs avec une passerelle satellite sont proposés pour permettre l'exploration de planète. Ainsi, le projet SWIPE (*Space Wireless sensor networks for Planetary Exploration*) a débuté en avril 2013 et dont la principale problématique liée à l'hybridation sera l'interconnexion entre des réseaux longue distance obtenus par un lien satellite et des réseaux sans infrastructure WSN déployé sur la Lune.

1.3.3.3.2 MANET avec une passerelle satellite

Les MANET n'ont pas rencontré le succès commercial escompté et sont réduits à des utilisations essentiellement militaires. Cependant, le recours à un MANET hybride a donné lieu récemment à des projets, en particulier pour la sécurité civile. En effet, la complémentarité entre les systèmes satellite et terrestres rend cette catégorie de réseaux hybrides très attractive. Premièrement, les propriétés de déploiement rapide du terminal satellite avec des systèmes mobiles ou portables rentrent en résonance avec les fonctionnalités d'auto-configuration du réseau sans infrastructure. Deuxièmement, la décentralisation de ce réseau le rend intéressant pour gérer des communications d'urgence. De manière générale, l'un des freins aux MANET est la difficulté à maintenir une qualité de service de bout en bout à cause du multi-bond des communications. Ce problème est une caractéristique générale des MANET et n'est pas spécifique à l'hybridation satellite du système.

Le fonctionnement décentralisé du MANET provoque une modification de bien des mécanismes de bout en bout, tels que la gestion de la sécurité et la gestion de la QoS. Ainsi, la passerelle satellite doit s'adapter au choix d'implantation de ces mécanismes, en particulier pour assurer une qualité de service cohérente. La gestion de la mobilité est réalisée par le protocole de routage dynamique entre les nœuds du réseau. La gestion de la table de routage au niveau de la passerelle du système satellite doit être modifiée pour minimiser la part des messages du protocole de routage passant par le lien satellite.

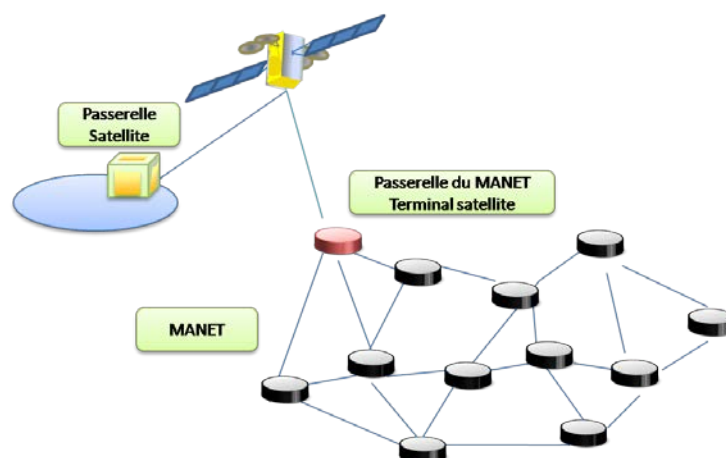


Figure 9 Réseau MANET instantané avec une passerelle satellite

Le projet DUMBO (DUMBO) est né en réaction au tsunami dans l'océan indien qui a détruit les côtes asiatiques en 2004. Il a permis la validation de la réalisation de communications d'urgence grâce à un système hybride MANET. Dans le but d'assurer l'interconnexion de différents MANET par le lien satellite et d'assurer la cohérence de l'adressage, les partenaires de ce projet ont opté pour l'utilisation du protocole Mobile-IP.

Le projet SAVION (Luglio, Monti, Roseti, Saitto, & Segal, 2007) est un projet italien qui propose un réseau hybride satellite MANET, principalement pour des communications de voix pour les pompiers et assurer l'interconnexion des pompiers entre eux, ainsi qu'avec des personnes éloignées du théâtre des opérations.

1.1.1.1.1 MANET avec une interface satellite interne

Comme le lien satellite se situe cette fois à l'intérieur du MANET, il doit simplement s'adapter aux protocoles utilisés dans le MANET et en particulier au protocole de routage. Cette configuration est plutôt rare, car l'utilisation d'un système satellite fournit une entité centralisée dans le réseau. Ainsi, la mise en place d'un système satellite entraîne de manière systématique l'utilisation d'une passerelle satellite. Les deux parties du MANET sont alors raccordées par l'intermédiaire de celle-ci. Même si l'utilisation du lien satellite de manière interne est possible, elle devient redondante.

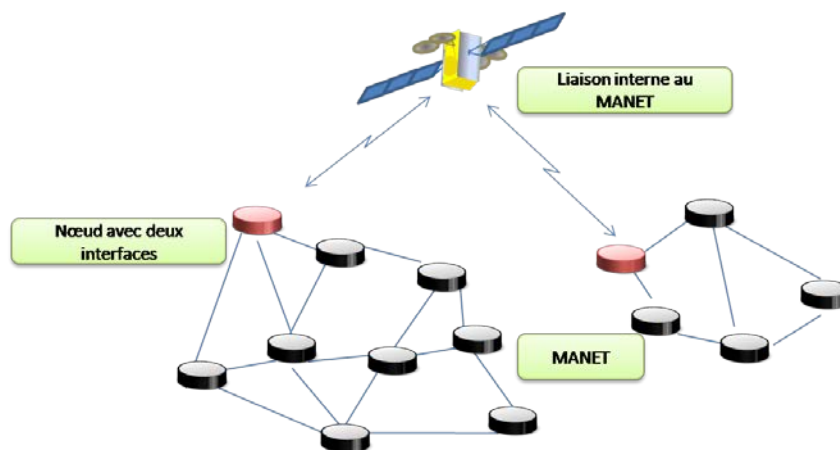


Figure 10 Réseau MANET instantané avec une interface satellite interne

1.4 Axes d'optimisations

La Figure 11 présente un récapitulatif de la classification des réseaux hybrides. Les réseaux hybrides intégrés induisent une complémentarité maximale entre les réseaux mobiles terrestres et satellite. Cependant, la complexité et le coût de déploiement de ce type de réseaux les rendent très risqués pour les opérateurs commerciaux. Seul le marché des Etats-Unis a les caractéristiques économiques et géographiques pour pouvoir héberger ce type de réseau. De plus, les clients éprouvent le désir que l'offre de terminaux mobiles soit la plus vaste possible. L'aspect dual-mode des terminaux des réseaux intégrés ne permet de proposer qu'un panel très limité, car la réutilisation des téléphones mobiles devient impossible. Cependant, ce dernier point n'entrave pas le déploiement de réseaux cellulaires pour des communications de sécurité civile, dont le besoin d'une couverture globale et de résistance aux catastrophes est plus prononcé. Le développement de réseaux intégrés repose sur

une meilleure intégration des segments terrestres et satellite en particulier au niveau du terminal avec la création potentielle d'un nouveau standard GMR-1 4G.

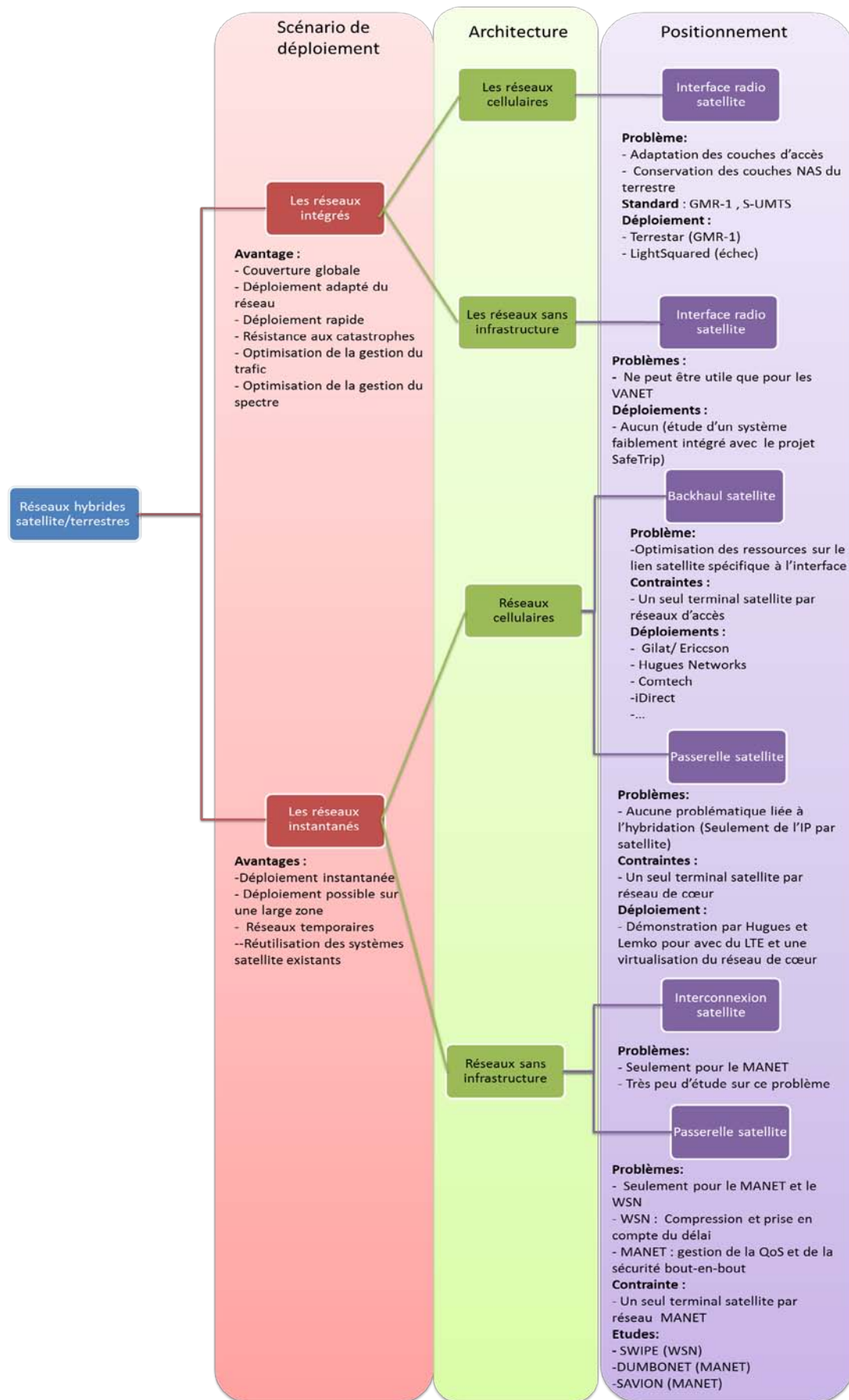


Figure 11 Récapitulatif de la classification

Pour les réseaux instantanés, le dimensionnement devient l'axe le plus contraignant. En effet, avec le succès technologique que rencontre la norme LTE, l'efficacité spectrale atteinte par une station de base est multipliée d'un facteur 3 par rapport à la 3,5G HSPA. Les améliorations des technologies satellite n'égale pas les progrès de l'interface radio terrestre. Le backhaul satellite devient limitant pour le déploiement d'une station de base. L'utilisation de ce même lien par plusieurs stations de base est encore plus inefficace. Pour créer un réseau instantané efficace, plusieurs terminaux satellite sont donc nécessaires.

De même, pour les réseaux instantanés MANET, l'utilisation d'un seul lien satellite ne fournit pas des performances optimales. Le déploiement d'autres terminaux satellite dans le MANET permet d'augmenter la capacité des passerelles. En outre, les systèmes satellite utilisés ont besoin d'être facilement transportables et utilisent, de manière préférentielle, les systèmes satellite portables ayant pour origine les systèmes mobiles. Leur débit n'atteint en moyenne que 512 kb/s, ce qui est insuffisant pour prendre en charge les applications de l'ensemble du réseau MANET. Deuxièmement, l'ajout de lien satellite multiplie le nombre de passerelles dans le MANET ce qui entraîne une limitation des congestions dans le MANET et augmente ainsi les performances des applications.

Dans les deux sous-catégories des réseaux instantanés, le problème de dimensionnement doit être résolu. La solution que nous proposons est simple ; elle consiste à ajouter des terminaux satellite pour augmenter la capacité fournie par le système satellite. Cependant, cette intention aboutit à de nouvelles difficultés technologiques.

Pour les réseaux cellulaires, deux mécanismes sont touchés :

- **La gestion de la mobilité.** Dans les scénarios précédents, les mécanismes de *handover* ne souffrent pas du lien satellite puisque les messages sensibles au délai ne l'empruntent pas. Ce constat n'est plus valable lors d'un *handover* entre des stations de base possédant un backhaul satellite différent.
- **La configuration des stations de base.** La configuration des stations de base devient plus complexe, car ces entités sont séparées l'une de l'autre par des liaisons satellite. En outre, les stations de base peuvent être déployées rapidement. L'architecture de ces réseaux est plus dynamique, ce qui entraîne des problèmes de configuration de l'interface radio en fonction des stations de base voisines.

Considérant les réseaux hybrides MANET, deux problèmes surviennent à cause de l'ajout de terminaux satellite :

- **L'élection de passerelle dans le MANET.** L'élection de passerelle permet de choisir les terminaux satellite qui doivent être activés pour satisfaire les besoins en ressources du MANET.
- **La sélection de la passerelle.** La sélection de la passerelle est exécutée par la suite. Elle est intimement liée au protocole de routage dynamique et définit la passerelle de destination d'un paquet.

Les chapitres suivants expliquent de manière plus poussée les problèmes qu'implique l'utilisation de plusieurs terminaux satellite pour un réseau instantané. Les chapitres 2, 3, 4 et 5 développeront les implications et les optimisations liées à l'utilisation de plusieurs terminaux satellite dans le

déploiement d'un réseau de stations de base LTE. Les chapitres 6 et 7 s'intéresseront plus spécifiquement à la sélection de passerelle dans un MANET pour des communications de sécurité civile.

Chapitre 2 : Réseaux hybrides LTE-satellite : Description et analyse du standard.

Nous nous intéressons, dans ce chapitre, à la catégorie des réseaux cellulaires hybrides avec un backhaul satellite. Ce type de réseaux est très utilisé par les opérateurs de téléphonie mobile pour couvrir des régions isolées en particulier dans les pays en voie de développement. Nous souhaitons développer ce type de scénarios en utilisant le nouveau standard LTE. Ce chapitre illustre les inadéquations entre les spécifications du standard et l'utilisation de plusieurs backhuls satellite dans un réseau LTE. Après la description des scénarios d'utilisation, différents éléments de la spécification du standard sont présentés pour nous permettre de choisir l'architecture adaptée à ce type de réseau, ainsi que d'analyser les conséquences de cette architecture sur le standard.

2.1 Scénarios étudiés

Les réseaux cellulaires hybrides avec un backhaul satellite sont destinés à des déploiements dans les pays en voie de développement. Les offres actuelles ne se concentrent que sur des services de voix fournis par des stations de base GSM (Anderson & Mavrakakis, 2009) même si l'évolution des stations de base vers des systèmes de troisième génération est déjà entamée (ComTech EF DATA, 2011). L'évolution naturelle de ce type de réseau est évidemment liée à la quatrième génération de la téléphonie mobile. *Long Term Evolution*, LTE, promu par le consortium 3GPP est le standard victorieux à la course pour la quatrième génération de communications mobiles. Il a rattrapé son retard sur le WiMAX mobile et a été adopté par presque tous les acteurs du marché de la téléphonie mobile. Le déploiement de ces réseaux et leur commercialisation a débuté fin 2009, avec Telionara en Suède et NTT DoCoMo au Japon. Par la suite, la majorité des opérateurs mobiles ont choisi le standard LTE et en déploient les infrastructures dans le monde entier.

L'évolution vers une technologie LTE pose un problème de dimensionnement du système. En effet, les besoins en capacité au niveau du backhaul sont multipliés par le passage de la troisième génération à la quatrième. Ainsi, un terminal satellite qui opère en tant que backhaul pour une station de base LTE représente déjà un aspect limitant. L'utilisation d'un lien satellite pour plusieurs stations de base aggrave ce problème qui devient alors très problématique. Afin de réduire la limitation des performances du système due au backhaul satellite, le déploiement d'un réseau composé de plusieurs stations de base ne possèdera pas qu'un backhaul satellite, mais chaque station de base utilisera son propre terminal satellite pour avoir une indépendance vis-à-vis de cette liaison.

En plus de l'augmentation de la capacité du système, ce type de déploiement améliore la rapidité d'installation des réseaux de cette nature. En effet, il est inutile de déployer des liens terrestres entre les stations de base et le terminal satellite. Ce scénario est rendu possible grâce à des mécanismes d'auto-configuration et d'auto-optimisation. Ces techniques, nommées SON (Self-Organized, Self-Optimized Networks), ont été largement développées par le consortium 3GPP dans le but de

diminuer les dépenses d'exploitation, OPEX (installation, configuration et gestion du réseau) (4G Americas, June 2011). Elles permettent de configurer plus rapidement et plus efficacement le réseau après son déploiement, ainsi que d'optimiser l'utilisation du réseau, grâce par exemple à une gestion automatisée des interférences (Hämäläinen, Sanneck, & Sartori, 2012).

En raison de cette plus grande flexibilité de déploiement, les scénarios d'utilisation de ce type de réseaux s'élargissent. Ils ne se concentrent plus que sur les régions isolées, mais ils s'étendent à des réseaux temporaires pour assurer la communication lorsque les infrastructures terrestres sont détruites ou endommagées, ainsi que lors d'événements où la capacité du réseau traditionnel devient insuffisante. Ces caractéristiques intéressantes pour les opérateurs commerciaux le sont aussi pour des réseaux de sécurité civile. Cette idée est renforcée par le fait qu'auréolé de sa victoire dans le monde commercial, LTE a été choisi ou est pressenti pour la nouvelle génération des systèmes de communication de la sécurité civile. En 2012, les organisations gouvernementales nord-américaines ont élu LTE/EPC comme la prochaine technologie pour leur réseau (Federal Communications Commission (FCC), January 2011). De même en Europe, LTE est devenu le standard favori.

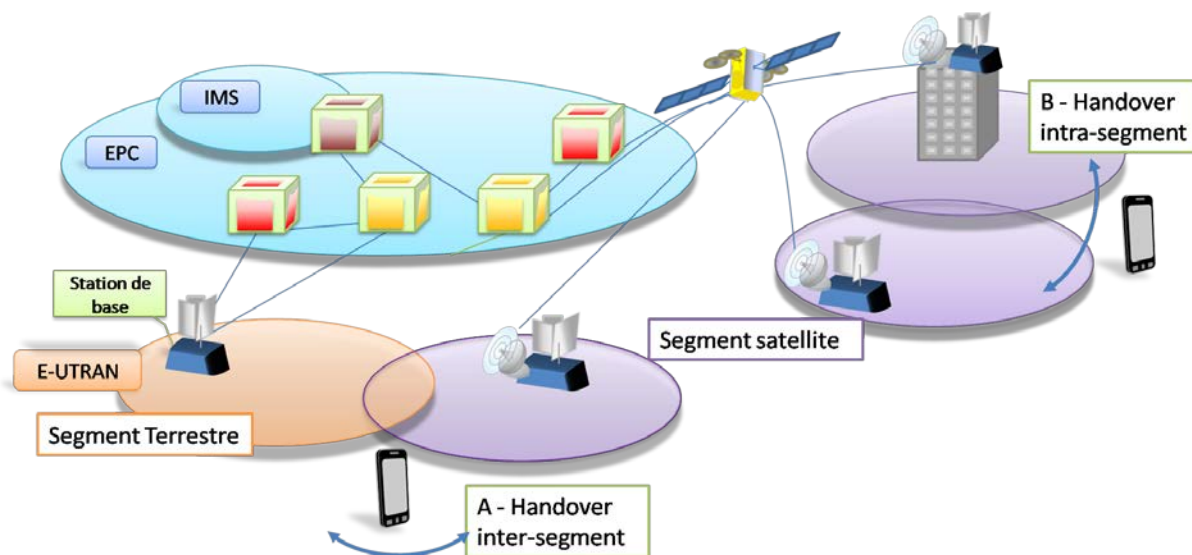


Figure 12 Nouveaux scénarios d'architecture LTE avec un backhaul satellite

Dans l'optique d'intégrer le segment satellite de la façon la plus transparente dans le réseau de l'opérateur, des optimisations en termes de mobilité sont nécessaires. En effet, les nouveaux scénarios d'utilisation de ce type de réseau induisent des besoins de modifications des mécanismes liés à trois nouveaux types de *handovers*:

- *Le handover inter-segments du terrestre vers le satellite ou inter-terrestre-satellite.* L'UE passe d'une station de base traditionnelle vers une station de base avec un backhaul satellite (Figure 12. A).
- *Le handover inter-segments du satellite vers le terrestre ou inter-satellite-terrestre.* L'UE passe d'une station de base avec un backhaul satellite vers une station traditionnelle (Figure 12. A).

- *Le handover intra-segment satellite ou intra-satellite.* L'UE reste dans le segment satellite avec un changement de station de base, en restant connecté à un eNB satellite (Figure 12. B).

2.2 Le standard LTE/SAE

Le standard de quatrième génération LTE est à l'origine d'un bouleversement aussi bien pour le réseau de cœur que pour le réseau d'accès. L'évolution du réseau de cœur est nommée *Evolved Packet Core*, EPC. le réseau d'accès, qui est réduit à la station de base, prend le nom d'*Evolved-UTRAN* (E-UTRAN). Deux projets distincts portent la responsabilité de ces évolutions dans le 3GPP : le projet LTE est en charge de l'E-UTRAN alors que SAE (*System Architecture Evolution*) se préoccupe de l'EPC. Le système complet est appelé *Evolved Packet System* (EPS). Cependant, la dénomination de LTE englobe de manière récurrente à la fois l'évolution des couches d'accès et du réseau de cœur. Par conséquent, dans cette thèse, LTE pourra englober le standard général. De la même façon, EPS et LTE/SAE auront la même signification. Les spécifications mises à disposition librement par le 3GPP (The 3rd Generation Partnership Project (3GPP)) sont décrites dans de nombreux ouvrages tels que (Sesia, Toufik, & Baker, 2011).

2.2.1 Architecture standard de LTE

L'architecture standard de LTE/SAE sert de socle pour la présentation des différentes architectures utilisées dans cette thèse. Du point de vue de l'interface radio, E-UTRAN, les spécifications seront grandement modifiées par rapport aux standards précédents, dans le but d'atteindre des débits très largement supérieurs. De même, le réseau de cœur a été modifié pour ne reposer que sur le protocole IP et devenir l'EPC.

L'architecture de base d'un réseau LTE/SAE se compose de diverses entités (The 3rd Generation Partnership Project (3GPP), September 2012). Le réseau d'accès, E-UTRAN se réduit à la station de base, intitulée *Evolved Node-B* (eNB) et le terminal de l'utilisateur est appelé UE (User Equipment). Quant au réseau de cœur, il comprend trois entités principales:

- la PGW (*Packet Data Network GateWay*) est la passerelle du réseau LTE vers l'extérieur ;
- la MME (*Mobility Management Entity*) est l'entité responsable du plan de contrôle ;
- la SGW (*Serving GateWay*) est en charge du plan utilisateur et elle est plus proche de la station de base que la PGW.

La Figure 13 représente une vision épurée du réseau. En effet, nous omettons l'architecture de service lié à l'*IP Multimedia Subsystem* (IMS) (The 3rd Generation Partnership Project (3GPP), December 2012) (The 3rd Generation Partnership Project (3GPP), December 2012), les entités responsables de la facturation et de l'authentification ainsi que le HSS (*Home Subscriber Server*) qui contient les informations générales sur les utilisateurs. En plus de ces entités, la figure représente aussi les différentes interfaces entre celles-ci.

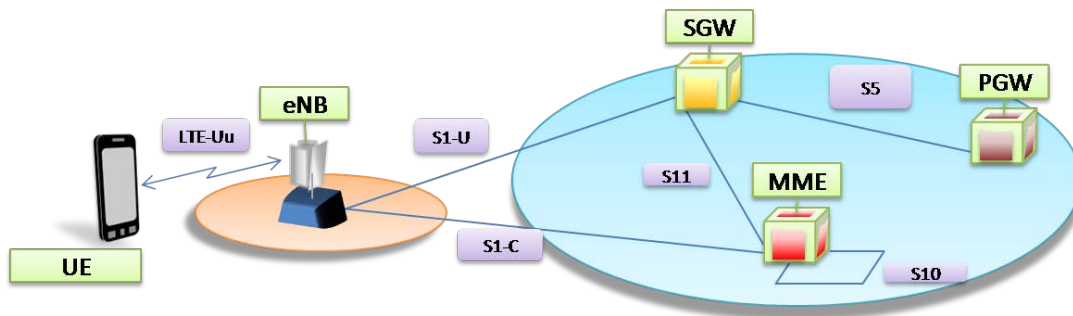


Figure 13 Architecture standard LTE

2.2.2 La pile protocolaire de LTE

Ce paragraphe décrit succinctement la pile protocolaire de LTE. La Figure 14 et la Figure 15 représentent les protocoles utilisés pour le plan de contrôle et le plan utilisateur. Les protocoles de l'interface radio LTE-Uu que sont PDCP (*Packet Data Convergence Protocol*) (The 3rd Generation Partnership Project (3GPP), September 2012), RLC (*Radio Link Control*) (The 3rd Generation Partnership Project (3GPP), September 2012), MAC (*Medium Access Control*) (The 3rd Generation Partnership Project (3GPP), September 2012) ainsi que la couche physique sont communs au plan utilisateur et au plan de contrôle. PDCP assure l'envoi ordonné des paquets de plus haut niveau ainsi que des fonctions de chiffrement et d'intégrité. Les protocoles RLC et MAC ont moins d'impact lors d'un *handover*, car ils sont remis à zéro pendant son exécution. Dans l'EPC, les données du plan utilisateur sont transmises par l'intermédiaire du protocole GTP-U (*GPRS Tunnelling Protocol for User plane*) (The 3rd Generation Partnership Project (3GPP), September 2012) au-dessus des protocoles UDP/IP.

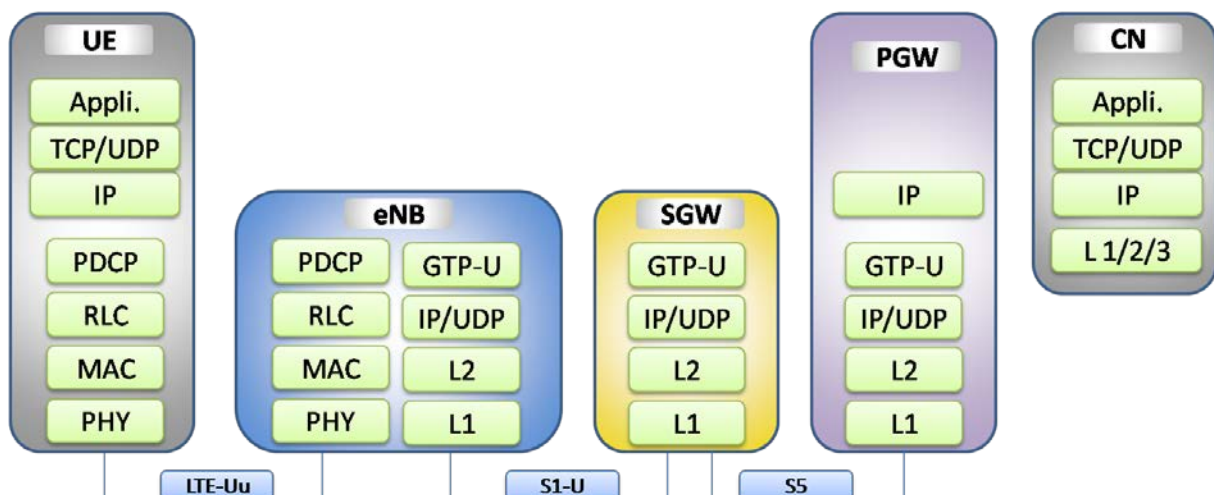


Figure 14 Piles protocolaires du plan utilisateur

Le plan de contrôle est partiellement différent. La MME en est l'entité principale. Dans le réseau de cœur, cette gestion est réalisée grâce au protocole GTP-C (*GPRS Tunnelling Protocol for Control plane*) (The 3rd Generation Partnership Project (3GPP), September 2012). L'interaction entre le

réseau de cœur et le réseau d'accès est réalisée par l'intermédiaire du protocole S1-AP (S1 Application Protocol) (The 3rd Generation Partnership Project (3GPP), September 2012). Le protocole SCTP (*Stream Control Transmission Protocol*) garantit alors la bonne réception des messages de ce protocole. S1-AP est responsable en particulier de la signalisation du *handover* et offre aussi la possibilité d'une transmission transparente des messages NAS (*Non-Access Stratum*) (The 3rd Generation Partnership Project (3GPP), September 2012). Ces messages permettent une signalisation directe entre la MME et l'UE. Ce protocole gère une partie des fonctions de mobilité de l'équipement ainsi que la gestion des sessions et de la sécurité. Sur l'interface air, LTE-Uu, le protocole principal du plan de contrôle est le protocole RRC (*Radio Resource Control*) (The 3rd Generation Partnership Project (3GPP), September 2012). Il occupe un rôle prépondérant vis-à-vis de la gestion des *handovers*.

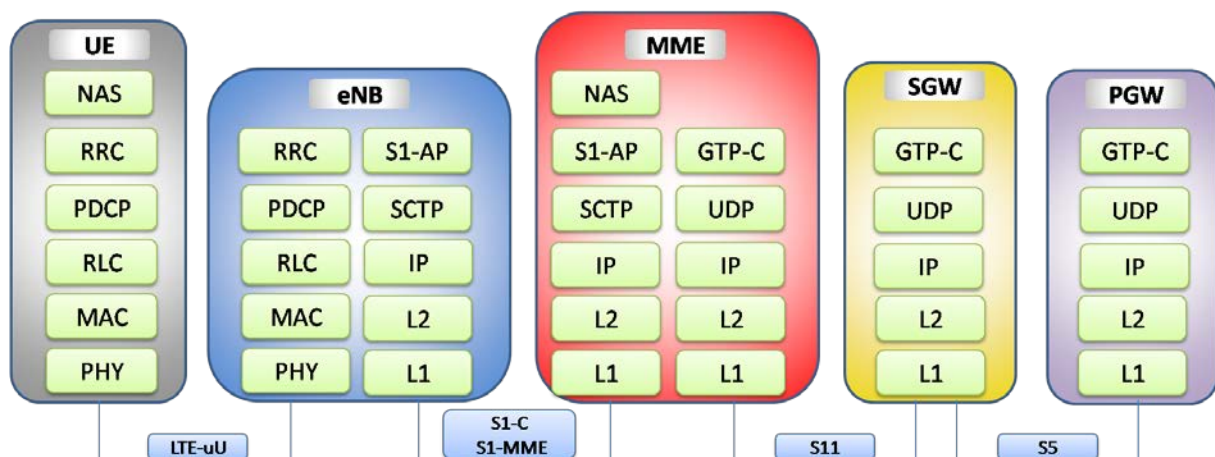


Figure 15 Piles protocolaires du plan de contrôle

2.2.3 La gestion des flux

La gestion des flux et de la qualité de service est liée au protocole GTP-U dans le réseau de cœur. Il permet une différenciation de la qualité de service grâce à la notion de canal, appelé *bearer* dans les spécifications. Deux types de *bearer* existent :

- Les *bearers* GBR (*Guaranteed Bit-Rate*) : ils offrent la possibilité de réserver des ressources. Ils ne sont donc établis par la MME qu'en fonction des besoins. Celle-ci gère l'établissement, la modification et la suppression des *bearers*.
- Les *bearers* non-GBR : ils ne garantissent aucun débit et ne réserve aucune ressource. Par conséquent, l'établissement de ces *bearers* peut être indépendant des flux qu'ils transportent ; ils ne consomment aucune ressource lorsqu'ils n'ont aucune donnée utile à transporter. Lors de la connexion du terminal au réseau de cœur, et pour chacune de ses adresses IP, un *bearer* par défaut est mis en place. Comme il est permanent, il sera nécessairement de type non-GBR.

La correspondance entre les flux et les *bearers* est effectuée aux extrémités du réseau de l'opérateur en coopération avec le réseau IMS. Chaque flux traverse la PGW ou l'UE qui, en fonction de filtres de correspondance, attribuent un *bearer* à chaque flux (Ekstrom, February 2009). Un *bearer* a la capacité de gérer plusieurs flux. Le nombre de *bearers* est limité à 8 par adresse IP à cause des numéros de *bearers* sur l'interface radio (The 3rd Generation Partnership Project (3GPP), September

2012). Le *bearer* est défini dans le réseau de cœur par un identifiant de tunnel (TEID, *Tunnel Endpoint Identifier*) dans l'en-tête du protocole GTP-U. Chaque *bearer* obtient une classe de service. Celle-ci est traduite en DSCP (*Differentiated Services Code Point*) par l'eNB et par la PGW pour permettre aux couches inférieures à GTP-U, dans le réseau de cœur, de respecter les exigences de la classe de service. Les objectifs de la qualité de service définis dans LTE se concentrent sur l'interface radio, LTE-Uu. En effet, le consortium 3GPP ne s'occupe en rien de la qualité de service dans la partie du réseau de cœur et ce pour deux raisons. Premièrement, il ne définit pas les couches basses du réseau de cœur. Deuxièmement, il émet l'hypothèse que celui-ci offre des performances de meilleure qualité que l'interface radio et, que ce lien est donc surdimensionné. De manière de plus en plus prononcée, le deuxième argument s'étiole face à l'augmentation exponentielle des ressources consommées au niveau du backhaul (Robson, July 2011). Ces liens connaissent des difficultés pour acheminer les débits demandés. Dans la définition des classes de QoS dans LTE (The 3rd Generation Partnership Project (3GPP), December 2012), les caractéristiques du réseau de cœur se voient attribuer une valeur moyenne constante. Par exemple la latence induite par le réseau de cœur est considérée comme étant de l'ordre de 20 ms.

2.2.4 La gestion de la mobilité

La gestion de la mobilité dans LTE est décomposée en deux catégories de procédures : les procédures *Idle Mode* et les procédures *Active Mode*. Ces deux modes se différencient par l'état de l'UE vis-à-vis du réseau d'accès. Ils sont respectivement définis par la connexion ou la non-connexion de l'UE à un eNB, du point de vue du protocole RRC.

2.2.4.1 *Idle Mode*

Si l'UE est en *Idle Mode*, aucune communication n'est activée. Dans ces conditions, les procédures liées à la mobilité regroupent :

- **La sélection et la re-sélection de cellule.** Le terminal utilisateur doit choisir une cellule où il réside. La cellule sélectionnée est celle auprès de laquelle l'UE doit se connecter lors de la demande d'établissement d'un service. La re-sélection intervient lors du mouvement de l'utilisateur, l'UE peut élire une autre cellule en fonction des caractéristiques du signal (The 3rd Generation Partnership Project (3GPP), September 2012).
- **Tracking Area et Paging.** Du point de vue du réseau de cœur, la localisation de l'UE n'est connue que de la MME avec la granularité des *Tracking Areas*. C'est un ensemble d'eNB où l'UE qui est en *Idle Mode* est localisé. Si l'établissement d'un service est nécessaire alors, la procédure de *paging* est exécutée. Le message est diffusé à tous les eNB appartenant à la *Tracking Area* et l'UE répondra à ce message. Cela permet de localiser la cellule où se situe l'UE et de déclencher la procédure d'activation de sa connexion au niveau du protocole RRC (The 3rd Generation Partnership Project (3GPP), December 2012).

Dans ce mode, peu d'optimisations sont nécessaires pour assurer le bon fonctionnement de ces mécanismes. Cependant, on différenciera tout de même les *Tracking Areas* du segment satellite de celles du segment terrestre. En outre, les procédures de sélection et de re-sélection de cellules privilégieront les cellules du segment terrestre.

2.2.4.2 Active Mode

Dans le mode connecté, la gestion de la mobilité est en grande partie représentée par les procédures de *handovers* que LTE propose. Les *handovers* dans le standard LTE sont de type *Break-before-Make* et sont contrôlés par le réseau, plus précisément, par l'eNB avec l'assistance de l'équipement utilisateur. En effet, c'est l'eNB qui décide l'instant de déclenchement ainsi que la cellule cible grâce aux mesures effectuées par l'UE. Le *handover* assure une procédure sans perte, grâce à un mécanisme de transfert des paquets de l'utilisateur pendant le *handover* entre l'eNB source et cible. C'est le mécanisme de *Forward*. En outre, la mise à jour du chemin du *handover*, qui correspond à la modification du chemin des données de l'ancienne à la nouvelle eNB, n'est réalisée qu'après l'exécution du *handover* ; il est nommé *Late Path Switch*. Ces procédures sont décrites dans (The 3rd Generation Partnership Project (3GPP), September 2012) et (The 3rd Generation Partnership Project (3GPP), December 2012).

Il existe de nombreux types de *handover* dans les spécifications LTE. Le premier critère de différenciation est lié à la présence ou non d'une interface directe entre les eNB source et cible. Si cette présence est avérée, la procédure se nomme *X2-handover*. Dans le cas contraire, on l'appelle *S1-handover*. Ce nom a pour origine l'interface de liaison entre les eNB. L'interface X2 est le nom de l'interface directe entre les deux eNB alors que l'interface S1 symbolise le backhaul. L'intérêt de l'interface X2 pendant le *handover* est la diminution du délai d'acheminement des messages entre l'eNB source et cible. Cependant, dans notre cas, il est rare, voire impossible, d'envisager une interface directe entre les eNB. Une technologie terrestre serait rédhibitoire, car cela irait en contradiction avec la flexibilité et la rapidité de déploiement que l'on veut obtenir. En conséquence, nous focaliserons notre étude sur le *S1-handover*.

Une autre caractéristique implique une différenciation des procédures du *handover* ; c'est la nécessité de changer de MME et de SGW pour l'UE. En effet, une MME et une SGW servent un groupe d'eNB. Si un *handover* se produit entre deux eNB qui ne dépendent pas des mêmes entités, la procédure doit tenir compte du besoin de transfert des contextes entre les MME et les SGW. Ce changement se nomme relocalisation ; elle peut s'exécuter de manière indépendante pour la SGW et la MME.

2.2.4.3 S1 handover avec et sans relocalisation de la MME et de la SSGW

La procédure du *handover* se décompose en trois phases. La première étape est la phase de préparation du *handover*. Elle débute par l'envoi de rapports de mesures de la part de l'UE. Ensuite, l'eNB source décide de l'exécution du *handover*, en fonction de ces mesures, et prépare les différentes entités du réseau à ce *handover*, en particulier l'eNB cible. La deuxième phase est l'exécution du *handover*. Elle comprend le détachement de l'UE par rapport à l'eNB source et son attachement à l'interface radio de l'eNB cible. La dernière étape est la phase de terminaison. Elle est responsable de la libération des ressources ainsi que du *Path Switch*.

2.2.4.3.1 Configuration des mesures

La décision qui déclenche l'exécution du *handover* repose sur deux mesures, RSRP (*Reference Signal Received Power*) et RSRQ (*Reference Signal Received Quality*) (The 3rd Generation Partnership Project (3GPP), September 2012). Ces deux mesures sont réalisées par l'UE. La première, RSRP donne une indication quant à la force du signal d'un eNB, alors que RSRQ représente la qualité du signal en

tenant compte des interférences. L'eNB fournit les paramètres de mesures qui peuvent déclencher le *handover*. Ainsi, si ces paramètres sont atteints, des rapports de mesures informeront l'eNB que l'événement en question s'est produit (dépassement de seuil, perte du signal...). Cette configuration est définie par l'intermédiaire du protocole RRC (The 3rd Generation Partnership Project (3GPP), September 2012). Les événements possibles sont définis dans le Tableau 5. La présence de quatre nouveaux paramètres permet d'affiner la précision et la pertinence de l'événement. Ce sont les seuils, l'hystérésis, l'*offset* et le TTT (*Time-To-Trigger*) (Tableau 4).

Paramètres	Description	Utilité
Seuils	Les seuils sont des valeurs constantes définies par l'UE.	Paramètres de la décision
Hystérésis	A la valeur du seuil qui déclenche un événement, on ajoute la valeur de l'hystérésis ; on la déduit pour la fin de l'événement.	Evite l'effet Ping-Pong
Time-To-Trigger	TTT représente la durée minimale de l'événement avant qu'un rapport de mesures ne soit envoyé.	Evite l'effet Ping-Pong
Offset cellule individuelle	Offset par cellule pour favoriser ou défavoriser des cellules.	Gestion plus précise des <i>handovers</i>

Tableau 4 Les caractéristiques des paramètres de mesure

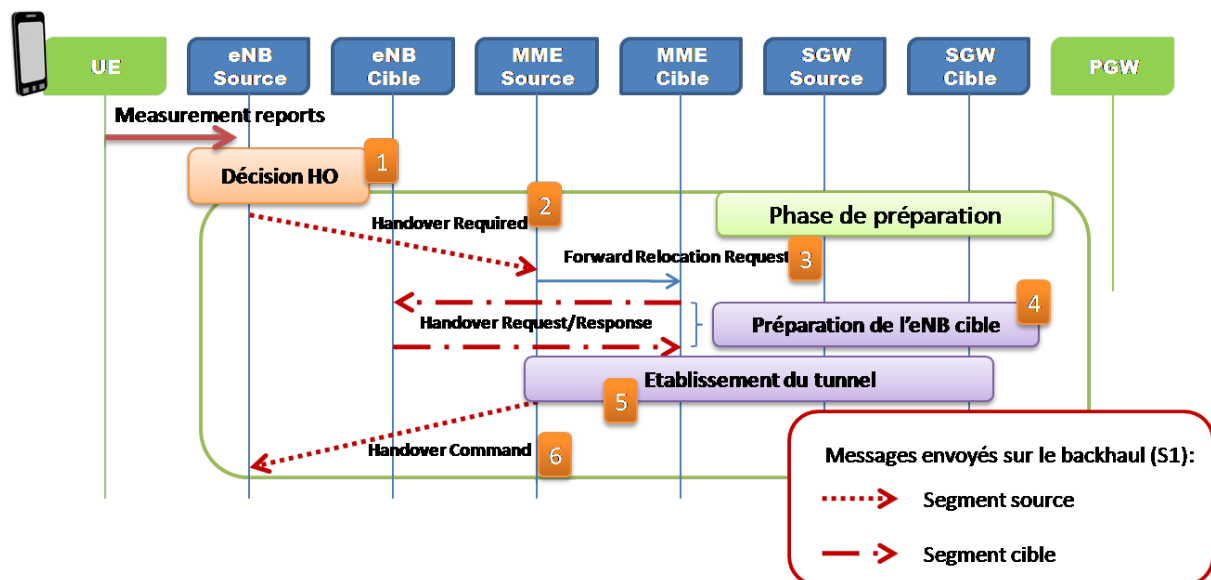
Événement	Description
A1	La mesure de la cellule courante (SCell) dépasse un certain seuil.
A2	La mesure de la cellule courante devient inférieure à un certain seuil
A3	La mesure de la cellule voisine devient meilleure que celle de la cellule courante (la cellule primaire depuis LTE-Advanced) à laquelle un <i>offset</i> additionnel est ajouté.
A4	La mesure de la cellule voisine dépasse un seuil.
A5	La mesure de la cellule courante (la cellule primaire depuis LTE-Advanced) passe sous le seuil 1 et la mesure de la cellule voisine au-dessus du seuil 2.
A6	la mesure de la cellule voisine avec un <i>offset</i> additionnel devient meilleure que la mesure de la cellule secondaire. Cet événement est apparu dans LTE

	Advanced avec l'agrégation de porteuse.
B1	Equivalent à A1 sauf que la cellule voisine est une cellule d'une autre technologie d'accès 3GPP
B2	Equivalent à A2 sauf que la cellule voisine est une cellule d'une autre technologie d'accès 3GPP

Tableau 5 Les différents événements paramétrables dans l'UE.

2.2.4.3.2 La phase de préparation

Cette étape a pour objectif de vérifier que l'eNB cible peut accepter le nouveau terminal, en réalisant du contrôle d'admission. Il prépare aussi le réseau, en particulier la SGW, pour permettre le transfert des paquets et éviter les pertes pendant la phase d'exécution, grâce à la création de tunnel de *Forward*. La Figure 16 représente les échanges de messages du plan de contrôle lors de la préparation.

Figure 16 La phase de préparation du *handover*

1. La décision du *handover* est prise par l'eNB source. Elle est fondée sur les rapports de mesures effectuées par l'UE. Ces mesures sont d'une importance capitale pour l'efficacité d'un *handover* (The 3rd Generation Partnership Project (3GPP), December 2012).
2. L'eNB source envoie le message *Handover Required* qui enclenche le début du *handover*. Il contient, entre autres, les causes du *handover* et l'identifiant de l'eNB cible.
3. Ensuite, la MME décide si une relocalisation de la MME et de la SGW est nécessaire. Si c'est le cas, la MME source transférera le contexte de l'UE à la MME cible. Sinon le message *Forward Relocation Request* ne sera pas transmis.
4. L'eNB cible reçoit les informations de l'UE grâce au message *Handover Request* et renvoie des paramètres de la nouvelle connexion destinée à l'UE. Lors de cette étape, l'eNB cible exécute le contrôle d'admission pour valider la procédure du *handover*.
5. Les tunnels du mécanisme de *Forward* sont établis pour assurer que le *handover* ne soit pas la cause de pertes de données.

6. La MME source envoie le message *Handover Command* à l'eNB source ; ce qui clôt la phase de préparation.

2.2.4.3.3 La phase d'exécution

Après la préparation du *handover*, l'eNB source va déclencher son exécution. Cette phase se sépare en deux parties : le changement d'eNB par l'UE (messages 1 à 3 de la Figure 17) et le transfert du contexte de l'UE entre les eNB par l'intermédiaire des MME (messages 4 à 6 de la Figure 17).

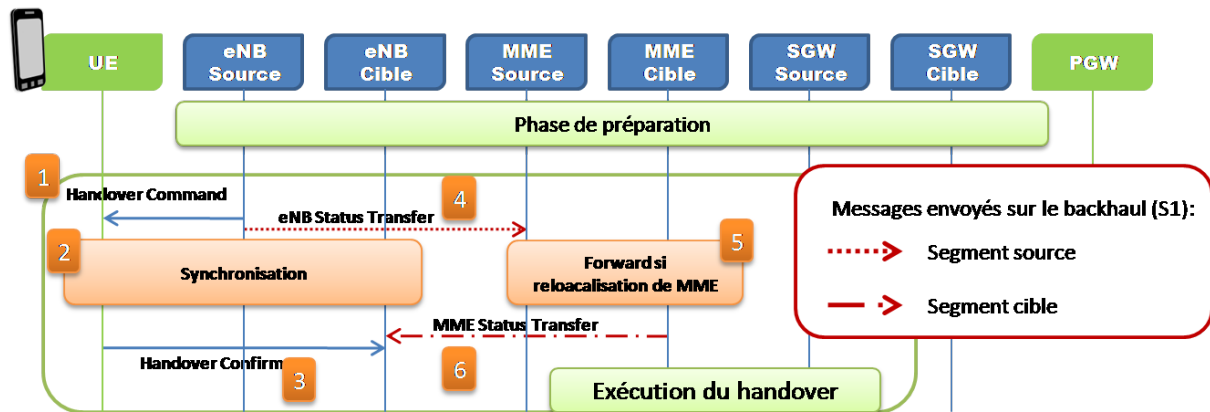


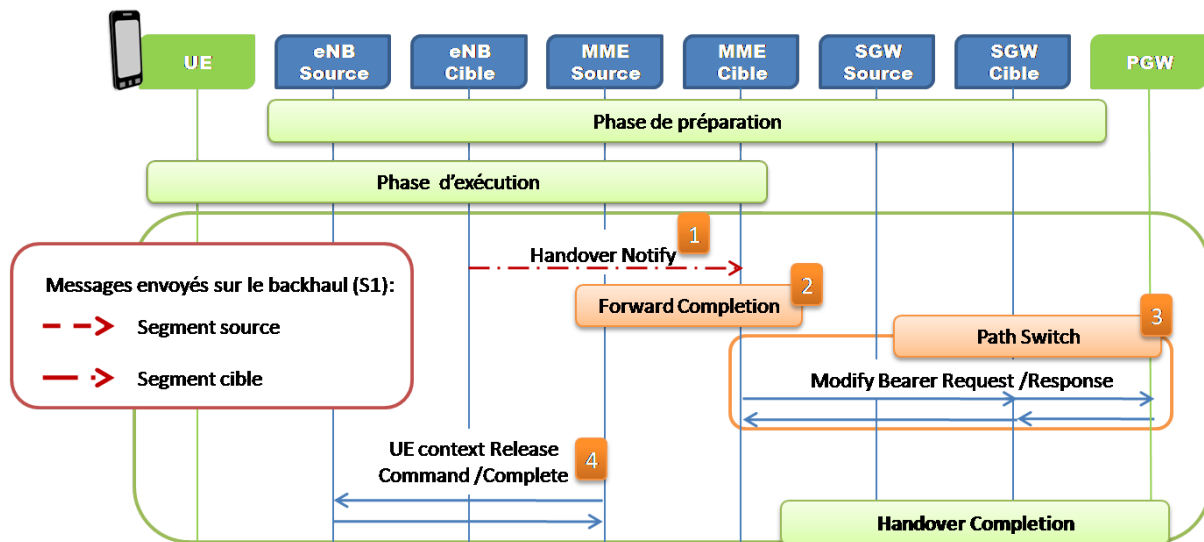
Figure 17 La phase d'exécution du *handover*

Le premier échange de messages de signalisation se déroule sur l'interface radio et le deuxième dans le réseau de cœur et sur le backhaul :

- Sur l'interface radio :
 - 1 L'eNB source envoie le message RRC correspondant au *Handover Command* pour signifier à l'UE de se détacher de cet eNB. Ce message contient des informations sur la nouvelle connexion à instaurer avec l'eNB cible.
 - 2 L'UE débute l'établissement de la connexion avec l'eNB cible en se synchronisant. L'eNB cible lui alloue des ressources sur le lien montant de l'interface radio.
 - 3 L'UE envoie le message *Handover Confirm* pour valider le bon déroulement de l'attachement à la nouvelle eNB. Ce message est aussi nommé *RRC Connection Reconfiguration Complete* par le protocole RRC.
- Dans le réseau de cœur :
 - 4 Pour réaliser un *handover* sans perte, le contexte du protocole PDCP de l'eNB source est transféré à l'eNB cible. Cela permet de ne pas remettre à zéro la couche PDCP dans l'UE et évite ainsi les pertes des données dans les files d'attente de PDCP.
 - 5 Les messages sont transférés par l'intermédiaire de la MME cible si une relocalisation est survenue.
 - 6 Le contexte est alors transmis de manière transparente à l'eNB cible.

2.2.4.3.4 La phase de terminaison

La phase de terminaison vise à mettre à jour le chemin des données pour le terminal, ainsi que de libérer les ressources dans l'eNB source.

Figure 18 La phase de terminaison du *handover*

Voici l'échange de message pendant la phase de terminaison du *handover* :

1. L'eNB cible envoie le message *Handover Notify* pour informer la MME du bon déroulement et de l'exécution du *handover*.
2. S'il y a eu relocalisation de MME, le message est transféré à l'ancienne MME (*Forward Completion*).
3. La modification du chemin des données est effectuée par le *Path Switch*. La SGW va modifier le routage des paquets à destination de l'UE. Ils seront routés, non plus vers l'eNB source, mais vers l'eNB cible. S'il y a eu relocalisation de la SGW alors le chemin entre la SGW et PGW est aussi modifié.
4. À l'expiration d'une temporisation, la MME source libèrera les ressources liées à l'UE dans l'eNB source.

2.2.4.3.5 Le mécanisme de *Forward*

Le mécanisme de *Forward* appartient au plan utilisateur. Il est influencé par les trois phases du plan de contrôle. Pendant la phase de préparation, des tunnels sont établis. Ce sont ces tunnels qui assureront le mécanisme de *Forward* de l'eNB source vers l'eNB cible. Pendant la phase d'exécution du *handover*, le *Forward* est activé. Pendant la phase de terminaison, le *Path Switch* représente la fin de ce mécanisme.

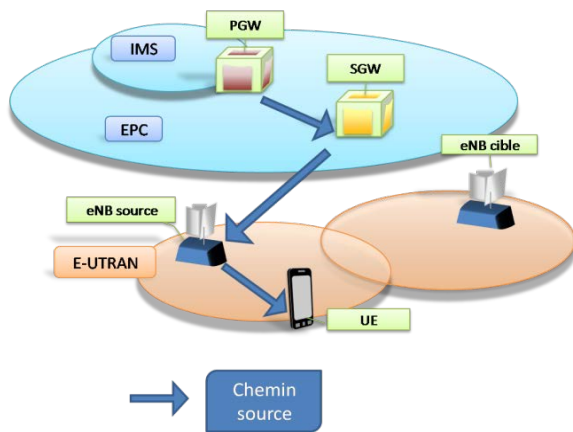
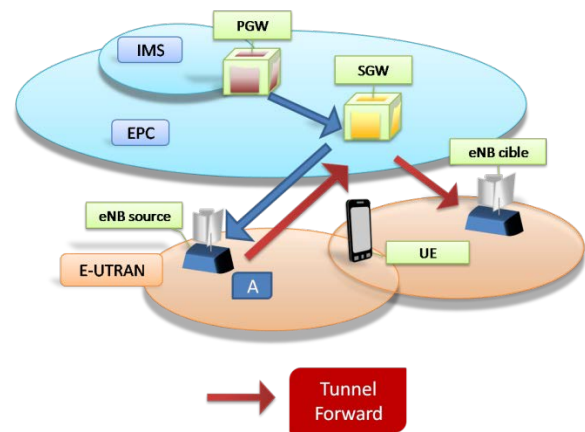
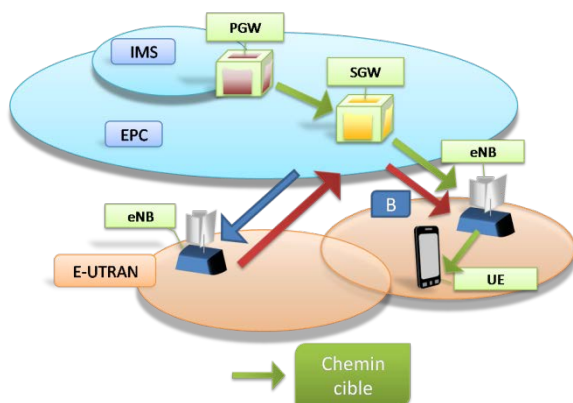
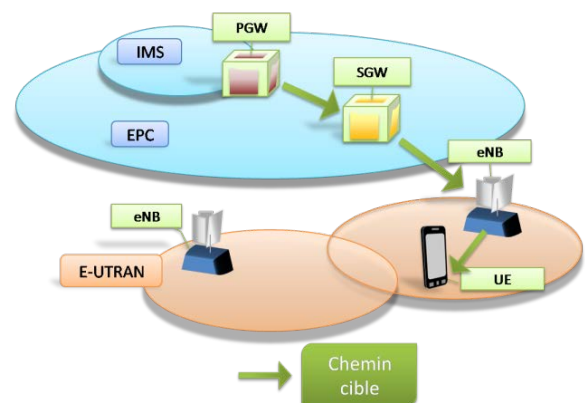


Figure 19 Le chemin source des données

Figure 20 Le chemin de *Forward*

Dans la Figure 19, les chemins des données avant et pendant la phase de préparation sont représentés. Les tunnels de *Forward* ainsi que le chemin cible sont construits, mais ils restent cependant inactifs et ne sont pas encore utilisés.

Dès le commencement de la phase de l'exécution du *handover*, les tunnels de *Forward*, représentés en rouge dans la Figure 20, sont activés. L'eNB source transfère les paquets vers l'eNB cible, en passant par le backhaul et par l'entité SGW (A dans la Figure 20). L'eNB cible conserve les paquets de données dans une file d'attente jusqu'à ce que le *handover* soit exécuté.

Figure 21 Le chemin des données après le *Path Switch*Figure 22 Fin du *Forward*

Pendant la phase de terminaison, le *Path Switch* est effectué. Par conséquent, le chemin cible est activé, alors que le chemin source désactivé. Cependant, durant un court laps de temps, des paquets se trouvent toujours en transit dans les tunnels de *Forward*, alors que de nouveaux paquets arrivent à l'eNB cible par le chemin cible (B dans la Figure 21). Ainsi, l'eNB cible doit les réordonner. Les données empruntant le chemin cible doivent rester en attente jusqu'à ce que tous les paquets provenant des tunnels de *Forward* soient transmis sur l'interface radio (Figure 22).

2.3 Les solutions de réseaux LTE hybrides avec un backhaul satellite

Dans un souci de convergence, le projet SAE a défini une pluralité de moyens de raccordement des réseaux d'accès en fonction de leur nature (LTE, technologie 3GPP, technologie 3GPP2, etc.). Ce large panel nous offre de nombreuses possibilités d'architecture pour intégrer un réseau d'accès LTE avec un backhaul satellite. Un des critères importants pour le choix de l'architecture réside dans les possibilités d'optimisation qui ne se traduiront pas par un changement du niveau d'impact. Les différents niveaux d'optimisations sont décrits dans le Tableau 6, alors que le Tableau 7 définit la classification des niveaux d'impact. L'architecture choisie répondra à ces deux critères. En fonction du niveau de modification requis pour optimiser les mécanismes du *handover*, celle qui induit le minimum d'impact sur les spécifications sera préconisée, puisque la modification des spécifications implique des coûts importants.

Niveau d'optimisation	Description
Niveau 0	Aucune optimisation n'est possible
Niveau 1	Adaptation des paramètres des spécifications
Niveau 2	Modification des mécanismes sur le segment satellite
Niveau 3	Modification des mécanismes inter-segments sur les deux segments

Tableau 6 Description des différents niveaux du critère d'optimisation

Niveau de l'impact	Description
Niveau 0	Aucun impact
Niveau 1	Impact sur la signalisation sur le segment satellite
Niveau 2	Impact sur la signalisation sur les interfaces inter-segments
Niveau 3	Impact sur les spécifications du segment terrestre

Tableau 7 Description des différents niveaux du critère d'impact sur la signalisation

2.3.1 Réseau LTE hybride traditionnel

L'utilisation de l'architecture traditionnelle du réseau LTE est la solution la plus simple (Figure 23). Cette simplicité induit une limitation quant aux possibilités d'optimisations des *handovers*. Deux variantes d'architectures sont possibles. La première se résume à intégrer le réseau d'accès au réseau de cœur terrestre, au travers d'un backhaul satellite. La MME et la SGW qui sont associées aux eNB satellite peuvent aussi servir pour des eNB terrestres. Certains mécanismes ne peuvent donc pas être optimisés et différenciés de ceux du backhaul terrestre. Dans ce cas, des modifications de la signalisation ne sont pas autorisées, car elles en induiraient sur la spécification du système LTE, y compris sur le segment terrestre. Par conséquent, les possibilités d'optimisation ne dépassent pas le niveau 1 ; en revanche l'impact est nul sur la spécification.

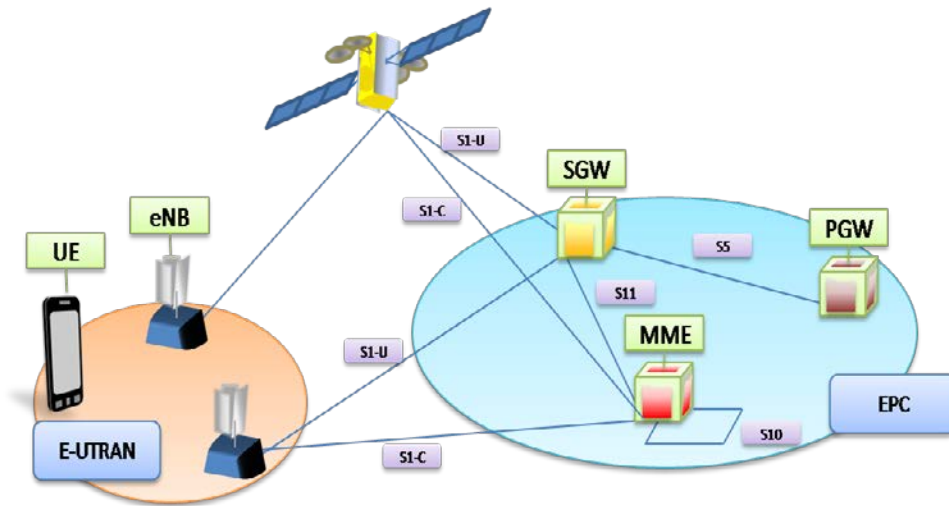


Figure 23 Architecture hybride LTE traditionnelle

La deuxième variante consiste à dédier une partie du réseau de cœur au segment satellite. Ainsi, une MME et une SGW sont réservées aux eNB reliés par un backhaul satellite. Cette architecture offre une plus grande flexibilité d'optimisation, car les entités et les interfaces directement en relation avec le backhaul satellite sont isolées du segment terrestre. Ainsi, les protocoles et les mécanismes du segment satellite peuvent être optimisés. Le niveau d'optimisation atteint est le niveau 2 avec une modification de la signalisation de niveau 1. En outre, cette architecture procure l'avantage d'intégrer les entités dédiées au segment satellite au plus près de la passerelle satellite. Elle offre donc la possibilité de co-localiser ces équipements.

2.3.2 Architecture avec un autre réseau d'accès 3GPP

Pour assurer la compatibilité avec les autres standards du 3GPP, les spécifications définissent de nouvelles interfaces, entre les entités du réseau de cœur de l'EPC et celles des anciennes générations. Les *handovers* sont admis entre les différentes technologies d'accès et ont été dénommés *handovers* inter-RAT (*inter Radio Acces Network*). La signalisation de ces *handovers* est adaptée pour chaque technologie d'accès. L'intérêt de considérer le segment satellite comme un autre réseau du 3GPP provient de la création de nouvelles interfaces spécifiques entre le segment terrestre et le segment satellite (Figure 24). Dans ce cas, le segment satellite se compose d'eNB, auxquels s'ajoutent une ou plusieurs MME ainsi qu'une ou plusieurs SGW qui lui sont spécifiquement dédiées. En conséquence, cette architecture possède les mêmes avantages que l'architecture traditionnelle avec des entités dédiées. Les nouvelles interfaces entre le segment terrestre et satellite et, en particulier, l'interface entre une MME du segment satellite et une MME du segment terrestre, que nous nommerons S10-Sat, permettent de modifier la signalisation des *handovers* inter-segments (Niveau 3), tout en restreignant l'impact aux mécanismes inter-segments (Niveau 2). De manière évidente, la définition de ces interfaces sera très similaire à celle de l'interface de l'architecture traditionnelle et seuls certains mécanismes seront légèrement affectés. Ainsi, l'interface S5-SAT entre la SGW-Sat et la PGW ne sera pas du tout modifiée.

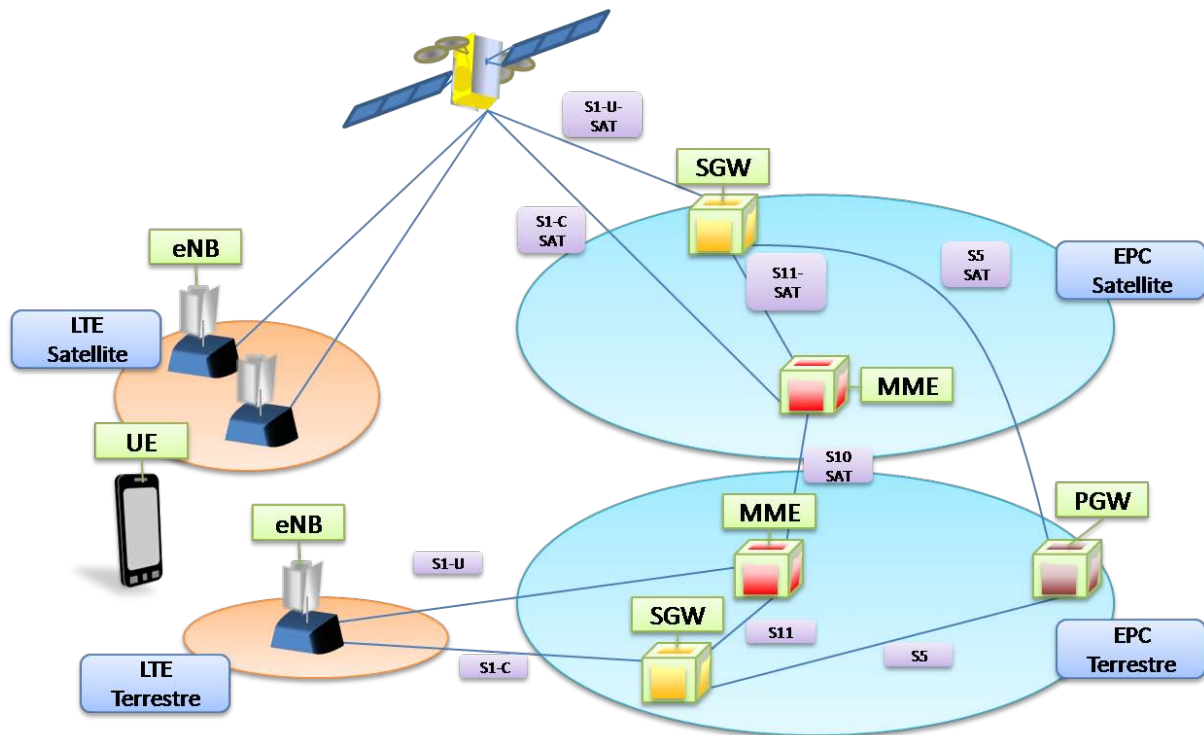


Figure 24 Architecture avec un autre réseau 3GPP

2.3.3 Architecture avec un Home eNB

Les Home-eNB définis par le 3GPP sont des stations de base simplifiées. Leur backhaul est mis en œuvre par l'intermédiaire d'un raccordement large bande de type fibre et ADSL. Ils possèdent une couverture peu étendue, égale à quelques dizaines de mètres. Deux raisons sont avancées pour le déploiement de ces *Home-eNB* (HeNB). La première découle de la limitation de la couverture dans les zones urbaines. En effet, un Home-eNB permet d'améliorer le signal radio à l'intérieur des bâtiments où la force du signal des macro-eNB reste limitée. Deuxièmement, la multiplication de ces *Home-eNB* augmente le nombre de cellules et, dans le même temps, la capacité totale du réseau d'accès. Ainsi, l'*Home-eNB* apporte une solution pour compenser la rareté du spectre fréquentiel dans le déploiement de LTE.

Cette architecture est intéressante, puisque les procédures pour le HeNB prennent en compte le fait que le backhaul n'est pas réalisé par une liaison dédiée au réseau mobile. En outre, dans cette architecture, les *Home-eNB* sont gérés dans le réseau de cœur par une unique entité, appelée Home-eNB GateWay (HeNB-GW) (Figure 12Figure 25). Cette solution permet d'isoler le segment satellite sans affecter les entités du réseau de cœur, MME et SGW, tout en autorisant la modification de la signalisation dans le segment satellite (niveau 2). Un autre aspect intéressant des Home-eNB réside dans des mécanismes d'auto-configuration, définis pour obtenir des fonctionnalités de *Plug-and-Play*, qui peuvent être réutilisés dans notre scénario. Cependant, les scénarios d'utilisation de cette architecture sont peu nombreux, en raison de sa faible couverture à la fois géographique et en termes d'utilisateurs. De surcroît, les modifications pour permettre une utilisation optimisée de cette architecture devront être financièrement à la charge de l'utilisateur et non de l'opérateur, puisque l'HeNB est localisé chez lui.

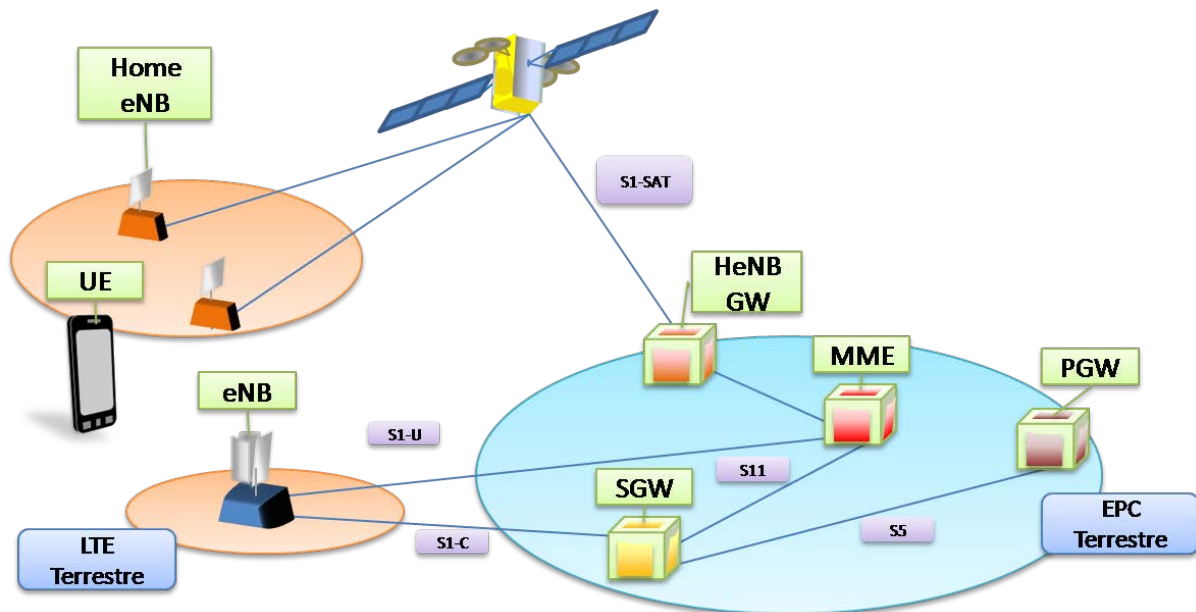


Figure 25 Architecture avec un Home eNB

2.3.4 Architecture avec un réseau extérieur

Le 3GPP a aussi développé des interfaces pour connecter des réseaux extérieurs au réseau de cœur LTE. Deux méthodes sont spécifiées. Premièrement, si la sécurité du réseau extérieur est considérée comme suffisante et comparable à celle assurée par LTE, ce réseau est directement raccordé par l'intermédiaire du PGW. Ce sera par exemple le cas pour un réseau WiMAX mobile. Si ce n'est pas le cas, ce réseau est connecté au PGW grâce à une entité supplémentaire qui assure le niveau de sécurité nécessaire (The 3rd Generation Partnership Project (3GPP), September 2012). Dans ce cas de figure, la gestion de la mobilité est sommaire et les spécifications ne prévoient que des *handovers* non-optimisés. Elle ne repose sur aucune technologie particulière et sur l'hypothèse que l'UE est multimode. La marge de modification se révèle donc infime.

2.3.5 Les choix d'architectures

Le Tableau 8 récapitule les différentes caractéristiques des architectures présentées. L'architecture traditionnelle, avec des entités qui ne sont pas dédiées au segment satellite, ainsi que l'architecture dédiée au raccordement des réseaux non-3GPP n'offrent que trop peu de possibilités d'optimisation. La solution, qui consiste à considérer les eNB satellite comme des Home-eNB, est évincée en raison de la faible capacité et de la faible portée de ces stations de base. Par conséquent, deux choix sont retenus, l'architecture traditionnelle LTE avec une MME et une SGW dédiées au segment satellite, et l'architecture où le segment satellite est considéré comme un autre réseau d'accès 3GPP. La première permet d'optimiser les procédures sur le segment satellite. Tant que ce type d'optimisation suffit, l'architecture traditionnelle sera adoptée. Dans le cas contraire, on s'orientera vers l'architecture inter-RAT.

Architectures	Optimisation	Impact	Commentaires
Traditionnelle	1	0	Optimisation limitée
Entités non dédiées			
Traditionnelle	2	1	Vraie séparation des segments satellite et terrestre dans l'EPC
Entités dédiées			
inter-RAT 3GPP	3	2	Création de nouvelles interfaces inter-segments
Home eNB	2	1	Faible capacité et couverture
Réseaux extérieurs	0	0	Aucune optimisation liée à la mobilité

Tableau 8 Récapitulatif des architectures potentielles

2.4 Analyses du standard LTE

Le délai induit par le backhaul satellite est la caractéristique la plus néfaste pour les procédures de *handover*. En effet, la principale contrainte liée à ce mécanisme réside dans la rapidité de sa réalisation. Son exécution doit être terminée avant que le signal en provenance de l'eNB source ne devienne trop faible. Dans LTE, cette procédure repose sur des liens de backhaul à haut-débit et à faible latence. Le lien satellite va à l'encontre de ces deux caractéristiques et met en péril le bon déroulement du *handover*. Ainsi, ce paragraphe met en relief deux dangers causés par le backhaul satellite : que sont le ralentissement de la phase de préparation du *handover* et l'inefficacité du mécanisme de *Forward* traditionnel.

2.4.1 Ralentissement de la phase de préparation du *handover*

Le premier effet néfaste du délai est causé par l'envoi de messages de signalisation sur le lien satellite ; ce qui a pour conséquence de retarder la phase de préparation du *handover*. Ainsi, suivant le type de *handover*, cette étape conduit à un ou deux allers-retours sur le lien satellite (Tableau 9), alors que la phase de préparation dure en moyenne 50ms pour un *handover* terrestre (Nokia, Nokia-Siemens Networks, May 2009).

<i>Handover</i>	Message envoyé sur le lien satellite	Délai ajouté
HO intra-satellite	Handover Required Handover Request Handover Response Handover Command	4 . Délai_{sat} ~1200 ms
HO inter-satellite-terrestre	Handover Required Handover Command	2 . Délai_{sat} ~600 ms
HO inter-terrestre-satellite	Handover Request Handover Response	2 . Délai_{sat} ~600 ms

Tableau 9 Impact du délai satellite sur la signalisation de la phase de préparation

Les paramètres des mesures permettent de compenser une partie du ralentissement de la phase de préparation. Par exemple, le TTT, qui définit la durée pendant laquelle la condition de l'événement doit être valide pour déclencher un rapport de mesures à l'eNB, est très important. Il permet d'éviter

des rapports et des *handovers* intempestifs. Une valeur faible du TTT pourrait ainsi compenser une partie du délai de la préparation. Cependant, cela augmenterait dans le même cas les risques de débiter des *handovers* inutilement.

Par conséquent, ce ralentissement de la phase de préparation induit une augmentation des risques d'échec du *handover*. En effet, la probabilité, que l'UE perde le signal avant que la phase de préparation soit terminée, devient beaucoup plus importante. La perte de ce signal prend le nom de *Radio Link Failure* (RLF) dans LTE. Les spécifications définissent différents types de RLF (Tableau 10). Ce tableau souligne les différences entre ces RLF et, en particulier, il précise que le *Too-Late For Measure* RLF (type 1) est la perte de connexion qui engendre l'effet néfaste le plus important. L'interruption due à ce type de RLF est allongée de 200 ms par rapport aux autres (NTT DoCoMo, Inc., March 2009).

RLF	Description	Commentaire
Type 1 Too-Late For Measure	<i>Radio Link Failure</i> survient avant le déclenchement du <i>handover</i> . L'interface radio tombe avant que le rapport de mesures ne puisse être envoyé (Figure 26).	Ce type de RLF est le plus néfaste, car aucune préparation ne va être réalisée. Ainsi le rétablissement de la connexion avec une nouvelle eNB est plus long.
Type 2 Too-Late For Command	<i>Radio Link Failure</i> avant le message <i>Handover Command</i> (RRC Connection Reconfiguration). L'interface radio tombe pendant la préparation du <i>handover</i> . Le <i>handover</i> ne peut être exécuté, car l'UE ne reçoit jamais la commande qui déclenche la phase d'exécution (Figure 27).	Le temps du rétablissement est diminué, car les probabilités que l'UE choisisse l'eNB cible déjà préparé comme eNB de récupération sont importantes.
Type 3 Too-Early	<i>Handover Failure</i> causé par l'échec de la synchronisation avec l'eNB cible. Ainsi, la phase d'exécution ne peut être menée à son terme (Figure 28).	Dans ce cas-là, l'UE va réaliser une tentative de rétablissement de la connexion avec l'eNB source.

Tableau 10 Description des différentes *Radio Link Failure*

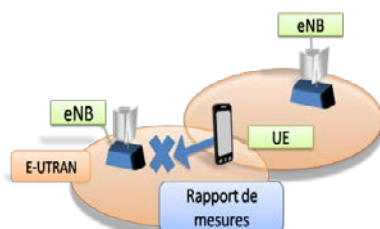


Figure 26 *Too-Late For Measure* RLF (type 1).

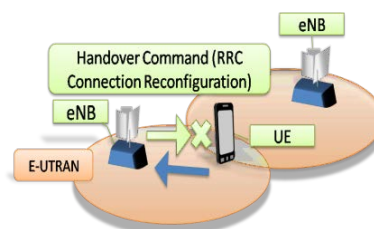


Figure 27 *Too-Late For Command* RLF (type 2).

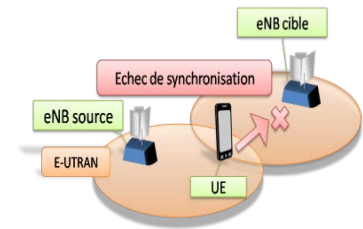


Figure 28 *Too-Early* RLF (type 3).

Les échecs de *handover* sont reconnus comme une nuisance en particulier lors de fortes mobilités et dans des environnements urbains denses (Ericsson, February 2009). Cet aspect est exacerbé par des paramètres de *handover* de plus en plus agressifs, en particulier le TTT, qui devient de plus en plus faible jusqu'à atteindre 0 ms. En effet, une faible valeur de TTT augmente de manière significative le nombre et la fréquence des *handovers*. Ce paramètre permet de diminuer les occurrences des *Too-Late For Measure* RLF, mais augmente les probabilités de subir des *Too-Late For Command* RLF et des *Too-Early* RLF.

2.4.2 Le mécanisme de *Forward* et d'ordonnancement inefficace

Le mécanisme de *Forward*, réalisé entre les eNB et la SGW pendant la phase de préparation, est inefficace. En effet, ces tunnels sont établis au niveau du backhaul dans le but de transférer les paquets de l'eNB source à l'eNB cible. Ce lien est assuré par un système satellite. Par conséquent, au lieu de subir une fois le délai d'une transmission satellite, les paquets utilisant les tunnels souffrent de trois émissions sur le lien satellite dans le cadre d'un *handover* intra-satellite et de deux pour les *handovers* inter-segments (Figure 29).

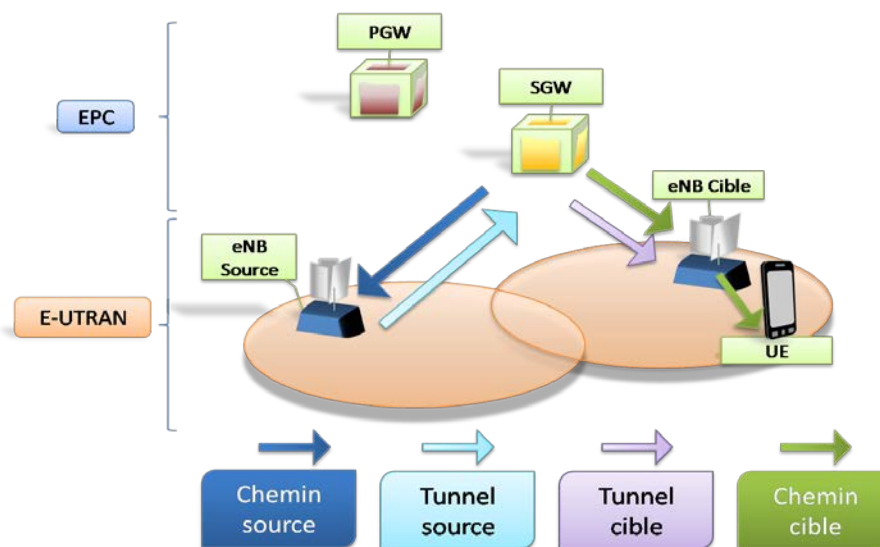


Figure 29 Différenciation entre le tunnel source et cible.

<i>Handover</i>	Tunnel satellite		Délai (ms) induit par le satellite		
	Source	cible	Avant	Pendant	Après
Inter-satellite-terrestre	✓	✗	300	600	0
Inter-terrestre-satellite	✗	✓	0	300	300
Intra-satellite	✓	✓	300	900	300

Tableau 11 Conséquences du backhaul satellite sur le mécanisme de *Forward*

Tout d'abord, ce délai est néfaste pour les protocoles de transport et applicatifs. En effet, les applications de transport de la voix et de la vidéo sont sensibles au délai de la transmission. C'est par ailleurs, la gigue qui engendre une baisse importante de la qualité de l'application. Les *buffers* dans ces applications peuvent compenser cette gigue. Mais dans notre cas, le *handover* cause une variation très importante du délai de 300ms à 600ms (Tableau 11). Cette fluctuation immédiate de la gigue entraîne des pertes et une diminution de la qualité de l'application. Pour les flux qui reposent sur le protocole de transport TCP, cette variation du délai a aussi des conséquences néfastes. Les

mécanismes implantés dans TCP dépendent en grande partie de l'estimation du RTT (Round Trip Time). Le changement brusque de cette valeur ne peut être pris en compte immédiatement ; ce qui peut causer des congestions ou bien encore des détections de congestion erronées.

En outre, ce mécanisme de *Forward* implique une surconsommation des ressources au niveau du lien satellite, dans le cadre de deux *handovers*. En effet, si le segment source est un segment satellite, le même paquet subit un aller-retour sur le lien satellite. Ce phénomène est acceptable dans le réseau terrestre, car le lien de backhaul permet d'écouler des débits très élevés, et le mécanisme de *Forward* est utilisé pendant une durée réduite. Ces deux caractéristiques ne sont plus vraies avec le segment satellite et, par conséquent, cette surconsommation a un impact plus important, qui affecte les performances des applications de l'UE qui subit le *handover* mais aussi ceux qui sont actifs à cet instant dans le segment satellite.

Le mécanisme d'ordonnancement est lui aussi affecté par le backhaul satellite. Il se fonde sur le transfert du contexte PDCP lors de l'exécution du *handover*. Ce contexte est transmis de l'eNB source à l'eNB cible par l'intermédiaire de la SGW et donc du backhaul satellite (2 dans la Figure 30). De la même manière que pour la phase de préparation, ce transfert est ralenti par le lien satellite, avec un aller et retour pour le *handover* intra-segment satellite et un aller simple pour le *handover* du segment satellite vers le terrestre. Cependant, la latence induite par le lien satellite transforme les données contenues dans le contexte en informations obsolètes. En effet, la phase d'exécution du *handover* est très rapide d'un point de vue de l'UE, de quelques dizaines de millisecondes au maximum. C'est à la fin de cette étape que l'eNB cible doit configurer sa couche PDCP (1 dans la Figure 30). L'eNB cible n'ayant pas encore obtenu le contexte PDCP, elle remet à zéro la couche PDCP et l'ordonnancement devient impossible.

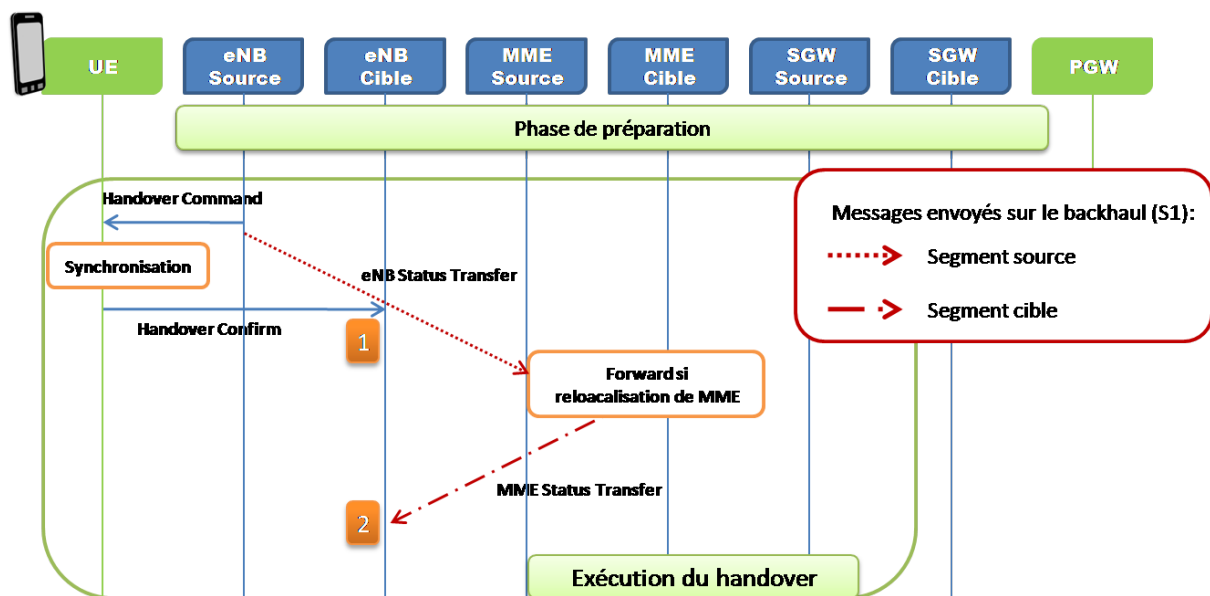


Figure 30 Transfert du contexte PDCP dans le cadre d'un *handover* intra-satellite

Ce chapitre met en relief les défauts des spécifications du *handover* LTE intégrant un backhaul satellite. Deux axes d'améliorations sont ainsi tracés, avec l'optimisation de la phase de préparation pour prendre en compte le délai induit par le satellite et la définition d'un nouveau mécanisme de *Forward* tenant aussi en compte le reséquencement.

Le chapitre suivant proposera des optimisations de la phase de préparation pour contrer son ralentissement par le lien satellite.

Chapitre 3 : Réseaux hybrides LTE-satellite : Optimisation de la phase de préparation

Dans ce chapitre, nous nous intéressons à l'optimisation de la préparation du *handover*, en vue de réduire ou d'évincer les problèmes dus au ralentissement de la phase de préparation mis la lumière dans le chapitre précédent. Dans cette optique, les axes potentiels d'étude sont :

- **L'élimination de la phase de préparation** ; cette option engendrerait de nombreuses modifications de la signalisation liée au *handover*. En effet, elle entraînerait une refonte complète de la procédure du *handover*. Par conséquent, elle ne serait être utilisable dans le cadre de *handover* inter-segment, car elle modifierait trop profondément le standard dans les entités du segment terrestre.
- **La minimisation de la durée de la préparation** ; en raison de l'absence d'interface directe entre eNB (interface X2) et du délai incompressible induit par le satellite, il nous semble impossible de réduire de manière significative la durée de cette étape.
- **L'optimisation de la décision** ; cette solution consiste à adapter et à modifier les paramètres de décision, pour que la phase de préparation puisse prendre en considération le délai dû au satellite. Elle est envisageable puisque les paramètres des mesures sont configurés par l'eNB qui est responsable de l'exécution de l'algorithme de décision.
- **L'amélioration de la récupération d'un *Radio Link Failure*** ; si les solutions réalisables, pour supprimer ou réduire les conséquences du ralentissement de la phase de préparation, sont limitées, le *handover* subira des RLF de manière plus fréquente. Pour contrer cela, les mécanismes de récupération de la perte de connexion doivent être optimisés et adaptés aux backhaul satellite.

Par la suite, nous nous concentrons sur l'optimisation de la décision du *handover* et sur l'amélioration des mécanismes de récupération après l'occurrence d'un RLF. En effet, seules ces deux solutions permettent de minimiser l'impact du ralentissement de la préparation, puisque l'élimination de cette étape et la minimisation de sa durée semble impossible, à moins de modifier en profondeur les spécifications du *handover*.

3.1 Etat de l'art

Les problèmes induits par le ralentissement de la préparation du *handover*, dans le cadre d'un réseau cellulaire avec un backhaul satellite, n'ont jamais fait l'objet d'étude à notre connaissance. Ceci s'explique par le fait, que les scénarios subissant ce ralentissement étaient inexistantes. En effet, comme nous l'avons vu dans le premier chapitre, les eNB déployés dans une même zone possédaient le même backhaul satellite. Ainsi, les *handovers* ne subissaient aucunement les problèmes liés au lien satellite. De nos jours, les réseaux cellulaires avec un backhaul satellite conservent cette caractéristique et ces systèmes connaissent une évolution de la 2G vers la 3G, pour lesquels plusieurs

liens satellite ne sont pas nécessaires. C'est seulement lors d'un futur passage à LTE que cette nécessité se fera ressentir.

De manière générale, l'exécution d'un *handover* dans un système hybride, où le lien satellite a un impact est rare. Seuls certains systèmes cellulaires intégrés gèrent les *handovers*. Cependant la gestion de la mobilité est plus faiblement affectée par le satellite. En effet, la taille des cellules du segment satellite s'étend sur plusieurs centaines de kilomètres et, par conséquent, la fréquence des *handovers* est largement inférieure à celle des systèmes terrestres. Dans ces systèmes, le lien satellite se situe au niveau de l'interface radio et ce sont les messages *Measurement Report* et *Handover Command* qui seraient retardés par la transmission satellite. Le standard GMR-1 3G adapte les mesures du terminal au lien satellite et ne modifie pas la procédure du *handover* dans le réseau de cœur (European Telecommunications Standards Institute (ETSI), July 2009). De nouvelles propositions de procédure tirent parti de la caractéristique bi-mode du terminal, comme dans le projet SINUS (Efthymiou, Hu, Sheriff, & Properzi, 1998) et proposent une signalisation éloignée du standard LTE. Le changement complet de procédure du *handover*, ainsi que l'utilisation de la caractéristique bi-mode du terminal est impossible dans notre système. En effet, seul le *hard handover* est implanté dans LTE, ce qui nous oblige à conserver ce type de procédure. En outre, l'UE ne possède qu'une interface radio LTE et ne peut être considéré comme bi-mode. Par conséquent, nous nous focalisons sur des optimisations proposées dans le cadre du standard LTE par le monde terrestre, qui peuvent être utilisées dans un réseau possédant un backhaul satellite. Ces solutions ont pour principal objectif de réduire les échecs dus à une mauvaise gestion des *handovers* et de réduire leurs conséquences néfastes.

3.1.1 Techniques d'optimisation de la gestion de la mobilité

Ces techniques font partie des mécanismes SON. Elles augmentent la robustesse des procédures de la gestion de la mobilité dans LTE et sont appelées MRO (*Mobility Robustness Optimization*). Leur but est de minimiser les échecs liés aux *handovers*. Elles récupèrent des informations sur les causes de ces échecs par l'intermédiaire des eNB et des UE, puis modifient les paramètres du *handover* pour y remédier. Par exemple, si un RLF survient, les eNB peuvent recevoir des informations sur la cause de l'échec grâce à une variable (*VarRLF-report* (The 3rd Generation Partnership Project (3GPP), September 2012)) envoyé par l'UE. Celui-ci indique si le *handover* a été déclenché de manière hâtive ou tardive, ou si la cellule cible est inappropriée. Des algorithmes interprètent ces rapports en termes de nouveaux paramètres. Ces algorithmes ne sont pas spécifiés par le standard et sont laissés libres d'implantation (The 3rd Generation Partnership Project (3GPP), December 2012). Cependant, le standard définit les messages de signalisation entre les différentes entités pour ces techniques SON, ainsi que trois types d'architectures SON :

- **Les architectures centralisées** ; une entité centrale récupère les informations d'un ensemble d'eNB et reconfigure les paramètres de ceux-ci.
- **Les architectures distribuées** ; chaque eNB exécute l'algorithme d'optimisation sans entité coordinatrice. Seuls des échanges d'informations directes entre eNB sont autorisés.
- **Les architectures hybrides** ; il y a à la fois des algorithmes SON dans chaque eNB et dans une entité coordinatrice.

Les algorithmes MRO ont la possibilité de régler plusieurs paramètres du *handover* : les valeurs de l'hystérésis, du TTT, des différents *offsets* et des seuils (The 3rd Generation Partnership Project (3GPP), September 2012). De nombreuses études se focalisent sur l'adaptation du TTT et de l'hystérésis. Ainsi, dans le projet de la commission européenne SOCRATES (Van Den Berg, et al., February 2008), l'algorithme analyse les différents événements néfastes (*handover* déclenché trop tôt ou trop tard, effet ping-pong...) et ajuste en conséquence le TTT et l'hystérésis. Les bénéfices de cet algorithme semblent limités, en particulier par le fait, que les deux paramètres ajustés sont dépendants de l'événement. Ils ne peuvent donc pas être spécifiques à un couple d'eNB (Bălan, Sas, Jansen, Spaey, & Demeester, 2011) (Jansen, Balan, Stefanski, Moerman, & Kurner, May 2011). Par conséquent, un échec entre deux eNB va influencer toutes les procédures, ayant pour source le même eNB. Pour remédier à ce problème, (Catovic, Garavaglia, & Patil, December 2009) et (Komine, Yamamoto, & Konishi, July 2012) propose d'utiliser les *offsets* spécifiques à deux eNB ; ce qui permet d'adapter précisément les mesures et les événements du *handover* à ce couple.

La convergence de ces algorithmes peut poser problème. Certains définissent un pas très grand qui, par exemple, peut être égal à 1000 *handovers* dans (Legg, Hui, & Johansson, September 2010). Cet aspect est important si le scénario d'utilisation nécessite le déploiement d'eNB temporaires. Bien que des propositions d'algorithme possèdent une convergence beaucoup plus rapide (Hui & Legg, August 2012), la configuration de départ doit être cohérente et prendre en compte le backhaul satellite, en particulier pour les *handovers* inter-segments.

L'adaptation automatique des paramètres de *handover* ne résout pas le problème du ralentissement de la préparation du *handover*. Il permet de réduire ses effets néfastes en minimisant les échecs liés à cette procédure, en ajustant les différents paramètres. Les paramètres par défaut doivent être choisis de manière cohérente avec l'environnement de déploiement de l'eNB, pour réduire le temps de convergence des algorithmes. Cependant, ces techniques ne sont pas suffisantes pour prendre en compte le délai du satellite.

3.1.2 Préparation multiple

La préparation multiple de station de base est autorisée dans les spécifications du 3GPP (Sesia, Toufik, & Baker, 2011). Elle permet de configurer plusieurs eNB pendant la même phase de préparation dans le but d'améliorer la récupération d'un échec (Figure 31). En effet, si le *handover* ne peut s'effectuer avec la cellule cible et que l'UE est obligé de se raccorder à un autre eNB non-préparé, la phase de reconnexion entraînera une longue période d'interruption de service. La préparation multiple maximise les probabilités que le nouvel eNB sélectionné possède le contexte de l'UE et ne passe pas par cette période d'interruption (Nortel, October 2007).

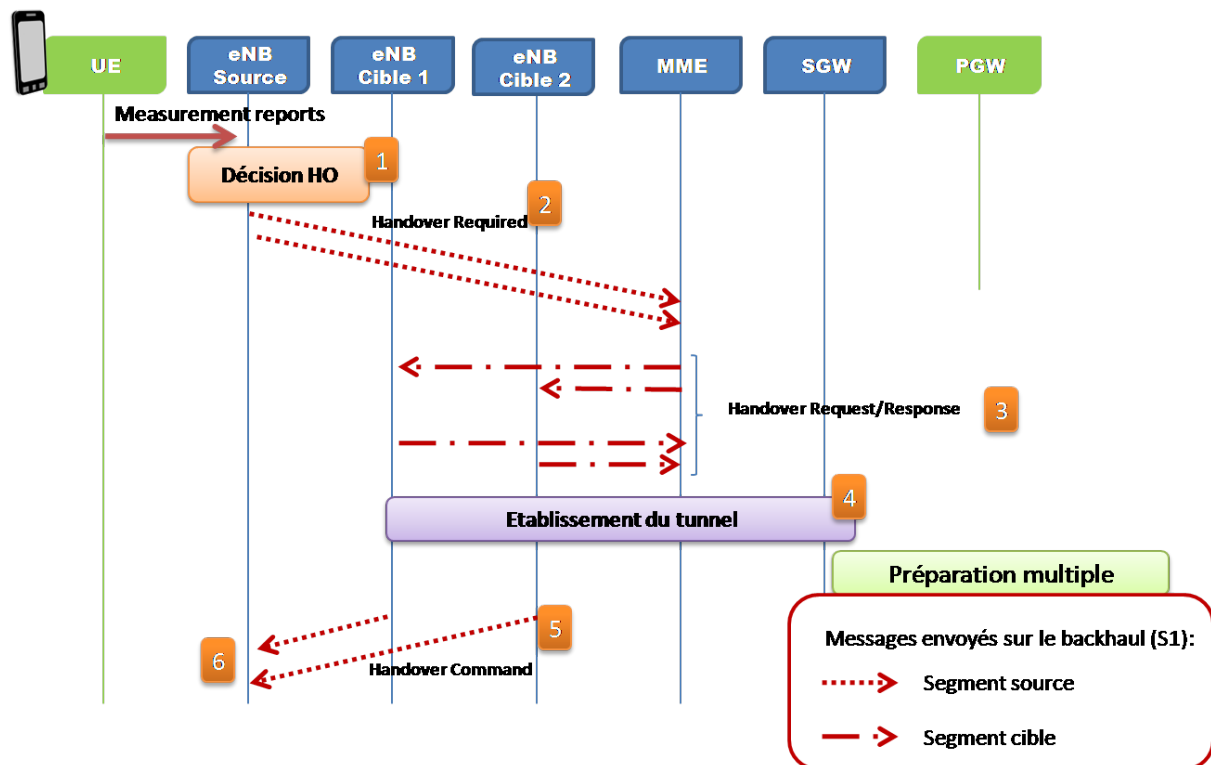


Figure 31 Préparation multiple

En plus des bénéfices lors d'un échec de *handover*, cette procédure permet de prendre partiellement en compte le ralentissement de la phase de préparation. Puisque cette phase dure longtemps, les mesures peuvent être très variables, ainsi l'eNB choisi au début de la phase de préparation (eNB cible 1 dans la Figure 31) peut être destitué de sa première place vis-à-vis de la force du signal. Grâce à la préparation multiple, l'eNB source peut modifier le choix de l'eNB cible juste avant l'envoi du message *Handover Command*. Ainsi, dans l'étape 6 de la Figure 31, l'eNB source peut choisir l'eNB cible 2 pour le *handover*.

Bien que cette procédure améliore de manière significative la phase de préparation du *handover*, elle ne permet pas de résoudre entièrement le problème du ralentissement. En effet, le choix du *handover* doit être effectué à la réception du *Handover Command* (étape 6 de la Figure 31) et la probabilité que le lien radio entre l'eNB source et l'UE soit déjà rompu avant cette étape est grande.

3.1.3 Forward handover

Le *Forward handover* est une proposition pour éviter la période d'interruption de service due à la sélection d'un eNB non-préparé lors d'un RLF (Panasonic, August 2006). Cette procédure autorise cet eNB, où l'UE se raccorde, à récupérer le contexte de l'UE par l'intermédiaire d'une nouvelle signalisation (Figure 32).

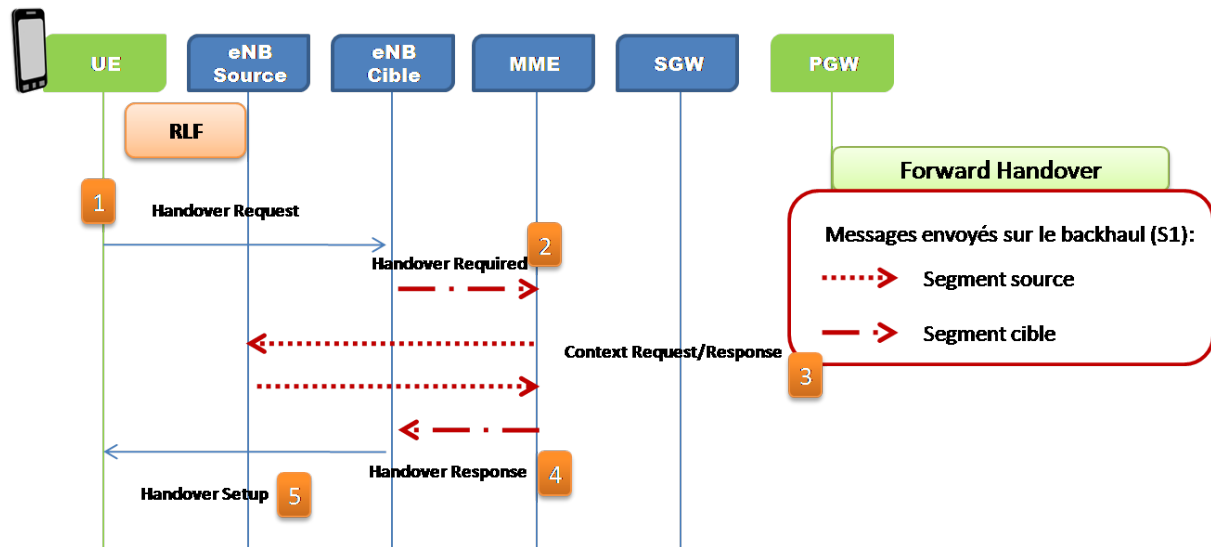


Figure 32 Forward Handover

Cependant, cette proposition n'est pas adoptée par le 3GPP (Ericsson, October 2006) et ne présente aucun avantage, car la récupération du contexte est obligatoirement ralentie par les transmissions satellite. Ceci entraîne aussi une longue période d'interruption de service.

3.2 Les optimisations

3.2.1 Optimisation de paramètres

Les différentes procédures et techniques utilisables en accord avec la spécification du 3GPP reposent toujours sur une configuration adaptée des paramètres de mesure. En effet, la prise en compte du délai de la phase de préparation est primordiale pour ne pas subir de nombreux RLF et pour assurer une convergence rapide des algorithmes SON.

3.2.1.1 Différenciation des handovers

Tout d'abord, on différencie les paramètres de mesure entre les trois types de *handovers* (inter-satellite-terrestre, intra-satellite, inter-terrestre-satellite). Cette distinction est importante, car la durée de la phase de préparation varie selon les segments impliqués. Pour la réaliser, la spécification du protocole RRC (The 3rd Generation Partnership Project (3GPP), September 2012) nous contraint d'utiliser deux événements différents. Cependant, l'UE n'a pas la capacité de différencier un eNB terrestre d'un eNB satellite. Par conséquent, les eNB terrestres et satellite peuvent répondre aux deux événements. Il nous reste deux possibilités pour éviter qu'un *handover* vers un eNB terrestre ne soit la conséquence d'un événement destiné à un *handover* vers un eNB satellite.

La première solution consiste à déléguer cette responsabilité à l'eNB source. L'UE envoie les rapports de mesures pour tous les événements même s'ils sont destinés à des eNB cible d'une autre nature. C'est l'eNB source qui va décider si l'eNB cible qui a déclenché ce rapport était du type voulu. Si c'est le cas, l'eNB source donnera suite à la procédure, sinon il ignorera ce rapport. En effet, l'eNB source connaît la nature du segment de l'eNB cible par l'intermédiaire du code du *Tracking Area* et il peut ainsi déclencher le *handover* en cas de segment cible satellite. Cette solution est relativement simple,

cependant elle induit une modification de l'algorithme de décision, mais aussi une augmentation inutile du nombre de rapports de mesures au niveau de l'interface radio.

La deuxième solution permet d'éviter cela, en utilisant la liste noire des cellules liée à un événement (*blackCellsToAddModList* dans (The 3rd Generation Partnership Project (3GPP), September 2012)). L'événement pour déclencher un *handover* vers le segment satellite se voit attribuer, par l'eNB source, une liste noire contenant toutes les cellules terrestres voisines. De la même manière, on isole les *handovers* entre des cellules d'un même eNB satellite puisque la préparation du *handover* n'est pas affectée par le satellite.

3.2.1.2 Choix des paramètres de mesure

Pour le *handover* intra-satellite, les eNB source et cible possèdent les mêmes contraintes dues au backhaul. C'est dans ce cas que la phase de préparation subit le plus important ralentissement. Le choix de l'événement A3 permet une comparaison de la qualité du signal entre eNB satellite. Dans ce cadre-là, on conserve l'*offset* général, qui représente la différence de la puissance du signal nécessaire pour déclencher le *handover*. L'adaptation se concentre sur une diminution du TTT pour anticiper le *handover*. La valeur de l'hystérésis reste similaire à un *handover* normal pour éviter l'effet ping-pong.

Le *handover* terrestre-satellite permet de réaliser la continuité de la connexion, lorsque l'eNB terrestre ne permet plus d'assurer le service. Cependant, ce *handover* implique une augmentation du délai subi par les applications de l'UE. Par conséquent, ce type de *handover* ne doit être réalisé qu'à condition qu'aucun eNB terrestre ne puisse être utilisé par l'UE. Ainsi, l'événement A5 est choisi pour déclencher ce *handover* (la cellule source devient inférieure à un seuil et la cellule cible supérieure à un autre seuil). Ce choix va permettre de défavoriser ce type de *handover* par rapport aux *handovers* terrestres. Par conséquent, les effets ping-pong sont très limités ; le TTT peut être grandement réduit pour diminuer l'effet du ralentissement de la préparation.

Au contraire, le *handover* satellite-terrestre doit être encouragé pour fournir une meilleure qualité de service à l'utilisateur. L'événement A4 est donc préféré, pour que le *handover* se produise dès que le niveau du signal de l'eNB terrestre permet un service acceptable. Pour cet événement, les valeurs de TTT et de l'hystérésis ne sont pas très significatives, car celui-ci n'est pas sujet à l'effet ping-pong. En outre, on crée un autre événement dans le cas où la force du signal de l'eNB terrestre est faible, mais la force du signal de l'eNB satellite (source) est inférieure à celle-ci. On utilise l'événement A3 avec un faible TTT et avec une hystérésis assez faible pour compenser le ralentissement de la préparation, car l'effet ping-pong est encore négligeable.

	Evt	TTT	Hyst	Seuil _{source}	Seuil _{cible}	Offset _{A3}
<i>Handover</i> inter-terrestre-satellite	A5	Faible	Faible	Seuil terrestre	Seuil satellite	×
<i>Handover</i> intra-satellite	A3	Faible	Normal	×	×	Normal
<i>Handover</i> Inter-satellite-terrestre(1)	A4	Faible	Faible	Seuil terrestre	×	×
<i>Handover</i> inter-satellite-terrestre(2)	A3	Faible	Normal	×	×	Normal

Tableau 12 Événement et paramètres des *handovers*

3.2.1.3 Limitations

La marge de manœuvre obtenue par la modification du TTT est relativement limitée comme cela est représenté dans le Tableau 13. De surcroît, plus la vitesse de l'UE augmente plus cette marge est réduite. Les *handovers* inter-segment peuvent jouer sur les seuils spécifiques au satellite pour anticiper le *handover*. Cependant, les possibilités d'optimisation qui reposent sur l'événement A3 sont plus réduites. En effet, les *offsets* de comparaison ont des plages de valeurs très faibles, dans la majorité des cas entre 0 et 3 dB (Tableau 13). En outre, ce sont les *handovers* les plus sensibles à l'effet ping-pong ; ce qui limite l'utilisation du TTT et de l'hystérésis comme moyen de compensation. Par conséquent, pour tirer le meilleur parti de ces paramètres, il est nécessaire d'utiliser les techniques SON similaires à celle décrite dans (Catovic, Garavaglia, & Patil, December 2009). Cette dernière permet d'ajuster les paramètres liés à l'événement (TTT, hystérésis, seuils, $offset_{A3}$) en cas d'échecs de *handover* répétés sur différents eNB cibles, et de modifier ceux qui sont dépendants d'un eNB voisin (CIO, *Cell Individual Offset*) en cas d'échecs de *handover* spécifiques à celui-ci.

Vitesse de l'UE :	3 km/h	50 km/h	120 km/h	250 km/h
TTT (ms)	320-1280	0-640	0-320	0-320
TTT (ms) (zones très urbaines)	0-320	0-320	0-320	0-320
Offset (dB)	0-3	1-3	2-4.5	2-4.5
Offset (dB) (zones très urbaines)	0-3	0-3	2-4.5	2-4.5

Tableau 13 Paramètres préconisés par le 3GPP dans (3GPP TSG-RAN WG1, March 2009)

L'optimisation des paramètres permet d'anticiper la préparation, cependant sa marge de manœuvre reste limitée et ne permet pas de compenser totalement le ralentissement de la phase de préparation. Ainsi, les probabilités de rencontrer des *Too-Late for Command* RLF restent importantes. En outre, la modification des paramètres, en particulier du TTT, augmente les occurrences de *Too-Early* RLF.

3.2.2 Préparation anticipée à double décision

3.2.2.1 Principe

Pour parfaire la phase de préparation, on crée un algorithme de décision à deux étapes. L'intérêt de cette procédure est d'anticiper plus efficacement la préparation de l'eNB cible, dans un premier temps, puis de rejouer un algorithme qui déclenchera la phase d'exécution si l'eNB cible est toujours le plus intéressant. Cette deuxième étape permet de prendre en compte les modifications des mesures qui sont survenues lors de la longue phase de préparation. Ainsi, la première étape permet une préparation de *handover* la plus tôt possible pour éviter des *Too-Late for Measure* RLF, alors que la deuxième va réduire le nombre de *Too-Late for Command* RLF et de *Too-Early* RLF.

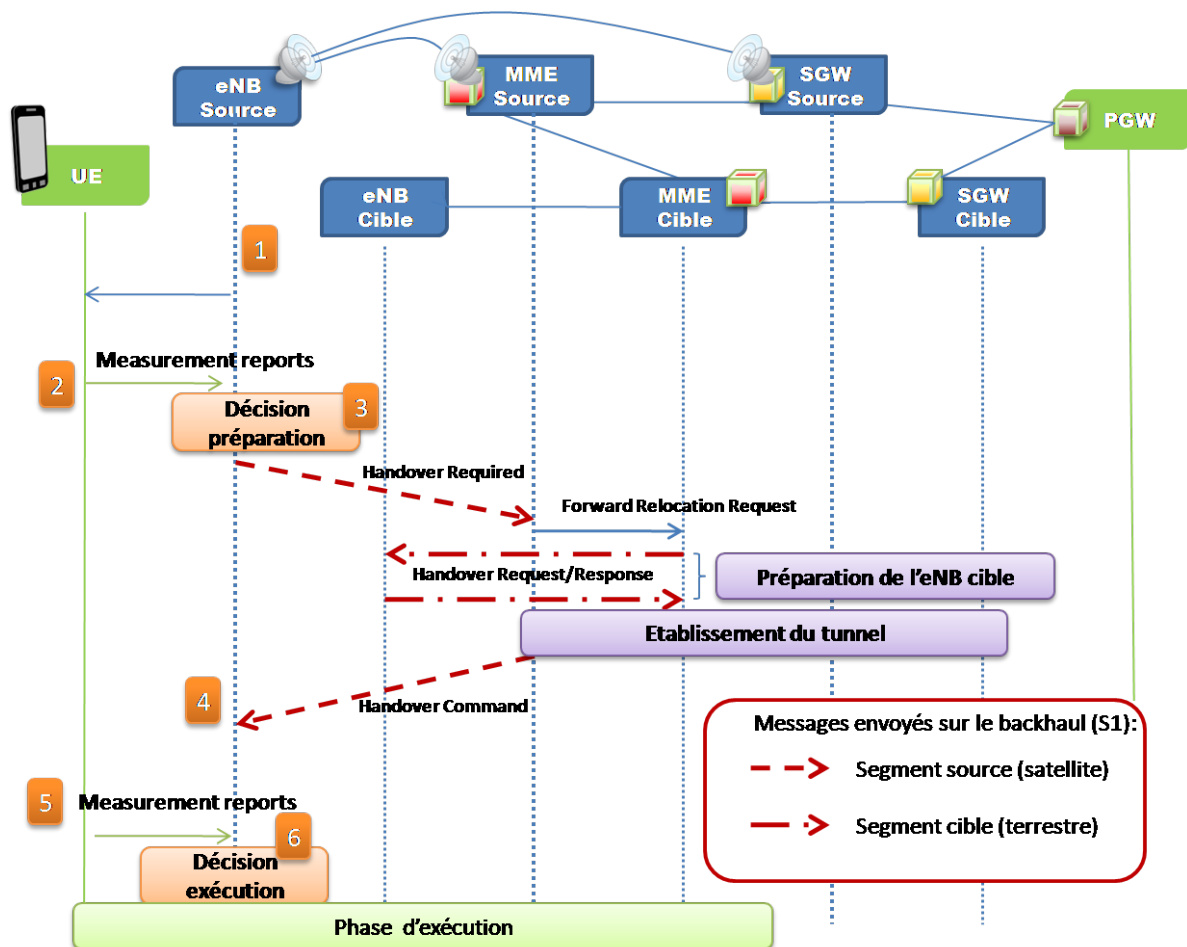


Figure 33 Préparation anticipée à double décision

La préparation anticipée à double décision est représentée pour le *handover* inter-satellite-terrestre. Cependant, elle est aussi valable pour le *handover* intra-satellite. Le déroulement de cette nouvelle préparation est décrit ci-après :

- 1 On configure les paramètres de deux événements. Le premier permet de déclencher de manière anticipée (2 sur la Figure 33) avec de plus faibles exigences en termes de qualité et de force du signal. Les valeurs des paramètres seront plus contraignantes pour le deuxième, car il signifiera le début immédiat de l'exécution du *handover* (5 sur la Figure 33).
- 2 Cette étape représente le déclenchement du premier événement qui va entraîner le début de la préparation du *handover*. Elle est anticipée par rapport à la procédure normale du standard LTE.
- 3 De la même manière que dans la procédure standard, cet algorithme déclenche la préparation du *handover* par l'eNB source, qui va établir les tunnels de *Forward* et créer le contexte de l'UE dans l'eNB cible.
- 4 L'eNB source reçoit le message *Handover Command*. A ce moment-là, il n'ordonne pas directement l'exécution du *handover* à l'UE par l'intermédiaire du message *RRC Connection Reconfiguration*, mais temporise cet ordre jusqu'à la réception du rapport de mesures du deuxième événement (étape 5).

- 5 Lorsque les valeurs des mesures de l'UE atteignent les exigences du deuxième événement défini dans l'étape 1, un rapport est transmis à l'eNB source.
- 6 C'est seulement à cet instant que l'eNB source va déclencher la phase d'exécution en envoyant le message *RRC Connection Reconfiguration* à l'UE.

3.2.2.2 Définition des événements

Cette nouvelle procédure induit des modifications au niveau de l'implantation de l'algorithme de décision dans l'eNB source. Dans le but de minimiser l'impact de cette procédure sur les entités du segment terrestre, la préparation anticipée à double décision est réservée au *handover* intra-satellite et inter-satellite-terrestre. En effet, pour ces derniers, l'eNB source fait partie du segment satellite. En outre, on ne se préoccupe que des événements A3. En effet, le *handover* inter-satellite-terrestre qui repose sur l'événement A4 (la cellule terrestre émet un signal dont la force surpasse un seuil) n'est pas contraint par le temps puisque la force du signal de l'eNB source est toujours acceptable.

De la même manière que pour l'optimisation des paramètres, l'événement qui déclenche la phase de préparation se voit attribuer un TTT faible. Cependant, il est possible de réduire les contraintes sur l'hystérésis, car la phase de préparation n'implique pas d'effet ping-pong. De la même manière, on anticipe cette phase en abaissant l'exigence de l'*offset* de l'événement A3. Pour compenser, le ralentissement du *handover* intra-satellite, qui est deux fois plus grand que pour le *handover* inter-satellite-terrestre, l'*offset*_{A3} a la possibilité d'être négatif (Figure 33). Quant à l'événement déclenchant la phase d'exécution, les paramètres sont conformes aux valeurs d'un *handover* standard (3GPP TSG-RAN WG1, March 2009).

<i>Handover</i> :	Intra-satellite		Inter-satellite-terrestre	
	Préparation	Exécution	Préparation	Exécution
TTTT	0	Normal	0	Normal
Hystérésis	Faible	Normal	Faible	Normal
Offset _{A3}	-2, 0	Normal	0	Normal

Tableau 14 Définition des événements de la préparation anticipée à double décision

3.2.2.3 Concaténation des multiples préparations.

La préparation anticipée du *handover* a l'inconvénient d'augmenter de manière significative le nombre de préparations de *handover*. Cette procédure implique une augmentation des rapports de mesures puisque les exigences de l'événement qui déclenche la préparation sont faibles. Cela a un impact limité sur l'interface radio, car dans notre scénario, seul le backhaul satellite limite les performances proposées à l'UE.

Pour rendre la procédure pleinement efficace, il faut combiner la préparation anticipée à la préparation multiple. En effet, l'intérêt de la double décision est de retarder la décision de l'exécution à la fin de la préparation, pour choisir le meilleur eNB cible au dernier moment. Pour cela, il faut avoir le choix entre plusieurs eNB déjà préparés. Une préparation de *handover* vers un eNB cible ne doit pas être bloquante pour la préparation d'un autre eNB. Si cela se produisait, cela entraînerait une augmentation des *Too-Late for Command* RLF. Dans la Figure 34, par exemple, l'eNB 1 déclenche en premier la phase de préparation. Lorsque la force du signal de l'eNB 2 dépasse celle de l'eNB source. Aucune phase de préparation n'est entamée pour l'eNB 2, puisque la force du signal de l'eNB 1 est supérieure à celle de l'eNB 2 (étape 2 dans la Figure 34). Dans l'étape 3, le signal de l'eNB 2 devient supérieur à celui de l'eNB 1, cependant aucun rapport de mesures n'est pas envoyé à

l'eNB. À la fin de la préparation de l'eNB 1, l'eNB source n'a pas le choix entre plusieurs eNB cibles et le seul qui est préparé n'est pas le plus approprié ; d'où la nécessité de la préparation multiple.

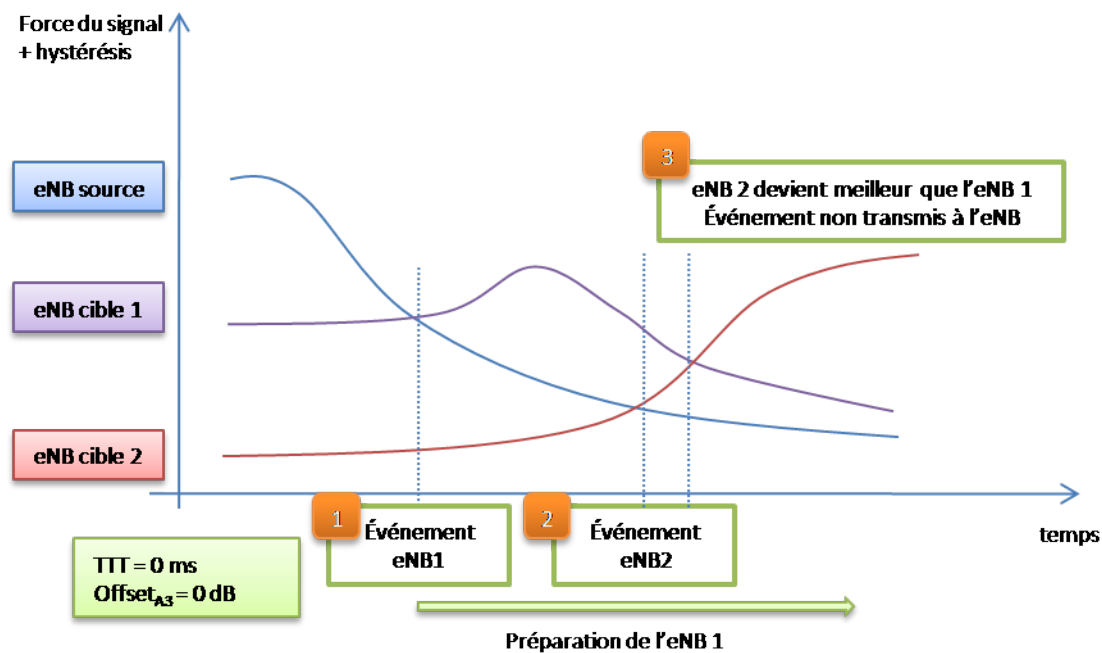


Figure 34 Événement pendant une préparation anticipée sans préparation multiple

La préparation anticipée qui est combinée avec la préparation multiple entraîne une augmentation de la signalisation sur le lien satellite. Contrairement à l'augmentation des rapports de mesures, celle-ci est néfaste pour les performances du système. Par conséquent, on propose de concaténer les messages *Handover Required* des préparations multiples. En effet, seule l'identité de l'eNB cible est modifiée. Les réponses des eNB cibles (*Handover Command*) peuvent difficilement être concaténées puisque les champs communs sont beaucoup plus réduits et la synchronisation de l'envoi de ces messages sur le lien satellite est difficilement réalisable.

3.3 Conclusion

Pour résoudre le problème du ralentissement de la phase de préparation, nous avons proposé deux optimisations. La première n'utilise qu'une configuration spécifique des événements déclenchant le *handover*, en prenant en compte le backhaul satellite. Cependant, la marge de manœuvre de cette optimisation reste restreinte. Cette constatation a encouragé la deuxième proposition. Celle-ci modifie la définition des événements mais aussi l'algorithme de décision à l'intérieur de l'eNB source. Cette amélioration permet de résorber les effets du backhaul satellite en compensant par une augmentation de la signalisation. Par conséquent, cette optimisation nécessite une gestion appropriée de la libération des ressources. Le tableau ci-après récapitule les différentes procédures de la phase de préparation abordées dans ce chapitre.

Procédure	Respect du standard	Niveau de la modification	RLF-1	Probabilité RLF-2	RLF-3	Commentaire
Standard	✓		Normale	Très élevée	Très élevée	
Préparation multiple	✓	Algorithme de décision	Faible	Élevée	Élevée	La récupération des <i>Too-Late-Command</i> RLF et Too-Early RLF sont plus efficaces que pour le <i>handover</i> standard
Optimisation des paramètres	✓	Configuration des événements	Faible	Élevée	Élevée	Les probabilités de <i>Too-Late-Command</i> RLF et Too-Early RLF restent élevées, car les mesures ont potentiellement évolué pendant la longue phase de préparation
Préparation anticipée	✓	Algorithme de décision et configuration des événements	Faible	Élevée	Normale	Si l'on ne peut pas préparer plusieurs eNB alors le risque de <i>Too-Late-Command</i> RLF reste important
Préparation anticipée et multiple	✓	Algorithme de décision et configuration des événements	Faible	Normale	Normale	Augmentation de la signalisation sur le lien satellite
Préparation anticipée Et multiple avec concaténation	✓(*)	Algorithme de décision et configuration des événements	Faible	Normale	Normale	Réduit la signalisation sur le lien satellite (*) le respect du standard dépend de l'implantation de la concaténation.

Tableau 15 Récapitulatif des différentes procédures pour la phase de préparation

La préparation à double décision combinée avec la préparation multiple est la solution la plus adaptée pour les handovers faisant intervenir le segment satellite. Elle permet, en effet, de contenir l'augmentation des RLF, dus au délai induit par le lien satellite, à des niveaux similaires aux *handovers* terrestres. Si aucune modification n'est possible, pour des raisons économiques par exemple, l'unique solution restante est l'adaptation des paramètres ; elle va permettre de prendre en compte les caractéristiques du lien de backhaul dans la phase de préparation.

Le chapitre 4 s'intéresse aux optimisations du mécanisme de *Forward*, qui permet d'assurer un *handover* sans perte ni duplication.

Chapitre 4 : Réseaux hybrides LTE-satellite : optimisation du mécanisme de *Forward*

Dans ce chapitre, nous proposons des améliorations du mécanisme de *Forward* lors d'un *handover*. En effet, dans le chapitre 2, une analyse de celui-ci a permis de mettre en relief son inefficacité lorsqu'une station de base possède un backhaul satellite. Elle est la conséquence de l'utilisation de tunnels de *Forward* entre l'eNB source et cible, qui fait varier de manière abrupte la latence des paquets de donnée pendant le *handover*. Cette augmentation de la gigue a des effets néfastes sur les applications reposant sur le protocole UDP ainsi que sur les mécanismes de TCP (Tableau 9 dans 2.4.1).

De plus, l'analyse soulève un autre problème : la surconsommation des ressources satellite, lorsque l'eNB source utilise un backhaul satellite. Le tunnel de *Forward* source est établi entre l'eNB source et la SGW. Ainsi, les paquets de donnée subissent un aller-retour qui semble inutile sur le lien satellite et qui est illustré dans la Figure 29 du chapitre 2.

4.1 Différenciation des *handovers*

Les trois types de *handover* ne subissent pas de manière équivalente les désagréments dus au backhaul satellite. Les optimisations du mécanisme de *Forward* doivent tenir compte de leurs différentes caractéristiques.

Handover inter-satellite-terrestre

Durant ce type de *handovers*, le trafic de l'UE souffre de la forte gigue que cause le mécanisme de *Forward*. Cette variation implique une baisse de la qualité de service aussi bien pour les flux UDP que TCP. Les attentes d'amélioration pour ce *handover* sont importantes puisque l'on passe du segment satellite au terrestre qui, théoriquement, offre de meilleures performances. Par conséquent, le mécanisme de *Forward* du *handover* inter-satellite-terrestre nécessite une optimisation.

Handover inter-terrestre-satellite

Ce *handover* ne subit pas de contraintes particulières lors du *handover*. En effet, la gigue induite par le mécanisme de *Forward* est simplement égale à l'ajout du délai satellite, car seul le backhaul de l'eNB cible est réalisé par un lien satellite. Il est donc impossible de réduire cette variation du délai lors du *handover*. En outre, les possibilités d'optimisation du mécanisme de *Forward* sont très faibles, car c'est le segment terrestre qui est à l'initiative du mécanisme de *Forward*. Nous souhaitons nous dispenser de modifications à l'intérieur de ce segment. Si l'on considère conjointement à cela que les améliorations sont minimales, nous avons décidé de ne proposer aucune optimisation.

Handover intra-satellite

Ce type de *handover* implique la plus importante gigue pendant le *Forward* des paquets, alors que les interfaces S1 possèdent des caractéristiques similaires et appartiennent toutes les deux au segment satellite. Tout comme pour le *handover* inter-satellite-terrestre, le segment source est fourni par un système satellite. Ainsi, l'optimisation du mécanisme de *Forward* est à la fois possible et nécessaire.

4.2 Contraintes

La contrainte du tunnel de *Forward* sur le lien satellite (C_{tunnel})

Le fondement de notre optimisation réside dans la suppression du tunnel source satellite qui engendre la surconsommation des ressources satellite et l'augmentation de la gigue. Cependant, cela induit la modification du point de départ du mécanisme de *Forward*. Celui-ci est habituellement l'eNB source, puisqu'il représente l'entité la plus proche de l'UE et, par conséquent, il détient la connaissance la plus précise, après l'eNB source, des paquets que l'UE a reçus et ceux qui doivent subir le mécanisme de *Forward*. Cette constatation implique donc l'utilisation d'une entité au-delà de la passerelle satellite comme point d'origine du *Forward*. Cela peut être la SGW ou la passerelle satellite.

La contrainte de l'échange de contexte sur le lien satellite (C_{contexte})

La deuxième contrainte se situe dans le mécanisme qui assure l'absence de perte de paquets lors du *handover*. L'analyse dans le chapitre 2 a montré l'impossibilité d'utiliser la procédure standard, qui est fondée sur la transmission du contexte PDCP de l'eNB source à l'eNB cible sur le lien de backhaul (Figure 30). La difficulté de ce mécanisme réside dans l'impossibilité pour l'eNB cible, de connaître le dernier paquet reçu par l'UE par l'intermédiaire du réseau de cœur. En effet, seul l'eNB source possède cette information et le backhaul satellite entrave la transmission de cette information pendant la phase d'exécution.

La contrainte de la préparation anticipée à double décision ($C_{\text{Prép2Dcision}}$)

La préparation anticipée à double décision est une optimisation proposée dans le chapitre 3 qui permet de prendre en compte la longue phase de préparation due au backhaul satellite. La particularité de la double décision implique la désolidarisation de la fin de la signalisation de la phase de préparation, du début de la phase de l'exécution. En effet, la réception du *Handover Command* n'entraîne plus automatiquement l'envoi du message *RRC Connection Reconfiguration*. Pour cette raison, la SGW est dans l'impossibilité de connaître l'instant où l'eNB source déclenche la phase d'exécution.

La contrainte de la préparation multiple ($C_{\text{PrépMultiple}}$)

La préparation anticipée à double décision se combine avec la préparation multiple. Cette préparation multiple peut poser problème dans le choix du mécanisme de *Forward*. Elle va permettre de configurer plusieurs eNB cibles. Cependant, la SGW ne possède aucun moyen de connaître l'eNB cible choisi par l'eNB source lors de la deuxième décision. Par conséquent, le mécanisme doit tenir compte de cette particularité pour un *handover* incluant un eNB avec un backhaul satellite.

Ces contraintes seront par la suite citées le plus souvent par leur abréviation listée dans le Tableau 16.

Nom	Information
C_{tunnel}	Le lien satellite rend inefficace l'utilisation des tunnels de <i>Forward</i> en particulier entre l'eNB source et la SGW source
C_{contexte}	Le lien satellite interdit l'échange du contexte PDCP entre eNB pendant la phase d'exécution
$C_{\text{Prép2Dcision}}$	La préparation anticipée à double décision désolidarise la fin de la signalisation de la phase de préparation et le début de l'exécution
$C_{\text{PrépMultiple}}$	Seul l'eNB source connaît l'eNB à partir de la fin de la phase de préparation.

Tableau 16 Contraintes sur le mécanisme contre la perte des paquets durant le *handover*

4.3 Etat de l'art

Le mécanisme de *Forward* permet d'assurer l'absence de perte pendant le *handover*. Cela réduit, en particulier, les effets néfastes du *handover* sur le protocole TCP. En outre, il est combiné à un mécanisme de réordonnancement, ainsi qu'à un mécanisme qui évite l'envoi dupliqué de paquets. Ce dernier permet de minimiser l'utilisation des ressources de l'interface radio (NTT DoCoMo, April 2006). Ces trois mécanismes sont très bénéfiques pour les trafics TCP. Ainsi, les auteurs de (Racz, Temesvary, & Reider, July 2007) ont étudié leurs impacts sur le protocole TCP et ont mis en évidence leur nécessité pour assurer de très hauts débits.

4.3.1 Les mécanismes de *Forward*

De la même manière que pour l'optimisation de la préparation, il n'y a aucun mécanisme de *Forward* et de réordonnancement, qui soit adapté à fois à LTE et au backhaul satellite, puisque les scénarios d'utilisation étaient inexistant. Ceux mis en place dans les systèmes intégrés ne sont pas affectés par les contraintes des tunnels (C_{tunnel}) et du contexte (C_{contexte}), puisque les tunnels de *Forward* n'empruntent pas l'interface satellite. Les études où le satellite influence le mécanisme de *Forward* sont très rares et, à notre connaissance, elles ne mettent en place que des procédures de type *handover* vertical. Dans ces cas-là, la mobilité s'effectue au niveau de la couche IP avec des protocoles tels que MIP, HMIP et FMIP. La connexion à un nouveau point d'accès se révèle alors très longue dans la majorité des cas, en comparaison des technologies cellulaires (Gayraud, Alphand, Berthou, Baudoin, & Duquerroy, September 2006). En outre, ces *handovers* sont de type *Make-Before-Break* pour lesquels la notion de *Forward* est moins contraignante voire inutile (Han, Han, & Lee, May 2008). Par conséquent, notre état de l'art se concentre sur les mécanismes de *Forward* liés aux technologies cellulaires et, en particulier LTE, bien que certains mécanismes des protocoles de la famille Mobile IP puissent être réutilisés (Montavont & Noel, 2002).

4.3.1.1 Les mécanismes pour éviter les pertes pendant les *handovers* intra-LTE

Les premières discussions entre les partenaires du 3GPP, qui traitent du mécanisme de *Forward*, ont permis de décider de l'utilisation d'un mécanisme qui supprime les pertes de paquets durant la procédure de *handover*. Ainsi, (Siemens, February 2006) regroupe trois mécanismes : le *Forward* entre eNB source et cible, le *bi-cast* et la réalisation conjointe du *handover* et du *Path Switch*.

Le *Forward* entre eNB source et eNB cible

C'est le mécanisme de *Forward* inter-eNB qui a été adopté par le standard. Il est inefficace dans le cadre d'un scénario mettant en jeu des eNB satellite et il est décrit dans le chapitre 2.

Le *bi-cast*

Pour cette proposition, les données sont envoyées à la fois à l'eNB source et à l'eNB cible par la SGW. Cette transmission *bi-cast* débute lors de la préparation du *handover* et prend fin pendant la phase de terminaison du *handover*. Cette technique permet d'éviter le tunnel satellite (C_{tunnel}), mais elle nécessite un mécanisme contre la duplication. En effet, l'eNB cible n'a pas connaissance de l'identité du dernier paquet reçu par l'UE. Par exemple, il ne peut savoir si les paquets 1 et 2, dans la Figure 35, ont été reçus par l'UE. Il est donc nécessaire d'implanter un mécanisme qui permette d'éviter que l'eNB cible ne transmette de nouveau les paquets déjà reçus par l'UE et acquittés pour l'eNB source.

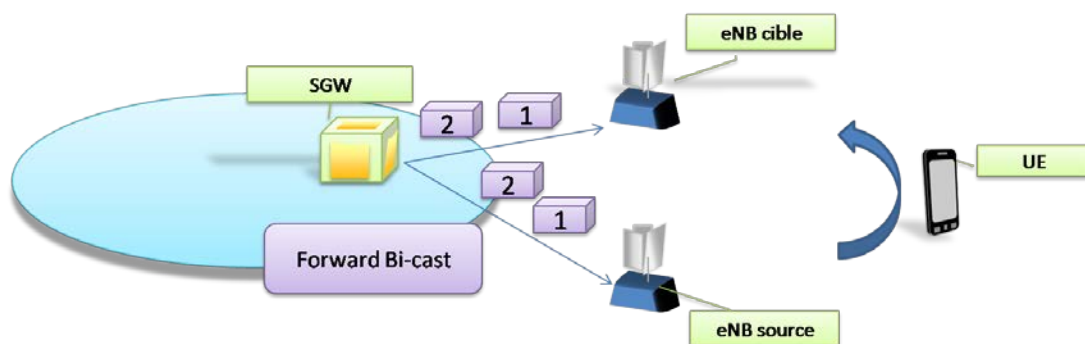
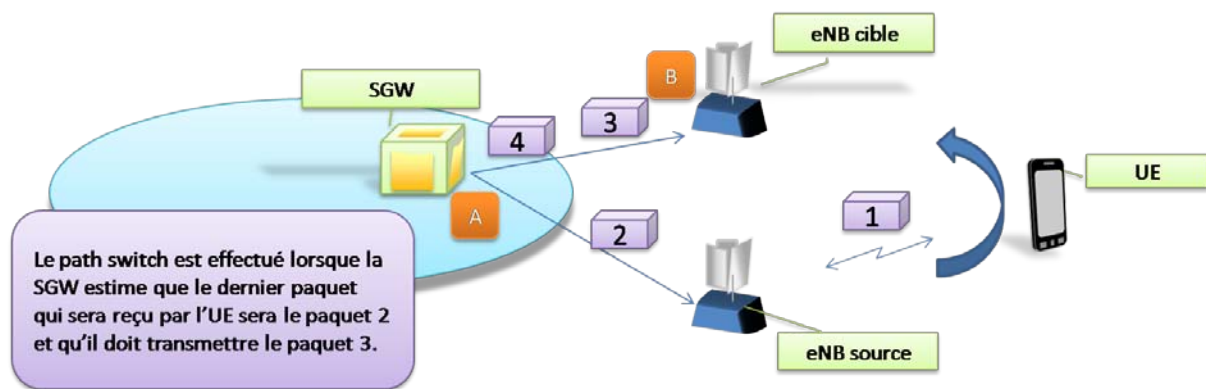


Figure 35 Forward par *bi-cast*

La réalisation conjointe du *handover* et du *Path Switch*

L'objectif est de réaliser le *Path Switch* conjointement avec le début de l'exécution du *handover*, pour que le transfert de paquets ne soit pas une nécessité (Figure 36). Cette technique a été réutilisée pour éviter que le *Forward* entre l'eNB source et cible ne fasse varier le RTT des paquets et ainsi, réduire les performances de TCP (Pacífico, Pacífico, Fischione, Hjalrmasson, & Johansson, December 2009). Bien que la gigue introduite par le lien X2 lors du *Forward* semble négligeable, cette technique pourrait être utile dans le cadre d'un backhaul satellite, puisque on s'affranchit des tunnels de *Forward* (C_{tunnel}). Cependant, elle repose sur la connaissance par la MME ou la SGW de l'instant précis où débute l'exécution du *handover* (A dans la Figure 36). Or, dans notre scénario, cette information ne peut être connue de ces entités en raison de l'utilisation de la préparation à double décision. En effet, il n'y a aucun lien temporel entre la fin de la signalisation de la préparation et le début de la phase d'exécution ($C_{\text{Prép2Décision}}$).

Figure 36 Réalisation conjointe du *handover* et du *Path Switch*.

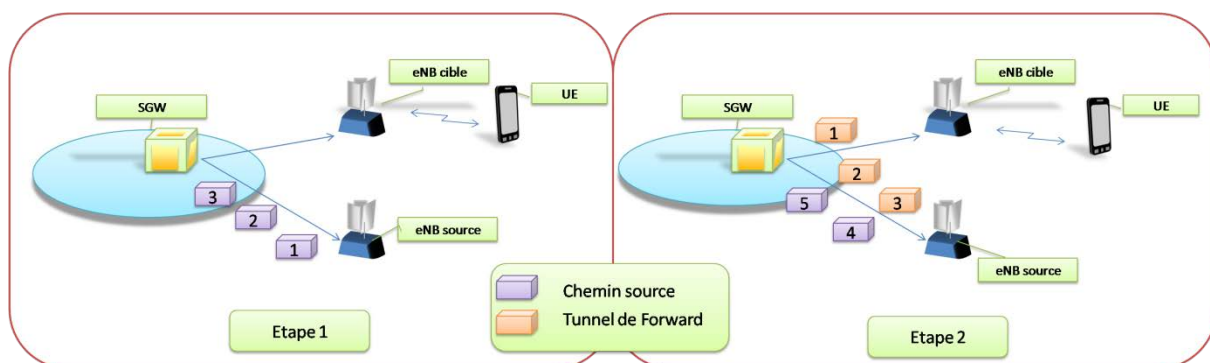
4.3.1.2 Les mécanismes pour éviter les pertes pendant les *handovers* inter-RAT

Le *handover* inter-RAT se déroule entre des stations de base n'utilise pas le même type de technologie d'accès. Les mécanismes de *Forward* doivent donc prendre en compte cette particularité ; ce qui peut potentiellement rendre les mécanismes de *Forward* intéressants dans notre scénario. Les mécanismes inter-RAT se fondant sur un *handover Make-Before-Break* ne sont pas utilisables dans notre scénario (Liu, Boukhatem, Martins, & Bertin, September 2010).

De nombreux mécanismes de *Forward* étudiés pour les *handovers* inter-RAT sont similaires à ceux du *handover* intra-LTE et sont décrits dans (NEC, January 2007) (Alcatel, November 2006), cependant deux nouvelles méthodes sont apportées par ces documents.

Backward handover

Ce mécanisme est simple. Les paquets ne sont pas transférés à l'eNB cible pendant la phase d'exécution (Etape 1 dans la Figure 37). L'eNB source stocke les paquets en attendant la fin de cette étape. Puis, pendant la phase de terminaison, l'eNB source va transférer les paquets stockés à l'eNB cible (Etape 2 dans la Figure 37). Ce mécanisme n'est pas intéressant avec des backhaul satellite, car il conserve les mêmes tunnels de *Forward*.

Figure 37 Backward *handover*

Forward par buffer

Pendant la phase de préparation, la SGW arrête d'envoyer les paquets à l'eNB source et elle utilise des *buffers* pour les stocker (Etape 1 dans la Figure 38). Lors de la phase de terminaison, elle recommence l'envoi des paquets mais, cette fois-ci, vers l'eNB cible (Etape 2 dans la Figure 38). La

complexité de cette solution réside dans la gestion du *buffer*. En outre, l'activation de ce *buffer* se déroule pendant la phase de préparation. Ainsi, l'eNB source ne reçoit plus de paquet dès que la signalisation est terminée. Ce comportement invalide l'aspect bénéfique de la préparation à double décision. En effet, l'intérêt de cette optimisation est d'anticiper la phase de préparation et de séparer la décision de préparation de celle d'exécution. Cependant, si lors de la préparation la SGW ne transmet plus les paquets à l'eNB source, il est nécessaire de déclencher l'exécution dès que la phase de préparation est finie, pour permettre à l'eNB de recevoir les paquets par l'intermédiaire de l'eNB cible.

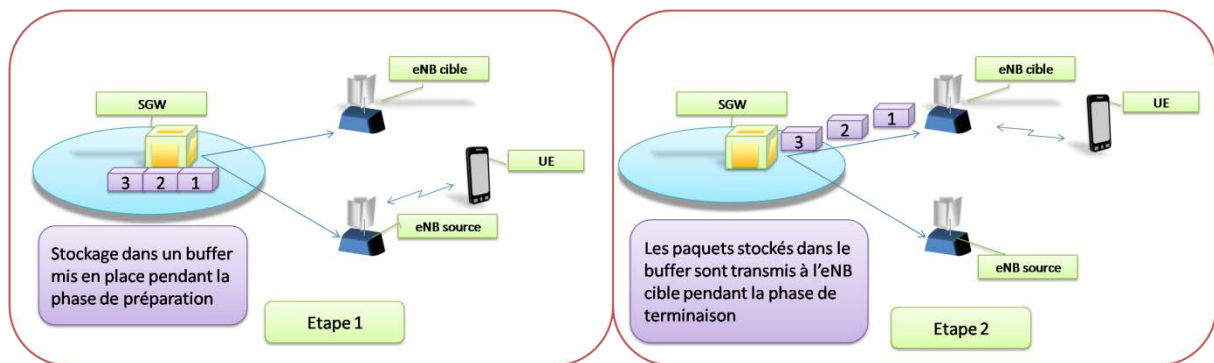


Figure 38 Forward par buffer

4.3.1.3 Les mécanismes pour éviter les pertes pendant un handover avec un eNB relais

Les eNB relais (RN, *Relay Node*) sont des stations de base dont le backhaul est réalisé par l'intermédiaire de l'interface radio d'une autre station de base (DeNB, *DonoreNodeB*). Cela offre la possibilité à l'opérateur d'étendre la couverture de son réseau sans devoir déployer des backhails filaires. La caractéristique commune entre notre scénario et l'architecture des RN est la limitation des ressources du lien S1. En effet, une partie de l'interface S1 est constituée de l'interface radio pour laquelle les ressources doivent être économisées (Figure 39).

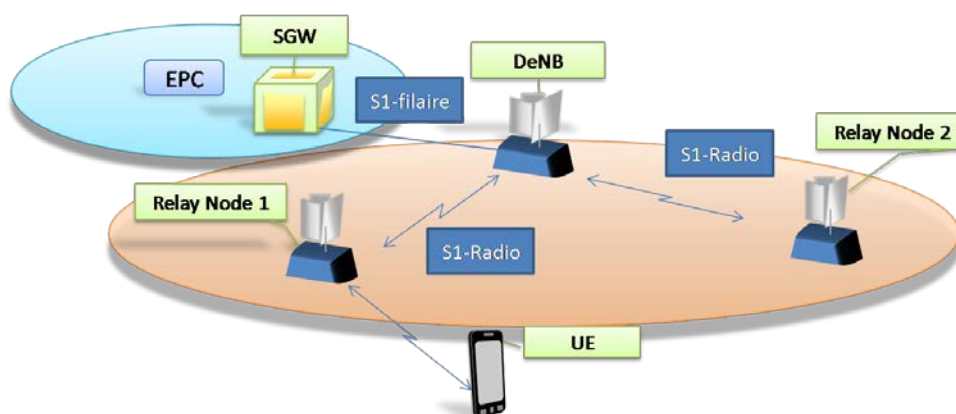


Figure 39 Architecture des eNB relais (RN)

Ainsi, (Research In Motion UK Limited, July 2009) propose d'utiliser la méthode de la réalisation conjointe de l'exécution du *handover* et du *Path Switch* proposé pour les *handovers* intra-LTE. De même, (LG Electronics Inc, May 2010) réutilise le principe du *Forward par buffer*. Seuls (Research In Motion UK Limited, July 2009) et (Motorola, May 2009) définissent un nouveau mécanisme. Dans

celui-ci, le DeNB source conserve une copie des paquets PDCP (A dans la Figure 40). Ces copies sont stockées dans un *buffer*. Des messages *PDCP Status Report* sont périodiquement envoyés entre DeNB et le RN source, pour ne conserver dans le *buffer* que les paquets dont la réception n'a pas été validée par l'UE (B dans la Figure 40). Cette méthode ne peut être réutilisée, car même si elle évite l'utilisation des tunnels, elle implique l'envoi d'une information entre l'eNB source et la SGW après la fin de la phase de préparation ce qui va à l'encontre de la contrainte $C_{Contexte}$.

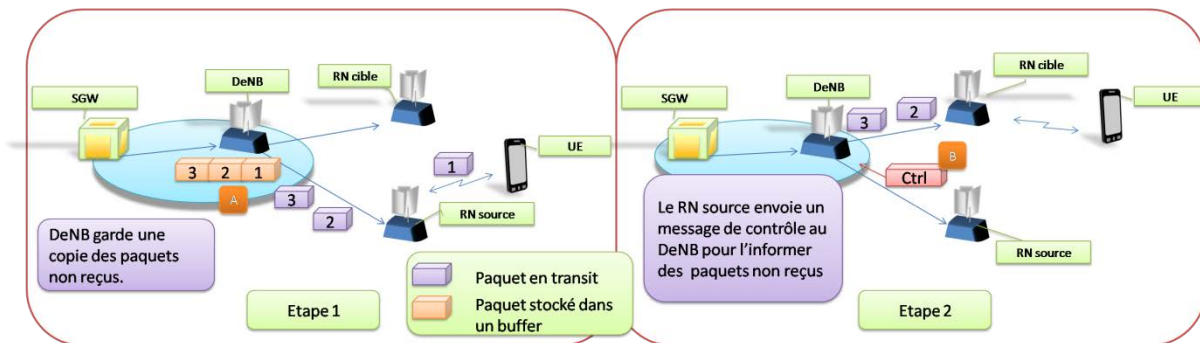


Figure 40 *Forward* pour les relais

4.3.2 Les mécanismes pour éviter les duplications dans le 3GPP

Lors de l'étape de standardisation, le choix du 3GPP s'orienta vers le *Forward* entre l'eNB source et cible des paquets de la couche PDCP (NTT DoCoMo, Inc., August 2007) (Ericsson, August 2007). Ce *Forward*, en particulier pour le *S1 handover*, fit l'objet de débats entre un *Forward* cumulatif ou sélectif. Le *Forward* sélectif permet de transférer uniquement les paquets, pour lesquels l'eNB n'a pas reçu d'accusé de réception (Etape 2-Sélectif dans la Figure 41). Le *Forward* cumulatif permet de s'abstenir de transférer le contexte PDCP, car, à partir du premier paquet PDCP non validé, tous les suivants sont transférés (Ericsson, August 2007) (Ericsson, October 2007). En contrepartie, il implique un certain nombre de retransmissions sur l'interface radio de paquets PDCP déjà reçus par l'UE (Etape 2-Cumulatif dans la Figure 41). Il offre la possibilité de ne pas transférer le contexte PDCP, avec comme conséquence, une faible probabilité de subir une réception légèrement désordonnée des paquets transférés. Bien qu'il réponde à la contrainte $C_{Contexte}$, il repose sur le transfert des paquets entre les eNB et va à l'encontre de C_{tunnel} .

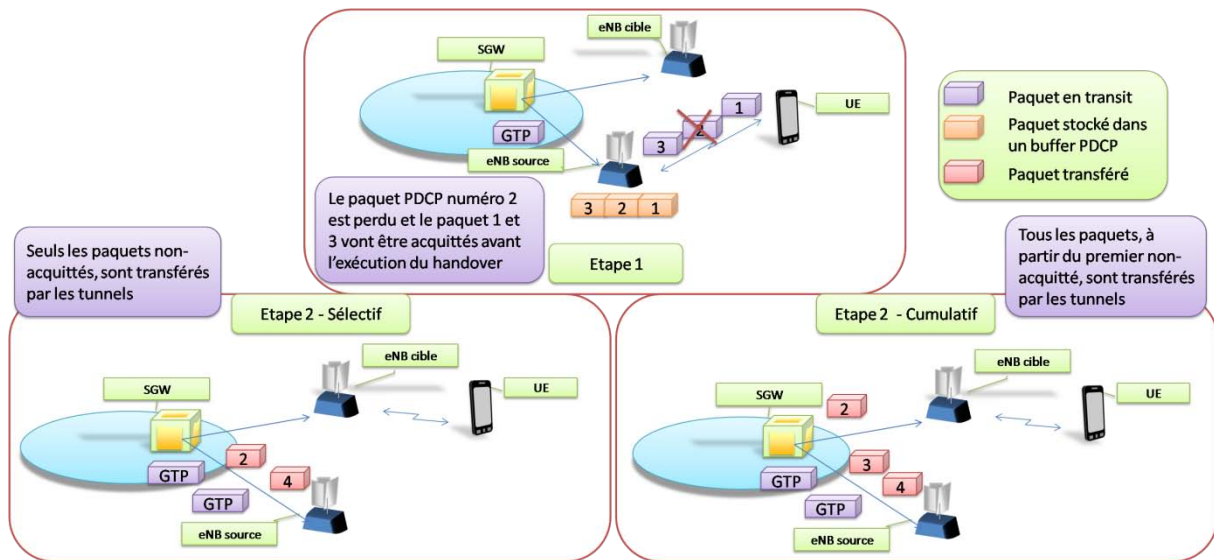


Figure 41 Mécanisme contre les duplications

4.3.3 Les mécanismes pour réaliser l'ordonnancement après le *Path Switch* dans le 3GPP

Après le *Path Switch*, l'eNB cible a besoin de réordonner les paquets, en provenance de l'eNB source par l'intermédiaire du *Forward* et ceux arrivant par le chemin cible. Plusieurs propositions ont été faites pendant les discussions du 3GPP (NTT DoCoMo, Inc., August 2007) (Ericson, August 2007) (NEC, August 2007). Elles sont classées en trois catégories.

Réordonnancement fondé sur une temporisation

Lorsque l'eNB cible reçoit le message *Handover Confirm*, une temporisation est déclenchée. Lorsque cette temporisation est en activité, l'eNB cible envoie les paquets qui proviennent du tunnel de *Forward* (A dans la Figure 42) et il emmagasine ceux venant du chemin cible. Puis, quand la temporisation arrive à son terme, il commence à transmettre les paquets contenus dans le *buffer* du chemin cible (B dans la Figure 42). En raison de la transmission des numéros de séquence PDCP dans les paquets transférés, l'eNB cible attribue des numéros de séquence aux paquets du chemin cible, en assurant leur continuité avec les derniers paquets transférés. Cette solution est la plus simple, mais elle implique que la durée du mécanisme de réordonnancement est constante. L'attribution des numéros de séquence est impossible à cause des contraintes C_{tunnel} et C_{Contexte} . Cependant, elle peut être évincée sans conséquence pour le réordonnancement et ainsi la seule difficulté réside dans le choix de la valeur de la temporisation, qui dépendra du mécanisme de *Forward* mis en place.

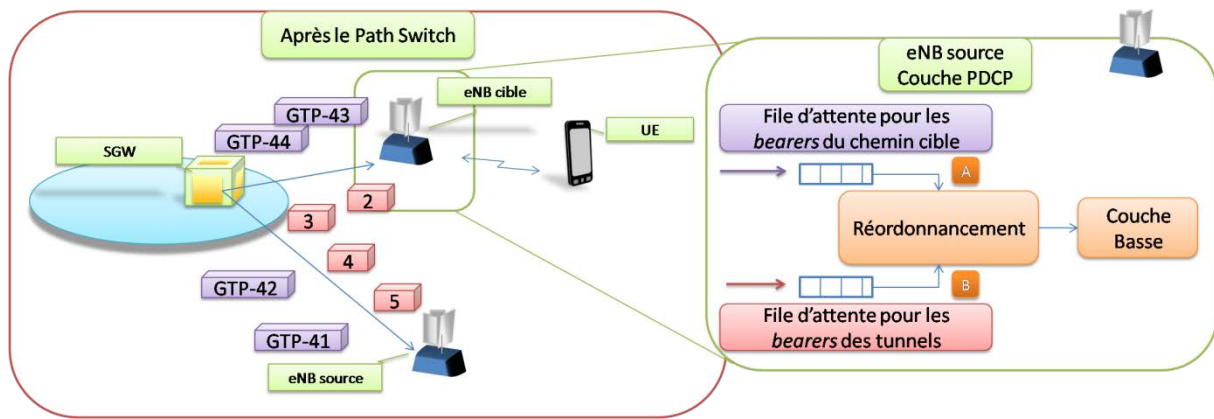


Figure 42 Réordonnement dans l'eNB cible

Réordonnement fondé sur le marquage de la fin des tunnels

Lors du *Path Switch*, un marquage va permettre à l'eNB cible de déterminer le dernier paquet qui a emprunté le chemin de *Forward*. À sa réception, l'eNB cible va pouvoir basculer de l'envoi des paquets transférés, à ceux en provenance du chemin cible. Le marquage peut être réalisé de plusieurs manières :

- **Marquage du dernier paquet GTP-U** empruntant le chemin *Forward* (Le paquet GTP-42 dans la Figure 42). Cette étape nécessite l'ajout d'une temporisation dans le cas où le dernier paquet est perdu.
- **Envoi d'un paquet GTP-U particulier (End Marker)** qui ne contient aucune donnée, mais signifie la fin des paquets transférés (Juste après le paquet GTP-42 dans la Figure 42). Cette étape nécessite aussi l'ajout d'une temporisation dans le cas où le dernier paquet est perdu.
- **Envoi de plusieurs paquets End Marker** pour minimiser l'impact d'une perte d'un paquet End Marker.

Tous les variantes de cette solution sont utilisables dans notre scénario, si la cohérence avec le mécanisme de *Forward* choisi est conservée.

Réordonnement fondé sur un *offset*

Cette méthode utilise un *offset* qui permet de prévenir l'eNB cible de la relation entre les numéros de séquence GTP et PDCP. Par exemple, dans la Figure 42, le paquet GTP numéro 41 aura comme numéro PDCP 6. Il est transmis pendant la préparation du *handover*. Pour connaître les derniers paquets PDCP reçus par l'UE sur la liaison descendante, et par l'eNB sur la liaison montante, des messages *PDCP Status Report* seront transmis sur l'interface radio. Grâce à cet *offset*, l'eNB cible reçoit les paquets du chemin cible. Le premier paquet qui arrive permet de définir le numéro de séquence du dernier paquet qui a suivi le chemin de *Forward* puis, à la réception de celui-ci, l'eNB cible peut basculer de l'envoi des paquets en provenance du chemin de *Forward*, à ceux en provenance du chemin cible. Tout comme la solution de marquage, la perte du dernier paquet transféré ou celui du premier paquet du chemin cible peut entraîner un blocage du mécanisme qui n'est résolu que grâce à une temporisation. Cette solution possède la particularité de ne pas avoir besoin du contexte PDCP puisqu'il est transféré entre l'eNB cible et l'UE après l'exécution du *handover* et, puisque le mécanisme de réordonnement peut se reposer uniquement sur les numéros de séquence GTP.

4.3.4 Les mécanismes intéressants

Seuls trois mécanismes remplissent les deux principales contraintes (C_{tunnel} , et C_{Contexte}). Le *Forward par buffer* et le *Forward* avec la réalisation conjointe de l'exécution du *handover* et du *Path Switch* rentrent en conflit avec l'optimisation de la phase de préparation à double décision. L'unique technique qui remplisse les critères est le *bi-cast*, bien qu'elle nécessite de mettre en place un mécanisme contre la duplication de paquets. Ceux définis par le 3GPP (cumulatif et sélectif) ne sont pas utilisables, puisque qu'ils reposent sur le transfert du contexte PDCP entre les eNB et l'utilisation des tunnels de *Forward*. L'utilisation de la préparation multiple impose de ne pas se limiter à un *Forward bi-cast*. En effet, la préparation multiple permet de préparer plusieurs eNB, soit dans l'optique de choisir après la phase de préparation le meilleur eNB, soit d'augmenter la rapidité de récupération en cas d'échec du *handover* sur la première eNB cible. Dans les deux cas, pour assurer une meilleure efficacité, l'utilisation d'un *Forward* multicast vers tous les eNB préparés est nécessaire pour qu'ils puissent recevoir les paquets transférés dès la phase d'exécution. Le Tableau 17 récapitule l'adaptation des différents mécanismes aux contraintes de notre scénario.

	C_{tunnel}	C_{Contexte}	$C_{\text{Prép2Dcision}}$	$C_{\text{PrépMultiple}}$	Commentaire
<i>Forward</i> inter-eNB	✗	✗	✓	✓	Utilise des tunnels de <i>Forward</i> sur le satellite
<i>Bi-cast</i>	✓	✓	✓	✓	Intéressant, mais nécessite un mécanisme contre les duplications
HO et path switch conjoint	✓	✓	✗	✗	La réussite de cette technique est rendue impossible à cause de la contrainte $C_{\text{Prép2Dcision}}$
<i>Backward HO</i>	✗	✗	✓	✓	Utilise des tunnels de <i>Forward</i> sur le satellite
<i>Forward buffer</i>	✓	✓	✗	✓	Mécanisme simple, mais ne peut être mis en place avec l'algorithme à double décision
RN ??	✓	✗	✓	✓	Impossible puisqu'il faut de nombreux rapports entre l'eNB source et la SGW.

Tableau 17 Mécanisme contre la perte de données pendant le *handover*

4.4 Optimisations

Dans ce paragraphe, nous proposons deux optimisations pour réaliser un *handover* induisant le minimum de pertes et de duplications de paquets. Le paragraphe précédent nous a permis de mettre en relief que seul un mécanisme de *Forward bi-cast* ou multicast semble correspondre aux attentes d'un *handover*, dont l'eNB source possède un backhaul satellite. Cependant, il ne peut être efficace

que s'il est exécuté conjointement avec un mécanisme contre les duplications de paquets. Dans un premier temps, nous adaptons ce mécanisme de *Forward* aux caractéristiques du *handover* inter-satellite-terrestre, puis nous reprendrons la réflexion pour le *handover* intra-satellite.

4.4.1 Mécanisme de *Forward bi-cast avec buffer* pour les *handovers* inter-satellite-terrestre

4.4.1.1 Principe

Le *handover* inter-satellite-terrestre possède deux particularités. Premièrement, ce *handover* implique un changement de segment ; les délais de transmission sur le chemin source (satellite) sont bien évidemment plus importants que sur le chemin cible (terrestre). Deuxièmement, le segment terrestre ne doit pas être affecté par l'optimisation, ou de la manière la plus faible possible. Le *Forward* multicast présente ainsi quelques difficultés additionnelles vis-à-vis de la gestion des duplications. En effet, le *bi-cast* doit être activé dans la SGW pendant la phase de préparation, entre l'eNB source et cible. Cependant, cette étape configure plusieurs eNB et par conséquent, cela nous oblige à utiliser du multicast, car nous n'avons aucune certitude quant au futur eNB cible. Ce multicast engendre alors une surconsommation des ressources dans le segment terrestre. Il implique des modifications de la gestion des files d'attente des entités terrestres, en particulier de l'eNB cible. Pour éliminer ces problèmes, nous proposons de combiner le *Forward bi-cast* avec celui par *buffer*. Le premier flux est toujours envoyé à l'eNB source, mais le deuxième flux est stocké dans un *buffer* à l'intérieur de la SGW satellite (Figure 43). Ce n'est que lors de la phase de terminaison que le *buffer* transmet les paquets stockés à l'eNB cible (Figure 44). Cette technique nous permet de nous affranchir de la transmission multicast, puisque lors de la phase de terminaison l'eNB cible choisi est connu. La gestion des files d'attente est déportée à l'intérieur du segment satellite dans la SGW source. L'attente de la phase de terminaison, pour envoyer les paquets stockés dans le *buffer*, entraîne une légère augmentation du temps d'interruption de service pendant le *handover*. Un mécanisme contre la duplication de paquets est toujours nécessaire.

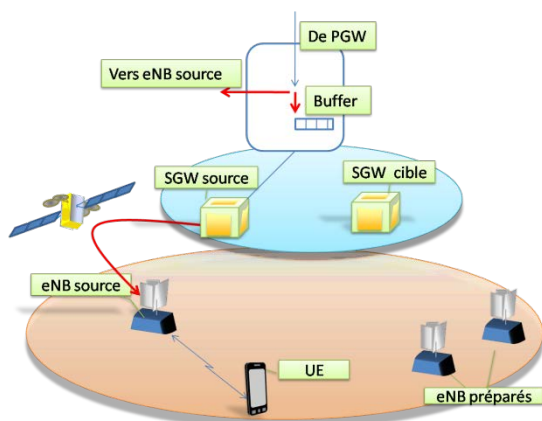


Figure 43 *Forward bi-cast avec buffer* avant la phase de terminaison

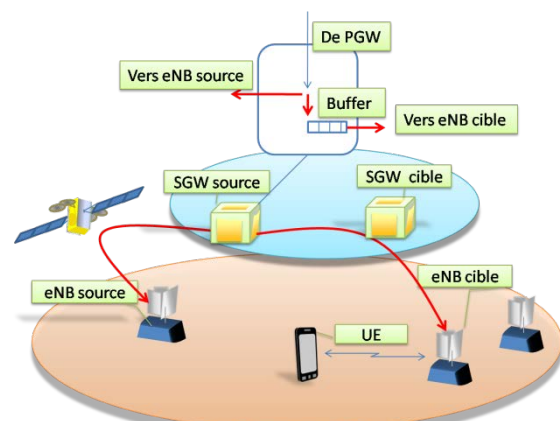


Figure 44 *Forward bi-cast avec buffer* après la phase de terminaison

4.4.1.2 La mise en place du buffer

Nous choisissons, dans un premier temps, de mettre en place le *buffer* pendant la création des tunnels (A dans la Figure 45). Cette solution nous permet de ne pas modifier la signalisation entre les différentes entités et d'informer la SGW satellite du prochain *handover*. Par conséquent, à la

réception du message *Create Indirect Data Forwarding Tunnel Request*, la SGW satellite va entamer la création du *buffer* et stocker les premiers paquets en direction de l'UE.

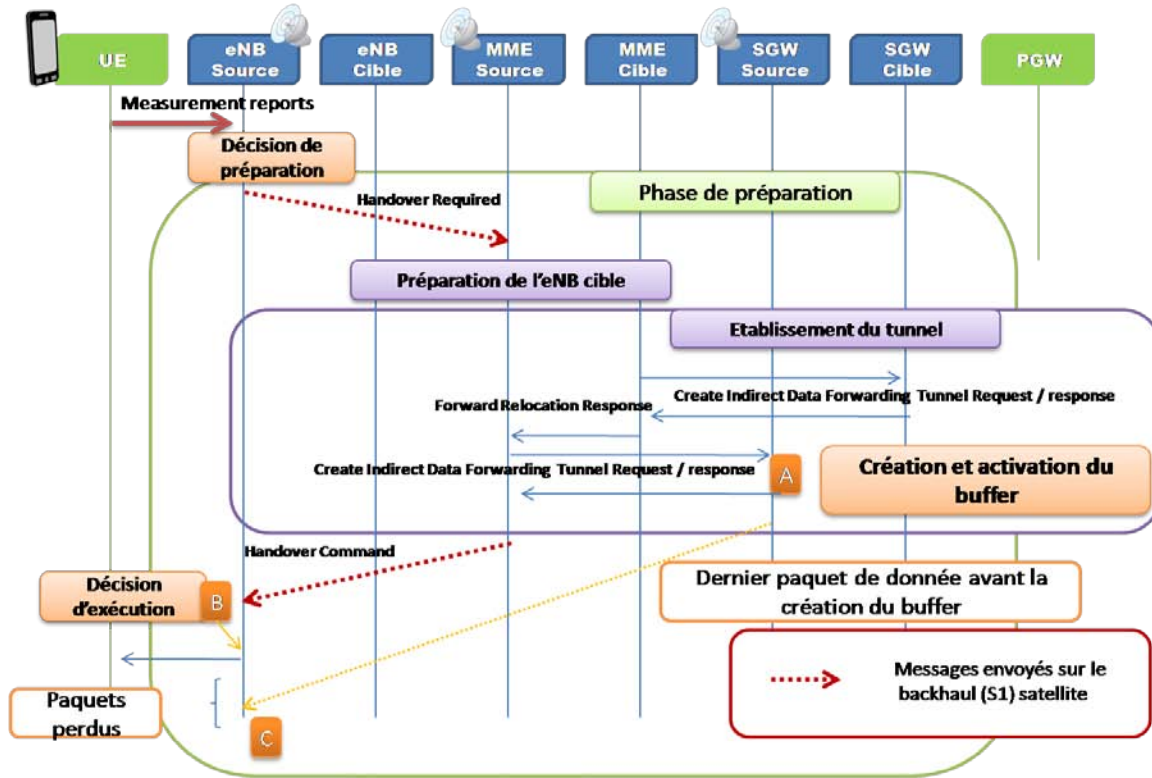


Figure 45 Mise en place du *buffer* de *Forward*

Malgré la mise en place du *buffer* pendant la phase de préparation, des pertes dues au *handover* peuvent survenir. En effet, le dernier paquet envoyé par la SGW source, qui n'est pas présent dans le *buffer* (P_{dernier}), est transmis en même temps que le message *Create Indirect Data Forwarding Tunnel Response*. Ainsi, si la fin de la transmission de paquets de l'eNB source vers l'UE, représenté par l'instant où l'eNB source décide d'exécuter le *handover* (B dans la Figure 45), se produit avant la réception du paquet P_{dernier} par l'UE (C dans la Figure 45), le mécanisme *bi-cast* avec *buffer* sera incapable d'assurer un *handover* sans perte de données. L'importance de ces pertes est proportionnelle à l'écart de temps entre la décision d'exécution ($T_{\text{exécution}}$ et B dans la Figure 45) et la réception, par l'eNB source, du dernier paquet de donnée non présent dans le *buffer* ($T_{\text{RxPdernier}}$ et C dans la Figure 45). Le calcul de cet écart (Équation 4.3) est donc obtenu par la soustraction de $T_{\text{RxPdernier}}$ (Équation 4.2) à $T_{\text{exécution}}$ (Équation 4.1). Les éléments de ces équations sont décrits dans le Tableau 18.

$$(4.1) \quad T_{\text{exécution}} = T_{\text{CreationBuffer}} + D_{\text{TxsGW} \rightarrow \text{MME}} + D_{\text{ProcMME}} + D_{\text{FileSatCTRL}} + D_{\text{Txsat}} + D_{\text{Décision}} + D_{\text{File-eNBCTRL}}$$

$$(4.2) \quad T_{\text{RxPdernier}} = T_{\text{CreationBuffer}} + D_{\text{FileSatUSR}} + D_{\text{Txsat}} + D_{\text{Décision}} + D_{\text{File-eNBUSR}}$$

$$(4.3) \quad E_{\text{exécution} \rightarrow \text{RxPdernier}} = (D_{\text{FileSatUSR}} - D_{\text{FileSatCTRL}}) + (D_{\text{FileeNBUSR}} - D_{\text{FileeNBCTRL}}) - (D_{\text{TxsGW} \rightarrow \text{MME}} + D_{\text{ProcMME}}) + D_{\text{Décision}}$$

Paramètres	Description
$T_{exécution}$	Instant où l'eNB décide de déclencher la phase d'exécution
$T_{RxPdernier}$	Instant où l'eNB source reçoit le dernier paquet qui n'est pas présent dans le <i>buffer</i>
$T_{CreationBuffer}$	Instant où la SGW source crée et active le <i>buffer</i>
$D_{TxSGW \rightarrow MME}$	Durée de la transmission entre la SGW et la MME
$D_{ProcMME}$	Durée du traitement du paquet de contrôle dans la MME
$D_{FileSat CTRL}$	Durée passée dans les files d'attente de la passerelle satellite par les paquets du plan de contrôle
$D_{FileSat USR}$	Durée passée dans les files d'attente de la passerelle satellite par les paquets du plan utilisateur
D_{TxSat}	Durée de la transmission sur le lien satellite
$D_{Décision}$	Durée entre la réception du <i>Handover Command</i> par l'eNB source et la décision de déclencher la phase d'exécution (B dans la Figure 45)
$D_{File-eNB CTRL}$	Durée passée dans les files d'attente de l'eNB source par les paquets du plan de contrôle
$D_{File-eNB USR}$	Durée passée dans les files d'attente de l'eNB source par les paquets du plan utilisateur

Tableau 18 Description des paramètres des équations

Pour obtenir cet écart, nous avons émis l'hypothèse qu'il n'y a aucun élément réseau entre la MME et la passerelle satellite, ainsi qu'entre la SGW et la passerelle satellite. Cette hypothèse est vraisemblable, puisque ces entités sont spécifiques au segment satellite et sont donc probablement co-localisées. Par conséquent, nous pouvons considérer que la durée $D_{TxSGW \rightarrow MME}$ est négligeable. De même, $D_{ProcMME}$ est négligeable, car il représente le temps de traitement de la réponse du message *Create indirect Data Forwarding Tunnel Response*.

Nous remarquons que l'écart est principalement dû à la différence de gestion dans les files d'attente entre les paquets appartenant au plan de contrôle et ceux du plan utilisateur. Comme, le lien satellite est le lien limitant, la grande majorité de cet écart est la conséquence de la gestion des files d'attente dans le lien satellite. Par conséquent, si la durée du retard induit par la décision du *handover* ($D_{Décision}$) est supérieur à $(D_{FileSat USR} - D_{FileSat CTRL})$, le mécanisme permet d'éviter les pertes. La valeur de $D_{Décision}$ est imprévisible. Elle peut être égale à 0 et ainsi engendrer des pertes. Cependant, l'optimisation de la préparation permet d'anticiper la phase de préparation et ainsi, cette durée est dans la plupart des cas non nulle.

Pour améliorer ce mécanisme, une solution consiste à anticiper l'activation du *buffer*. Ceci aurait pour conséquence d'ajouter des messages de signalisation en amont des messages de création des tunnels de *Forward*. Cette technique réduirait le problème mais ne le supprimerait pas, tout en ajoutant des modifications importantes au standard ; la marge d'anticipation est trop faible.

La deuxième solution propose de configurer le *buffer* de manière permanente dans la SGW. Ainsi, il stockerait dans un *buffer* annexe, tous les paquets à transmettre sur le lien satellite, pendant une

durée permettant de compenser le temps passé dans les files d'attente de la passerelle satellite. Lors de la création du *buffer*, la SGW remplit celui-ci par des paquets à destination de l'UE qui sont contenus dans le *buffer* annexe.

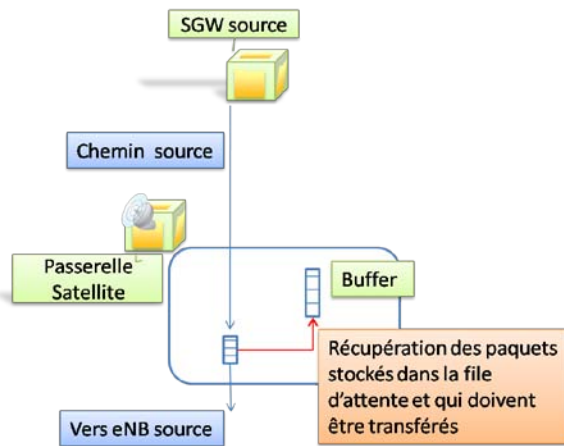


Figure 46 Activation du *buffer* de *Forward* dans la passerelle satellite

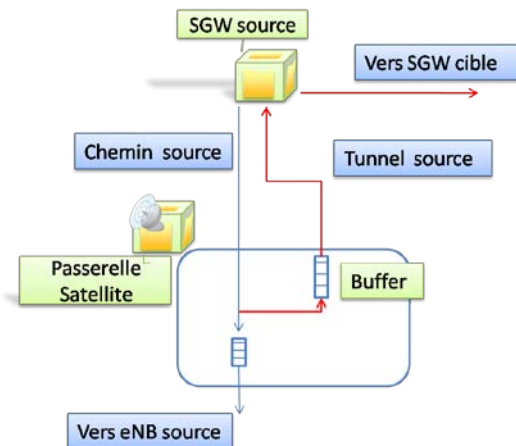


Figure 47 Buffer de *Forward* dans la passerelle satellite après le déclenchement du *Forward*

La troisième solution permet de récupérer les paquets qui sont présents dans les files d'attente de la passerelle satellite. On déporte le *buffer* de *Forward* dans cette entité. À l'activation du *buffer*, les paquets, présents dans les files d'attente, correspondant à l'UE sont copiés et envoyés vers le *buffer* de *Forward* (Figure 46). Puis, lorsque le *Forward* est déclenché lors de la phase de terminaison, la passerelle satellite transfère ces paquets par l'intermédiaire du tunnel source (Figure 47). Dans cette solution, le mécanisme est déporté vers la passerelle satellite pour créer le *buffer*, il faut donc connaître l'identifiant du tunnel source. Cette information est contenue dans le message *Handover Required*. La création du *buffer* est donc mise en œuvre lors du transfert de ce message par la passerelle satellite.

Pour minimiser les pertes, la solution de *buffer* dans la passerelle satellite semble la plus intéressante, puisqu'elle permet de s'affranchir de toute modification du standard dans la phase de préparation. Cependant, la passerelle satellite a besoin d'informations sur le *handover* en cours, comme les numéros de *bearer* à transférer ainsi que ceux du tunnel source. Ces informations peuvent être obtenues si la passerelle satellite parvient à décoder les messages de contrôle sur l'interface S1, en particulier le message *Handover Required*.

4.4.1.3 La gestion du *buffer* avant le déclenchement du *Forward*

Le stockage des paquets dans le *buffer* de *Forward* a besoin d'être optimisé. Le déclenchement du *Forward*, lorsque que le *buffer* commence à se vider, ne survient que lors de la phase de terminaison. De sa mise en œuvre, pendant la phase de préparation jusqu'au début de la transmission des paquets de *Forward*, il peut donc devoir gérer beaucoup de paquets. En outre, la préparation à double décision permet d'anticiper la phase de préparation, mais, en contrepartie, le début de l'exécution n'est plus corrélé avec la signalisation de la préparation et, dans la plupart des cas, il va allonger la durée de la période de stockage du *buffer*. Par conséquent, la gestion du *buffer* doit prendre en compte ces paramètres et éliminer les paquets qui sont reçus depuis longtemps par l'UE, dans le but de minimiser sa taille. Pour ce faire, nous proposons une estimation grossière du temps de transmission d'un paquet entre la passerelle satellite et l'UE ($T_{TxSat \rightarrow UE \text{ Est}}$). On ajoute à celle-ci

la durée entre le début de la phase d'exécution et le début du déclenchement du *Forward* qui est défini dans le paragraphe suivant. L'estimation générale ($T_{EstGénéral}$) doit être légèrement supérieure à la valeur réelle pour éviter que ce mécanisme n'induisse des pertes. Le *buffer* va supprimer les paquets dont le temps de stockage dépasse cette estimation ce qui permet de conserver une taille du *buffer* acceptable.

4.4.1.4 Le déclenchement du Forward avec buffer

4.4.1.4.1 L'instant de déclenchement

Le déclenchement du transfert des paquets contenus dans le *buffer* de *Forward* est effectué lors de la phase de terminaison. Cependant, les échanges de messages pendant celle-ci ne permettent pas d'informer la passerelle satellite du succès du *handover* et de l'identité de l'eNB cible, parmi les différents eNB préparés. Dans ces conditions, la création d'un message est obligatoire. Dans l'optique d'affecter le moins possible le segment terrestre, ce message que l'on nomme *Forward Activation* sera envoyé entre la MME source et la passerelle satellite (A dans la Figure 48).

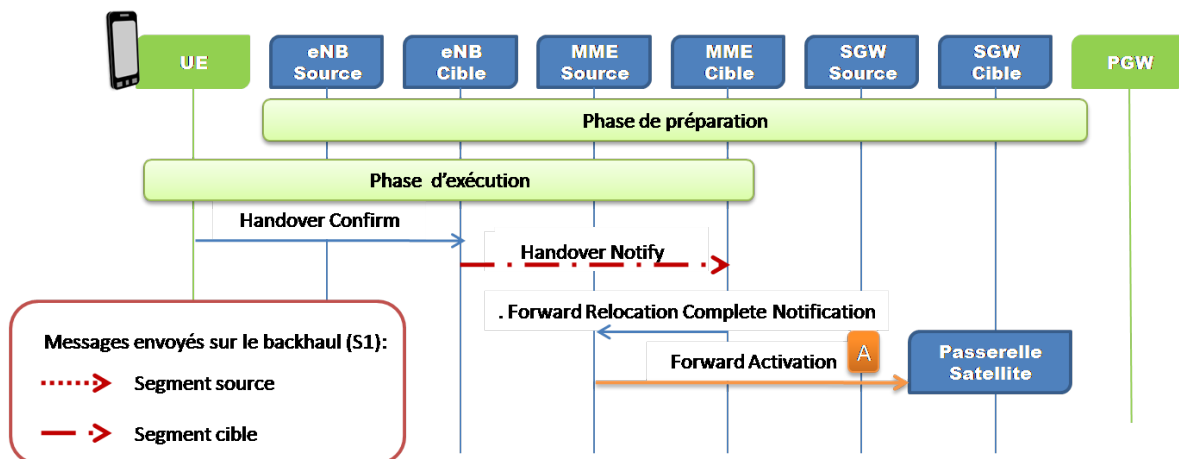


Figure 48 Déclenchement du Forward par le message *Forward Activation*

4.4.1.4.2 La duplication

La duplication est un problème important pour notre mécanisme de *Forward*. En effet, la passerelle satellite n'a aucune connaissance sur le dernier paquet envoyé à l'UE. Il faut donc développer une méthode pour l'en informer. Nous proposons trois solutions :

- **Une estimation précise de l'instant, où le dernier paquet reçu par l'UE a été envoyé par la passerelle satellite.** Cette estimation représente le même paramètre $T_{EstGénéral}$; c'est-à-dire la durée de transmission d'un paquet de la passerelle satellite à l'UE, à laquelle nous ajoutons la durée de la phase d'exécution et le temps entre l'envoi du *Handover Confirm* et la réception du *Forward Activation* par la passerelle satellite (Temps entre l'instant A et B dans la Figure 49). Ainsi, tous les paquets dont le temps de stockage est supérieur à cette estimation sont supprimés, à la réception du message *Forward Activation* et juste avant le début du transfert vers l'eNB cible. La précision de cette estimation, combinée avec le débit des flux, va donc influencer le nombre de pertes ou de duplications de paquets. Si cette estimation est inférieure à la valeur réelle, cela engendre des pertes et, au contraire, si elle est supérieure, elle implique des

transmissions de paquets déjà reçus par l'UE. Il est impossible d'avoir une estimation parfaite. Cependant, des algorithmes SON peuvent permettre d'améliorer la précision de cette estimation.

- **Un container transparent transmis par l'UE.** Puisque l'on ne peut pas envoyer le numéro de séquence du dernier paquet reçu par l'UE par l'intermédiaire du réseau de cœur, nous proposons de le transmettre par l'interface radio. Ainsi, nous ajoutons un container transparent au message *RRC Connexion Reconfiguration* qui n'est pas traité par l'UE, mais transmis à l'eNB cible dans le message *Handover Confirm* (Figure 49). Cette technique modifie de manière importante le standard, puisqu'à la fois la signalisation sur le réseau et les spécifications pour l'UE sont modifiées.
- **L'utilisation d'un *offset* entre les numéros de séquence PDCP et GTP.** Pour cette solution, on récupère le principe proposé dans les discussions du 3GPP (NTT DoCoMo, Inc., August 2007) (Ericson, August 2007) (NEC, August 2007) pour le réordonnancement des paquets après le *Path Switch* (Paragraphe 4.3.3). L'eNB source transmet donc la valeur de l'écart entre les numéros de séquence pendant la phase de préparation. Après l'exécution du *handover*, l'UE envoie le contexte PDCP par l'intermédiaire du message *PDCP Status Report* présent dans le standard (The 3rd Generation Partnership Project (3GPP), September 2012). Ainsi, l'eNB cible retrouve le numéro de séquence GTP du dernier paquet reçu par l'UE et transmet cette information, par l'intermédiaire de la signalisation de la phase de terminaison et du message *Forward Activation*. Cette méthode engendre l'ajout d'informations pendant la phase de préparation, avec l'*offset* entre l'eNB source et l'eNB cible, mais aussi pendant la phase de terminaison au niveau des messages entre l'eNB cible et la passerelle satellite, avec le numéro de séquence du dernier paquet reçu par l'UE.

Pour assurer un mécanisme sans perte, l'unique solution qui nous semble envisageable est donc l'utilisation d'un *offset* entre les numéros de séquence PDCP et GTP ; seul celui-ci assure un *handover* sans perte tout en conservant le standard sur l'interface radio. Cependant, si aucune modification de la signalisation n'est permise dans le segment terrestre, la solution de l'estimateur sera privilégiée avec comme inconvénient d'engendrer quelques pertes ou duplications.

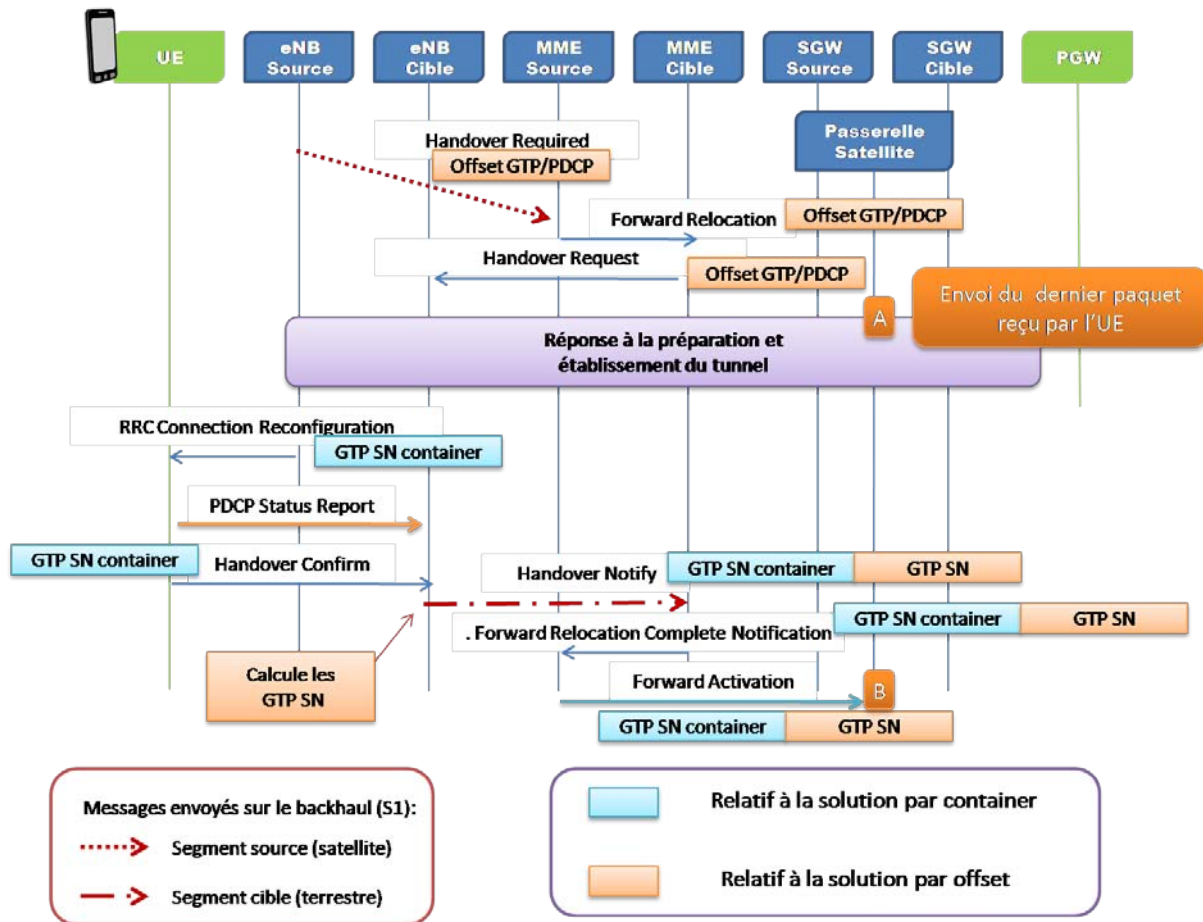


Figure 49 Mécanismes contre la duplication de paquets

4.4.1.5 La gestion du buffer et le Path Switch

Le dernier problème soulevé par le mécanisme de *Forward bi-cast par buffer* se situe dans la gestion du *buffer* pendant son *Forward* et, plus particulièrement, juste après le *Path Switch*. La Figure 50 représente le mécanisme de *Forward* avant le *Path Switch* et la Figure 51 après. L'eNB cible va recevoir des paquets, à la fois en provenance du *buffer* par les tunnels mais aussi par le chemin cible (Figure 51). La nouveauté, comparée à un *handover* terrestre, est la quantité de données qui doit être transmise par les tunnels. L'eNB cible qui appartient au segment terrestre n'est pas configuré pour recevoir un aussi grand nombre de paquets transférés. Si le débit de transmission de la passerelle satellite sur le tunnel source est égal à celui de réception sur le chemin source, la durée mise pour envoyer tous les paquets sera supérieure au délai induit par le satellite. La gestion du réordonnement par l'eNB cible, expliqué dans la Figure 42, devient beaucoup plus longue que pour un *handover* terrestre.

Dans notre procédure, le *Path Switch* est consécutif au début du *Forward*, puisque tous les deux sont déclenchés pendant la phase de terminaison. Par conséquent, après le *Path Switch*, le *buffer* transmet les paquets par les tunnels et l'eNB cible va stocker les paquets en provenance du chemin cible dans une file d'attente ($F_{CheminCible}$), et ce jusqu'à ce que la file d'attente des paquets du tunnel cible ($F_{TunnelCible}$) soit vide et que le mécanisme de réordonnement décide de basculer vers les paquets du chemin cible. Dans notre cas, la période pendant laquelle $F_{CheminCible}$ se remplit

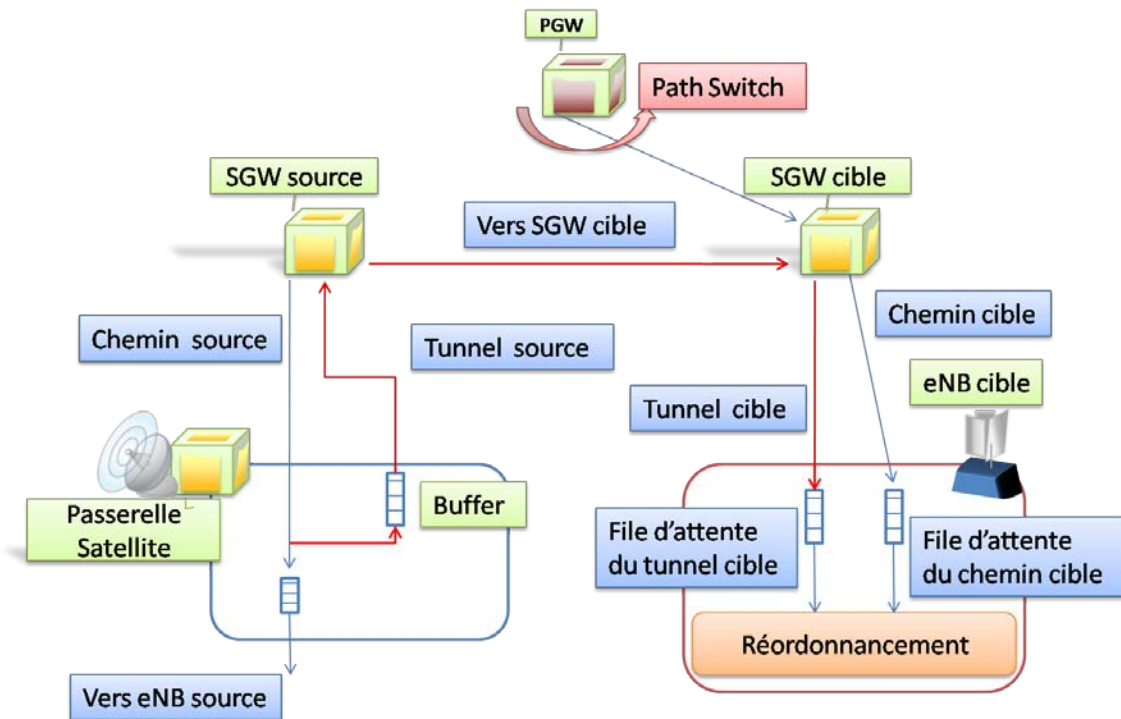
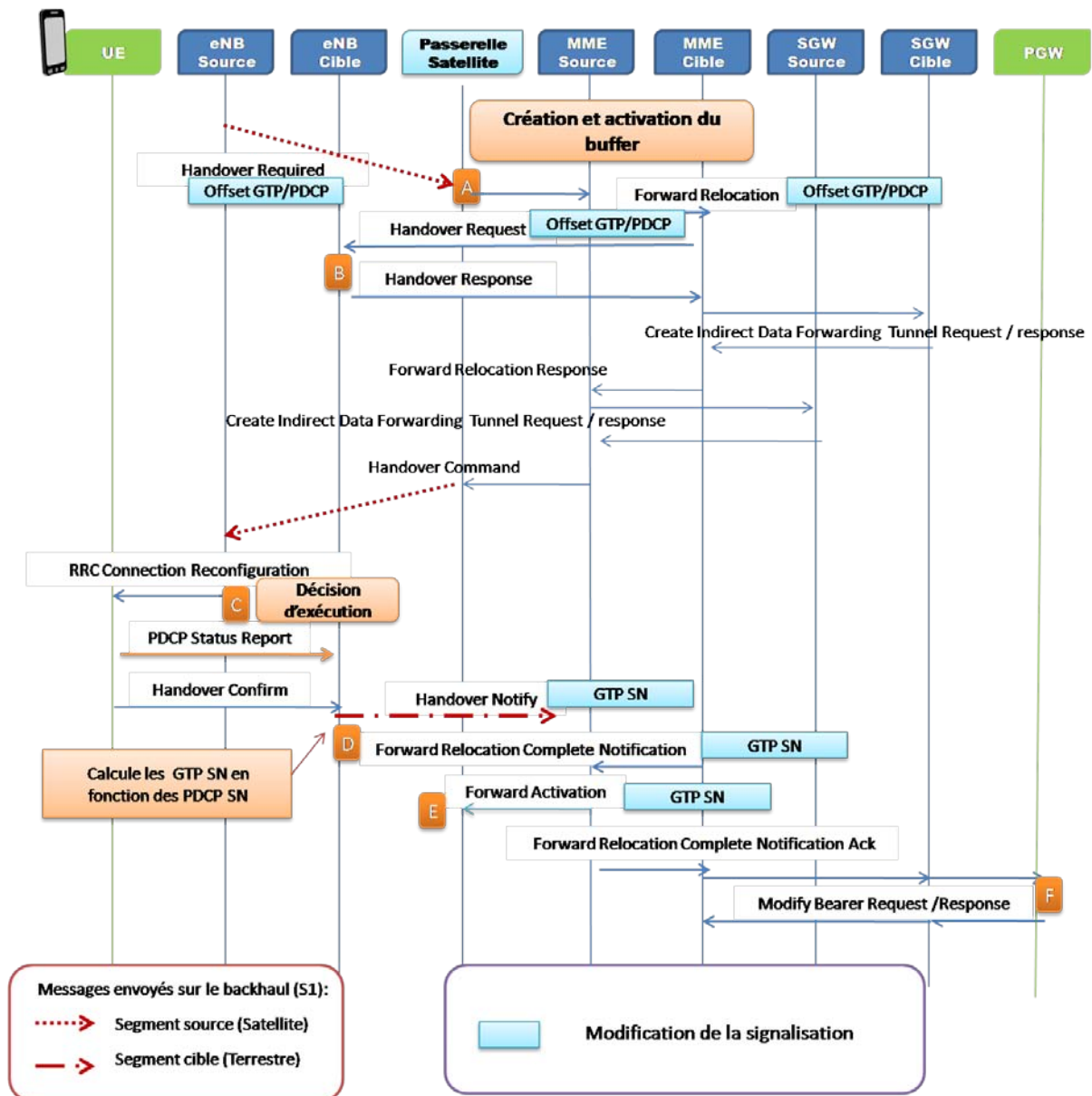


Figure 51 Le mécanisme de Forward après le Path Switch

4.4.1.6 Récapitulatif de la procédure optimisée

La Figure 52, ci-dessous, agrémentée de son explication permet d'avoir un récapitulatif du mécanisme de Forward proposé pour un *handover* inter-satellite-terrestre.

Figure 52 Procédure complète de l'optimisation du *handover* inter-segment-satellite

- La passerelle satellite reçoit le message *Handover Required*. Il contient les identifiants du tunnel source, ce qui permet à celle-ci de créer le *buffer* de *Forward*. Elle effectue une copie des paquets contenus dans ses files d'attente et les stockent dans le *buffer* de *Forward*. Dans le même temps, elle remplace l'adresse IP destination par l'adresse de la SGW ainsi que l'identifiant du *bearer*, en mettant celui du tunnel source. La gestion des paquets du *buffer* débute après cela.
- Le message *Handover Required* transporte l'*offset* entre les numéros de séquence GTP et PDCP. Il est transféré jusqu'à l'eNB cible par les messages *Forward Relocation* et *Handover Request*. Le reste de la procédure n'est pas modifié.
- La préparation à double décision permet d'attendre que l'UE soit toujours apte à changer de station de base. Juste avant la fin de la phase d'exécution, l'UE envoie le *PDCP Status Report* pour informer l'eNB cible des derniers paquets reçus par l'UE.

- D. L'eNB cible calcule les numéros de séquence GTP des derniers paquets reçus par l'UE grâce à l'*offset* et au *PDCP Status Report*. Il transmet cette information jusqu'à la MME source par l'intermédiaire des messages *Handover Notify* et *Forward Relocation Complete Notification*.
- E. La création du message *Forward Activation* permet à la passerelle satellite d'activer le *Forward* des paquets contenus dans son *buffer*. Nous profitons de ce message, pour transmettre les numéros de séquence GTP des derniers paquets reçus pour chaque *bearer* ; cela va permettre de supprimer les paquets dupliqués. Le débit auquel nous envoyons les paquets du *buffer* est calculé de manière à minimiser le temps de *Forward*.
- F. La PGW réalise le *Path Switch* normalement.

4.4.2 Optimisation du *handover* intra-satellite par multicast

4.4.2.1 Principe

Cette solution s'adresse aux *handovers* intra-segment-satellite pour lesquels aucune relocalisation de SGW et MME n'est nécessaire. Son principe repose sur la capacité intrinsèque des systèmes satellite à transmettre du trafic *broadcast* et *multicast*. Premièrement, les segments source et cible engendrent des délais similaires, puisqu'ils sont transmis tous les deux par le système satellite. Deuxièmement, la surconsommation des ressources, provoquée par la combinaison du *bi-cast* et de la préparation multiple, n'est plus vraie grâce à la prédilection du satellite pour le *broadcast* et le *multicast*. Ces deux contraintes, présentes pendant le *handover* inter-satellite-terrestre et absentes dans le *handover* intra-satellite, impliquent que le *Forward* par *buffer* n'est plus nécessaire et que seul un *Forward multicast* est suffisant et adapté pour le *handover* intra-satellite.

Le mécanisme de *Forward* transmet les paquets, sujets au *Forward* en *multicast*, sur le lien satellite et vers tous les eNB préparés. Ce *Forward multicast* est intéressant, puisqu'il optimise aussi la récupération de connexion après un RLF. En effet, la reconnexion se fera très probablement avec un eNB déjà préparé qui reçoit de manière continue les paquets transférés. Ainsi, ce mécanisme de *Forward* minimise les pertes de données dues à un RLF. La principale contrainte, à laquelle nous nous astreignons pour toutes les optimisations, est la minimisation de l'impact des modifications sur le standard et l'implantation des algorithmes dans les entités.

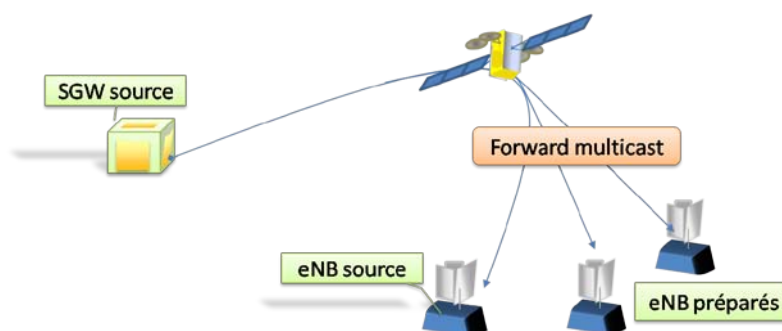


Figure 53 Principe du *Forward multicast* pendant un *handover* intra-satellite

4.4.2.2 Mise en place du mécanisme de Forward

4.4.2.2.1 La mise en place simple

Pendant l'établissement du mécanisme de *Forward*, l'opération principale est l'attribution des adresses *multicast*. Cette étape est réalisée pendant la préparation, en particulier, à la réception du message *Handover Required* par la MME (A dans la Figure 53). En effet, celle-ci a une connaissance de toutes les adresses *multicast* déjà utilisées dans d'autres procédures de *handover*. Pour les mêmes raisons, elle va choisir les nouveaux identifiants des *bearers* (TEID) pour le flux *multicast*. Puis, elle transmet ces informations à l'eNB cible par l'intermédiaire du message *Handover Request*. Le *Forward multicast* est activé lors de la création des tunnels de *Forward* (B dans la Figure 53).

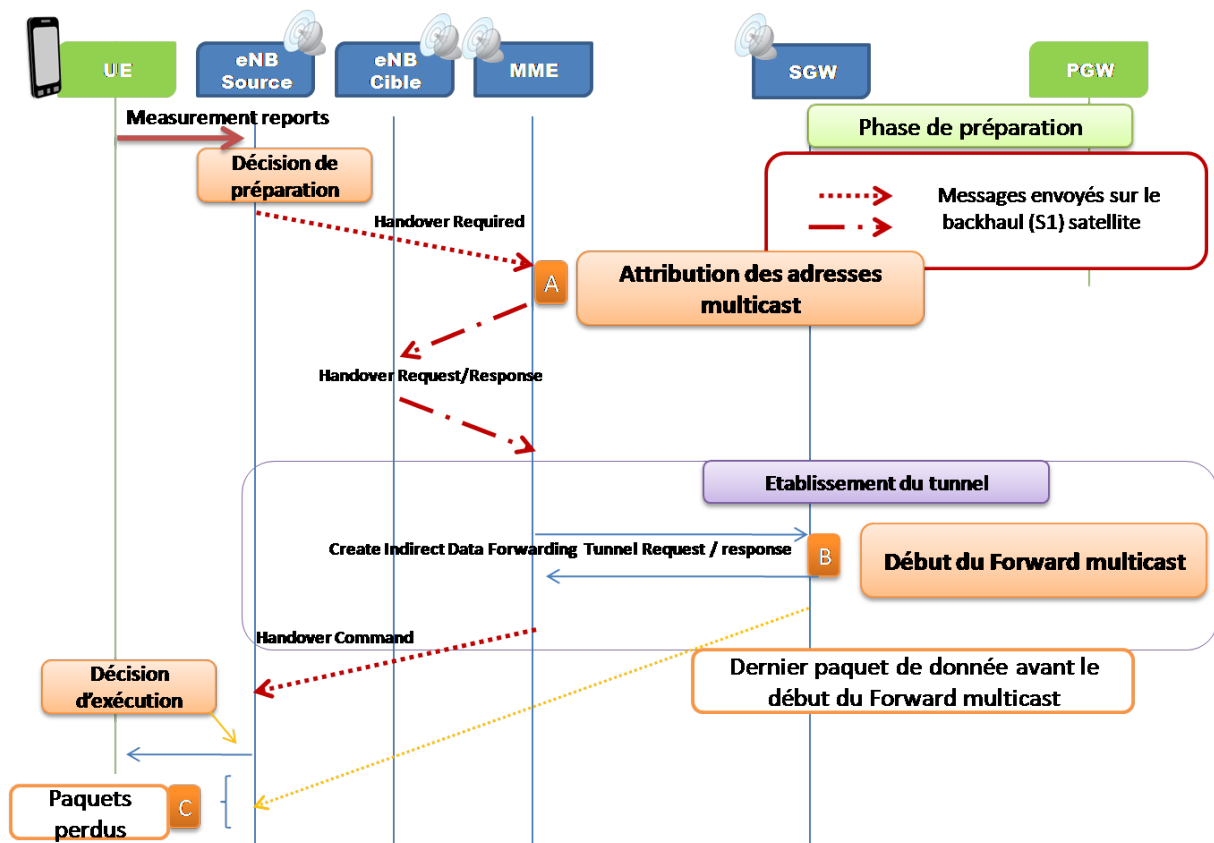


Figure 54 Mise en place simple du mécanisme de Forward

Le même problème, que l'on retrouve dans le cas du *Forward* bi-cast par Buffer (paragraphe 4.4.1.2), survient dans ce scénario, puisque la mise en place du *multicast* est trop tardive. Ceci peut entraîner des pertes des paquets contenus dans les files d'attente de la passerelle satellite. La solution choisie pour le mécanisme par *buffer* se fonde sur l'application du mécanisme de *Forward* sur les paquets contenus dans la file d'attente satellite. Cependant, dans le *handover* intra-satellite, une solution plus simple peut être proposée, grâce à l'anticipation du début du *Forward multicast*. Dans le mécanisme par *buffer*, l'intérêt de cette solution était mince puisque le gain possible était très faible. Pour le *handover* intra-satellite, le gain peut être bien plus significatif puisque le segment cible est réalisé par le système satellite. Par conséquent, le flux *multicast* va débuter lors de la réception du message *Handover Required*.

Cependant, aucun message de signalisation n'est transmis à la SGW pour lui ordonner de commencer le *Forward multicast*. La solution que nous proposons permet de s'affranchir de l'ajout de messages de signalisation et repose sur l'utilisation de la passerelle satellite comme point de départ des flux *multicast*. Cette entité transmet le message *Handover Required* à la MME. Ainsi, à cet instant, elle peut commencer à envoyer les flux *multicast*.

La configuration des eNB doit être réalisée avant l'activation du *Forward multicast* pour qu'il puisse décoder le flux *multicast*. Ainsi, nous proposons deux solutions. La première attribue à chaque eNB une adresse *multicast* statique. La configuration *multicast* est donc toujours valable. Cette solution ne permet pas de différencier les *handovers* de différentes UE par l'adresse *multicast*. Ainsi les eNB préparés par le même eNB source, mais pour un UE différent vont avoir la même adresse *multicast*. Pour résorber ce problème, la deuxième solution fournit un groupe d'adresses *multicast* à un eNB.

Dans le but de minimiser l'impact de ce nouveau mécanisme sur l'eNB, ce sont les terminaux satellite qui vont prendre en charge les flux *multicast*. Ils transposeront les flux *multicast* en flux unicast à destination des eNB. Ainsi, la passerelle va transformer le flux en direction de l'eNB source, en flux *multicast* destiné à l'eNB source et aux eNB préparés (A dans la Figure 55). Les terminaux satellite, quant à eux, vont transformer les flux *multicast* en flux unicast. Dans ce cas, le TEID reste inchangé pour l'eNB source (B dans la Figure 55) et sa valeur est modifiée pour les eNB préparés (C dans la Figure 55). Ce changement permet d'adapter le TEID à chaque eNB, en mettant celui du tunnel de Forward cible. Le TEID va permettre à l'eNB d'identifier l'UE de destination et le type de QoS attribué à ce *bearer*. Dans notre cas, les paquets sont envoyés en *multicast* et donc la valeur du TEID est la même pour les paquets à destination de l'eNB source et de l'eNB cible. Pour éviter les conflits de TEID sur les eNB, le TEID choisi pour le paquet *multicast* est celui du *bearer* source. C'est le terminal satellite de l'eNB cible qui, grâce au couple adresse *multicast* et TEID source, peut remplacer le TEID source par le TEID cible (C dans la Figure 55).

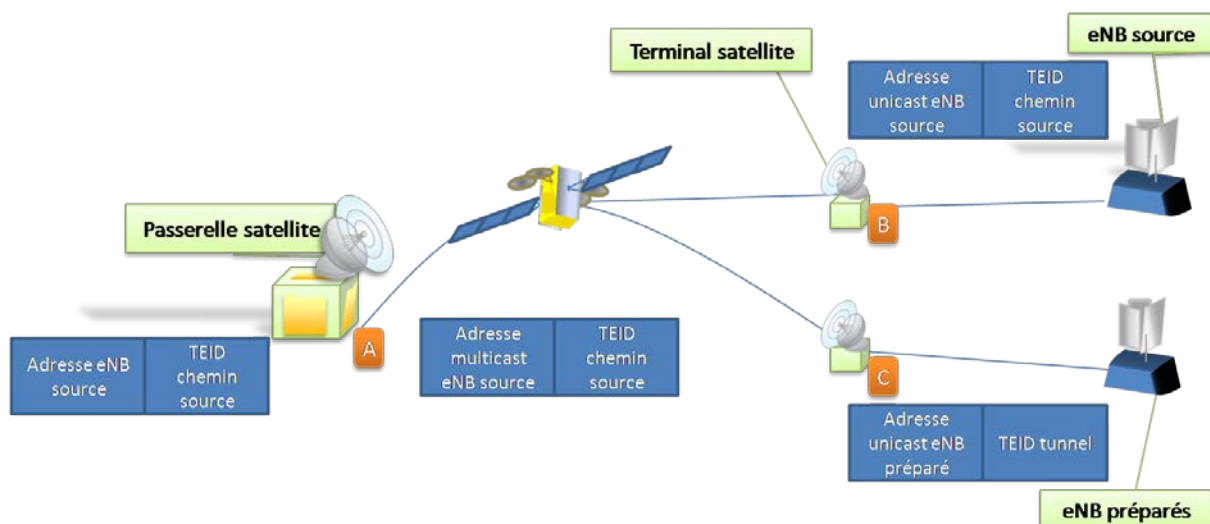


Figure 55 Implications de la passerelle et des terminaux satellite dans le mécanisme de Forward

4.4.2.3 La gestion du Forward pendant avant la fin de la phase d'exécution

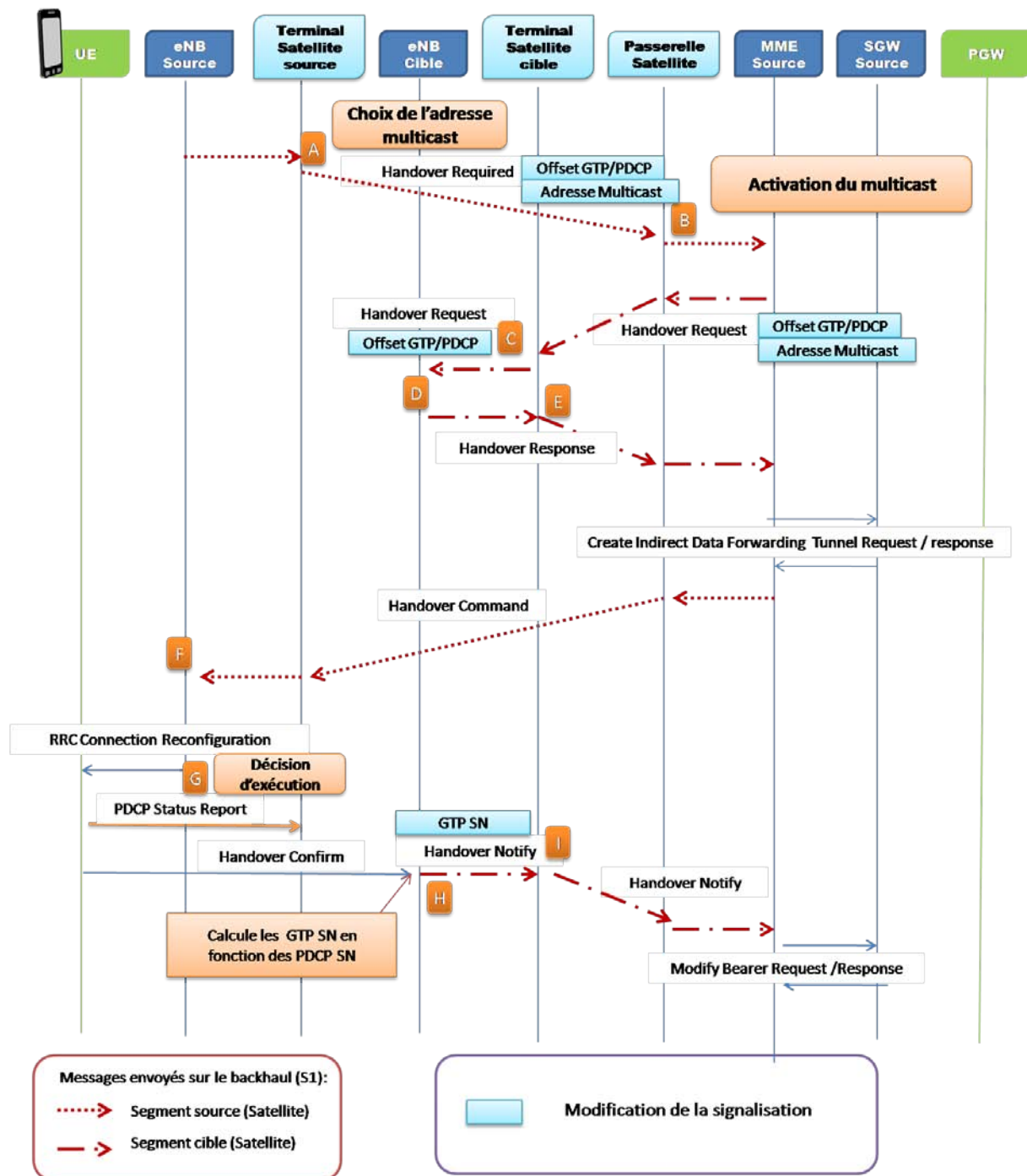
Les paquets sont envoyés à tous les terminaux satellite des eNB préparés. Il est nécessaire qu'ils gèrent l'arrivée des paquets de manière appropriée. Nous considérons donc qu'ils doivent supprimer les paquets en provenance du *multicast* jusqu'à l'attachement de l'UE à la nouvelle eNB. Au début de la phase de terminaison, le terminal de l'eNB cible va interrompre la suppression et transmettre les paquets à leur eNB. Cette solution va donc engendrer des pertes des paquets pendant la phase d'exécution, lorsque l'UE n'est connecté ni à l'eNB source ni à l'eNB cible. Pour contrecarrer ces pertes, chaque terminal met en place un *buffer* qui conserve les paquets pendant une période correspondant à la durée moyenne de la phase d'exécution. Cette période est obtenue par une estimation grossière et légèrement supérieure de la durée réelle, dans le but d'éviter les pertes de paquets. L'utilisation de cette technique implique des duplications de paquets.

4.4.2.4 La duplication de paquets

Les trois solutions proposées dans le cadre du *Forward bi-cast* par *buffer* sont transposables dans le cadre du mécanisme de *Forward multicast*. Le transfert de l'*offset* pendant la phase de préparation, combiné à l'envoi de *PDCCP Status Report*, permet d'avoir un mécanisme précis pour lutter contre les duplications de paquets, grâce à l'ajout de l'information d'*offset* dans les messages de la phase de préparation. Cette technique n'affecte que les messages envoyés au travers du segment satellite. Le transfert d'un container transparent, contenant les numéros de séquence GTP des derniers paquets reçus par l'UE, est tout aussi gênant, car il implique une modification du standard pour l'UE. La solution fondée sur l'estimation engendre moins de difficultés que pour le *handover* inter-satellite-terrestre, car celle-ci ne concerne que la durée de la phase d'exécution. Cependant, elle reste toujours imparfaite et ne peut assurer, de manière absolue, l'absence de perte ou de duplication.

4.4.2.5 Récapitulatif de la procédure optimisée

La Figure 56, ci-dessous, agrémentée de son explication, permet d'avoir un récapitulatif du mécanisme de *Forward* proposé.

Figure 56 Procédure complète de l'optimisation du *handover* intra-satellite

- Lorsque le terminal satellite de l'eNB source reçoit le message *Handover Required*, il choisit l'adresse *multicast* qui sera utilisée par la suite. Il ajoute cette adresse au message de signalisation, qui contient déjà l'information sur l'écart entre les numéros de séquence PDCP et GTP fournis par l'eNB source.
- La passerelle satellite reçoit à son tour le message *Handover Required* et en profite pour configurer le flux *multicast*. Elle configure le flux *multicast* à partir de cet instant. Aucun autre eNB n'est à ce moment-là capable de recevoir ces flux.
- Le terminal satellite de l'eNB cible reçoit le message *Handover Request* qui contient l'adresse *multicast*. A partir de cet instant, le terminal peut recevoir les paquets transférés.

- D. L'eNB cible récupère l'*offset* et définit l'identifiant du tunnel cible.
- E. L'identifiant du tunnel cible est envoyé dans le *Handover Response* ; il est lu par le terminal satellite cible pour paramétrer le *Forward* vers l'eNB cible. Il va stocker les flux transférés dans un *buffer*, qui conserve les paquets pendant une durée égale à l'estimation de celle de la phase d'exécution.
- F. La signalisation de la phase de préparation se termine.
- G. La phase de préparation à double décision permet de choisir le meilleur eNB pour exécuter le *handover*. Après l'établissement de la connexion avec le nouvel eNB, l'UE transmet les numéros de séquence PDCP des derniers paquets qu'il a reçus.
- H. L'eNB cible calcule les numéros de séquence GTP des derniers paquets reçus par l'UE.
- I. Les numéros de séquence GTP de ces paquets sont transmis au terminal satellite de l'eNB cible pour éliminer les paquets dupliqués.

4.4.3 Minimisation de l'impact des solutions sur la signalisation

Certaines solutions proposées dans ce chapitre impliquent d'ajouter des informations aux messages de signalisation. Nous pouvons citer en exemple, les valeurs d'*offset* entre les numéros de séquence GTP et UDP ainsi que l'adresse *multicast* pour le mécanisme de *Forward* pendant le *handover* intra-satellite (Figure 56). Dans l'intention de minimiser l'impact de ces solutions, nous souhaitons garder une compatibilité entre les équipements implantant ou non notre optimisation. La première solution fut de concentrer au maximum les modifications et les mécanismes à l'intérieur du système satellite et, en particulier, dans ses équipements : le terminal et la passerelle satellite.

La deuxième solution consiste à utiliser les extensions des messages de signalisation qui sont offertes par les protocoles du 3GPP. En effet, GTP-C offre la possibilité d'ajouter des champs d'information. En particulier, nous pouvons utiliser une *Private Extension* (Figure 57).

	Bits							
Octets	8	7	6	5	4	3	2	1
1	Type = 255 (decimal)							
2 to 3	Length = n							
4	Spare				Instance			
5 to 6	Enterprise ID							
7 to (n+4)	Proprietary value							

Figure 57 Private extension du protocole GTP-C

Elle est spécifique à une entreprise et elle est optionnelle. Par conséquent, si l'équipement ne prend pas en charge notre optimisation, le *handover* ne sera donc pas invalidé et pourra continuer sans anicroches.

Après cette phase de définition et de choix des mécanismes, nous allons procéder à leur évaluation.

Chapitre 5 : Réseaux hybrides LTE-satellite : Étude des optimisations des mécanismes de *Forward*

Ce chapitre a pour objectif d'estimer l'intérêt du mécanisme de *Forward* que nous proposons. Notre solution possède, indéniablement, l'avantage de préserver les ressources du lien satellite, en supprimant l'utilisation du tunnel de *Forward* entre l'eNB source et la SGW. En outre, dans le cadre du mécanisme de *Forward* proposé pour le *handover* intra-satellite, l'utilisation du *multicast* accentue l'économie de ressources satellite par rapport au mécanisme de *Forward* standard. Cependant, un investissement significatif est nécessaire. Ce chapitre a pour but de comparer les performances des applications pendant un *handover*, lorsque le mécanisme de *Forward* est établi et lorsqu'il est inexistant. En effet, l'utilisation du mécanisme de *Forward* dépendra du rapport entre l'efficacité et le coût de cette technique.

Cette étude porte uniquement sur le mécanisme de *Forward*. Elle n'est pas orientée vers l'évaluation des optimisations de la phase de préparation ; la préparation à double décision (Chapitre 3) n'est pas simulée, en raison d'une complexité trop importante des paramètres qu'il faut prendre en compte et en particulier ceux relatifs à la couche physique. En effet, ce sont, principalement, la force et la qualité du signal, conjointement avec les paramètres de mesures qui sont responsables du *handover*. Les équipements nécessaires pour obtenir des représentations acceptables de ces paramètres sont trop onéreux. Le mécanisme de *Forward* a un impact sur les couches hautes en particulier les couches transport et application ; les couches basses ne l'affectent que très légèrement, puisque le mécanisme de *Forward* ne s'intéresse qu'à la couche PDCP et non à celles en dessous.

Dans un premier temps, nous avons souhaité effectuer des simulations de la procédure du S1 *handover*, du déclenchement de celle-ci jusqu'à sa terminaison. Le but de ces simulations est d'observer le comportement du protocole TCP lors du *handover* et, en particulier, sa réaction face aux différents mécanismes de *Forward* possibles. Dans le cas du *handover* inter-satellite-terrestre, nous comparons l'utilisation du mécanisme de *Forward* par *buffer* proposé dans le chapitre 4 avec l'absence de mécanisme de *Forward*. Dans le cas du *handover* intra-satellite, nous comparons deux mécanismes de *Forward* proposés dans le chapitre 4 : le *Forward multicast* sans perte ni duplication et le *Forward multicast* avec de légères pertes. Pour ce dernier, ces pertes sont une conséquence de l'interruption de la connexion, lors de la phase d'exécution. Pour observer l'impact de ces mécanismes, nous souhaitons effectuer des simulations pour les différents états de TCP tels que le *slow-start* et la *congestion avoidance*.

La première partie met en avant le manque de simulateur du réseau de cœur qui nous a menés à implanter, dans le simulateur Network Simulator 3, les protocoles et procédures nécessaires à la simulation du *handover*. Ensuite, nous nous sommes rendu compte que l'utilisation de NS3 ne nous fournissait pas des résultats significatifs, à cause de la différence importante entre l'implantation des protocoles TCP dans le simulateur et sur des équipements réels. Ceci nous a donc conduit à effectuer de l'émulation, en utilisant la vraie pile protocolaire de Linux pour émuler les flux applicatifs et en réutilisant notre implantation du réseau de cœur dans NS3 en temps réel. Nous décrivons la

première implantation effectuée. En revanche, les résultats qui sont présentés dans la dernière partie de ce chapitre sont le résultat de cette émulation.

5.1 Les simulateurs

Très peu de simulateurs sous licence libre étaient disponibles au début de ce travail. En outre, les seuls qui sont actuellement présents n'en étaient qu'à leurs balbutiements. Nous les présentons succinctement ci-dessous.

LTE-SIM

LTE-SIM (LTE simulator) est un simulateur événementiel spécifique à LTE (LTE-SIM) (Piro, Grieco, Boggia, Capozzi, & Camarda, 2011). Tout d'abord centré sur l'interface radio, différentes fonctions ont été ajoutées successivement telles que la gestion de la QoS et la mobilité de l'utilisateur. La première version stable fut disponible fin 2010. L'EPC n'est pas complètement implanté, car ce simulateur se concentre sur l'UE et l'eNB. Une entité qui contrôle le réseau est implantée et représente la MME/SGW.

Simulateurs LTE de l'Université technologique de Vienne.

L'Université technologique de Vienne propose plusieurs simulateurs (Université Technologique de Vienne). Ils reposent sur MATLAB et proposent des simulateurs représentant l'interface radio. Ils séparent la partie lien montant et descendant et n'intègrent pas de fonctionnalités de gestion des *handovers* (Mehlführer, Wrulich, Ikuno, Bosanska, & Rupp, August 2009) (Ikuno, Wrulich, & Rupp, May 2010).

Projet LENA

Le projet LENA (LTE-EPC Network Simulator) est un projet qui a pour but d'implanter les principales entités définies par le 3GPP (LENA Project). Il s'intègre au simulateur à événement discret Network Simulator 3 (NS3) (Baldo, Miozzo, Requena, & Guerrero, July 2011). Il repose en particulier sur les deux précédentes implantations de simulateur LTE. La couche physique se fonde sur celle décrite par LTE-SIM, alors que le modèle d'erreur de la couche physique est similaire à celui mise en œuvre par l'Université de Vienne. Ce simulateur est le plus complet et permet de réutiliser toutes les implantations d'autres protocoles déjà réalisées dans NS3, en particulier pour les couches transport et application. La première version fut disponible en mars 2011. C'est uniquement au début de l'année 2013 que la procédure du X2-*handover* a été implantée.

Le mécanisme de *Forward* est proposé dans le but, de préserver les ressources satellites et d'améliorer les performances des applications lors d'un *handover*. L'efficacité de notre proposition est intimement liée à la couche de transport TCP. En particulier, les mécanismes qui permettent d'éviter les pertes sont mis en place pour optimiser les performances de TCP et pour s'adapter à sa gestion des pertes de paquets. En outre, l'utilisation de TCP sur des liens satellite pose problème en raison du long délai induit par la transmission satellite. Par conséquent, les deux premiers simulateurs ne pouvaient être considérés comme une solution, pour simuler notre mécanisme de *Forward*. Le troisième n'étant pas disponible, nous nous sommes tournés vers une implantation simplifiée de

l'EPC et de l'interface radio que nous avons nous-même effectuée. Nous avons donc utilisé NS3 dans le but de récupérer les couches transport et application.

5.2 Implantation

5.2.1 La simulation

L'objectif de ces simulations est d'observer le comportement des flux TCP, lors de *handovers* inter-satellite-terrestre ainsi qu'intra-satellite et, en particulier, leur interaction avec le mécanisme de *Forward*. Les implications des couches basses nous semblent insuffisantes pour justifier dans un premier temps leur implantation. Nous avons donc décidé de grandement simplifier l'interface radio de LTE et de ne garder qu'une version simplifiée du protocole PDCP. Les protocoles RLC, MAC et la couche physique sont ainsi ignorés. Bien que ces protocoles aient une réelle importance lors du *handover* et en particulier sur son déclenchement, l'impact de ces couches nous semble négligeable sur le mécanisme de *Forward*, en raison de la supériorité de l'ordre de grandeur entre le délai du lien satellite et celui de l'interface radio.

5.2.1.1 Le plan utilisateur

Le plan utilisateur est représenté dans la Figure 58. Il est constitué principalement du protocole GTP-U qui a été implanté en adéquation avec la spécification du 3GPP. Seuls certains champs spécifiques ne l'ont pas été, car ils ne sont pas utilisés lors du *handover*, tels que les extensions *Echo Request* et *Echo Response* qui permettent de vérifier si le lien est toujours actif.

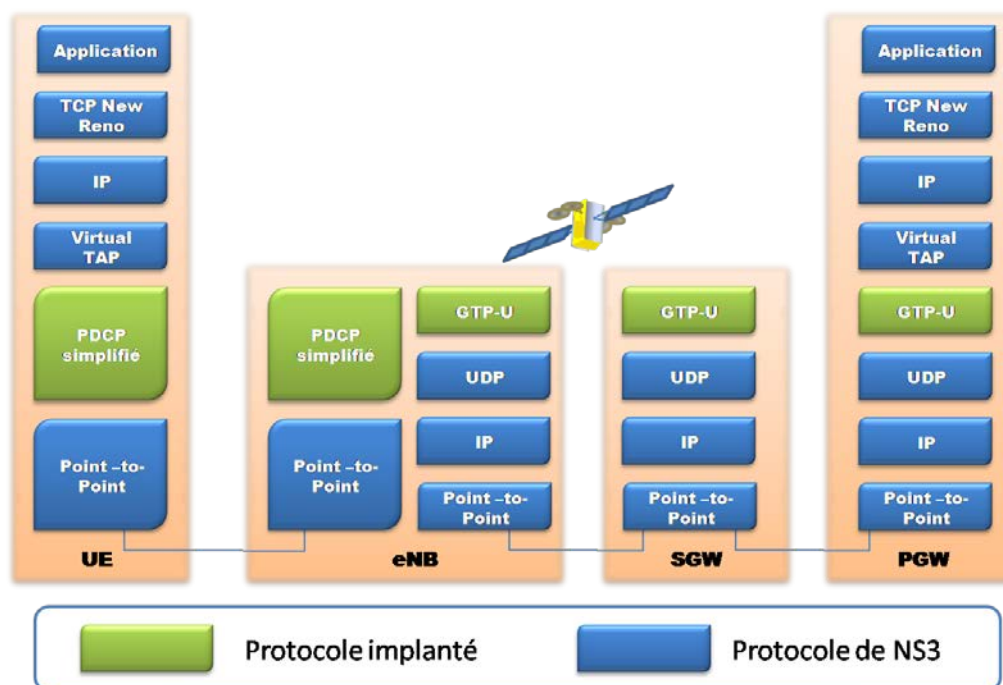


Figure 58 Implantation des couches protocolaires sur le plan utilisateur

Le protocole PDCP est très largement simplifié puisque les couches basses ne sont pas implantées. Il permet de transporter l'information avec des numéros de séquence et contient l'identifiant du *radio bearer*. Il n'y a aucun mécanisme de sécurité mis en place.

Le *Virtual Net Device* est une couche qui permet d'ajouter une pile protocolaire TCP/IP, au-dessus du protocole PDCP dans l'UE et au-dessus de GTP-U dans la PGW. C'est une couche spécifique à NS3 qui sert juste de passerelle entre PDCP et IP ou GTP-U et IP. Elle n'a pas d'impact sur le plan utilisateur ni sur le plan de contrôle.

Le protocole de transport utilisé est TCP New Reno implanté " nativement " dans NS3. L'application, quant à elle, est un transfert de fichier.

Le réseau de cœur est modélisé par des liaisons point à point avec différents paramètres en fonction de la nature du lien : satellite, filaire et radio. Ils sont définis dans le Tableau 19. Au-dessus, les entités sont interconnectées grâce à une pile UDP/IP. Ces capacités sont inférieures à celles utilisées dans la réalité, le étant but de prendre en compte le partage des ressources du réseau avec d'autres utilisateurs. Nous considérons donc qu'une dizaine d'utilisateurs se partagent équitablement ces ressources.

	Réseau de cœur (Filaire)	Interface Radio	Satellite (Backhaul)
Délai	2ms	5ms	300ms
Capacité	100Mbps	10Mbps	1Mbps

Tableau 19 Paramètres des liaisons point à point

5.2.1.2 Le plan de contrôle

L'implantation du plan de contrôle est détaillée dans la Figure 59. L'architecture décrite pour le plan de contrôle est conservée. Seuls les protocoles S1-AP, RRC et GTP-C diffèrent. Nous ne tenons compte que des messages qui interviennent lors d'un *handover*. Les procédures d'établissement des *bearers* ne sont effectuées qu'en amont de la simulation et n'induisent aucun échange de messages. La maintenance des *bearers* pendant le *handover* n'est pas considérée comme faisant partie de ces procédures puisqu'elle fait partie intégrante de la procédure de *handover* avec le *Path Switch*. Elle est donc implantée.

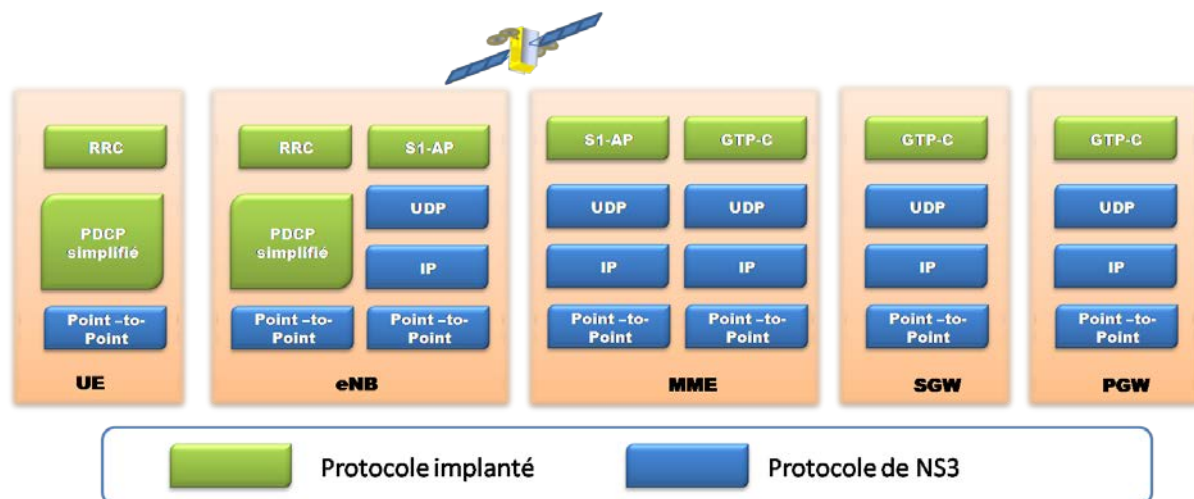


Figure 59 Implantation des couches protocolaires sur le plan utilisateur

Le *handover* n'est pas déclenché en fonction de la force ou de la qualité du signal sur l'interface radio. Il est déclenché par l'utilisateur de la simulation. Lorsque la phase de préparation est terminée, le changement physique d'eNB est simulé par la désactivation de la liaison entre l'UE et l'eNB source, puis l'activation de celle entre l'UE et l'eNB cible. L'interruption de service lors de ce *handover* est

simulée par une temporisation entre la désactivation et l'activation de 20ms. Nous choisissons une valeur constante, car cette valeur est dépendante de la couche physique que nous n'implantons que de manière très simplifiée. En outre, cette valeur est le plus souvent négligeable vis-à-vis du délai induit par le satellite.

5.2.2 Problème d'implantation du protocole TCP dans NS3

Les versions du protocole TCP disponibles dans NS3 sont seulement au nombre de 4 : RFC 793, Tahoe, Reno et New Reno. Sur la majorité des équipements, les versions de TCP utilisées sont TCP Cubic, qui est configuré par défaut dans Linux, et Compound pour Windows. Leur comportement est différent vis-à-vis de ceux proposées dans NS3. Le principal inconvénient des implantations TCP dans NS3 est l'absence des extensions qui sont essentielles pour les réseaux très haut-débit, ainsi que pour les réseaux avec des latences importantes. Elles sont décrites dans la RFC 1323 (Borman, Braden, & Jacobson, May 1982). Les trois principales extensions qui sont essentielles aux communications haut-débit par satellite sont le *Window Scaling*, le *Selective Acknowledgement* et les *Timestamps*.

5.2.2.1 Extensions TCP

WindowScaling

Le *WindowScaling* permet d'élargir la taille de la fenêtre pour atteindre des débits plus importants. En effet, le nombre de segments qui n'ont pas été acquittés est limité à la taille maximale de la fenêtre de congestion. La représentation de ce champ a une taille maximale de 16 bits. La limitation de ce champ est atteinte rapidement dans le cas de liaisons à très haut-débit et des réseaux à forte latence, comme les liaisons satellite. L'extension permet d'augmenter considérablement le débit maximal d'une connexion TCP.

Le Selective Acknowledgement

L'acquittement cumulatif utilisé dans NS3 ne permet que de réaliser une retransmission pendant une durée correspondante au RTT. Cette limitation est très pénalisante lors d'une transmission satellite puisque le RTT est très important, de l'ordre de 600ms. L'accusé de réception sélectif permet de réaliser plusieurs retransmissions pendant un RTT.

Les Timestamps

Les *Timestamps* horodatent les segments TCP pour permettre une mesure précise du RTT qui est primordiale pour le calcul du RTO. Ces extensions peuvent aussi être utilisées pour réordonner les segments TCP.

5.2.2.2 Implantation Linux

Gestion des fenêtres de congestion

L'implantation de TCP dans Linux est légèrement différente de celle définie dans les RFCs. Premièrement, la gestion de la fenêtre de congestion ne prend pas en compte les mêmes paramètres. La fenêtre de congestion TCP est comparée au nombre de segments qui sont en cours de transfert alors que, dans les RFC, elle est comparée au nombre de segments qui ne sont pas acquittés (Sarolahti & Kuznetsov, June 2002).

Lors de la détection d'une duplication ou d'une perte, TCP rentre dans le mode réordonnancement. Dans ce mode, la taille de la fenêtre de congestion est conservée constante, et donc, à chaque réception d'un nouvel accusé de réception un nouveau segment est envoyé. La durée de ce mode est calculée en fonction des retransmissions inutiles qui ont été observées et elle est définie par un seuil d'acquittements dupliqués. La valeur par défaut dans Linux est de trois acquittements dupliqués. Ce mode n'est donc pas visible sur la Figure 60. Ensuite, TCP linux rentre dans le mode Récupération. Dans ce mode, la fenêtre de congestion est réduite de 1 segment tous les deux acquittements. Lorsque l'on atteint le seuil, la fenêtre de congestion se stabilise, puis repart en *congestion avoidance*.

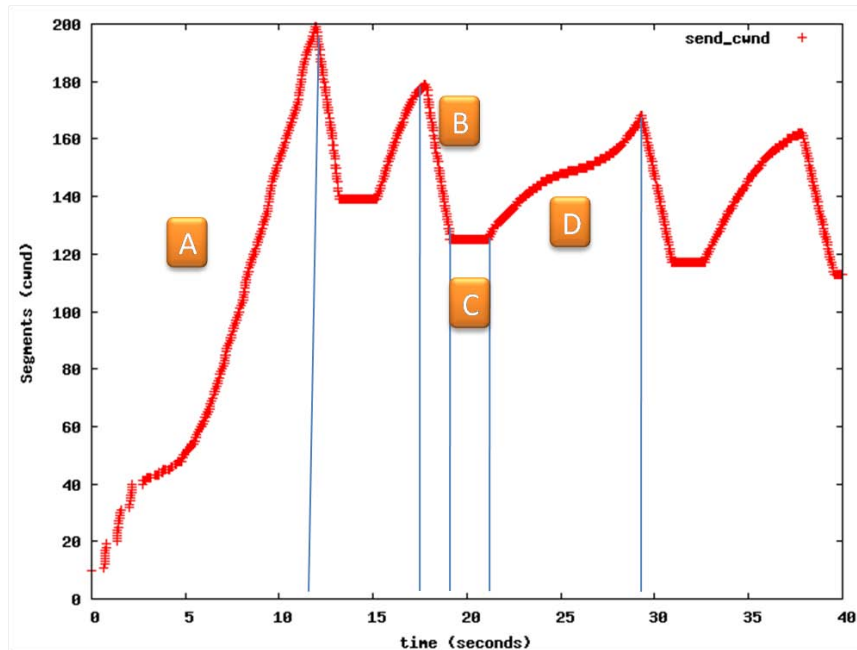


Figure 60 Fenêtre de congestion de TCP

Dans la partie A, la fenêtre de congestion est en *slow-start*. Les irrégularités dans la première partie sont dues à la gestion des *buffers* par linux. Elle se termine lorsqu'un accusé de réception dupliqué est reçu. TCP rentre en mode d'ordonnancement, cependant il ne dure que trois accusés de réception dupliqués et il est donc invisible sur ce graphique. Ensuite, TCP rentre dans le mode de récupération. Contrairement aux RFCs, l'implantation TCP n'envoie un nouveau segment que lorsqu'il reçoit deux nouveaux accusés. C'est pour cela que la taille de la fenêtre connaît une décroissance. Lorsque la fenêtre de congestion atteint un seuil défini grâce au maximum atteint dans la phase précédente, celle-ci se stabilise (phase C dans la Figure 60). Pendant cette phase, TCP envoie un paquet à chaque fois qu'un nouvel accusé de réception est reçu. Puis dans la phase D, TCP est dans la phase de *congestion avoidance* où la fenêtre de congestion évolue en suivant la fonction cubique.

Estimation du RTO

L'estimation du RTO est améliorée pour prendre en compte de manière plus efficace les brusques changements de RTT comme lors d'un *handover* inter-satellite-terrestre. Le RTO minimum est de 200ms alors que dans la RFC, il est de 1s.

Gestion des *burst*

Lorsqu'une grande partie de la fenêtre de congestion est acquittée, il peut survenir un envoi de nombreux paquets de données, comme par exemple, lors d'un *handover* inter-satellite-terrestre. Ce *burst* peut entraîner des congestions dans le réseau. Pour l'éviter, TCP linux limite le nombre de segments TCP envoyés à trois par acquittement.

5.2.3 L'émulation

Comme l'efficacité du mécanisme de *Forward* est intimement liée au comportement du protocole TCP, il est nécessaire de réaliser la simulation avec une version de TCP actuelle. C'est pour cette raison que l'architecture de la simulation a été modifiée.

La partie applicative est réalisée par la véritable implantation de la pile TCP/IP de linux. Pour permettre l'interconnexion de cette pile protocolaire avec l'architecture de simulation LTE déjà développée, nous utilisons des machines virtuelles Linux qui sont raccordées grâce à un élément de NS3 appelé *TAP Bridge Device* (Figure 61). L'application est toujours un transfert de fichier qui est réalisé par le protocole FTP.

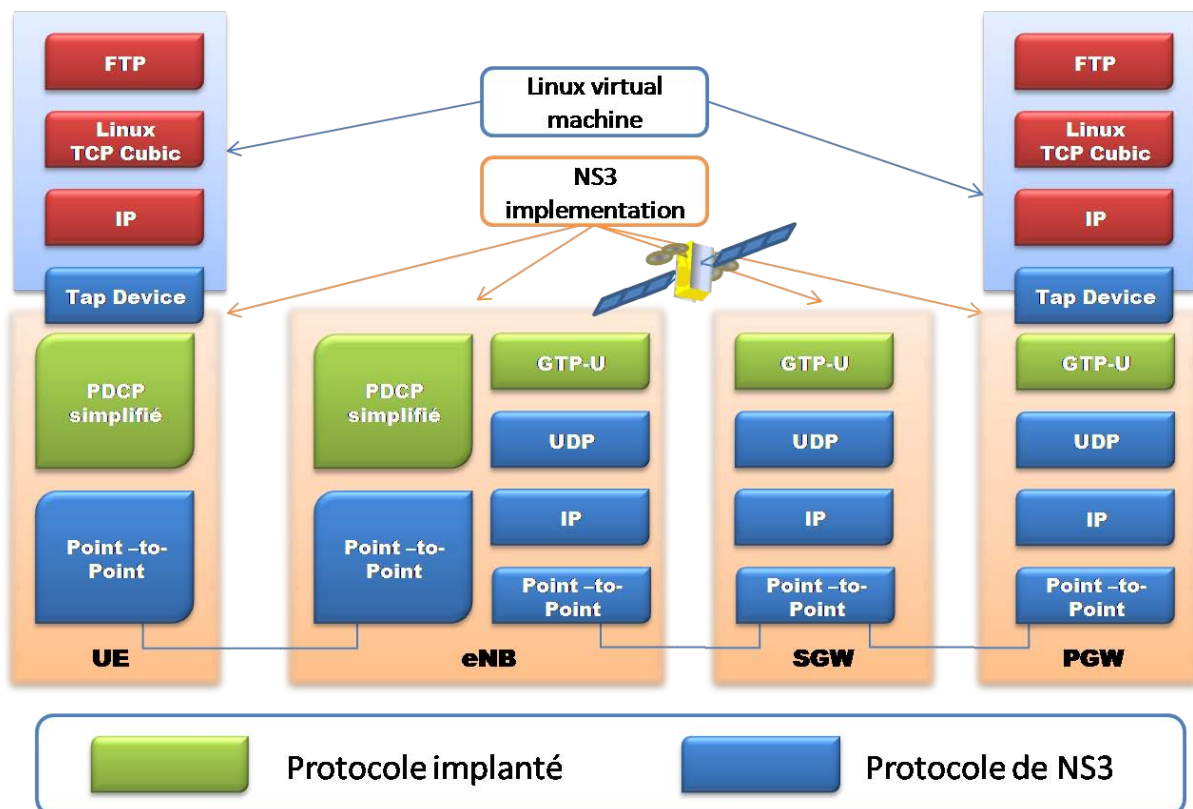


Figure 61 Architecture de simulation avec la pile protocolaire TCP/IP linux

Le tableau ci-après contient les différentes options principales de TCP qui sont utilisées par l'implantation de TCP Cubic dans linux.

Options TCP Cubic	États ou valeurs
Duplicate SACK	✓
Explicit Congestion Notification	✗
Forward Acknowledgement	✓
Forward RTO-Recovery	✗
Nombre d'accusés dupliqués pour terminer la phase de réordonnement	3
Selective Acknowledgement	✓
Slow Start after Idle	✓
Timestamps	✓
Windows Scaling	✓
Mémorisation	✓ ⁽¹⁾

(1) Dans les simulations de la section 5.4, la mémorisation de certains paramètres est désactivée.

Tableau 20 Paramètres de TCP cubic sur Linux

5.3 Émulation du mécanisme de *Forward* pour le *handover* inter-satellite-terrestre

Pour cette émulation, nous avons utilisé TCP Cubic de linux et FTP en tant que protocole applicatif. L'UE va télécharger un fichier de 50 Mo. Nous allons comparer les différentes fenêtres de congestion pour observer le comportement de TCP lors du *handover*. Pour observer la fenêtre de congestion de TCP, nous utilisons la sonde linux *tcpprobe*. Le mécanisme mis en œuvre est le mécanisme de *Forward* décrit dans le paragraphe 4.4.1.

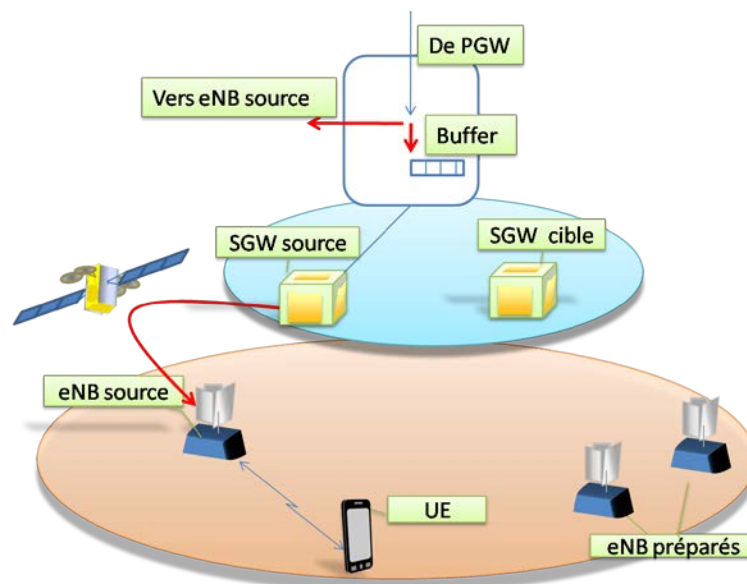


Figure 62 Principe du mécanisme de *Forward* par buffer

La Figure 62 permet de rappeler le principe du mécanisme mis en place pour un *handover* inter-satellite-terrestre. Pendant la préparation du *handover*, la SGW source va créer un *buffer*, qui va stocker des copies des paquets à destination de l'UE. Lorsque le *handover* est réalisé, la SGW source va transmettre à l'eNB cible les paquets non-reçus qui sont stockés dans son *buffer*.

Sans mécanisme de *Forward*

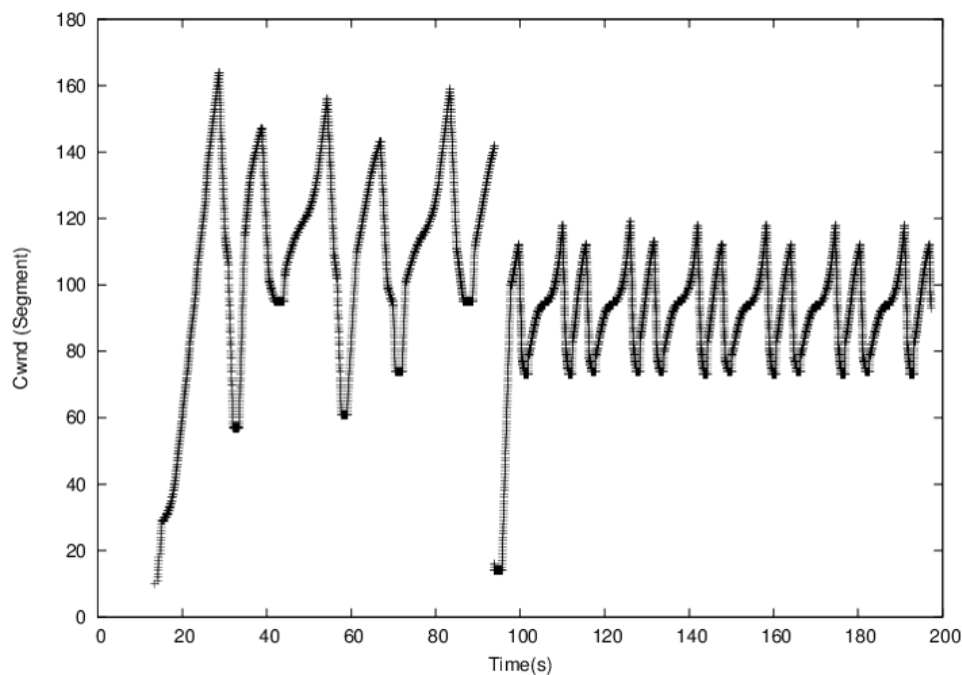


Figure 63 Fenêtre de congestion lors d'un *handover* sans *Forward* déclenché lors de la phase normale

La Figure 63 représente l'évolution de la fenêtre de congestion lors d'un *handover* inter-satellite-terrestre. On remarque très nettement l'instant où le *handover* est déclenché, 80 secondes après le début du flux TCP. La coupure survient lors de la phase normale de TCP avec une évolution suivant la fonction cubique. Nous remarquons qu'elle survient sur la fin de cette phase. Les pertes de paquets provoquent une coupure, cependant, TCP Cubic récupère très rapidement de celle-ci lorsque l'UE est connecté à l'eNB avec un backhaul terrestre.

Avec mécanisme de *Forward*

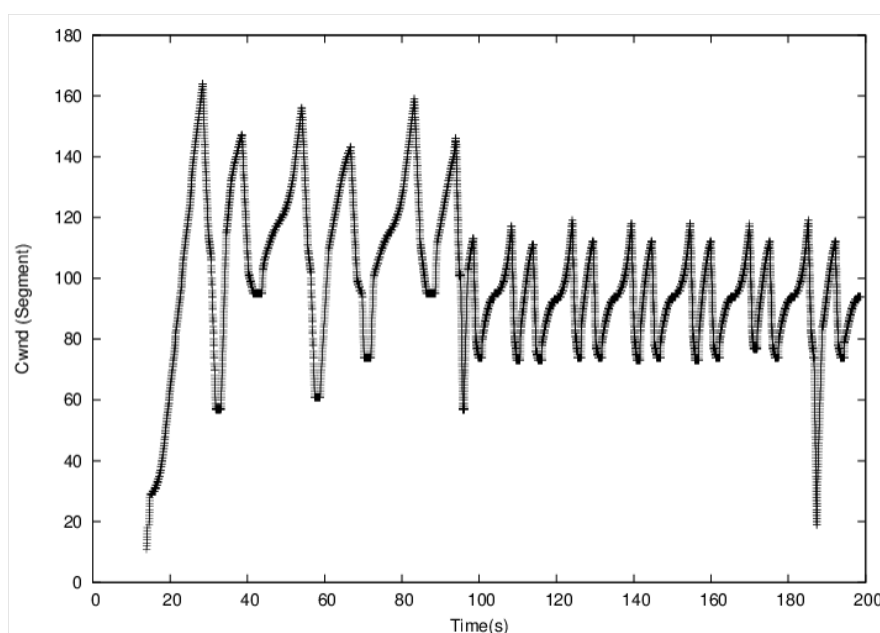


Figure 64 Fenêtre de congestion lors d'un *handover* avec *Forward* déclenché lors de la phase normale

Dans la Figure 64, la fenêtre de congestion n'est pas affectée par le *handover*. En effet, le mécanisme de *Forward* permet d'assurer qu'aucune perte ni duplication ne surviennent pendant le *handover*. Cependant, on peut remarquer que le gain reste globalement faible entre les deux courbes.

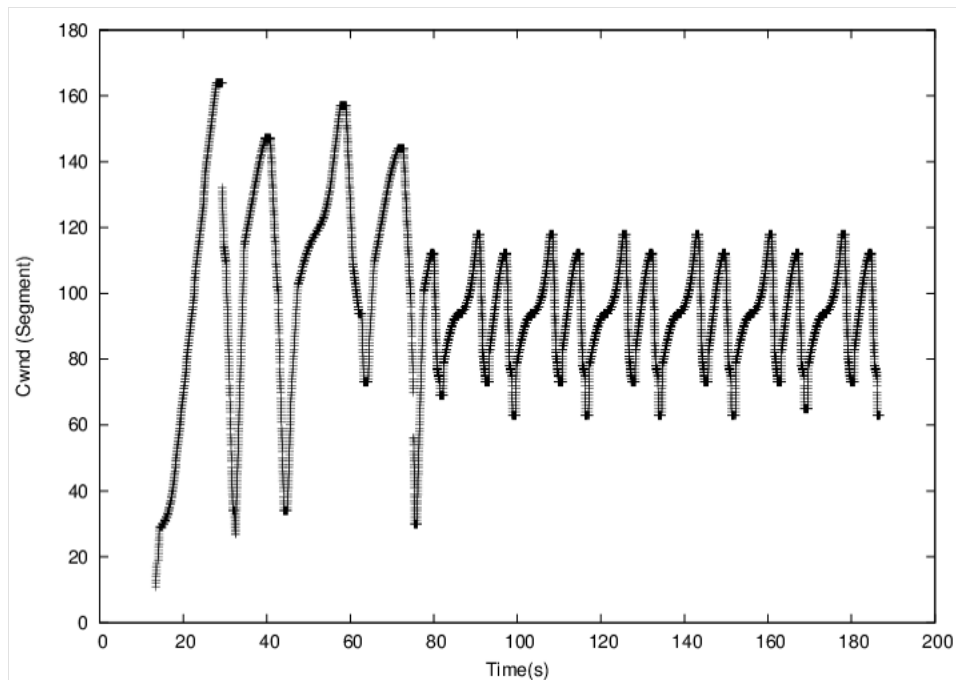


Figure 65 Fenêtre de congestion lors d'un *handover* avec *Forward* déclenché lors de la phase de récupération

La Figure 65 représente l'évolution de la fenêtre de congestion lors d'un *handover*, qui est déclenché pendant la phase de récupération, 62 secondes après le début de la connexion. Nous remarquons premièrement que les variations de la fenêtre de congestion sont plus importantes que dans les deux premiers cas. En effet, avant le déclenchement du *handover*, la fenêtre de congestion ne descend pas en dessous de 50 segments dans la Figure 63 et la Figure 64 alors que dans la Figure 65, elle atteint 30 segments. Ceci est la conséquence des effets de mémorisation de TCP. En effet, ce dernier conserve certains paramètres de la précédente connexion pour améliorer ses performances. Ces effets sont donc indépendants du mécanisme que nous proposons. Notre mécanisme permet encore une fois d'obtenir une fenêtre de congestion sans interruption due à des pertes ou des duplications de paquets. L'interruption que nous observons est simplement la résultante de l'acquittement d'une large partie de la fenêtre de congestion.

En effet, après l'exécution du *handover*, le premier accusé de réception, qui est transmis par l'UE, sera reçu avant les précédents, qui sont en cours de transfert sur le lien satellite. Il va donc accuser réception d'une grande partie de la fenêtre de congestion. Si celui-ci acquitte un nouveau segment, il pourra survenir un *burst* d'envoi. Cependant, l'implantation de TCP Cubic linux limite le nombre de transmissions à trois pour un même acquittement. Contrairement à la Figure 64, le *handover* est déclenché lors de la phase de récupération dans la Figure 65. Lors du *handover*, cet accusé de réception est un accusé de réception dupliqué. Cependant, grâce à l'acquittement sélectif, il valide tout de même une large partie de la fenêtre de congestion. Cela cause l'interruption dans la fenêtre de congestion.

L'intérêt de ce mécanisme est réel pour les flux TCP. Cependant, le nouveau mécanisme de *Forward* implique des modifications dans le segment terrestre, pour transmettre les numéros de séquences GTP, dans le but d'éviter les pertes et les duplications. En outre, la gestion du *buffer* est complexe, en particulier la gestion de l'envoi de paquets contenus dans le *buffer* après l'exécution du *handover*. Les bénéfices que procure le mécanisme de *Forward* nous semblent pour l'instant insuffisants pour compenser la complexité du mécanisme de *Forward*. L'intérêt pour celui-ci pourrait être exacerbé en cas de fréquences élevées des *handovers* entre les systèmes satellite et terrestres. Cependant, nous avons privilégié une minimisation des occurrences des *handovers* du système terrestre vers le satellite par l'intermédiaire des optimisations de la phase de préparation. Par conséquent, le changement fréquent d'un UE entre le système satellite et terrestre est limité.

5.4 Émulation du mécanisme de *Forward* pour le *handover* intra-satellite

Dans cette partie, nous nous intéressons au *handover* intra-satellite. Ce type de *handover* survient de manière plus fréquente que celui précédemment exposé, puisqu'il fait intervenir des eNB du même segment. En effet, les optimisations de la phase de préparation n'ont pas pour objectif de privilégier certains eNB à l'intérieur d'un même segment. Nous avons utilisé la même méthodologie que pour l'émulation des *handovers* inter-satellite-terrestre à la seule exception près que nous focalisons notre émulation sur une période plus courte en émulant le téléchargement d'un fichier plus petit qui fait 5Mo, pour nous concentrer sur l'instant du *handover*.

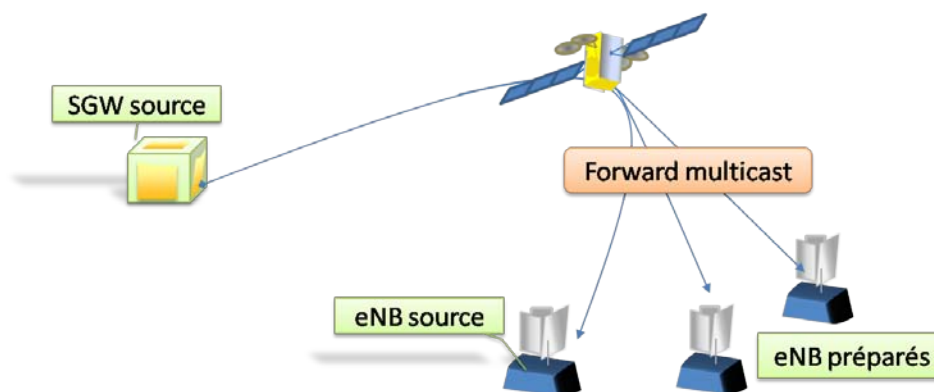


Figure 66 Principe du mécanisme de *Forward* du *handover* intra-satellite

La Figure 66 remémore le principe du mécanisme de *Forward* du *handover* intra-satellite alors qu'aucune interface X2 n'est disponible. L'objectif de celui-ci est de transmettre en *multicast* les données à destination de l'UE, vers les différents eNB, qui ont subi la phase de préparation. Deux variantes vont être émulées. La première ne se préoccupe pas de la gestion des pertes. Ainsi, la durée d'interruption du service pendant la phase d'exécution va engendrer quelques pertes. La deuxième solution assure l'absence de pertes et de paquets dupliqués en mémorisant les paquets qui sont susceptibles d'être perdus lors du *handover*. Nous considérons dans notre simulation que les terminaux satellite sont co-localisés avec les eNB.

Sans mécanisme de *Forward*

Les courbes suivantes représentent les évolutions de la fenêtre de congestion lorsqu'aucun mécanisme de *Forward* n'a été mis en œuvre.

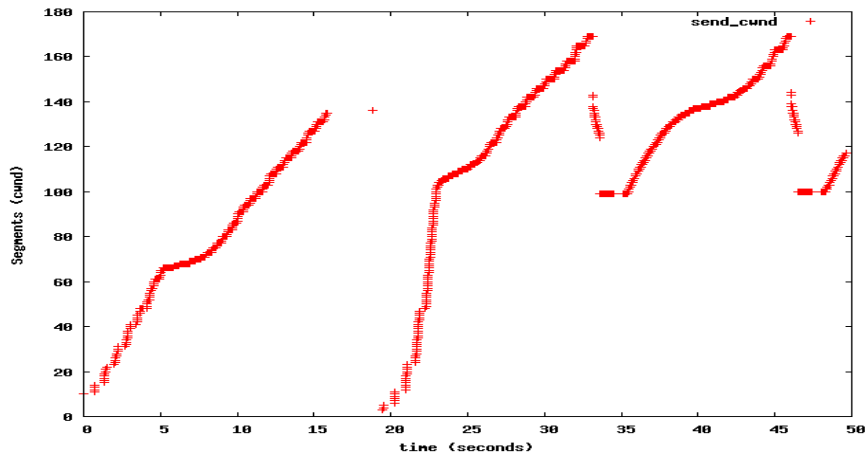


Figure 67 Fenêtre de congestion lors d'un *handover* intra-satellite sans mécanisme de *Forward* et déclenché après 15s

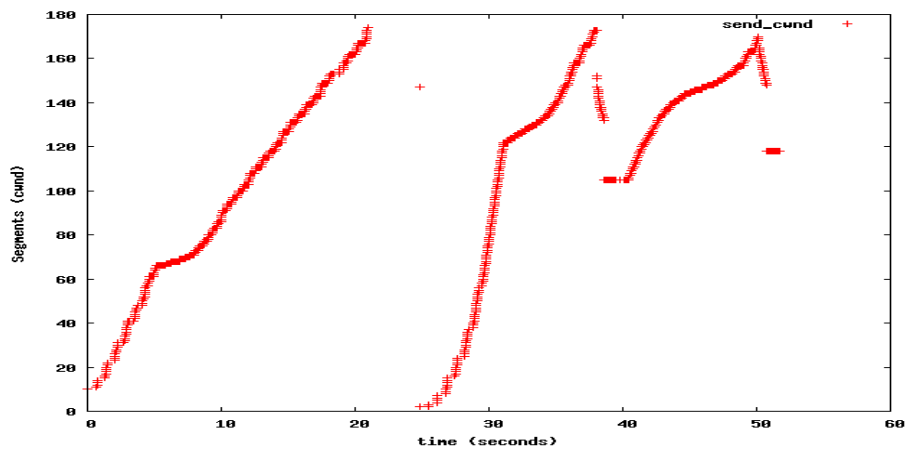


Figure 68 Fenêtre de congestion lors d'un *handover* intra-satellite sans mécanisme de *Forward* et déclenché après 20s

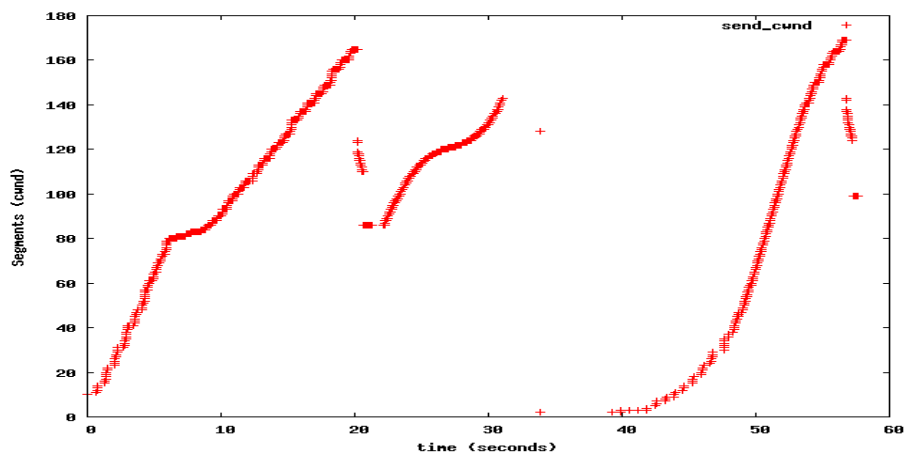


Figure 69 Fenêtre de congestion lors d'un *handover* intra-satellite sans mécanisme de *Forward* et déclenché après 30s

La Figure 67 et la Figure 68 représentent l'évolution de la fenêtre de congestion lorsque le *handover* survient pendant la phase de *slow-start*. Cependant, dans la Figure 67, il advient au milieu de cette phase alors que dans la Figure 68, il se produit à la fin de cette phase. Dans la Figure 69, le *handover* est déclenché lors de la phase de *congestion avoidance*. Nous remarquons que, lorsque le *handover* se produit lors de la phase de *slow-start*, TCP recommence la phase de *slow-start* plus rapidement que lorsqu'il survient pendant la phase de *congestion avoidance*. La récupération est beaucoup plus longue et entraîne un retard important. En effet, si nous comparons les durées des téléchargements entre la Figure 69 et les Figure 67 et Figure 68, nous remarquons que la différence atteint presque une dizaine de secondes.

Avec mécanisme de *Forward* et pertes

Les figures ci-après représentent l'évolution de la fenêtre de congestion, lors d'un *handover* avec un mécanisme de *Forward* qui ne gère pas les pertes dues à l'interruption de connexion pendant la phase d'exécution. Ainsi une dizaine de paquets sont perdus durant le *handover*.

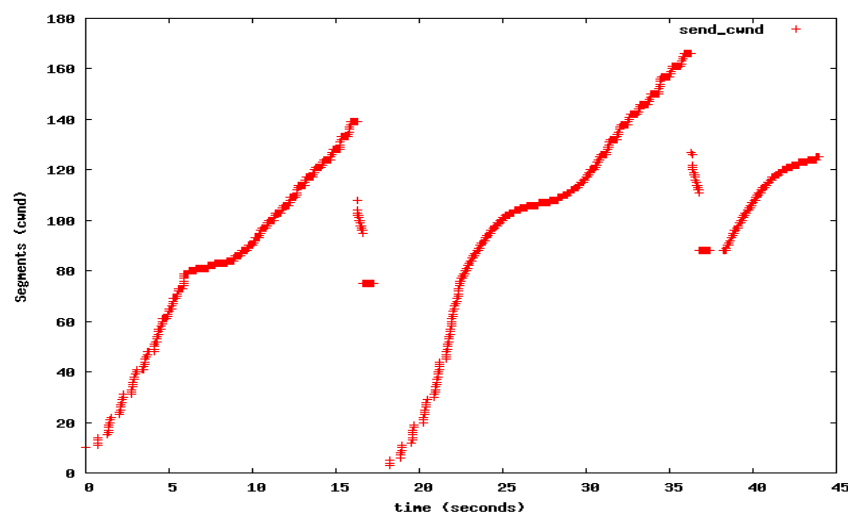


Figure 70 Fenêtre de congestion lors d'un *handover* intra-satellite avec mécanisme de *Forward* et des pertes et déclenché après 15s

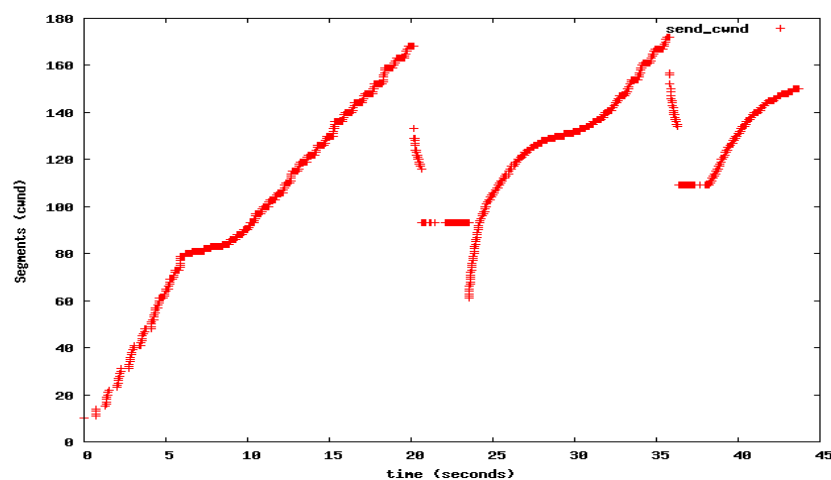


Figure 71 Fenêtre de congestion lors d'un *handover* intra-satellite avec mécanisme de *Forward*, sans gestion des pertes et déclenché après 20s

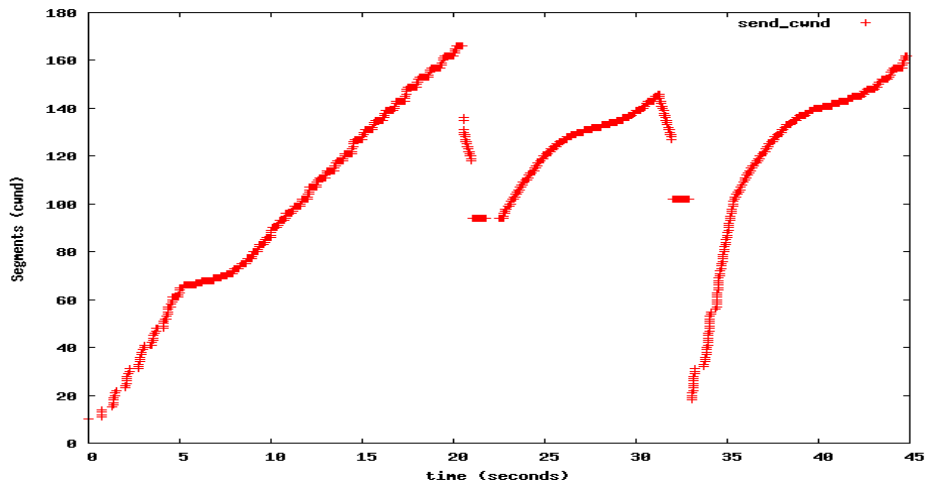


Figure 72 Fenêtre de congestion lors d'un *handover* intra-satellite avec mécanisme de *Forward*, sans gestion des pertes et déclenché après 30s

Les performances du mécanisme de *Forward* sont honorables. Il permet de réduire la durée du transfert de 5 secondes en comparaison du *handover* sans mécanisme de *Forward*. Nous observons, cependant, nettement l'influence de quelques pertes lors du *handover* et son impact sur la fenêtre de congestion.

Avec mécanisme de *Forward* sans pertes

La Figure 71 représente l'évolution de la fenêtre de congestion lors d'un *handover* avec un mécanisme de *Forward* qui assure un *handover* sans perte ni duplication de paquets.

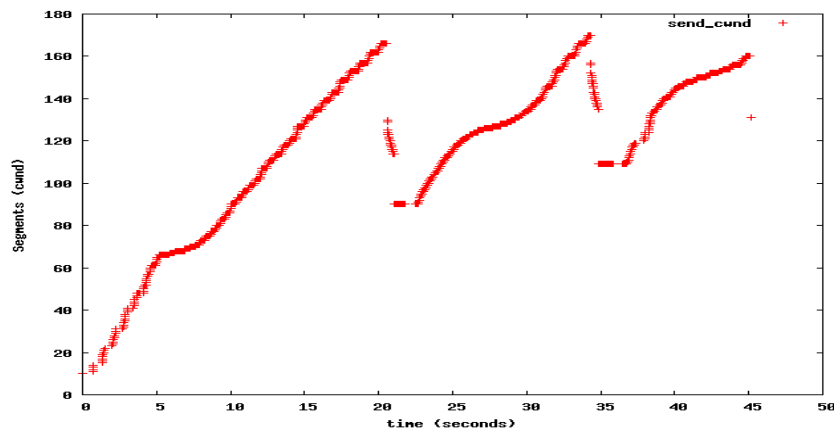


Figure 73 Fenêtre de congestion lors d'un *handover* intra-satellite avec mécanisme de *Forward* et sans perte

La fenêtre de congestion ne diffère pas en fonction de l'instant de déclenchement du *handover*. C'est pour cette raison qu'une seule figure est proposée. Nous remarquons que le mécanisme de *Forward* réalise parfaitement son objectif de rendre le *handover* invisible pour le protocole TCP. L'absence de perte et de duplication nous assure ainsi les meilleures performances pour les flux TCP pendant le *handover*. Le bénéfice en comparaison du mécanisme de *Forward* avec perte est non négligeable mais pas décisif. En effet, nous gagnons entre 1 et 2 secondes sur la durée de transfert. Cependant, une fréquence élevée des *handovers* ferait plus largement perdre en efficacité le mécanisme de *Forward* avec perte, alors que la deuxième solution ne serait pas atteinte par ce paramètre. La

fréquence des *handovers* est un paramètre lié à la phase de préparation qui n'est pas simulée dans ce chapitre.

5.5 Conclusion

Nous avons proposé des mécanismes de *Forward* qui permettent de remplacer celui défini dans le standard LTE. La modification était nécessaire, puisque celui-ci imposait une surconsommation des ressources satellite et n'était pas efficace, en raison de l'important délai induit par le satellite.

Dans un premier temps, nous avons optimisé le mécanisme de *Forward* pour un *handover* inter-satellite-terrestre. Pour vérifier l'intérêt de ce mécanisme, nous avons réalisé ces émulations, en comparant le comportement de TCP avec notre mécanisme de *Forward* et sans aucun mécanisme de *Forward*. Elles ont révélé qu'au regard de la complexité du mécanisme et des bénéfices qu'il apporte, la balance ne lui est pas favorable. En effet, il implique des modifications trop importantes, en particulier dans le réseau de cœur du segment terrestre pour être réellement intéressant en l'état.

Dans un second temps, nous avons proposé une optimisation du mécanisme de *Forward* dans le cadre d'un *handover* intra-satellite. Deux variantes ont été proposées. Elles permettent toutes deux de transférer les paquets en *multicast*, entre toutes les eNB qui sont potentiellement cibles. La différence repose sur la gestion des pertes de paquets lors de la phase d'exécution. La première ne va pas se préoccuper de ces paquets et va alors occasionner un certain nombre de pertes. Celles-ci affectent les paquets qui arrivent aux eNB, tandis que l'UE vient de se détacher de l'eNB source et n'est pas encore connecté à l'eNB cible. La deuxième solution, quant à elle, assure un mécanisme de *Forward* sans perte. Les émulations ont mis à jour un réel intérêt de ces deux mécanismes de *Forward*. Ils sont à l'origine de bénéfices indéniables, en comparaison des performances obtenues en l'absence de mécanisme de *Forward*. Cependant, le choix entre les deux mécanismes reste difficile. De nombreux paramètres peuvent être exploités comme la fréquence des *handovers* qui est dépendante des paramètres de la couche physique et de la phase de préparation, qui ne sont pas abordées dans cette thèse.

Les deux mécanismes de *Forward* pour le *handover* intra-satellite sont une réelle réussite. En effet, ils augmentent significativement les performances des applications qui se fondent sur TCP tout en assurant une gestion optimisée des ressources sur le système satellite. La cohérence entre ces mécanismes de *Forward* et la phase de préparation adaptée au lien satellite nous permet d'envisager une augmentation supérieure des performances pour un système complet.

Nous allons maintenant nous concentrer sur une autre catégorie de réseau hybride, les réseaux instantanés MANET-satellite.

Chapitre 6 : MONET : Présentation de MONET et du mécanisme de sélection de passerelle.

Le développement de réseaux MANET avec un backhaul satellite devient très intéressant et primordial dans le cadre de déploiement rapide et temporaire. Dans le chapitre 1, nous avons fait remarquer qu'aucun projet, à notre connaissance, n'a étudié de réseau MANET possédant plusieurs passerelles satellite. Pour combler cette absence, nous avons examiné ce scénario, puis nous avons proposé une optimisation d'un des principaux problèmes soulevés par l'hybridation satellite terrestre qu'est la sélection de passerelle. Cette optimisation a été réalisée dans le cadre du projet Européen MONET (MONET Project).

6.1 Le contexte et le projet MONET

Le projet MONET repose sur la constatation que l'hybridation d'un réseau MANET avec un système satellite est grandement bénéfique grâce à la complémentarité des deux technologies. En effet, la couverture très étendue du satellite alliée aux mécanismes d'auto-configuration des réseaux MANET permet de réaliser un déploiement rapide et temporaire sur une très large zone.

6.1.1 Les scénarios d'utilisation

Les scénarios d'utilisation ciblés par le projet MONET ont été définis par l'ensemble du consortium et sont séparés en trois catégories : les communications d'urgence pour la sécurité civile, les communications de contrôle d'un aéroport ainsi que des communications d'entreprises (MONET, 2010). Cependant, le scénario le plus probable et le plus prometteur est l'utilisation de ce réseau hybride pour la sécurité civile. En effet, l'utilisation d'un réseau hybride devient essentielle, lorsque le théâtre d'intervention se situe dans une zone où les infrastructures terrestres sont inexistantes, ainsi que lorsqu'elles sont détruites, par exemple, en conséquence d'un raz-de-marée ou d'un séisme. Le projet a extrait en particulier deux principaux cas d'utilisation d'urgence : les feux de forêts et le secours en zone montagneuse. Le projet MONET a pour objectif de définir une architecture de réseau hybride satellite/terrestre qui réponde aux attentes des communications d'urgence. Il doit proposer et réaliser des optimisations pour résoudre les problèmes posés par l'hybridation. Les intervenants sur le théâtre de l'opération sont séparés en trois catégories : le quartier général, les postes de commandement et les premiers intervenants (Figure 74). Le quartier général est une entité fixe et éloignée du théâtre des opérations. Il est raccordé aux différents acteurs par le segment satellite. Les postes de commandement dirigent chacun leur unité sur le terrain et les premiers intervenants sont composés de tous les autres acteurs ; policiers, pompiers, service médical et équipe de secours.

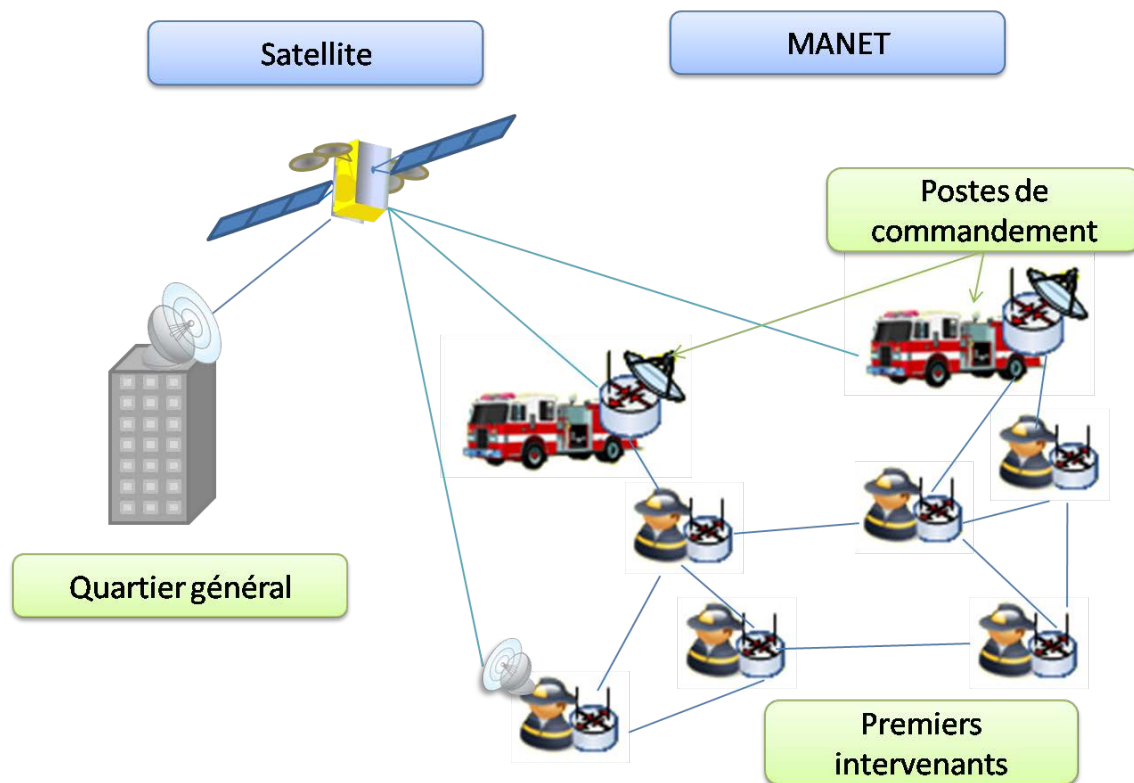


Figure 74 Scénario d'utilisation

6.1.2 Architecture

Le réseau hybride MONET est composé de deux segments : le segment satellite et le segment terrestre (MONET, 2010). De la même manière que pour la définition des scénarios, l'architecture a été construite par l'ensemble des partenaires du projet. Le lien satellite possède deux fonctions. Dans un premier temps, il permet de réaliser le raccordement d'un réseau MANET avec les réseaux extérieurs tels qu'Internet. Il permet aussi d'interconnecter les réseaux MANET disjoints, lorsque la distance les séparant devient trop grande. Il englobe tout le segment contenu entre le quartier général et les points d'accès satellite réalisant l'interface avec le MANET (Figure 75). Les postes de commandement disposent obligatoirement d'un terminal satellite et d'une interface MANET. Ils sont la plupart du temps implantés dans un véhicule. Ainsi, le terminal satellite qui doit être transportable peut être un dérivé des systèmes satellite fixes qui fournissent des débits bien plus importants que ceux provenant des systèmes mobiles. *A contrario*, les premiers intervenants peuvent potentiellement posséder un accès satellite, tels qu'un terminal BGAN d'Inmarsat (BGAN-France) avec de plus faibles performances.

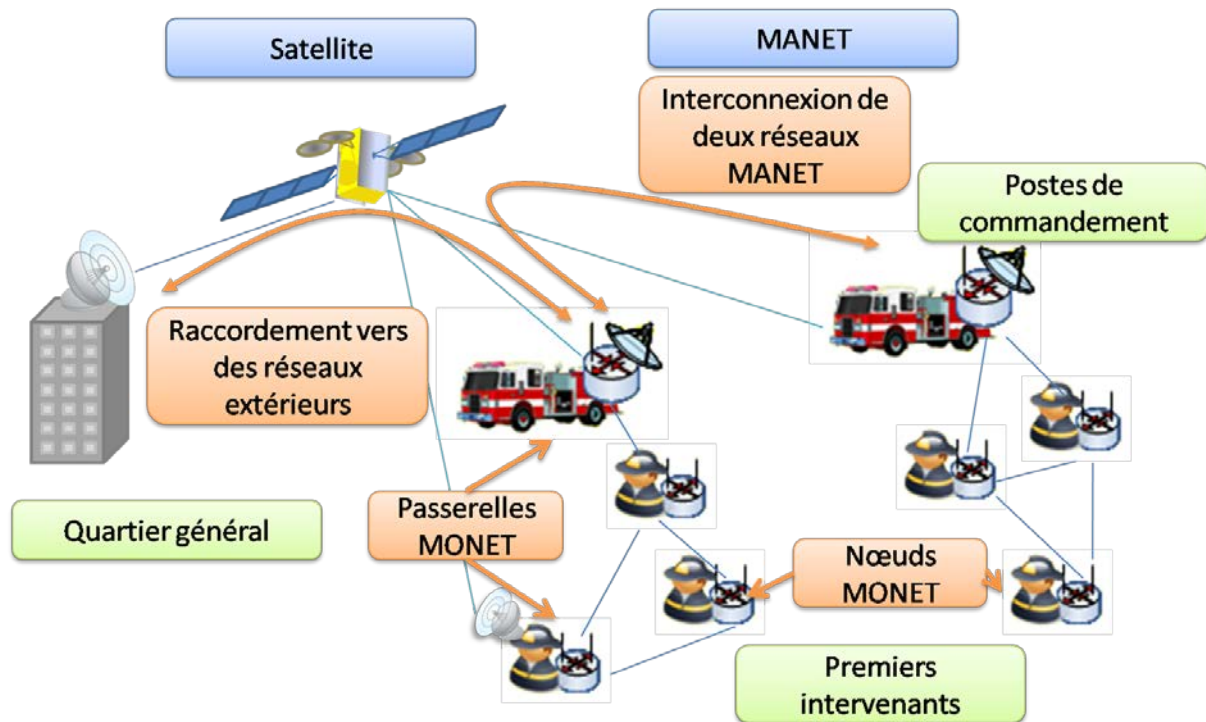


Figure 75 Architecture du réseau MONET

Le segment terrestre, quant à lui, est composé de nœuds MANET qui correspondent à la fois aux postes de commandement et aux premiers intervenants. Les différents partenaires du projet ont opté pour le protocole de routage OLSR (*Optimized Link State Routing*) (Clausen & Jacquet, October 2003). Il fait partie de la famille des protocoles de routage proactifs. Ce choix a été effectué en raison des besoins applicatifs de la sécurité publique. En effet, les communications d'urgence possèdent d'importantes contraintes en termes de délai. Les applications les plus importantes et les plus critiques sont des services de voix. Il est donc primordial de minimiser le délai des données liées à ces applications. Les protocoles proactifs tels qu'OLSR permettent de réduire le délai de bout-en-bout en comparaison aux protocoles de routage réactifs. En outre, OLSR utilise un mécanisme qui réduit la charge provoquée par la diffusion des messages de contrôle, qui se trouve être l'un des points faibles des protocoles proactifs. Ce mécanisme est appelé MPR (*Multi-Point Relay*).

Ma contribution a consisté à étudier un problème d'interconnexion particulier, la sélection de passerelle. Le réseau MONET possède plusieurs passerelles dotées d'un terminal satellite. Ces terminaux peuvent avoir différentes caractéristiques en fonction du système satellite auquel ils appartiennent. En effet, les terminaux satellite installés dans le poste de commandement fournissent des débits importants alors qu'un terminal BGAN offre une capacité limitée. Il s'agit alors de développer un algorithme de sélection de passerelle en fonction des caractéristiques du réseau MANET ainsi que des caractéristiques des terminaux satellite.

Bien que l'étude du placement des passerelles dans le réseau MANET ait de nombreuses conséquences sur les performances applicatives, nous n'étudions pas de mécanismes de placement. En effet, les terminaux satellite sont onéreux et sont donc présents en faible nombre sur le théâtre des opérations. Ainsi, un mécanisme de placement des passerelles n'apporterait qu'un avantage qui nous semble infime. Le positionnement de ces passerelles est, de toute façon, essentiellement lié à des contraintes opérationnelles.

6.2 La sélection de passerelle

Les études sur la sélection de passerelle dans un réseau MANET ont été nombreuses ces dix dernières années. Trois points se sont révélés primordiaux pour assurer leur bon fonctionnement et leur interconnexion avec les réseaux extérieurs : la gestion de la mobilité du réseau par rapport aux réseaux extérieurs, la méthode utilisée pour la découverte de la passerelle et les métriques permettant le choix de la passerelle (Kumar, Sarje, & Misra, 2010) (Nordström, Gunningberg, & Tschudin, 2011) (Le-Trung, Engelstad, Skeie, & Taherkordi, November 2008).

6.2.1 L'adressage IP et la mobilité du réseau

Un réseau MANET qui utilise plusieurs passerelles pose le problème de la cohérence entre le routage hiérarchique IP et la mobilité des nœuds du MANET. En effet, pour assurer une continuité du routage, lorsqu'un nœud change de passerelle, celui-ci devrait changer d'adresse IP. Dans le même temps, il a besoin de conserver son ancienne adresse IP si la continuité des applications doit être assurée. Par exemple, dans la Figure 76, le nœud A se déplace de la passerelle 1 vers la passerelle 2. Il doit donc changer de sous-réseau IP pour pouvoir choisir la passerelle 2 comme passerelle pour son trafic. Cependant, ce changement d'adresse entraîne la rupture des applications.

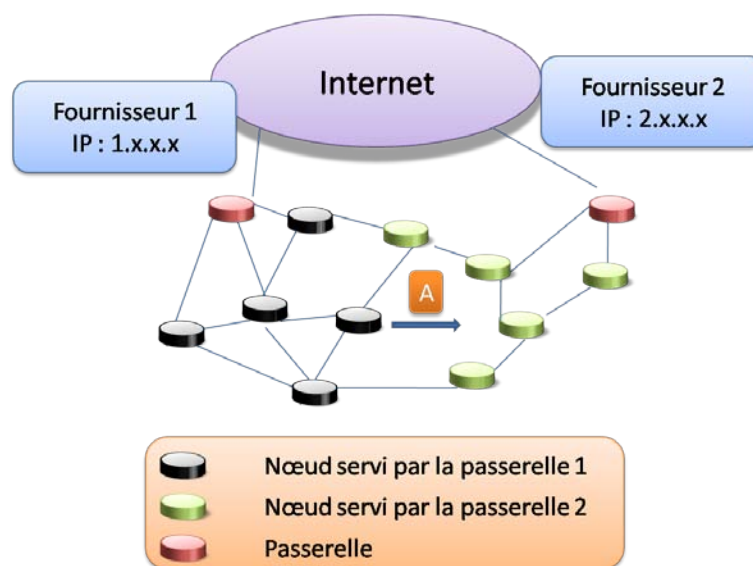


Figure 76 Problème de gestion de la mobilité dans les MANET à plusieurs passerelles

Pour résoudre ce problème plusieurs solutions ont été définies. La majorité de celles-ci reposent sur les protocoles de la famille Mobile IP. Dans ces solutions, les passerelles font office de *Foreign Agent* (FA). Ainsi, (Jönsson, Alriksson, Larsson, Johansson, & Maguire, August 2000) propose un mécanisme reposant sur le protocole de routage réactif AODV et MIPv4 appelé MIPMANET. Il permet d'adapter MIP à un protocole de routage à la demande. C'est l'une des solutions les plus répandues dans la littérature. Se fondant sur la même association, des propositions plus simples sont proposées dans (Broch, Maltz, & Johnson, June 1999) (Sun, Belding-Royer, & Perkins, 2002). Différentes études ont permis d'adapter MIPv4 à des protocoles de routage proactifs telles que (Engelstad, Tønnesen, Hafslund, & Egeland, 2004) (Benzaid, Minet, Al Agha, Adjih, & Allard, 2004).

D'autres articles traitent de ce problème sans MIP avec en particulier l'utilisation de tunnels et de traduction d'adresse (Engelstad, Tønnesen, Hafslund, & Egeland, 2004) (Engelstad & Egeland, NAT-based Internet Connectivity for On Demand MANETs, January 2004). Cependant, ces techniques s'avèrent peu efficaces à cause de la lourdeur des tunnels.

Dans notre scénario, la solution est beaucoup plus simple grâce à la place centrale qu'occupe le quartier général dans l'architecture MONET. Le réseau MANET peut ainsi être considéré comme un seul sous-réseau sans problème de cohérence avec le routage hiérarchique IP. Ainsi le nœud A dans la Figure 77 peut conserver son adresse IP s'il change de passerelle de raccordement.

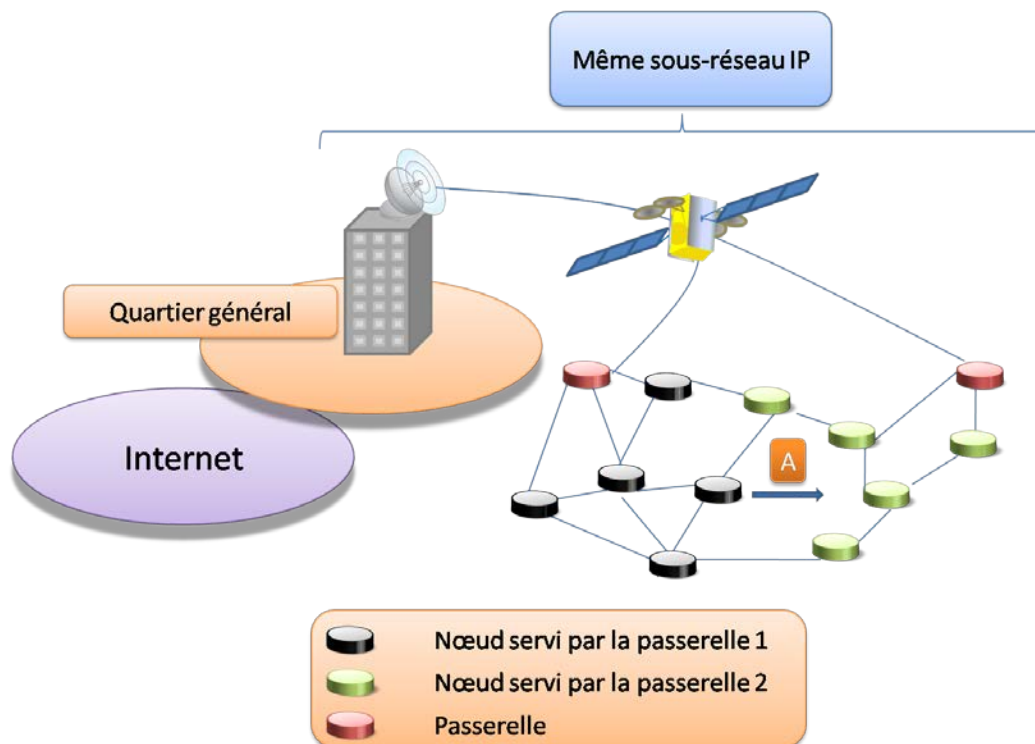


Figure 77 La gestion de la mobilité dans MONET

6.2.2 Les méthodes pour découvrir les passerelles.

La découverte de passerelle permet aux nœuds de déterminer les passerelles disponibles dans le réseau MANET. Dans le même esprit que pour les protocoles de routages MANET, les méthodes pour découvrir les passerelles sont séparées en trois catégories : réactive, proactive et hybride.

La méthode réactive

La méthode réactive est plus complexe que sa pendante proactive. Lorsque le nœud destination se situe à l'extérieur du MANET, le nœud source diffuse une requête vers une passerelle, génériquement appelée GW-SOL (*Gateway Solicitation*). Lorsqu'une passerelle reçoit ce message, elle répond au nœud source. Cette méthode est, bien sûr, le plus souvent utilisée conjointement avec un protocole de routage réactif. Dans ce cas, le nœud source doit tout d'abord connaître l'emplacement du nœud destination (Le-Trung, Engelstad, Skeie, & Taherkordi, November 2008). En effet, si celui-ci se situe à l'intérieur du MANET, la procédure est réalisée par le protocole de routage

sinon elle est effectuée grâce au mécanisme de découverte de passerelle. La caractérisation de la localisation du nœud destination peut être source d'un important retard. En effet, si le MANET est géré par un protocole de routage réactif, le nœud source envoie une requête de la même manière que si celui-ci se trouvait dans le MANET. C'est seulement si cette étape échoue que l'on considère que le nœud destination se trouve à l'extérieur de ce réseau et que l'on enclenche la découverte de la passerelle.

La méthode proactive

La méthode proactive est plus simple, la passerelle diffuse un message périodique, appelé génériquement *Gateway Advertisement* (GW-ADV) et plus spécifiquement Host Network Association (HNA) dans les spécifications du protocole OLSR. Ainsi, tous les nœuds du MANET ont connaissance du nœud qui fait office de passerelle. Cette méthode présente le même inconvénient que les protocoles de routage proactifs puisqu'elle implique un *overhead* important à cause de la diffusion de messages de contrôle.

La méthode hybride

La méthode hybride permet de diminuer l'*overhead* induit par la diffusion de messages, en réduisant la portée des GW-ADV à un certain nombre de sauts. La technique la plus répandue est la limitation du TTL (Time-To-Live) des GW-ADV au rayon de la zone proactive voulue. Les nœuds qui ne se situent pas dans cette zone doivent obtenir le chemin vers la passerelle sur demande. Différentes études (Zhuang, Liu, & Liu, June 2009), optimisent ainsi la taille de cette zone en fonction de différents paramètres tels que la vitesse de déplacement de la passerelle.

Dans notre scénario, l'utilisation d'une méthode proactive semble une évidence du fait de l'utilisation du protocole de routage proactif OLSR. De plus l'étude (Ghassemian, Hofmann, Prehofer, Friderikos, & Aghvami, March 2004) montre que la méthode proactive semble offrir de meilleures performances, en particulier en termes de délai. En outre, l'inconvénient de l'*overhead* est réduit grâce au mécanisme des MPR d'OLSR.

6.2.3 Les algorithmes de sélection de passerelle

Le choix de la passerelle est réalisé grâce à un algorithme qui nous permet de trouver le chemin le plus court. La sélection se fonde sur la même technique que le choix du chemin, à l'intérieur du réseau MANET par les protocoles de routage. Dans la majorité des cas, la sélection d'une passerelle repose entièrement, ou en grande partie, sur la comparaison des coûts entre les différents chemins menant aux différentes passerelles. En particulier, les méthodes proactives utilisent les algorithmes de plus court chemin tels que Dijkstra et Bellman-Ford. Le monde de la recherche MANET s'est donc attelé à la définition de la métrique la plus performante pour le routage. Les métriques utilisées pour la sélection de la passerelle sont généralement les mêmes que celles utilisées par les protocoles de routage.

Premièrement, une métrique, m , doit être isotonique. Cette propriété est nécessaire pour assurer l'absence de boucle infinie dans le routage lorsque la méthode proactive est utilisée (Yang, Wang, & Kravets, September 2005). Elle assure la propriété suivante : si à deux chemins (ab) et (cb), on ajoute

un même troisième chemin (bd) alors la relation d'ordre entre leur métrique est conservée. Par conséquent, si $m(ab) < m(cb)$ alors $m(abd) < m(cbd)$.

Les auteurs de (Yang, Wang, & Kravets, September 2005) établissent que seulement quatre métriques sont réellement significatives dans le choix du meilleur chemin dans un réseau ad-hoc; la distance, la qualité du lien, la capacité du lien et le niveau d'interférence physique. La métrique de distance la plus utilisée dans les réseaux MANET est le nombre de sauts. La qualité du lien est représentée par le nombre de pertes que subit celui-ci. Les interférences physiques ont lieu lorsque deux nœuds transmettent des données dont les signaux se superposent sur le support entravant ainsi la bonne réception du paquet. Les auteurs écartent la charge, car elle induit des instabilités de chemins en raison de sa grande variabilité. Ainsi, les métriques les plus populaires prennent en considération ces paramètres. Par exemple, la métrique *Expected Transmission Count* (ETX) (De Couto, Aguayo, Bicket, & Morris, 2005) prend en compte le nombre de sauts du chemin ainsi que la qualité du lien. La qualité est mesurée à l'aide de messages sondes qui sont échangés entre les nœuds puis l'on calcule le taux de paquets reçus. La métrique utilisée dans le récent standard 802.11s (Hiertz, et al., 2010), appelée *AirTime Link Metric*, est très similaire à ETX. Dans les réseaux MANET, la mobilité des nœuds entraîne de nombreux changements de topologie qui affectent les performances des réseaux. Par conséquent, l'une des métriques les plus efficaces serait la stabilité des liens. Cependant, l'obtention de cette métrique nécessite une connaissance de la topologie courante et de son évolution que nous sommes incapables d'obtenir pour notre scénario.

L'architecture MONET comporte plusieurs passerelles vers le réseau MANET avec des caractéristiques très différentes. En particulier, les différents terminaux satellite n'offrent pas les mêmes capacités. Il est donc primordial de tenir compte de la capacité ou de la charge sur les passerelles.

Le principal problème des mécanismes de partage de charge est leur instabilité. C'est pour cette raison que la charge des liens du MANET est le plus souvent écartée des métriques intervenant dans le protocole de routage. En outre, le partage de charge sur plusieurs chemins, entre deux nœuds du MANET ne conduit qu'à pas ou peu de gain. En effet dans (Pearlman, Haas, Sholander, & Tabrizi, August 2000) (Nasipuri & Das, October 1999), les auteurs expliquent que les chemins ne sont pas assez distants les uns des autres pour qu'une répartition entre eux ne supprime les congestions ou même ne les diminue sensiblement. En effet, la zone d'interférence des deux chemins se superpose la plupart du temps. Ceci est principalement vrai pour des nœuds utilisant un unique canal. Cependant, l'inefficacité du partage de charge n'est plus vraie dans le cadre de la sélection de passerelle. En effet, les chemins menant aux passerelles sont suffisamment disjoints pour permettre une réelle amélioration des performances. Cette constatation permet de réduire les congestions qui surviennent à l'intérieur du MANET, mais le partage de charge va permettre aussi de diminuer les pertes dues aux congestions sur l'interface satellite. Cependant, le problème d'instabilité du routage reste entier.

6.2.4 Le partage de charge

Le partage de charge se découpe en trois étapes : la phase de mesure, la décision du partage de charge puis son exécution. La phase de mesure permet de calculer les différentes métriques nécessaires au mécanisme de sélection de passerelle. La phase de décision va permettre de définir le

moment où le changement de passerelle est nécessaire, en fonction des valeurs fournies par la phase de mesure. La dernière étape représente la véritable mise en œuvre de la modification.

6.2.4.1 Phase de mesure et les métriques

La phase de mesure de la charge peut prendre plusieurs formes. En effet, les moyens pour estimer la charge du réseau vont principalement dépendre du type de réseau (filaire ou sans fil) ainsi que du type de trafic que l'on veut estimer.

6.2.4.1.1 L'estimation indépendante du trafic

L'une des métriques les plus simples pour estimer la charge se réduit à mesurer soit le nombre de nœuds, soit le nombre de flux servis par une passerelle. Cette estimation ne prend aucunement en compte la disparité du trafic. La densité des nœuds est utilisée dans des articles théoriques tels que (Le-Trung, Engelstad, Skeie, & Taherkordi, November 2008). Le nombre de flux permet de faire du partage de charge dans certaines études, tout particulièrement lorsque les flux utilisent comme protocole de transport TCP (Galvez, Ruiz, & Skarmeta, 2012). En effet, l'utilisation des ressources par ces flux ne peut être que difficilement prédite, car ils possèdent comme objectif d'utiliser toute la bande passante. Une mesure instantanée du débit n'a donc que peu d'intérêt.

6.2.4.1.2 La mesure instantanée du débit

Les mesures instantanées de débit ne sont pas très significatives pour mesurer la charge sur l'interface radio d'un MANET. En effet, il est très difficile de réellement estimer l'étendue des congestions dans le MANET. Elles sont le plus souvent utilisées pour connaître la charge sur la passerelle sur son lien filaire ou dans des études théoriques.

Capacité résiduelle

La charge résiduelle (Ancillotti, Bruno, & Conti, June 2010) est obtenue en mesurant le volume de données reçues par la passerelle pendant un intervalle de temps défini. Puis on retranche la valeur de la capacité totale du lien.

Proportion de capacité consommée

De nombreux papiers reposent sur la proportion de la capacité du lien qui est consommée par les différents flux. Dans ce cas, les passerelles mesurent le débit grâce au volume de paquets traités par la passerelle, pendant une durée prédéfinie. Puis, elles divisent cette valeur par la capacité totale du lien.

6.2.4.1.3 La charge mesurée dans le MANET

Dans ce paragraphe, on prend en compte les paramètres qui sont réellement représentatifs de la charge à l'intérieur du MANET.

La durée de contention

La durée de contention représente la proportion du temps passé à attendre que le medium soit libre. Cette solution permet de prendre en compte la charge induite par les nœuds situés dans le voisinage (Nandiraju, Santhanam, Nandiraju, & Agrawal, October 2006) (Lee & Riley, March 2005) (Huang, Lee, & Tseng, 2004).

Les files d'attente

L'une des méthodes les plus répandues pour mesurer la charge dans un MANET utilise le taux de remplissage des files d'attente au niveau de la couche MAC. Cette donnée symbolise les phénomènes de congestion dans le MANET qui sont une cause importante de pertes de paquets.

La mesure du RTT

Des estimations de la charge peuvent être réalisées grâce à la variance du délai entre l'émission de la requête d'une sélection de passerelle réactive jusqu'à la réception de la réponse comme dans (Prasad & Wu, 2006) (Brännström, Ahlund, & Zaslavsky, November 2005) (Ramachandran, Buddhikot, Chandranmenon, Miller, Belding-Royer, & Almeroth, September 2005). Cette technique prend en compte le délai induit par l'attente dans les files d'attente sur l'ensemble du chemin. Par conséquent, si les délais de ces messages augmentent alors la charge sur le chemin est plus importante.

6.2.4.2 *La décision du partage de charge*

La décision du partage de charge peut se décomposer en deux familles, le partage de charge continu et le partage de charge événementiel.

Le partage de charge continu

Ce type de mécanisme a pour objectif de répartir la charge de la manière la plus équitable possible entre les passerelles. Cette technique ne repose donc que sur les algorithmes de minimisation du coût du chemin emprunté pour atteindre les passerelles. Elle a comme inconvénient de modifier les chemins très fréquemment. En outre, ces changements peuvent être néfastes puisque le partage de charge peut se faire à l'encontre de la métrique de la distance ; ce qui peut entraîner une dégradation des performances. C'est pour cela que le paramètre de longueur du chemin est souvent intégré au coût. Si ce n'est pas le cas, des algorithmes réduisent la portée du partage de charge à une zone plus ou moins large de nœuds limitrophes (Pham V. , Larsen, Kure, & Engelstad, November 2010). Ceci a pour objectif de limiter cette caractéristique néfaste sans l'éliminer (paragraphe 6.2.4.3.3).

Le partage de charge événementiel

Les algorithmes ne sont déclenchés que lorsqu'un événement survient, en particulier l'apparition d'une surcharge sur un chemin ou sur une passerelle. L'avantage de cette technique est de réduire les changements intempestifs de passerelle. Cependant, la définition de l'événement est très importante pour suffisamment anticiper la congestion et ne pas subir ses effets néfastes. Il est le plus souvent déclenché par le dépassement d'un seuil, par exemple, lorsque le débit mesuré sur une passerelle excède 90% de sa capacité totale (Park, Lee, Lee, & Shin, 2006).

6.2.4.3 *L'exécution du partage de charge*

6.2.4.3.1 Les entités responsables de l'exécution

Le partage de charge centralisé

Le partage de charge centralisé offre le plus de possibilités. En effet, il permet de concentrer ainsi toutes les métriques nécessaires pour l'exécution du partage de charge. Il est cependant peu utilisé dans la littérature puisque il est difficile d'obtenir une entité centrale dans un réseau MANET. Il est principalement utilisé dans un contexte théorique tel que dans les articles (Le-Trung, Engelstad, Skeie, & Taherkordi, November 2008) (Hoffmann & Medina, June 2009) (Tokito, Sasabe, Hasegawa, & Nakano, March 2009).

Le partage de charge distribué

Dans la majorité des cas, l'exécution de l'algorithme est réalisée par le nœud source situé à l'intérieur du MANET. Il exécute ainsi des algorithmes pour trouver le plus court chemin. La passerelle peut être partiellement incluse dans la phase d'exécution, en modifiant certains paramètres comme la valeur de certaines métriques. Par exemple, dans (Shin, Lee, Na, Park, & Kim, September 2004), la passerelle va retarder la réponse à une GW-SOL pour réduire la probabilité qu'elle ne soit choisie.

6.2.4.3.2 La granularité du partage de charge

La granularité du partage de charge peut être très variée. Il peut être effectué en considérant les paquets, les flux, les nœuds et les domaines.

Par paquet

Le partage de charge par paquet est très rare puisqu'il va entraîner une réception désorganisée des paquets d'un même flux. Le déséquilibrage des flux TCP lui est très néfaste. En effet, si les déséquilibrages ne sont pas gérés par une version améliorée de TCP, les performances des flux seront diminuées. L'une des seules propositions qui, à notre connaissance, propose de réaliser un partage de charge par paquet sur chaque nœud est (Huang, Lee, & Tseng, 2004). Tous les nœuds vont donc envoyer une proportion de trafic équivalente aux différences de capacité vers toutes les passerelles. Cette solution ne prend aucunement en compte la métrique de distance et va occasionner des congestions dans le MANET. Ceci suppose que les nœuds transmettent simultanément vers toutes les passerelles. Pour assurer que le choix de la passerelle n'est réalisé que par le nœud source, les auteurs utilisent des tunnels entre le nœud source et la passerelle (Figure 78). Ces tunnels ne sont utiles que pour le lien montant puisque dans le sens descendant, le routage entre la passerelle et le nœud source est réalisé uniquement par le protocole de routage du MANET.

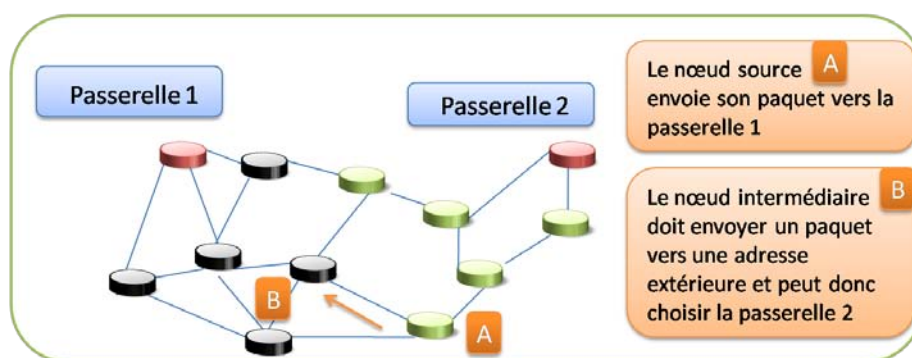


Figure 78 Incohérence de la sélection de passerelle

Par flux

Le partage de charge par flux reste assez limité jusqu'à présent. En effet, de même que pour celui par paquet, l'utilisation de tunnel entre la source et la passerelle semble obligatoire pour assurer la continuité du routage du flux vers la même passerelle. Le surplus d'*overhead* en fait une technique très peu utilisée, excepté lorsque le trafic que l'on veut partager est constitué de flux TCP. En outre, la majorité de ces solutions repose sur un partage de charge centralisé (Galvez, Ruiz, & Skarmeta, 2012) (Ancillotti, Bruno, & Conti, June 2010).

Par nœud

L'ajout du tunnel IP est toujours obligatoire et, dans la majorité des cas, ce partage de charge est centralisé (Hoffmann & Medina, June 2009). Le nombre de ces solutions reste faible. La granularité de ce partage de charge est plus grande par rapport au partage par flux, cependant les algorithmes et les mesures sont légèrement plus simples.

Par domaine

Cette solution est la plus répandue (Le-Trung, Engelstad, Skeie, & Taherkordi, November 2008) (Tada & Yamamoto, December 2009). En effet, elle permet de s'abstenir d'utiliser des tunnels, pour assurer le routage vers la passerelle choisie par le nœud source. En effet, elle assure que chaque nœud sur le chemin envoie les paquets vers la même passerelle. Cette caractéristique est obtenue grâce à la propriété d'isotonie des métriques. Ainsi, tous les mécanismes de sélection de passerelle qui sont fondés sur des métriques isotoniques peuvent être classés dans le partage de charge par domaine.

6.2.4.3.3 La zone du partage de charge

Dans la majorité des cas, les propositions effectuées dans la littérature proposent un mécanisme de partage de charge qui est réalisable par l'ensemble des nœuds du réseau MANET. Cependant, certaines études réduisent le champ d'action de ce mécanisme en instituant une zone.

Zone de partage de charge fixe

Celle-ci est le plus souvent définie comme la zone intermédiaire entre les domaines de deux passerelles (Figure 79). Cette solution permet d'éviter des chemins trop longs, obtenus lors d'un déséquilibre entre les passerelles tel que présenté dans (Tada & Yamamoto, December 2009). Elle est principalement proposée dans les cas où la métrique ne prend pas en compte directement la distance. Les limites de ce domaine ne dépendent pas de la charge et sont définies par les métriques de distance. Par exemple, (Pham V. , Larsen, Kure, & Engelstad, November 2010) propose de n'effectuer le partage de charge que sur une zone de deux sauts à la frontière des différents domaines des passerelles.

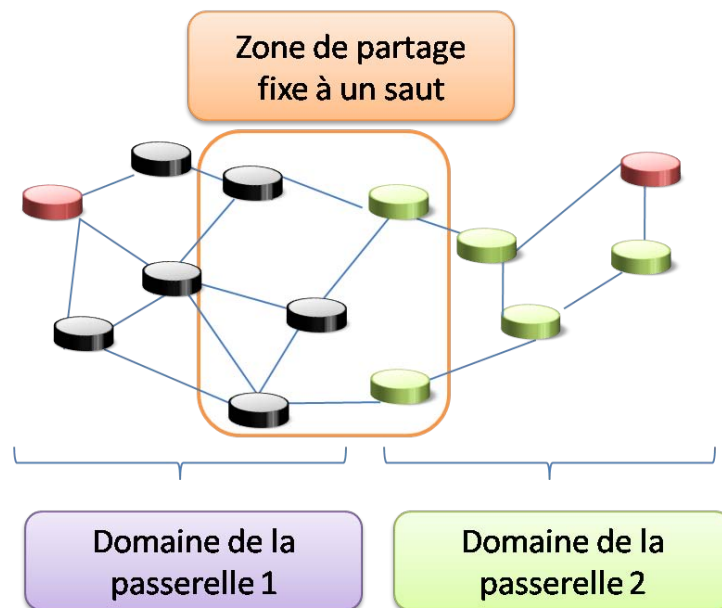


Figure 79 Zone de partage fixe à un saut

Zone de partage de charge progressive

Dans certains cas, la zone du partage de charge est toujours localisée dans le voisinage de la frontière du domaine. Cependant, cette frontière va se déplacer à chaque itération de l'algorithme en fonction des nœuds qui intègrent ou qui partent du domaine (Figure 80). L'exécution du partage de charge est donc progressive. À chaque itération de l'algorithme, le domaine peut s'agrandir et ainsi, de nouveaux nœuds sont intégrés à la zone de partage de charge. Par exemple, seuls les nœuds à la frontière sont autorisés à changer de passerelle. Ceci peut être obtenu par un message GW-ADV qui ne serait relayé que par les nœuds appartenant au domaine de la passerelle dont le message est originaire (Xie, Yu, Kumar, & Agrawal, December 2006).

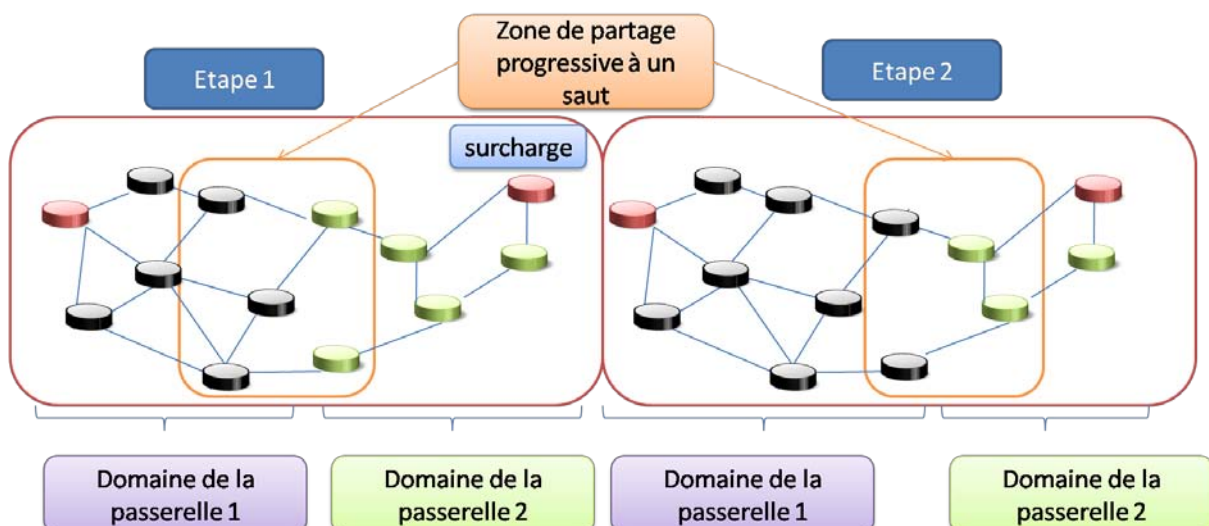


Figure 80 Zone de partage progressive à un saut

6.2.5 Les problèmes récurrents des mécanismes de partage de charge

De nombreux articles portant sur le partage de charge entre différentes passerelles comportent des défauts qui biaisent l'appréciation des algorithmes. Ils en existent deux qui nous paraissent très gênants : l'absence de mécanisme contre les oscillations de passerelle et le minimalisme en termes de diversité de scénarios de trafic dans les simulations.

6.2.5.1.1 Les problèmes d'oscillation

Le phénomène d'oscillation va se matérialiser par une alternance très fréquente dans le choix de la passerelle par un nœud. Il survient lorsque l'on détermine la passerelle de manière périodique. Ce sont donc les sélections de type proactif qui sont principalement affectées par ce phénomène. La Figure 81 présente un exemple d'oscillation de passerelle. Dans l'étape 1, la passerelle 1 se retrouve surchargée à cause des flux en provenance des nœuds A, B et C. Le mécanisme de partage de charge va remédier à cela en faisant basculer une partie du trafic vers la passerelle 2. Ensuite, pendant l'étape 2, les nœuds B et C transmettent les paquets à destination de la passerelle 2. Ces nouveaux flux entraînent une surcharge sur celle-ci. De la même manière que lors de la première étape, le partage de charge va impliquer un retour à l'étape 1. La sélection de passerelle va donc osciller à chaque itération de l'algorithme.

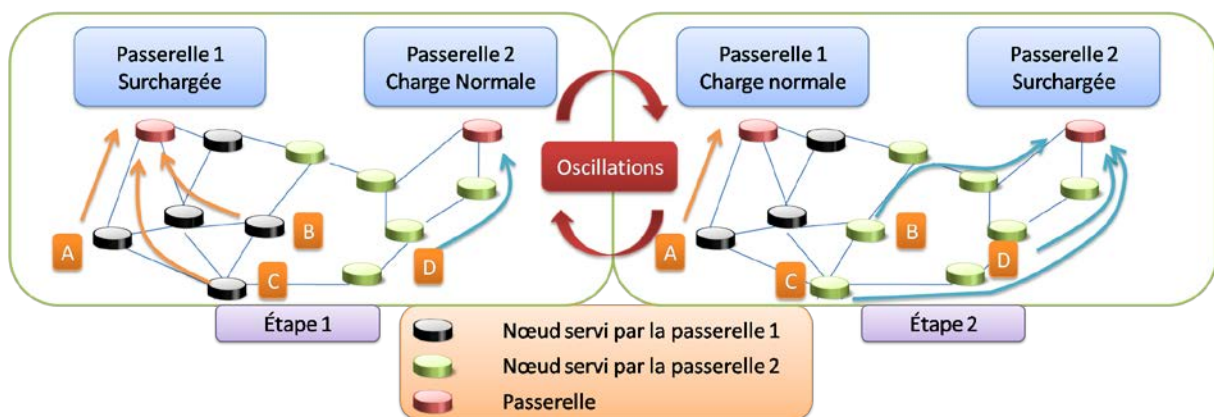


Figure 81 Exemple d'oscillation de passerelle

Le problème de l'oscillation de passerelle est, dans une grande partie des cas, ignoré ou sous-estimé. Quatre techniques sont mises en place, à notre connaissance, pour y remédier. Certaines études utilisent conjointement plusieurs techniques pour une meilleure efficacité.

Une temporisation à chaque sélection de passerelle

La première solution temporise les changements répétitifs de passerelle. La modification du choix de la passerelle entraîne une temporisation pendant laquelle aucune autre sélection n'est autorisée. Le principal défaut de cette technique survient si le changement de passerelle n'est pas opportun. Il implique à ce moment-là une baisse des performances. Cet état désavantageux durera alors tout le temps de la temporisation. Dans la Figure 81, si l'étape 2 est plus néfaste que l'étape 1 en termes de pertes, la temporisation va tout de même éviter qu'il ne retourne dans l'étape 1.

Une valeur d'hystérésis sur la charge

La deuxième solution repose sur la définition d'une hystérésis. La comparaison des charges entre les chemins vers les différentes passerelles n'entraîne de modification que si la différence entre les charges est supérieure à la valeur de l'hystérésis (Pham V. , Larsen, Kure, & Engelstad, November 2010). Cette technique permet de réduire le nombre d'oscillations de passerelle. Cependant, l'efficacité de ce mécanisme dépend de la quantité de la charge qui va être déviée vers la nouvelle passerelle. Si elle est largement supérieure à la valeur de l'hystérésis, cette solution se révélera inefficace. Il est très délicat de connaître cette valeur de manière décentralisée. Si, par exemple, dans la Figure 81, les flux en provenance des nœuds A et D ont des débits similaires, la différence de charge entre les passerelles n'est la conséquence que des trafics des nœuds B et C. Si ce sont eux qui entraînent une charge supérieure à la valeur de l'hystérésis, l'oscillation se reproduira.

Réduction de la granularité du partage de charge

Ce type de mécanismes a été proposé dans (Maurina, Riggio, Rasheed, & Granelli, September 2009) et (Nandiraju, Santhanam, Nandiraju, & Agrawal, October 2006). Les passerelles vont isoler les nœuds qui sont potentiellement la cause de la surcharge. Puis, elles vont envoyer un message en unicast à ces nœuds, pour leur intimer de changer de passerelle. Ces techniques permettent de diminuer le nombre potentiel de nœuds qui changent de passerelle et ainsi réduire les occurrences du phénomène d'oscillation.

Détection et suppression

Cette dernière technique met en place des mécanismes de détection d'oscillation de passerelle pour ensuite bloquer le système dans l'état que l'algorithme juge le plus bénéfique. Cette technique est, à notre connaissance, uniquement utilisée par des mécanismes de sélection de passerelle centralisés. En effet, ceux-ci permettent d'avoir une connaissance de toutes les métriques des différents chemins. Ils peuvent donc connaître la part de trafic qui va être déviée vers une autre passerelle et ainsi anticiper l'occurrence de phénomène d'oscillation. Ils choisiront entre les deux solutions, celle qui implique le moins de surcharge. Cette solution est bien sûr la plus efficace, mais elle nécessite une centralisation de l'algorithme de sélection.

6.2.5.2 *Les problèmes de simulation*

Il est très difficile de voir la réelle efficacité d'un mécanisme de sélection de passerelle à cause du manque d'exhaustivité des simulations. De nombreuses propositions ne sont accompagnées que de simulations, pour lesquelles le scénario est précisément créé pour mettre en valeur les bénéfices de l'algorithme de partage de charge. En particulier, l'un des points les plus importants pour vérifier l'efficacité de l'algorithme se situe dans la diversité des scénarios de trafics que l'on a simulés. Cela permet de vérifier tout particulièrement l'impact du phénomène d'oscillation.

Dans l'article (Pham V. , Larsen, Engelstad, & Kure, October 2009), les auteurs soulèvent ce point important. En effet, l'efficacité du mécanisme de partage de charge n'est efficace que dans certaines conditions. Ils utilisent des mécanismes à la fois fondés sur la distance en nombre de sauts et la charge sur la passerelle, pour établir la dépendance du mécanisme vis-à-vis du scénario de trafic, comme par exemple l'inefficacité du partage de charge lorsque la topologie des nœuds est uniforme.

Le gain n'est en moyenne que de 1%. Ils concluent en définissant les paramètres qui vont influencer son efficacité.

Besoin de capacité

Le partage de charge n'est utile que lorsqu'un chemin vers une passerelle rencontre un problème causé par une surcharge. En effet, il n'est pas nécessaire de modifier le chemin si aucune des deux passerelles ne rencontre de surcharge.

La distance entre les passerelles

La distance entre les passerelles est très significative quant à l'efficacité du partage de charge. En effet, si les passerelles sont proches, les chemins vers les deux passerelles vont l'être également. Les zones d'interférences vont alors se chevaucher. Ainsi le partage de charge entre les passerelles va avoir les mêmes défauts que le multi-chemin à l'intérieur du MANET.

La répartition du trafic

La répartition du trafic et la topologie du MANET sont très importantes puisque qu'elles vont accentuer ou diminuer l'asymétrie de charge entre les différentes passerelles et ainsi, augmenter l'efficacité du partage de charge. La topologie du réseau va intervenir dans la répartition du trafic.

6.2.6 Descriptions des principaux mécanismes partage de charge.

Le mécanisme proposé dans (Shin, Lee, Na, Park, & Kim, September 2004) est réactif. Les auteurs vont retarder la réponse à la GW-SOL en fonction de la charge sur la passerelle. Ce délai va permettre de pénaliser de manière indirecte cette passerelle puisque les réponses des autres passerelles sont favorisées ; cela va aussi compenser leur éloignement. Ce mécanisme n'a pas besoin de technique pour éviter l'oscillation de passerelle, car il n'y a pas de répartition de charge périodique. Dans ce cas, le choix est fait durant la découverte de la passerelle. Si aucune mise à jour de la passerelle n'est effectuée, le choix fait à la première et seule demande risque de ne plus correspondre à l'état réel du réseau.

(Le-Trung, Engelstad, Skeie, & Taherkordi, November 2008), (Galvez, Ruiz, & Skarmeta, 2012) et (Hoffmann & Medina, June 2009) sont tous trois des études théoriques d'un mécanisme de partage de charge centralisé. L'un des intérêts de la centralisation provient de la connaissance générale que possède cette entité. Cela permet entre autre de pouvoir supprimer les effets d'oscillation. (Le-Trung, Engelstad, Skeie, & Taherkordi, November 2008) se fonde sur trois métriques : la distance en nombre de sauts, la charge à destination de l'extérieur, qui correspond dans cette étude au nombre de nœuds que sert une passerelle, et la charge à l'intérieur du MANET, qui est obtenue grâce à la densité des nœuds sur le chemin. Il ne prend donc pas en compte la disparité des ressources consommées par les différents trafics. (Hoffmann & Medina, June 2009) prend en compte les débits réels dans le réseau. Cependant, les considérations théoriques ne sont pas applicables de manière réaliste dans un réseau MANET puisque l'algorithme cherche à optimiser l'utilisation des ressources dans le MANET en programmant la transmission de tous les nœuds et le choix de leur passerelle.

Très peu de mécanismes s'intéressent au partage de charge lorsque des flux TCP sont à prendre en considération. En effet, leur consommation de ressources n'est pas constante et le partage de charge ne peut être réalisé en se fondant sur des mesures instantanées. Ainsi, dans (Galvez, Ruiz, & Skarmeta, 2012), les auteurs proposent une architecture centralisée, dont le partage de charge repose sur le nombre de flux. Ils proposent d'utiliser des algorithmes utilisés dans l'ordonnancement des tâches dans les systèmes parallèles pour résoudre le problème du choix de passerelle.

(Tokito, Sasabe, Hasegawa, & Nakano, March 2009) est aussi une étude théorique, cependant elle n'est pas centralisée. Elle tient compte de la différence générale entre la charge des différentes passerelles. Une zone de partage de charge fixe est mise en place. Ainsi, les nœuds situés dans cette zone peuvent-ils choisir leur passerelle avec pour objectif de diminuer les différences de charge.

De la même manière, (Pham V. , Larsen, Kure, & Engelstad, November 2010), (Tada & Yamamoto, December 2009) et (Iqbal & Kabir, December 2011) proposent des algorithmes qui comportent une zone fixe de partage de charge. (Pham V. , Larsen, Kure, & Engelstad, November 2010) utilise le temps de contention comme estimation de la charge. Le coût du lien est la valeur maximale de cette métrique parmi tous les liens du chemin. Lorsqu'un nœud se situe dans la zone de partage de charge, il choisit comme passerelle celle dont le chemin possède le coût le plus faible. Dans le cas contraire, le choix est seulement sujet au nombre de sauts. Pour éviter le phénomène d'oscillation, les auteurs mettent en place un seuil de charge. C'est seulement quand la différence de coût des chemins est supérieure à ce seuil que le nœud est autorisé à changer de passerelle. (Tada & Yamamoto, December 2009) définit une métrique fondée sur le nombre de nœuds raccordés à une passerelle. Le coût lié à celle-ci est égal à la somme des coûts des chemins actifs qui sont reliés à cette passerelle. Ainsi, il est dépendant du nombre de nœuds actifs, mais aussi de la longueur des chemins. Tout comme le précédent article, il définit une zone pour interdire le partage de charge sur des chemins de longueur trop importante. (Iqbal & Kabir, December 2011) a créé une métrique, reposant sur le nombre de sauts, la charge mesurée par l'intermédiaire des files d'attente et la densité des nœuds autour du chemin (Équation 6.1).

$$(6.1) \quad \text{Coût} = \text{nombre}_{\text{sauts}} + \frac{\text{Charge}_{\text{Filesd'attente}}}{\text{Charge}_{\text{Filesd'attente}}+1} + \frac{\text{Densité}_{\text{nœuds}}}{\text{Densité}_{\text{nœuds}}+1}$$

La particularité des deux derniers paramètres vient du fait qu'ils ne peuvent dépasser 1. L'influence de la densité ainsi que de la charge ne peut excéder une zone équivalente à deux sauts. Aucun mécanisme n'est mis en place pour éviter les oscillations de passerelle.

Très peu d'algorithmes retiennent en place ce type de mécanisme. (Fu, Kwang-Mien, Tan, & Boon-Sain, May 2006) ignore de la même manière que les précédents algorithmes le phénomène d'oscillation. Il met en place une méthode proactive où il introduit la notion de sauts virtuels. La charge est traduite en termes de sauts virtuels pour éloigner la passerelle surchargée (Équation 6.2). Le coût des chemins vers l'extérieur est égal à la distance du chemin entre la source et la passerelle à laquelle le nombre de sauts virtuels est ajouté.

$$(6.2) \quad \text{Saut}_{\text{virtuel}} = K \text{Charge}_{\text{fileAttente}}$$

(Maurina, Riggio, Rasheed, & Granelli, September 2009) et (Nandiraju, Santhanam, Nandiraju, & Agrawal, October 2006) prennent en compte ce problème. Le principe est le même dans les deux articles. La passerelle décide de se séparer des nœuds qui lui posent le plus de problèmes. Par

conséquent, elle transmet, en unicast et à ces nœuds, le conseil de trouver une autre passerelle. La différence entre les deux études se situe dans le choix des nœuds qui doivent changer de passerelle. Dans (Nandiraju, Santhanam, Nandiraju, & Agrawal, October 2006), ce choix est uniquement fondé sur la charge induite par les nœuds. Elle est mesurée grâce aux files d'attente alors que (Maurina, Riggio, Rasheed, & Granelli, September 2009) y ajoute la métrique ETX pour prendre en compte la longueur du chemin et sa qualité.

(Nandiraju, Santhanam, Nandiraju, & Agrawal, October 2006) et (Park, Lee, Lee, & Shin, 2006) sont les seuls algorithmes présentés à réaliser une exécution événementielle. En effet, le partage de charge n'est activé que lorsque la passerelle est surchargée alors que les autres algorithmes tendent à obtenir une répartition équitable permanente de la charge entre les passerelles. Les deux mécanismes vont rediriger le trafic lorsque les passerelles sont surchargées. (Park, Lee, Lee, & Shin, 2006) utilise la métrique de la capacité résiduelle. Si elle est inférieure à 5% de la capacité du lien de la passerelle alors le partage de charge est déclenché.

(Xie, Yu, Kumar, & Agrawal, December 2006) et (Huang, Lee, & Tseng, 2004) ont la particularité de réaliser un partage de charge progressif. En effet, à chaque itération du protocole seuls les nœuds en bordure du domaine peuvent changer de passerelle. Dans (Xie, Yu, Kumar, & Agrawal, December 2006), on ajoute la possibilité de se connecter à deux passerelles en même temps.

(Ancillotti, Bruno, & Conti, June 2010) est le seul article que nous citons qui propose une étude expérimentale de son mécanisme. Il définit deux métriques. La première est égale au maximum de la capacité résiduelle sur les liens extérieurs. La seconde reprend la même métrique à l'exception que celle-ci est divisée par le coût du chemin fourni par la métrique ETX. Les auteurs ont retenu peu de scénarios, en raison de la difficulté à en reproduire de nombreux lors de tests en grandeur nature.

Cependant, peu d'articles fournissent des simulations assez complètes pour estimer la pertinence de l'algorithme proposé. En effet, les scénarios de trafic sont la plupart du temps très réduits et sont restreints à ceux qui mettent en exergue le gain obtenu. Au-delà de la maigre variété de scénarios de trafic, peu d'articles montrent les variations obtenues lors de plusieurs simulations possédant des paramètres aléatoires. Par exemple, (Huang, Lee, & Tseng, 2004) et (Hoffmann & Medina, June 2009) utilisent tous deux une distribution uniforme du trafic entre les nœuds. Il n'y a donc aucun changement de trafic lors des simulations. Dans (Le-Trung, Engelstad, Skeie, & Taherkordi, November 2008) et (Nandiraju, Santhanam, Nandiraju, & Agrawal, October 2006), un seul scénario prédéfini est choisi. (Nandiraju, Santhanam, Nandiraju, & Agrawal, October 2006) définit statiquement 3 nœuds sources et (Le-Trung, Engelstad, Skeie, & Taherkordi, November 2008) définit statiquement 3 sources TCP et 3 sources UDP. Dans (Tada & Yamamoto, December 2009), les sources sont choisies aléatoirement, mais elles dépendent du domaine de la passerelle. Ainsi, 8 nœuds sources sont élus dans le domaine de la première passerelle et 2 dans l'autre.

Le tableau ci-dessous récapitule les différentes caractéristiques des principaux mécanismes de partage de charge sur les différentes passerelles du MANET.

	Intégration à interne t	Méthode de	Métrique de charge	Autre métrique	Oscillations de passerelle	Simulations	Particularités	Centralisé ou Décentralisé
(Shin, Lee, Na, Park, & Kim, September 2004)	MIP	Réactive	File d'attente	-	Non	?	Retarde les réponses	D

							aux GW-SOL en fonction de la charge	
(Le-Trung, Engelstad, Skeie, & Taherkordi, November 2008)	MIP	-	Nombre de nœuds servis	Sauts + densité des nœuds au voisinage	Non	Un seul scénario	Théorique	C
(Ancillotti, Bruno, & Conti, June 2010)	NAT	Réactive	Charge résiduelle	Sauts	Non	Expérimental	Pour les flux TCP	C
(Pham V. , Larsen, Kure, & Engelstad, November 2010)	-	Proactif	Taux de contention	Sauts	Hystérésis	30 scénarios Intervalle de confiance 95%	Zone de partage fixe	D
(Pham V. , Larsen, Kure, & Engelstad, November 2010)	MIP	Réactive	Nombre de nœuds servis	Nombre de sauts	Non	1 simulation	Zone de partage fixe	D
(Hoffmann & Medina, June 2009)	-	-	Proportion de charge	Nombre de sauts	Central	Intervalle de confiance 95%	Théorique	C
(Iqbal & Kabir, December 2011)	-	Proactive	File d'attente	Sauts + Nombre de nœuds voisins	Non	10 simulations par scénario	Zone de partage fixe	D
(Fu, Kwang-Mien, Tan, & Boon-Sain, May 2006)	-	Proactive	File d'attente + contention	Nombre de sauts	Non	-		D
(Maurina, Riggio, Rasheed, & Granelli, September 2009)	MIP	Réactive	File d'attente	ETX	Limitation des nœuds affectés par le partage	1 scénario		D
(Galvez, Ruiz, & Skarmeta, 2012)	-	-	Nombre de flux	Longueur du chemin	temporisation	5 scénarios Intervalle de confiance 95%	Théorique	D
(Tokito, Sasabe, Hasegawa, & Nakano, March 2009)	-	-	Débit instantané	Nombre de sauts	Centralisé	Intervalle de confiance 95%	Théorique	C
(Tada & Yamamoto, December 2009)	MIP	Réactive	Nombre de nœuds	Nombre de sauts	Non	1 scénario	Zone de partage fixe	D
(Nandiraju, Santhanam, Nandiraju, & Agrawal, October 2006)	-	Réactive	File d'attente	-	Limitation des nœuds affectés par le partage	1 scénario	événementiel	D
(Park, Lee, Lee, & Shin, 2006)	MIP	Hybride	Charge résiduelle	Nombre de	Non	plusieurs	événementiel	D

	elle	sauts	scénario					
(Xie, Yu, Kumar, & Agrawal, December 2006)	MIP	-	?	Nomb re de sauts	Non	2 scénario s	Zone progress ive	D
(Huang, Lee, & Tseng, 2004)	MIP	Proac tive	Débit instant ané	-	Temporis ation + seuil de charge	1 scénario	Zone progress ive	D

Tableau 21 Récapitulatif des algorithmes de sélection de passerelle selon la charge

6.3 Les spécificités de MONET pour la sélection de passerelle

L'architecture définie dans le projet MONET induit de nombreuses conséquences dans la mise en place d'un mécanisme de sélection de passerelle, en particulier, en raison de l'hybridation satellite/MANET.

Besoin de partage de charge

Le partage de charge est primordial dans le réseau MONET. En effet, le bénéfice qu'il apporte est avéré, dans le cadre de la sélection de passerelle, pour limiter les congestions dans le MANET. En outre, il permet de gérer les disparités de flux entre les différentes passerelles. L'architecture de MONET augmente l'importance de ce mécanisme, car les capacités des passerelles sont très diverses. Les liens satellite pour relier les postes de commandement au quartier général atteignent des débits supérieurs à un mégabit par seconde alors que les technologies telles que le BGAN proposent des capacités variant de 384 kb/s à 492 kb/s pour la liaison montante et de 240 kb/s à 492 kb/s pour la liaison descendante.

Séparation des mécanismes de sélection de passerelle sur le lien montant et descendant

Les liens satellite sont asymétriques. Par conséquent, la charge sur le lien montant et descendant des passerelles sont différents. La séparation des mécanismes sur le lien montant et descendant semble nécessaire.

Besoin d'un échange d'information de routage vers le quartier général

Le protocole de routage à l'intérieur du quartier général nécessite l'obtention d'informations de la part du réseau MANET. Elles sont essentielles lorsqu'un MANET se sépare en deux réseaux disjoints. Par exemple, dans la Figure 82, le réseau MANET se sépare. Le quartier général routait les paquets à destination du nœud A par l'intermédiaire de la passerelle 1. Après la séparation, le nœud A ne peut être atteint que grâce à la passerelle 2. Pour pouvoir assurer la continuité du service, le quartier général doit être informé que le nœud A n'est maintenant atteignable que par la passerelle 2. Comme le routage ne peut être hiérarchique, on ne peut différencier deux sous-réseaux pour les deux MANET disjoints, les passerelles doivent obligatoirement transmettre les adresses des différents nœuds pour lesquels elles possèdent un chemin.

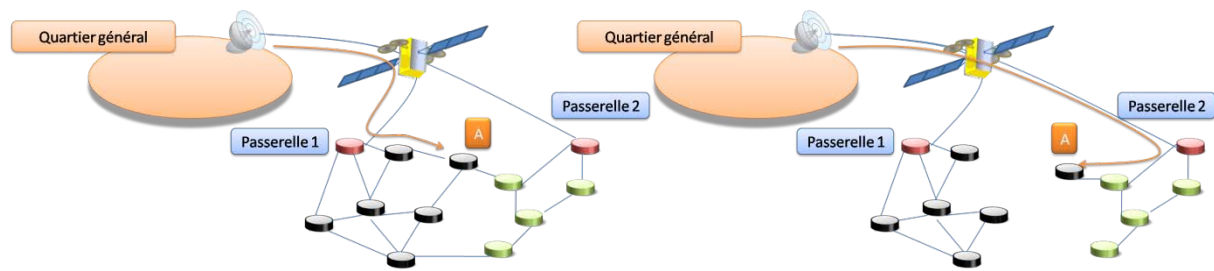


Figure 82 Séparation d'un MANET

6.4 Conclusion

De nombreux algorithmes ont été proposés dans la littérature, cependant aucun ne rassemble toutes les caractéristiques qui nous semblent les plus performantes. Dans un premier temps, l'utilisation d'OLSR nous pousse à utiliser une sélection de passerelle proactive. Celle-ci doit se fonder à la fois, sur la métrique du chemin à l'intérieur du MANET ainsi que sur la charge sur la passerelle. En effet, la première métrique nous permet d'éviter l'apparition de congestions à l'intérieur du MANET alors que la seconde a pour utilité de réduire les pertes sur le lien satellite, essentiellement sur un lien BGAN.

Dans un second temps, la décision devrait être événementielle. En effet, un partage de charge continu a pour objectif d'équilibrer les charges entre les passerelles, mais l'équilibre n'est pas directement lié au nombre de pertes. Si aucune des passerelles n'est surchargée et qu'aucune perte dans les liens satellite ne survient, il est alors inutile d'égaliser le trafic. En outre, cela va à l'encontre des métriques du MANET, comme par exemple le nombre de sauts, et cela peut augmenter les congestions à l'intérieur de celui-ci.

Le système satellite nous oblige à séparer le partage de charge sur le lien montant et sur le lien descendant.

L'exécution du partage de charge sur le lien descendant est le moins contraignant grâce à la présence d'une entité centrale dans le quartier général. Les granularités par flux, par nœuds et par domaines sont possibles sans causer d'*overhead* supplémentaire. Celle par paquets est écartée, car elle induit de trop nombreux déséquilibrages.

L'exécution du partage de charge sur le lien montant est plus complexe. En effet, il est obligatoirement distribué. La granularité par domaine est donc préférable dans l'optique de minimiser l'*overhead*. En outre, un mécanisme efficace contre les oscillations de passerelle doit être effectué alors qu'il est souvent absent ou incomplet, dans la majeure partie des solutions proposées dans la littérature.

Ces contraintes sur le partage de charge vont être prises en compte dans la définition de notre algorithme et son évaluation qui sont décrites dans le chapitre suivant.

Chapitre 7 : MONET : Optimisation du partage de charge parmi les passerelles satellite d'un MANET

Dans ce chapitre, nous décrivons le mécanisme du partage de charge sur les différentes passerelles satellite que nous avons proposé. La définition d'un mécanisme traitant du phénomène d'oscillation de passerelle a reçu une attention particulière. Dans la majorité des cas, il est considéré avec trop de légèreté, alors qu'il peut avoir des effets néfastes importants. L'algorithme proposé prend en compte les différentes contraintes de l'hybridation qui ont été mises en exergue dans le chapitre précédent. Dans une seconde partie, l'efficacité de cette solution est étudiée à l'aide de simulations.

7.1 L'algorithme d'optimisation

7.1.1 Principe

L'un des principaux objectifs est de définir un mécanisme de partage de charge qui prenne en compte le phénomène d'oscillation de passerelle. Nous nous focalisons, dans un premier temps, sur le mécanisme pour le trafic montant. Elle doit être mis en œuvre de manière distribuée comme nous venons de le voir. Dans ce cas, les entités les plus à même de détecter et gérer ce phénomène sont les passerelles, puisque ce sont elles qui possèdent le plus d'informations sur la topologie et la distribution du trafic. En effet, elles ont une connaissance complète de la topologie du MANET, grâce au protocole de routage MANET, et elles surveillent les flux des nœuds qu'elles servent. Par conséquent, nous désignons les passerelles comme responsables de la détection des oscillations de passerelle. Elles doivent donc aussi avoir une influence sur l'exécution du partage de charge.

Le principe que nous proposons est relativement simple. La passerelle même influence le choix de passerelle, en jouant sur les métriques à l'intérieur du MANET. Le partage de charge ne commencera que lorsqu'elle est surchargée ou sous-chargée. Notre algorithme est donc événementiel. Elle va augmenter ou diminuer le coût des chemins dont elle est la destination, en fonction de son état de charge et, par là même, augmenter ou réduire l'étendue de son domaine de manière progressive (Figure 83). Grâce à ce principe, la passerelle a la mainmise sur la phase d'exécution du partage de charge ; ce qui lui offre la possibilité de mettre en place un système contre l'oscillation. En outre, cette solution permet de définir une zone de partage de charge progressive qui diminue les effets néfastes de l'oscillation. Elle prend aussi en compte la longueur du chemin à l'intérieur du MANET, puisque l'augmentation de l'étendue du domaine se fait de manière progressive, par rapport à une frontière des domaines référence, définie par la métrique interne au MANET.

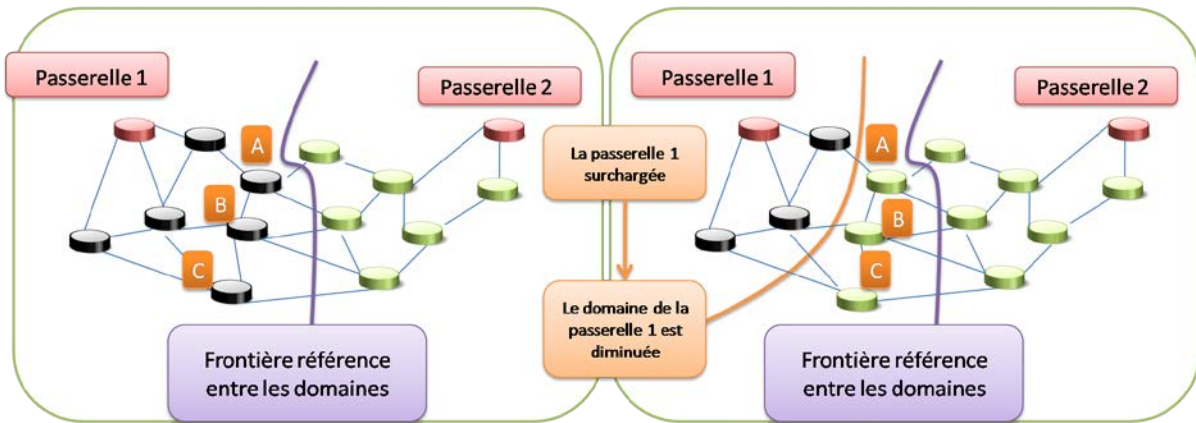


Figure 83 Principe du partage de charge proposé

Nous conservons ce principe pour le lien descendant pour maîtriser le coût des chemins dans le MANET et ainsi, réduire le nombre de congestions internes. Les mécanismes mis en place pour le lien descendant sont similaires à ceux utilisés pour le lien montant. Les deux principales différences se situent au niveau de la récupération des métriques internes au MANET présentée dans la partie 7.1.3, ainsi que la centralisation de toutes les informations. Ce dernier point évite en particulier certains messages de contrôle entre les passerelles.

7.1.2 La phase d'exécution

Nous présentons tout d'abord la phase d'exécution, car elle a un impact important sur la définition de la phase de mesure et de décision.

Le changement du coût du chemin par la passerelle

La phase d'exécution permet de réaliser le principe précédemment exposé. Nous souhaitons augmenter ou diminuer la taille des domaines servis par les passerelles en cas de sous-charge ou de surcharge. La solution que nous proposons influence le coût des chemins à destination d'une passerelle, en le corrélant à l'état de celle-ci. À chaque fois que la phase d'exécution est enclenchée, nous allons ajouter ou soustraire au coût du chemin une valeur fixe. Cette valeur est le pas de notre zone progressive. Les états de surcharge et de sous-charge sont définis par la phase de décision. À chaque itération, si la passerelle est surchargée, le domaine servi par celle-ci diminue de ce pas. Par conséquent, une nouvelle variable qui va représenter toutes les variations est créée. Elle est appelée correction, $Corr_{passerelle}$. Elle définit le surcoût induit par chaque passerelle (Équation 7.1). Lorsque celle-ci décide d'exécuter l'algorithme pour cause de surcharge, la valeur de $Corr_{passerelle}$ est incrémentée du pas (Équation 7.2). Le coût pour atteindre cette passerelle est alors augmenté (Figure 84) et elle sera donc moins sollicitée par la suite.

$$(7.1) \text{ Coût}_{passerelle} = \text{Coût}_{interne} + Corr_{passerelle}$$

$$(7.2) \text{ Corr}_{passerelle} = Corr_{passerelle} + Pas$$

Ce pas doit avoir la même valeur pour chaque passerelle dans le but d'éviter une progression différente des domaines entre elles. En effet, il est impossible de connaître, *a priori*, la charge qui va être apportée par la modification des valeurs de corrections, puisqu'une passerelle ne possède

d'informations que sur le trafic produit par les nœuds de son domaine. C'est pour cette même raison que ce pas reste fixe dans le temps.

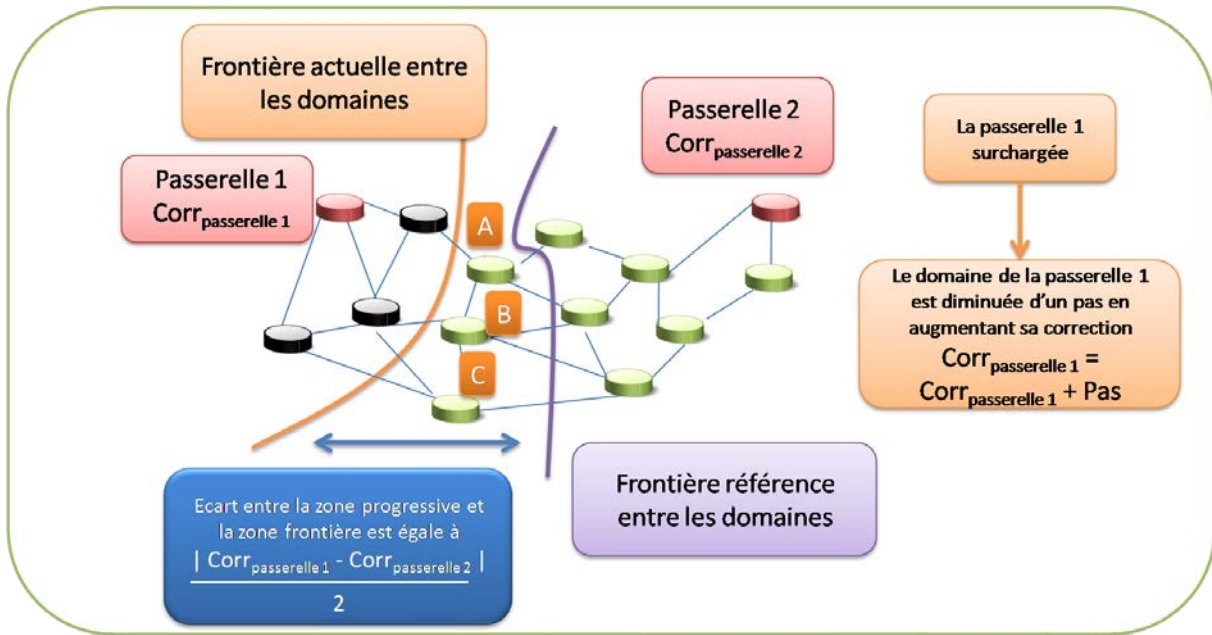


Figure 84 Gestion des domaines par la correction

Dans la Figure 84, nous utilisons le nombre de sauts comme métrique interne au MANET. La passerelle 1 est surchargée. Elle augmente donc la valeur de sa correction d'un saut pour restreindre son domaine. Ainsi, les nœuds A, B et C sont intégrés au domaine de la passerelle 2.

Nous conservons le protocole OLSR pour le routage dans le MANET. Tous les nœuds sur le chemin d'un paquet doivent donc le router et effectuer le choix de la passerelle. Par conséquent, ils doivent avoir connaissance du surcoût décidé par chaque passerelle, puis exécuter l'algorithme du plus court chemin. La passerelle diffuse la valeur de sa correction par l'intermédiaire du message de GW-ADV, appelé HNA dans OLSR. La périodicité du partage de charge coïncide donc avec la périodicité de l'envoi des messages HNA.

La limite de la valeur de correction

L'évolution des frontières entre les domaines des passerelles est dépendante de la différence de corrections entre deux passerelles (Figure 84). En effet, les nœuds comparent les coûts des chemins vers les différentes passerelles, la frontière référence est ainsi définie par l'égalité entre ces coûts. La variation de la frontière ($\Delta_{frontière}$) est donc égale à la différence entre les corrections des passerelles divisée par deux (Équation (7.3)). En effet, pour augmenter d'un saut, la différence de correction doit compenser le fait que le nœud se situe à un saut de plus, mais aussi qu'il est à un saut de moins de l'autre passerelle.

$$(7.3) \quad \Delta_{frontière} = \frac{|Corr_{passerelle\ 1} - Corr_{passerelle\ 2}|}{2}$$

La valeur de cette différence peut être faible tout en ayant des valeurs de correction élevées. Cependant, la valeur d'une correction est limitée par la taille du champ qui lui est attribuée dans

l'évolution du message HNA. Il est donc nécessaire de restreindre la valeur de la correction. Si elle devient égale à la valeur maximale, le mécanisme que nous proposons ne sera plus opérationnel.

Pour limiter l'augmentation de la valeur de la correction, les solutions reposent sur le fait que seule la différence des corrections a un impact sur le mécanisme (Équation (7.3)). Par conséquent, nous souhaitons conserver les différences de corrections, tout en minimisant les valeurs des corrections. Après la modification des corrections par l'algorithme, en fonction de la charge, la première solution va s'assurer que la correction la plus faible soit égale à zéro tout en conservant les différences entre les corrections de passerelle. Par exemple, dans la Figure 85, l'exécution de l'algorithme va augmenter la correction de la passerelle 2, puis la rectification va automatiquement soustraire la correction la plus faible à toutes les corrections de passerelle, pour obtenir une correction minimale de 0 et conserver ainsi des différences entre les corrections.

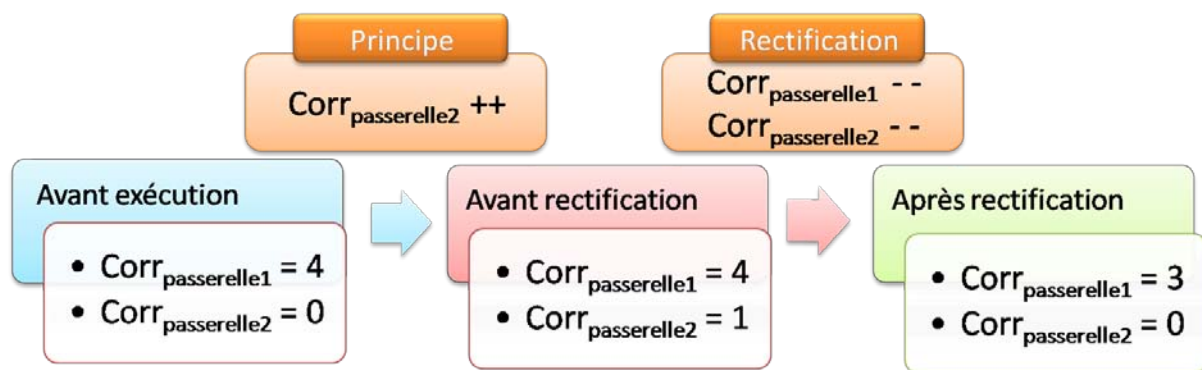


Figure 85 Exemple de rectification

Cependant, cette technique implique une gestion complexe. En effet, dès que la seule passerelle possédant une correction égale à zéro est surchargée, sa correction n'est pas incrémentée, mais celles des autres passerelles doit être décrétementée. Ce mécanisme met un certain temps à être mis en place, puisque il doit, dans un premier temps, ordonner aux autres passerelles le décrétement, puis celles-ci doivent alors disséminer la nouvelle correction. En outre, cette opération peut être très répétitive.

Pour contrecarrer son inconvénient, nous ne nous occupons de ce problème que lorsque l'une des corrections atteint la valeur maximale autorisée par le message HNA. Dans ce cas-là, les corrections sont recalculées pour que l'une d'elles soit égale à zéro, tout en conservant les écarts entre celles-ci. Par conséquent, un champ est ajouté dans le message HNA pour informer tous les nœuds du MANET de la nécessité de réaliser la rectification de toutes les valeurs de corrections. Dans chaque nœud y compris les passerelles, la valeur minimale des corrections est déduite de toutes les corrections. Cette solution est aussi utilisée lorsqu'une nouvelle passerelle apparaît. En effet, si la rectification des corrections n'était pas exécutée, la nouvelle passerelle utiliserait une correction égale à zéro et la différence entre les corrections serait trop importante. Par conséquent, le domaine de départ serait très étendu pour la nouvelle passerelle.

7.1.3 La phase de mesure

La phase de mesure est relativement simple. Deux types de métriques sont utilisés dans notre mécanisme de partage de charge. Le premier type concerne la charge des liens satellite et le second

est la métrique utilisée par le protocole de routage à l'intérieur du MANET. La charge va définir les variations du domaine, alors que la métrique interne permet de prendre en compte le coût du chemin interne et caractérise le domaine initial.

La charge sur le lien satellite

Les passerelles sont responsables des mesures de la charge sur le lien montant et le quartier général l'est pour le lien descendant. En effet, seules ces entités ont la possibilité d'effectuer les mesures de charge sur le lien satellite. Nous les effectuons à chaque extrémité du système satellite dans l'optique de minimiser l'*overhead*; nous évitons ainsi les transmissions inutiles d'information de charge.

La charge sur le lien satellite peut être calculée grâce au débit et à la capacité du satellite. Le débit se mesure en évaluant le volume de données pendant une fenêtre de temps coulissante de la capacité totale du lien satellite. Pour cette mesure, nous nous intéressons plus particulièrement aux flux UDP. En effet, nous considérons que les flux les plus prioritaires et possédant les contraintes les plus fortes sont ceux qui transportent les flux de voix et de vidéo. Les flux TCP sont dans la majorité des cas des flux qui possèdent une QoS *Best-Effort* dans un scénario de sécurité civile. En outre, comme expliqué dans (Galvez, Ruiz, & Skarmeta, 2012), les mesures de débits effectuées sur des flux TCP ne peuvent être utilisées pour du partage de charge puisque ces flux sont élastiques et tendent à utiliser toutes les ressources disponibles. Par conséquent, les mesures de débits ne sont focalisées que sur les flux UDP bien que les flux TCP soient bien évidemment affectés par l'algorithme.

La métrique à l'intérieur du MANET

La métrique à l'intérieur du MANET va permettre de définir les frontières de référence des domaines des différentes passerelles. De nombreuses métriques ont été définies dans le cadre des protocoles de routage MANET. Le mécanisme que nous proposons se veut le plus indépendant possible de cette métrique. Les principales contraintes sur celle-ci sont la propriété d'isotonie, pour assurer la cohérence du domaine, et la prise en compte de la longueur du chemin dans cette métrique. La métrique la plus simple reste donc le nombre de sauts. Cependant, d'autres métriques plus complexes peuvent être utilisées telles que ETX, ETT.

Pour le lien montant, la métrique interne est déjà échangée dans le cadre du protocole de routage interne au MANET. Par conséquent, les passerelles ont connaissance de ces métriques, puisqu'elles font partie intégrante du MANET. Pour le lien descendant, ce n'est pas le cas et le transfert de ces informations est donc nécessaire. Les passerelles doivent transmettre leur table de routage MANET au quartier général. Ce message est nécessaire pour assurer un routage cohérent lorsqu'un MANET se sépare (Paragraphe 6.3). Nous lui ajoutons la métrique correspondant au meilleur chemin entre la passerelle et le nœud destination. Nous nommons ce message GW-ADV-EXT (*GateWay Advertisement Extension*). Pour éviter de trop nombreuses transmissions en cas de topologies très dynamiques, nous définissons un temps minimum entre deux envois de GW-ADV-EXT. La Figure 86, ci-dessous, récapitule différents points de la mesure.

D'autres méthodes pour réduire l'*overhead* peuvent être utilisées telles que l'utilisation d'un codage spécifique des adresses IP pour éviter la répétition du sous-réseau commun aux nœuds MANET ou limiter les informations à transmettre aux modifications de la table de routage de la passerelle.

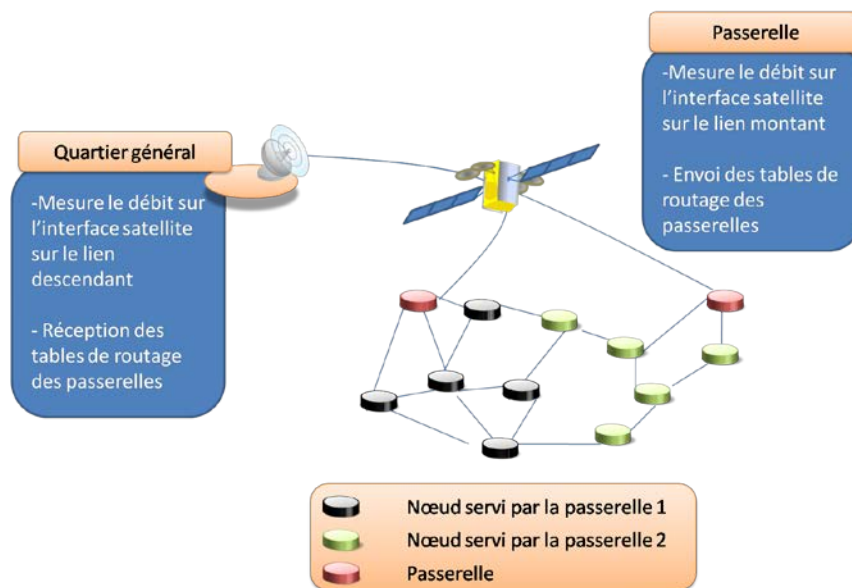


Figure 86 Récapitulatif de la phase de mesure

7.1.4 La phase de décision

Notre proposition se fonde sur deux événements : la surcharge et la sous-charge. Cette méthode a l'intérêt de ne modifier le choix de la passerelle que lorsque cela se révèle nécessaire. Ces événements sont déclenchés, lorsque la charge de la passerelle excède le seuil de surcharge ($\text{Seuil}_{\text{surcharge}}$) ou à l'instant où celle-ci devient inférieure au seuil de sous-charge ($\text{Seuil}_{\text{sous-charge}}$).

Ces événements n'occasionnent pas instantanément la phase d'exécution. En effet, cette phase contient toutes les vérifications nécessaires au bon fonctionnement du mécanisme, dont le plus important est l'évitement du phénomène d'oscillation de passerelle. Ce phénomène, présenté dans le chapitre précédent, est néfaste pour les flux le subissant. Notre volonté est de proposer un mécanisme qui détecte et annule son effet.

7.1.5 Mécanisme contre l'oscillation de passerelle

L'évitement de l'oscillation de passerelle se décompose en trois étapes. Ce sont la détection du phénomène d'oscillation, sa prise en charge par le mécanisme, puis la reprise de l'algorithme lorsque le phénomène n'est plus d'actualité. Elles sont décrites dans la Figure 87.

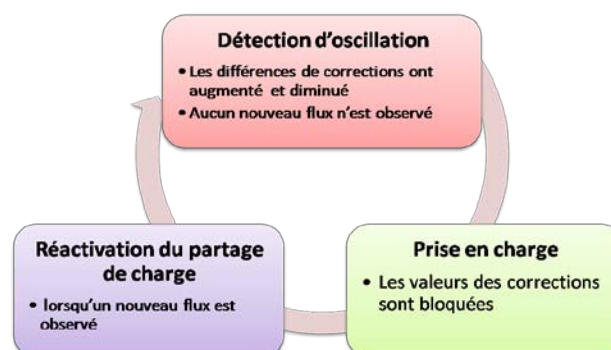


Figure 87 Mécanisme contre l'oscillation

Détection

Une solution pour détecter l'oscillation de passerelle se fonde sur le changement répétitif d'état d'une passerelle. Cela se traduira par le fait qu'une passerelle alternera à chaque étape entre les différents états représentant sa charge. Cela signifie que des flux de certains nœuds sont ballotés entre deux passerelles. Cependant, cette méthode nécessite plusieurs itérations pour détecter l'oscillation ; cela retarde sa prise en charge. Si l'on réduit ce nombre d'itérations pour déclencher plus promptement le mécanisme, il peut potentiellement détecter de manière erronée un phénomène d'oscillation. En effet, une oscillation peu répétitive peut être la simple conséquence de l'apparition de nouveaux flux.

Par conséquent, la solution choisie va prendre en compte, l'apparition ou la disparition de nouveaux flux, ainsi que le changement d'état de charge des passerelles. Cette solution nous permet de réduire le nombre d'itérations à un minimum de trois. En effet, il suffit de détecter une augmentation puis une diminution des différences de corrections, pour en déduire l'occurrence d'un phénomène d'oscillation (Figure 88). Pendant ce temps, ces passerelles ne décèlent aucune modification du trafic et donc seule l'oscillation de passerelle peut être responsable des changements d'état.

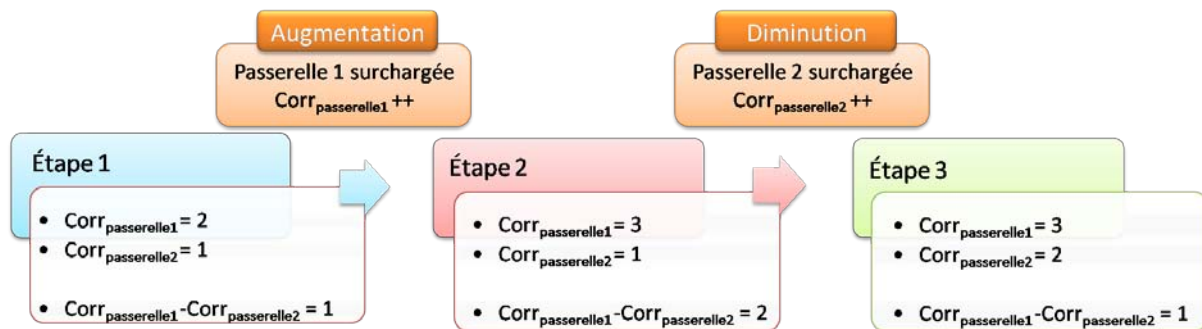


Figure 88 Exemple de détection du phénomène d'oscillation

Pour déterminer si de nouveaux flux apparaissent et disparaissent, les passerelles vont conserver des traces des flux lors des mesures. Ainsi, si un nœud débute un nouveau flux, il est enregistré lors de son passage dans la passerelle. Nous concluons à la disparition lorsqu'aucun nouveau paquet correspondant au flux n'est plus observé par la passerelle pendant une durée prédéfinie ($T_{FluxDisparition}$). Lors de la première itération, la passerelle sort de l'état de surcharge ainsi certains flux sont transférés vers une autre passerelle. L'autre passerelle va détecter les flux transférés en les considérant comme de nouveaux flux. À la nouvelle itération, les deux passerelles ont déjà reconnu tous les flux qui subissent l'oscillation.

L'utilisation d'une connaissance par flux dans le mécanisme d'oscillation peut poser des problèmes de passage à l'échelle. Cependant, le nombre de flux est assez limité dans notre scénario, ce qui réduit ce potentiel problème.

Nous considérons que tout nouveau flux dans le MANET peut modifier la répartition de la charge. Dès qu'une passerelle met à jour un changement de trafic, il en informe tous les nœuds et, en particulier, toutes les autres passerelles par l'intermédiaire du message HNA. Un champ informe de la détection d'un nouveau flux (Figure 89).

Prise en considération de l'oscillation

La passerelle qui détecte en premier le phénomène d'oscillation va informer les autres passerelles de sa détection. De la même manière que pour diffuser un changement de trafic, l'information est transmise par l'intermédiaire du message HNA (Figure 89). Dès qu'une passerelle détient l'information de la présence d'une oscillation, l'exécution du partage de charge est stabilisée jusqu'à ce que l'on détecte un changement de trafic. La stabilisation des valeurs de correction peut suivre deux politiques. La première favorise la stabilisation avec un domaine plus étendu pour la passerelle qui a la capacité la plus importante. La seconde permet de minimiser le coût moyen des chemins en privilégiant la plus faible distance.

La réactivation du partage de charge

Le partage de charge est réactivé lorsque l'une des passerelles observe un nouveau flux ou reconnaît la perte d'un autre. La diffusion de la découverte d'un nouveau flux est par la suite effectuée grâce au message HNA (Figure 89).

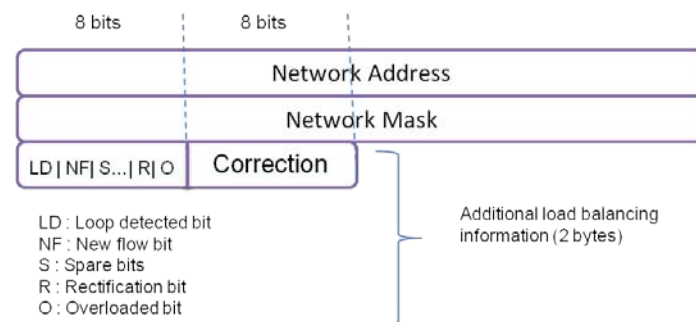


Figure 89 Exemple de message HNA

7.1.5.1 Limitation de l'étendue de la zone progressive

Limitation due au coût entre deux passerelles

Nous introduisons une limitation à l'augmentation de la couverture du domaine, servi par une passerelle, en fonction du coût du chemin à l'intérieur du MANET. Cette limitation permet d'éviter qu'une passerelle ne récupère tout le trafic de sa concurrente. Par conséquent, la différence de correction ne doit pas être supérieure au coût du chemin entre les deux passerelles. Ceci aurait pour conséquence de virtuellement avoir une passerelle dont le coût du chemin pour atteindre les réseaux extérieurs serait inférieur en passant par une autre passerelle que par elle-même.

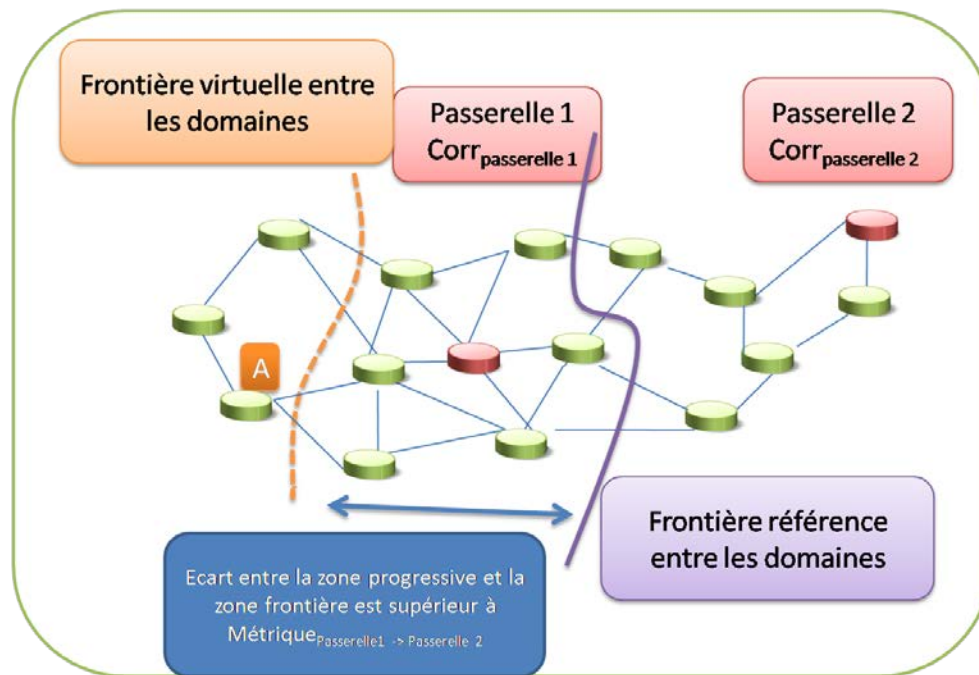


Figure 90 Différence de correction excessive

Dans la Figure 90, nous pouvons observer une frontière virtuelle qui dépasse la passerelle concurrente. Dans ces conditions, tout le domaine de la passerelle 1 est englobé par le domaine de la passerelle 2. Si l'implantation de la hiérarchie des protocoles de routage privilégie le protocole de routage OLSR, la passerelle va transmettre à la passerelle 2, les paquets qu'elle reçoit qui ont pour destination des réseaux extérieurs. Dans le cas contraire, les paquets sont transmis sur son interface satellite. Dans ces deux cas, une augmentation de la correction n'aura plus d'effet puisque le domaine est déjà à sa taille maximale.

Limitation due à un coût maximal

Nous ajoutons une limitation dans le but d'éviter l'obtention de chemin ayant des coûts trop importants. Ceci est mis en place pour restreindre le champ d'action du partage de charge à une zone limitée. En particulier, elle permet de limiter l'augmentation de la congestion dans le MANET et la baisse des performances due à des chemins trop longs. Nous définissons donc un coût maximal ($C_{\max}$), pour lequel nous considérons que le gain obtenu par le partage de charge ne compensera plus les pertes induites par l'augmentation du coût. Cette valeur est donc définie de manière statique.

7.1.5.2 Blocage du mécanisme

Lorsque toutes les passerelles sont surchargées

La phase de décision bloque le mécanisme de partage de charge lorsque toutes les passerelles sont surchargées. En effet, il nous est impossible d'améliorer le partage de trafic dans ces conditions. Le premier effet néfaste est une certaine instabilité, puisque les deux passerelles vont essayer d'augmenter l'étendue de leur domaine conjointement, la différence restera inchangée. L'instabilité est en fait provoquée parce que les mises à jour des corrections ne sont pas synchronisées. En effet, les nœuds ne vont pas recevoir les messages HNA au même instant. Pour y remédier, les passerelles ajoutent un champ, grâce auquel, elles informent les autres passerelles qu'elles sont surchargées.

Ainsi, lorsque toutes le sont, l'incrémentation de la valeur de la correction est abandonnée jusqu'à un nouveau changement d'état de l'une d'elles.

Lorsque le nombre de nœuds source est insuffisant

La phase de décision bloque le mécanisme de partage de charge, lorsque le nombre de nœuds sources devient insuffisant pour atteindre des performances minimales. En effet, comme la granularité est le pas d'incrémentation de la taille du domaine, plusieurs nœuds à la limite de celui-ci vont potentiellement changer de passerelles. Si ce faible nombre de nœuds est à l'origine d'une grande part du nombre de flux, la modification de la passerelle n'a que peu d'impact sur la bonne répartition du trafic et cela engendre des phénomènes d'oscillations. En outre, l'apparition de nouveaux nœuds sources va complètement modifier la répartition du trafic ce qui peut provoquer une phase assez longue avant que le mécanisme de partage de charge ne trouve son point d'équilibre.

7.1.6 La sous-charge

Contrairement à la surcharge, la sous-charge n'induit pas de pertes de paquets et peut donc être traitée avec plus de parcimonie. Le principe du mécanisme de partage de charge établit que la valeur du pas doit être déduite de la correction. Cependant, il est inutile d'augmenter le coût des chemins s'il n'y a pas un besoin de la part des passerelles voisines. Par conséquent, les passerelles sous-chargées ne décrémentent leur correction qu'à la condition que cette action diminue le coût des chemins. Cela survient lorsque la passerelle qui a la valeur maximale des corrections est surchargée.

7.2 Évaluation par simulation

Dans le but d'évaluer le mécanisme de partage de charge, nous avons réalisé des simulations grâce au logiciel Network Simulator 3 (NS3) (Network Simulator 3). Dans un premier temps, nous décrivons le modèle et les paramètres de simulations qui sont utilisés par défaut. Nous analysons ensuite les résultats.

7.2.1 Modèle par défaut

7.2.1.1 Le segment MANET

Le réseau MANET est composé de 49 nœuds. Leur position initiale forme une grille (7x7) dont les nœuds sont espacés de 40m. L'interface WiFi est représentée par le modèle WiFi YANS (Lacage & Henderson, October 2006) développé dans NS3. Nous utilisons le standard 802.11g avec un débit théorique de 24 Mb/s en mode ad-hoc. Le modèle du canal repose sur un modèle de propagation en log-distance avec un exposant de valeur 2,25 qui représente une utilisation en extérieur. En outre, nous utilisons le modèle d'erreur NIST (Pei & Henderson).

Le protocole de routage OLSR a été modifié pour prendre en compte le partage de charge lors de la sélection de passerelle. En particulier, nous avons ajouté aux messages HNA des champs utilisés pour transmettre les valeurs de corrections, ainsi que ceux nécessaires à la phase de décision.

7.2.1.2 Le segment satellite

Deux passerelles sont simulées et sont initialement choisies aux extrémités de la grille ; elles sont diagonalement opposées. Les liens satellite sont définis à l'aide d'une double liaison point-à-point avec un délai de transmission de 300ms. Les capacités des satellites sont asymétriques. Les liens satellite possèdent, respectivement, une liaison descendante de 1Mb/s et de 512 kb/s et une liaison descendante de 256 kb/s pour les deux liens.

7.2.2 Les scénarios de trafic

Les simulations tendent à représenter les applications réellement utilisées lors d'intervention par la sécurité civile. Pour ce faire, quatre types de trafics sont représentés : les services de voix, de vidéoconférence, des transferts de fichiers et des rapports de position. Le dernier type de trafic simule des transmissions de la position des intervenants sur le théâtre de l'opération jusqu'au quartier général. Par conséquent, tous les nœuds transmettent périodiquement ce message. En ce qui concerne les services de voix, nous représentons les trafics de voix par 16 nœuds qui utilisent ce service. Ils sont aléatoirement choisis et débutent la transmission aléatoirement dans une plage de 600 ms. Les nœuds sources de flux de vidéoconférence sont au nombre de deux et sont choisis de la même manière que les nœuds utilisant le service de voix. Des informations supplémentaires sur les paramètres des applications sont rassemblées dans le tableau ci-après.

	Voix	Vidéoconférence	Transfert de fichier
Nombre de nœuds-source	16	2	5
Durée de l'application	60s	180s	(fichier = 1MB)
Périodicité de l'application	600s	600s	600s
Débits	64 kb/s	256kb/s	-
Taille des paquets	160o	160o	1000o
Bidirectionnel	Oui	Oui	Non
Couche transport	UDP	UDP	TCP New Reno

Tableau 22 Paramètres de simulations des applications

Nous nous intéressons à l'amélioration du taux d'erreurs, obtenu grâce à l'algorithme de partage de charge qui a été décrit dans ce chapitre. Nous comparons ainsi les taux de perte entre simulations où les scénarios de mobilité et de trafic sont les mêmes. Seul l'activation et la désactivation du partage de charge change.

7.2.3 Première phase de simulation

Ces premiers résultats vont permettre d'observer l'efficacité du partage de charge. Dans ces simulations, nous ajoutons des variantes, en utilisant des nœuds fixes et des nœuds mobiles, et la présence ou non de trafic entre nœuds du MANET. Pour obtenir des données statistiques, en fonction de différentes distributions des applications, 70 itérations de la simulation sont effectuées. Les paramètres spécifiques au mécanisme de partage de charge sont définis dans le tableau ci-après.

Paramètres	Valeurs
Périodicité des messages Hello	1s
Périodicité des messages HNA	1s
Périodicité des messages TC	3s
Seuil de surcharge	0.9
Seuil de sous-charge	0
Nombre de nœud source min	1
Nombre de sauts max	4
Périodicité du partage de charge sur le lien descendant	1s

Tableau 23 Paramètres d'OLSR et du partage de charge

Réseau Ad-Hoc fixe sans trafic interne

La Figure 91 représente la fonction de répartition des différentes itérations des simulations en fonction de l'amélioration du taux de pertes entre un scénario avec et sans partage de charge pour les flux de vidéoconférence. La Figure 92 procure les mêmes résultats pour les flux de voix. L'axe des abscisses est la différence entre les taux de perte.

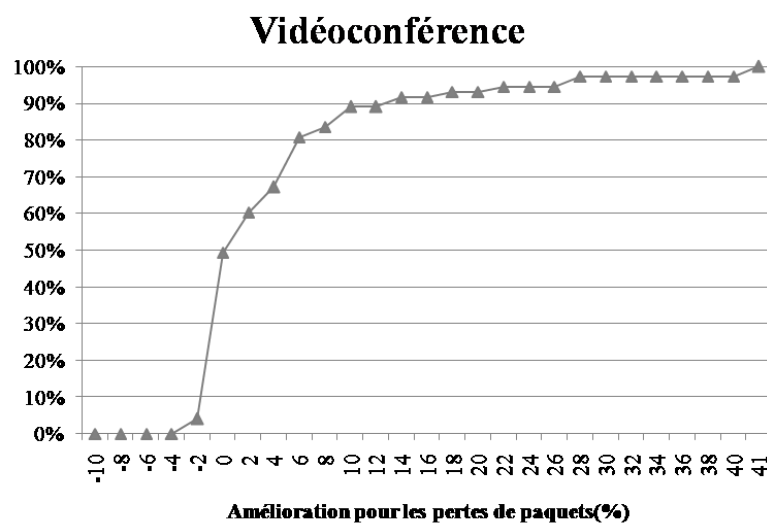


Figure 91 Répartition de l'amélioration des pertes de paquets pour des flux de vidéoconférence

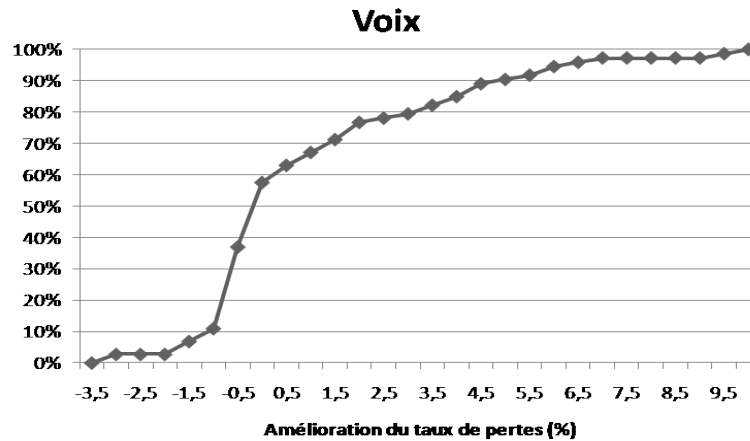


Figure 92 Répartition de l'amélioration des pertes de paquets pour des flux de voix

Le gain apporté par le mécanisme est relativement limité puisque, en moyenne, il est de 4,7%. Cependant, dans le meilleur des cas, nous avons gagné 41%. Ces améliorations de performances importantes surviennent lorsque le trafic n'est pas du tout bien réparti et, en particulier, lorsque c'est la passerelle avec la capacité la plus faible qui est la plus surmenée. Dans de nombreux cas, le besoin de partage de charge n'est pas flagrant. Il n'y a pas de pertes significatives dans les passerelles et, dans ces cas, le partage de charge ne s'impose pas. En outre, pour certains scénarios, il détériore légèrement les performances des flux de vidéoconférence avec un minimum de -2,5%. Une étude de ces cas a permis d'avancer la conclusion suivante. Les pertes sont principalement causées par une mauvaise gestion du mécanisme contre l'oscillation de passerelle. En effet, la reprise du partage de charge après la détection d'une oscillation ne s'effectue que lorsqu'un trafic apparaît ou disparaît. En raison des rapports de position qui sont envoyés périodiquement, les passerelles interprètent cela comme un nouveau trafic et réactivent le partage de charge. Par conséquent, l'oscillation se reproduit causant des pertes.

Pour les services de voix (Figure 92), le gain moyen est de 1.4% et la valeur maximale atteint 10%.. Nous avons remarqué que les pertes maximales pour les trafics de voix surviennent lorsque le gain est maximal pour ceux de vidéoconférence. Ainsi, les phénomènes de pertes sont la conséquence d'un flux de vidéoconférence pour lequel le partage de charge a modifié son choix de passerelle. Cet inconvénient peut être résorbé par l'utilisation de mécanisme de QoS pour privilégier les flux de voix qui sont prioritaires dans un scénario de sécurité civile.

Impact du trafic interne et de la mobilité

Dans ce paragraphe, nous observons l'impact du trafic entre deux nœuds du MANET ainsi que l'effet de la mobilité des nœuds sur l'efficacité du partage de charge.

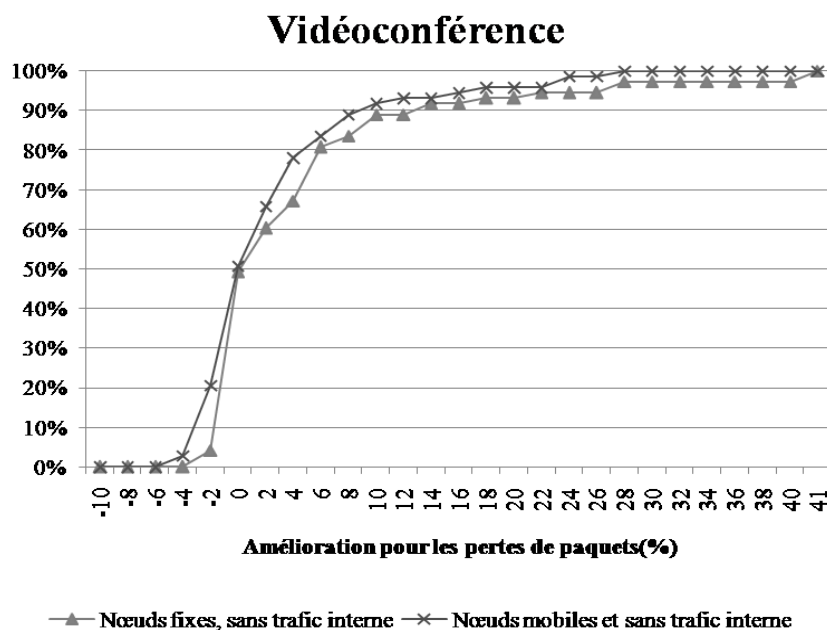


Figure 93 Répartition de l'amélioration des pertes de paquets pour des flux de vidéoconférence en fonction de la mobilité

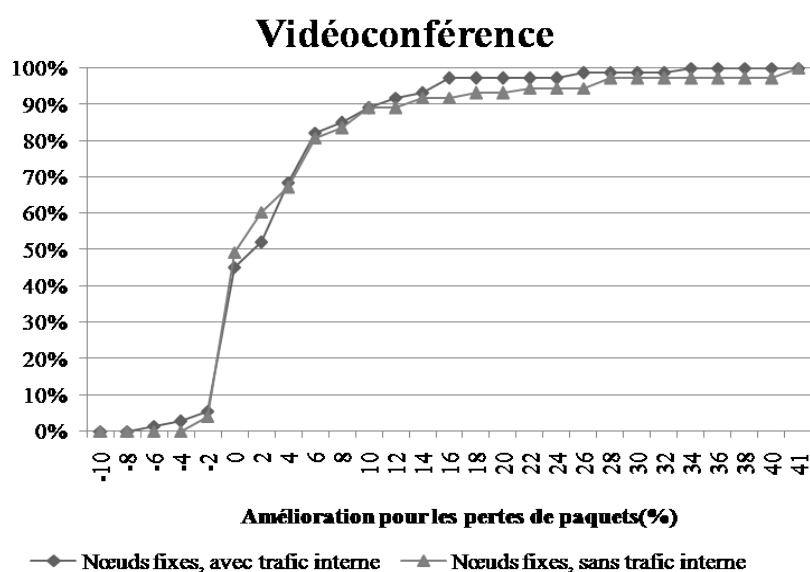


Figure 94 Répartition de l'amélioration des pertes de paquets pour des flux de vidéoconférence en fonction du trafic interne

Voix	Sans trafic interne et nœuds fixes	Sans trafic interne et nœuds mobiles	Avec trafic interne et nœuds fixes	Avec trafic interne et nœuds mobiles
Moyenne	1,4	1	0,6	0,3
Minimum	-3	-4	-1	-1,7
Maximum	10	7,4	3,9	2,8
Intervalle de confiance à 95% sur la moyenne	0,6	0,5	0,2	0,2

Tableau 24 Données statistiques des flux de voix

Vidéoconférence	Sans trafic interne et nœuds fixes	Sans trafic interne et nœuds mobiles	Avec trafic interne et nœuds fixes	Avec trafic interne et nœuds mobiles
Moyenne	4,7	3	3,9	2
Minimum	-2,5	-3,7	-5,5	-10
Maximum	41	28	34	25,5
Intervalle de confiance à 95%	2	1,4	1,5	1,2

Tableau 25 Données statistiques des flux de vidéoconférence

La mobilité et la présence de trafic à l'intérieur du MANET influence un peu le partage de charge. L'amélioration apportée par notre mécanisme est largement influencé par les cas de figure dans lesquels le trafic est très déséquilibré. La mobilité des nœuds va contribuer à réduire les durées pendant lesquelles ces phénomènes sont observés.

Les gains obtenus sont donc sensiblement du même ordre de grandeur, mais légèrement réduits. L'apparition de trafic interne va avoir deux impacts : le premier sera que la prise en compte de la surcharge d'une passerelle ne tient pas compte des congestions internes potentielles dans le MANET. Le partage de charge peut alors, dans certains cas, conduire à une augmentation du taux de perte. Par ailleurs, et comme dans l'analyse précédente, la prise en compte du trafic interne augmente le nombre de flux et réduit alors potentiellement l'apparition des cas très favorables qui conduisent aux améliorations les plus fortes.

7.2.4 Deuxième phase de simulation

Modification

Pour ces simulations, nous avons modifié le mécanisme contre les oscillations. En effet, pour remédier au problème posé par l'apparition et la disparition de flux périodiques tels que les rapports de position, la condition pour retourner dans l'état actif du partage de charge n'est valable que pour les flux qui sont susceptibles de modifier la répartition de la charge. Ainsi, leur débit doit dépasser un seuil pour influencer le partage de charge. Grâce à cela, les rapports de mesures sont ignorés. Le débit théorique du WiFi est maintenant de 6 Mb/s et nous réalisons 50 tirages aléatoires.

Influence du seuil de sous-charge

Nous faisons varier le seuil de sous-charge. Il prend les valeurs 0, 0.05, 0.1 et 0.2.

Voix	0	0.05	0.1	0.2
Moyenne	1.5	0.15	0.18	0.1
Minimum	-4.24	-5.6	-5.6	-6.8
Maximum	6	13	13	12
Intervalle de confiance à 95%	0.8	0.9	0.9	0.9

Tableau 26 Données statistiques des flux de voix en fonction du seuil de sous-charge

Vidéoconférence	0	0.05	0.1	0.2
Moyenne	5.2	2.8	2.8	2.6
Minimum	-4.5	-13.7	-13.7	-13
Maximum	40.7	35.7	35.7	40.2
Intervalle de confiance à 95%	2.4	2.5	2.5	2.6

Tableau 27 Données statistiques des flux de vidéoconférence en fonction du seuil de sous-charge

L'influence du seuil de sous-charge est faible puisque les résultats sont très similaires lorsqu'il prend les valeurs de 0.05, 0.1 et 0.2. Lorsqu'il est égal à zéro, la partie du mécanisme en charge de la sous-charge est alors désactivée. L'efficacité du partage de charge est presque doublée. Nous remarquons que cette augmentation est principalement due à la minimisation du nombre de scénarios où notre mécanisme entraîne une dégradation du taux de pertes. En effet, la dégradation la plus forte est de 4.5% alors qu'elle est de 13.7% lorsque le mécanisme pour la sous-charge est activé. La raison de cette amélioration provient du fait que l'activation de la surcharge augmente les pertes, en raison de l'arrêt des variations des corrections dans des conditions défavorables, lors de la détection de l'oscillation. Par la suite, nous ignorerons donc la sous-charge des passerelles.

Étude du mécanisme contre les oscillations

Dans cette étude, nous faisons varier les opérations réalisées au moment de la stabilisation du partage de charge après la détection d'une oscillation. En effet, dans les précédents résultats nous avançons la passerelle possédant la plus forte capacité. La passerelle qui détectait le phénomène d'oscillation incrémentait la valeur de sa correction si elle possédait la capacité la plus importante et la décrérait dans le cas où elle possédait la plus faible.

Dans cette nouvelle simulation, nous ne favorisons plus la passerelle possédant la capacité la plus forte mais nous favorisons le plus court chemin. Ainsi, lors de la détection de l'oscillation, nous choisissons les valeurs de correction qui minimisent l'écart entre de correction des passerelles parmi celles obtenues précédemment.

Nous proposons une troisième configuration. En plus d'observer les débits sur l'interface satellite, nous prenons en compte la charge sur le lien WiFi de la passerelle satellite. Nous utilisons la même technique de mesure que dans (Pham V. , Larsen, Kure, & Engelstad, November 2010) (Équation (7.4)).

La dernière simulation nous offre la possibilité de comparer nos performances au mécanisme proposé dans (Pham V. , Larsen, Kure, & Engelstad, November 2010). Il ne met en place qu'une valeur d'hystérésis sur la modification de la charge pour éviter les oscillations de passerelle. Sur une zone de partage de charge fixe égale à 2 sauts, il privilégie le chemin menant à la passerelle la moins chargée. Cette métrique est calculée grâce au taux d'occupation des interfaces WiFi avec une prise en compte des mesures précédentes (Équation (7.4)). Le coût total du chemin est égal à la charge maximale entre les nœuds (Équation 131(7.5)(7.5)). Nous ajoutons au mécanisme décrit dans (Pham V. , Larsen, Kure, & Engelstad, November 2010), la prise en considération dans le chemin de la charge de l'interface satellite. Les paramètres utiles sont décrits dans le Tableau 28.

$$(7.4) \quad Charge_{noeud\ i+1} = \alpha Charge_{noeud\ i} + \frac{(1-\alpha) T_{occupé}}{T_{total}}$$

$$(7.5) \quad Charge_{chemin} = \max(Charge_{noeud})$$

Paramètre	Valeur
α (coefficient de mémorisation)	0.05
Hystérésis	3
Périodicité	1

Tableau 28 Paramètres de l'algorithme de (Pham V. , Larsen, Kure, & Engelstad, November 2010)

Voix	Capacité la plus forte	Distance la plus faible	Distance la plus faible et prise en compte de la charge du wifi	Mécanisme de (Pham V. , Larsen, Kure, & Engelstad, November 2010)
Moyenne	1.5	1,7	1,8	1,5
Minimum	-4.24	-2,1	-2,8	-2,2
Maximum	6	15,4	15,6	9,3
Intervalle de confiance à 95%	0.8	0,9	0,9	0,7

Tableau 29 Données statistiques des flux de voix avec l'influence de l'oscillation

Vidéoconférence	Capacité la plus forte	Distance la plus faible	Distance la plus faible et prise en compte de la charge du wifi	Mécanisme de (Pham V. , Larsen, Kure, & Engelstad, November 2010)
Moyenne	5.2	6.0	5,9	5,7
Minimum	-4.5	-2,75	-3,25	-11,8
Maximum	40.7	39,3	41,5	26
Intervalle de confiance à 95%	2.4	2,8	2,8	2,2

Tableau 30 Données statistiques des flux de vidéoconférence avec l'influence de l'oscillation

Durée d'un Flux TCP (s)	Capacité la plus forte	Distance la plus faible	Distance la plus faible et prise en compte de la charge du wifi	Mécanisme de (Pham V. , Larsen, Kure, & Engelstad, November 2010)
Moyenne	0	2,0	4,9	-5,6
Minimum	-33	-22,7	-22,3	-53,3
Maximum	42	43,9	64,6	61,7
Intervalle de confiance à 95%	3.8	4,5	4,9	6,6

Tableau 31 Données statistiques sur la durée des flux de TCP avec l'influence de l'oscillation

Nous notons une amélioration des résultats, lorsque nous utilisons le mécanisme contre les oscillations fondé sur la distance la plus faible vis-à-vis de celui fondé sur la plus forte capacité. Cela s'explique par l'incapacité des passerelles à connaître l'état des corrections qui implique le moins de pertes de paquets. Si l'on favorise une passerelle et que c'est celle-ci qui provoque le plus de pertes, cette mésaventure durera plus longtemps, jusqu'à un nouveau changement de trafic. Dans ces cas-là, le système contre l'oscillation est néfaste. La réduction au maximum l'écart de correction permet de diminuer les congestions à l'intérieur du MANET en réduisant la distance des chemins. La prise en compte de la charge sur le WiFi n'influence que très légèrement les résultats. Le mécanisme proposé dans (Pham V. , Larsen, Kure, & Engelstad, November 2010) que nous avons amélioré obtient, tout de même, des résultats du même ordre de grandeur que notre proposition bien qu'ils soient

légèrement plus faible. En examinant de plus près, les différences entre les deux mécanismes, nous avons remarqué que les performances maximales n'étaient pas atteintes pour les mêmes répartitions de trafic. Le phénomène d'oscillation est très marqué dans la simulation de (Pham V. , Larsen, Kure, & Engelstad, November 2010), mais les effets néfastes de cette oscillation sont amoindries par le fait que les deux liens satellite induisent exactement le même délai. *A contrario*, les flux TCP sont beaucoup plus affectés par cette oscillation. En effet, l'algorithme de (Pham V. , Larsen, Kure, & Engelstad, November 2010) augmente, en moyenne, la durée des flux TCP de 5s alors que pour les deux versions de l'algorithme que nous proposons améliore ce résultat.

7.2.1 Conclusion

Les résultats du partage de charge offre des résultats intéressants et permettent d'obtenir un gain moyen d'à peu près 6%. Il n'offre cependant pas des résultats suffisants pour être réellement bénéfiques pour les applications du MANET. En effet, l'intérêt est très important au vue des gains maximaux obtenus lors de scénarios spécifiques (jusqu'à plus de 45%). Cependant, les pertes lors de certains scénarios sont trop importantes pour faire de cette technique une réussite. Ces pertes sont principalement dues à la difficulté de gestion du phénomène d'oscillation de passerelle. La complexité pour anticiper tous les cas néfastes est trop importante.

En outre, certains paramètres sont favorables à l'utilisation d'un mécanisme de partage de charge, comme l'asymétrie des liens satellite et l'égalité des délais des terminaux satellite. Nous pouvons ajouter la minimisation de la présence de flux entre nœuds MANET qui peut être beaucoup plus importante lors de déploiement d'un réseau MANET sur un théâtre d'intervention. Si l'on considère conjointement ses nouveaux paramètres néfastes pour le partage de charge et les résultats mitigés que nous avons obtenus, nous pouvons conclure que les gains qu'apportent le partage de charge ne sont pas assez intéressants pour justifier l'implantation de ce mécanisme. Il peut tout de même être utile dans le cadre de passerelles à très faible débit et lorsque par exemple seul des BGAN sont utilisés.

Conclusion et ouverture

Conclusion

Les systèmes satellite de diffusion de la télévision constituent le principal marché des télécommunications par satellite et les systèmes fixes et mobiles sont cantonnés à des marchés de niche. Dans le but de diversifier les offres de ces derniers, nous avons proposé dans cette thèse d'étendre les perspectives de ces réseaux, en étudiant l'hybridation de ceux-ci avec des réseaux terrestres. Cette vision nous permet de nous affranchir d'un aspect très limitant des systèmes satellite qu'est la concurrence avec le monde terrestre. Pour cela, nous avons choisi deux familles de réseaux hybrides mobiles qui nous paraissaient prometteurs. Nous avons sélectionné les réseaux mobiles avec un backhaul satellite comme étant la première famille de réseaux hybrides que nous souhaitons promouvoir. En effet, le déploiement de stations de base GSM pour la téléphonie est un marché porteur, en particulier dans les pays en voie de développement. Avec le développement du standard de quatrième génération, LTE, de nouvelles contraintes surviennent. En effet, le besoin de déployer plusieurs terminaux satellite en backhaul devient primordial pour raccorder plusieurs stations de base LTE (eNB) d'un même site en raison des capacités très importantes que requièrent ces eNB. Cette nouvelle contrainte sur l'architecture du réseau implique de nouveaux problèmes, dont le principal se situe dans les procédures de gestion de la mobilité. Pour supprimer cette difficulté qui peut être rédhibitoire pour le déploiement de ce type de réseau, nous proposons différentes optimisations, pour prendre en compte les caractéristiques du lien satellite dans la gestion de la mobilité. Ces optimisations ont pour objectif de réduire les échecs de *handovers*, qui seraient la résultante du backhaul satellite, et d'optimiser la gestion des données sur le plan utilisateur, pour réduire l'impact du *handover* sur les applications de l'utilisateur.

Au vu de ces objectifs nous avons défini deux familles d'optimisations possibles. La première se concentre sur la phase en amont du *handover*. Nous avons modifié les paramètres de cette phase de préparation avec l'intention de prendre en compte le délai induit par le satellite qui la retarde ; ceci nous a permis de minimiser le nombre d'échecs du *handover*. En outre, nous prôtons l'utilisation d'une phase de préparation à double décision, qui permet de ne pas conclure tout déclenchement de la phase de préparation par l'exécution du *handover*. Cette solution offre donc la possibilité de prendre en compte un changement de condition sur l'interface radio, durant la longue période induite par la phase de préparation sur le backhaul satellite. Pour être réellement efficace, elle doit être combinée avec la préparation de multiples eNB, qui sont potentiellement intéressants pour élargir le champ des eNB cibles. La préparation multiple permet donc d'augmenter les chances de trouver un eNB qui regroupe les qualités qui réduisent les risques d'échecs.

La deuxième phase d'optimisation s'attarde sur la gestion des paquets du plan utilisateur pendant l'interruption du service lors du changement de station de base. Nous avons proposé une optimisation pour chaque scénario qui nous a semblé le plus souffrir du *handover*. Le premier scénario correspond au *handover* depuis un eNB avec un backhaul satellite, vers un eNB avec un backhaul terrestre, alors que le deuxième optimise la procédure du *handover* entre deux eNB avec un backhaul satellite. Ces deux modifications reposent sur l'envoi en *multicast* des données de l'utilisateur vers l'eNB source et cible, dans le but de minimiser les pertes de paquets pendant le

handover. La différence entre les deux mécanismes se situe dans la gestion du transfert *bi-cast* avec, dans le premier cas, la prise en compte des pertes de paquets durant l'exécution du *handover*, pendant lequel l'UE n'est plus attaché à un eNB. Des émulations ont été effectuées pour observer l'impact réel de nos mécanismes sur le comportement de TCP. Elles mettent en valeur leur réussite. Les bénéfices dans le cadre du *handover* inter-satellite-terrestre sont pour l'instant insuffisants pour justifier les modifications que nous proposons. *A contrario*, ces émulations ont montré que le mécanisme prévu dans le cas du *handover* intra-satellite est bénéfique et nécessaire.

La deuxième famille de réseaux hybrides, que nous avons choisie d'étudier, se compose d'un réseau MANET et de passerelles satellite qui permettent une interconnexion vers des réseaux extérieurs. Ce type de réseaux hybrides regroupe des qualités très intéressantes dans le cadre de communications d'urgence. En effet, la combinaison du segment satellite et du segment terrestre permet le déploiement rapide d'un réseau sur le théâtre d'intervention de manière indépendante vis-à-vis des installations terrestres. Dans ce contexte, nous avons étudié l'une des procédures cruciales affectées par l'hybridation satellite : la sélection de la passerelle satellite dans le MANET. En effet, plusieurs terminaux satellite sont déployés. Chaque nœud du réseau doit donc choisir la passerelle satellite qu'il est préférable d'utiliser, en fonction de paramètres qui ont pour origine soit le MANET soit le segment satellite. Nous avons privilégié la charge sur ce dernier. Nous avons analysé les inconvénients des solutions existantes, pour en déduire que la plupart néglige un point qui nous semble important : l'oscillation de passerelle. Dans le chapitre, nous avons donc développé un mécanisme prenant en compte ce phénomène. Cette thèse a permis de mettre en évidence que les propositions de la littérature ignorent le plus souvent ce risque d'oscillation de la sélection de passerelle et que leurs simulations ne permettent pas d'avoir une vision globale du comportement de leur solution.

Le réseau hybride LTE avec un backhaul satellite nous semble le plus prometteur pour une commercialisation prochaine. En effet, ce standard connaît un engouement impressionnant de la part de nombreux acteurs du monde des télécommunications que ce soit dans le monde la téléphonie commerciale ou des communications d'urgence. Nous espérons avoir affermi la position technique de cette solution, en proposant des optimisations des mécanismes de mobilité qui pourraient être un frein au développement de ce marché et en montrant leur bon comportement.

Ouverture

Cette thèse n'est bien sûr qu'une pierre à l'édifice, de nombreuses améliorations sont encore nécessaires au développement des scénarios étudiés dans cette thèse.

Les réseaux hybrides MANET satellite sont très complexes et par conséquent peuvent être sujet à de nombreuses optimisations. Dans le domaine de la sélection de passerelle, le mécanisme que nous proposons ne permet pas de totalement appréhender l'un des problèmes qui nous semble le plus néfaste, l'oscillation de passerelle :

- L'utilisation d'une métrique qui offre une granularité plus fine que celle du nombre de sauts pourrait réduire certains inconvénients de ce phénomène, bien qu'elle n'en supprimerait pas son apparition. En effet, une métrique qui permet un pas de l'algorithme plus ténu ne ferait potentiellement subir l'oscillation que sur un nombre de nœuds plus restreints.
- L'oscillation de passerelle est très difficile à prendre en compte sans une complète connaissance de la topologie du réseau MANET ainsi que de l'état du trafic dans le réseau. Ceci nous amène à penser qu'en l'absence d'une gestion centralisée de ce mécanisme, le phénomène d'oscillation reste difficilement contrôlable.

Les réseaux mobiles LTE avec un backhaul satellite font l'objet d'une attention particulière pour des communications d'urgence. De nouvelles problématiques voient le jour dans ce contexte. Dans le but de fournir une réelle application des optimisations que nous avons proposées, de nouvelles études sont nécessaires que ce soit dans le domaine commercial ou dans celui des communications d'urgence :

- Les optimisations de la phase de préparation nécessitent d'être validées par des simulations. En particulier, il est important d'étudier la consommation des ressources de ces mécanismes. En effet, le nombre de préparations augmente significativement en raison des optimisations que nous avons réalisées. Bien que le nombre d'utilisateurs soit plus restreint que sur le segment terrestre, une attention particulière doit être apportée à la gestion des contextes utilisateurs pendant la phase de préparation.
- Nos mécanismes de *Forward* doivent, de la même manière, subir des tests plus avancés, en particulier dans le cas du *handover* intra-satellite. Par exemple, notre sélection de la méthode de *Forward* permet d'observer le comportement des flux lors du *handover*, avec une considération plus importante pour le lien satellite que pour l'interface radio. En effet, l'impact du délai satellite est le plus significatif pour ce mécanisme. La prise en compte de l'interface radio permettra de terminer la validation de l'optimisation que l'on propose, en réalisant des études plus poussées et en y intégrant tous les éléments de LTE.
- L'une des contraintes les plus importantes pour l'optimisation de ces mécanismes est l'impact de cette solution sur le standard ainsi que sur les équipements du monde terrestre. En effet, toute modification entraîne un coût très important. Nous avons donc favorisé, tout au long de cette thèse, les solutions ayant un impact nul ou très faible sur le standard, autorisant uniquement, par exemple, la modification transparente de certains messages par l'intermédiaire des équipements satellite. Cette notion se doit d'être développée de manière plus systématique dans le but d'être le plus compétitif et crédible possible.

- L'amélioration de la phase de préparation et des mécanismes de *Forward* est nécessaire, en raison de l'implication du satellite dans ces procédures ainsi que par l'absence de lien direct (X2) entre eNB. Pour compléter notre panel d'optimisation, il serait intéressant d'étudier l'établissement d'un lien X2 entre les eNB en utilisant les ressources de l'interface radio LTE. Cette solution semble possible lorsque les eNB sont à portée les uns des autres, alors que notre solution serait privilégiée dans le cas contraire. La consommation des ressources de l'interface radio pour les opérations du *handover* n'aurait qu'un faible impact sur les utilisateurs, dans la mesure où les ressources sur l'interface radio seraient en excès, car elles dépassent les capacités du backhaul satellite.

L'intérêt d'un réseau hybride LTE satellite est une réalité. Ce type de réseau offre encore de belles perspectives d'optimisations que ce soit dans la gestion de la mobilité ou dans d'autres procédures telles la gestion de la sécurité, la gestion des *bearers* et les applications de sécurité civile comme le push-to-talk.

Bibliographie

- [1] Satellite Industry Association (SIA), *"State of the Satellite Industry Report"*, Mai 2012.
- [2] International Telecommunication Union - Radiocommunication sector (ITU-R), *"M.2134, Requirements related to technical performance for IMT-Advanced radio interface(s)"*, November 2008.
- [3] The 3rd Generation Partnership Project (3GPP), *"TR 36.913 v11.0.0, Requirements for further advancements for Evolved Universal Terrestrial Radio Access (E-UTRA) (Release 11)"*, September 2011.
- [4] Oliveira A. et al., *"Internetworking of satellite and wireless ad hoc networks for emergency and disaster relieve services"*, International Journal of Satellite Communications Policy and Management (IJSCPM), vol. 1, no. 1, pp. 1-14, April 2011.
- [5] Stüber G. L. , *"Principle of Mobile Communication"*, 3rd ed., Springer, 2011.
- [6] Sesia S. , Toufik I. , and Baker M. , *"LTE, The UMTS Long Term Evolution, From Theory to Practice"*, 2nd ed., John Wiley & Sons, 2011.
- [7] The 3rd Generation Partnership Project (3GPP). [Online]. <http://www.3gpp.org>
- [8] The 3rd Generation Partnership Project 2 (3GPP2). [Online]. <http://www.3gpp2.org>
- [9] Ericsson, *"Traffic and Market report, On the pulse of the network society"*, White Paper, 2012. [Online]. <http://www.ericsson.com/res/docs/2012/ericsson-mobility-report-november-2012.pdf>
- [10] IETF MANET Working Group. [Online]. <http://datatracker.ietf.org/wg/manet/charter/>
- [11] F Vallejo A. Y. , *"AmerHis: Triple Play over an OBP-based DVB-RCS Satellite Platform"*, the 23rd AIAA International Communications Satellite, September 2005.
- [12] Holma H. and Toskala A. , *"LTE for UMTS : Evolution to LTE-Advanced"*, 2nd:adp, Ed., John Wiley & Sons, 2011.
- [13] Motion Picture Experts Group. [Online]. <http://mpeg.chiariglione.org/>
- [14] European Telecommunications Standards Institute (ETSI), *"EN 300 421 V1.1.2, Digital Video Broadcasting (DVB);Framing structure, channel coding and modulation for 11/12 GHz satellite services"*, August 1997.
- [15] European Telecommunications Standards Institute (ETSI), *"EN 302 307 V1.2.1, Digital Video Broadcasting (DVB); Second generation framing structure, channel coding and modulation systems for Broadcasting, Interactive Services, News Gathering and other broadband satellite applications (DVB-S2)"*, April 2009.

- [16] Morello A. and Mignone V. , *"DVB-S2: The Second Generation Standard for Satellite Broad-Band Services"*, Proceedings of the IEEE, vol. 94, no. 1, pp. 210-227, January 2006.
- [17] Akyildiz I. F. and Jeong S. H. , *"Satellite ATM networks: a survey"*, IEEE Communications Magazine, vol. 35, no. 7, pp. 30-43, July 1997.
- [18] Vasquez-Castro M. A. , Cardoso A. , and Rinaldo R. , *"Encapsulation and framing efficiency of DVB-S2 satellite systems"*, the IEEE 59th Vehicular Technology Conference (VTC 2004-Spring), vol. 5, pp. 2896- 2900, May 2004.
- [19] Hobaya F. et al., *"Encapsulation Requirements for Return Links and Mesh Systems over Satellite"*, the IEEE 72nd Vehicular Technology Conference Fall (VTC 2010-Fall), pp. 1-5, September 2010.
- [20] iDirect. (2008, Juin) iDirect Evolution DVB-S2/ACM. [Online].
http://www.idirect.net/Company/Resource-Center/Collateral-Library/~media/Files/Evolution%20Campaign/iDirect_Evolution_Tech_Overview.ashx
- [21] Telecommunications Industry Association (TIA);, *"TIA-1008-B, IP over Satellite"*, 2012.
- [22] Vidal O. et al., *"Next generation High Throughput Satellite system"*, pp. 1-7, October 2012.
- [23] Caini C. , Firrincieli R. , and Lacamera D. , *"PEPsal: A Performance Enhancing Proxy for TCP Satellite Connections [Internetworking and Resource Management in Satellite Systems Series]"*, the IEEE Magazine on Aerospace and Electronic Systems, vol. 22, no. 8, pp. B-9-B-16, August 2007.
- [24] iDirect, *"XHO One-way Hub-Optimizations"*, 2010. [Online].
http://www.idirect.net/Company/Resource-Center/Collateral-Library/~media/Files/Spec%20Sheets/TIB_XHO_011010.ashx
- [25] Pratt S. R. , Raines R. A. , Fossa JR. C. E. , and Temple M. A. , *"An Operational and Performance Overview of the IRIDIUM Low Earth Orbit Satellite System"*, IEEE Communications Survey& Tutorials, vol. 2, no. 2, pp. 2-10, Second Quarter 1999.
- [26] Globalstar. [Online]. <http://ca.globalstar.com>
- [27] European Telecommunications Standards Institute (ETSI), *"TS 101 376-1-3 V3.1.1, GEO-Mobile Radio Interface Specifications (Release 3); Third Generation Satellite Packet Radio Service; Part 1: General specifications; Sub-part 3: General System Description GMR-1 3G 41.202"*, July 2009.
- [28] Deslandes V. , *"Analyse et optimisation du partage de spectre dans les systèmes mobiles intégrés satellite et terrestre"*, Institut National Polytechnique de Toulouse, Toulouse, Thèse de doctorat, 2012.

- [29] International Telecommunication Union - Radiocommunication sector (ITU-R), "*Annex 9 to Document 4B/148 (Working document towards a preliminary draft Report ITU-R [M/S.TERM] – Terminology used for networks using both satellite and terrestrial links)*", May 2009.
- [30] Evans B. G. et al., "*Service scenarios and system architecture for satellite UMTS IP based network (SATIN)*", the 20th AIAA International Communications Satellite Systems Conference, May 2002.
- [31] Chini P. , Giambene G. , and Kota S. , "*A survey on mobile satellite systems*", John Wiley & Sons International Journal of Satellite Communications and Networking, vol. 28, no. 1, pp. 29-57, August 2010.
- [32] SafeTrip. <http://www.safetrip.eu/>.
- [33] Scalise S. et al., "*S-MIM: A novel radio interface for efficient messaging services over satellite*", IEEE International Conference on Communications (ICC), pp. 6953-6958, June 2012.
- [34] Picard M. , Oularbi M. R. , Flandin G. , and Houcke S. , "*An adaptive multi-user multi-antenna receiver for satellite-based AIS detection*", the 6th Advanced Satellite Multimedia Systems Conference (ASMS) and 12th Signal Processing for Space Communications Workshop (SPSC), pp. 273 - 280, September 2012.
- [35] Rodriguez C. G. , "*Manet Routing Assisted by Satellites*", Télécom Bretagne, Thèse de doctorat, 2011.
- [36] Anderson T. and Mavrakakis D. , "*Cellular Backhaul over Satellite*", iDirect White Paper, 2009. [Online]. <http://www.idirect.net/Applications/Cellular-Backhaul.aspx>
- [37] Gilat, "*Efficient GSM Backhaul over Satellite*", 2010. [Online]. http://www.gilat.com/dynimages/t_brochures/files/Gilat-Ericsson%20Solution%20Brochure%202010-05.pdf
- [38] Hugues Networks and Lemko Corporation, "*Hughes and Lemko Corporation Demonstrate High-Speed Wireless 4G/LTE Video Calls over Satellite Backhaul*", 2012. [Online]. http://defense.hughes.com/defense/press_releases/hughes-and-lemko-corporation-demonstrate-high-speed-wireless-4glte-video-calls-over-satellite-backhaul
- [39] DUMBO. [Online]. <http://www.interlab.ait.ac.th/dumbo/>
- [40] Luglio M. , Monti C. , Roseti C. , Saitto A. , and Segal M. , "*Interworking between MANET and satellite systems for emergency applications*", John Wiley & Sons International Journal of Satellite Communications and Networking, vol. 25, no. 5, pp. 551–558, September 2007.
- [41] ComTech EF DATA, "*Challenges & Opportunities for 3G Backhaul over Satellite*", White Paper, 2011. [Online]. http://www.google.fr/url?sa=t&rct=j&q=&esrc=s&source=web&cd=1&ved=0CDMQFjAA&url=http%3A%2F%2Fwww.comtechefdata.com%2Ffiles%2Farticles_papers%2FWP-Challenges-Opportunities-for-3G-Backhaul-over-Satellite.pdf&ei=N0x4UfPNGNSChQethYHgCQ&usg=AFQjCNFi_6przurfOD1

- [42] 4G Americas, *"Self-Optimizing Networks: The Benefits of SON in LTE"*, White Paper, June 2011. [Online]. <http://www.4gamericas.org/documents/Self-Optimizing%20Networks-Benefits%20of%20SON%20in%20LTE-July%202011.pdf>
- [43] Hämäläinen S. , Sanneck H. , and Sartori C. , *"LTE Self-Organising Networks (SON): Network Management Automation for Operational Efficiency"*, John Wiley & Sons, 2012.
- [44] Federal Communications Commission (FCC), *"FCC TAKES ACTION TO ADVANCE NATIONWIDE BROADBAND COMMUNICATIONS FOR AMERICA'S FIRST RESPONDERS"*, January 2011.
- [45] The 3rd Generation Partnership Project (3GPP), *"TS 36 300 v11.3.0, Evolved Universal Terrestrial Radio Access (E-UTRA) and Evolved Universal Terrestrial Radio Access Network (E-UTRAN); Overall description; Stage 2 (Release 11)"*, September 2012.
- [46] The 3rd Generation Partnership Project (3GPP), *"TS 22.228 v12.4.0, Service requirements for the Internet Protocol (IP) multimedia core network subsystem (IMS); Stage 1 (Release 12)"*, December 2012.
- [47] The 3rd Generation Partnership Project (3GPP), *"TS 23.228 v11.7.0, IP Multimedia Subsystem (IMS); Stage 2 (Release 11)"*, December 2012.
- [48] The 3rd Generation Partnership Project (3GPP), *"TS 36.323 v11.0.0, Evolved Universal Terrestrial Radio Access (E-UTRA); Packet Data Convergence Protocol (PDCP) specification (Release 11)"*, September 2012.
- [49] The 3rd Generation Partnership Project (3GPP), *"TS 36.322 v11.0.0, Evolved Universal Terrestrial Radio Access (E-UTRA); Radio Link Control (RLC) protocol specification (Release 11)"*, September 2012.
- [50] The 3rd Generation Partnership Project (3GPP), *"TS 36.321 v11.0.0, Evolved Universal Terrestrial Radio Access (E-UTRA); Medium Access Control (MAC) protocol specification (Release 11)"*, September 2012.
- [51] The 3rd Generation Partnership Project (3GPP), *"TS 29.281 v11.4.0, General Packet Radio System (GPRS) Tunnelling Protocol User Plane (GTPv1-U) (Release 11)"*, September 2012.
- [52] The 3rd Generation Partnership Project (3GPP), *"TS 29.274 v11.4.0, 3GPP Evolved Packet System (EPS); Evolved General Packet Radio Service (GPRS) Tunnelling Protocol for Control plane (GTPv2-C); Stage 3 (Release 11)"*, September 2012.
- [53] The 3rd Generation Partnership Project (3GPP), *"TS 36.413 v11.1.0, Evolved Universal Terrestrial Radio Access Network (E-UTRAN); S1 Application Protocol (S1AP) (Release 11)"*, September 2012.
- [54] The 3rd Generation Partnership Project (3GPP), *"TS 24.301 v11.4.0, Non-Access-Stratum (NAS) protocol for Evolved Packet System (EPS); Stage 3 (Release 11)"*, September 2012.
- [55] The 3rd Generation Partnership Project (3GPP), *"TS 36.331 v11.1.0, Evolved Universal Terrestrial Radio Access (E-UTRA); Radio Resource Control (RRC); Protocol specification (Release 11)"*, September 2012.

- [56] Ekstrom H. , *"QoS Control in the 3GPP Evolved Packet System"*, IEEE Communication Magazine, vol. 47, no. 2, pp. 76-83, February 2009.
- [57] Robson J. , *"Guidelines for LTE Backhaul Traffic Estimation"*, NGMN Alliance, White Paper, July 2011.
- [58] The 3rd Generation Partnership Project (3GPP), *"TS 23.203 v11.8.0, Policy and charging control architecture (Release 11)"*, December 2012.
- [59] The 3rd Generation Partnership Project (3GPP), *"TS 36.304 v11.1.0, Evolved Universal Terrestrial Radio Access (E-UTRA); User Equipment (UE) procedures in idle mode (Release 11)"*, September 2012.
- [60] The 3rd Generation Partnership Project (3GPP), *"TS 23 401 v11.4.0, General Packet Radio Service (GPRS) enhancements for Evolved universal Terrestrial Radio Access Network (Release 11)"*, December 2012.
- [61] The 3rd Generation Partnership Project (3GPP), *"TS 36.214 v11.0.0, Physical layer; Measurements (Release 11)"*, September 2012.
- [62] The 3rd Generation Partnership Project (3GPP), *"TS 23.402 v11.4.0, Architecture enhancements for non-3GPP accesses (Release 11)"*, September 2012.
- [63] Nokia, Nokia-Siemens Networks, *"R1-092092, Analysis of RLF Reasons in Manhattan Environment for Rel'8 Mobility"*, 3GPP TSG-RAN Working Group 1 (Radio) meeting #57, San Francisco, May 2009.
- [64] NTT DoCoMo, Inc., *"R2-092433, Evaluation model for Rel-8 mobility performance"*, 3GPP TSG-RAN WG2 Meeting #65bis, Seoul, March 2009.
- [65] Ericsson, *"R1-090913, Considerations on Mobility Enhancements for Release 9"*, 3GPP TSG RAN WG1 Meeting #56, Athens, February 2009.
- [66] European Telecommunications Standards Institute (ETSI), *"TS 101 376-4-13 V3.1.1, GEO-Mobile Radio Interface Specifications (Release 3) ; Third Generation Satellite Packet Radio Service;Part 4: Radio interface protocol specifications; Sub-part 13: Radio Resource Control (RRC) protocol ; lu Mode ; GMR-1 3G 44.118"*, July 2009.
- [67] Efthymiou N. , Hu H. F. , Sheriff R. E. , and Properzi A. , *"Inter-segment handover algorithm for an integrated terrestrial/satellite-UMTS environment"*, The 9th IEEE International Symposium on Personal, Indoor and Mobile Radio Communications, vol. 2, pp. 993-998, Septembre 1998.
- [68] The 3rd Generation Partnership Project (3GPP), *"TS 28.627 v11.0.0, Technical Specification Group Services and System Aspects;Telecommunication Management;Self-Organizing Networks (SON) Policy Network Resource Model (NRM) Integration Reference Point (IRP);Requirements (Release 11)"*, December 2012.
- [69] The 3rd Generation Partnership Project (3GPP), *"TS 32.522 v11.4.0, Technical Specification Group Services and System Aspects;Telecommunication management;Self-Organizing Networks (SON) Policy Network Resource Model (NRM) Integration Reference Point (IRP);Information Service (IS) (Release 11)"*, September 2012.

- [70] Van Den Berg J. L. et al., "*SOCRATES: Self-Optimisation and self-ConfigurATion in wirelESs networks*", COST 2100 (TD(08)422), February 2008.
- [71] Bălan I. M. , Sas B. , Jansen T. , Spaey K. , and Demeester P. , "*An enhanced weighted performance-based handover parameter optimization algorithm for LTE networks*", EURASIP Journal on Wireless Communications and Networking, pp. 98-109, September 2011.
- [72] Jansen T. , Balan I. , Stefanski S. , Moerman I. , and Kurner T. , "*Weighted Performance Based Handover Parameter Optimization in LTE*", the IEEE 73rd Vehicular Technology Conference (VTC-2011-Spring), pp. 1-5, May 2011.
- [73] Catovic A. , Garavaglia A. , and Patil S. S. , "*METHOD AND APPARATUS FOR AUTOMATIC HANDOVER OPTIMIZATION*", Patent 20090323638, December 2009.
- [74] Komine T. , Yamamoto T. , and Konishi S. , "*A proposal of cell selection algorithm for LTE handover optimization*", IEEE Symposium on Computers and Communications (ISCC), pp. 37-42, July 2012.
- [75] Legg P. , Hui G. , and Johansson J. , "*A Simulation Study of LTE Intra-Frequency Handover Performance*", the IEEE 72nd Vehicular Technology Conference Fall (VTC 2010-Fall), pp. 1-5, September 2010.
- [76] Hui G. and Legg P. , "*Soft Metric Assisted Mobility Robustness Optimization in LTE Networks*", International Symposium on Wireless Communication Systems (ISWCS), pp. 1-5, August 2012.
- [77] Nortel, "*R3-071906, Multi-eNB Handover preparation for Radio Link Failure Recovery*", 3GPP TSG-RAN WG3 Meeting #57bis, Sophia Antipolis, October 2007.
- [78] Panasonic, "*R2-062146, Necessity of forward handover*", 3GPP TSG RAN WG2 Meeting #54, Tallinn, August 2006.
- [79] Ericsson, "*R2-062863, On the issue of forward handover in LTE*", 3GPP TSG-RAN WG2 Meeting #55, Seoul, October 2006.
- [80] 3GPP TSG-RAN WG1, "*LS answer to R1-091127 = R2-091997 on mobility evaluation*", 3GPP TSG RAN WG2 Meeting #65bis, Seoul, March 2009.
- [81] NTT DoCoMo, "*R3-060411, In-sequence delivery for lossless/no duplication handover*", 3GPP TSG-RAN WG3 Meeting #51bis, Sophia Antipolis, April 2006.
- [82] Racz A. , Temesvary A. , and Reider N. , "*Handover Performance in 3GPP Long Term Evolution (LTE) Systems*", the 16th IST Mobile and Wireless Communications Summit, pp. 1-5, July 2007.
- [83] Gayraud T. , Alphand O. , Berthou P. , Baudoin C. , and Duquerroy L. , "*Mobility Architectures for DVB-S/RCS Satellite*", the 2nd Ka band Broadband Communications Conference (Ka-Band), September 2006.
- [84] Han M. H. , Han K. S. , and Lee D. J. , "*Fast IP Handover Performance Improvements Using Performance Enhancing Proxys between Satellite Networks and Wireless LAN Networks for High-Speed Trains*", the IEEE 61th Vehicular Technology Conference, 2008 (VTC-Spring 2008), pp. 2341-2344, May 2008.

- [85] Montavont N. and Noel T. , *"Handover management for mobile nodes in IPv6 networks"*, IEEE Communications Magazine, vol. 40, no. 8, pp. 38-43, August 2002.
- [86] Siemens, *"R3-060279, Intra-LTE-Access Mobility Support for UEs in LTE_ACTIVE as discussed at RAN3#50"*, 3GPP TSG RAN WG3 Meeting #51, Denver, February 2006.
- [87] Pacifico D. , Pacifico M. , Fischione C. , Hjalrmasson H. , and Johansson K. H. , *"Improving TCP Performance During the Intra LTE Handover"*, IEEE Global Telecommunications Conference (GLOBECOM), pp. 1-8, December 2009.
- [88] Liu B. , Boukhateem N. , Martins P. , and Bertin P. , *"Multihoming at layer-2 for inter-RAT handover"*, the IEEE 21st International Symposium on Personal Indoor and Mobile Radio Communications (PIMRC), pp. 1173 - 1178 , September 2010.
- [89] NEC, *"LGW-070024, Options for user plane handling in LTE-GERAN handovers"*, 3GPP Workshop LTE GSM, Sophia Antipolis, January 2007.
- [90] Alcatel, *"Comparison of Inter-3GPP RAT handovers, R3-061802"*, 3GPP TSG-RAN WG3 Meeting #54, Riga, November 2006.
- [91] Research In Motion UK Limited, *"R2-093735, Joint PDCP protocols on Uu and Un interfaces to improve type-1 relay handover"*, 3GPP TSG RAN WG2 Meeting #66bis, Los Angeles, July 2009.
- [92] LG Electronics Inc, *"R3-101447, Consideration on Data Forwarding for Relay based Handover Procedure"*, 3GPP TSG-RAN WG3 Meeting #68, Montreal, May 2010.
- [93] Motorola, *"R2-093207, Handovers involving Type-1 Relay Node"*, 3GPP TSG-RAN-WG2 Meeting #66, San Francisco, May 2009.
- [94] NTT DoCoMo, Inc., *"R3-071352, PDCP SN continuation for DL Data Forwarding during handover"*, 3GPP TSG-RAN WG3 Meeting #57, Athens, August 2007.
- [95] Ericsson, *"R3-071500, PDCP SN Handling at Data Forwarding"*, 3GPP TSG-RAN WG3 Meeting #57, Athens, August 2007.
- [96] Ericsson, *"R3-071826, Forwarding and Sequence Number Handling at S1 Handovers"*, 3GPP TSG-RAN WG3 Meeting #57-bis, Sophia Antipolis, October 2007.
- [97] Ericson, *"R3-071747, Agreements and open issues on SN during data forwardin"*, 3GPP TSG RAN WG3 Meeting #57, Athens, August 2007.
- [98] NEC, *"R3-071897, Proposed Data forwarding solution"*, 3GPP TSG RAN WG3 Meeting #57, Athens, August 2007.
- [99] LTE-SIM. [Online]. <http://telematics.poliba.it/index.php/en/lte-sim>
- [100] Piro G. , Grieco L. A. , Boggia G. , Capozzi F. , and Camarda P. , *"Simulating LTE Cellular Systems: an Open Source Framework"*, IEEE Transaction Vehicular Technologies, vol. 60, no. 2, pp. 498-513, February 2011.
- [101] Université Technologique de Vienne. LTE Simulators. [Online]. <http://www.nt.tuwien.ac.at/about-us/staff/josep-colom-ikuno/lte-simulators/>

- [102] Mehlführer C. , Wrulich M. , Ikuno J. C. , Bosanska D. , and Rupp M. , "*Simulating the Long Term Evolution Physical Layer*", the 17th European Signal Processing Conference (EUSIPCO 2009), August 2009.
- [103] Ikuno J. C. , Wrulich M. , and Rupp M. , "*System level simulation of LTE networks*", the IEEE 71st Vehicular Technology Conference (VTC-2010-Spring), pp. 1-5, May 2010.
- [104] LENA Project. [Online]. http://iptechwiki.cttc.es/LTE-EPC_Network_Simulator_%28LENA%29
- [105] Baldo N. , Miozzo M. , Requena M. , and Guerrero J. N. , "*An Open Source Product-Oriented LTE Network Simulator based on ns-3*", the 14th ACM International Conference on Modeling, Analysis and Simulation of Wireless and Mobile Systems (MSWIM 2011), pp. 293-298, July 2011.
- [106] Borman D. , Braden R. , and Jacobson V. , "*TCP Extensions for High Performance*", RFC 1323, May 1982.
- [107] Sarolahti P. and Kuznetsov A. , "*Congestion control in Linux TCP*", the FREENIX Track: 2002 USENIX Annual Technical Conference, pp. 49-62, june 2002.
- [108] MONET Project. [Online]. <http://monet.tekever.com>
- [109] MONET, "*D2.1- MONET concept and use-cases*", European Commission Deliverable, 2010.
- [110] MONET, "*D3.1 MONET Architectural Desgin*", European Commission Deliverable, 2010.
- [111] BGAN-France. [Online]. <http://www.bgan-france.com/>
- [112] Clausen T. and Jacquet P. , "*Optimized Link State Routing Protocol (OLSR)*", RFC 3626, October 2003.
- [113] Kumar R. , Sarje A. K. , and Misra M. , "*Review Strategies and Analysis of Mobile Ad Hoc Network-Internet Integration Solutions*", International Journal of Computer Science Issues (IJCSI), vol. 7, no. 4, 2010.
- [114] Nordström E. , Gunningberg P. , and Tschudin C. , "*Robust and flexible Internet connectivity for mobile ad hoc networks*", Elsevier Ad Hoc Networks, vol. 9, no. 1, pp. 1-15, January 2011.
- [115] Le-Trung Q. , Engelstad P. E. , Skeie T. , and Taherkordi A. , "*Load-balance of intra/inter-MANET traffic over multiple internet gateways*", the 6th ACM International Conference on Advances in Mobile Computing and Multimedia (MoMM'08), pp. 50--57, November 2008.
- [116] Jönsson U. , Alriksson F. , Larsson T. , Johansson P. , and Maguire G. Q. , "*MIPMANET - Mobile IP for Mobile Ad Hoc Networks*", the 1st ACM international symposium on Mobile ad hoc networking & computing (MobiHoc), pp. 75–85, August 2000.
- [117] Broch J. , Maltz D. A. , and Johnson D. B. , "*Supporting hierarchy and heterogenous interfaces in multi-hop wireless ad hoc network*", the 4th International Symposium on Parallel Architectures, Algorithms, and Networks (I-SPAN '99), pp. 370–375, June1999.

- [118] Sun Y. , Belding-Royer E. M. , and Perkins C. E. , *"Internet Connectivity for Ad hoc Mobile Networks"*, Springer, International Journal of Wireless Information Networks, vol. 2, no. 9, pp. 75–88, April 2002.
- [119] Engelstad P. E. , Tønnesen A. , Hafslund A. , and Egeland G. , *"Internet Connectivity for Multi-Homed Proactive Ad Hoc Networks"*, IEEE International Conference on Communications (ICC'2004), vol. 7, pp. 4050–4056, June 2004.
- [120] Benzaid M. , Minet P. , Al Agha K. , Adjih C. , and Allard G. , *"Integration of Mobile-IP and OLSR for a Universal Mobility"*, Wireless Networks, pp. 377–388, July 2004.
- [121] Engelstad P. E. and Egeland G. , *"NAT-based Internet Connectivity for On Demand MANETs"*, Proceedings of the 1st IFIP TC6 Working Conference (WONS 2004), pp. 4050–4056, Janvier January 2004.
- [122] Zhuang L. , Liu Y. , and Liu K. , *"Hybrid Gateway Discovery Mechanism in Mobile Ad-Hoc for Internet Connectivity"*, the 2009 ACM International Conference on Wireless Communications and Mobiles Computing : Connecting the World Wirelessly (ICWCMC), pp. 143-147, june 2009.
- [123] Ghassemian M. , Hofmann P. , Prehofer C. , Friderikos V. , and Aghvami H. , *"Performance Analysis of Internet Gateway Discovery Protocols in Ad Hoc Networks"*, IEEE Wireless Communications and Networking Conference (WCNC 2004), pp. 120-125, March 2004.
- [124] Yang Y. , Wang J. , and Kravets R. , *"Designing Routing Metrics for Mesh Networks"*, IEEE Workshop on Wireless Mesh NEtworks (WiMesh), September 2005.
- [125] De Couto D. S. , Aguayo D. , Bicket J. , and Morris R. , *"A High-Throughput Path Metric for Multi-hop Wireless Routing"*, Journal Wireless Networks - Special issue: Selected papers from ACM MobiCom 2003, vol. 11, no. 4, pp. 419 - 434 , July 2005.
- [126] Hiertz G. R. et al., *"IEEE 802.11s: the WLAN Mesh Standard"*, IEEE Wireless Communications, vol. 17, no. 1, pp. 104-111, February 2010.
- [127] Pearlman M. R. , Haas Z. J. , Sholander P. , and Tabrizi S. S. , *"On the impact of alternate path routing for load balancing in mobile ad hoc networks"*, the 1st Annual Workshop on Mobile and Ad Hoc Networking and Computing (MobiHOC), pp. 3-10, August 2000.
- [128] Nasipuri A. and Das S. , *"On-demand multipath routing for mobile ad hoc networks"*, the 8th International Conference on Computer Communications and Networks (ICCCN), pp. 64–70, October 1999.
- [129] Galvez J. J. , Ruiz P. M. , and Skarmeta A. F. , *"Responsive on-line gateway load-balancing for wireless mesh networks"*, Elsevier Ad Hoc Networks, vol. 10, no. 1, pp. 46-61, January 2012.
- [130] Ancillotti E. , Bruno R. , and Conti M. , *"Load-balanced routing and gateway selection in wireless mesh networks: Design, implementation and experimentation"*, IEEE International Symposium on a World of Wireless Mobile and Multimedia Networks (WoWMoM), pp. 1-7, June 2010.

- [131] Nandiraju D. , Santhanam L. , Nandiraju N. , and Agrawal D. P. , "*Achieving load balancing in wireless mesh networks through multiple gateways*", the IEEE International Conference on Mobile Adhoc and Sensor Systems (MASS), pp. 807-812, October 2006.
- [132] Lee Y. J. and Riley G. F. , "*A workload-based adaptive load-balancing technique for mobile ad hoc networks*", IEEE Wireless Communications and Networking Conference (WCNC), pp. 2001-2007, March 2005.
- [133] Huang C. , Lee H. , and Tseng Y. , "*A Two-Tier Heterogeneous Mobile Ad hoc Network Architecture and Its Load Balance Routing Problem*", the ACM journal Mobile Networks and Applications, vol. 9, no. 4, pp. 379–391, August 2004.
- [134] Prasad R. and Wu H. , "*Gateway deployment optimization in cellular Wi-Fi mesh networks*", Journal of Networks, vol. 1, no. 3, pp. 31–39, July 2006.
- [135] Brännström R. , Ahlund C. , and Zaslavsky A. , "*Maintaining Gateway Connectivity in Multi-hop Ad hoc Networks*", The 30th IEEE Conference on Local Computer Networks (LCN), pp. 682-689, November 2005.
- [136] Ramachandran K. et al., "*On the design and implementatation of infrastructure mesh networks*", the IEEE workshop on Wireless Mesh Networks (WiMesh), September 2005.
- [137] Pham V. , Larsen E. , Kure, and Engelstad P. E. , "*Gateway load balancing in future tactical networks*", Military Communications Conference (MILCOM), pp. 1844-1850, November 2010.
- [138] Park B. N. , Lee W. , Lee C. , and Shin C. K. , "*QoS-Aware Adaptive Internet Gateway Selection in Ad Hoc Wireless Internet Access Networks*", Elsevier Computer Communications, vol. 30, no. 2, pp. 369-384, January 2006.
- [139] Hoffmann F. and Medina D. , "*Optimum Internet Gateway Selection in Ad Hoc Networks*", IEEE International Conference on Communications (ICC '09), pp. 1-5, June 2009.
- [140] Tokito H. , Sasabe M. , Hasegawa G. , and Nakano H. , "*Routing Method for Gateway Load Balancing in Wireless Mesh Networks*", the IEEE 8th International Conference in Networks (ICN), pp. 127-132, March 2009.
- [141] Shin J. , Lee H. , Na J. , Park A. , and Kim S. , "*Load Balancing among Internet Gateways in Ad Hoc Networks*", the IEEE 60th Vehicular Technology Conference (VTC2004-Fall), pp. 2883-2886, September 2004.
- [142] Tada K. and Yamamoto M. , "*Load-Balancing Gateway Selection Method in Multi-hop Wireless Networks*", IEEE Global Telecommunications Conference (GLOBECOM), pp. 1-6, December 2009.
- [143] Xie B. , Yu Y. , Kumar A. , and Agrawal D. P. , "*Load-balancing and inter-domain mobility for wireless mesh networks*", IEEE Global Telecommunications Conference (GLOBECOM), pp. 409-414, December 2006.
- [144] Maurina S. , Riggio R. , Rasheed T. , and Granelli F. , "*On tree-based routing in multi-gateway association based wireless mesh networks*", the IEEE 20th International Symposium on Personal, Indoor and Mobile Radio Communications (PIMRC), pp. 1542-1546, September 2009.

- [145] Pham V. , Larsen E. , Engelstad P. , and Kure Ø. , "*Performance analysis of gateway load balancing in ad hoc networks with random topologies*", the 7th ACM international symposium on Mobility management and wireless access (MobiWAC'09), pp. 66-74, October 2009.
- [146] Iqbal S. M. and Kabir H. , "*Internet gateway discovery and selection scheme in Mobile Ad Hoc Network*", the 14th International Conference on Computer and Information Technology (ICCIT 2011), pp. 44-49, December 2011.
- [147] Fu Y. , Kwang-Mien C. , Tan K. S. , and Boon-Sain Y. , "*Multi-Metric Gateway Discovery for MANET*", the 63rd IEEE Vehicular Technology Conference (VTC 2006-Spring), pp. 777-781, May 2006.
- [148] Network Simulator 3. [Online]. <http://www.nsam.org>
- [149] Lacage M. and Henderson T. R. , "*Yet another network simulator*", Workshop on NS-2 : the IP network simulator, October 2006.
- [150] Pei G. and Henderson T. R.. Validation of OFDM error rate model in ns3. [Online]. <http://www.nsam.org/~pei/802.11ofdm.pdf>

Liste des publications

Conférences internationales

Crosnier M. , Dhaou R. , Planchou F. , and Beylot A.-L. , *"TCP Performance Optimization for Handover Management for LTE Satellite/Terrestrial Hybrid Networks"*, International IEEE AESS European Conference on Space and Satellite Telecommunications (ESTEL 2012), pp. 1-5, October 2012.

Crosnier M. , Dhaou R. , Planchou F. , and Beylot A.-L. , *"A Cluster-Based Load Balancing Between Satellite Gateways in a MANET"*, IEEE Advanced Satellite Multimedia Systems Conference and Signal Processing for Space communications Workshop (ASMS/SPSC 2012), pp. 303-307, September 2012.

Kretschmar V. , Crosnier M. , Monier M. , and Deslandes V. , *"Optimized hybrid MANET-Satellite networks for public safety"*, IAA International Communications Satellite Systems Conference (ICSSC 2011), December 2011.

Crosnier M. , Dhaou R. , Planchou F. , and Beylot A.-L. , *"Handover Management Optimization for LTE Terrestrial Network with Satellite Backhaul "*, IEEE Vehicular Technology Conference (VTC 2011), pp. 1-5, May 2011.

Participation à l'écriture des livrables de projets européens

MONET, *"D3.1 MONET Architectural Design"*, European Commission Deliverable, 2010.

MONET, *"D3.2 Evaluation Metrics and Test Plan"*, European Commission Deliverable, 2010.

MONET, *"D4.1 Algorithms and Mechanisms for System Optimization"*, European Commission Deliverable, 2010.