



Autour de la réservation de bande passante dans les réseaux ad hoc

Claude Chaudet

► To cite this version:

Claude Chaudet. Autour de la réservation de bande passante dans les réseaux ad hoc. Réseaux et télécommunications [cs.NI]. INSA de Lyon, 2004. Français. NNT: . tel-00007706

HAL Id: tel-00007706

<https://theses.hal.science/tel-00007706>

Submitted on 19 Dec 2004

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.

N° d'ordre : 2004ISAL0053
Année 2004

Thèse

Autour de la réservation de bande passante dans les réseaux *ad hoc*

Présentée devant
l'Institut National des Sciences Appliquées de Lyon

Pour obtenir
le grade de docteur

Ecole doctorale : Informatique et Information pour la Société
Spécialité : Réseaux et télécommunications

par
Claude CHAUDET

Soutenue le 28 septembre 2004 devant la Commission d'examen

Jury

M. André-Luc BEYLOT	Professeur	Rapporteur
M. Michel DIAZ	Directeur de recherche	Président
Mme. Isabelle GUÉRIN LASSOUS	Chargée de recherche	Directeur de thèse
M. Philippe JACQUET	Directeur de recherche	Rapporteur
M. David SIMPLOT-RYL	Professeur	Examineur
M. Stéphane UBÉDA	Professeur	Directeur de thèse

Travaux effectués dans l'équipe Inria Ares, au sein du laboratoire Citi de l'Insa de Lyon.

Table des matières

1	Équité de l'accès au médium : étude d'un cas pathologique	13
1.1	Le protocole IEEE 802.11	13
1.2	Éléments de propagation radio	14
1.3	Description du protocole d'accès au médium	17
1.3.1	Principes de base	17
1.3.2	Évitement de collisions	19
1.3.3	Gestion des communications distantes	20
1.3.4	Résumé	22
1.4	État de l'art	22
1.4.1	Capacité du protocole d'accès au médium	23
1.4.2	Équité d'accès	25
1.5	Modélisation d'un scénario particulier	27
1.5.1	Scénario étudié	27
1.5.2	Description de la modélisation	29
1.5.3	Résultats	39
1.5.4	Influence d'une différence dans les tailles de trames	42
1.5.5	Influence de l'EIFS	45
1.6	Conclusion	45
2	Un protocole de réservation de bande passante : BRuIT	49
2.1	Routage au mieux dans les réseaux <i>ad hoc</i>	49
2.1.1	Routage proactif	49
2.1.2	Routage réactif	50
2.1.3	Autres approches	51
2.2	Qualité de Service et réseaux <i>ad hoc</i>	51
2.2.1	Architectures de qualité de service	51
2.2.2	Différentiation de services	53
2.2.3	Mécanismes de routage avec qualité de service	56
2.2.4	Mécanismes de réservation de ressources	57
2.2.5	Évaluation des ressources disponibles	59
2.2.6	Synthèse	60
2.3	Le protocole BRuIT	60
2.3.1	Une approche réactive	61
2.3.2	Contrôle d'admission	64
2.3.3	Limitations de ce protocole	67
2.3.4	Dépassement de la capacité du médium	72
2.4	Évaluation	73
2.4.1	Mécanisme de réservation	73
2.4.2	Routage	75
2.4.3	Évaluation du coût de ce mécanisme	79
2.4.4	Dégradation des flux	79
2.5	Conclusion	79

3	Co-existence des trafics privilégiés et au mieux	85
3.1	Description formelle du problème	85
3.1.1	Une suite d’allocations admissibles	87
3.1.2	Propriétés de la suite	88
3.2	Algorithme distribué dérivé de la suite d’allocations	91
3.2.1	Simulation de l’algorithme	92
3.3	Vitesse de convergence	100
3.4	Gestion de la mobilité	103
3.4.1	Simulation de l’algorithme mobile	108
3.5	Conclusion	110
4	Réseaux hybrides	111
4.1	Introduction	111
4.2	L’architecture Mobile IP et la micromobilité	111
4.2.1	Mobile IP	111
4.2.2	Micromobilité	112
4.2.3	Architecture considérée	114
4.2.4	Interactions entre le routage d’infrastructure et le routage <i>ad hoc</i>	117
4.3	Stratégies de transmission des notifications de mobilité	117
4.4	Optimisations diverses	120
4.4.1	Remplissage gratuit de cache ARP	120
4.4.2	Longueur des files d’attente	122
4.5	Optimisation de la transmission de notifications de mobilité	123
4.5.1	<i>Differential Route Update</i>	123
4.5.2	<i>Nack Route</i>	124
4.5.3	<i>Nack Only</i>	124
4.5.4	Résultats de simulation	124
4.6	Conclusion	126
A	Spécification des paquets de contrôle de BRuIT	131
A.1	Paquets <i>Hello</i>	131
A.2	Paquets de recherche de route	132
A.3	Paquets de confirmation de réservation	132
A.4	Paquets de libération et d’erreurs	133
B	Liste de publications	135
B.1	Journaux nationaux	135
B.2	Conférences internationales avec comité de lecture	135
B.3	Conférences nationales avec comité de lecture	135

Introduction

Histoire succincte des télécommunications

On entend par télécommunications toute transmission, émission ou réception de signes, de signaux, d'écrits, d'images, de sons ou de renseignements de toutes natures, par fil, radio-électricité, optique ou autres systèmes électromagnétiques.

Constitution de l'Union Télégraphique
Internationale – 1907

La communication est le fondement de toute société humaine. Très tôt dans son histoire, l'homme a souhaité dépasser les limites imposées par la portée de sa voix et par sa perception. Les premiers réseaux de télécommunications au sens de la définition donnée par l'Union Télégraphique Internationale (future Union Internationale des Télécommunications) remontent à l'Antiquité. En Grèce, quatre siècles avant notre ère, des brasiers allumés au sommet de tours permettaient la transmission de signaux lumineux et donnèrent naissance aux phares côtiers. Les signaux sonores tels que les cloches des églises ou les chants tyroliens furent aussi utilisés durant des siècles afin de transmettre des informations. La portée de ces moyens de communication primitifs reste cependant limitée. Transmettre un message sur de longues distances nécessite le déploiement d'une infrastructure de relais. L'absence de codification évoluée limite la complexité des messages transmis.

Le Français Claude Chappe (1763 – 1805), mettant à profit les avancées réalisées en optique par Isaac Newton et Christiaan Huygens, conçoit en 1790 le télégraphe optique. Le 22 mars 1792, il présente à la tribune de l'Assemblée Législative son projet reposant sur un réseau de sémaphores, c'est-à-dire de tours habituellement établies sur les côtes et servant à faire des signaux de la terre aux navires et réciproquement, espacés de quelques dizaines de kilomètres utilisant des longues-vues pour couvrir de longues distances. La première transmission reliant Belleville à Saint-Martin-du-Tertre sur une distance de trente-cinq kilomètres est couronnée de succès le 12 juillet 1793. La première ligne reliant Paris à Lille est installée en 1794 et représentera un atout stratégique dans le contexte troublé de l'époque post-révolutionnaire.

En 1820, le Français André-Marie Ampère (1775 – 1836) invente l'électro-aimant et propose un système de communication à distance utilisant cet outil. En 1832, à Saint-Petersbourg en Russie, le baron Paul Schilling Von Canstadt (1786 – 1837) présente au Tsar Nicolas *I^{er}* le premier télégraphe électrique. Un télégraphe utilisant l'électricité représente une véritable révolution. Il s'agit du premier moyen de communication rapide ne faisant appel ni au son ni à la lumière, utilisant ainsi un médium de communication inintelligible pour quiconque ne dispose que de ses sens. En 1833, les Allemands Karl Friedrich Gauss (1777 – 1855) et Wilhelm Weber (1813 – 1894) relient l'université de Göttingen et son observatoire astronomique distants d'un kilomètre en utilisant un système similaire. En 1839, les Anglais Sir William Fothergill Cooke (1806 – 1879) et Sir Charles Wheatstone (1802 – 1875) installent une ligne de télégraphie électrique longue de vingt kilomètres le long d'une voie de chemin de fer reliant Londres à sa banlieue. En 1837, le peintre américain Samuel Finley Morse (1791 – 1872) conçoit, lui aussi, un télégraphe utilisant un électro-aimant. Il élabore, en 1838, le code qui portera son nom, permettant ainsi la transmission d'un message quelconque par le biais d'impulsions électriques. Après un certain nombre de déboires financiers, Morse inaugure en mai 1844 la première liaison télégraphique interurbaine de Washington à Baltimore et transmet avec succès le message *What hath God wrought?* sur soixante

kilomètres. Si Morse ne peut revendiquer la paternité du télégraphe électrique, il a cependant été le premier à convaincre ses contemporains de financer une telle entreprise. En France, dès 1850, le télégraphe optique sera abandonné au profit du télégraphe électrique.

Le physicien écossais James Clerk Maxwell (1831–1879) publie en 1873 dans son *Traité d'électricité et de magnétisme* les célèbres équations différentielles décrivant la nature des champs électromagnétiques. Il jette ainsi, sans le savoir, les bases des télécommunications modernes. Heinrich Rudolph Hertz (1857 – 1894), physicien allemand, démontre expérimentalement l'existence des ondes de Maxwell en 1877 au moyen d'un mécanisme, appelé par la suite *oscillateur de Hertz*, fabricant des ondes électromagnétiques en produisant des étincelles électriques dans un éclateur, appareil entre les pièces duquel on fait jaillir des étincelles électriques. En 1890, le professeur Edouard Branly, médecin français, découvre les propriétés de la limaille de fer et invente le *cohéreur*, appareil détectant les ondes électromagnétiques. Le premier récepteur d'ondes électromagnétique est constitué d'un tube rempli de limaille de fer, devenant conducteur lorsqu'il est traversé par une onde électromagnétique. Le physicien anglais Oliver Lodge (1851 – 1940) parvient grâce au principe de *syntonie*, c'est-à-dire en accordant un émetteur et un récepteur sur la même fréquence, à améliorer sensiblement l'efficacité de ce système.

En 1893, l'inventeur croate Nikola Tesla (1856 – 1943) montre que des courants à haute fréquence peuvent être utilisés pour établir des communications à distance. Il sera reconnu après sa mort comme étant l'inventeur de la radiocommunication. L'ingénieur russe Aleksandr Stepanovitch Popov (1859 – 1906) découvre le principe de l'antenne qu'il appliqua au cohéreur de Branly en 1894, améliorant ainsi sa sensibilité. Il parviendra de cette manière à transmettre un message de télégraphie sans fil en alphabet Morse en mars 1896. Le physicien italien Guglielmo Marconi (1874 – 1937) étudie les travaux de Branly et de Lodge. En 1895, il décide alors de concevoir une expérience sur la propagation des ondes en extérieur dans la propriété de ses parents à Bologne. En 1896, il met au point la Télégraphie sans fil (TSF). Au fur et à mesure des années, la portée des communications qu'il réalise ne cesse de s'accroître. En décembre 1901, il relie Terre-neuve à la Cornouaille, et atteint la portée record de 3 400 km. Les télécommunications modernes sont nées.

En 1904, l'ingénieur inspecteur des télégraphes et romancier français Edouard Estaunié emploie pour la première fois le mot *télécommunication*. L'utilisation de ce terme sera officialisée par l'Union Télégraphique Internationale en 1907. La TSF ne cessera alors de prendre de l'importance. En 1912, le signal de détresse émis par le navire Titanic est perçu par le navire Carpathia. Ce dernier se détourne alors de sa route, sauvant sept cent cinq passagers. Quelques mois plus tard, la TSF devient le média privilégié des opérations de secours et une conférence internationale impose aux armateurs d'équiper leurs navires d'émetteurs-récepteurs. À la veille de la Première Guerre Mondiale, tous les grands pays disposent d'un réseau de télégraphie sans fil. Durant la Grande Guerre, la TSF représente un réel intérêt stratégique, les autres infrastructures de communication étant beaucoup plus vulnérables au sabotage. À l'issue de cette guerre, les différents gouvernements réalisent que le spectre radio est une ressource limitée et légifèrent sur l'utilisation des différentes bandes de fréquences.

En 1920, la première liaison de radiotélégraphie reliant la France et les États-Unis est ouverte au public et utilise une antenne de quarante-huit hectares pour transmettre des signaux d'une longueur d'onde de vingt-trois kilomètres et demi au moyen d'un émetteur à arc d'une puissance d'un mégawatt. Dès 1920, la radio sera utilisée comme média de diffusion d'informations. Entre 1921 et 1925, les techniques de transmission d'images par ondes radio se perfectionnent. En 1933, Edwin Armstrong (1890 – 1954) invente la modulation de fréquence, plus résistante aux phénomènes de brouillage que la modulation d'amplitude utilisée jusqu'alors.

Le développement des réseaux téléphoniques évincera momentanément l'avancée des technologies de communication radio en dehors de la diffusion radiophonique et télévisuelle. Parallèlement, l'informatique entame la phase moderne de son histoire. Les réseaux de données se développent. En 1955, IBM réalise SABRE (*Semi Automated Business Related Environment*), réseau reliant mille deux cents téléscripteurs de la compagnie American Airlines à travers les États-Unis. En 1961, Leonard Kleinrock, chercheur au *Massachusetts Institute of Technology* publie une théorie prônant l'utilisation de la commutation de paquets pour le transfert de données, établissant ainsi les fondations des télécommunications modernes. Le réseau ARPANET apparaît en 1969, reliant Stanford, l'université de Los Angeles (UCLA), l'université de Santa Barbara et l'université de Salt Lake City.

En 1970, Norman Abramson et son équipe de l'université d'Hawaii conçoivent Aloha [Abr70], protocole assurant des transmissions en mode paquet par ondes radio ultra-hautes fréquences, afin de relier les différentes îles de l'archipel par le biais de satellites. En 1972, une évolution de ce protocole, *Slot-*

ted Aloha [Rob75] améliore les performances des transmissions radiofréquences. Le développement des réseaux en mode paquet ne cessera de prendre de l'ampleur, donnant naissance en 1980 à la norme Ethernet [MB76] encore largement utilisée à ce jour.

Communications militaires

En un mot, la conduite des troupes demande des attentions continuelles de la part d'un général. Sans quitter de vue l'armée des ennemis, il faut sans cesse éclairer la vôtre ; sachez lorsque le nombre des ennemis augmentera, soyez informé de la mort ou de la désertion du moindre de vos soldats.

SUN TZU - l'art de la guerre - Ve siècle av.
J.C.

Les forces armées ont toujours eu besoin de moyens de communication fiables, sûrs et rapides. La transmission radio est, à bien des égards, un média attractif pour les transmissions militaires. Elle est, en effet, invisible aux sens humains, rapide et capable de traverser de nombreux obstacles.

En 1973, l'agence militaire états-unienne DARPA (*Defense Advanced Research Projects Agency*), instigatrice du projet ARPANET, finance un projet appelé PRNet (*Packet Radio Networks*) afin d'étudier les communications radiofréquences en mode paquet dans des réseaux autonomes. À cette époque, les équipements radio sont encombrants et requièrent beaucoup d'énergie pour effectuer des transmissions et les émetteurs sont destinés à être embarqués dans des véhicules. Le réseau doit être capable de prendre en compte dynamiquement les changements dans sa topologie, afin à la fois de réagir à une certaine mobilité ainsi que pour pallier les éventuelles pannes et destructions d'équipement. Les nœuds de ce réseau doivent pouvoir servir de relais afin de rendre possibles les communications entre terminaux hors de portée radio. Un nœud PRNet est composé d'un émetteur-récepteur radio et d'un ou plusieurs terminaux associés reliés par une interface filaire. L'interface radio fournit les fonctions des niveaux physique, liaison et routage du modèle OSI. Ce projet bénéficiera pleinement des avancées en matière de routage dans les réseaux filaires et en matière de gestion des accès concurrents au médium.

Les accès concurrents au médium sont gérés par le protocole CSMA [TK75a], approprié aux médias de transmission diffusants tels que la radiofréquence. Ce protocole repose sur une écoute du canal de transmission couplée à une attente aléatoire avant émission afin de réduire la fréquence des émissions simultanées de paquets rendant souvent la réception de ces paquets impossible. Ce mécanisme simple et distribué permet d'obtenir des performances nettement supérieures au protocole ALOHA. Il reste cependant vulnérable aux problèmes de type station cachée [TK75b]. Deux nœuds hors de portée radio l'un de l'autre voulant transmettre un message simultanément à un voisin commun provoqueront une collision au niveau du récepteur.

Le routage dans ces réseaux est de type vecteur de distance. Chaque nœud maintient une table contenant pour chaque destination son adresse, la distance la séparant du nœud ainsi que l'adresse du voisin à qui transmettre tout paquet adressé à cette destination. Chaque nœud transmet régulièrement à tous ses voisins le contenu de cette table par le biais d'un paquet de contrôle nommé PROP (*Packet Radio Organization Packet*). La bonne transmission des paquets tout au long de la route est assurée par un mécanisme d'acquittements passifs. Chaque nœud captera la transmission du suivant sur la route, pour peu que le lien radio soit bidirectionnel. Le destinataire ne retransmettant pas le paquet devra acquitter explicitement la bonne réception du message. En cas d'échec de transmission, un mécanisme de retransmission est mis en place.

Le projet PRNet aboutira, en 1978 à la conception de réseaux pouvant comporter jusqu'à cent trente-huit mobiles, auto-configurés, robustes et offrant des débits de l'ordre de 100 kbit/s à 400 kbit/s [KGBK78, JT87].

Dans les années 1980, la DARPA crée le projet SURAN (*Survivable Radio Networks*) afin de résoudre les problématiques de sécurité, de passage à l'échelle et de gestion des ressources qui avaient été laissées en suspens lors du projet PRNet. L'armée de terre des États-Unis utilisera les concepts de ce type de réseaux dans son infrastructure de communication à la fin des années 1980 en concevant le système

Single Channel Ground-Airborne Radio System (SINCGARS) encore utilisé de nos jours. L'armée de l'air, via l'expérience *Strategic Command and Control Communications* (C3), et la marine, via le réseau *Intra-Taskforce* (ITF), adapteront elles aussi ces concepts à leurs problématiques particulières.

Dans les années 1990, l'armée états-unienne développe *Tactical Internet*, réseau comportant potentiellement un millier de nœuds mobiles, véhicules et fantassins. Ce réseau se basant sur des versions améliorées des protocoles d'Internet, comme OSPF, permettra essentiellement de démontrer que ce type de réseaux est à même de satisfaire les exigences de la défense.

En 1994, la DARPA crée le projet GloMo (*Global Mobile Information System*) afin d'étudier la possibilité d'adapter les concepts d'Internet en pleine expansion à des utilisateurs mobiles. Le but de ce projet est de créer des réseaux de terminaux peu onéreux, portables, disposant d'une capacité de traitement suffisante pour permettre l'utilisation d'algorithmes robustes sur des terminaux hétérogènes. Les applications doivent pouvoir s'adapter aux changements de topologie ainsi qu'aux changements de qualité du réseau.

La plupart de ces projets auront été réalisés par des universitaires sous l'impulsion de la DARPA. Les avancées réalisées dans le domaine des télécommunications par ondes radiofréquences auront contribué au développement des communications sans fil destinées au grand public.

Communications civiles

Peut-être les robots écoutaient-ils tout ce qui se passait et étaient-ils au courant de ce qu'un humain pouvait désirer à un moment donné. Et si le robot nécessaire n'était pas appelé personnellement pour la tâche en question, le réseau radio qui reliait entre eux tous les robots entraînait en action, convoquant aussitôt le robot voulu à pied d'œuvre.

ISAAC AZIMOV - Face aux feux du soleil -
1957

L'avènement de la microélectronique et la miniaturisation des équipements ont permis la fusion du terminal et de l'interface de communication radio en une seule entité, développant la mobilité potentielle des utilisateurs. Parallèlement, la réduction continue des coûts de production a rendu accessible à chacun les technologies de communication radio. La recherche civile en communications mobiles a très vite saisi l'intérêt commercial pouvant résider dans un moyen de transmission sans fil permettant aussi bien aux êtres humains qu'aux équipements informatiques, et pourquoi pas *domotiques*, de communiquer.

Depuis la création du GSM en 1992, la téléphonie mobile a connu un succès fulgurant dans les pays industrialisés. En France, au troisième trimestre 2003, l'Autorité de Régulation des Télécommunications dénombrait quarante millions de téléphones mobiles pour soixante millions d'habitants. Selon l'institut Médiamétrie, à peine plus de dix ans après la création des réseaux de deuxième génération, 66 % des Français disposent d'un téléphone mobile alors que le taux de pénétration des micro-ordinateurs est estimé, à la même époque, à 42,6 % de la population et celui de l'accès à Internet à 27 %.

Les communications par paquets sans fil n'ont pas été négligées pour autant. Dès 1991, des groupes de travail sont créés à travers le monde afin de fédérer et de dynamiser les recherches en matière de réseaux radiofréquences.

L'*ATM Forum Working Group* publie en 1996 WATM [Dea96b, Dea96a], transposition dans le monde radio de l'architecture *Asynchronous Transfer Mode* (ATM). Cette technologie, intégrant bon nombre de concepts du monde cellulaire, permettra d'effectuer des transmissions à 25 Mbit/s dans la bande de fréquences libre située au-delà de 5 GHz. Cette technologie, intégrée à certains routeurs, est principalement destinée aux opérateurs.

Au Japon, le *Multimedia Mobile Access Communication Systems Promotion Council* (MMAC-PC) est créé en 1996. Ce groupe de travail définira une norme de communications utilisant à la fois la bande de fréquences des 5 GHz et une bande de fréquences située entre 30 GHz et 300 GHz afin de proposer une interface de communication sans fil à très haut débit, destinée à être le dernier lien d'un réseau de fibres optiques.

Le projet BRAN (*Broadband Radio Access Networks*) de l'*European Telecommunications Standards Institute* (ETSI) publie en 1996 la norme Hiperlan/1 [ETS98] (*High Performance Radio Local Area Network*). Ce standard européen définit les couches physique et liaison destinées à permettre une communication sans fil à un débit physique de 20 Mbit/s dans la bande de fréquences située entre 5,1 GHz et 5,3 GHz.

Le groupe 802.11 de l'*Institute of Electrical and Electronics Engineers* (IEEE), quant à lui, finalise en 1997 la première version de la norme qui portera son nom [IEE97]. Cette spécification autorise des transmissions allant jusqu'à un débit physique de 2 Mbit/s dans la bande de fréquence libre située au-delà de 2,4 GHz. L'industrie ne tarde pas à concevoir et à commercialiser les premières cartes d'interface basées sur cette norme. Parallèlement, la norme ne cesse d'évoluer et en 1999, une première révision, nommée IEEE 802.11b [IEE99a], voit le jour. Cette version, proposant des débits allant jusqu'à 11 Mbit/s dans la même bande de fréquences est à ce jour celle dont le succès auprès du grand public a été le plus important. La *Wireless Ethernet Compatibility Alliance* (WECA¹) voit le jour afin d'assurer l'interopérabilité des différentes cartes d'interfaces basées sur cette norme. Dès le mois de mars 2000, la certification Wi-Fi s'impose. La révision 802.11a [IEE99b] du standard est, elle aussi, publiée en 1999. Elle s'appuie sur OFDM, technique de modulation datant des années 1960 remise au goût du jour par les travaux sur les DSL, permettant d'atteindre un débit physique de 54 Mbit/s dans la bande de fréquences située autour de 5 GHz.

En avril 2000 Hiperlan/2 [ETS00a, ETS00d, ETS00e, ETS00b, ETS00c, ETS00f, ETS00g], successeur d'Hiperlan/1, voit le jour. Ce standard utilise lui aussi une modulation OFDM et est physiquement compatible avec IEEE 802.11a et le standard de MMAC-PC. Cette norme est destinée à la fois aux réseaux de troisième génération utilisant une infrastructure UMTS, aux réseaux de type ATM et aux réseaux IP. Elle propose par ailleurs une gestion de qualité de service, destinée aux applications présentant des contraintes temps réel telles que la voix ou la vidéo.

Parallèlement à toute cette agitation, le consortium industriel *Bluetooth Special Interest Group* (SIG), publie en 1999 un protocole [Blu99, BCG02] destiné à assurer une liaison radio entre un micro-ordinateur et ses périphériques associés tels que claviers, imprimantes, etc. Le principal objectif de ce consortium est de fournir un protocole de communication entre périphériques commun à tous les constructeurs, c'est pourquoi il emprunte son nom à Harald Blaataand (910 – 986, littéralement *Harald à la dent bleue*) unificateur du Danemark et de la Norvège. Cette technologie, opérant dans la même bande de fréquence que IEEE 802.11, offre un débit de l'ordre de 1 Mbit/s pour une portée de l'ordre du mètre. Les réseaux d'appareils *Bluetooth* sont organisés en *piconets*, architecture hiérarchique comportant un maître et jusqu'à sept esclaves. Il est possible d'interconnecter plusieurs *piconets* pour former un *scatternet*, certains équipements jouant à la fois le rôle de maître et d'esclave. Ce standard n'a pas été conçu pour faire communiquer des micro-ordinateurs. Cependant, les informaticiens ayant la manie d'utiliser l'intégralité des possibilités d'une technologie, les premiers réseaux Bluetooth n'ont pas tardé à voir le jour.

Même si les publicistes ont parfois forcé la main au grand public quant à l'adoption de ces technologies, les réseaux sans fil ont connu en ce début de millénaire un essor rapide. Les solutions proposées actuellement dans le commerce reposent sur l'utilisation de stations de base qui, à la manière des réseaux cellulaires, jouent parfois le rôle d'arbitre dans l'accès au canal radio, parfois celui de relais de communication.

Ainsi naquirent les réseaux *ad hoc*

La nature ne m'a point dit : *ne soit point pauvre* ; encore moins : *sois riche* ; mais elle me crie : *sois indépendant*.

SÉBASTIEN ROCH NICOLAS, dit NICOLAS
DE CHAMFORT - Maximes et pensées -
1794

La nécessité d'utiliser une station de base pour permettre à des micro-ordinateurs de communiquer ne pouvait résister longtemps au désir de s'affranchir de toute contrainte. Aussi, les résultats des recherches

¹<http://www.weca.net>

militaires sur les PRNet ont très vite été repris et adaptés à ces nouvelles technologies pour donner naissance aux réseaux *ad hoc* tels que nous les connaissons aujourd'hui.

Les réseaux *ad hoc* sont des réseaux sans fil, mobiles, spontanés, ne nécessitant la présence d'aucune infrastructure fixe et capables de s'organiser dynamiquement sans intervention de l'utilisateur. La simple présence de terminaux équipés d'une interface radio, qu'il s'agisse de micro-ordinateurs, d'assistants personnels, de dispositifs embarqués dans des véhicules ou encore de capteurs, suffit à créer un réseau *ad hoc*.

La recherche dans le domaine de ces réseaux polymorphes n'a jamais été aussi active qu'aujourd'hui. Sans doute est-ce dû en partie à la démocratisation des équipements sans fil, sans doute est-ce aussi dû au succès quasi-planétaire d'Internet et à l'engouement suscité par la communication électronique. Peut-être est-ce dû à l'intérêt purement intellectuel suscité par l'abondance de problèmes que peuvent poser des réseaux si génériques qu'il est impossible d'en dégager une application privilégiée.

Les utilisations potentielles en sont en effet multiples et sont un sujet de recherche en soi, comme l'atteste le document [YCG+03]. L'utilisation militaire est bien sûr l'application qui vient tout d'abord à l'esprit, la défense nationale ayant toujours besoin de réseaux mobiles capables de se reconfigurer dynamiquement lorsqu'un élément vient à être détruit. On peut aussi imaginer utiliser ce type de réseaux chaque fois qu'il est impossible, trop long ou trop onéreux de déployer une infrastructure filaire. Par exemple, dans des situations d'urgence, les réseaux *ad hoc* peuvent être utilisés pour organiser les secours. Il est aussi envisageable de déployer de cette manière un réseau dans un monument historique pour organiser les visites, d'offrir un accès à Internet à quelques habitations disséminées sans défigurer le paysage, ou encore de reconstruire rapidement une infrastructure détruite par une catastrophe. Les réseaux *ad hoc* peuvent par ailleurs représenter une solution avantageuse en présence de multiples terminaux tels que des véhicules capables de communiquer des informations de trafic, des capteurs de mesures sismiques, de feux de forêts.

Toutefois, les différences séparant le monde filaire et le monde des réseaux *ad hoc* sont multiples. Il s'agit de réseaux radio et le signal véhiculant l'information, à ce titre, soumis à toutes les contraintes de la propagation radio dans l'air. En comparaison des réseaux filaires actuels, la bande passante offerte par les réseaux sans fil est faible, en partie à cause des régulations empêchant tout un chacun d'utiliser la totalité du spectre radio. De plus, ce débit est partagé entre mobiles voisins, sans qu'il soit possible d'isoler simplement des nœuds. L'atténuation d'un signal radio dans l'air est rapide et il n'est pas possible, en raison de la différence de puissance entre un signal émis et un signal reçu, d'effectuer de la détection de collisions à l'image d'Ethernet, ce qui augmente le surcoût moyen du protocole ainsi que la latence des liens radio. Par ailleurs, ces réseaux sont mobiles, leur topologie n'est *a priori* absolument pas maîtrisée et les protocoles usuels d'Internet ne sont pas toujours adaptés à des changements fréquents de topologie. Il faut aussi prendre en compte des problématiques particulières de sécurité ou encore de consommation d'énergie. Ces réseaux sont donc, à bien des égards, fondamentalement différents des réseaux filaires avec lesquels ils doivent pouvoir s'interconnecter et leurs performances potentielles sont encore aujourd'hui méconnues. Le fonctionnement des réseaux *ad hoc* n'est pas non plus à rapprocher de celui des réseaux cellulaires, ces derniers reposant entièrement sur la présence de stations de base assurant la connexion et le routage par le biais d'un réseau filaire.

Les travaux réalisés sur les PRNet sont d'une grande utilité, grâce à l'analogie évidente entre ces types de réseaux. Toutefois, les réseaux *ad hoc* diffèrent de leurs prédécesseurs à certains égards. L'utilisation d'IP ainsi que l'interopérabilité avec les réseaux publics actuels est une contrainte forte compte tenu de l'omniprésence d'Internet aujourd'hui. La multiplicité des terminaux potentiels dans des scénarios à large échelle doit aussi être prise en compte, ainsi que la diversité des applications pouvant être déployées sur ce type de réseaux.

C'est pourquoi, dès 1995 l'IETF crée le groupe de travail MANET² (*Mobile Ad Hoc Networks*) afin de définir un protocole de routage IP point à point standard dans les réseaux *ad hoc*, tels qu'ils sont définis par la RFC 2501 [CM99]. La tâche de ce groupe de travail n'est pas aisée, chacune des propositions qu'il a eu à examiner étant, bien évidemment, la meilleure. En 2003, quatre protocoles seulement ont été retenus : *Optimized Link-State Routing* (OLSR) [CJ03], *Ad Hoc On Demand Distance Vector* (AODV) [PBRD03], *Dynamic Source Routing* (DSR) [JMH03] et *Topology Dissemination Based on Reverse-Path Forwarding* (TBRPF) [OTG04]. OLSR, AODV et TBRPF ont d'ores et déjà franchi une étape vers la normalisation en devenant *experimental RFC*.

²<http://www.ietf.org/html.charters/manet-charter.html>

Si la normalisation du routage est en bonne voie, de nombreuses problématiques restent ouvertes telles que l'adressage, la sécurité ou encore la qualité de service. L'*Internet Research Task Force* (IRTF) a créé en 2003 le groupe *Ad hoc Network Scaling* afin d'étudier le comportement de ces réseaux et afin d'examiner d'autres problématiques telles que l'interaction entre les différentes couches, l'auto-configuration des mobiles, les problèmes de passage à l'échelle des différents protocoles de routage ainsi que le routage avec qualité de service.

Qualité de service et réseaux *ad hoc*

Nous avons choisi la qualité parce que la chance était devenue trop chère.

JEAN ABRAHAM

Le terme qualité de service peut regrouper une multitude de concepts distincts. Dans le domaine des réseaux et télécommunications, il désigne tout mécanisme permettant d'adapter le comportement du réseau aux besoins des applications. Cette notion englobe les mécanismes permettant d'allouer une proportion des ressources du réseau à un flux de données, de lui garantir, par exemple, un délai borné, un taux de pertes limité. Si le réseau est incapable d'assurer ce niveau de service, il est alors chargé d'en avertir l'application demandeuse.

L'importance de la notion de qualité de service ne s'est pas fait sentir à la création d'Internet, faute d'applications nécessitant de telles garanties. Il a fallu attendre la récente démocratisation des liaisons particulières à haut débit pour voir se développer réellement la transmission de voix et de vidéo sur des réseaux IP. Aujourd'hui, les majors de l'industrie musicale et les grands éditeurs d'œuvres cinématographiques sont à la recherche de nouveaux modes de diffusion permettant de réduire les coûts et délais de distribution, afin de proposer une solution de remplacement au téléchargement illégal. De nombreux opérateurs proposent par ailleurs des offres comprenant accès à Internet, téléphonie IP et diffusion télévisuelle. Chacun de ces services est sensible à l'apparition de congestions dans le réseau résultant d'une occupation du médium supérieure à sa capacité. Pour pallier ce problème, il est possible, dans une certaine mesure, d'augmenter la capacité du cœur de réseau et des réseaux d'accès contournant ainsi temporairement le besoin de solutions de qualité de service.

Parallèlement, les ventes d'ordinateurs portables, de téléphones cellulaires ou encore d'assistants personnels ont connu un essor fulgurant. Le développement du GPRS, l'hypothétique déploiement de l'UMTS et l'apparition dans de nombreux lieux publics tels que gares, aéroports ou cafés de connexions sans fil à Internet préfigurent l'ère de l'homme connecté. Il n'est pas irréaliste d'imaginer, dans un futur proche, un réseau *ad hoc* hétérogène composé de véhicules, de téléphones, de points d'accès offrant services et connexion à Internet à tout le monde, en tout lieu et à tout moment.

Le médium radio, moyen de communication privilégié de l'utilisateur mobile, ne présente cependant pas les mêmes caractéristiques qu'un médium filaire. Le spectre radio est une ressource rare et fortement réglementée. Henry Nyquist en 1924 puis Claude Shannon en 1948 ont fixé les limites du débit d'informations transmissibles sur un canal de bande passante donnée. Or, il ne sera pas possible d'accroître indéfiniment la bande passante allouée aux communications radio contrairement au monde filaire. Il est donc certain que le besoin de solutions de qualité de service performantes et adaptées aux spécificités de ce type de réseaux se fera rapidement sentir.

L'objectif de cette thèse est d'explorer une partie du vaste domaine de la qualité de service dans les réseaux *ad hoc*, c'est-à-dire les problématiques de réservation de bande passante. Nous avons choisi de nous intéresser dans un premier temps à cette métrique plutôt qu'au délai de bout en bout ou à la gigue car contrôler l'utilisation de la bande passante permet, en sus d'offrir des garanties de qualité de service, de limiter l'apparition de congestions dans le réseau. Identifier, prévoir et réagir face à l'apparition de congestion aura une influence positive sur le bon fonctionnement du réseau et par conséquent sur les délais et les taux de pertes des différents flux.

Le premier chapitre de ce manuscrit étudie le caractère équitable du mode de partage des ressources entre mobiles utilisant la norme IEEE 802.11. Au travers de la modélisation et de la simulation d'un scénario particulier, nous mettons en lumière le type de problème pouvant survenir dans un réseau dont la topologie n'est pas maîtrisée. Nous avons opté pour une modélisation théorique des performances des différents mobiles de ce scénario afin de s'affranchir des problèmes liés à la mise en place et à l'absence de

contrôle de l'environnement survenant lors d'expérimentations. Les résultats obtenus par simulation sont par ailleurs fortement dépendants de l'implantation du protocole réalisée. Il ne s'agit pas d'une étude exhaustive des problèmes de performance pouvant survenir dans ce type de réseaux, mais l'existence de ce type de configurations démontre l'existence de limites sur ce qu'il est possible de réaliser lorsqu'on souhaite évaluer la capacité résiduelle d'un canal radio en présence de transmissions distantes et, par conséquent, lorsqu'on souhaite offrir des garanties dans des réseaux radio.

Le deuxième chapitre présente BRuIT, un protocole de réservation de bande passante destiné aux réseaux *ad hoc*. Ce protocole intègre à un protocole de recherche de route avec qualité de service à la demande un mécanisme d'estimation de la capacité résiduelle du canal radio basé sur la transmission régulière d'informations entre nœuds voisins. Ce type de mécanisme permet, lorsque la topologie du réseau s'y prête, de résoudre les problèmes tels que ceux qui sont mentionnés au premier chapitre de ce manuscrit. Ce protocole, adopte une approche *a priori* afin de limiter l'apparition de congestions, à l'inverse des propositions actuelles. L'étude de ce protocole a été réalisée par simulation et permet d'une part de dégager certaines propriétés relatives à son fonctionnement et d'autre part d'identifier les limites de ce mécanisme.

En présence de trafics privilégiés, il est nécessaire de s'assurer que l'existence de trafics au mieux n'invalide pas les réservations de bande passante, ces deux types de flux utilisant la même ressource. Il est donc nécessaire de limiter le débit des trafics au mieux. Toutefois, il est primordial de déterminer la proportion de bande passante à allouer à ces trafics afin de ne pas sous-utiliser le réseau, et ce tout en conservant un comportement équitable. Le troisième chapitre présente un algorithme distribué d'allocation équitable de bande passante pour les réseaux *ad hoc*. Le but de cet algorithme est d'indiquer à chaque mobile la proportion de la capacité du canal qu'il est en mesure d'utiliser sans provoquer de congestion dans le réseau. L'algorithme cherche à la fois à minimiser les écarts entre les bandes passantes allouées à chacun des nœuds et à maximiser l'utilisation des ressources disponibles. Cet algorithme peut être utilisé dans le cadre de BRuIT, ou de façon autonome afin de compenser les problèmes d'équité pouvant survenir dans des topologies déséquilibrées. Les propriétés de cet algorithme sont étudiées de façon théorique et par simulation. Deux algorithmes sont présentés dans ce chapitre. Le premier algorithme présenté est destiné à être utilisé dans des réseaux statiques et le second algorithme est dérivé du premier afin de prendre en compte la mobilité des nœuds.

Enfin, dans des réseaux radio, de nombreuses collisions peuvent survenir lorsque la charge du réseau est importante. Ce phénomène peut avoir un impact sur la bonne transmission des paquets de contrôle nécessaires au bon fonctionnement de BRuIT, ces derniers étant transmis en diffusion. Les trames diffusées ne sont pas acquittées et il est impossible pour un émetteur de s'assurer de la bonne réception des informations transmises. Nous avons donc souhaité étudier la mise en œuvre d'une diffusion acquittée. Les réseaux hybrides, composés d'une infrastructure non mobile mais utilisant un médium sans fil afin d'offrir une architecture à un réseau *ad hoc* sous-jacent, présentent une organisation hiérarchique. La présence de cette structure offre un cadre privilégié pour une première étude des performances d'une diffusion acquittée. Le quatrième chapitre de ce document étudie l'impact du mode de transmission des notifications de mobilité dans des réseaux hybrides. Au travers de simulations, nous montrons l'impact de la charge du réseau sur les performances du protocole de micromobilité assurant le routage dans les nœuds d'infrastructure et la localisation des mobiles. Cette étude montre que sur un médium contraint comme l'est le médium radio, lorsqu'on souhaite améliorer le fonctionnement du réseau par le biais de la transmission régulière d'informations, il est important de déterminer le bon compromis entre la fiabilité de ces informations et le volume de trafic de contrôle nécessaire pour atteindre cette fiabilité.

Finalement, le dernier chapitre rappelle les points principaux développés tout au long de ce document, conclut cette thèse et présente les perspectives et implications relatives aux résultats obtenus. Les travaux contenus dans ce document ont été publiés dans plusieurs conférences et journaux dont la liste est disponible en annexe B.

Équité de l'accès au médium : étude d'un cas pathologique

Le protocole IEEE 802.11 et ses déclinaisons ont été conçus à l'origine afin de constituer des réseaux locaux sans fil administrés par une ou plusieurs stations de base. Grâce à son succès commercial, il est très rapidement devenu incontournable dans le monde des réseaux *ad hoc*. L'écrasante majorité des publications scientifiques dans ce domaine suppose maintenant que le protocole d'accès au médium sous-jacent est celui de la norme IEEE 802.11 ou du moins qu'il a le même comportement. En conséquence, les performances de ce protocole ont été largement étudiées. De très nombreuses simulations ont été réalisées au moyen de simulateurs reconnus, quelques expérimentations ont été menées et le comportement de ce protocole a été modélisé de différentes manières. La plupart des études cherchent à évaluer les performances brutes de ce protocole, telles que le débit maximal pouvant être atteint, le délai moyen, la stabilité des liens, l'influence du nombre de mobiles en concurrence dans le réseau sur ces paramètres ou encore l'impact de perturbations sur les performances du réseau. Si le comportement de ces réseaux est maintenant prévisible dans de nombreux cas, il est possible de dénombrer quelques cas pathologiques mettant en défaut ce protocole.

Dans ce chapitre, nous nous intéresserons à l'un de ces scénarios dans lequel un problème de performance lié à un déséquilibre dans la topologie du réseau affecte l'un des émetteurs. Nous avons choisi d'adopter une approche théorique afin d'évaluer les performances de cette configuration. En effet, lors d'expérimentations, il est difficile de maîtriser tous les paramètres de l'environnement et les simulations fournissent parfois des résultats erronés. L'approche consistant à modéliser ce scénario nous a donc paru être la plus à même de refléter précisément le comportement du protocole d'accès au médium tel qu'il est décrit dans le standard tout en s'affranchissant des problèmes liés à la propagation radio dépendants de la couche physique utilisée.

Ce chapitre présente tout d'abord en section 1.2 quelques éléments de propagation radio, ainsi que le protocole d'accès au médium commun aux différentes révisions de la norme IEEE 802.11 en section 1.3. Un état de l'art des différentes modélisations du comportement de ce protocole est présenté en section 1.4. Le scénario étudié est ensuite décrit et modélisé sous forme de chaîne de Markov en temps discret en section 1.5. Les performances ainsi déterminées sont alors étudiées et comparées à des résultats obtenus par simulation et par expérimentation.

1.1 Le protocole IEEE 802.11

Le protocole 802.11 de l'*Institute of Electrical and Electronics Engineers* (IEEE) [IEE97], parfois nommé Wi-Fi, définit plusieurs couches physiques et une couche d'accès au médium pour les réseaux locaux sans fil (*Wireless Local Area Networks* — WLAN). Si, dans sa première version définie en 1997, des transmissions infrarouges étaient envisagées, les versions les plus récentes du standard telles que IEEE 802.11b [IEE99a], IEEE 802.11g [IEE03] ou encore IEEE 802.11a [IEE99b] sur la base desquelles sont construites l'essentiel des cartes d'interface commercialisées, s'adressent principalement à des transmissions radiofréquences.

Les différentes couches physiques définissent différents codages permettant d'assurer une transmission sans fil fiable et un multiplexage de plusieurs canaux de transmission. Elles rendent possible des transmissions à des puissances limitées dans les bandes de fréquences libres, en particulier la bande de fréquences dédiée aux mondes industriel, scientifique et médical (ISM) située aux alentours de $2,4\text{ GHz}$,

dans laquelle à la fois les États-Unis, l'Europe et le Japon autorisent des transmissions sans licence. La première déclinaison de cette norme définissait, en sus des transmissions infrarouges, les modalités de transmission dans cette bande de fréquences allant de 2 400 MHz à 2 495 MHz. Elle proposait d'utiliser différentes techniques d'étalement de spectre.

L'étalement par saut de fréquences (*Frequency Hopping Spread Spectrum* — FHSS) utilisé par la norme IEEE 802.11 divise la bande de fréquence en 75 canaux d'une largeur de 1 MHz. Une allocation aléatoire d'une séquence de fréquences permet de multiplexer efficacement plusieurs transmissions et offre une bonne résistance aux interférences en bande, c'est-à-dire aux interférences causées par des signaux émis dans la même bande de fréquences que le signal utile. En effet, un récepteur FHSS large bande peut être considéré comme un ensemble de récepteurs en bande étroite changeant de fréquences régulièrement et de façon aléatoire ou pseudo-aléatoire. La probabilité pour que deux signaux utilisent la même bande étroite simultanément est donc réduite. Cependant, le débit résultant est relativement faible et ce mode de transmission n'est utilisé que pour des transmissions à 1 Mbit/s ou 2 Mbit/s. L'étalement de spectre par séquence directe (*Direct Sequence Spread Spectrum* — DSSS) transmet chaque bit en utilisant onze changements d'état du signal. L'émission de chaque bit correspond à la transmission d'une séquence appelée séquence de Barker composée de onze *chips*. Cette technique autorise des transmissions à une vitesse de 1 Mbaud permettant, elle aussi d'obtenir un débit de 1 Mbit/s en utilisant une modulation DBPSK (*Differential Binary Phase Shift Keying*) ainsi qu'un débit de 2 Mbit/s en utilisant une modulation DQPSK (*Differential Quadrature Phase Shift Keying*). La bande de fréquences est alors divisée en 14 canaux d'une largeur de 22 MHz se recouvrant partiellement. Ce type de techniques permet en outre d'aboutir à des durées d'initialisation du système plus rapides que les systèmes FHSS.

L'étalement de spectre à haut débit par séquence directe (*High Rate Direct Sequence Spread Spectrum* — HR-DSSS) permet, par l'ajout d'un encodage de type CCK (*Complementary Code Keying*), d'atteindre des débits de 11 Mbit/s et est à la base de la déclinaison la plus utilisée à ce jour du standard, IEEE 802.11b. Enfin, le multiplexage orthogonal en répartition de fréquences (*Orthogonal Frequency Division Multiplexing* — OFDM) utilise des techniques proches de celles qui sont utilisées par l'ADSL dans la bande de fréquences allant de 5150 MHz à 5350 MHz et permet d'atteindre des débits physiques de 54 Mbit/s. Cette technique, utilisée dans IEEE 802.11a et compatible avec le standard européen Hyperlan/2 [ETS00a, ETS00d, ETS00e, ETS00b, ETS00c, ETS00f, ETS00g], présente quelques propriétés intéressantes, comme une bonne immunité aux problèmes d'évanouissement multitrajets, c'est-à-dire aux problèmes liés aux interférences causées par plusieurs instances du même signal ayant suivi des chemins de longueur différentes et arrivant déphasés au niveau du récepteur. Ces multiples signaux se combinent alors, pouvant s'annuler s'ils sont en opposition de phase. Le principe de HR-DSSS est robuste face à ce type de perturbations, mais limite la portée de transmission.

La plupart des protocoles pour réseaux *ad hoc* n'utilisent cependant pas les possibilités de multiplexage offertes par ces techniques. Affecter un unique canal à un réseau *ad hoc* particulier permet de faire coexister plusieurs réseaux *ad hoc* indépendants dans la même zone géographique. Un changement de canal au sein d'un même réseau poserait en outre des problèmes liés à la dynamique du réseau.

Au-dessus de ces différentes couches physiques, la norme définit un unique protocole d'accès au médium, pouvant fonctionner dans deux modes distincts, afin de gérer les accès concurrents à un même médium partagé. Ce protocole fait partie de la famille des protocoles de gestion des accès multiples par détection de porteuse avec évitement de collisions (*Carrier Sense Multiple Access with Collision Avoidance* — CSMA/CA) [KT75]. Il associe un mécanisme de détection de porteuse avant transmission à un mécanisme d'attente aléatoire permettant de limiter le nombre et l'impact des collisions.

Quelle que soit la technique de transmission mise en œuvre, les signaux émis sont soumis aux lois de la propagation radio. Que ces problèmes apparaissent sur l'unique canal utilisé par le réseau ou sur un canal particulier dédié à la signalisation, ils n'en demeurent pas moins présents et peuvent affecter les performances du réseau, ne serait-ce que de façon transitoire. Afin de définir le niveau de qualité de service pouvant être raisonnablement offert dans ces réseaux, il est indispensable de connaître précisément le comportement des couches sous-jacentes.

1.2 Éléments de propagation radio

La transmission par ondes infrarouges étant devenue anecdotique pour le standard IEEE 802.11, L'essentiel des émissions de trames au moyen de tels matériels a lieu dans des bandes de fréquences de

type micro-ondes. Au niveau physique, un terminal souhaitant transmettre une trame émet sur le canal radio un signal à une certaine puissance dont la limite est fixée légalement et est dépendante de chaque pays. Le tableau 1.1 résume les puissances autorisées en France (à l'exception de certains départements d'outremer) en juillet 2004.

Bande de fréquences	en extérieur	en intérieur
2 400 MHz – 2 454 MHz	100 mW	100 mW
2 454 MHz – 2 483,5 MHz	10 mW	100 mW
5 150 MHz – 5 250 MHz	interdit	200 mW
5 250 MHz – 5 350 MHz	interdit	200 mW ¹ ou 100 mW ²
5 470 MHz – 5 725 MHz	interdit	interdit

TAB. 1.1 – Puissances isotropes rayonnées équivalentes (PIRE) autorisées en France (source : Autorité de Régulation des Télécommunications)

Le signal se propage alors dans l'air et est atténué et modifié par divers phénomènes avant de parvenir à chaque station dans une zone géographique. La taille et la forme de cette zone de réception est définie par les antennes de l'émetteur et des récepteurs potentiels. La figure 1.1(a) représente en coordonnées polaires le résultat d'une mesure réalisée par Guillaume Villemaud dans la chambre anéchoïque¹ de l'université de Limoges de la somme des puissances reçue en polarisations horizontale et verticale en faisant varier la position du récepteur sur un cercle à une distance constante de l'émetteur. Cette mesure a été réalisée en utilisant une carte Wi-Fi standard installée dans un ordinateur portable. La puissance mesurée est loin d'être constante et la zone de couverture est loin d'être circulaire. À partir de ces mesures, il est possible de déduire une mesure du gain en émission de l'antenne d'un récepteur Wi-Fi standard et donc d'obtenir l'allure de la zone de couverture. Cette allure est représentée en figure 1.1(b) dans le cas d'un modèle de propagation des signaux radio simple (espace libre). Cette mesure n'est valable que pour un ordinateur donné équipé d'une carte donnée et dans un environnement particulier. Toutefois, on peut remarquer que la zone de couverture est nettement plus étendue dans la direction de l'antenne que dans les autres, confirmant ainsi les résultats d'expérimentations obtenus par Dominique Dhoutaut et présentés dans [Dho02].

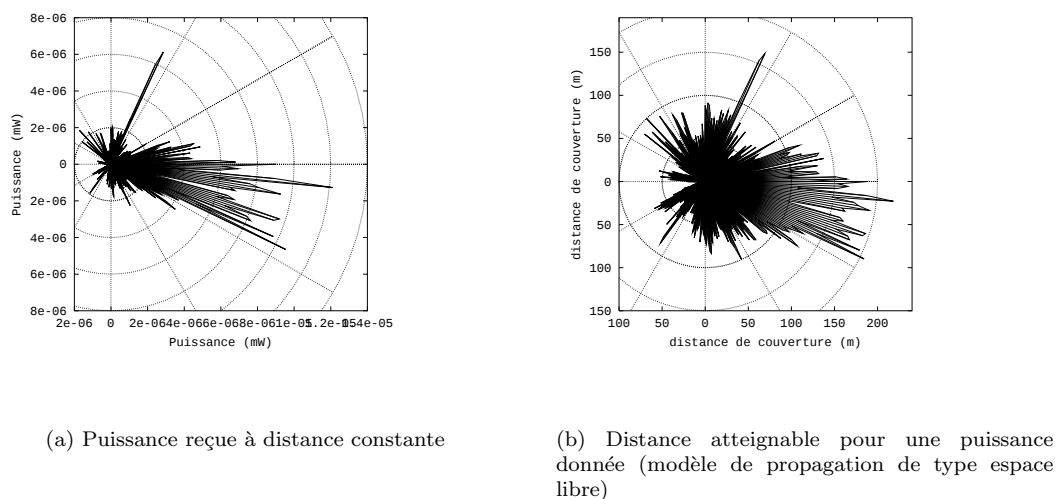


FIG. 1.1 – Caractéristiques d'une antenne de carte Wi-Fi standard (Orinoco / WaveLan)

¹Avec mécanisme de contrôle de puissance et sélection dynamique de fréquences

²Sans mécanisme de contrôle de puissance mais avec sélection dynamique de fréquences

¹Une chambre anéchoïque est une pièce dont les murs absorbent les ondes radio, empêchant ainsi toute réflexion et permettant ainsi de mesurer les caractéristiques d'un signal non-perturbé.

La qualité de réception du signal radio va directement affecter le taux d'erreurs bit. Si le nombre d'erreurs est suffisamment faible pour autoriser la correction des erreurs, le décodage des informations est possible et les données sont transmises avec succès. La spécification des cartes d'interface Avaya indique qu'un taux d'erreur-bit inférieur à 10^{-5} est nécessaire au décodage des informations.

Divers phénomènes peuvent affecter la qualité du signal transmis. En premier lieu, les signaux radio subissent un affaiblissement fonction de la distance qu'ils parcourent et des milieux qu'ils traversent. La propagation radio est, dans le cas général, un phénomène complexe. En espace libre, c'est-à-dire sans obstacles tels que des bâtiments, on peut considérer qu'un signal émis avec une puissance P_t (en Watts) sera reçu à une distance d (en mètres) avec une puissance P_r (en Watts) donnée par l'équation de Friis en espace libre :

$$P_r(d) = G_t \cdot G_r \cdot \frac{\lambda^2}{L \cdot (4 \cdot \pi \cdot d)^2} \cdot P_t,$$

où G_t et G_r sont les gains des antennes émettrice et réceptrice, L le facteur de perte hors propagation dans le système et λ est la longueur d'onde (en mètres) du signal. La puissance du signal décroît en fonction du carré de la distance. Ce modèle est quelque peu simpliste et ne rend pas vraiment compte de l'affaiblissement pouvant survenir dans un environnement urbain ou à l'intérieur de bâtiments.

En environnement réel, un signal se propageant dans toutes les directions de l'espace va se réfléchir contre différents obstacles, comme des bâtiments ou tout simplement le sol. Ces différentes instances du signal parcourront des distances différentes avant d'atteindre la destination, et pourront ainsi arriver déphasées, et atténuer la qualité du signal initial. Considérer une simple réflexion sur le sol permet d'obtenir une bonne approximation à grande distance. Cette approximation consiste à considérer que l'on se trouve en espace libre jusqu'à une distance de $20 \cdot h_t \cdot h_r / \lambda$, h_r et h_t représentant les hauteurs des antennes réceptrice et émettrice. Au-delà de cette distance, la puissance du signal s'exprime par :

$$P_r = G_t \cdot G_r \cdot \frac{h_t^2 \cdot h_r^2}{d^4} \cdot P_t.$$

Cette formulation n'est toujours qu'une approximation des conditions réelles de propagation radio et ne rend compte que d'une réflexion sur le sol. Certains phénomènes comme la diffraction au passage des ouvertures, l'absorption par différents matériaux ou les réflexions contre les différents objets de l'environnement sont dépendants du scénario d'utilisation et ne peuvent être généralisés. La propagation en environnement complexe est communément approchée en considérant un affaiblissement en fonction de l'inverse de la distance à la puissance α . α est une constante dépendant de l'environnement. Elle vaut 2 en espace libre, entre 2 et 5 en environnement urbain et varie de 1,6 à 6 en intérieur.

Le signal, en plus de l'affaiblissement dû au milieu, est aussi perturbé par différentes sources de bruit. Dans une bande de fréquence de largeur $B\text{ Hz}$, à une température de $T\text{ K}$, on considère qu'à cause de l'agitation aléatoire des électrons dans les circuits du récepteur, un bruit thermique blanc de puissance $P_N = k \cdot T \cdot B$ perturbe tout signal, k étant la constante de Boltzmann ($k = 1,379 \cdot 10^{-23} \text{ W} \cdot \text{Hz}^{-1} \cdot \text{K}^{-1}$). Cependant, les principales perturbations pouvant altérer un signal sont dues aux autres signaux émis dans la même bande de fréquence. Les signaux émis simultanément par plusieurs stations vont s'entremêler et devenir difficiles voire impossibles à décoder. Une analogie commune permettant de conceptualiser les transmissions radio est celle de la parole. Une parole est analogue à un signal radio ; son émission est le plus souvent diffuse et atteint tous les correspondants dans un certain voisinage. Si plusieurs interlocuteurs parlent simultanément, les mots s'entremêlent et la conversation est difficile voire impossible à suivre.

Le critère sur le signal permettant un décodage correct des informations doit donc prendre en compte non seulement la qualité du signal reçu mais aussi les perturbations dues aux signaux interférents. Il s'exprime par le rapport entre la puissance du signal reçu et la puissance du bruit, ce bruit prenant en compte l'ensemble des perturbations :

$$SNR = \frac{P_R}{P_N + \sum P_{\text{autres signaux}}}.$$

Les cartes d'interface basées sur la norme IEEE 802.11 permettent actuellement de décoder des signaux dont le rapport signal sur bruit dépasse un certain seuil, fonction de la modulation utilisée. Les données fournies par le constructeur des cartes Orinoco indiquent qu'un rapport signal sur bruit supérieur à 16 dB est nécessaire pour décoder des signaux à 11 Mbit/s . Ce seuil vaut 11 dB pour un

débit de $5,5 \text{ Mbit/s}$, 7 dB pour un débit de 2 Mbit/s et 4 dB pour un débit de 1 Mbit/s . Cette valeur nommée seuil du rapport signal sur bruit, qui sera notée par la suite $SNR_{threshold}$, limite la réutilisation spatiale, c'est-à-dire la possibilité d'effectuer deux transmissions simultanées sans qu'elles n'interfèrent, c'est-à-dire sans collision. Si l'on néglige le bruit thermique et si l'on se place en espace libre, ceci se traduit par l'impossibilité pour un récepteur situé à une distance d d'un émetteur de décoder le signal si un autre terminal émet un signal à la même puissance dans un rayon de $\sqrt{SNR_{threshold}} \cdot d$, soit jusqu'à plus du triple de la distance entre l'émetteur et le récepteur dans notre cas. Dans un modèle d'affaiblissement en $1/d^\alpha$, ce rayon devient $\sqrt[\alpha]{SNR_{threshold}} \cdot d$. Lorsque plusieurs sources d'interférences s'ajoutent, cette distance est évidemment accrue.

À l'image des réseaux filaires, les collisions entre plusieurs trames rendent impossible le décodage des informations. Toutefois, la propagation des ondes hertziennes dans l'air est sensiblement différente de la propagation d'une onde électrique dans un câble ou d'une onde lumineuse dans une fibre optique. En conséquence les mécanismes permettant d'assurer une transmission relativement fiable au niveau liaison doivent être adaptés. En particulier, la couche d'accès au médium définie par le standard IEEE 802.11, bien que basée sur les principes d'Ethernet, diffère sensiblement des protocoles d'accès au médium filaires.

1.3 Description du protocole d'accès au médium

La norme IEEE 802.11 définit deux modes d'accès au médium adaptés aux transmissions radio : le mode centralisé (*Point Coordination Function* — PCF) peut être utilisé lorsque les communications sont gérées par une station de base fixe et le mode distribué (*Distributed Coordination Function* — DCF) est utilisé à la fois pour les communications via une station de base et pour les communications directes de mobile à mobile. C'est ce dernier mode qui sera utilisé dans le cas des réseaux *ad hoc*.

Les collisions, dans un environnement tel que celui décrit à la section précédente, conduisent le plus souvent à la perte des deux trames incriminées. Il semble donc peu judicieux de laisser les différents terminaux émettre comme bon leur semble. L'évaluation du protocole d'accès au médium ALOHA [Abr70], dont le principe est de laisser les terminaux émettre dès qu'ils ont une trame à transmettre, a en effet montré qu'avec des trafics suivant une loi de Poisson de paramètre λ , le débit utile résultant était égal à $\lambda e^{-2\lambda}$. Les travaux de Kleinrock et Tobagi [TK75a] montrent que ce débit, fonction directe de la fréquence de génération des trames n'excède pas 18 % de la capacité du médium.

Laisser les terminaux transmettre à leur guise ne conduit donc pas à une utilisation de la bande passante efficace. Plusieurs améliorations sont envisageables afin d'utiliser au mieux la bande passante en réduisant les collisions. Tout d'abord, il est possible de ne pas émettre si le médium est occupé. Les terminaux attendent alors que le canal se libère avant d'émettre un signal. Ce principe simple constitue la base des protocoles dits à détection de porteuse (*Carrier Sense Multiple Access* — CSMA). Cette famille regroupe entre autres des protocoles tels qu'Ethernet [MB76] ou le protocole IEEE 802.3 [IEE85] dont l'efficacité peut atteindre 95 %.

1.3.1 Principes de base

Ethernet ajoute à ce mécanisme un procédé de détection de collision. Il peut être classé parmi les protocoles de type CSMA/CD (*Carrier Sense Multiple Access with Collision Detection*). Chaque station, au fur et à mesure de la transmission d'une trame, compare le signal présent sur le médium à celui qu'elle émet. Une différence indique alors une collision et la transmission est arrêtée, libérant le canal pour d'autres transmissions. Ce procédé ne peut cependant pas être appliqué dans un contexte radio. En effet, le fort affaiblissement des signaux radio rend impossible la détection d'un signal perturbateur par l'émetteur. La puissance perçue d'un tel signal serait en effet beaucoup trop faible en comparaison de la puissance du signal émis. De surcroît, les collisions ont un impact à la réception et non à l'émission. La présence d'un signal interférant au niveau de l'émetteur ne fournit aucune indication sur le niveau de bruit au niveau du récepteur. La DCF du protocole IEEE 802.11 met donc en œuvre un certain nombre de mécanismes visant non pas à détecter les collisions mais à les éviter. Elle fait à ce titre partie de la famille CSMA/CA (*Carrier Sense Multiple Access with Collision Avoidance*).

Compte tenu de l'impossibilité pour les émetteurs de mesurer la qualité du signal au niveau du récepteur, chaque récepteur doit acquitter toute trame qui lui est explicitement destinée et cet acquittement constituera la seule indication sur le succès de la transmission dont disposera l'émetteur. Ce

mode de fonctionnement est très proche de celui du protocole d'accès au médium MACAW [BDSZ94], dont la DCF de la norme IEEE 802.11 s'inspire très largement. Lorsqu'un terminal reçoit une trame de données, il procède à une détection d'erreurs au moyen d'un CRC standard IEEE sur 32 bits. Si la trame ne contient pas d'erreur, il renvoie à l'émetteur un acquittement. L'intervalle de temps séparant la fin de la réception de la trame de données et le début de l'émission de l'acquittement est égal à une valeur constante *SIFS* (*Short Inter Frame Spacing*). Il n'est évidemment pas possible d'acquitter les trames diffusées (*broadcast*), du fait de l'impossibilité pour l'émetteur de recevoir plusieurs acquittements simultanés. Aussi n'y a-t-il aucune garantie sur la bonne réception de ce type de trames.

Lorsqu'un terminal désire transmettre une trame, il s'assurera tout d'abord que le médium est libre durant un temps constant *DIFS* (*DCF Inter Frame Spacing*) plus long que *SIFS* afin de donner une priorité absolue aux acquittements. Le cas échéant, il effectue la transmission, puis attend l'acquittement correspondant de la part du récepteur. L'absence de réception de cet acquittement provoque la retransmission de la trame et ce processus sera répété jusqu'au succès de l'opération ou jusqu'à atteindre le nombre maximal de retransmissions autorisé. Dans ce dernier cas, la trame est détruite.

La détection de porteuse permet d'éviter certains cas de figure dans lesquels deux émissions simultanées provoqueraient une collision au niveau d'un récepteur. Cependant, il est impossible de distinguer par ce biais les situations dans lesquelles deux émissions simultanées ne provoqueraient pas de collision, à l'image du scénario représenté en figure 1.2. Dans cette configuration, les émetteurs *B* et *C* se trouvent en zone de détection de porteuse, c'est-à-dire que les émissions de l'un bloque le mécanisme de détection de porteuse de l'autre. Cependant, les deux récepteurs sont suffisamment éloignés des émetteurs perturbateurs pour autoriser la simultanéité des deux communications. Ce problème, connu comme le problème de la station exposée conduit à une sous-utilisation de la capacité du canal radio.

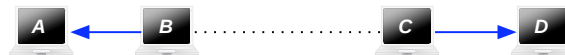


FIG. 1.2 – Problème de la station exposée

Si l'émetteur constate que le médium est déjà occupé lorsqu'il souhaite émettre, il reporte sa transmission jusqu'à la libération du médium. Toutefois, si plusieurs stations sont en attente de la fin d'une même transmission, elles ne doivent pas commencer à émettre au moment où cette transmission cesse, sans quoi une collision surviendrait irrémédiablement. C'est pourquoi, lorsque le canal radio se libère, tout émetteur désirant accéder au médium attend un temps aléatoire en plus d'un intervalle *DIFS*. Chaque émetteur potentiel tire de façon uniforme un nombre aléatoire (appelé *backoff*) dans un intervalle appelé fenêtre de contention. Cette valeur est ensuite décrémentée d'une unité à chaque intervalle de temps passé sans que le médium ne soit occupé. La première station à atteindre la valeur 0 émet alors sa trame. Les autres stations suspendront le processus qui sera repris dès la fin de la transmission. Un nœud voulant émettre plusieurs trames en séquence devra passer par une procédure d'attente aléatoire entre deux trames afin de ne pas monopoliser le canal radio.

L'exemple de la figure 1.3 met en scène trois stations à portée de communication. À la date t_0 , le nœud 1 est en train d'émettre, les deux autres attendent la libération du canal. À la date t_1 , la transmission est terminée, les trois émetteurs patientent un temps *DIFS* avant de commencer à décrémenter leur *backoff*. Le nœud 3 possède le plus petit nombre aléatoire et gagne la contention à t_2 . Le processus de décrémentement est alors suspendu pour les deux autres et reprendra à t_3 . Le nœud 3 sera alors le seul à tirer une valeur aléatoire et le processus reprendra.

Ce mécanisme ne permet évidemment pas de supprimer les collisions entre trames. Si deux émetteurs tirent la même valeur aléatoire, ils émettront au même instant. À l'image d'Ethernet, l'attente aléatoire sera tout d'abord tirée dans un intervalle de valeurs faible. En cas de collision, la taille de cet intervalle sera doublée pour les stations incriminées. Si une collision se reproduit, cette valeur sera encore doublée, jusqu'à atteindre une limite définie par la norme. Une transmission réussie avec succès réinitialise la fenêtre de contention. Cette augmentation exponentielle de l'attente aléatoire permet à la fois de ne pas ajouter un surcoût trop élevé dans les réseaux comportant peu de terminaux et de permettre un certain passage à l'échelle en terme de nombre d'émetteurs. Toutefois, une collision ne pouvant être détectée, l'envoi d'une trame ne sera jamais interrompu et occupera le canal radio jusqu'à la fin de la transmission.

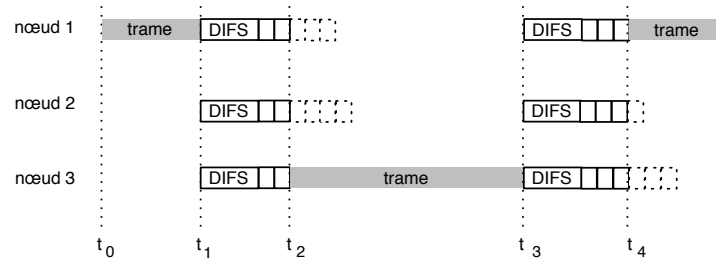


FIG. 1.3 – Exemple d'accès au médium pour 3 mobiles

C'est pourquoi la taille initiale de la fenêtre de contention est plus importante que pour Ethernet (32 unités de temps à la place de 2).

1.3.2 Évitement de collisions

À l'image du mécanisme introduit par le protocole d'accès au médium radio MACA [Kar90], il est possible de précéder l'envoi de chaque trame de données par un échange de messages courts. L'émetteur envoie au récepteur une requête d'émission (*Request To Send* — RTS). Le récepteur, si le canal radio est disponible, autorise l'émetteur à transmettre par une confirmation (*Clear To Send* — CTS). À la réception de l'autorisation, l'émetteur transmet la trame de données. L'intervalle de temps séparant la réception d'une des trames de cet échange (RTS, CTS, données et acquittement) et l'émission de la suivante est égal à *SIFS* afin d'empêcher l'interruption de ce mécanisme par une autre trame. Tout mobile à portée radio de l'émetteur ou du récepteur captera l'une de ces trames contenant la durée de l'envoi de la trame correspondante. Ces voisins s'abstiendront alors de transmettre jusqu'à la fin de cette trame afin de ne pas provoquer de collision. Ce mécanisme permet de réduire l'impact des collisions puisqu'elles n'arriveront essentiellement que sur des trames courtes. Toutefois, cet échange ajoute un surcoût à chaque trame. Ces trames étant transmises à un débit de 2 Mbit/s afin de garantir une certaine compatibilité avec la première version du standard, la durée de cet échange est de $540\mu s$, soit le temps nécessaire pour transmettre 5 940 bits à un débit de 11 Mbit/s. Il n'est donc pas rentable de l'utiliser systématiquement. Les cartes d'interface proposent en général de n'utiliser ce mécanisme que pour des trames excédant une taille paramétrable.

Ce mécanisme permet en outre de résoudre les situations comme celle qui est représentée en figure 1.4. Dans ce scénario, appelé problème de la station cachée [TK75b], deux émetteurs *A* et *C* souhaitent émettre une trame en direction du même récepteur *B*. *A* et *C* ne sont pas à portée radio et ne détectent donc pas les émissions de l'autre. Sans échange RTS-CTS préalable à la transmission, les deux trafics engendreraient régulièrement des collisions. Avant de transmettre une trame, *A* envoie un message RTS à *B*. *B* autorise la transmission en répondant par un message CTS à destination de *A*. Le médium radio étant par nature diffusant, ce message atteindra *C* qui sera alors informé que le médium sera occupé durant une durée correspondant à l'émission de la trame. *C* n'émettra alors pas durant cette période et ne provoquera pas de collision au niveau de *B*. Ce principe de réservation du médium est appelé détection de porteuse (*Virtual Carrier Sense*) et la période de réservation est appelée vecteur d'allocation du réseau (NAV — *Network Allocation Vector*).

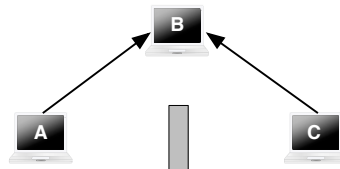


FIG. 1.4 – Problème de la station cachée

Le mécanisme de RTS-CTS ne permet cependant pas de résoudre tous les cas de stations cachées. Par exemple, considérons le scénario représenté en figure 1.5. Si C émet un RTS à destination de D au moment où B émet un CTS à destination de A , le CTS ne sera pas compris par C et la transmission entre C et D pourra avoir lieu, provoquant une collision au niveau de B . Ce type de situation survient toutefois rarement puisqu'il est nécessaire que les émissions du RTS de C et du CTS de B débutent simultanément.

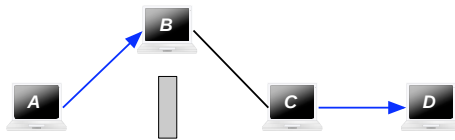


FIG. 1.5 – Situation mettant en défaut le mécanisme RTS-CTS

1.3.3 Gestion des communications distantes

Le problème de la station cachée est un des nombreux problèmes pouvant survenir dès lors que tous les participants du réseau ne sont pas à portée radio. En effet, ne pas connaître la totalité de la topologie du réseau, ou du moins l'ensemble des mobiles avec qui l'on est en compétition pour l'accès au médium ne facilite pas la conception d'un protocole efficace.

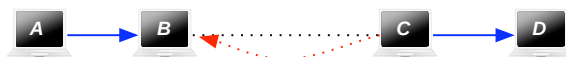


FIG. 1.6 – Collisions entre trafics distants

Prenons par exemple le cas représenté en figure 1.6. Dans ce scénario, les mobiles B et C ne sont pas à portée de communications, mais sont suffisamment proches pour que la réception d'une trame par B soit perturbée par le signal résultant de l'envoi d'une trame par C . Dans ce cas, il est impossible pour B de réserver le canal radio en envoyant une trame CTS, puisque C ne la comprendra pas. C'est pour palier ce type de situations que les cartes d'interfaces basées sur la norme IEEE 802.11 définissent un seuil de détection de porteuse inférieur au seuil de communication. En d'autres termes, la présence d'un signal distant sur le médium empêche une station de transmettre une trame. Il existe donc autour de chaque mobile deux zones distinctes, comme représenté en figure 1.7. Une première zone est définie par la valeur minimale de puissance nécessaire au décodage d'un signal et est directement fonction de la performance de la modulation utilisée, du taux d'erreurs bit résultant et de la puissance du code correcteur d'erreurs utilisé. Une seconde zone, plus étendue, est définie par la valeur de puissance minimale pour que la carte d'interface considère le médium comme occupé. [DGL03] présente des résultats de mesures indiquant que le rapport entre le rayon de ces deux zones semble être approximativement de 2. En termes de puissance, pour des cartes d'interface Avaya, le seuil de réception varie de -83 dBm à 11 Mbit/s à -94 dBm à 1 Mbit/s .

Considérons maintenant la situation représentée par la figure 1.8. Dans cette configuration, les mobiles B et C ne sont pas à portée de communication, mais la distance est telle que si A et C émettent simultanément une trame, le rapport signal sur bruit résultant au niveau de B ne permet pas le décodage des données. En considérant le modèle de propagation en espace libre présenté en section 1.2, ce type de situation peut survenir dès lors que la distance entre B et C est inférieure au triple de la distance entre A et B mais est supérieure à la portée de communication des mobiles. Plaçons nous dans le cas où la distance entre A et C serait suffisamment grande pour qu'une trame émise par A ne se traduise pas par un niveau de signal suffisamment élevé au niveau de C pour que C considère le médium comme occupé. Ce type de situation survient si les distances répondent aux conditions exprimées ci-dessous (R représentant la distance de communication et CS la distance de détection de porteuse). Dans le cas d'une zone de détection de porteuse de rayon double de la zone de communication et pour un modèle de propagation de type espace libre, l'allure de la zone où peut se situer le nœud C est représentée par la zone non grisée sur la figure 1.9.

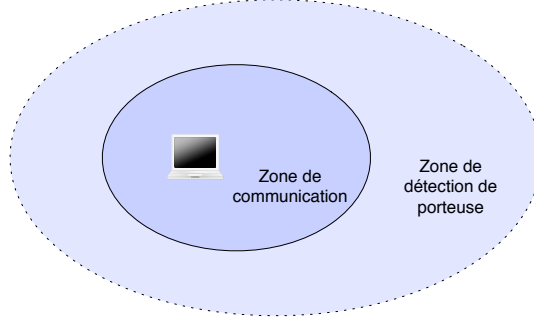


FIG. 1.7 – Zones de communication et de détection de porteuse



FIG. 1.8 – Collisions entre données et acquittements

$$\begin{cases} d(A, B) & \leq R; \\ d(B, C) & > R; \\ d(A, C) & > CS = 2 \cdot R; \\ \left(\frac{d(B, C)}{d(A, B)} \right)^\alpha & \leq SNR_{threshold}. \end{cases}$$

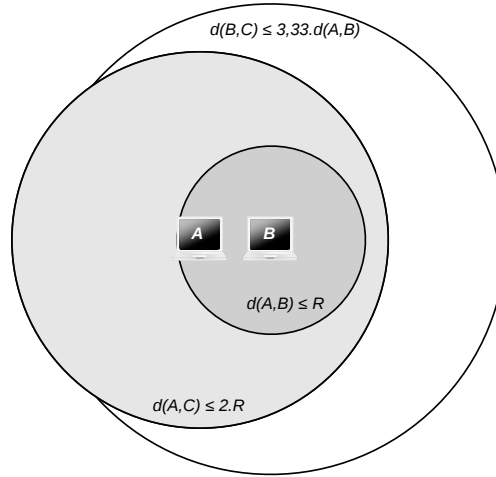


FIG. 1.9 – Zone dans laquelle un émetteur peut provoquer une collision au niveau de B avec les acquittements de A

Dans le scénario de la figure 1.8, si B transmet une trame à destination de A , C considérera que le médium est occupé. Toutefois, la transmission de l'acquittement de A vers B sera invisible pour C qui pourra alors commencer à transmettre des données à destination de D avant la fin de l'acquittement de A . Dans ce cas, l'acquittement de A sera brouillé par les données de C et B conclura que la transmission précédente a échoué et retransmettra ses données. Le scénario est bien sûr symétrique. C'est pour palier à ce genre de situations que le standard IEEE 802.11 impose aux émetteurs détectant sur le canal une trame qu'ils ne peuvent pas décoder de différer l'émission de leurs données d'un temps supérieur à $DIFS$ appelé $EIFS$ dimensionné pour permettre la transmission d'un acquittement.

Le standard spécifie que ce temps allongé doit être utilisé lorsque la somme de contrôle d'une trame de donnée est erronée. Or, les en-têtes physiques des trames IEEE 802.11 (*PLCP Header* et *PMD header*) sont transmises en utilisant une modulation DBPSK à un débit de 1 *Mbit/s*. Cette partie des trames peut être décodée correctement à une distance supérieure aux données transmises à un plus haut débit. Il existe donc une zone autour de chaque mobile dans laquelle un émetteur potentiel peut décoder les en-têtes physiques des trames et identifier le signal comme étant une trame de données mais non le contenu. Cet émetteur utilisera donc automatiquement le temps d'attente *EIFS* s'il désire transmettre une trame.

1.3.4 Résumé

Les différentes valeurs des paramètres définis par la norme IEEE 802.11b sont résumées dans le tableau 1.2. La figure 1.10 représente les différentes étapes de la transmission d'une trame entre un émetteur *E* et un récepteur *R* sans échange RTS-CTS et la figure 1.11 représente une transmission avec échange RTS-CTS.

Paramètre	Durée	Taille
SIFS	10 μs	n/a
Backoff slot	20 μs	n/a
DIFS	50 μs	n/a
EIFS	364 μs	n/a
Entête physique longue (PLCP)	192 μs	192 bits
Entête physique courte (PLCP)	96 μs	120 bits
Entête MAC + somme de contrôle (H)	24,7 μs	34 octets
RTS (hors en-tête physique)	80 μs	20 octets
CTS (hors en-tête physique)	56 μs	14 octets
Acquittement (hors en-tête physique)	56 μs	14 octets

TAB. 1.2 – Valeurs caractéristiques de la norme IEEE 802.11b (mode 11 Mbit/s - HR-DSSS)

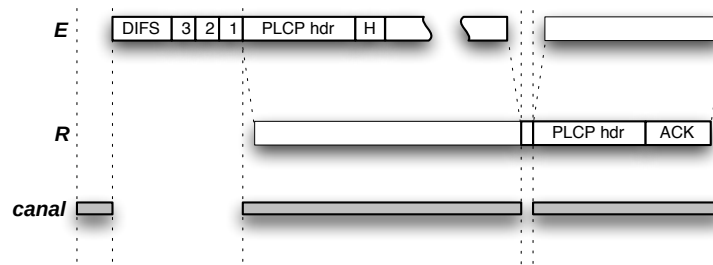


FIG. 1.10 – Détail de la transmission d'une trame

Le fonctionnement de ce protocole est complexe. Les mécanismes tels que le *backoff*, l'EIFS ou l'échange RTS-CTS permettent de résoudre bon nombre de problèmes pouvant survenir couramment dans des réseaux radio, mais leur utilisation représente un surcoût non négligeable. Les performances des solutions IEEE 802.11 semblent donc très liées au scénario d'utilisation et il est difficile de les déterminer *a priori*.

1.4 État de l'art

L'évaluation des performances du protocole IEEE 802.11 a été la source de nombreuses publications. Un certain nombre d'évaluations théoriques analysent le surcoût introduit par le mode d'accès au médium, la pertinence de l'échange RTS-CTS ou plus simplement le débit espéré ou le délai des trames. Toutefois,

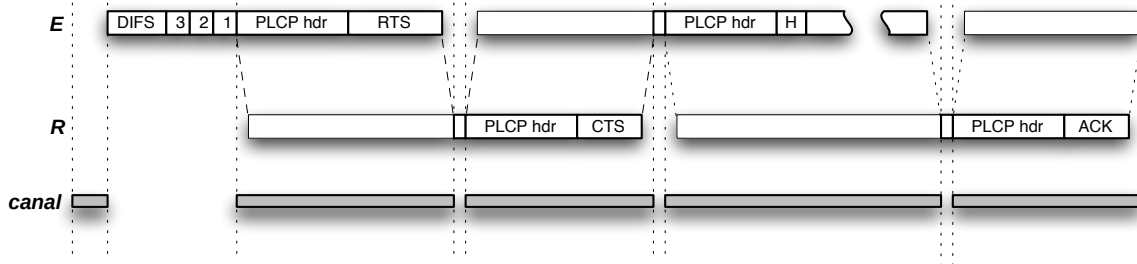


FIG. 1.11 – Détail de la transmission d’une trame avec échange RTS/CTS

l’évaluation de ce protocole dans un contexte *ad hoc* multisauts est complexe, du fait du caractère aléatoire de la topologie et assez peu de travaux s’y sont intéressés.

1.4.1 Capacité du protocole d’accès au médium

Dans [CCG98], Calì et al. évaluent l’efficacité du protocole IEEE 802.11 dans un réseau à un saut, en saturation, sans mécanisme RTS-CTS lorsque les tailles des trames suivent une loi géométrique. L’analyse en moyenne présentée dans cet article permet d’obtenir une expression du rapport entre le temps nécessaire pour envoyer un certain volume de données et le temps effectivement passé pour envoyer ces données, évaluant ainsi le surcoût lié au protocole d’accès au médium. Cette valeur semble proche des résultats de simulation obtenus. Les auteurs montrent que le surcoût du protocole a un impact important sur les trames de petite taille et que le nombre d’émetteurs en concurrence a un impact négatif sur les performances du protocole. À partir d’une évaluation de la limite supérieure de la capacité du médium, ils proposent en outre une modification de l’algorithme de détermination de la fenêtre de contention adaptée au nombre de nœuds en concurrence. Ils montrent que tant que ce nombre de stations reste compris entre la moitié et le double du nombre de stations ayant été utilisé pour le réglage de la taille de la fenêtre de contention, les performances du protocole restent proches de leur limite théorique.

Giuseppe Bianchi, dans [Bia00], propose une modélisation du mécanisme d’attente aléatoire sous forme de chaîne de Markov en temps discret afin de déterminer le débit total en saturation de la DCF. Pour un réseau comportant n stations en concurrence, chacune ayant toujours une trame à émettre, ce modèle suppose que la probabilité de collision est constante pour chaque trame. L’auteur modélise tout d’abord le comportement d’une station isolée afin de déterminer la probabilité qu’un terminal émette une trame dans un intervalle de temps donné ainsi que la probabilité de collision. Le débit de saturation maximal peut, dans cette situation être approché par $1 / (n \cdot \sqrt{T_c^*/2})$ où T_c^* est la durée moyenne d’une collision exprimée en nombre d’unités de *backoff* et est fonction de la longueur des trames. L’efficacité d’un réseau ainsi déterminée est fonction du nombre de mobiles en concurrence et de la probabilité de transmission et approche 85 % dans le meilleur cas.

Tay et Chua, dans [TC01], évaluent le débit en saturation d’un réseau à un saut, c’est-à-dire synchronisé, sans bruit et sans mécanisme RTS-CTS, cet article montre que la probabilité de collision ne dépend que de la taille des fenêtres de contention initiale et maximale et du nombre de stations en concurrence. Ils montrent par ailleurs que le débit de saturation ne dépend que du rapport la taille de la fenêtre de contention initiale et le nombre de stations. En conséquence, diviser par deux la taille de la fenêtre de contention initiale est équivalent en terme d’utilisation du canal à doubler le nombre de stations en concurrence. Enfin, ils montrent que pour utiliser au mieux la capacité d’un tel réseau, il est nécessaire de choisir la taille de la fenêtre de contention en fonction de la racine carrée de la taille de paquet. Les analyses présentées dans cet article, ainsi que les expressions sont validées en utilisant le simulateur de Giuseppe Bianchi.

Vishnevsky et Lyakhov, dans [VL02a], s’intéressent à l’impact d’une perturbation du canal radio par un bruit homogène sur les performances du protocole IEEE 802.11 dans le cas d’un réseau à un saut. L’ajout de cette perturbation au niveau radio aura pour conséquence un accroissement du taux d’erreurs-bits sur chaque trame. Les auteurs modifient le modèle proposé par Bianchi en ajoutant une probabilité, fonction de la longueur de trame, pour chaque trame de subir une perturbation. Ils aboutissent à une

estimation du débit résultant et de la probabilité de rejet d'une trame fonction du nombre de stations dans le réseau et de l'intensité du bruit perturbateur. Ces résultats analytiques sont confrontés à des résultats obtenus par simulation et l'écart entre ces deux estimations est de l'ordre de 5 %.

Dans [VL02b], les mêmes auteurs modifient les études effectuées par Bianchi et par Calì *et al.* afin de prendre en compte le fait que, lorsque de nombreuses collisions surviennent dans un réseau, c'est-à-dire dès que le réseau est saturé, les tailles des fenêtres de contention des différents émetteurs concurrents augmentent. Ainsi, un émetteur ayant réussi une transmission réinitialisant sa fenêtre de contention aura une probabilité plus élevée que ses concurrents de réémettre une trame. Ce phénomène conduit à des comportements inéquitables puisque les différents émetteurs transmettent par rafales. Les résultats découlant de cette analyse se rapprochent des résultats de simulations effectuées au moyen du simulateur GPSS (*General Purpose Simulation System*), auxquels ils sont comparés. La précision de l'évaluation du débit à saturation et de la probabilité de perte d'une trame est améliorée dans les cas où la taille de la fenêtre de contention initiale est faible, c'est-à-dire lorsque le phénomène considéré a le plus d'impact.

Dans [HG01], Heindl et German proposent une modélisation de la DCF de 802.11 en utilisant des réseaux de Petri stochastiques. Ce formalisme permet de décrire des mécanismes définis par le standard IEEE 802.11 généralement laissés de côté comme l'utilisation d'*EIFS* ou la fonction synchronisation (TSF), mécanisme de synchronisation des horloges des mobiles basé sur l'émission périodique par une station de base de trames particulières. Les auteurs décrivent un modèle de station puis étudient les performances de réseaux composés de trois et dix mobiles sans multisautes. Le canal est supposé parfait et la génération de trames par les différents mobiles suit une loi de Poisson. Les résultats obtenus montrent que l'utilisation de la fonction de synchronisation n'a que peu d'impact sur les performances et que l'utilisation de l'*EIFS*, déclenché lors de collisions, peut se révéler profitable dans des cas où le réseau est moyennement chargé. Toutefois, les valeurs déterminées par cette étude sont sujettes à caution car la durée de l'*EIFS* considérée est plus du triple de la durée définie par la norme pour ce paramètre. Si les considérations générales ne pâtissent sans doute pas de cette erreur, la quantification de l'impact réel de ce mécanisme ou la valeur à partir de laquelle l'*EIFS* représente une perte de performances devraient sans doute être recalculées. Les débits obtenus se situent aux alentours de 80 % de la capacité du médium dans le meilleur cas sans utilisation du mécanisme RTS-CTS et avoisinent les 90 % avec ce mécanisme, ce qui peut paraître surprenant compte tenu du faible nombre de stations (dix) en concurrence.

La première analyse du protocole d'accès au médium du standard IEEE 802.11 original dans un contexte réellement *ad hoc*, c'est-à-dire lorsque tous les mobiles ne sont pas à portée de communication les uns des autres, a été publiée par Chhaya et al. dans [CG97]. Cet article présente une étude analytique du débit offert par ce protocole. Cette étude ne considère cependant ni interférences, ni bruit perturbateur. Elle s'applique dans des réseaux faiblement étendus présentant une mobilité limitée. Les auteurs calculent la probabilité de succès d'une transmission et en déduisent le débit d'une station et le débit global du réseau comme somme des débits individuels. La pertinence du mécanisme RTS-CTS est évaluée et cette étude montre que si la probabilité de succès d'une transmission est accrue grâce à ce mécanisme, il n'en est pas toujours de même pour le débit total. Le surcoût introduit par ce mécanisme limite en effet le gain de performance en fonction, bien sûr, de la longueur de trame mais aussi des caractéristiques de trafic et de la topologie avoisinante des mobiles. Une étude d'équité montre que les mobiles en bordure de réseau ont une probabilité de succès accrue par rapport aux mobiles en compétition avec plus de nœuds cachés. L'intérêt présenté par le mécanisme RTS-CTS est en outre accru si les émetteurs utilisent des longueurs de trames différentes.

Gupta et Kumar, dans [GK00], déterminent une borne sur la capacité des réseaux *ad hoc* multi-sauts mettant en œuvre un protocole de type CSMA/CA. Lorsqu'un tel réseau, composé de n nœuds disposés aléatoirement, est étendu, c'est-à-dire lorsque le nombre de sauts séparant chaque source de chaque destination croît, et en l'absence d'interférences, le débit de bout en bout disponible pour chaque communication est de l'ordre de $\Theta(W/\sqrt{n})$, W étant la capacité du canal radio, et ce dans le cas d'un positionnement bien choisi des nœuds avec une portée de transmission elle aussi bien choisie et dans le cas d'un profil de trafic approprié. Dans le cas d'un placement aléatoire, le réseau ne pourra offrir à chaque nœud qu'une capacité de l'ordre de $\Theta(W/\sqrt{n \cdot \log(n)})$. Ces bornes ont été comparées à des résultats d'expérimentation dans [GGK01]. Les tests réels, menés sur des réseaux de 2 à 12 nœuds transmettant à un débit binaire de 2 Mbit/s, se révèlent plus restrictifs encore puisque le débit disponible pour chaque nœud est de l'ordre de $\Theta(W/n^{1.68})$.

Li et al., dans [LBDC⁺01] étudient par simulation et analytiquement la capacité de réseaux *ad hoc* utilisant le protocole IEEE 802.11. Au travers de l'évaluation du débit de bout en bout que peuvent

espérer obtenir les émetteurs de différents scénarios tels qu’une chaîne de nœuds, une grille régulière ou un scénario aléatoire, ils montrent qu’il est possible de s’approcher de la borne théorique de $W/\sqrt{n}$ déterminée par Gupta et Kumar. Cependant, dans le cas général, les performances de réseaux entièrement aléatoires sont bien en deçà de cette limite. Le passage à l’échelle des réseaux *ad hoc* ainsi formés semble donc difficile. Cependant, ce problème de performances est fortement réduit lorsque l’on considère que les trafics gardent une certaine localité, c’est-à-dire lorsque la distance en nombre de sauts entre les émetteurs et leurs récepteurs associés reste en moyenne faible, évitant la saturation du cœur du réseau et améliorant ainsi la répartition de la charge dans le réseau.

Jun et Sichitiu, dans [JS03] s’intéressent eux aussi à la capacité d’un réseau sans fil multisauts mais dans le contexte un peu particulier où toutes les communications sont à destination d’un nœud central, une passerelle. Ce type de réseaux, appelés *Mesh Networks*, est destiné, par exemple, à fournir un accès à Internet sans fil. Dans ce contexte, un simple calcul montre que, en fonction de la longueur des routes, si l’on souhaite conserver un comportement équitable dans le sens où les nœuds proches de la passerelle doivent pouvoir être en mesure de transmettre autant de trafic que les nœuds lointains, les performances du réseau sont fortement affectées par la présence d’un tel point de concentration des trafics.

Ces différentes études semblent montrer que le moindre paramètre du protocole d’accès au médium peut affecter les performances du protocole d’accès au médium. Ces différentes études sont cependant difficilement comparables puisque les valeurs des paramètres utilisés diffèrent souvent. Certains utilisent les valeurs définies dans la version initiale de la norme IEEE 802.11, d’autres celles de la révision b de cette norme. De plus, les simulateurs utilisés afin de valider les modélisations sont presque aussi nombreux que les modélisations elles-mêmes. Aucun travail n’a, à ce jour, été proposé afin de comparer les différents résultats obtenus à des résultats d’expérimentations. Chacune de ces études améliore tout de même la connaissance du comportement théorique du protocole d’accès au médium. Il semble impossible *a priori* de définir les détails d’un protocole distribué d’accès au médium conduisant aux performances optimales quel que soit le nombre d’émetteurs en concurrence, les profils de trafic, etc. En effet, toutes les études s’accordent pour conclure que l’utilisation du canal radio dépend énormément du nombre de stations en contention et du volume et de la répartition du trafic engendré par les émissions et le routage.

1.4.2 Équité d’accès

En dehors de l’évaluation des performances du protocole d’accès au médium, certains travaux se sont intéressés à l’équité de ce protocole. Dans une configuration simple, un réseau à un saut, le protocole permet à chaque émetteur d’accéder au médium avec la même probabilité. Le protocole fournit donc une certaine équité en terme de nombre de trames émises. Il n’y a aucune garantie cependant sur le volume de données effectivement transmis. En effet, la taille des trames n’entre pas en ligne de compte dans la probabilité d’accès au médium. Par ailleurs, dans un contexte multisauts, des différences dans la topologie telles que l’existence de zones plus denses que d’autres ou encore la présence d’interférences entre différents trafics peuvent avoir un impact important sur la probabilité de transmission avec succès des différents émetteurs.

Un problème d’inégalité dans les chances d’accès au médium lié au comportement des protocoles utilisant un délai de rétention exponentiel (*Binary Exponential Backoff*) a été mis en lumière dès 1994 par Bharghavan *et al.* lors de la conception du protocole d’accès au médium MACAW [BDSZ94]. Dans le cas où les émetteurs en concurrence tirent avant chaque émission un *backoff* dans une fenêtre croissant avec le nombre de collisions subies et dont la taille est remise à zéro lors d’une transmission réussie, les différents émetteurs auront un comportement par rafale. En effet, transmettre une trame avec succès implique que le prochain *backoff* sera tiré aléatoirement dans un intervalle de taille plus faible statistiquement que les émetteurs concurrents.

Ce comportement par rafale est étudié de façon théorique par Li *et al.* dans [LÊG04b]. Cet article modélise le comportement des émetteurs dans une situation de station cachée. Dans ce scénario, l’utilisation de l’échange RTS-CTS ne réduit pas le nombre de collisions mais en limite l’impact. En conséquence, l’accroissement des fenêtres de contention des deux émetteurs concurrents n’est pas limité et l’on constate le même type de comportement par rafale. Les auteurs, au moyen d’une chaîne de Markov en temps discret, confirment l’existence de ce phénomène et déterminent le nombre moyen de trames constituant une rafale et le temps moyen durant lequel un émetteur souffre de famine.

Dans [LÊG04a], les mêmes auteurs étudient la pertinence du mécanisme d’EIFS défini par la norme IEEE 802.11 et déclenché lorsqu’une trame erronée est détectée sur le médium radio. L’article montre que

la valeur choisie pour cette temporisation peut se révéler trop grande ou trop petite selon les scénarios étudiés. Afin de résoudre ce type de problèmes, les auteurs proposent un mécanisme adaptatif de réglage de ce délai d'attente basé sur une mesure du temps d'occupation du médium permettant de déterminer si le signal correspond à une trame de données ou à une trame de signalisation. Ils montrent par simulation que ce procédé permet de rendre l'accès au médium plus équitable dans les scénarios présentés.

Dans [NKG00], Nandagopal et al. proposent une évaluation de l'équité de protocoles d'accès au médium. Cette étude permet d'identifier par simulation quelques situations intrinsèquement peu équitables. Ils ramènent le problème de l'évaluation de l'équité à un problème de maximisation d'une fonction concave et dérivable sous un ensemble de contraintes représentant l'impossibilité pour deux émetteurs voisins de transmettre une trame simultanément avec succès. Le choix de la fonction à maximiser détermine le type d'équité recherchée : équité max-min [BM01, BG87], équité proportionnelle [KMT98], etc. Le premier scénario étudié présente un déséquilibre dans la topologie. Les émetteurs en contention dans les zones du réseau les moins denses obtiendront le débit le plus élevé. La seconde situation oppose l'équité en terme de paquets et l'équité en terme de flux. Deux mobiles ayant la même probabilité d'accès au médium ne constituent pas forcément un scénario équitable lorsque l'un des deux doit retransmettre plus de flux que l'autre. Si cette étude dégage quelques scénarios problématiques en terme d'équité, les paramètres utilisés lors des simulations ne rendent pas compte de phénomènes tels que les interférences au-delà de la zone de communication des mobiles.

Dans [BWK00], Bensaou *et al.* proposent une solution au problème d'équité lié à l'accroissement de la fenêtre de contention lors d'une collision et à sa réinitialisation lors d'une transmission réussie. Ce phénomène tend à provoquer une équité à court terme en favorisant statistiquement les stations venant d'émettre une trame. Ils définissent et étudient un algorithme d'ajustement de la fenêtre de contention basé sur une estimation locale par les nœuds de l'équité de l'accès au médium. Lorsqu'un mobile estime avoir dépassé la proportion du médium qui lui est accordée, il augmente sa fenêtre de contention et *vice-versa*. Cette modification permet, selon les résultats de simulation présentés, d'obtenir une équité au niveau de chaque émetteur aussi bien qu'une équité au niveau de chaque flux. [WB01] étend cet algorithme afin de prendre en compte des tailles de paquets variables et de ne plus présupposer l'utilisation du mécanisme RTS-CTS lors de l'estimation de la proportion du canal utilisé par chaque nœud. Enfin, [FBW02] étudie de façon théorique les performances de cet algorithme et indique qu'il est probablement auto-stabilisant, sans toutefois le démontrer.

Dans ces articles, les auteurs mettent en lumière un scénario simple présentant un problème d'équité dans l'accès au médium représenté en figure 1.12(a). Les lignes représentent le fait que deux mobiles peuvent communiquer entre eux. L'absence d'un tel lien représente une indépendance totale entre les deux nœuds. Dans ce scénario, la paire (E_1 , R_1) ne peut accéder au médium correctement du fait de collisions au niveau de R_1 entre les trames de E_1 et celles de E_2 .

Un second scénario est représenté dans [LÊG04a] et représenté en figure 1.12(b). Dans ce scénario, les lignes en pointillés représentent le fait que les deux mobiles ne peuvent communiquer mais sont en zone de détection de porteuse l'un de l'autre. Ce scénario mettant en jeu des délais d'attente longs (*EIFS*) au niveau de l'émetteur E_1 . En conséquence, cet émetteur patientera en moyenne plus longtemps que E_2 avec lequel il est en concurrence. Ce dernier pourra alors accéder au médium plus souvent, bloquant le mécanisme de détection de porteuse de E_1 et accentuant le phénomène.

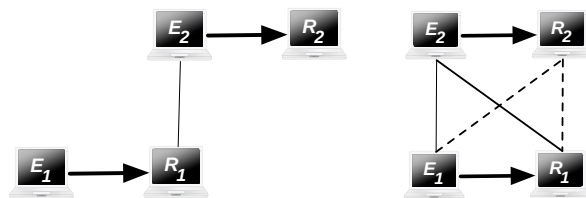


FIG. 1.12 – Scénarios inéquitables

Dans [DGL02], un troisième scénario dans lequel un émetteur souffre d'un problème de performances lié à l'asymétrie de la topologie est étudié par simulation. Ce scénario, qui est celui que nous modélisons dans ce chapitre, sera décrit dans la section suivante.

En résumé, le partage du médium effectué par le protocole IEEE 802.11 semble prévisible et équitable lorsque le réseau n'est pas saturé. Toutefois, l'apparition de congestion modifie radicalement le mode de partage du canal radio et ce de façon imprévisible car dépendante de la situation. Éviter l'apparition de congestions semble donc être un objectif prioritaire si l'on souhaite maîtriser le comportement et les performances du réseau.

1.5 Modélisation d'un scénario particulier

1.5.1 Scénario étudié

Le protocole IEEE 802.11 a été conçu pour fournir un accès au médium équitable aux différentes stations, c'est-à-dire que chaque émetteur a une probabilité égale de transmettre une trame. Or, [Dho02], a mis en lumière par simulation un scénario simple conduisant à une grave inégalité dans l'accès au médium. Dans cette configuration, représentée en figure 1.13, trois couples de mobiles sont en concurrence pour l'accès au médium. Chaque paire est composée d'un émetteur fonctionnant en saturation, c'est-à-dire ayant toujours une trame à émettre, et d'un récepteur proche de l'émetteur. Afin de simplifier la modélisation du problème, les trois émetteurs utiliseront des trames de longueurs égales et constantes.

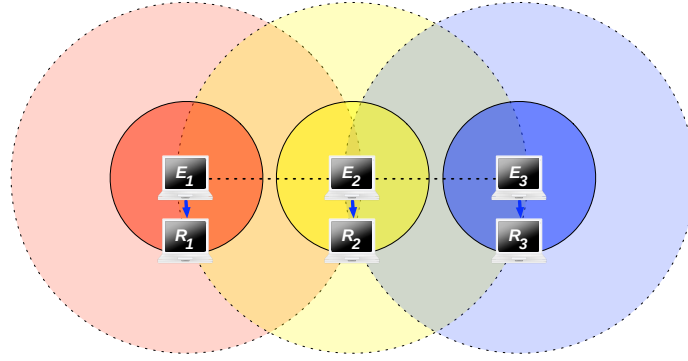


FIG. 1.13 – Scénario considéré : trois couples émetteur-récepteur sont en zone de détection de porteuse de leurs voisins directs et émettent continuellement des trames.

Les émetteurs sont à une distance suffisamment importante l'un de l'autre pour ne pas pouvoir communiquer directement entre eux, mais les trames d'un émetteur empêchent ses voisins immédiats de transmettre au même moment. E_1 et E_2 ne peuvent émettre simultanément, E_2 et E_3 non plus mais E_1 et E_3 n'ont aucune conscience de leur existence réciproque. Par ailleurs, les récepteurs sont suffisamment proches de leurs émetteurs associés pour permettre à deux émetteurs voisins de transmettre une trame simultanément sans provoquer de collision. Cette contrainte permet de ne pas modéliser l'augmentation de la fenêtre de contention due aux collisions. Cette hypothèse n'est pas irréaliste si l'on considère les valeurs du seuil de détection de porteuse et du seuil du rapport signal sur bruit. Si l'on note R la portée de transmission des nœuds, CS le rayon de la zone de détection de porteuse et $SNR_{threshold}$ le rapport signal sur bruit nécessaire au décodage d'une trame, et si l'on considère un affaiblissement en $1/d^\alpha$, cette configuration peut se traduire par :

$$\begin{cases} \forall i \in \{1; 2; 3\}, d(E_i, R_i) \leq R; \\ \forall i \in \{1; 2\}, d(E_i, E_{i+1}) \in]R; CS]; \\ d(E_1, E_3) > CS; \\ \forall i \in \{1; 2; 3\}, \forall j \in \{i-1; i+1\} \cap \{1, 2, 3\}, \left(\frac{d(E_j, R_i)}{d(E_i, R_i)} \right)^\alpha > SNR_{threshold}. \end{cases}$$

Ce scénario combine deux causes d'inégalité. Tout d'abord, la topologie est déséquilibrée. L'émetteur central est en compétition avec deux autres émetteurs alors que chacun des deux autres n'est en compétition qu'avec l'émetteur central. Dans un tel scénario, on pourrait espérer que la paire centrale obtienne un débit diminué de moitié ou de deux tiers par rapport à chacune des deux paires extérieures. De plus, la

distance entre les émetteurs étant supérieure à la portée de communication, lorsqu'un émetteur transmet une trame, le mécanisme d'*EIFS* est activé pour tous ses voisins directs.

Ce scénario représente la configuration la plus simple faisant intervenir un déséquilibre dans la topologie. Le mécanisme d'accès au médium est tel que l'émetteur central désirant accéder au médium devra attendre que son *backoff* atteigne la valeur 0. Or, il ne pourra décrémenter cette valeur que lorsque les deux émetteurs extérieurs font silence simultanément. En effet, l'intervalle de silence *SIFS* marquant les différentes phases de l'émission d'une trame est trop court pour permettre une quelconque décrément du médium. Du point de vue du mécanisme de détection de porteuse, la séquence RTS, CTS, données et acquittement pourra être considérée comme une longue émission de trame. Les deux émetteurs extérieurs, en revanche, sont libres d'accéder au médium pour peu que l'émetteur central fasse silence. Ceci aura pour conséquence que le système peut être séparé en deux sous-systèmes différents : la paire centrale d'une part et les deux paires extérieures de l'autre. Lorsqu'un des deux sous-systèmes est actif (*i.e.* en émission), l'autre est en attente de libération du médium. Il faut noter que dans ce scénario, l'utilisation de l'échange RTS-CTS n'a pas de réel intérêt et ne représente qu'un surcoût. Toutefois, ce mécanisme n'étant pas déclenché dynamiquement en fonction de la situation, il semble pertinent de l'étudier.

Les deux émetteurs extérieurs étant indépendants l'un de l'autre leurs périodes de silence respectives peuvent se recouvrir totalement, partiellement ou pas du tout. Dans ce dernier cas, l'émetteur central ne percevra pas le canal libre. Le recouvrement des deux périodes de silence évoluera en fonction des tirages aléatoires des *backoffs*. La figure 1.14 représente une telle évolution. Entre les dates t_0 et t_1 , les deux émetteurs extérieurs sont silencieux simultanément. Ils terminent la décrément de leur *backoff*, émettent leur trame et tirent une nouvelle valeur. Comme E_1 commence à émettre à la date t_2 avant même que E_3 n'ait terminé la trame précédente à la date t_3 , l'émetteur central ne voit pas le médium se libérer. Le tirage suivant permet un grand recouvrement des périodes de silence des deux paires extérieures et le canal est perçu comme libre par E_2 entre les dates t_4 et t_5 .

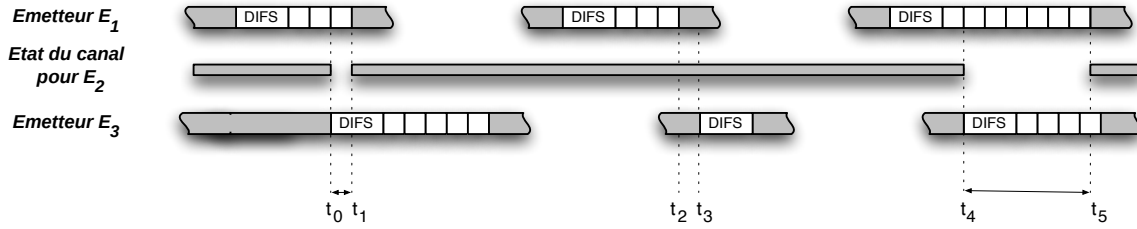


FIG. 1.14 – Désynchronisation des deux émetteurs extérieurs

En revanche, lorsque la paire centrale est active, les deux émetteurs extérieurs sont en attente. La paire extérieure alterne alors périodes d'émission et périodes de contention durant lesquelles les deux paires extérieures ont toujours l'opportunité de décrémenter simultanément leurs *backoffs*. Tôt ou tard l'une des deux paires extérieures terminera la décrément de son *backoff*, émettra une trame bloquant ainsi la paire centrale et laissant tout loisir à l'autre paire extérieure de terminer la décrément de son *backoff*.

Enfin, il faut noter que les émetteurs étant à une distance supérieure à la portée de communication les uns des autres, lorsqu'un sous-système est actif, les émetteurs de l'autre sous-système doivent patienter un temps *EIFS* et non *DIFS* avant de pouvoir décrémenter d'une seule unité tout *backoff*. Ce qui signifie que pour que la paire centrale gagne la contention, il ne suffit pas que les paires extérieures fassent silence simultanément, il est aussi nécessaire que cette période de silence commune soit suffisamment longue.

Ainsi, deux types de situations peuvent survenir. Dans la situation représentée sur la figure 1.15(a), l'émetteur central a l'accès au médium et les deux émetteurs extérieurs patientent. Le système est alors synchrone. À chaque fois que l'émetteur central est en période de contention, les deux émetteurs extérieurs décrémentent leur *backoff* d'un nombre d'unités uniquement dépendant du tirage de la paire centrale, et ce après avoir patienté un temps *EIFS*. La figure 1.15(b) représente la situation lorsque les paires extérieures sont actives. Lorsque leurs périodes de silence se recouvrent suffisamment, la paire centrale peut décrémenter son *backoff* d'un nombre de slots dépendant de l'alignement entre les deux paires extérieures et des tirages des *backoffs* des deux émetteurs extérieurs.

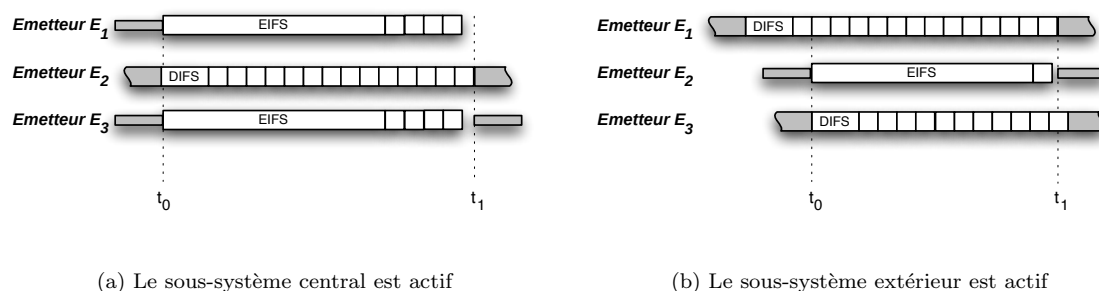


FIG. 1.15 – Situations de contention dans l'accès au médium

Compte tenu de ce fonctionnement du système, afin de déduire la proportion de bande passante que peut espérer obtenir l'émetteur central, nous allons déterminer le nombre de trames émises par chacun des deux sous-systèmes. En fonction de l'outil de modélisation choisi, il conviendra de définir les instants définissant les états du système de manière à minimiser la taille de la modélisation.

1.5.2 Description de la modélisation

Nous avons choisi de modéliser ce système au moyen d'un processus stochastique en temps discret. Nous cherchons à obtenir des informations sur la proportion de paquets envoyés par chacun des différents émetteurs ou encore la taille des rafales. Nous faisons l'hypothèse d'un canal parfait, c'est-à-dire que toute trame émise sera reçue correctement. Cette hypothèse n'est pas irréaliste car le protocole de niveau physique est raisonnablement fiable pour des récepteurs proches de leurs émetteurs associés. D'autre part, nous faisons l'hypothèse que les collisions ne surviennent pas. En conséquence, la séquence d'envoi d'une trame est parfaitement déterminée. Lorsqu'un émetteur arrive à accéder au canal radio, il transmettra une trame avec succès. D'autre part, compte tenu du mode d'accès au médium proposé par IEEE 802.11, le seul phénomène aléatoire pouvant survenir dans cette configuration est le tirage d'un *backoff*. En conséquence, il semble naturel de ne représenter par des états du processus stochastique que les instants correspondant aux périodes de contention. Par ailleurs, une étude basée sur une description de l'état des paires à chaque microseconde serait trop fine et engendrerait une modélisation de taille trop importante. En revanche, s'il est possible de faire correspondre une transition à une transmission de trame, il sera dès lors aisé de déduire de ce modèle les mesures de performances recherchées.

Élection d'une paire de référence

Comme nous l'avons vu auparavant, le système peut se trouver dans deux types de situations distinctes, en fonction du sous-système actif. Lorsque l'émetteur central est actif, une transition du système correspondra naturellement à une émission de trame de cet émetteur. Les états du système correspondront alors aux périodes de contention de la paire centrale. Lorsque les paires extérieures sont actives, compte tenu de la symétrie de la topologie, nous élisons une paire de référence qui définira de façon analogue les états du système. À chaque période de contention de cet émetteur, les transitions possibles ainsi que leurs probabilités associées seront définies par les tirages des *backoffs* des deux paires extérieures et par l'alignement des périodes de silence des deux émetteurs extérieurs. Nous aurons donc à chaque instant une paire de référence élue dynamiquement définissant les états et transitions du système. La figure 1.16 illustre cette élection. Au départ, l'émetteur E_1 est la référence et définit l'état du système à la date t_0 . Compte tenu des tirages de *backoff* des deux paires extérieures, à la date t_1 , la paire centrale parvient à émettre une trame alors qu'il reste 4 unités de *backoff* à E_1 et 2 à E_3 . Elle devient alors paire de référence puisqu'elle est la seule à émettre. Durant la période de contention située entre les dates t_2 et t_3 , l'émetteur E_3 gagne la contention et devient dès lors émetteur de référence. Il le restera jusqu'à ce que E_2 gagne l'accès au médium à nouveau.

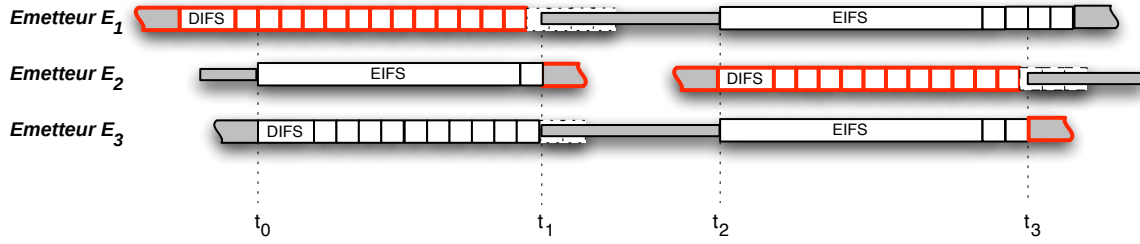


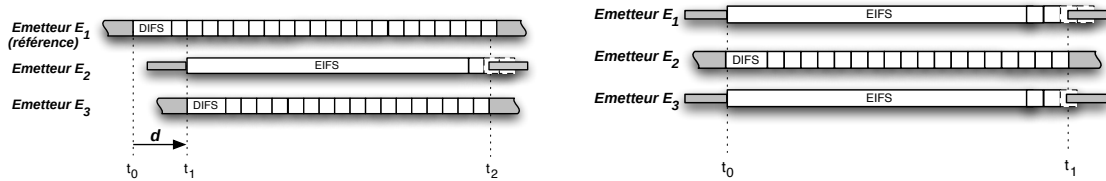
FIG. 1.16 – Élection dynamique de la paire de référence

Représentation des états

Compte tenu de ces observations, nous pouvons maintenant déterminer quelles informations seront nécessaires à la représentation d'un état du système. Un état devra tout d'abord contenir l'identité du sous-système actif. Il devra par ailleurs contenir la valeur des *backoffs* restant des émetteurs en attente puisque cette valeur est conservée tant qu'elle n'est pas consommée. Enfin, lorsque les paires extérieures sont actives, un état doit rendre compte du décalage existant entre les deux émetteurs extérieurs. Nous avons choisi de représenter un état du système par le triplet $(W_e; W_c; d)$ où W_e représente le *backoff* restant de la paire extérieure de référence, W_c le *backoff* restant de la paire centrale et d le décalage entre les deux paires extérieures.

Lorsque les paires extérieures sont actives, W_e vaut 0 et W_c contient le nombre d'unités de *backoff* restant à décrementer à la paire centrale. d contient le nombre de microsecondes entre le début de la période de silence de la paire de référence et le début de la période de silence la plus proche de l'autre paire extérieure. La figure 1.17(a) représente une telle situation. L'état du système à la date t_0 est $(0; 3; d)$. Une des évolutions possibles du système correspond à la situation dans laquelle l'émetteur E_1 tire un *backoff* de 21 unités et l'émetteur E_3 17 unités. Dans ce cas, la paire centrale a l'opportunité de décrementer d'une unité son *backoff* et le décalage évolue de $(17 - 21) \times Slot.time \mu s$. Cette transition aboutira donc à l'état $(0; 2; d - 80)$.

Lorsque la paire centrale émet, c'est W_c qui est nul et W_e contient le minimum des deux *backoffs* restant à décrementer des paires extérieures. Le décalage d contient de manière analogue la différence en microsecondes entre les deux *backoffs* des paires extérieures. Par exemple, la figure 1.17(b) représente le passage de l'état $(3; 0; 20)$ à la date t_0 à l'état $(1; 0; 20)$ à la date t_1 .



(a) Passage de l'état $(0; 3; d)$ à $(0; 2; d - 80)$

(b) Passage de $(3; 0; 20)$ à $(1; 0; 20)$

FIG. 1.17 – Transitions dans le système

Ces trois informations suffisent à décrire le système. À partir d'un état donné, il est possible de considérer toutes les possibilités de tirages de *backoff*, et de déterminer en fonction du résultat de ces tirages les états accessibles.

Caractérisation de l'ensemble d'états

Compte tenu de cette représentation, l'ensemble des états peut être séparé en deux sous-ensembles, comme sur la figure 1.18. Le premier sous-ensemble contient les états $(0; W_c; d)$ et correspond aux cas où

la paire centrale est en attente. Le second sous-ensemble correspond aux états $(W_e; 0; d)$ correspondant aux situations dans lesquelles la paire centrale est active. Des transitions seront présentes à l'intérieur de chaque sous-ensemble et entre les deux sous-ensembles.

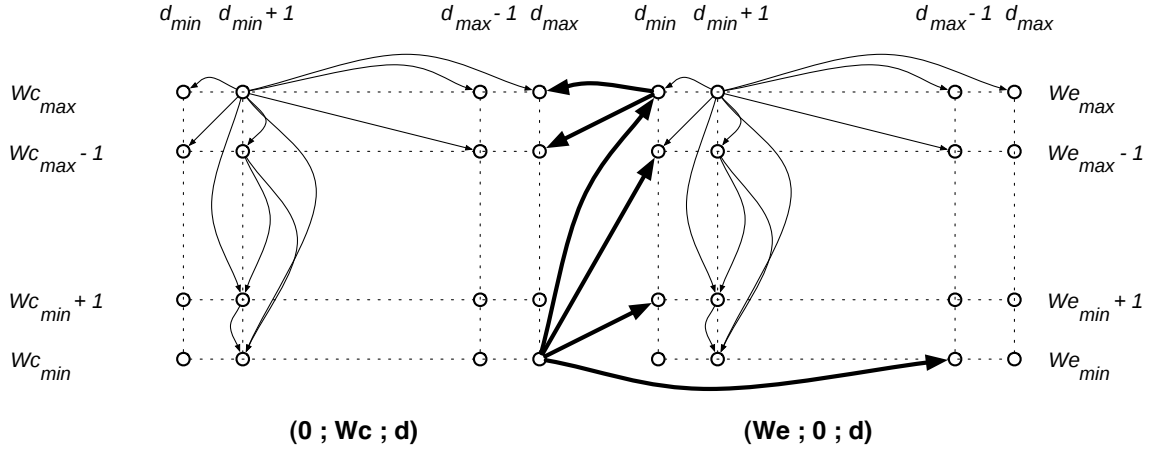


FIG. 1.18 – Séparation de l'ensemble d'états en deux sous-ensembles

Comme les collisions entre deux trames sont négligées dans notre modélisation, l'accroissement exponentiel de la fenêtre de contention est impossible. Les valeurs W_e et W_c sont donc bornées par la valeur CW_{min} définie par la norme IEEE 802.11. D'autre part, W_c et W_e contiennent le nombre d'unités de *backoff* qu'il restait à décrémenter au sous-système actif lorsque celui-ci a perdu l'accès au médium. Or, avant que le sous-système en attente ne puisse décrémenter son *backoff* d'une seule unité, ce dernier a dû patienter un temps *EIFS* durant lequel le sous-système actif a pu décrémenter un nombre d'unités de *backoff* correspondant à un temps $EIFS - DIFS$. En conséquence, W_c et W_e ne pourront excéder $CW_{min} - \lfloor \frac{EIFS - DIFS}{Slot_time} \rfloor - 1$, soit une valeur de 15 en considérant les paramètres définis par la norme IEEE 802.11.

Le décalage lui aussi peut être borné. Premièrement, quand la paire centrale est active, ce décalage contient la différence entre deux valeurs de *backoff*. Il est donc dans ce cas forcément multiple de la durée d'une unité de *backoff*. De plus, chacun des deux *backoffs* résiduels des paires extérieures est borné par $CW_{min} - \lfloor \frac{EIFS - DIFS}{Slot_time} \rfloor - 1$. Comme W_e correspond au minimum des deux *backoffs* des paires extérieures, l'autre *backoff* résiduel sera compris dans l'intervalle $[W_e; CW_{min} - \lfloor \frac{EIFS - DIFS}{Slot_time} \rfloor - 1]$, lui aussi ayant bénéficié d'une décrémentatation préalable d'un nombre d'unités correspondant à la différence entre *EIFS* et *DIFS*. En conséquence, le décalage, pour une valeur de W_e donnée ne pourra pas excéder $CW_{min} - \lfloor \frac{EIFS - DIFS}{Slot_time} \rfloor - 1 - W_e$. En considérant les valeurs définies par le standard IEEE 802.11b, le sous-ensemble d'états correspondant sera composé de 120 états et aura l'allure représentée en figure 1.19.

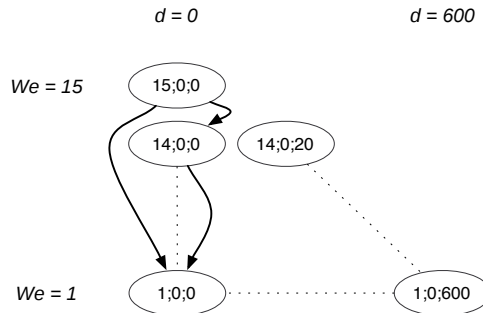


FIG. 1.19 – Sous-ensemble d'états correspondant aux situations où la paire centrale est active

Lorsque les paires extérieures sont actives, le décalage évolue au fur et à mesure des tirages successifs de *backoffs*. Faire évoluer indéfiniment cette valeur de cette façon conduirait à un nombre infini d'états puisque, à chaque étape, il est toujours possible de tirer plusieurs couples de *backoffs* conduisant à une augmentation du décalage. De plus, cette stratégie reviendrait à toujours considérer le décalage entre la n^e trame de la paire de référence et la n^e trame de la seconde paire extérieure. Ce qui rendrait complexe la détection des recouvrements des périodes de silence des deux paires extérieures. Par exemple, dans la situation représentée en figure 1.20, le décalage devient trop important et la période de silence correspondant à la n^e trame de la paire de référence recouvre partiellement la période de silence correspondant à la $n+1^e$ trame de la seconde paire extérieure. De façon analogue, dans la situation représentée en figure 1.21, la période de silence correspondant à la $n+1^e$ trame de la paire de référence recouvre partiellement la période de silence correspondant à la n^e trame de la seconde paire extérieure.

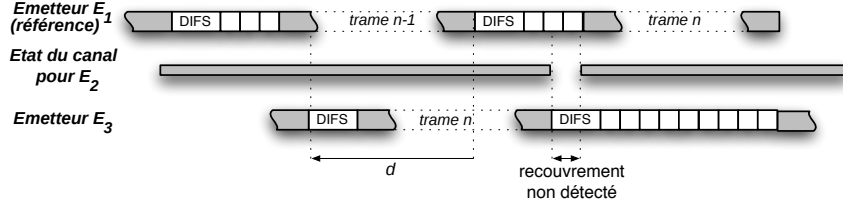


FIG. 1.20 – Un décalage trop petit rend complexe la détection des périodes de silence communes

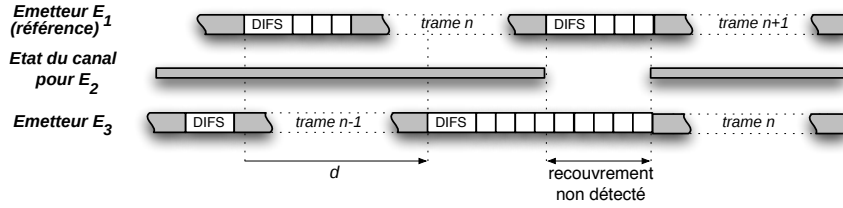


FIG. 1.21 – Un décalage trop grand rend complexe la détection des périodes de silence communes

Afin de ne pas rendre complexe voire impossible la détection des périodes de silence communes, le décalage représentera plutôt le décalage entre la n^e trame de la paire de référence et la trame la plus proche de la seconde paire extérieure. Ce qui signifie que le décalage doit être maintenu dans un intervalle $[d_{min}; d_{max}]$. Afin de résoudre la situation de la figure 1.21, lorsque le décalage devient trop grand, la seconde paire extérieure est en avance par rapport à la paire de référence. Comme seules les émissions de la paire de référence sont importantes pour l'évaluation des performances, nous considérerons que la seconde paire extérieure n'émettra pas de trame durant cette transition et attendra la paire de référence. Cette opération revient à ne considérer qu'un seul tirage de *backoff* pour calculer les probabilités de transition à partir de certains états. En d'autres termes, si l'on note L la durée minimale (avec *backoff* nul) de transmission d'une trame, lorsque le décalage dépasse une valeur $d_{seuil-transition}$, la prochaine émission de trame l'amènera dans l'intervalle $[d - L - CW_{min} \cdot Slot.time; d - L]$ au lieu de $[d - CW_{min} \cdot Slot.time; d + CW_{min} \cdot Slot.time]$. Cette opération, que nous appellerons *translation*, permet en outre de s'assurer que le décalage ne dépasse jamais une valeur $d_{max} = d_{seuil-transition} + CW_{min} \cdot Slot.time$.

De façon analogue, lorsque le décalage devient trop petit, la paire de référence a trop d'avance sur la seconde paire extérieure et devrait patienter. Cependant, une transition doit correspondre à une émission de trame de la paire de référence. Celle-ci ne peut donc en aucun cas ne pas émettre durant une transition. Il aurait été envisageable dans cette situation d'inverser les rôles des deux paires extérieures avant d'effectuer une opération comme précédemment. Nous avons choisi de permettre à la seconde paire extérieure d'émettre deux trames pendant que la paire de référence n'en émet qu'une. Les transitions à partir de certains états dépendront alors de trois tirages de *backoff*, un pour la paire de référence et deux pour l'autre paire extérieure. En d'autres termes, lorsqu'une transition conduit dans un état tel que le décalage est inférieur à un certain seuil $d' < d_{min}$, on ajoute à cette transition une transmission de trame amenant le nouveau décalage dans un intervalle $[d' + L; d' + L + CW_{min} \cdot Slot.time]$.

Les valeurs d_{min} et $d_{seuil-transmission}$ déclenchant une opération de translation doivent être choisies avec soin afin de s'assurer qu'une opération de translation aboutit toujours dans un état admissible. D'autre part, afin de simplifier l'écriture du modèle, il serait préférable de n'effectuer ces opérations dans des intervalles de décalages pour lesquels le recouvrement des deux périodes de silence des paires extérieures ne permet aucune décrémentation du *backoff* de la paire centrale.

Tout d'abord, identifions l'ensemble $[d_{decr-min}; d_{decr-max}]$ des décalages tels que la paire centrale peut décrémenter au moins une unité de *backoff*. Le décalage $d_{decr-max}$ maximum autorisant une décrémentation de *backoff* est tel que :

$$DIFS + CW_{min} \cdot Slot_time - d_{decr-max} = EIFS + Slot_time.$$

En conséquence,

$$d_{decr-max} = DIFS - EIFS + (CW_{min} - 1) \cdot Slot_time.$$

En utilisant les valeurs définies par la norme IEEE 802.11, $d_{decr-max} = 286 \mu s$. Le décalage minimum autorisant une décrémentation est, quant à lui, tel que :

$$DIFS + CW_{min} \cdot Slot_time + d_{decr-min} = EIFS + Slot_time.$$

En conséquence,

$$d_{decr-min} = EIFS - DIFS - (CW_{min} - 1) \cdot Slot_time.$$

En utilisant les valeurs définies par la norme IEEE 802.11, $d_{decr-min} = -286 \mu s$.

Calculons maintenant les décalages provoquant une translation. Compte tenu de la définition de nos opérations de translation, il nous faut déterminer deux valeurs d_{min} et $d_{seuil-transmission}$ telles qu'une translation nous amène dans l'intervalle des décalages admissibles et telles que les intervalles déclenchant une opération de translation et permettant une décrémentation du *backoff* de l'émetteur central sont disjoints.

En d'autres termes :

$$\begin{cases} \forall d \in [d_{seuil-transmission}; d_{max}], \forall k \in [0; CW_{min}], d - L - k \cdot Slot_time \in [d_{min}; d_{max}]; \\ \forall d \in [d_{min} - L - CW_{min} \cdot Slot_time; d_{min}], \forall k \in [0; CW_{min}], d + L + k \cdot Slot_time \in [d_{min}; d_{max}]; \\ d_{min} < d_{decr-min}; \\ d_{seuil-transmission} > d_{decr-max}. \end{cases}$$

En d'autres termes,

$$\begin{cases} d_{seuil-transmission} - L - CW_{min} \cdot Slot_time \geq d_{min}; \\ d_{min} + L + CW_{min} \cdot Slot_time \leq d_{max} + 1; \\ d_{min} < d_{decr-min}; \\ d_{seuil-transmission} > d_{decr-max}. \end{cases}$$

Nous choisirons par conséquent les valeurs suivantes qui vérifient les équations précédentes et qui minimisent la taille de l'intervalle $[d_{min}; d_{max}]$:

$$\begin{cases} d_{min} &= d_{decr-min} - CW_{min} \cdot Slot_time; \\ d_{decr-min} &= EIFS - DIFS - (CW_{min} - 1) \cdot Slot_time; \\ d_{decr-max} &= DIFS - EIFS + (CW_{min} - 1) \cdot Slot_time; \\ d_{seuil-transmission} &= d_{min} + L + CW_{min} \cdot Slot_time; \\ d_{max} &= d_{seuil-transmission} + CW_{min} \cdot Slot_time. \end{cases}$$

Soit, numériquement, en utilisant les valeurs définies par la norme IEEE 802.11b :

$$\begin{cases} d_{min} &= -906; \\ d_{decr-min} &= -286; \\ d_{decr-max} &= 286; \\ d_{seuil-transmission} &= -286 + L; \\ d_{max} &= 314 + L. \end{cases}$$

Une dernière observation devrait nous permettre de réduire de façon considérable dans certains cas le nombre d'états du processus stochastique modélisant ce scénario. En effet, les modifications du décalage

sont le résultat d'une différence entre les *backoffs* des deux émetteurs extérieurs ou d'une des deux opérations de translation décrites précédemment. En conséquence, le décalage ne peut être modifié que d'un multiple du temps d'une unité de *backoff* ou de la durée minimale d'une de trame. Ce qui veut dire que si le plus grand commun diviseur de la longueur minimale d'une trame et de la taille d'une unité de *backoff* n'est pas 1, il n'est pas possible d'atteindre l'état $(0; W_c; d+1)$ à partir de l'état $(0; W_c; d)$ sans passer par un état correspondant à une situation dans laquelle la paire centrale est active. Or, lorsque les paires extérieures gagnent l'accès au médium, le décalage est un multiple de la durée d'une unité de *backoff*.

En résumé, si l'on note G le PGCD de la longueur minimale d'une trame et de la durée d'une unité de *backoff*, à partir d'un état $(0; W_c; d)$ seuls les états $(0; W_c; d+k \cdot G)$, $k \in \mathbb{N}$ (en respectant la condition $d+k \cdot G \in [d_{min}; d_{max}]$) sont accessibles et à partir des états $(W_e; 0; d)$, seuls les états $(0; W_c; k \cdot G)$, $k \in \mathbb{N}$ sont accessibles. Le sous-ensemble d'états correspondant aux états $(0; W_c; k \cdot G)$, $k \in \mathbb{N}$ est donc absorbant. On peut vérifier qu'il s'agit bien du seul sous-ensemble absorbant en examinant la matrice de transitions, ce qui a été réalisé pour chacune des matrices considérées ci-après. En conséquence, il est possible de réduire le processus stochastique à cette sous-ensemble lors du calcul de la probabilité stationnaire. Cette observation indique qu'en choisissant certaines longueurs de trames, il est possible de diviser le nombre d'états par la durée d'une unité de *backoff*, c'est-à-dire 20 en utilisant les paramètres de la norme IEEE 802.11b.

En prenant en considération toutes ces remarques, l'allure du processus en temps discret modélisant ce système est représentée (sans transitions) par la figure 1.22. Il reste maintenant à expliciter les transitions et les probabilités associées.

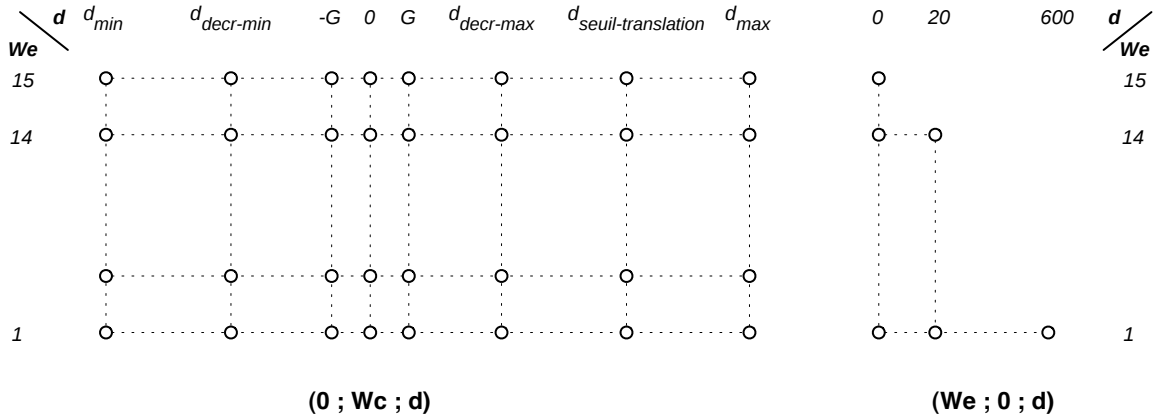


FIG. 1.22 – Allure du processus stochastique sans transitions

Transitions et probabilités

Afin de déterminer les transitions et les probabilités qui leur sont associées, il suffit, dans un état donné, de considérer les différents tirages de *backoff* possibles, ce tirage étant le seul phénomène aléatoire dans le système. À cette fin, il est possible de dégager cinq grandes classes d'états correspondant à cinq situations précises :

- $(W_c = 0; W_e > 0)$: la paire centrale est active, les paires extérieures sont en attente du médium ;
- $(W_e = 0; W_c > 0; d \in [d_{decr-min}; d_{decr-max}])$: les paires extérieures sont actives et leurs périodes de silence sont suffisamment bien alignées pour permettre à la paire centrale une décrémentation de son *backoff* ;
- $(W_e = 0; W_c > 0; d \in [d_{decr-max}; d_{seuil-translation}])$: les paires extérieures sont actives et la paire centrale ne peut décrémenter son *backoff* à cause d'un mauvais alignement. Le décalage est dans un intervalle ne nécessitant aucune translation ;
- $(W_e = 0; W_c > 0; d \in [d_{seuil-translation}; d_{max}])$: les paires extérieures sont actives et la paire centrale ne peut décrémenter son *backoff* à cause d'un mauvais alignement. Le décalage est trop grand et une translation est nécessaire ;

- $(W_e = 0 ; W_c > 0 ; d \in [d_{min} ; d_{decr-min}])$: les paires extérieures sont actives et la paire centrale ne peut décrémenter son *backoff* à cause d'un mauvais alignement. Le décalage est trop petit et une translation est nécessaire.

Ces cinq situations sont décrites plus précisément ci-dessous par ordre de complexité croissante.

Premier cas : la paire centrale est active À partir d'un état $(W_e ; 0 ; d)$, la seule opération possible est la décrémentement conjointe des deux *backoffs* des paires extérieures. En effet, le système est entièrement synchronisé sur les périodes de silence de la paire centrale. Lorsque celle-ci entre en période de contention, les deux paires extérieures ont la possibilité de décrémenter leurs *backoff* pour peu que l'émetteur central patiente un temps supérieur à *EIFS* ; en d'autres termes, si l'émetteur central tire un *backoff* supérieur ou égal à $W_{decr} = \lceil \frac{EIFS + Slot_time - DIFS}{Slot_time} \rceil$. Pour que les émetteurs extérieurs décrémentent leur *backoff* d'exactly k unités, il faut que l'émetteur central tire une valeur de *backoff* W telle que $DIFS + Slot_time \cdot W \in [EIFS + k \cdot Slot_time ; EIFS + (k + 1) \cdot Slot_time[$. Dans le cas où $DIFS + Slot_time \cdot W \geq EIFS + W_e \cdot Slot_time$, les paires extérieures gagnent l'accès au médium. Dans ce cas, le décalage est conservé et il reste à l'émetteur central $W - W_{decr} - W_e$ unités de *backoff* à décrémenter. Chaque tirage de *backoff* de l'émetteur central permettant aux émetteurs extérieurs de décrémenter une unité mène dans un état distinct. Les probabilités de transition sont explicitées ci-dessous et représentées (en utilisant les valeurs de paramètres de la norme IEEE 802.11b) en figure 1.23.

$$\begin{cases} P((W_e ; 0 ; d) \rightarrow (W_e ; 0 ; d)) = \frac{W_{decr}}{CW_{min} + 1} ; \\ \forall k \in [1 ; W_e], P((W_e ; 0 ; d) \rightarrow (W_e - k ; 0 ; d)) = \frac{1}{CW_{min} + 1} ; \\ \forall W_c \in [1 ; CW_{min} + 1 - \lceil \frac{EIFS - DIFS}{Slot_time} \rceil - W_e], P((W_e ; 0 ; d) \rightarrow (0 ; W_c ; d)) = \frac{1}{CW_{min} + 1}. \end{cases}$$

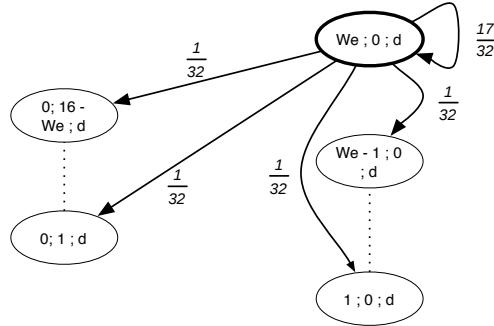


FIG. 1.23 – Transitions possibles à partir d'un état $(W_e ; 0 ; d)$

Deuxième cas : la paire centrale est en attente et ne peut pas décrémenter son *backoff*

Lorsque le décalage entre les deux périodes de silence des paires extérieures est compris strictement entre les valeurs $d_{decr-max}$ et $d_{seuil-translation}$, la seule opération possible est une modification du décalage. Cette modification est due à la différence entre les *backoffs* tirés par les deux paires extérieures. À partir d'un état $(0 ; W_c ; d)$, le tirage par la paire extérieure de référence d'une valeur W_1 et par la seconde paire d'une valeur W_2 conduit dans l'état $(0 ; W_c ; d + (W_1 - W_2) \cdot Slot_time)$. Les *backoffs* étant tirés de manière uniforme dans l'intervalle $[0 ; CW_{min}]$, les transitions possibles sont explicitées ci-dessous et représentées en utilisant les valeurs de paramètres de la norme IEEE 802.11b en figure 1.24 :

$$\forall k \in [-CW_{min} ; CW_{min}], P((0 ; W_c ; d) \rightarrow (0 ; W_c ; d + Slot_time \cdot k)) = \frac{CW_{min} - |k|}{(CW_{min} + 1)^2}.$$

Troisième cas : la paire centrale est en attente et le décalage est trop grand et nécessite une translation Dans ce cas, le décalage a dépassé la valeur $d_{seuil-translation}$. Seule la paire extérieure de

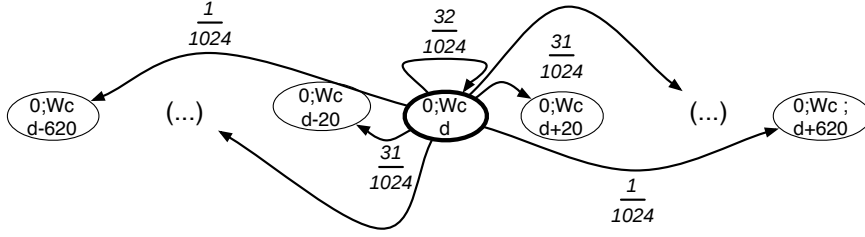


FIG. 1.24 – Transitions possibles à partir d'un état $(0; W_c; d); d \in]d_{decr-max}; d_{seuil-translation}[$

référence est autorisée à émettre une trame. Un seul *backoff* définit les probabilités des transition. La figure 1.25 représente le modèle modifié suite à l'opération de translation décrite plus haut. Si l'on note L la durée minimale de transmission d'une trame, surcoût du protocole inclus, les transitions possibles sont alors :

$$\forall k \in [0; CW_{min}], P((0; W_c; d) \rightarrow (0; W_c; d - L - k \cdot Slot_time)) = \frac{1}{CW_{min} + 1}.$$

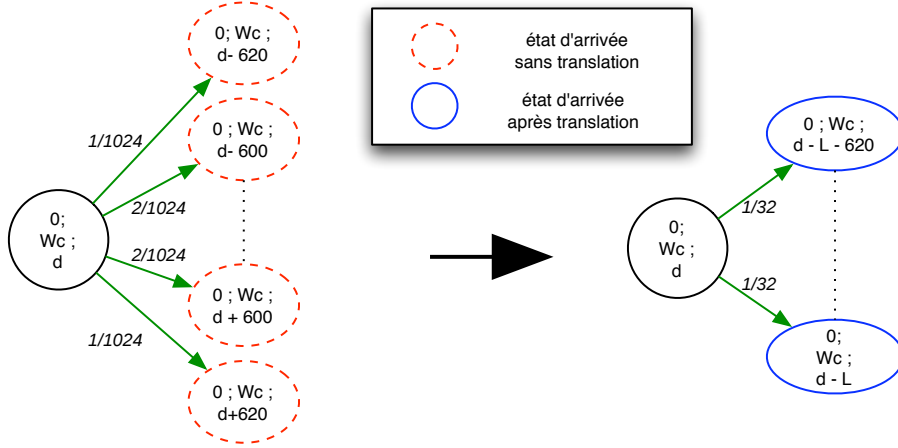


FIG. 1.25 – Modifications du modèle afin de limiter l'accroissement du décalage

Quatrième cas : la paire centrale est en attente et le décalage est trop petit et nécessite une translation Dans l'intervalle de décalages $[d_{min}; d_{decr-min}[$, une modification trop importante du décalage peut conduire hors de l'intervalle $[d_{min}; d_{max}]$. Toutes les transitions incriminées seront remplacées de la manière décrite précédemment alors que les autres seront tirées normalement. La figure 1.26 représente ces modifications. Si l'on note $A = \lceil \frac{d_{min}-1-d}{Slot_time} \rceil$ la valeur seuil telle que $d - Slot_time \cdot A < d_{min}$ toute modification du décalage supérieure ou égale à $A + 1$ unités conduit dans un état ne nécessitant pas de translation. Dans le cas contraire, les états tels que le décalage appartient à l'intervalle $[d - CW_{min} \cdot Slot_time; d + Slot_time \cdot A]$ seront remplacés par l'ensemble des états $d - 620 + L + k \cdot Slot_time; k \in [0; A + CW_{min}]$. La seconde paire extérieure émettant une trame supplémentaire durant cette transition, il faut considérer un tirage de *backoff* supplémentaire.

Il est possible d'exprimer sous forme matricielle la relation liant les probabilités de transition avant translation et les probabilités admissibles de la façon suivante :

$$\begin{pmatrix} P((0; W_c; d) \rightarrow (0; W_c; d + L - CW_{min} \cdot Slot_time)) \\ \vdots \\ P((0; W_c; d) \rightarrow (0; W_c; d + L + (CW_{min} - A) \cdot Slot_time)) \end{pmatrix} = \\
\frac{1}{CW_{min} + 1} \cdot \begin{pmatrix} 1 & 0 & \dots & 0 \\ \vdots & 1 & \ddots & \vdots \\ 1 & & \ddots & 0 \\ 0 & 1 & & 1 \\ \vdots & \ddots & \ddots & \vdots \\ 0 & \dots & 0 & 1 \end{pmatrix} \times \begin{pmatrix} P((0; W_c; d) \rightarrow (0; W_c; d - CW_{min} \cdot Slot_time)) \\ \vdots \\ P((0; W_c; d) \rightarrow (0; W_c; d - A \cdot Slot_time)) \end{pmatrix}.$$

Les autres probabilités étant calculées comme dans le deuxième cas :

$$\forall k \in [A + 1; CW_{min}], P((0; W_c; d) \rightarrow (0; W_c; d + k \cdot Slot_time)) = \frac{CW_{min} - |k|}{(CW_{min} + 1)^2}.$$

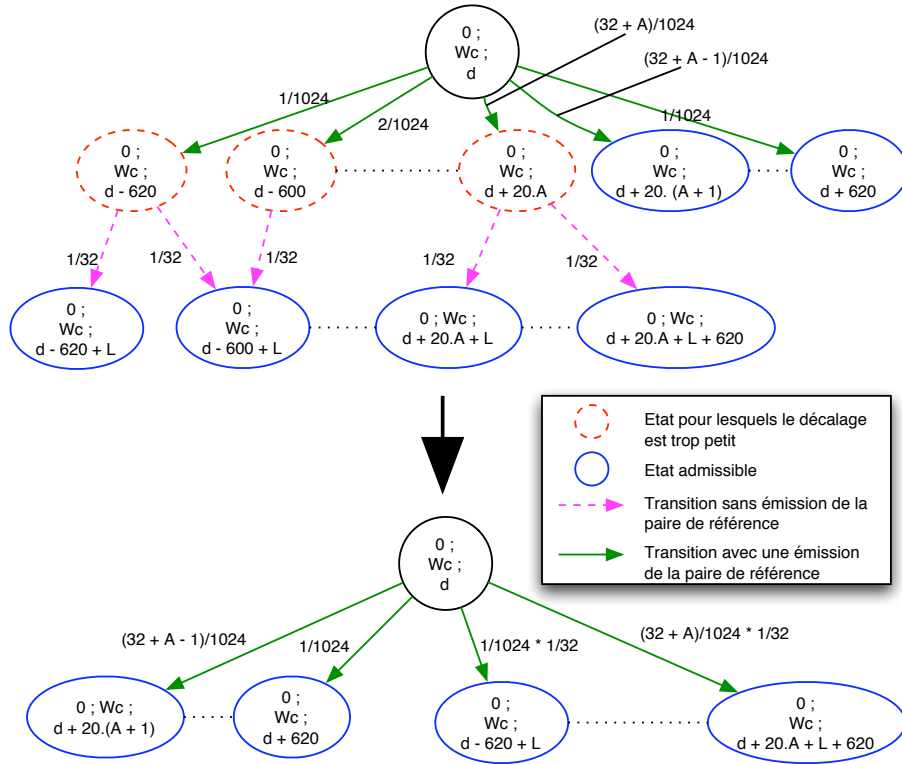


FIG. 1.26 – Modifications du modèle afin de limiter la décroissance du décalage

Cinquième cas : la paire centrale est en attente et peut décrémenter son *backoff* Cette situation correspondant à des décalages appartenant à l'intervalle $[d_{decr-min}; d_{decr-max}]$ est la plus complexe. En effet, c'est à partir de ces états, et uniquement de ceux-ci que la paire centrale peut décrémenter son *backoff* et éventuellement gagner l'accès au médium. Examinons les transitions possibles à partir d'un état $(0; W_c; d)$. Pour un couple $(W_1; W_2)$ de *backoffs* tirés par la paire de référence et la seconde paire extérieure, la période de silence commune correspondra à l'intersection des intervalles

$[0; DIFS + Slot_time \cdot W_1]$ et $[d; d + DIFS + Slot_time \cdot W_2]$. Elle aura une durée de :

$$\begin{aligned} T_{recouvrement} &= \min(DIFS + Slot_time \cdot W_1; d + DIFS + Slot_time \cdot W_2) - \max(0; d) \\ &= DIFS + \min(Slot_time \cdot W_1; d + Slot_time \cdot W_2) - \max(0; d). \end{aligned}$$

Si cette durée est inférieure à $EIFS + Slot_time$, seule une modification du décalage d'une valeur $Slot_time \cdot (W_1 - W_2)$ est possible.

Si cette durée est dans l'intervalle $[EIFS + Slot_time; EIFS + W_c \cdot Slot_time]$, la paire centrale décrémente son *backoff* sans parvenir à émettre. Dans ce cas, le système passe de l'état $(0; W_c; d)$ à l'état $(0; W_c - \lfloor \frac{T_{recouvrement} - EIFS}{Slot_time} \rfloor; d + Slot_time \cdot (W_1 - W_2))$. Chaque couple de *backoffs* tirés par les paires extérieures conduisant dans l'un de ces états est le seul à amener à cet état, la probabilité associée à chacune de ces transitions est $1/(CW_{min} + 1)$.

Preuve. Supposons que deux couples $(W_1; W_2)$ et $(W'_1; W'_2)$ conduisent dans le même état, cela signifie que :

$$\begin{cases} W_1 - W_2 = W'_1 - W'_2; \\ \lceil \frac{T_{recouvrement}}{Slot_time} \rceil = \lceil \frac{T'_{recouvrement}}{Slot_time} \rceil. \end{cases}$$

De la première condition, on peut déduire que $W_1 - W'_1 = W_2 - W'_2$. En conséquence :

$$\begin{aligned} \min(Slot_time \cdot W_1; d + Slot_time \cdot W_2) &= Slot_time \cdot W_1 \\ \Leftrightarrow \min(Slot_time \cdot W'_1; d + Slot_time \cdot W'_2) &= Slot_time \cdot W'_1. \end{aligned}$$

Or, l'état initial est le même pour les deux transitions. On peut vérifier que $\exists k \in \mathbb{Z}, T_{recouvrement} = T'_{recouvrement} + k \cdot Slot_time$. Le système ci-dessus est donc équivalent à :

$$\begin{cases} W_1 - W_2 = W'_1 - W'_2; \\ \lceil \frac{T_{recouvrement}}{Slot_time} \rceil = \lceil \frac{T'_{recouvrement}}{Slot_time} \rceil + k \cdot Slot_time. \end{cases}$$

k est nécessairement nul et $T_{recouvrement} = T'_{recouvrement}$. L'état initial étant le même pour les deux transitions, on a nécessairement $Slot_time \cdot W_1 = Slot_time \cdot W'_1$ ou $d + Slot_time \cdot W_2 = d + Slot_time \cdot W'_2$. En d'autres termes, $W_1 = W'_1$ ou $W_2 = W'_2$ et $W_1 - W_2 = W'_1 - W'_2$. Donc, nécessairement $W_1 = W'_1$ et $W_2 = W'_2$. ■

Enfin, si la durée de recouvrement des deux périodes de silence est supérieure ou égale à $EIFS + W_c \cdot Slot_time$, la paire centrale gagne l'accès au médium. Supposons que l'on se place dans la situation où la paire extérieure de référence a tiré un *backoff* de W_1 unités et la seconde paire extérieure de W_2 unités.

Dans le cas où le décalage aurait été positif, $d \geq 0$, il restera à l'ancienne paire de référence un nombre d'unités de *backoff* non consommées égal à :

$$W'_1 = \lceil \frac{DIFS - EIFS - d}{Slot_time} \rceil + W_1 - W_c.$$

La seconde paire extérieure, quant à elle aura un *backoff* résiduel égal à :

$$W'_2 = \lceil \frac{DIFS - EIFS}{Slot_time} \rceil + W_2 - W_c.$$

Dans le cas où le décalage était négatif, $d < 0$, la situation est inversée et les *backoff* résiduels seront :

$$W'_1 = \lceil \frac{DIFS - EIFS}{Slot_time} \rceil + W_1 - W_c;$$

$$W'_2 = \lceil \frac{DIFS - EIFS + d}{Slot_time} \rceil + W_2 - W_c.$$

Dans tous les cas, l'état d'arrivée sera $(\min(W'_1, W'_2); 0; |W'_1 - W'_2|)$.

Propriétés du processus stochastique La modélisation décrite ci-dessus vérifie les propriétés usuelles des chaîne de Markov en temps discret. En effet, le système est sans mémoire puisque la probabilité d'être dans un état donné ne dépend que de l'état précédent. La chaîne est homogène, apériodique et irréductible, ce qu'il est possible de vérifier, et a été vérifié, sur chacune des matrices de transitions générées. En conséquence, le vecteur des probabilités stationnaires existe et est unique et indépendant de l'état initial du système.

1.5.3 Résultats

La chaîne de Markov décrite précédemment a été résolue en utilisant la bibliothèque d'algèbre linéaire MUMPS [ADLK01] pour différentes tailles de trames (de 700 octets à 1500 octets de charge utile, soit de 5,6 *kbit* à 12 *kbit*), différents débits (11 *Mbit/s* qui est le débit maximal de 802.11b et 2 *Mbit/s* qui est le débit maximal de 802.11 dans sa première version ainsi que le débit des trames émises en diffusion) et en activant ou non le mécanisme RTS-CTS. En fonction de ces différents paramètres, la taille de la chaîne de Markov varie de 1800 états et 155000 transitions à 4500 états et 455000 transitions si l'on ne considère que des longueurs de trames telles que la durée minimale de transmission est un multiple de la taille d'une unité de *backoff*. Quand le PGCD de ces deux valeurs est plus petit, le nombre d'états est de l'ordre de 100000.

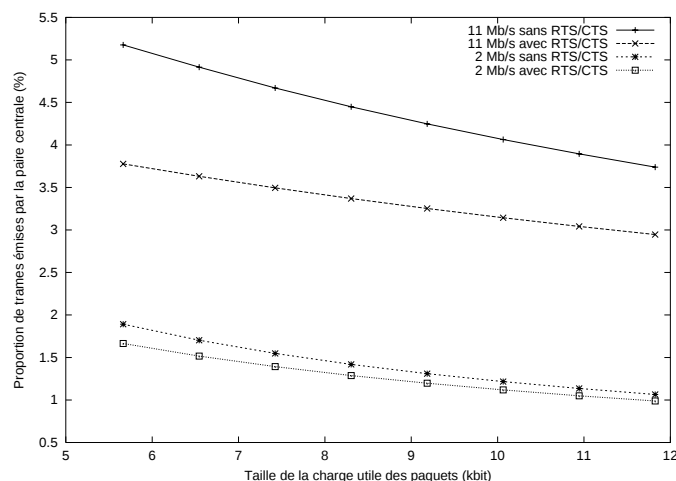


FIG. 1.27 – Proportion de trames émises par la paire centrale en fonction de la taille de trame

La figure 1.27 représente le pourcentage de trames que la paire centrale peut espérer émettre. Ce pourcentage varie de 5,2 % pour de petites trames transmises à 11 *Mbit/s* sans échange RTS-CTS à 1 % pour de grandes trames émises à 2 *Mbit/s* avec échange RTS-CTS. Ces chiffres sont bien en deçà de ce qui pourrait être attendu d'un protocole d'accès au médium équitable duquel la paire centrale pourrait espérer obtenir au moins 33 % de la capacité du médium, selon la définition de l'équité adoptée.

Ces résultats montrent en outre que la proportion de trames émises par la paire centrale décroît quand la taille de ces trames augmente, quand le surcoût du protocole augmente ou quand le débit binaire diminue. En effet, modifier l'un de ces paramètres signifie que la durée de l'échange RTS-CTS-données-acquittement sera allongée. Dans ce cas, la probabilité que les deux périodes de silence des paires extérieures soient bien alignées diminue puisque la taille de ces périodes de silence reste constante.

Comme dans notre configuration un échange RTS-CTS-données-acquittement ne peut être interrompu, il est possible de représenter le pourcentage de trames émises par la paire centrale en fonction de la durée en microsecondes d'un tel échange. Cette comparaison est représentée en figure 1.28. On peut remarquer une absence de points de mesure dans l'intervalle entre 2000 μs et 3500 μs . Les bornes de cet intervalle correspondent au temps nécessaire à l'envoi d'une taille de longueur maximale, avec échange RTS-CTS à un débit de 11 *Mbit/s* et au temps nécessaire à l'envoi d'une trame de la longueur mini-

male considérée, sans échange RTS-CTS à un débit de 2 *Mbit/s*. Ces résultats présentent une certaine régularité et peuvent être interpolés par la fonction $2365,37 \cdot \text{Duree_trame}^{-0,8736}$ avec un coefficient de détermination égal à 99,83 %.

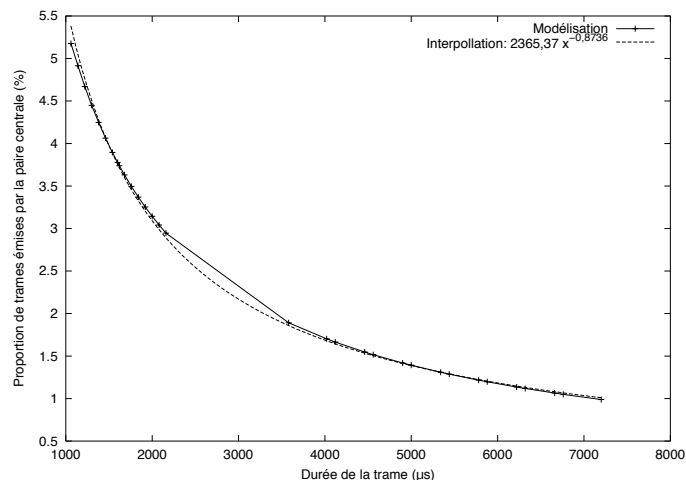


FIG. 1.28 – Proportion de trames émises par la paire centrale en fonction de la durée de transmission d’une trame

Afin de valider cette modélisation, les résultats obtenus ont été comparés avec les résultats de simulations réalisées en utilisant le simulateur NS-2 dans sa version 2.27. Les résultats de la modélisation sont comparés en figures 1.29 et 1.30 à la moyenne des performances obtenues sur 30 simulations réalisées en utilisant des graines différentes pour alimenter le générateur de nombres aléatoires.

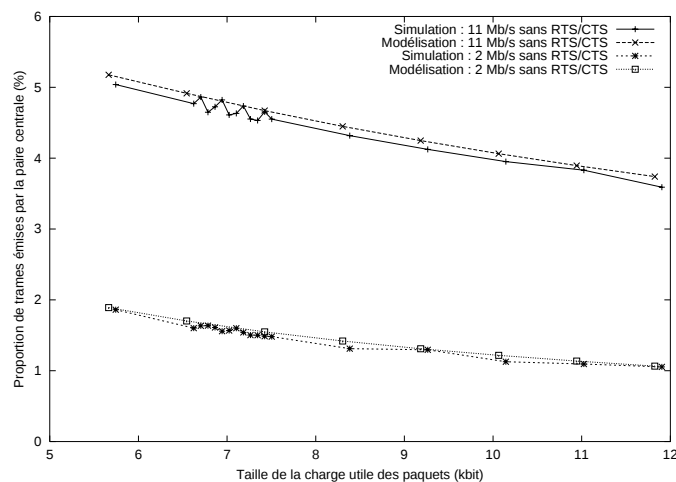


FIG. 1.29 – Comparaisons entre modélisation et simulations – sans échange RTS-CTS

Les premières simulations réalisées donnaient des résultats en deçà des résultats obtenus par modélisation. Cette différence était due à la manière dont NS-2 traitait l’attente *EIFS*. En effet, un mobile patientait un temps *DIFS* en plus d’un temps *EIFS* avant de commencer à décrémenter son *backoff* alors que la norme IEEE 802.11 spécifie, page 93 que *All backoff slots occur following a DIFS period during*

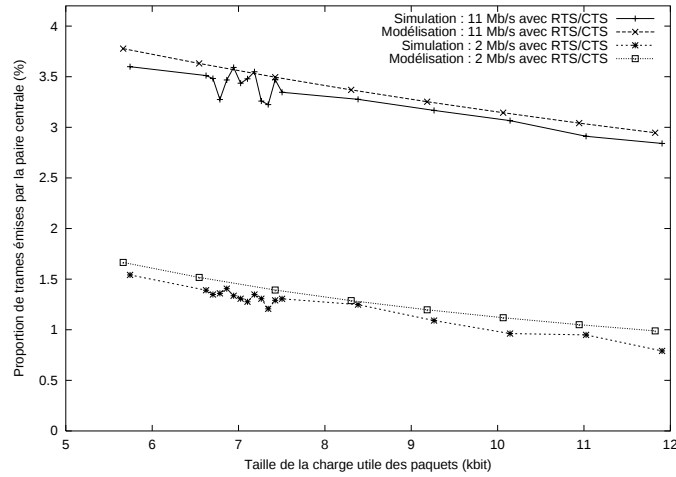


FIG. 1.30 – Comparaisons entre modélisation et simulations – avec échange RTS-CTS

which the medium is determined to be idle for the duration of the DIFS period, or following an EIFS period during which the medium is determined to be idle for the duration of the EIFS period following detection of a frame that was not received correctly. Toute décrémentation du *backoff* est effectuée après une période DIFS durant laquelle le médium doit être perçu libre, ou après une période EIFS durant laquelle le médium doit être perçu libre si la trame précédente n'avait pas été reçue correctement. Les résultats obtenus après correction du simulateur sont très proches des résultats obtenus par la modélisation et la différence entre deux valeurs excède rarement 5 %. Cette différence croît alors que le pourcentage de trames émises par la paire centrale décroît. Elle atteint 10 % pour une transmission de trames de 12 kbit à 2 Mbit/s en utilisant l'échange RTS-CTS.

Les résultats de modélisation ont été obtenus pour des tailles de trames telles que la durée de transmission d'une trame est un multiple de la durée d'une unité de *backoff* afin de minimiser le nombre d'états de la chaîne. Les simulations, quant à elles ont été réalisées avec de nombreuses tailles de trames. On peut observer certaines irrégularités dans les résultats obtenus pour des tailles de trames dont le PGCD est plus faible. Ceci signifie que ce PGCD a une influence sur la synchronisation des émetteurs du fait de la synchronisation occasionnelle des émetteurs extérieurs lorsque la paire centrale est active.

Une expérimentation réelle de ce scénario en intérieur avait été réalisée dont les résultats sont présentés dans [Dho02]. Les performances obtenues sont en réalité un peu plus équitables vis-à-vis de la paire centrale comme le montre la figure 1.31. Dans cette expérience, chaque émetteur envoyait des paquets de 1000 octets (8 kbit) à son récepteur associé de manière à saturer le médium radio dont la capacité avait été limitée à 2 Mbit/s. Les différentes courbes représentent les débits respectifs de chacune des trois paires de mobiles. Hors de l'intervalle situé entre les secondes 800 et 870, un signal extérieur provenant d'une station de base a perturbé l'expérience, handicapant la seconde paire extérieure et favorisant ainsi la paire centrale. À l'intérieur de cet intervalle, le pourcentage de bande passante obtenu par la paire centrale avoisine les 7 % du débit obtenu par les deux autres paires. Ce résultat est légèrement supérieur au débit déduit de la modélisation de celui mesuré lors des simulations. Ceci peut s'expliquer par les grandes variations dans les niveaux de signal. En effet, des mesures de signal réalisées à un emplacement constant indiquent un écart-type du niveau de signal de l'ordre de 3 dB. Par conséquent, en limite de zone de couverture, l'émission d'une trame par un mobile ne déclenchera pas systématiquement le mécanisme de détection de porteuse des émetteurs potentiels voisins. Toutefois, le problème d'équité souligné dans ce chapitre est bien présent.

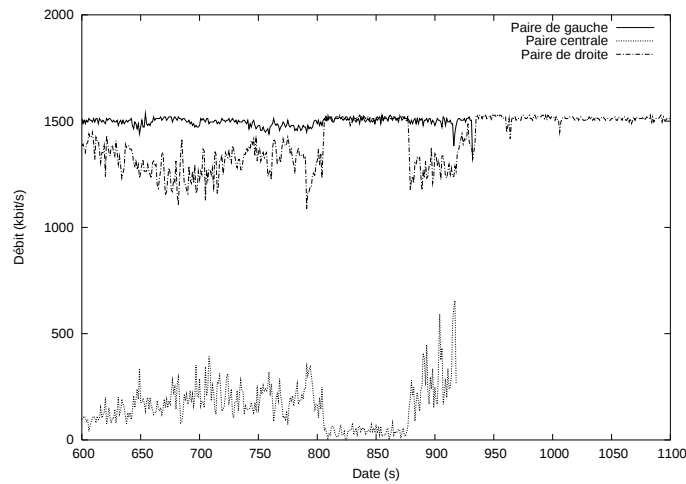


FIG. 1.31 – Débit mesuré en environnement réel dans le scénario des trois paires (source : Dominique Dhoutaut [Dho02])

1.5.4 Influence d’une différence dans les tailles de trames

Utiliser des tailles de trames différentes semble améliorer l’équité du protocole d’accès au médium dans cette situation. Cependant, cette opération a un coût puisque le surcoût lié au protocole d’accès au médium est, lui, constant. Par exemple, sur un médium de bande passante 11 Mbit/s sans mécanisme RTS-CTS, la transmission de trames avec une charge utile de 1000 octets (8 kbit) résulte en un débit au niveau application d’environ $4,52\text{ Mbit/s}$. Avec des trames contenant une charge utile de 500 octets (4 kbit), ce débit n’est plus que de $2,85\text{ Mbit/s}$. Les figures 1.33 et 1.32 représentent le débit utile au niveau application obtenu par la paire centrale et par une paire extérieure en fonction de la taille de trame. Ces courbes dérivent directement des proportions d’accès au médium déterminées par la modélisation. On constate que réduire la taille de trame conduit à une meilleure équité en terme de nombre de paquets transmis mais à des résultats similaires en terme de débit obtenu par la paire centrale. La réduction de la taille des trames ne conduit donc qu’à une perte de performance globale du réseau. Ceci confirme le résultat de simulation présenté par Wang et Bensaou dans [WB01].

Si diminuer la taille des trames des émetteurs consommant une proportion importante de la bande passante ne conduit qu’à une perte de performances au niveau global, il est peut-être possible d’améliorer l’équité du réseau à moindre coût en augmentant la taille des trames de l’émetteur lésé. La figure 1.34 représente le débit au niveau application obtenu par simulation pour la paire centrale en fonction de la taille de ses trames lorsque les paires extérieures utilisent une longueur de trame fixe de 938 octets ($7,5\text{ kbit}$). Cette valeur correspond à l’une des tailles de trames considérées lors de la modélisation et située au centre de l’intervalle de mesure. Dans ces différentes situations, la probabilité d’accès de la paire centrale reste constante. En effet, cette probabilité ne dépend que de la taille des trames des paires extérieures puisque la taille de trame de la paire centrale n’aura qu’une influence sur le temps passé dans un état mais non sur les probabilités de transition, les paires extérieures étant synchronisées lorsque la paire centrale transmet. Seule la taille des trames des paires extérieures a une influence sur les bornes du décalage, et donc sur l’ensemble d’états, et sur les probabilités de transition.

Ces résultats indiquent que le débit utile de la paire centrale peut être amélioré d’environ 50 %. Parallèlement, le débit de chacune des deux paires extérieures décroît d’environ 4 %. Une solution partielle à ce type de problème d’équité pourrait résider dans une politique locale visant à pénaliser les émetteurs monopolisant le canal radio en les forçant à diminuer la taille maximale des paquets émis (MTU — *Maximum Transfer Unit*), provoquant ainsi une fragmentation des segments au niveau TCP. À l’inverse, les nœuds s’estimant lésés pourraient augmenter progressivement cette valeur. Toute la difficulté dans

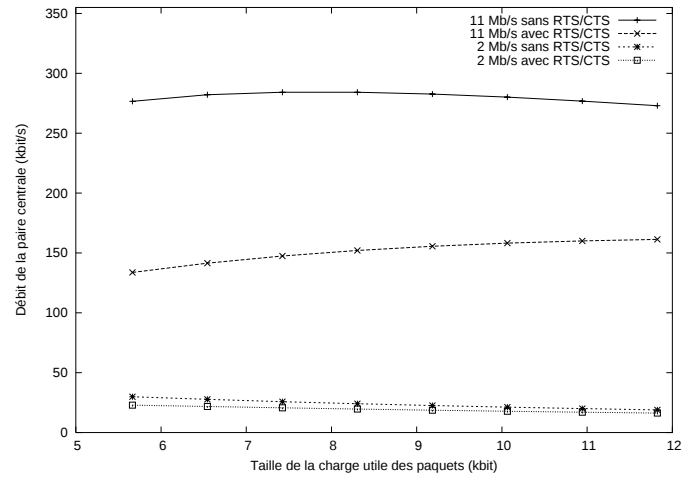


FIG. 1.32 – Débit au niveau application (modélisation) – paire centrale

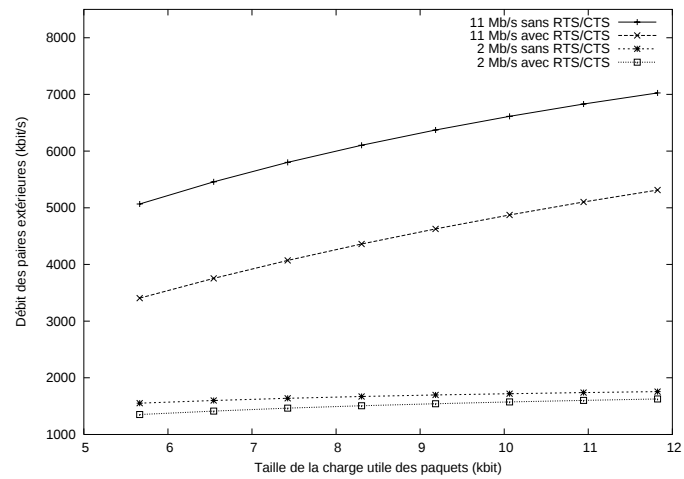


FIG. 1.33 – Débit au niveau application (modélisation) – paires extérieures

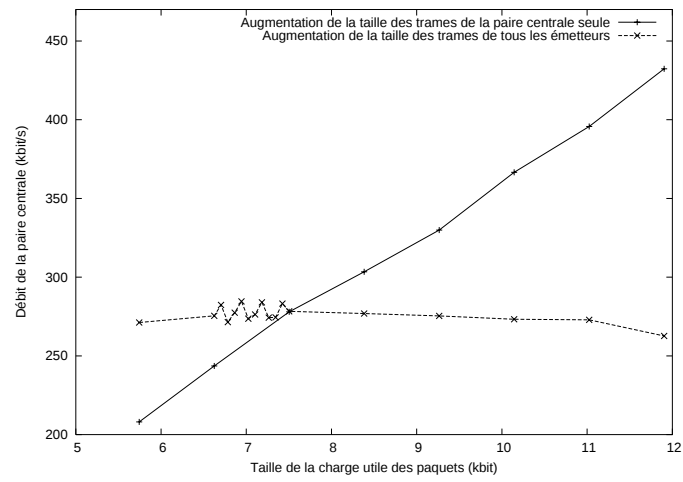


FIG. 1.34 – Débit au niveau application lorsque la taille de trame des paires extérieures est fixe – 938 octets (7,5 kbit) (simulation) – paire centrale

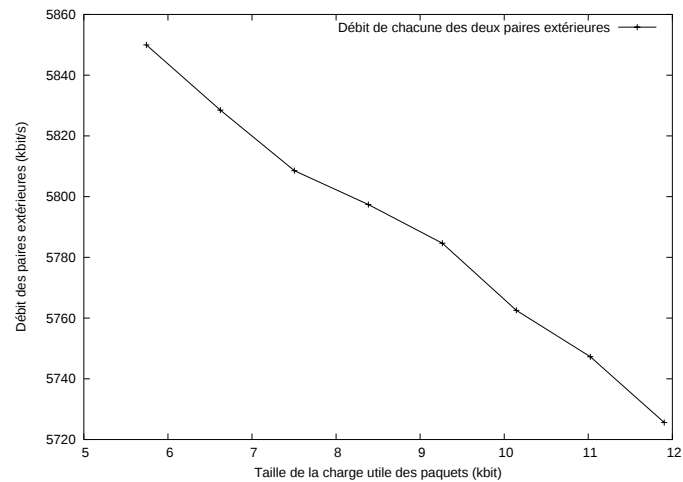


FIG. 1.35 – Débit au niveau application lorsque la taille de trame des paires extérieures est fixe – 938 octets (7,5 kbit) (simulation) – paires extérieures

cette approche réside dans la détermination de seuils déclenchant ces variations afin de ne pas introduire un coût important en terme de performance.

1.5.5 Influence de l'EIFS

Ce scénario mêle deux phénomènes différents. Si l'asynchronisme entre les deux paires extérieures diminue la probabilité d'accès au médium de l'émetteur central, l'utilisation de l'EIFS accentue ce phénomène en allongeant la durée nécessaire avant décrémentation du *backoff*. C'est pourquoi la modélisation précédente a été adaptée afin d'évaluer les performances sans utilisation de ce mécanisme. La figure 1.36 compare les résultats obtenus précédemment en utilisant l'EIFS avec les résultats obtenus en remplaçant EIFS par un simple DIFS. Ces résultats ont une allure similaire aux précédents. La courbe peut être interpolée par la fonction $5447,26 \cdot \text{Durée_trame}^{-0,8034}$ avec un coefficient de détermination valant 99,73%. Remplacer l'EIFS par un simple DIFS résulte, dans cette situation, en une amélioration de l'équité au niveau paquet d'un facteur 2,3. Les simulations réalisées dans cette configuration confirment les résultats de modélisation et sont présentées en figure 1.37.

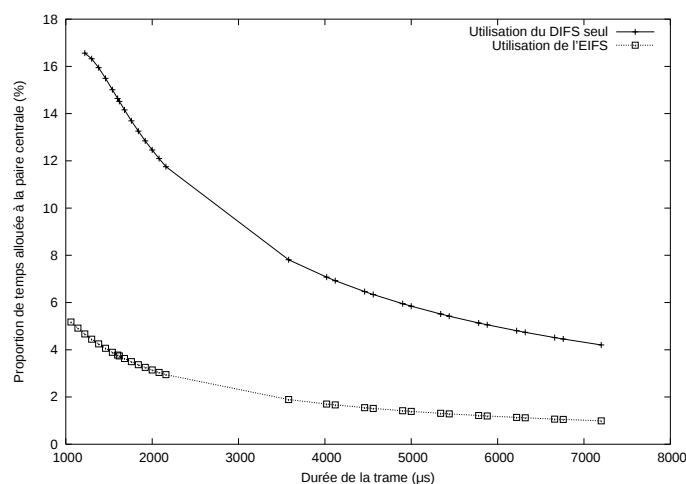


FIG. 1.36 – Proportion de trames émises par la paire centrale sans EIFS – gain en équité en nombre de trames

L'équité en terme de nombre de trames est sensiblement améliorée par cette modification dans cette configuration. Les figures 1.38 et 1.39 présentent ces résultats en terme de débit utile. Les résultats sont semblables. Si le débit utile de chacune des paires extérieures est diminué de 10%, le débit utile de la paire centrale est multiplié par 3.

Une conclusion de ces résultats pourrait être que l'EIFS accroît de façon significative les inégalités d'accès au médium. Cependant, le mécanisme d'EIFS n'est pas utile dans cette situation et il est normal qu'il représente un surcoût inutile. Afin d'évaluer l'intérêt de ce mécanisme, considérons le scénario de la figure 1.40. Dans cette situation, l'EIFS est utile afin d'éviter que les trames émises par le nœud *B* n'entrent en collision au niveau de *C* avec les acquittements émis par *D*. La figure 1.41 compare par simulation les débits globaux de ce réseau en utilisant des trames de 1000 octets (8 kbit) sans échange RTS-CTS, avec et sans utilisation de l'EIFS. La perte de performances due aux collisions entre trames et acquittements est partiellement compensée par la suppression du surcoût de transmission lié à l'utilisation de l'EIFS et la différence de performances est de l'ordre de 10%.

1.6 Conclusion

Ce chapitre présente la première modélisation théorique du scénario *ad hoc* le plus simple présentant un déséquilibre dans l'accès au médium. Ce scénario, mettant en scène trois couples émetteurs et

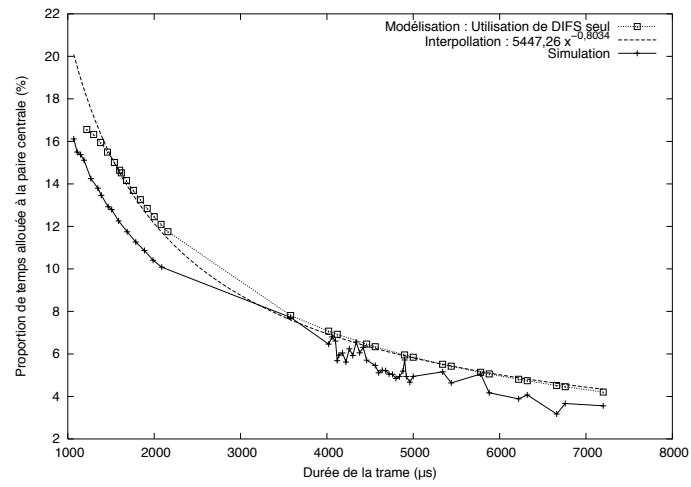


FIG. 1.37 – Proportion de trames émises par la paire centrale sans *EIFS* – comparaison entre modélisation, simulation et interpolation

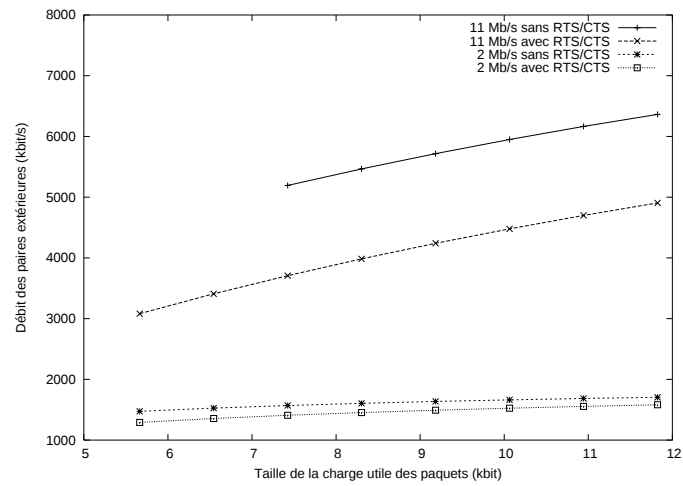


FIG. 1.38 – Débit au niveau application sans *EIFS* (modélisation) – paires extérieures

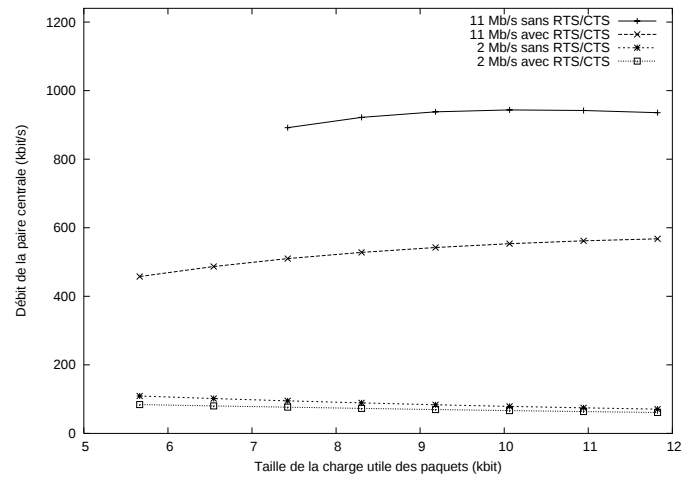


FIG. 1.39 – Débit au niveau application sans *EIFS* (modélisation) – paire centrale



FIG. 1.40 – Scénario dans lequel l'utilisation de l'*EIFS* empêche des collisions entre acquittements et données en *B*

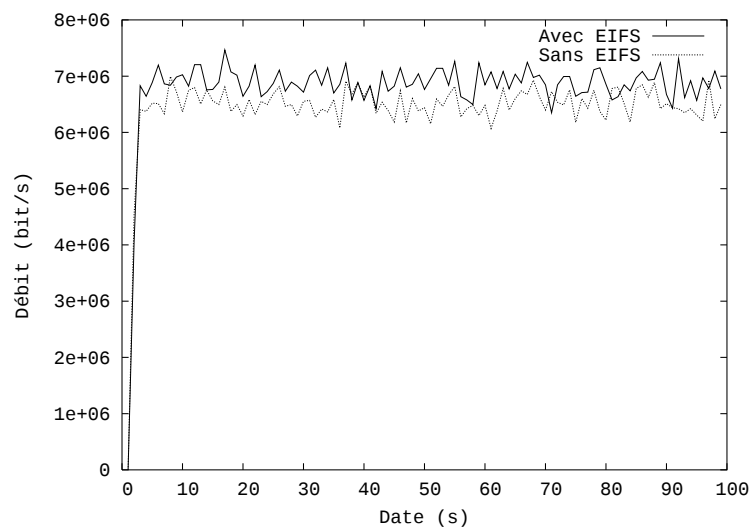


FIG. 1.41 – Impact sur les performance de l'utilisation de l'*EIFS*

récepteurs tels que les émetteurs sont en zone de détection de porteuse de leurs voisins affiche des performances fortement déséquilibrées. L'émetteur central ne peut accéder au médium qu'au plus 5,2 % du temps. Ce problème de performance est dû en grande partie à l'asymétrie du scénario, un émetteur étant en compétition avec deux autres émetteurs indépendants entre eux. Les mécanismes mis en œuvre dans la norme IEEE 802.11 accentuent par ailleurs le déséquilibre constaté. Au regard de ces résultats, il est légitime de s'interroger sur la pertinence du mécanisme d'EIFS. Sa présence représente un gain de performances de l'ordre de 10 % dans un scénario destiné à démontrer son utilité, le gain en terme de retransmission de paquets qu'il occasionne étant en partie compensé par le surcoût qu'il représente.

Cette étude ne concerne bien évidemment qu'un scénario isolé. Toutefois, dans un réseau *ad hoc* dont la topologie est *a priori* quelconque, ce type de situation peut survenir fréquemment et à différents endroits du réseau, tout comme les configurations dans lesquelles l'EIFS représente un gain de performances. Il est difficile de conclure sur la pertinence du mécanisme et le but de ce chapitre n'est pas de juger le choix qui a été fait lors de la conception du protocole IEEE 802.11. Il est toutefois important d'identifier précisément les cas pathologiques pouvant générer un problème de performances dans un réseau afin de définir des mécanismes permettant d'y remédier.

Le modèle présenté ici met en outre l'accent sur l'un des problèmes survenant lors de l'évaluation de performances d'un réseau *ad hoc*. En effet, la modélisation de phénomènes tels que l'asynchronisme pouvant exister entre différents mobiles est ardue et conduit à un modèle complexe. La chaîne de Markov présentée dans ce chapitre semble difficilement réductible et comporte pourtant un nombre d'états nettement supérieur à celui des modélisations de réseaux synchrones.

Une conclusion plus générale à tirer de cette étude est qu'il est extrêmement difficile dans un contexte *ad hoc* de prévoir la bande passante disponible afin d'offrir des garanties de qualité de service. En effet, dans une telle situation, l'émetteur central n'est pas capable d'identifier précisément la source du brouillage l'affectant. Les trames des deux perturbateurs se recouvrant mutuellement, il ne peut déduire un profil de trafic d'une écoute de ces trames. Il semble donc important, lorsque l'on souhaite offrir des garanties dans un réseau *ad hoc*, de se donner les moyens d'identifier le plus précisément possible les caractéristiques des trafics contre lesquels chaque émetteur aura à concourir pour l'accès au médium. Cette problématique constitue la base du protocole BRuIT présenté dans le chapitre suivant.

Un protocole de réservation de bande passante : BRuIT

Offrir aux applications multimédia ou temps réel des garanties sur la disponibilité des ressources du réseau est un problème ardu. Dans les réseaux filaires, de nombreux travaux ont conduit à la définition de solutions de qualité de service et l'étude de ces solutions a permis de déterminer leurs avantages et leurs limites. Les réseaux *ad hoc* introduisent de nombreuses contraintes liées essentiellement à l'utilisation du médium radio, à la mobilité et à l'absence de coordinateur central, compliquant ainsi l'évaluation de la capacité du réseau et, en conséquence, les solutions de qualité de service.

Dans ce chapitre, nous présentons un protocole de réservation de bande passante nommé BRuIT pour *Bandwidth reservation under interferences influence*. Ce protocole adopte une approche *a priori* afin de limiter l'apparition de congestion dans le réseau et ainsi de conserver un certain niveau de précision lors de l'évaluation des ressources disponibles. La transmission régulière d'informations sur le volume de trafic transmis par chacun des nœuds permet aux mobiles d'améliorer leur perception de l'état de saturation du réseau et par conséquent d'effectuer un contrôle d'admission plus précis. BRuIT est basé sur un protocole de routage avec qualité de service adoptant un fonctionnement réactif. Ce protocole utilise les informations collectées afin de déterminer si une requête cherchant à réserver un certain volume de bande passante peut être satisfaite ou non.

Ce chapitre débute par un état de l'art des solutions de qualité de service, au sens large du terme, qui ont été proposées dans un contexte *ad hoc*. Le protocole BRuIT est alors décrit en section 2.3 puis évalué par simulation en section 2.4.

2.1 Routage au mieux dans les réseaux *ad hoc*

Depuis la création du groupe de travail MANET à l'IETF, la recherche en matière de réseaux *ad hoc* a principalement concentré ses efforts sur la définition d'un protocole de routage point à point pour ces réseaux. L'abondance de propositions examinées par ce groupe de travail en témoigne. Sur plus de trente protocoles, quatre ont été retenues et aucun de ces quatre protocoles n'a su tirer son épingle du jeu en toutes situations. Ces quatre protocoles, héritiers des techniques utilisées dans le monde filaire, ainsi que la plupart de leurs concurrents, peuvent être classés en deux grandes catégories : les protocoles réactifs et les protocoles proactifs. Si ces deux grandes catégories intègrent la plupart des propositions faites à ce jour, il convient de remarquer l'existence d'une approche hybride ainsi que de quelques approches atypiques.

2.1.1 Routage proactif

L'approche proactive est certainement la plus proche des protocoles de routage actuellement utilisés dans les réseaux filaires. Chaque nœud mobile, routeur potentiel, dispose d'une table de routage indiquant, pour chaque destination dans le réseau, le routeur suivant sur le chemin. Si ce type d'approche nécessite le stockage d'informations dans chaque nœud du réseau, les routes sont immédiatement disponibles.

Dans Internet et les réseaux filaires actuels en général, la maintenance des tables est assurée une combinaison de protocoles de routage. À l'intérieur d'une même zone, ou domaine, un protocole de type *état de lien* tel qu'*Open Shortest Path First* (OSPF) [Moy89, Moy91] est, le plus souvent, chargé de la maintenance des tables de routage. Chaque routeur détermine la qualité du lien le reliant à ses voisins directs. Cette qualité est généralement exprimée comme l'inverse de la bande passante du lien considéré.

Cette information est alors transmise à tous les autres routeurs du domaine. Chaque routeur disposant construit ainsi un graphe pondéré représentant la cartographie complète du réseau. Il est alors à même de calculer, au moyen de l'algorithme de Dijkstra [Dij59], l'arbre des plus courts chemins le reliant à chacune des destinations potentielles du réseau relativement à la métrique considérée. Entre domaines, le routage est assuré par le protocole *Border Gateway Protocol* (BGP) [RL94], de type vecteur de distance. Chaque routeur transmet à ses voisins le chemin qu'il préconise pour atteindre chaque domaine ainsi que le coût associé. Fort de ces informations, chaque routeur détermine alors la meilleure route pour joindre chacun des sous-réseaux.

L'utilisation conjointe de ces deux protocoles a fait ses preuves dans Internet. Cependant, compte tenu du temps requis pour la mise à jour des tables de routage lorsqu'un changement dans le réseau survient, il était nécessaire d'adapter ces protocoles à la forte mobilité des réseaux *ad hoc*.

Le protocole *Optimized Link-State Routing* (OLSR) [CJ03] adapte le concept du routage à état de lien aux réseaux *ad hoc*. Ce protocole repose sur l'élection d'un sous-ensemble des nœuds, appelés *Multipoints Relais* (MPR) formant un ensemble dominant du graphe formé par les mobiles. Afin de limiter le coût des diffusions dans le réseau, chaque nœud sélectionne son ensemble de MPR parmi ses voisins directs de telle façon qu'un message diffusé retransmis par ces nœuds seuls atteint tout voisin à une distance de deux sauts de l'origine. La diffusion des informations topologiques permettant le routage est effectuée par les Multipoints Relais seuls et la fréquence de diffusion des informations de voisinage et de routage est adaptée à la dynamique des réseaux *ad hoc*.

Le protocole *Topology Dissemination Based on Reverse-Path Forwarding* (TBRPF) [OTG04] est, lui aussi, un protocole à état de lien maintenant un arbre de routage pour chaque nœud du réseau *ad hoc*. Les arbres de plus courts chemins sont déterminés au moyen d'un algorithme de Dijkstra adapté, et ce, grâce aux informations topologiques transmises par les nœuds du réseau. Chaque nœud transmet périodiquement à tous ses voisins directs, d'une part la liste de ces voisins directs et d'autre part l'arbre de routage qu'il a construit. À partir de ces informations, chaque mobile peut construire itérativement la topologie à une distance de deux sauts de lui ainsi qu'un arbre de plus courts chemins. TBRPF préconise en outre de ne transmettre non pas la totalité des informations dans chaque paquet de contrôle, mais plutôt les différences par rapport au dernier paquet émis, générant ainsi un volume de trafic relativement faible.

2.1.2 Routage réactif

Le routage réactif fonctionne à la demande. Aucune information n'est stockée dans les routeurs sur des destinations vers lesquelles le nœud concerné n'a pas de route active. Lorsqu'une application désire contacter un correspondant, et seulement à cet instant, une route est recherchée. Ce type de protocoles présente l'avantage de ne pas surcharger le réseau en trafic de contrôle et de ne requérir des routeurs qu'une capacité de stockage minimale. Toutefois, le délai d'établissement des communications peut être élevé lorsque le réseau est chargé ou lorsque le correspondant est à une distance importante de l'émetteur.

Le protocole *Ad Hoc On Demand Distance Vector* (AODV) [PBRD03] est un protocole réactif de type vecteur de distance. Lorsqu'un mobile désire envoyer un flux de paquets à destination d'un autre mobile et qu'il ne dispose d'aucune route pour joindre son correspondant, il diffuse un paquet de requête de route (RREQ). Ce paquet se propage par inondation à travers tout le réseau. Lorsque la destination est atteinte, elle envoie un paquet de réponse (RREP) qui traversera le chemin inverse de la requête correspondante, provoquant l'activation de la route dans les nœuds intermédiaires. Lorsque ce paquet arrive à la source du trafic, le transfert peut commencer. Pour palier la mobilité des nœuds intermédiaires, un mécanisme de reconstruction locale des routes est mis en place.

Le protocole *Dynamic Source Routing* (DSR) [JMH03] effectue, quant à lui un routage par la source. La recherche de routes est effectuée de la même manière qu'avec AODV, par inondation du réseau. Cependant, aucune information n'est stockée dans les routeurs intermédiaires. Les paquets de requête sont, quant à eux modifiés de saut en saut et contiennent la liste intégrale des routeurs qu'ils ont traversés. Pour chaque requête de route reçue, la destination envoie une réponse. Ces réponses décrivent l'intégralité du chemin emprunté par la requête correspondante. Lors de la communication, la route que chaque paquet de données doit emprunter est contenue dans l'en-tête du paquet, ce qui permet à la source disposant d'autant de routes qu'elle a reçu de réponses d'effectuer un équilibrage de charge ou de disposer de routes de secours pour palier la mobilité.

2.1.3 Autres approches

La plupart des protocoles proactifs ou réactifs sont, finalement, une adaptation de protocoles filaires éprouvés au contexte mobile des réseaux *ad hoc*. Déterminer quelle philosophie ou quel protocole est le plus approprié est un problème épineux, compte tenu de l'absence de limite quant à l'utilisation des réseaux *ad hoc*. En fonction de la dynamique du réseau, de sa densité, de son étendue ou du type d'applications utilisées, tel ou tel protocole va se révéler le meilleur en terme d'utilisation du réseau, d'équilibrage de charge, etc. Un certain nombre d'approches qualifiées d'hybrides, telles que *Zone Routing Protocol* (ZRP) [HPS02] proposent d'utiliser un protocole proactif pour le routage à destination de mobiles proches et un protocole réactif pour les longues distances.

D'autres approches ont par ailleurs été examinées. Un routage géographique, par exemple, est envisageable compte tenu du fort lien existant entre la topologie du réseau et la géographie de la zone concernée. La plupart des protocoles dits géographiques, par exemple [HBBW04] utilisent un système de localisation par satellite tel que *Global Positioning System* ou *Galileo* et reposent sur la transmission périodique d'informations de localisation. Un algorithme de routage détermine alors la route en prenant en compte ce paramètre.

Beaucoup de stratégies ont été examinées pour résoudre le problème du routage dans ce type de réseaux. Toutefois bon nombre de problématiques ont été examinées de façon nettement moins exhaustive par la communauté scientifique. Les problématiques de qualité de service ont suscité un certain intérêt sans toutefois parvenir à une solution satisfaisante, malgré le nombre croissant de propositions. Si la problématique du routage avec qualité de service fait partie des axes de recherche privilégiés par le sous-groupe de recherche *Ad hoc Network Systems* de l'IRTF, l'étude en reste encore embryonnaire par certains aspects.

2.2 Qualité de Service et réseaux *ad hoc*

Dans un état de l'art de la recherche en matière de qualité de service pour réseaux *ad hoc* [WH01], Wu et Harms introduisent une classification des solutions de qualité de service pour les réseaux *ad hoc*. Cette classification sépare les solutions proposées dans ce domaine en quatre catégories. Les modèles de qualité de service regroupent les définitions d'architectures destinées à assurer une certaine qualité de service. Les protocoles de signalisation définissent un ensemble de messages de contrôle, destinés par exemple à provoquer la réservation de ressources dans les routeurs. Les protocoles de routage avec qualité de service sont chargés de la recherche de routes répondant à certains critères. Enfin, les protocoles d'accès au médium avec qualité de service fournissent un ensemble d'outils permettant de mettre en œuvre certaines règles de qualité de service. Ces différentes catégories ne sont évidemment pas disjointes et il est souvent difficile de classer les solutions proposées dans une catégorie particulière. C'est pourquoi, dans la suite, nous préférons modifier quelque peu cette terminologie. Les catégories modèles de qualité de service et protocoles de routage avec qualité de service seront conservées intactes. Au terme protocoles d'accès au médium avec qualité de service, nous substituerons le terme mécanismes de différenciation de service. En effet, la différenciation de service peut, comme nous le verrons, être effectuée à d'autres niveaux. Enfin, la dernière catégorie que nous considérerons sera celle des mécanismes de réservation, le terme signalisation étant trop vague, la mise en œuvre de tout protocole nécessitant une certaine signalisation ou trafic de contrôle.

2.2.1 Architectures de qualité de service

La conception de première architecture de qualité de service définie par l'*Internet Engineering Task Force* (IETF) date de la fin des années 1980 ; il s'agit de l'architecture *Integrated Services* [BCS94], souvent appelée *IntServ*. Ce modèle définit une extension à l'architecture IP destinée à assurer aux différents flux de données des garanties sur le délai de bout en bout, le débit, etc. Tout flux de données, identifié par l'adresse de sa source, l'adresse de sa destination et son port de destination, peut bénéficier de garanties individuelles pour peu que le réseau puisse assurer le niveau de service demandé. Deux types de services sont définis. *Guaranteed Service* [SPG97] offre des garanties de délai strictes pour des applications temps réel telles que la voix sur IP ou la vidéo. Les routeurs sont alors chargés de borner le délai passé par les paquets de tels flux dans les files d'attente. *Controlled Load* [Wro97] offre aux applications sensibles à une montée en charge du réseau un service similaire à celui d'un réseau non

chargé, c'est-à-dire qu'une proportion importante de paquets est transmise et un délai faible est garanti. IntServ est, la plupart du temps utilisé conjointement avec le protocole de réservation RSVP [BZB⁺97] décrit plus bas. L'architecture ainsi définie est orientée connexion, ce qui requiert un volume de stockage et une capacité de traitement importants au niveau de chaque routeur. Il s'agit là de la principale limitation de ce modèle, Internet comportant aujourd'hui une multitude de nœuds, il n'est pas possible de déployer un mécanisme si coûteux dans le cœur du réseau.

L'architecture *Differentiated Services* (DiffServ) [BBC⁺98] répond à ce besoin d'un mécanisme léger de qualité de service dans Internet. L'architecture proposée définit plusieurs classes de trafic, les différents flux s'intégrant à une de ces classes afin de bénéficier des garanties correspondantes. On ne parle plus alors de garanties par flux mais de garanties par agrégat de flux. Par exemple, la classe *Premium service* ou *Expedited Forwarding* (EF) [JNP99] assure un délai, une gigue et un taux de pertes faibles ainsi qu'une bande passante minimale garantie. Elle est destinée aux applications temps réel telles que la vidéo ou la voix. La classe *Assured Forwarding* (AF) ou *Assured Service* [HBWW99], permet aux applications de définir finement le taux de pertes admissible. Quatre sous-classes sont définies, chacune séparée en trois niveaux de priorité. À chaque sous-classe est allouée une proportion de la bande passante et lorsque le débit à émettre dans une classe dépasse l'allocation, les paquets sont détruits par ordre de priorité. La philosophie de DiffServ consiste à reporter le maximum de traitements en bordure de réseau. Chaque paquet d'un flux appartenant à une classe est marqué par le routeur d'entrée au moyen d'un champ particulier de l'en-tête IP [NBBB98]. Les routeurs traversés par le paquet appliquent alors des politiques de mise en file d'attente (ou *Per Hop Behavior* - PHB) particulières et dépendantes de la classe de trafic associée au paquet.

Ni le modèle IntServ, ni le modèle DiffServ ne représente la solution parfaite en matière de qualité de service pour Internet. IntServ est en général jugé comme passant mal à l'échelle et DiffServ ne permet pas aisément de privilégier un flux particulier, fût-il le plus important du réseau local. Dans beaucoup de réseaux, par ailleurs, il est difficile de s'assurer de la coopération des différents routeurs afin de traduire au niveau de l'accès au médium la priorité des différentes classes ou des différents flux. Le groupe de travail *Integrated Services over Specific Link Layer* (ISSLL¹) de l'IETF préconise l'utilisation de DiffServ dans le cœur du réseau en raison de son faible coût. En effet, il est rare que le cœur du réseau soit un goulet d'étranglement. IntServ pourrait en revanche être utilisé en bordure de réseau, où les technologies sont hétérogènes et où il est souvent nécessaire de déployer des mécanismes particuliers afin de s'assurer de maîtriser le partage des ressources. Le travail principal de ce groupe consiste à définir les techniques requises pour la mise en œuvre d'IntServ au-dessus de diverses technologies de niveau liaison, par exemple ATM. Ce groupe de travail définit la correspondance entre les différentes classes de services définies par IntServ et les paramètres des couches liaison considérées. Toutefois, aucun document de ce groupe de travail ne concerne à ce jour des réseaux *ad hoc* basés sur la norme IEEE 802.11.

Le modèle de qualité de service proposé par Xiao *et al.*, *a Flexible Quality of Service Model for Mobile Ad Hoc Networks* (FQMM) [XSLC00] définit une architecture hybride adaptée à des réseaux *ad hoc* de taille moyenne, soit une cinquantaine de mobiles. Ce modèle combine réservation statistique et réservation statique en définissant plusieurs classes de trafic, l'une de ces classes étant dédiée aux réservations explicites de bande passante. FQMM définit l'ensemble des composants nécessaires à la mise en place de l'architecture préconisée. Le routage est assuré par un protocole de routage au mieux et une vérification *a posteriori* du respect des contraintes est effectuée.

Le modèle *Two-Layered Quality of Service Model for Reactive Routing Protocols for Mobile Ad Hoc Networks* (2LQoS) [NBMR02] proposé par Nikaein *et al.* considère deux types de métriques de qualité de service afin de définir des classes de trafic. Les métriques en rapport avec le bon fonctionnement du réseau, telles que le nombre de sauts des routes, le niveau de batteries des mobiles routeurs, ou encore la stabilité des routes sont utilisées lors de la découverte de chemin. L'utilisation de ce type de métriques permet de ne pas considérer certaines routes lors de l'exploration du réseau afin, par exemple, d'effectuer un équilibrage de charge. La recherche d'une route vers une destination s'effectue de façon réactive et ce processus sélectionne *a priori* plusieurs routes vers une même destination. Les applications spécifient alors un certain nombre de critères portant sur le délai, la gigue ou encore le débit de la route demandée, permettant de sélectionner la route la plus convenable parmi l'ensemble des routes découvertes. Les différentes classes définies par IntServ ou DiffServ peuvent alors être traduites en combinaison de ces métriques.

¹<http://www.ietf.org/html.charters/issll-charter.html>

Si la définition d'un modèle de qualité de service semble indispensable, il est difficile d'évaluer ou de comparer les différentes approches proposées. En effet, ces descriptions restent souvent conceptuelles et ne proposent ni mécanisme d'évaluation des ressources disponibles, ni de mécanisme permettant une différenciation effective de service. Toutefois, ces tâches sont très dépendantes des protocoles sous-jacents. Effectuer une différenciation de services sur un médium radio partagé ne pourra pas être réalisé au moyen de mécanismes similaires à ceux qui sont utilisés sur une fibre optique.

2.2.2 Différenciation de services

Les protocoles de différenciation cherchent à mettre en œuvre des priorités entre différents flux ou différents terminaux. Au sein d'un même mobile, il est possible de définir des priorités entre plusieurs flux émis ou routés au moyen de files d'attente dont le fonctionnement est plus souple que la simple file FIFO.

Un premier type de file d'attente évolué est la file à priorité (*Priority Queuing*). Cette file d'attente est composée de plusieurs files FIFO, une par classe, et les files sont vidées par ordre de priorité. Tant qu'il reste des paquets dans une sous-file, les files de priorité moindre ne sont pas vidées. Ce type de file d'attente définissant une priorité absolue, il est aisé de provoquer une famine des flux de faible priorité en saturant les files privilégiées ou en multipliant le nombre de files.

Une évolution moins discriminatoire de la file précédente consiste à ne pas vider en priorité absolue les files privilégiées, mais de prendre un paquet tour à tour dans chacune des files. Cette stratégie, appelée tourniquet (*Round Robin*) permet de définir plusieurs classes de trafic équivalentes au sein d'un même nœud.

Il est par ailleurs envisageable de déséquilibrer ce mécanisme en autorisant à chaque tour les files prioritaires à émettre un volume de données plus important que les files de faible priorité. Les files de type (*Weighted Fair Queueing*) appliquent ce principe. Un classificateur trie les paquets, en général au moyen d'une fonction de hachage considérant divers critères tels que le protocole de transport utilisé, le type de trafic transporté ou plus simplement la destination. Les différentes files ainsi remplies sont vidées tour à tour d'un volume de données dépendant de leur priorité.

Une multitude de politiques de gestion de files d'attentes sont envisageables et plus encore ont été proposées. Cependant, lorsqu'il s'agit de définir des priorités d'accès entre différents terminaux ou entre différents flux routés par des terminaux différents, il n'est pas possible de n'utiliser que des files, aussi évoluées soient-elles. Il est nécessaire de modifier le comportement du protocole d'accès au médium. La plupart des protocoles gérant les accès concurrents à un médium partagé sont, dans la mesure du possible, équitables et diverses techniques ont été proposées afin de déséquilibrer ce mécanisme. Nous ne présenterons pas ici les solutions centralisées pour lesquelles un coordinateur indique à chaque nœud du réseau lorsqu'il est autorisé à émettre. En effet, ces solutions sont peu adaptées aux réseaux *ad hoc*.

Dans [DC99], Deng et Chang proposent d'intégrer à la fonction DCF de la norme IEEE 802.11 un système de priorité entre stations basé sur l'utilisation de différents temps d'attente avant décrémentation du *backoff*. Cet intervalle, auparavant constant et égal à *DIFS*, est alors fonction de la priorité de la station concernée, les stations de plus haute priorité utilisant une valeur plus faible et commençant la décrémentation de leur *backoff* plus rapidement que les autres.

Dans [AC01], puis dans [AC03b], Aad *et al.* étudient plusieurs techniques permettant d'effectuer une différenciation entre stations basées sur la fonction de coordination distribuée de la norme IEEE 802.11 utilisée en présence d'un point d'accès. Plusieurs paramètres de ce protocole peuvent, en effet, être adaptés afin de donner la priorité à certains flux par rapport à d'autres. Il est possible d'adapter la taille initiale de la fenêtre de contention, de modifier l'accroissement de cette fenêtre lors de collisions, de limiter la taille maximale des trames ou encore de définir des temps d'attente avant décrémentation du *backoff* différents en fonction de la priorité des stations. L'étude analytique ainsi que les simulations de l'influence de ces quatre adaptations sur les flux TCP et UDP montrent que remplacer le mode d'accroissement de la taille de la fenêtre de contention peut conduire à une bonne différenciation des flux UDP mais s'accompagne d'un accroissement significatif de la gigue et une variation importante des débits. L'impact sur les flux TCP, en revanche, est plus faible. Cette méthode de différenciation présente en outre l'inconvénient de n'être active qu'en présence de collisions. La qualité de la différenciation de flux TCP obtenue en utilisant différentes tailles minimales de fenêtres de contention est fortement liée à la priorité du point d'accès retransmettant les acquittements. Limiter la longueur maximale des trames des stations en fonction de leur priorité permet d'obtenir une bonne différenciation, mais peut avoir un

fort surcoût lié au fonctionnement du protocole IEEE 802.11. La différenciation basée sur différents temps d'attente avant décrémentation du *backoff* conduit, elle aussi à une différenciation efficace et permet en outre de donner à des flux TCP la priorité sur des flux UDP. Toutefois, la qualité de ces différentes techniques vis-à-vis de TCP est très liée à la priorité du point d'accès retransmettant les acquittements TCP. Cette observation plaide en faveur d'une différenciation par flux plutôt que d'une différenciation par terminal, à plus forte raison dans un contexte *ad hoc*.

Le groupe 802.11 de l'IEEE est en cours d'élaboration d'une évolution du protocole de transmission sans fil portant son nom en vue d'une prise en charge de notions de qualité de service. Cette norme, IEEE 802.11e, définit deux méthodes d'accès au canal radio. L'une, HCF, est centralisée et destinée à être utilisée en coordination avec un point d'accès. L'autre, totalement distribuée, appelée *Enhanced Distributed Coordination Function* (EDCF) est une évolution de la DCF incorporant certaines des techniques présentées ci-dessus afin de définir un mécanisme d'accès au médium totalement distribué avec différenciation de service. Une présentation de l'état de ces travaux, ainsi qu'une rapide analyse a été présentée par Mangold *et al.* dans [MCM⁺02]. Une évaluation des performances de ce protocole a par ailleurs été présentée dans [GN02].

Trois méthodes de différenciation ont été retenues. Tout d'abord, deux stations peuvent utiliser différentes tailles minimales de fenêtre de contention. Il est aussi possible d'utiliser différents schémas d'augmentation de cette fenêtre lors de collisions entre trames. Dans la norme IEEE 802.11 initiale, la taille de la fenêtre de contention est doublée à chaque collision. Dans cette évolution, le facteur par lequel est multipliée la fenêtre de contention, appelé facteur de persistance, est modifiable. Enfin, il est possible de définir des priorités dans l'accès au canal par l'utilisation de différents intervalles précédant la décrémentation du *backoff*. Cet intervalle se nomme, dans cette norme, *AIFS* (*Arbitration Inter Frame Spacing*) et ne peut descendre en deçà de la valeurs de *DIFS* définie par la norme IEEE 802.11.

La couche d'accès au médium de chaque terminal comporte huit files d'attente distinctes appelées catégories de trafic définissant huit niveaux de priorité. À chacune de ces files d'attente sont associées des valeurs des trois paramètres précédents. Au sein d'un nœud, ces files d'attentes sont vidées par ordre de priorité. La méthode EDCF permet donc de définir pour chaque trame un niveau de priorité au sein du nœud émetteur, ainsi qu'un niveau de priorité lors de son émission sur le médium.

L'évaluation de ce protocole par simulation indique que la différenciation s'opère bien. Les flux de faible priorité peuvent être victimes de famine, il faut donc être prudent lors de l'utilisation de différentes classes de trafic dans un environnement surchargé. Un mécanisme de limitation des débits des différents flux pourrait être nécessaire pour pallier ce phénomène. La coexistence de mobiles utilisant la norme IEEE 802.11 dans le réseau limite la marge de manœuvre offerte par l'EDCF. En effet, la limite inférieure de l'AIFS étant fixée à DIFS, rendre des flux privilégiés prioritaires par rapport aux flux au mieux ne peut mettre en jeu que deux paramètres, la taille initiale de la fenêtre de contention et son mode d'accroissement, ce dernier paramètre n'étant utilisé que lorsque des collisions affectent les trames des flux privilégiés. Faire cohabiter des flux au mieux et des flux privilégiés limite donc la possibilité de différencier les flux privilégiés entre eux.

Cette évolution du standard modifie par ailleurs le comportement du mode d'accès au médium en autorisant des rafales de trames au niveau MAC. Ce mécanisme, nommé *Contention Free Burst* (CFB) permet aux terminaux d'émettre successivement plusieurs trames séparées non pas par une période de contention usuelle mais par un intervalle *SIFS*, à la manière de différents fragments d'une même trame. Si ce mode de fonctionnement peut se révéler utile dans le cas de certains types de trafic comme les trafics ON-OFF alternant émissions en rafales et périodes d'inactivité, il a un impact négatif sur la gigue des flux.

À l'image de la fonction de coordination distribuée de la norme IEEE 802.11, la taille des fenêtres de contention des différents flux sont statiques et ne prennent pas en compte la dynamique du réseau. Une taille de fenêtre trop élevée engendre un surcoût inutile, mais l'utilisation de tailles de fenêtres trop faibles accroît le nombre de collisions dans le réseau. Lorsqu'une collision frappe une trame, la taille de cette fenêtre est doublée pour la DCF et multipliée par le facteur de persistance pour l'EDCF. Lorsqu'une transmission est réussie, cette valeur est réinitialisée. Or, transmettre une trame avec succès ne signifie pas que la situation de congestion ayant provoqué l'accroissement de la fenêtre de contention est résolue. Lorsque le réseau est chargé, on constate alors des oscillations dans la taille des fenêtres qui ne sont pas souhaitables. Dans [NABT03, AC03a], Aad *et al.* proposent, dans le cadre de la fonction DCF, de ne pas réinitialiser la fenêtre de contention après une transmission réussie mais de diviser sa taille par un facteur. Cette technique conduit à une amélioration de performances sensible en terme d'utilisation

globale de la capacité du canal radio. Romdhani *et al.* proposent d'utiliser la même technique pour les nœuds utilisant la fonction EDCF. Cette amélioration du protocole d'accès au médium est présentée dans [RNT03] et appelée *Adaptive Enhanced Distributed Coordination Function* (AEDCF). Dans des réseaux saturés, cette approche semble conduire à de meilleurs résultats en terme de débit, de délai et de taux de pertes que la fonction EDCF telle qu'elle est définie à ce jour.

Dans [ZCYM04], Zhu *et al.* proposent *Enhanced Distributed Coordination Function with Dual-Measurement* (EDCF-DM), une évolution supplémentaire par rapport à AEDCF consistant à réduire et augmenter conjointement les fenêtres de contention de toutes les files d'attente internes à un nœud. En effet, l'accroissement des fenêtres de contention est généralement lié à la présence de congestion dans le réseau, et si une classe de flux subit des collisions, il est probable que les autres classes de trafic seront soumises aux mêmes problèmes. Cette modification conduit à une amélioration des délais de transmission et à une différenciation encore plus prononcée que EDCF et AEDCF.

Dans [BCV01], Barry *et al.* proposent un mécanisme de différenciation de services pour réseaux locaux sans fil reposant sur l'ajustement des valeurs minimales et maximales des fenêtres de contention des différents mobiles. L'utilisation d'une couche d'accès au médium virtuelle simulant l'envoi de trames permet en outre de mesurer les paramètres du réseau tels que le délai, le taux de collision, etc. Dans cette couche MAC virtuelle, tout se passe comme si la trame allait effectivement être envoyée. Elle concourt pour l'accès au médium et lorsque la trame peut être émise, elle ne l'est pas réellement afin de ne pas utiliser de ressources réseau inutilement. On utilise pour mesurer le délai de transmission des trames des paquets tests passés à cette couche, et l'on utilise une probabilité de collision afin de s'adapter à la charge du réseau. Cette technique permet une évaluation correcte des ressources du réseau mais est assez consommatrice en terme de puissance de calcul.

Un certain nombre d'autres propositions [PM01, MLGTR03] visant à offrir une différenciation de service au protocole 802.11 de l'IEEE ont été publiées. Les mécanismes de différenciation de service mis en œuvre dans ces solutions sont globalement identiques à ceux qui ont été choisis pour la conception de la norme IEEE 802.11e.

Ces différentes modifications sont cependant sensibles aux problèmes de type stations cachées. Assurer une priorité entre plusieurs flux dans un tel contexte est difficile, comme le souligne [XV02]. Les auteurs proposent un mécanisme nommé *Busy Tone Priority Scheduling* (BTPS) reposant sur l'utilisation de deux porteuses en bande étroite afin de signaler la présence de paquets à émettre dans les stations de haute priorité en marge de la bande passante destinée aux transmissions. L'utilisation de deux porteuses de signalisation permet d'obtenir un comportement analogue à celui des mécanismes RTS-CTS.

Toutes les solutions présentées jusqu'ici pour effectuer une différenciation de services dans des réseaux sans fil se plaçaient au niveau de la couche d'accès au médium. Dans SWAN [ACVS03, ACVS02], Ahn *et al.* proposent un mécanisme permettant de différencier deux classes de trafic et se positionnant au-dessus d'une couche d'accès au médium au mieux. Les débits du trafic au mieux et du trafic temps réel sont limités au moyen de filtres et le débit du trafic au mieux est adapté en fonction des besoins du trafic temps réel. L'objectif de ce mécanisme, en outre de la différenciation de service, est de maintenir le délai de transmission le plus bas possible tout en conservant un débit élevé en évitant un remplissage trop important des files d'attente. Lors de la transmission d'une trame, les routeurs sont chargés de mesurer les délais d'accès au médium en observant les instants de mise en file d'attente de chaque trame et les instants de réception des acquittements de niveau MAC correspondants. Lorsque ce délai s'allonge au-delà d'une limite fixée, les routeurs positionnent le bit de notification explicite de congestion (ECN) dans les trames suivantes afin d'avertir les destinations de l'augmentation de la charge dans une zone du réseau. Les destinations sont alors chargées d'avertir les sources correspondantes afin qu'elles prennent les mesures appropriées. Divers mécanismes sont mis en place afin d'éviter les phénomènes d'oscillations pouvant survenir lorsque toutes les sources limitent leur trafic simultanément puis l'augmentent simultanément par la suite.

Dans [LAS01a], [LAS01b] puis dans [LAS03], Lindgren, Almquist et Schelén évaluent les performances de différentes solutions permettant d'effectuer une différenciation de services dans des réseaux locaux sans fil. Ces articles comparent essentiellement Blackburst [SK96], la fonction EDCF de la norme IEEE 802.11e, DFS (*Distributed Fair Scheduling*) [VBG00] de Gupta et Vaidya et la fonction PCF définie par la norme IEEE 802.11. En terme d'utilisation des ressources radio, Blackburst semble obtenir les meilleures performances, puisqu'il n'engendre aucune collision. Ce protocole permet par ailleurs d'obtenir la différenciation la plus nette entre deux classes de trafic distinctes. La fonction EDCF conduit elle aussi à une différenciation efficace entre deux classes de trafic, mais la présence de collisions nombreuses limite

les performances globales du protocole. Même si Blackburst engendre un surcoût important lorsque la charge du réseau augmente, cette solution reste la meilleure. La solution préconisée par Deng et Chang obtient, elle aussi d'assez bons résultats. En termes de délais de transmission, Blackburst reste la solution la plus performante, la période de brouillage précédant l'émission de chaque trame étant compensée par l'absence de collisions dans le réseau, et donc par l'absence d'accroissement du *backoff* et l'absence de retransmissions. La différenciation offerte par Blackburst est cependant très marquée et peut parfois conduire à une famine touchant les flux de faible priorité.

Un grand nombre de mécanismes de différenciation de services ont été proposés et analysés pour les réseaux locaux sans fil basés sur la norme IEEE 802.11. Le travail du groupe IEEE 802.11e incorpore la plupart de ces propositions. La solution qui en découle est prometteuse et devrait être suffisamment souple pour permettre une réelle intégration d'une architecture de différenciation de services telle que DiffServ dans les réseaux locaux sans fil. Toutefois, les performances de l'EDCF dans un réseau *ad hoc* multisauts seront soumises aux mêmes contraintes que la DCF et feront sans doute couler beaucoup d'encre.

2.2.3 Mécanismes de routage avec qualité de service

La différenciation de services permet de s'assurer dans une certaine mesure du bon fonctionnement des applications multimédias au moyen de mécanismes souvent légers en termes de puissance de calcul, d'espace de stockage et de signalisation. Toutefois, certaines applications multimédias nécessitent des garanties plus quantitatives. Transmettre un flux audio dans un réseau en respectant des contraintes de délai n'est pas toujours possible. Lorsque la charge du réseau est importante, il peut être nécessaire, avant d'autoriser une application à utiliser des ressources réseau, de s'assurer que les besoins de cette application pourront être satisfaits. Comparer ces besoins et les ressources disponibles dans le réseau est la tâche de mécanismes de contrôle d'admission.

D'autre part, dans un réseau étendu, il est probable que certaines zones subiront une plus forte charge que d'autres. Dans ce cadre, un routage plus court chemin usuel pourrait conduire à une mauvaise répartition de la charge. Sélectionner des routes ne correspondant pas à un critère de plus court chemin peut conduire à un meilleur équilibrage de la charge. En présence d'un mécanisme de contrôle d'admission, cette politique peut permettre à un plus grand nombre de flux d'utiliser les ressources du réseau. Cependant, utiliser une politique de routage basée sur la disponibilité de ressources a un coût, puisqu'il est nécessaire de connaître l'état du réseau. La mesure de la capacité d'un réseau, et plus particulièrement d'un réseau *ad hoc* est un problème pouvant être complexe, comme l'a montré le chapitre 1.

Le contrôle d'admission, l'équilibrage de la charge du réseau ainsi que la recherche de routes répondant aux critères des applications sont en général les tâches incombant à un protocole de routage avec qualité de service. Ce type de protocole a pour but de trouver des routes satisfaisant les contraintes des applications, tout en s'assurant que l'utilisation du réseau reste conforme à certains critères.

Dans [MBAAP02], Munaretto *et al.* proposent une adaptation du protocole de routage OLSR afin de ne plus rechercher des plus courts chemins en terme de nombre de sauts mais en terme de délai, de bande passante ou de n'importe quelle métrique de qualité de service. Les mobiles sont chargés de mesurer les paramètres de qualité de service afin de transmettre ces informations aux autres routeurs du réseau. Les simulations réalisées montrent en particulier que le routage plus court chemin dans des réseaux *ad hoc*, malgré le coût de l'accès au médium, ne sélectionnent pas systématiquement la route présentant le meilleur délai de bout en bout.

Ramanathan et Steenstrup, dans [RS98], proposent de déployer une architecture hiérarchique appelée *Multimedia support for Mobile Wireless Networks* (MMWN) destinée à des réseaux *ad hoc* afin de faciliter le routage avec qualité de service. Un sous-ensemble dominant des mobiles est sélectionné afin de créer un cœur de réseau autour duquel les autres mobiles s'organisent en cellules. Les nœuds du cœur sont interconnectés par le biais d'autres nœuds passerelles. Sur cette abstraction du réseau, on détermine à nouveau des cellules, itérant le processus jusqu'à construire une vue hiérarchique du réseau. Au-dessus de cette architecture, les auteurs proposent de déployer un service de localisation des mobiles ainsi qu'un algorithme de routage de type état de lien. Chaque nœud du cœur est chargé de la surveillance des liens qui lui sont affiliés. Le routage est effectué au moyen de l'algorithme de plus court chemin de Dijkstra [Dij59] appliqué à un graphe pondéré en fonction des besoins des applications. L'algorithme

choisit la route présentant les meilleures caractéristiques par rapport à la métrique considérée puis, s'il y a plusieurs chemins équivalents, choisit le plus court en nombre de sauts.

Le protocole CEDAR (*a Core Extraction Distributed Ad hoc Routing algorithm*) [SSB99] proposé par Sinha *et al.* repose, lui aussi sur l'élection d'un cœur de réseau formant un ensemble dominant dans le graphe de connectivité associé. Ce cœur est établi de manière distribuée et en utilisant uniquement des informations locales son cardinal n'est en conséquence pas toujours minimum. Le but de CEDAR est de permettre un routage basé sur des informations de disponibilité de bande passante. Les nœuds du cœur sont chargés de propager les informations sur l'état des liens qui les entourent. La stabilité des liens ainsi que la bande passante disponible détermine la distance de propagation des caractéristiques de chaque lien. La recherche de route est effectuée en inondant le cœur de réseau en mode point à point. N'impliquer que les nœuds du cœur dans cette recherche de route permet de réduire considérablement le surcoût en terme d'utilisation des ressources réseau par rapport à une inondation totale.

Réduire le coût des recherches de routes répondant à un critère de qualité de service est aussi la préoccupation de Chen et Nahrstedt, auteurs de *Ticket Based Probing* (TBP) [CN99]. La solution distribuée proposée dans cet article cherche à résoudre le problème de la découverte d'un chemin de coût minimum satisfaisant une contrainte de délai ainsi que le problème de la découverte d'un chemin de coût minimum satisfaisant une contrainte de bande passante. La fonction de coût considérée ici peut être le nombre de sauts de la route ou n'importe quelle métrique additive, c'est-à-dire pour laquelle la valeur d'une route est la somme des valeurs des liens la composant. Déterminer une route minimisant une métrique additive et répondant à une contrainte sur une autre métrique additive, telle que le délai, est un problème NP-complet. En revanche, déterminer une route minimisant une métrique additive et répondant à une contrainte sur une métrique concave, c'est-à-dire pour laquelle la valeur d'une route est le minimum des valeurs des liens sur la route, telle que la bande passante est un problème résoluble en temps polynomial.

La découverte de chemins admissibles dans TBP est effectuée à la demande, par l'envoi au travers du réseau d'une sonde à laquelle sont associés un certain nombre de jetons appelés tickets et représentant la possibilité d'explorer plusieurs chemins distincts dans le réseau. Chaque routeur recevant cette requête peut décider de la transmettre à un voisin ou de la diviser si elle contient plus d'un ticket. Les sondes résultant de la division d'une requête se partagent le nombre de tickets de cette dernière et le processus continue jusqu'à découvrir une ou plusieurs routes admissibles quant au critère de qualité de service considéré. Lorsque les sondes découvrent plus d'un chemin admissible, la source choisit dans cet ensemble le chemin de coût minimal.

Concernant la recherche de route satisfaisant une contrainte de délai, une heuristique est proposée consistant à chercher parallèlement des routes présentant des délais faibles, de manière à s'assurer du succès de la démarche et des routes présentant en priorité des coûts faibles de manière à rechercher des routes admissibles de coût peu élevé.

Le nombre de tickets associés à une requête est déterminé par la source en fonction de la difficulté présumée de la recherche de route. Dans [RHZ00], Raju *et al.* proposent d'utiliser des règles de logique floue afin de déterminer ce nombre de tickets.

La comparaison de ce protocole avec un algorithme de plus court chemin et un algorithme d'inondation montre que les routes obtenues sont en général proches des meilleures et que la recherche est un processus peu coûteux en terme d'occupation du réseau.

Ce protocole n'est pas *a priori* lié à une technologie particulière. Cependant, il est nécessaire pour le mettre en œuvre, d'être à même d'estimer la bande passante et le délai des liens, ce qui est une tâche ardue avec les protocoles de type CSMA/CA. Laio *et al.*, dans [LTWS01, LTWS02] adaptent ce mécanisme à la recherche de chemins multiples dans un réseau radio fonctionnant sur un médium de type TDMA.

2.2.4 Mécanismes de réservation de ressources

Afin de ne pas être vulnérable aux variations de charge du réseau, il est parfois nécessaire de réserver une part des ressources dudit réseau. Cette réservation, afin d'être effective, doit avoir un impact sur le mécanisme de contrôle d'admission et sur le mécanisme de routage. Les différents composants du réseau devront alors œuvrer au respect de cette réservation.

Il n'est pas rare que ces protocoles de routage avec réservation réalisent un routage par flux plutôt que par destination, accroissant par là même la complexité du mécanisme de routage ainsi que la taille

des informations devant être stockées dans les différents routeurs. En conséquence, la critique qui est souvent faite à ces solutions est qu'elles ne passent pas à l'échelle et leur utilisation se limite souvent à des réseaux locaux.

Resources Reservation Protocol (RSVP) [BZB⁺97] est le protocole standard de réservation de ressources pour Internet. Généralement associé à IntServ, il s'agit à l'origine d'un protocole destiné à assurer des réservations sur des arbres *multicast*, le mode point à point en étant un cas particulier. Les demandes de réservation sont à l'initiative du récepteur qui envoie une requête sous la forme d'un message *Path* à destination de la source. Ce paquet de contrôle traverse le réseau, collectant diverses informations. La source, à la réception de ce message, répond par un message *Resv* qui, au fur et à mesure de son trajet vers la destination, va être soumis à un contrôle d'admission dans chaque routeur traversé afin de s'assurer de la disponibilité des ressources. En cas de réussite, ce message provoque une réservation des ressources avant de continuer sa route. Il ne s'agit pas d'un protocole de routage par inondation malgré son fonctionnement de type requête-réponse. RSVP emprunte les routes définies par le protocole de routage en vigueur dans le réseau et n'est qu'un protocole de mise en place et de suppression de réservations. Il est particulièrement adapté pour les applications de type voix.

Dans des réseaux radio multisauts disposant d'un mécanisme de multiplexage temporel des communications (*Time Division Multiple Access* — TDMA), c'est-à-dire d'un réseau dans lesquels la bande passante du canal radio est divisée en unités de temps, de nombreux travaux semblables ont été publiés [CGT97, LL99, HL00, LL00, Lin01, GS02, ZC02]. La plupart de ces propositions supposent qu'un mécanisme de multiplexage de codes (*Code Division Multiple Access* – CDMA) est opérationnel dans le réseau et qu'un processus d'allocation de codes orthogonaux est déployé. Ainsi, deux communications simultanées peuvent être réalisées sans risque de collisions à moins que deux émetteurs n'envoient en même temps une trame au même récepteur. Dans ce cas, le récepteur concerné est tout de même supposé capable de décoder la première trame reçue. Si ce type de collisions ne conduit pas à la perte des deux trames incriminées, elles ne sont pas souhaitables, aussi faudra-t-il s'assurer que deux transmissions voisines n'ont pas lieu simultanément. Le problème consistant à déterminer une allocation d'unités de temps maximale respectant cette contrainte sur une route donnée, et ce tout en calculant la bande passante disponible sur cette route peut être réduit en un problème de satisfaisabilité, connu pour être NP-complet. Plusieurs mécanismes de découverte de routes disposant d'une bande passante disponible suffisante et de réservation des unités de temps correspondantes sont proposés, adoptant une approche réactive ou proactive.

Toutefois, ce type de mécanisme ne peut pas aisément être adapté à des réseaux *ad hoc* reposant sur une couche MAC de type CSMA/CA. En effet, les systèmes TDMA sont des systèmes synchrones. L'évaluation de la bande passante disponible au niveau d'un mobile y est équivalente à un problème d'intersection d'ensembles et l'ordonnancement des trames peut y être planifié de façon distribuée. Le comportement d'un protocole asynchrone tel que IEEE 802.11 est nettement moins prévisible. Adapter un mécanisme distribué de division de la bande passante en unités de temps nécessite dès lors la présence d'une horloge globale au niveau du réseau.

Dans [LG97], Lin et Gerla proposent un mécanisme de réservation élémentaire de bande passante adapté aux réseaux CSMA/CA appelé *Multiple Access Collision Avoidance with Piggyback reservation* (MACA/PR). Le but de ce protocole n'est pas d'assurer une réservation de bout en bout mais de permettre aux flux soumis à des contraintes temps réel de réserver le canal radio pour une rafale de paquets. Dans cette situation, l'émetteur commence par un échange RTS-CTS puis le récepteur indique dans l'acquittement que l'échange n'est pas terminé. Les nœuds voisins maintiennent des tables contenant l'ensemble des réservations ainsi effectuées dans leurs voisinages. L'émetteur transmet alors dès réception de cet acquittement la trame suivante de l'échange, et ce processus est répété jusqu'à ce que la rafale soit achevée ou jusqu'à la première collision.

Le protocole Insignia [LAZC99, LAZC00, LAC01] proposé par Lee *et al.* définit un ensemble de messages destinés à offrir aux réseaux *ad hoc* un mécanisme de réservation de bande passante léger et réactif aux variations de disponibilité de ressources. Afin de ne pas surcharger le réseau de paquets de contrôle, la plupart des informations nécessaires à l'établissement de routes avec réservation sont contenues dans les paquets de données, sous la forme d'une option IP. Une application désirant réserver un certain volume de bande passante spécifie dans l'en-tête des paquets de données deux niveaux de bande passante, le niveau souhaité dans le meilleur cas et un niveau dégradé. Un algorithme de recherche de route répondant aux critères spécifiés doit alors entrer en jeu afin de déterminer une route admissible. Chaque paquet appartenant à un flux privilégié contiendra cet en-tête spécifique qui pourra être modifiée

par les routeurs tout au long de la route, afin d'avertir la destination de la disponibilité des ressources, à la manière du mécanisme *Explicit Congestion Notification* d'IP. Ainsi, lorsque le débit d'un flux ne peut plus être assuré, la destination est chargée d'avertir la source afin qu'elle prenne les mesures adéquates. Lorsqu'un flux transite à débit réduit dans le réseau, la disponibilité de nouvelles ressources est signalée à la destination qui, encore une fois, avertit la source explicitement. Il s'agit là du seul message en tant que tel généré par ce protocole. Insignia est indépendant du protocole de routage, du mécanisme de contrôle de disponibilité des ressources ainsi que du mécanisme de contrôle d'admission et de réservation de bande passante.

Dans [XG03], Xue et Ganz proposent un mécanisme de réservation de bande passante dont le fonctionnement est similaire au protocole RSVP. Lorsqu'une application désire obtenir un chemin répondant à certains critères sur la bande passante et le délai de bout en bout, elle émet une requête qui se propage à travers le réseau jusqu'à atteindre la destination. Lors du passage dans chaque routeur, un contrôle d'admission est effectué, basé sur les informations sur la charge du réseau transmise par leurs voisins directs. La destination répond à chacune des instances de cette requête par un message traversant le réseau en suivant le chemin inverse. La réservation des ressources n'est effective que lorsque le premier paquet de données emprunte la route sélectionnée par la source parmi les différentes instances lui parvenant. Si, dans le principe, cette approche se rapproche de celle que nous avons sélectionnée pour BRuIT, le contrôle d'admission est ici basé sur une mesure optimiste de l'occupation du canal radio. D'autre part, le délai de bout en bout est considéré comme symétrique, ce qui n'est pas toujours le cas dans ce type de réseaux.

2.2.5 Évaluation des ressources disponibles

Les protocoles de réservation de bande passante, ainsi que les protocoles de routage dans une moindre mesure, ont besoin, afin d'être en mesure d'accepter ou de refuser des requêtes, d'une estimation la plus précise possible de la capacité résiduelle du canal radio compte tenu de la topologie et des trafics transitant sur le réseau. L'intégralité des protocoles présentés dans le paragraphe précédent suppose l'existence d'une couche d'accès au médium capable d'estimer la proportion de bande passante résiduelle. En réalité, l'estimation de la bande passante disponible est un problème difficile, d'une part à cause des interférences pouvant survenir dans un contexte multisaits, et d'autre part à cause des variations de topologie liées à la mobilité des nœuds et des variations de charge du réseau liées à l'apparition et à la disparition des flux. Il n'y a aucun lien évident entre bande passante disponible instantanée et bande passante disponible en moyenne.

Shah *et al.*, dans [SCN03], présentent un mécanisme d'estimation de la bande passante disponible pour des réseaux basés sur le protocole IEEE 802.11. Ce mécanisme repose sur le fait que le débit auquel est envoyé un paquet peut être mesuré en calculant le rapport entre la taille de la charge utile de la trame et la différence entre la date à laquelle l'émetteur a reçu l'acquittement correspondant et la date de début de la transmission de cette trame. En normalisant ce débit par rapport à une taille de paquet de référence, il est possible d'estimer la bande passante d'un lien. D'autre part, les auteurs de ce mécanisme supposent implicitement que le degré de contention de tous les nœuds du réseau est identique, ce qui n'est *a priori* pas le cas dans un réseau *ad hoc*.

Dans [KGL01], Kazantzidis *et al.* relient la bande passante disponible sur un lien au débit émis sur ce lien en l'exprimant comme $(1 - u)$ fois ce débit, $u \in [0, 1]$ mesurant l'utilisation du lien. Le débit est mesuré sur une fenêtre de paquets empruntant ce chemin, et est exprimé comme la moyenne de la taille des paquets divisée par le temps nécessaire à leur transmission. L'utilisation du lien quant à elle est le quotient du temps durant lequel le médium a été perçu comme étant libre par l'émetteur et du temps nécessaire à l'acheminement de tous les paquets de la fenêtre de mesure. Dans [KG02], Kazantzidis et Gerla proposent d'utiliser comme mesure de l'utilisation du lien le niveau de remplissage de la file d'attente de niveau MAC. Cette métrique semble cependant pouvoir aisément être faussée. En effet, en ne considérant qu'un couple émetteur et récepteur, si lorsque la file d'attente est à demi pleine l'émetteur commence à envoyer des paquets exactement au débit supporté par le canal, la bande passante disponible sur le lien sera nulle alors que la file ne sera pas remplie.

Ces mécanismes d'estimation ne permettent qu'une évaluation *a posteriori* de la capacité des liens. L'apparition de nouveaux trafics invalidera momentanément l'estimation réalisée. Il est donc nécessaire de déterminer la période d'échantillonnage permettant de ne pas considérer les variations instantanées de la capacité des liens tout en conservant une certaine réactivité et il s'agit là d'un problème ardu.

2.2.6 Synthèse

En résumé, un grand nombre de solutions ont été proposées pour garantir une certaine qualité de service pour les réseaux *ad hoc*. Le domaine de la différenciation de services a reçu une attention toute particulière, du fait du succès de cette philosophie dans les réseaux filaires en comparaison des mécanismes offrant des garanties fortes. La norme IEEE 802.11e devrait bientôt voir le jour, offrant une solution de différenciation dans les réseaux locaux sans fil. Les domaines du routage avec qualité de service et des mécanismes de réservation, souvent liés, n'ont pas connu le même succès. Bon nombre de propositions ont été publiées cependant, mais beaucoup de contraintes non encore maîtrisées liées aux spécificités des réseaux *ad hoc* doivent être appréhendées afin d'aboutir à une solution réellement satisfaisante. Enfin, l'évolution du domaine de la qualité de service dans ces réseaux ne permet pas encore de définir des modèles réellement adaptés. L'évaluation des ressources disponibles dans les réseaux *ad hoc* est le domaine qui nécessite le plus de travail. Peu de propositions ont vu le jour car il est difficile de prendre en compte l'intégralité des contraintes limitant la bande passante disponible. Une telle évaluation suppose une parfaite compréhension des phénomènes radio et de leur impact au niveau des protocoles de niveaux supérieurs.

Dans la section suivante, nous présentons BRuIT, un protocole chargé d'assurer un routage avec qualité de service et un mécanisme de mise en place de réservations dans les réseaux *ad hoc*. BRuIT constitue une première étape vers un protocole de niveau routage prenant en compte les caractéristiques des couches protocolaires sous-jacentes, et en particulier le mode de partage du médium inhérent aux réseaux radios.

2.3 Le protocole BRuIT

[BGLV00] puis [BCGLV01] présentent une étude de la complexité d'un mécanisme d'allocation de bande passante cherchant à maximiser l'utilisation du réseau sur un médium partagé de la sorte. Considérons un ensemble de mobiles $(m_i)_{i \in [1, n]}$ disposant de bande passantes résiduelles $(b_i)_{i \in [1, n]}$ *a priori* différentes. Chaque mobile m_i désire utiliser un volume de bande passante r_i . On notera $N(m_i)$ l'ensemble des mobiles avec lesquels m_i partage le médium, c'est-à-dire qu'un mobile appartenant à cet ensemble ne peut émettre une trame simultanément à m_i . Afin de modéliser le problème de l'allocation de bande passante, définissons un ensemble de variables binaires $(x_i)_{i \in [1, n]}$. x_i sera égal à 1 si l'on accepte la requête du mobile m_i et égal à 0 sinon. S'assurer du respect des capacités des différents nœuds se traduit par l'ensemble de contraintes suivant :

$$\forall i \in [1, n], \quad \sum_{j \in N(m_i) \cup m_i} x_j \cdot r_j \leq b_i.$$

Il est dès lors possible de déterminer l'ensemble des requêtes maximisant l'utilisation du réseau, c'est-à-dire la bande passante totale allouée :

$$\max \sum_{i=1}^n r_i \cdot x_i, \text{ sous contraintes } \forall i \in [1, n], \quad \sum_{j \in N(m_i) \cup m_i} x_j \cdot r_j \leq b_i.$$

Il est aussi envisageable de maximiser le nombre de requêtes acceptées :

$$\max \sum_{i=1}^n x_i, \text{ sous contraintes } \forall i \in [1, n], \quad \sum_{j \in N(m_i) \cup m_i} x_j \cdot r_j \leq b_i.$$

Ces deux problèmes sont des programmes linéaires en $\{0,1\}$ (*Maximum bounded $\{0,1\}$ -linear programs*) et peuvent se réduire en un problème de sac à dos multidimensionnel en $\{0,1\}$. Or, dans [BS92], Berman et Schnitger montrent que les problèmes de programmation linéaire en $\{0,1\}$ appartiennent à la classe des problèmes d'optimisation non déterministes polynomiaux dont la fonction objective est bornée par un polynôme fonction du nombre de variables considérées (NPO-PB). Le problème de décision associé est *a fortiori* NP-complet. En d'autres termes, déterminer de façon centralisée, à partir d'un ensemble de requêtes de bande passante formulées par des mobiles appartenant à un réseau *ad hoc*, quelles requêtes accepter et quelles requêtes rejeter de manière à maximiser soit l'utilisation globale du réseau, soit le nombre de requêtes honorées est un problème NP-complet.

Notons A^* une solution optimale à ce problème et A une solution obtenue au moyen d'une heuristique gloutonne consistant à accepter les requêtes par ordre décroissant de volume de bande passante. Dans ce cas, l'algorithme glouton conduit à une allocation telle que $\sum_{i \in A^*} r_i \leq 2 \cdot \sum_{i \in A} r_i \cdot \deg(i)$, où $\deg(i)$ représente le nombre de mobiles gênés par la requête i .

Cette étude ne considère cependant qu'une partie des contraintes, c'est-à-dire les exclusions mutuelles en émission entre différents mobiles. Dans [GJM04, GJM03], Georgiadis *et al.* étudient la complexité d'un schéma de réservation de bande passante en présence de brouilleurs. Cette fois, on considère qu'autour de chaque mobile, existe une zone de rayon H_I dans laquelle aucun nœud ne doit être en train de recevoir une trame lorsque le mobile au centre de la zone émet une trame. Réciproquement, aucun mobile ne doit émettre une trame dans cette même zone lorsque le mobile central est en réception. Les auteurs montrent, en réduisant ce problème à un problème de chemin avec paire interdites, que la simple détermination de l'existence d'un chemin d'une source à une destination répondant à une exigence en terme de bande passante disponible est un problème NP-complet, et ce même en ne considérant qu'un partage du médium entre voisins directs, sans interférences provenant de mobiles hors de portée de communication. Lorsque l'on considère un rayon d'interférence supérieur à la portée de communication, ce problème reste bien sûr NP-complet et la détermination d'un plus court chemin sous ces conditions est un problème d'optimisation NP-complet NPO-PB-complet.

Si le problème de réservation de bande passante dans des réseaux *ad hoc* est un problème complexe, il est cependant envisageable de proposer un protocole distribué offrant un certain niveau de garanties. Du fait de la mobilité des nœuds et de la versatilité du médium radio, il serait illusoire d'espérer fournir aux applications des garanties fortes et durables. Le protocole BRuIT (*Bandwidth Reservation under Interferences influence*) présenté ci-après constitue un premier pas vers un protocole distribué de réservation de bande passante. Son but est de fournir un contrôle d'admission le plus réaliste possible afin de donner aux applications des indications sur la disponibilité des ressources et donner un sens aux réservations de bande passante.

La prise en compte de l'occupation du réseau dans un voisinage étendu permet aux mobiles de limiter l'acceptation de réservations qui ne pourraient être honorées dans le réseau et qui rendent le partage du médium imprévisible. Nous adoptons ici une approche *a priori* afin d'éviter l'apparition de congestions dans le réseau alors que la plupart des solutions actuelles reposent sur une approche *a posteriori* traitant les congestions une fois qu'elles sont apparues. Il n'est pas exclu de fusionner les deux approches afin de tirer parti des avantages de chacune des deux philosophies.

2.3.1 Une approche réactive

L'utilisation de protocoles de routage à la demande est en général considérée comme présentant un coût d'établissement de route trop élevé pour des réseaux tels qu'Internet. Toutefois, dans des réseaux locaux bénéficiant d'un mode de transmission en diffusion locale tels que des réseaux *ad hoc*, leur utilisation est tout à fait envisageable. Ces protocoles réactifs ne sont bien évidemment pas adaptés à toutes les situations, et l'opposition entre cette philosophie et la philosophie proactive, plus proche du mode de fonctionnement d'Internet actuel, fait couler beaucoup d'encre. Toutefois, nous cherchons ici à fournir des garanties à des applications, indépendamment de leur source et de leur destination, à la manière de RSVP. Deux applications provenant de la même source et à destination du même mobile pourront sélectionner deux chemins différents en fonction de la disponibilité des ressources. L'établissement de tels circuits virtuels est proche de l'établissement de routes de manière réactive, c'est pourquoi BRuIT adopte une telle approche.

Lorsqu'une application souhaite obtenir une route répondant à une exigence de qualité de service, elle émet une requête de réservation, dont le format exact est décrit en annexe A.2 et expliqué par la suite. Cette requête va se propager, de proche en proche, inondant le réseau jusqu'à atteindre la destination. À chaque fois qu'un tel message parvient à un mobile, celui-ci effectue un contrôle d'admission, dont les détails constituent la particularité de BRuIT et seront explicités en section 2.3.2. Ce processus de contrôle d'admission a pour but de s'assurer que l'introduction d'un nouveau trafic dans le réseau dont le volume correspond à la requête de bande passante formulée ne provoque aucun dépassement de la capacité du médium. Si cette condition est vérifiée, la requête est réémise et le nœud routeur potentiel conserve les caractéristiques du flux, adresses source et destination, numéro de flux et bande passante demandée. Il conserve par ailleurs l'adresse du nœud duquel il a reçu la requête. Si les ressources ne sont pas disponibles, la requête est simplement détruite.

Afin d'éviter les cycles dans le routage, une requête n'est examinée par un mobile que lors de sa première réception. Un numéro de séquence déterminé de façon unique pour un émetteur donné par incrément permet, en conjonction avec l'adresse de l'émetteur de la requête, de détecter les demandes dupliquées. Examiner de multiples instances de la même requête peut cependant être envisagé en s'assurant que les routes ainsi formées ne comportent pas de cycles.

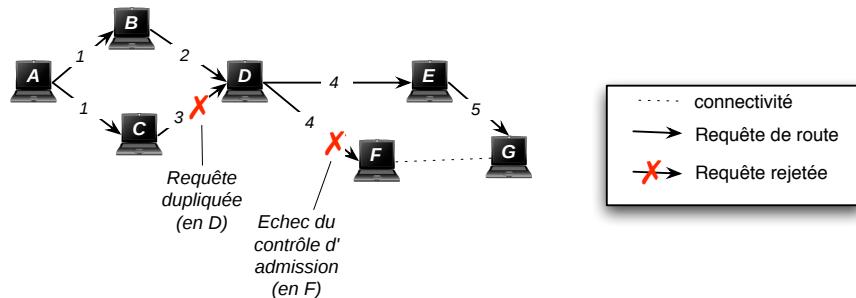


FIG. 2.1 – Exemple de propagation de requête de route

Par exemple, sur la figure 2.1, le mobile *A* désire rechercher une route à destination de *G*, il envoie une requête qui atteint simultanément *B* et *C*. Chacun va alors la réémettre à son tour, et *D* ne réémettra que la première instance qu'il aura reçue de cette requête. Le contrôle d'admission pourra échouer par exemple au niveau de *F* pour faute de disponibilité de ressources. *F* ne retransmettra donc pas ce message et finalement, la demande parviendra à *G* par l'intermédiaire de *E*.

Ainsi, si l'occupation du réseau le permet, une ou plusieurs instances de la requête atteindront leur destination, déterminant une ou plusieurs routes admissibles selon les critères de qualité de service spécifiés. Dans notre approche, la première requête à parvenir à la destination définira la route à utiliser. Plusieurs arguments semblent en effet aller dans ce sens plutôt que dans celui de l'attente de multiples requêtes permettant de sélectionner parmi l'ensemble des routes découvertes la meilleure. En effet, afin que plusieurs requêtes parviennent à destination, plusieurs chemins disjoints admissibles doivent être déterminés, puisque chaque routeur ne retransmet qu'au plus une fois chaque requête. La probabilité de déterminer plusieurs chemins admissibles est donc réduite. De plus, la requête parvenant la première à destination est potentiellement la requête correspondant au chemin le plus court en terme de délai de bout en bout. L'examen d'un seul paquet de contrôle ne permet toutefois en aucun cas de s'assurer ce fait. Il ne s'agit là que d'une considération *a priori*. Enfin, une réponse rapide à une requête de bande passante permet de réduire la probabilité que les ressources ne soient plus disponibles lors de l'envoi d'une réponse confirmant la réservation des ressources.

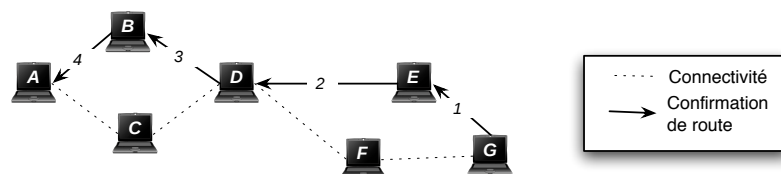


FIG. 2.2 – Exemple de transmission de confirmation de route

La bande passante n'est effectivement réservée que lors du passage dans les routeurs du message de confirmation émis en réponse à la requête par la destination. Le format de ce message est explicité en annexe A.3. Il n'est pas possible d'effectuer une réservation lors du passage de la requête initiale car les nœuds ne connaissent pas *a priori* l'état de l'ensemble des mobiles en aval sur la route. La destination émet donc, après avoir effectué, elle aussi un contrôle d'admission, un message de confirmation. Ce message, à l'inverse du message de recherche de route, est émis en mode point à point puisqu'il ne s'agit plus d'explorer le réseau. Cette confirmation emprunte la route inverse, comme le montre l'exemple de la figure 2.2, provoquant la réservation effective des ressources après une nouvelle vérification de la disponibilité des ressources. En cas de réussite, le routeur conserve les caractéristiques du flux ainsi que

l'adresse des nœuds voisins immédiats sur la route. Lorsque cette confirmation parvient à destination, la route est établie et les ressources sont réservées.

La conservation dans chaque routeur des adresses du mobile en amont et du mobile en aval sur la route permet d'utiliser la route déterminée dans les deux sens pour peu que les applications marquent les paquets correspondant à ce flux en accord avec l'identifiant de la route. Ce procédé permet indifféremment à une application qui sera la source d'un flux ou à une application qui sera la destination d'un flux de réserver des ressources. D'autre part, il est possible d'utiliser la route ainsi déterminée dans les deux sens simultanément, ce qui permet aux acquittements TCP de parcourir le réseau en bénéficiant du même niveau de service que le flux auxquelles elles correspondent. Enfin, il est possible de mettre en œuvre un mécanisme de surveillance de l'état de la route basé sur la transmission régulière de trames dans le sens inverse du flux de données apportant les informations collectées le long de la route à la source du flux.

Une fois que le flux de données est transmis, la source comme la destination peut libérer explicitement les ressources au moyen d'un message de contrôle particulier. Cependant, une libération explicite des ressources n'est pas toujours possible par exemple en cas de panne de l'application, ou lorsque la mobilité des routeurs empêche le message de libérer les ressources de bout en bout. Aussi les routes et les réservations qui leur sont associées expirent au bout d'un délai lorsque aucun paquet de donnée n'emprunte la route. Afin de ne pas pénaliser les applications ayant un profil d'émission irrégulier, la source pourra, de façon optionnelle, émettre régulièrement des paquets vides qui parcourront la route à la manière du processus de maintien des sessions TCP (*keep-alive*).

Lorsqu'une route disparaît du fait de la mobilité d'un routeur, les nœuds en aval du point de cassure n'ont plus de paquets à transmettre, qu'il s'agisse de paquets de données effectifs ou de paquets de rafraîchissement de route. En revanche, les routeurs en amont du point de cassure continueront à retransmettre les paquets du flux, ceux-ci étant perdus au niveau du dernier routeur avant le point de cassure. C'est ce dernier routeur qui pourra avertir la source de la cassure par l'envoi d'un message explicite, dont le format exact est explicité en annexe A.4 qui, traversant les routeurs, provoquera la libération des ressources. À la réception d'un tel message, la source devra réeffectuer une recherche de route. Conserver une route de secours ne semble pas approprié du fait de la mobilité des nœuds et des variations de disponibilité de ressources.

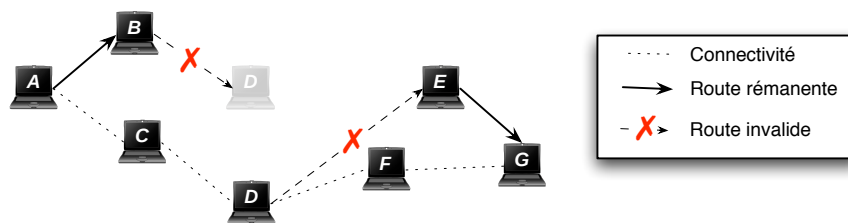


FIG. 2.3 – Exemple de cassure de route

Par exemple, dans la situation représentée en figure 2.3, le mobile *D* se déplace jusqu'à sortir de la zone de couverture de *B* et *E*. Le mobile *B* ne recevant plus d'acquittements de la part de *D* constatera le départ de ce dernier et en avertira *A*. Le mobile *E* ne recevra plus de paquets à retransmettre et la route expirera bientôt. Enfin, le mobile *D* sera dans les deux situations à la fois et éliminera la route invalide.

Le processus de réservation ainsi défini est semblable au protocole RSVP et, par là même, souffrira sans doute des mêmes limitations particulièrement concernant les problèmes liés au passage à l'échelle, du fait de la nécessité de maintenir des états pour chaque flux dans chaque routeur. BRuIT ne sera sans doute pas adapté à des réseaux *ad hoc* à large échelle tels que, par exemple, ceux qui sont définis par le projet Terminodes². Toutefois, BRuIT peut être utilisé tel quel dans des réseaux *ad hoc* de taille moyenne. À titre indicatif, les tables de filtrage d'un mobile fonctionnant sous le système Linux (avec NetFilter disponible à partir de la version 2.4) peuvent distinguer jusqu'à 256 flux en ne se basant que sur le champ TOS (*type of Service*) de l'entête IP. Ce nombre est à multiplier par le nombre de couples source-destination possibles dans le réseau. Par ailleurs, il est tout à fait envisageable de n'utiliser ce

²<http://www.terminodes.org>

protocole que pour le routage des trafics multimédias, en conjonction avec un protocole de routage *ad hoc*.

2.3.2 Contrôle d'admission

Tout mécanisme de réservation repose sur un mécanisme de contrôle d'admission qui se doit d'être le plus fiable possible. En effet, il est indispensable, lorsqu'on souhaite offrir des garanties à des flux d'évaluer de façon précise la disponibilité des ressources. Dans les réseaux filaires, mesurer la disponibilité des ressources revient à évaluer la capacité résiduelle des liens du réseau.

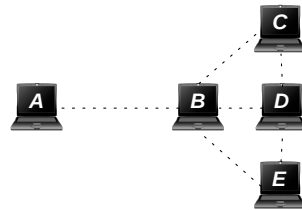


FIG. 2.4 – Réseau multi-sauts simple

Le médium radio, en revanche, est un médium partagé. Il n'est pas possible d'isoler des liens. Lorsqu'une trame est émise, elle est reçue par l'ensemble des mobiles dans une certaine zone géographique, la couche MAC déterminant alors si la trame doit être transmise à la couche IP ou tout simplement ignorée. Ce mode de partage se rapproche d'une certaine façon de celui d'un bus, sans pour autant être identique. En effet, si les émissions d'un mobile atteignent l'ensemble de ses voisins, elles n'atteignent pas l'ensemble des nœuds du réseau. Un nœud peut avoir deux voisins ne pouvant communiquer entre eux. Par exemple, si sur la figure 2.4 les traits représentent la connectivité entre les mobiles et si l'on suppose que les mobiles ne partagent l'accès au médium qu'avec leurs voisins directs, *A* doit partager le médium avec *B* alors que *B* est lui en concurrence avec tous les nœuds du réseau. Considérer le lien (*A,B*) n'a plus réellement de sens. Par exemple, le délai du lien étant fonction du temps d'accès au médium de l'émetteur, le délai de *A* vers *B* sera *a priori* très différent du délai de *B* vers *A*. De même, la bande passante disponible est fonction du niveau de contention et est donc fonction du mobile émetteur plutôt que du lien considéré.

Les réseaux filaires sont fidèlement représentés par des graphes, chaque sommet représentant un routeur et chaque arête un lien de communication. Associer une capacité aux arêtes permet, par une simple opération, de déterminer si une requête parvenant à un routeur peut être acceptée et, le cas échéant, sur quelle interface la réémettre. L'interdépendance entre les liens introduite par l'utilisation du médium radio dans les réseaux *ad hoc* rend cette comparaison complexe. Il ne s'agit plus de considérer les capacités des arêtes du graphe mais plutôt d'associer une capacité à chaque clique dans le graphe représentant la connectivité dans un tel réseau et de considérer ces cliques lors du contrôle d'admission. Ceci implique que la décision d'accepter ou de refuser une requête peut difficilement être prise par un nœud sans information sur l'état des mobiles dans son voisinage.

Toutefois, même une vision en cliques sur le graphe de connectivité peut parfois se révéler insuffisante. Nous avons vu au chapitre 1, avec le problème des trois paires, que deux mobiles qui ne peuvent pas communiquer entre eux ne peuvent parfois pas émettre simultanément du fait du mécanisme de détection de porteuse.

Dans la suite de ce document, nous ferons deux approximations. La première concerne la taille de cette zone de détection de porteuse. Un mobile considère le canal radio comme étant occupé si le niveau de signal sur ce canal est supérieur à un certain seuil. Or, même si on ne considère qu'une simple source, le niveau de signal présente des discontinuités dans l'espace, particulièrement en intérieur compte tenu des phénomènes de réflexion, de diffraction et des combinaisons de plusieurs instances du même signal. La taille réelle de la zone de détection de porteuse sera donc variable, parfois dépassant à peine la zone de communication, parfois allant deux, trois voire quatre fois plus loin. Afin de ne pas introduire une trop grande complexité dans le mécanisme d'estimation de la capacité du canal, nous considérerons que le rayon de la zone de détection de porteuse est double du rayon de la zone de communication. Nous considérerons par ailleurs que si les émissions d'un mobile bloquent le mécanisme de détection de

porteuse d'un autre, la réciproque est vraie. Ces approximations semblent réalistes au regard des résultats d'expérimentations présentés dans [Dho02].

Il serait possible de définir dans chacune de ces zones de partage, c'est-à-dire au niveau de chaque mobile, un ordonnancement précis des différentes trames, de manière à provoquer un partage du médium correspondant aux requêtes des flux privilégiés. Ce mode de fonctionnement correspondrait à réserver des unités de temps, à la manière des systèmes TDMA. Ce type de problème peut être ramené à un problème de coloriage de graphes. Toutefois, la maintenance d'un tel ordonnancement est peu aisée. En effet, il est possible, dans un réseau statique, de concevoir un protocole autorisant chaque mobile à émettre une trame tour à tour, l'immobilité de la topologie permettant de planifier un ordre dans les transmissions tout en tenant compte des problèmes de réutilisation spatiale. Un système si parfaitement ordonné est fragile en environnement mobile, l'ordre d'émission des trames devant être recalculé pour le réseau entier à chaque changement de topologie.

Il semble donc plus aisé de se reposer sur le protocole d'accès au médium qui, s'il est loin d'être parfait dans un cadre *ad hoc*, a un comportement prévisible tant que la capacité du canal radio n'est pas dépassée. Garantir du mieux que l'on peut des réservations de bande passante nécessite par conséquent de minimiser la fréquence d'apparition des situations de dépassement de la capacité du médium. Le contrôle d'admission que nous réaliserons devra simplement s'assurer, que nulle part dans le réseau, la capacité du médium n'est dépassée. En d'autres termes, pour chaque mobile du réseau, il convient de s'assurer que la somme des émissions de chaque nœud et des nœuds appartenant à sa zone de détection de porteuse ne dépasse pas la capacité du médium.

Compte tenu des approximations que nous avons faites, il s'agit d'apporter à chaque mobile la connaissance du volume de trafic émis dans son voisinage à deux sauts. Dans ce but, chaque mobile transmet régulièrement à tous ses voisins un paquet de contrôle que nous appellerons paquet *Hello*, contenant la liste de ses voisins directs ainsi que le volume de bande passante que chacun de ces voisins utilisera et la bande passante disponible au niveau de chacun de ces voisins, cette bande passante disponible étant la différence entre la capacité du canal radio et la somme des trafics émis dans la zone de détection de porteuse du mobile. Le format exact de ce paquet est décrit en annexe A.1. Ces paquets permettent de reconstruire la topologie à deux sauts autour du mobile et ces informations peuvent être utilisées à des fins de routage.

Les requêtes de réservation de bande passante pourront alors être examinées et leur acceptation sera conditionnée au fait que le mobile routeur effectuant ce contrôle d'admission puisse réémettre ce trafic sans perturber aucun voisin. En d'autres termes, il devra s'assurer d'une part que le débit de ce trafic additionné aux débits des trafics présents dans la zone de détection de porteuse n'excède pas la capacité du médium et, d'autre part, que le débit de ce nouveau trafic ne provoque aucun dépassement de capacité au niveau d'un autre mobile et ce dans l'intégralité de son voisinage à deux sauts.

Réservations et routage multisauts

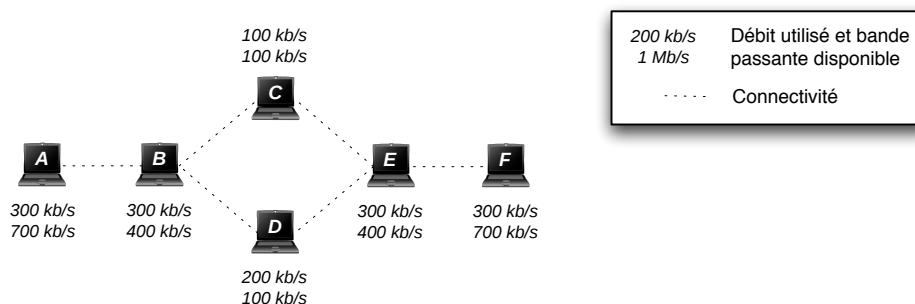


FIG. 2.5 – Un exemple de topologie et de l'état de l'occupation du médium (capacité du médium : 1,6 Mbit/s ; zone de détection de porteuse : 2 sauts)

Considérons la situation décrite par la figure 2.5. Un certain nombre de réservations ont été acceptées, conduisant à la situation représentée. Les bandes passantes disponibles sont calculées en supposant que le débit utile du canal dédié au trafic privilégié est de 1,6 Mbit/s et que la zone de détection de porteuse

a un rayon double de la zone de communication. Supposons que A désire transmettre, à destination de F un flux de 50 kbit/s. La requête de route se propagera dans tout le réseau et les différents contrôles d'admission réussiront puisque tous les mobiles disposent d'au moins 50 kbit/s de bande passante disponible.

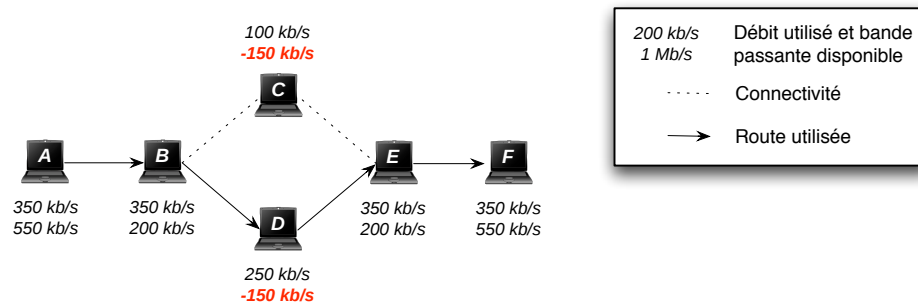


FIG. 2.6 – Dépassement de capacité dû à un contrôle d'admission incorrect (capacité du médium : 1,6 Mbit/s ; zone de détection de porteuse : 2 sauts)

Toutefois, l'introduction du flux dans le réseau conduit à la situation représentée en figure 2.6. En effet, effectuer une simple comparaison du débit demandé par l'application et de la bande passante disponible ne prend pas en compte le fait que ce flux sera routé par des mobiles voisins du nœud effectuant le contrôle d'admission.

Ce problème de sous-estimation de l'impact d'un flux peut être résolu de plusieurs manières. Il est possible, lors du contrôle d'admission, de considérer la distance séparant le mobile concerné de la source et de la destination. Si le mobile n'a pas dans sa liste de voisins directs le mobile source, le flux devra lui être transmis par deux voisins avant de lui parvenir. De même, si la destination n'est pas dans sa liste de voisins à deux sauts, il sait que le flux devra être retransmis deux fois dans son voisinage avant de parvenir à destination.

BRuIT prend en considération ce phénomène en multipliant, lors du contrôle d'admission, la bande passante demandée par le nombre de réémissions qu'il subira. Le nombre de ces réémissions en amont sur la route peut être déterminé en examinant le contenu du paquet de requête et en le comparant avec la liste des voisins maintenue au sein du nœud concerné. Le nombre de réémissions en aval ne peut qu'être estimé lors de l'examen de la requête en recherchant dans cette liste de voisins la destination. Si elle n'y figure pas, le routeur saura qu'il subira probablement deux réémissions supplémentaires dans son voisinage étendu. Le nombre de réémissions réellement subies pourra être déterminé lors du passage de la confirmation de réservation. Prendre en considération ce phénomène est nécessaire, mais ne suffit cependant pas, comme nous le verrons au paragraphe suivant.

Débits multiples et échelle commune

La norme IEEE 802.11 autorise la transmission de données à différents débits. En présence d'une station de base comme en mode *ad hoc*, le débit utilisé est déterminé en fonction de la qualité du lien reliant le nœud et la station. Ce mécanisme de sélection automatique du débit (*Auto-rate fallback* – ARF) engendre des problèmes de performances tels que celui signalé dans [HRBSD03]. Lorsqu'il s'agit d'offrir des réservations de bande passante, ce type de phénomène pourra modifier dynamiquement la capacité du canal et invalider des réservations effectuées. Il n'est pas possible de prévoir les mouvements des mobiles et un protocole doit pouvoir réagir à ce type de dégradation de la qualité du canal. Le mécanisme d'adaptation permettant de gérer les variations dans la capacité du canal que nous avons choisi sera présenté en section 2.3.4.

Considérons pour l'instant une image du réseau à un instant fixé. Les capacités des différents liens sont *a priori* hétérogènes. Les applications ne sont pas, et n'ont pas à être, conscientes du débit réel du lien radio entre le mobile émetteur et le prochain saut sur la route. Lorsqu'une application indique qu'elle désire utiliser un certain volume de bande passante, il s'agit du volume de données qu'elle émettra par unité de temps et c'est aux couches sous-jacentes de traduire cette requête en terme de proportion de la capacité du canal.

Nous supposons qu'il est possible de connaître le débit utilisé par la couche MAC pour joindre tel ou tel voisin. Nous supposons par ailleurs que ce débit est symétrique, c'est-à-dire que si A choisit d'émettre les trames à destination de B à un débit de 2 Mbit/s, B émet aussi les trames à destination de A au même débit. Afin d'effectuer le contrôle d'admission, tous les mobiles utiliseront la même échelle correspondant au débit maximal du canal. Sur la figure 2.7, par exemple, un flux de 100 kbit/s occupera le médium pour passer de B à D durant un temps double de celui qui est requis pour passer de A à B . Le poids de cette requête pour le nœud D est donc au minimum de quatre fois le débit demandé par l'application. Si la requête était parvenue par l'intermédiaire du nœud C , le coût de ce trafic aurait été de cinq fois ce débit.

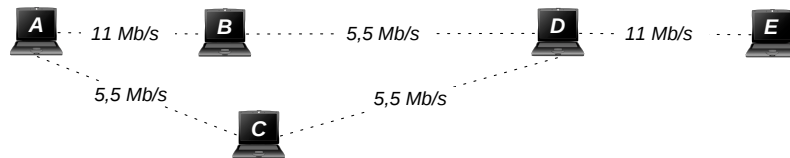


FIG. 2.7 – Un exemple de topologie présentant des liens hétérogènes

Ceci signifie qu'un nœud, afin d'effectuer un contrôle d'admission réaliste doit connaître non seulement le nombre de réémissions dans son voisinage induites par ce nouveau trafic mais aussi le débit de ces réémissions. La solution la plus simple à ce problème consiste à inclure aux paquets d'établissement de route l'ensemble des caractéristiques de la route empruntée jusque-là. Ainsi un mobile pourra évaluer le coût de l'introduction du trafic correspondant dans le réseau.

Grâce à la présence de l'ensemble de la route dans les paquets de recherche de routes, il est dès lors possible pour un routeur de sélectionner une route différente entre la source et lui-même si celle-ci présente un coût plus faible. Ici le coût peut être une métrique quelconque comme le délai, pour peu que l'on dispose d'un mécanisme permettant de le mesurer précisément, le nombre de sauts de la route ou encore le nombre de réémissions subies par le nœud routeur.

2.3.3 Limitations de ce protocole

Identifier l'ensemble de ses brouilleurs potentiels représente une nécessité lorsque l'on souhaite effectuer un contrôle d'admission précis. Toutefois, il ne s'agit pas d'une tâche aisée et deux phénomènes principaux, en sus de la mobilité, peuvent limiter la précision du protocole présenté ici. Premièrement, l'identification des brouilleurs potentiels dans BRuIT ne peut s'effectuer que si l'on peut communiquer avec eux ou s'il existe un nœud intermédiaire pour relayer les informations. Deuxièmement, nous ne considérons lors du contrôle d'admission que les blocages du mécanisme de détection de porteuse et l'impact des interférences à la réception, conséquence de la présence d'une multitude de signaux lointains est négligée. Dans ce paragraphe, nous tâcherons de quantifier l'importance de ces deux limitations et examinerons les palliatifs potentiels.

Estimation erronée à cause d'une faible densité du réseau

Le mécanisme d'identification des brouilleurs présenté plus haut n'est efficace que lorsque deux nœuds en exclusion mutuelle pour l'accès au médium sont séparés par un nœud intermédiaire afin de relayer les informations de l'un à l'autre. Par exemple, le problème des trois paires présenté au chapitre 1 ne peut être résolu s'il n'y a pas chaque couple d'émetteurs un nœud transmettant les informations nécessaires.

Considérons un graphe géométrique aléatoire tel que ceux définis dans [Pen03]. Ce type de graphe correspond bien à la représentation d'un réseau *ad hoc* quelconque. Il est en effet formé par le placement aléatoire de nœuds dans une zone géographique fixée. Il existe une arête entre deux nœuds s'ils sont à une distance inférieure à un certain seuil, paramètre du graphe. Nous noterons par la suite $G(n, \rho)$ le graphe formé par n nœuds disposés aléatoirement et uniformément dans un carré unité lorsque la portée des différents nœuds est ρ .

Déterminons, dans un tel graphe, la probabilité pour que le mécanisme de connaissance du voisinage de BRuIT soit inopérant, c'est-à-dire la probabilité pour que deux nœuds soient en zone de détection

de porteuse l'un de l'autre sans être à portée et qu'il n'y a aucun nœud entre eux pour relayer les informations de l'un à l'autre. Si l'on considère la situation représentée en figure 2.8, cela correspond au fait qu'il n'y a aucun nœud dans la zone grisée.

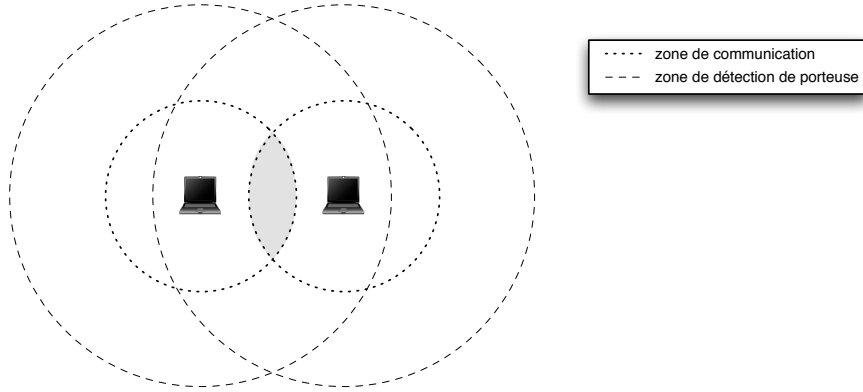


FIG. 2.8 – Voisinage à un saut commun entre deux nœuds

Dans un graphe formé dans un carré unité, si l'on considère une zone quelconque d'aire A , la probabilité pour qu'aucun nœud ne soit dans cette zone est $e^{-n \cdot A}$. Si l'on considère deux nœuds séparés d'une distance $d \in [\rho, 2 \cdot \rho]$, la probabilité qu'il n'y ait aucun nœud entre eux est la probabilité qu'il n'y ait aucun nœud dans l'intersection des deux cercles de rayon ρ centrés sur chacun des deux nœuds. Pour calculer cette aire, il suffit de multiplier par deux l'aire entre la corde définie par les points d'intersection des deux cercles et le périmètre du cercle, soit :

$$\forall d \in [0; 2 \cdot \rho], 2 \cdot \left(\arccos \left(\frac{d}{2 \cdot \rho} \right) \cdot \rho^2 - \frac{d}{2} \cdot \sqrt{\rho^2 - \left(\frac{d}{2} \right)^2} \right).$$

Considérons maintenant l'anneau compris entre les cercles de rayon r et $r + dr$ autour d'un nœud. L'aire de cet anneau est $\pi \cdot (dr^2 + 2 \cdot r \cdot dr)$, soit approximativement $2 \cdot \pi \cdot r \cdot dr$. Les nœuds étant distribués uniformément dans le réseau, il y aura en moyenne $2 \cdot \pi \cdot n \cdot r \cdot dr$ nœuds dans cet anneau.

Dans une zone quelconque d'aire A , la probabilité pour qu'aucun nœud ne se trouve dans cette zone est $e^{-n \cdot A}$. En conséquence, pour un nœud donné, le nombre de nœuds appartenant à sa zone de détection de porteuse, mais pas à sa zone de communication et tels qu'il n'y ait aucun nœud entre les deux pour relayer les informations peut s'exprimer par :

$$p(n, \rho) = \int_{r=\rho}^{2 \cdot \rho} 2 \cdot \pi \cdot n \cdot r \cdot e^{-n \cdot \left(2 \cdot \left(\arccos \left(\frac{d}{2 \cdot \rho} \right) \cdot \rho^2 - \frac{d}{2} \cdot \sqrt{\rho^2 - \left(\frac{d}{2} \right)^2} \right) \right)} \cdot dr.$$

La résolution numérique de cette intégrale, pour différentes valeurs de n et de ρ conduit à la figure 2.9. On constate sur cette figure que ce nombre de brouilleurs non détectés augmente avec la portée de transmission, l'aire de la zone de détection de porteuse croissant en conséquence, et avec la densité du réseau.

Transmettre dans les paquets *Hello* des informations sur le voisinage direct ne permet donc pas de prendre en compte toutes les émissions interférentes lors du contrôle d'admission. Il est alors légitime de se demander si accroître la vision des nœuds en ne transmettant plus uniquement le voisinage direct mais le voisinage à deux ou trois sauts dans chaque paquet *Hello* permettrait d'améliorer la perception qu'ont les mobiles de leur environnement. Il n'est cependant pas possible pour un mobile de savoir si un émetteur situé à trois sauts dans la topologie est un brouilleur potentiel ou non. Des informations de localisation telles que celles qui peuvent être fournies par un système de type GPS ne seront pas d'un grand secours puisqu'il n'y a, en réalité, pas de lien direct entre la qualité d'un lien et la position des mobiles. Cette qualité est fonction de l'environnement. En conséquence, un mobile collectant des informations sur son voisinage à trois ou quatre sauts dans la topologie ne sera pas à même de distinguer les voisins brouilleurs des autres. Il sera dans l'obligation de considérer les émissions de tous ces mobiles, ce qui conduira à une surestimation de la gêne occasionnée par les émetteurs concurrents.

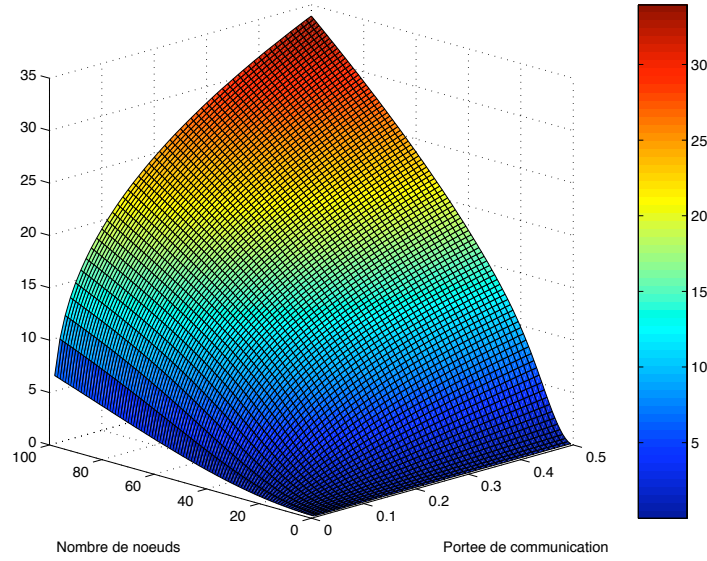


FIG. 2.9 – Nombre de brouilleurs non détectés en fonction de la portée d’émission et du nombre de nœuds dans le réseau (graphes géométriques aléatoires)

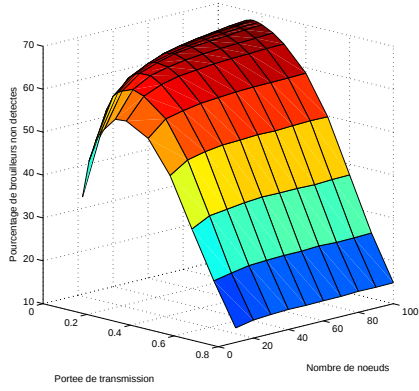
La figure 2.10 présente les résultats de simulations réalisées sur des graphes géométriques aléatoires afin d’évaluer la qualité de la perception qu’ont les mobiles de leur environnement lorsqu’ils considèrent leur voisinage à un, deux, trois et quatre sauts. Les courbes représentées sont la moyenne des résultats obtenus sur mille graphes géométriques aléatoires. Les figures 2.10(a), 2.10(b), 2.10(c) et 2.10(d) représentent le rapport entre le nombre de brouilleurs détectés en considérant le voisinage à k sauts et le nombre réel de brouilleurs. Les figures 2.10(e) et 2.10(f) représentent le rapport entre le nombre de voisins dont on tient compte dans le contrôle d’admission mais qui ne sont pas des brouilleurs et le nombre de brouilleurs réels lorsque l’on considère le voisinage à trois et quatre sauts comme l’ensemble des brouilleurs. Pour une distance inférieure à trois sauts, ce nombre est nul.

En ne considérant que son voisinage direct, jusqu’à 65 % des voisins brouilleurs potentiels sont ignorés. Ce pourcentage varie avec la portée d’émission et avec le nombre de nœuds. Il est maximal pour une portée de 0,3. En considérant le voisinage à deux sauts, ce rapport ne dépasse plus 49 %. Accroître encore la distance de transmission des informations ne conduit pas à une identification des brouilleurs sensiblement meilleure. En revanche, si le nombre de nœuds situés à une distance supérieure au double de la portée de transmission considérés à tort comme des brouilleurs est nul lorsqu’on considère le voisinage à un ou deux sauts, il croît rapidement alors que l’on augmente la distance de transmission des informations. Accroître cette distance représente donc une perte de pertinence vis-à-vis de l’identification des brouilleurs.

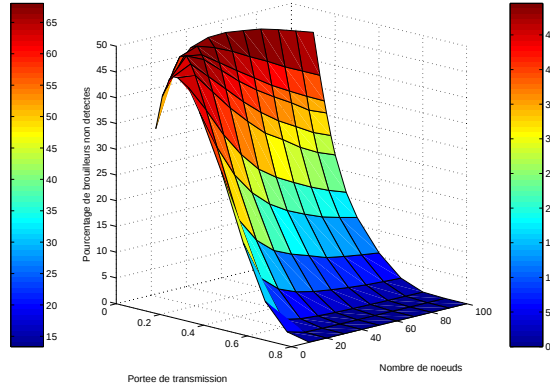
En définitive, considérer comme l’ensemble des brouilleurs potentiels l’ensemble des nœuds situés à deux sauts ne rend pas compte avec exactitude de l’ensemble des nœuds avec qui l’on doit partager le médium. Toutefois, il semble que dans un réseau *ad hoc* quelconque, il s’agisse du meilleur compromis entre sous-évaluation et surévaluation de la capacité résiduelle du médium.

Interférences à la réception

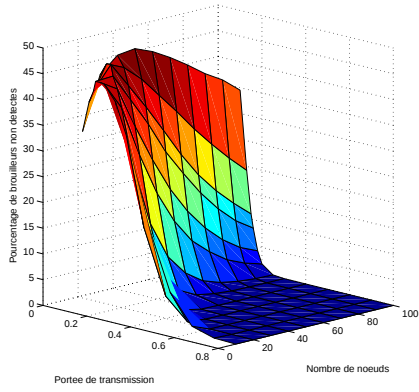
Le mécanisme de contrôle d’admission de BRuIT vise à s’assurer que le volume d’émissions n’excède pas la capacité du médium mais n’offre *a priori* aucune garantie sur la bonne réception des trames. En réception, s’assurer du non-dépassement de la capacité du médium est difficile. En effet, la simple vérification de la disponibilité des ressources n’est pas suffisante, un mauvais ordonnancement des communications pouvant provoquer un grand nombre de collisions, même si la somme des trafics environnants laisse suffisamment de place à la communication en cours pour être reçue. La principale différence entre la perturbation de la détection de porteuse et les interférences à la réception réside dans le fait qu’un récepteur ne peut pas bloquer un brouilleur potentiel pour la durée de la communication. Les mécanismes de requête-réponse (RTS-CTS) permettent de réserver l’utilisation du médium dans une certaine mesure



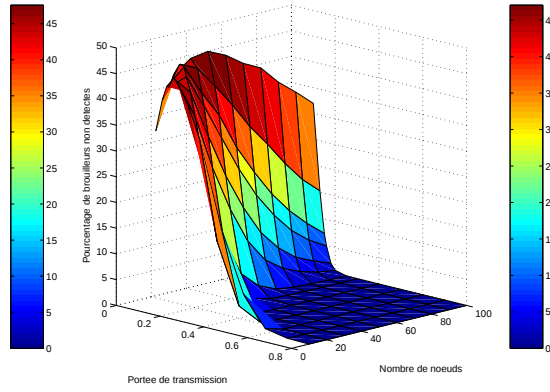
(a) Pourcentage de brouilleurs non détectés en considérant le voisinage à un saut



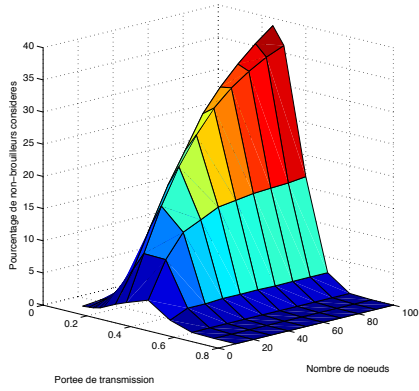
(b) Pourcentage de brouilleurs non détectés en considérant le voisinage à deux sauts



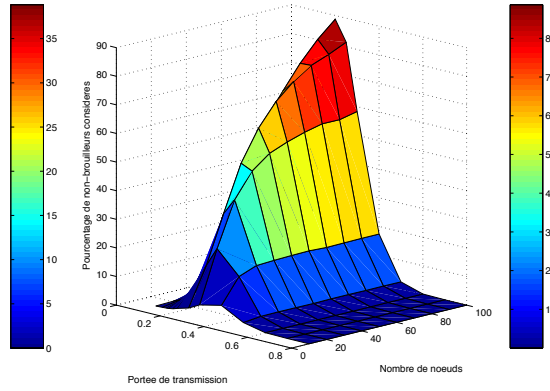
(c) Pourcentage de brouilleurs non détectés en considérant le voisinage à trois sauts



(d) Pourcentage de brouilleurs non détectés en considérant le voisinage à quatre sauts



(e) Rapport entre le nombre de non-brouilleurs considérés et le nombre de brouilleurs réel en considérant le voisinage à trois sauts



(f) Rapport entre le nombre de non-brouilleurs considérés et le nombre de brouilleurs réel en considérant le voisinage à quatre sauts

FIG. 2.10 – Erreurs de perception lorsque les mobiles considèrent leur voisinage à k sauts

au niveau des mobiles à portée de communication du récepteur. Toutefois, la bonne réception d'une trame impose un rapport signal sur bruit pouvant atteindre les 16 dB selon les spécifications de certaines cartes d'interface, ce qui signifie que le signal doit être reçu avec une puissance quarante fois supérieure au bruit engendré par des communications perturbatrices.

Considérons un modèle d'affaiblissement en l'inverse de distance séparant émetteur et récepteur à la puissance α . Supposons que l'on a, pour une émission à une puissance donnée, à une distance de référence d_0 une puissance reçue égale à P_0 . La puissance à une distance r quelconque de l'émetteur peut s'écrire $P(r) = P_0 \cdot d_0^\alpha / r^\alpha$.

Si l'on se place à nouveau dans un graphe géométrique aléatoire pour lequel n nœuds sont disposés de manière uniforme dans un carré unité, entre les cercles de rayons r et $r + dr$ autour d'un récepteur quelconque, nous trouverons en moyenne un nombre de nœuds égal à $2 \cdot \pi \cdot r \cdot n \cdot dr$. La puissance totale rayonnée par les mobiles situés dans une telle zone sera alors :

$$P(r, d + dr) = \frac{2 \cdot \pi \cdot n \cdot P_0 \cdot d_0^\alpha \cdot r \cdot dr}{r^\alpha}.$$

Pour évaluer l'influence cumulée des mobiles situés à une distance supérieure à d , il suffit alors de calculer³ :

$$\int_d^{+\infty} \frac{2 \cdot \pi \cdot n \cdot P_0 \cdot d_0^\alpha \cdot r}{r^\alpha} \cdot dr.$$

Cette intégrale diverge pour $\alpha \leq 2$ et vaut, si $\alpha > 2$:

$$\frac{2 \cdot \pi \cdot n \cdot P_0 \cdot d_0^\alpha}{(\alpha - 2) \cdot d^{\alpha-2}}.$$

Prenons comme point de référence pour la définition de P_0 et d_0 l'émetteur du flux. Nous ne considérons que des brouilleurs plus éloignés que cette distance d_0 car, d'une part les modèles de propagation en $1/d^\alpha$ ne sont pas valides en champ proche et d'autre part, la présence d'un brouilleur plus proche du récepteur que l'émetteur rendrait la communication impossible dans un tel modèle. Le rapport entre la puissance des brouilleurs à une distance supérieure à d et la puissance des brouilleurs à une distance supérieure à d_0 peut s'écrire :

$$\frac{2 \cdot \pi \cdot n \cdot P_0 \cdot d_0^\alpha}{(\alpha - 2) \cdot d^{\alpha-2}} \cdot \frac{(\alpha - 2) \cdot d_0^{\alpha-2}}{2 \cdot \pi \cdot n \cdot P_0 \cdot d_0^\alpha} = \left(\frac{d_0}{d} \right)^{\alpha-2}.$$

La figure 2.11 représente l'évolution de ce rapport en fonction du rapport des distances entre l'émetteur et le récepteur pour différentes valeurs de α . Le fait que l'intégrale permettant d'aboutir à ce résultat diverge pour une valeur de $\alpha = 2$ signifie que dans un modèle d'espace libre parfait, il est théoriquement impossible de négliger l'influence des brouilleurs, et ce quelle que soit leur distance au récepteur. Cette figure montre que plus le coefficient d'affaiblissement est faible, plus l'influence des brouilleurs lointains aura une influence importante. Il est nécessaire de se placer en environnement présentant un fort affaiblissement, par exemple en intérieur, pour pouvoir négliger l'influence des brouilleurs distants de plus de deux ou trois sauts comparativement à l'influence des brouilleurs plus proches.

En conclusion, un nœud sera d'une part soumis, lorsqu'il souhaite émettre, à l'influence des autres mobiles du réseau situés dans sa zone de détection de porteuse et dont il n'aura pas connaissance. D'autre part, en réception le signal pourra être compromis par d'autres transmissions situées à une distance très supérieure à deux sauts radio. Il est difficile, voire impossible, d'identifier précisément l'ensemble des interférences subies par une communication sans fil. En conséquence, il semble important de prévoir un mécanisme permettant de pallier les faiblesses du contrôle d'admission.

³ Ici, nous faisons l'approximation que la portée des mobiles est suffisamment faible devant la taille de la zone considérée pour qu'il soit équivalent d'intégrer la fonction jusqu'aux limites de la zone et jusqu'à l'infini. Si R est le rayon du cercle inscrit au carré délimitant la zone considérée. Une condition nécessaire est que la puissance totale rayonnée par les mobiles situés dans un rayon supérieur à R soit négligeable devant la puissance totale rayonnée par les mobiles situés dans un rayon inférieur à R . Ceci se traduit par $R^{2-\alpha} \ll d^{2-\alpha} - R^{2-\alpha}$. Par exemple, lorsque $\alpha = 3$, pour que le rapport soit inférieur à 10 %, il faut que $R = 11 \cdot d$.

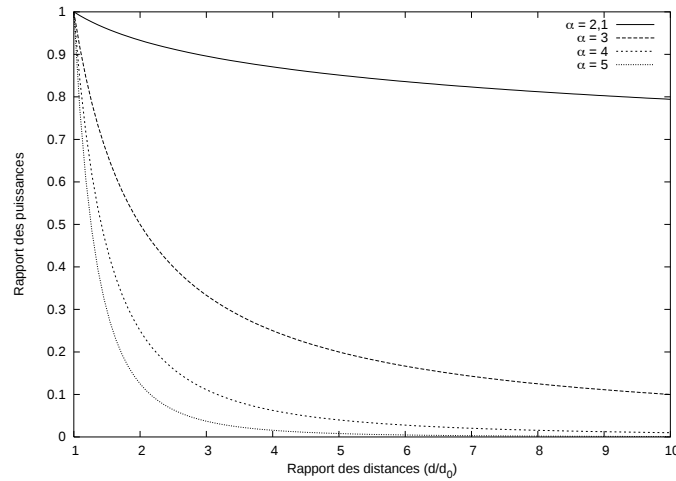


FIG. 2.11 – Influence comparée des brouilleurs situés à une distance supérieure à la portée d’émission pour différentes valeurs du paramètre d’affaiblissement

2.3.4 Dépassement de la capacité du médium

Il existe donc, dans un réseau dont la topologie n’est pas connue *a priori* et évolue au cours du temps, de nombreuses situations mettant en défaut le mécanisme de contrôle d’admission présenté. Un mécanisme apportant aux mobiles une connaissance parfaite de l’ensemble de l’état du réseau permettrait, bien sûr d’obtenir de meilleurs résultats. Toutefois, la mise en œuvre d’un tel mécanisme est extrêmement coûteuse puisqu’il requiert la transmission d’un volume d’information important tout en alliant fiabilité et réactivité. De plus, le travail de Georgiadis *et al.* montre que le problème considéré ici est NP-complet [GJM04, GJM03].

Des erreurs surviendront donc de temps à autre, malgré tous les efforts déployés. Un mauvais contrôle d’admission se traduira par l’acceptation de trafics qu’il n’est pas possible de router. Par ailleurs, la mobilité de certains nœuds modifiera l’état des ressources disponibles d’autres nœuds alors que ces derniers sont en train de retransmettre des flux acceptés au préalable. Chacune de ces deux situations se traduira par un dépassement local de la capacité du médium, provoquant des pertes de paquets, un accroissement des délais de transmission et un remplissage des files d’attente important.

En effectuant une surveillance de ces différents paramètres, les mobiles sont en mesure d’effectuer en permanence une surveillance des différents trafics privilégiés qu’ils retransmettent. Lorsque le débit de réémission de ces trafics est inférieur au débit reçu par le saut précédent, cela signifie que l’occupation du canal est trop importante. Quelle politique appliquer alors pour résoudre le problème ?

Toutes les applications n’ont pas les mêmes contraintes. Certaines applications, comme le transfert de données pourront être ralenties de façon quelconque sans que cela ait un impact sur le bon déroulement des opérations. Les applications présentant de fortes contraintes de délais, comme la voix, ne devraient pas subir de dégradation du fait de congestion. Peut-être vaut-il mieux les re-router dans certains cas et si ce processus présente un coût raisonnable. Enfin, certaines applications dites adaptatives, pourront modifier leur débit d’émission en fonction des conditions du réseau pour peu qu’elles soient averties des changements.

Un mode de routage par flux est un avantage pour le réseau lors de la résolution de collisions. En effet, s’il est possible de connaître pour chaque flux le comportement approprié à adopter face à une congestion, il est alors envisageable d’adapter les mesures destinées à résoudre le problème en accord avec les spécificités de chaque flux.

C’est pourquoi, dans l’architecture présentée, chaque requête de route avec qualité de service ne contiendra pas uniquement une demande de bande passante mais aussi un profil de dégradation. Ce profil sera constitué de deux valeurs : un incrément et un seuil. Lorsqu’un routeur constate une dégradation des

performances, il examine les différents profils de dégradation des flux et retranchera à la bande passante allouée à chaque flux une fois l'incrément de ce flux. Si l'application émettrice l'a spécifié dans la requête de route, elle sera avertie de ce changement, lui permettant d'adapter son débit d'émission aux nouvelles conditions du réseau. Lorsque décrémenter le débit d'un flux conduirait à dépasser le seuil prescrit par l'application émettrice, la route est considérée comme cassée et un message est envoyé à la source afin qu'elle effectue une nouvelle demande de route. Un message de libération des ressources est envoyé en parallèle à la destination.

Ce type de fonctionnement peut être rendu aussi réactif que nécessaire puisqu'il n'est basé que sur une estimation locale de l'état du canal. Il est par ailleurs souple puisqu'il permet aux applications de spécifier précisément le comportement à adopter face à un problème. Toutefois, ce mécanisme requiert des volumes de stockage et de traitement supplémentaires de la part des différents routeurs. D'autre part, ce mécanisme coopératif peut avoir pour effet d'accroître les inégalités dans l'accès au médium. Si l'on considère à nouveau le cas des trois paires, ce mécanisme aura tendance à réduire le débit de la paire centrale puisque les paires extérieures ne constatent que rarement une décroissance de leur débit. Ce mécanisme permettra cependant à la paire centrale d'être avertie de la situation.

2.4 Évaluation

Les performances du protocole BRuIT ont été évaluées au moyen du simulateur de réseaux NS-2 en version 2.26. Ce simulateur, pour des raisons de performance, effectue de nombreuses simplifications au niveau de la modélisation du canal et de la propagation radio. Les interférences ne sont pas additives, c'est-à-dire que pour déterminer le rapport signal sur bruit affectant une trame, les puissances des signaux interférents sont comparées séparément avec la puissance du signal utile. Les modèles de propagation radio sont simples. NS propose un modèle de propagation en espace libre et un modèle avec réflexion sur le sol présentant une atténuation en $1/d^2$ jusqu'à une certaine distance à partir de laquelle le signal décroît en $1/d^4$. NS propose par ailleurs un modèle de propagation probabiliste (*shadowing*) mais n'intégrant aucune notion de corrélation temporelle ou spatiale. La bonne réception d'une trame est uniquement dépendante de l'affaiblissement subi par le signal, la puissance de celui-ci étant comparée à un seuil permettant de déterminer si un signal peut être décodé ou non, et de la présence d'interférents. Aucune estimation du taux d'erreurs bit en fonction de la modulation utilisée n'est mise en œuvre. Enfin, les valeurs des paramètres utilisées par le simulateur correspondent à un débit binaire de 2 Mbit/s et non de 11 Mbit/s comme le permet la norme IEEE 802.11b. Les résultats présentés dans cette section sont donc à considérer avec précaution. S'ils permettent de valider l'aspect algorithmique de BRuIT et d'en déterminer quelques caractéristiques, seule une implantation réelle de ce protocole permettrait d'en déterminer les limites réelles.

2.4.1 Mécanisme de réservation

Afin de valider le mécanisme de réservation, plusieurs simulations simples ont été réalisées. Par exemple, considérons la topologie représentée en figure 2.12 comportant douze nœuds. Dans cette situation, le mobile *A* souhaite envoyer à destination de *B* un flux à un débit constant de 160 kbit/s, ce qui correspond par exemple à un trafic audio de bonne qualité. Vingt secondes après le début de la simulation, le nœud *C* cherche à envoyer un flux présentant les mêmes caractéristiques à destination du nœud *D*.



FIG. 2.12 – Un scénario simple

La figure 2.13 représente les débits obtenus avec le protocole de routage sans qualité de service AODV. Lorsque le second flux est présent sur le réseau, les débits deviennent irréguliers, même si en moyenne

sur toute la durée de la simulation, l'équité entre les flux semble respectée. La route doit par ailleurs régulièrement être reconstruite, ce qui se traduit par les périodes durant lesquelles un flux a un débit nul.

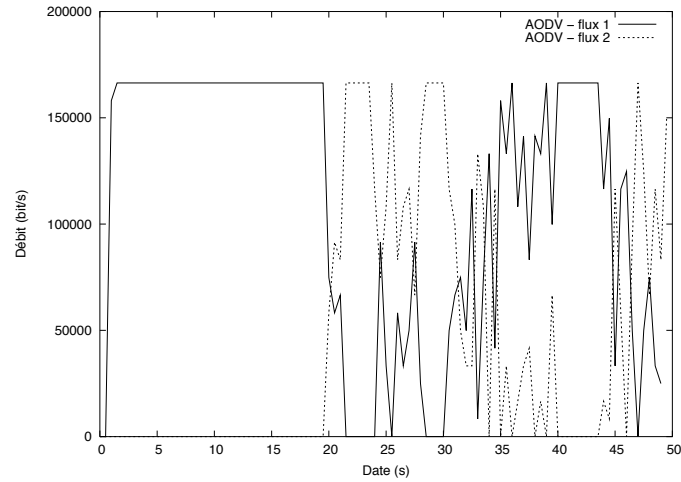


FIG. 2.13 – Débits des deux flux concurrents avec le protocole de routage AODV

En utilisant le protocole BRuIT, les débits obtenus par les deux flux sont représentés en figure 2.14. Le premier flux à entrer sur le réseau est autorisé à effectuer une réservation de bande passante correspondant au débit qu'il requiert. En revanche, lorsque le second flux émet une requête de route similaire, aucune réponse ne lui parvient. Il n'effectue donc pas de réservation et émet le flux en mode au mieux. Le débit du trafic au mieux étant limité par un filtre, afin de garantir que les réservations ne seront pas perturbées par les trafics non privilégiés, il n'utilise que la part disponible des ressources, c'est-à-dire environ 40 kbit/s.

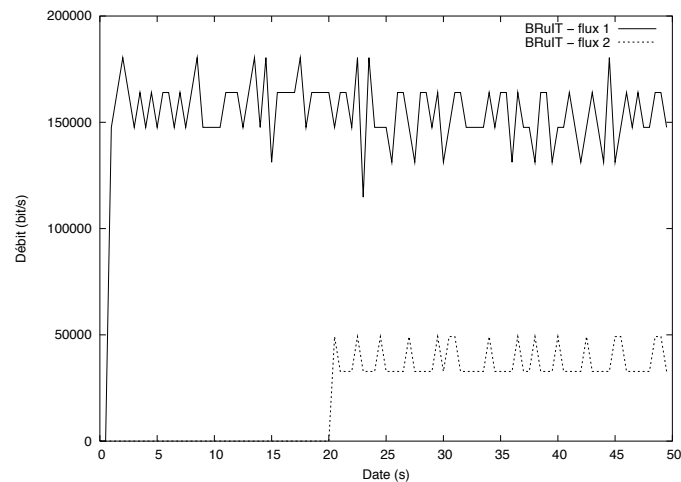


FIG. 2.14 – Débits des deux flux concurrents avec le protocole BRuIT

BRuIT refusera donc les réservations qu'il considère comme pouvant surcharger le réseau. La figure 2.15(a) représente, pour des graphes géométriques aléatoires dont le rayon d'émission est fixé à

0,16⁴, le taux d'acceptation des requêtes émises. Cette figure représente l'évolution de ce taux d'acceptation au fur et à mesure que le nombre de nœuds croît et pour une charge variant de 5 à 30 flux de 80 kb/s. Comme l'on pouvait s'y attendre, le taux d'acceptation décroît alors que le nombre de flux augmente, la capacité du réseau demeurant constante. En revanche, le taux d'acceptation croît avec la densité du réseau. En effet, plus un réseau est dense, plus la probabilité qu'il existe une route d'une source à une destination donnée est importante. En conséquence, la probabilité qu'une route admissible existe augmente aussi avec le nombre de nœuds en présence. Pour 5 flux, ce taux atteint 61,2 % et pour 30 flux, il atteint 20,4 %. Le rapport entre le nombre de flux pour lesquels une route sans contrainte de qualité de service peut être trouvée par le protocole AODV et le nombre total de flux est représenté en figure 2.15(b). Ce taux est plus élevé puisqu'il atteint 95,2 % pour 5 flux et 42,5 % pour 30 flux. La différence entre les résultats obtenus par BRuIT et les résultats obtenus par AODV permet d'évaluer l'impact de l'ajout de contraintes de qualité de service sur le routage. Le contrôle d'admission de BRuIT accepte entre une et trois fois moins de flux qu'AODV. Il faut cependant noter que la figure 2.15(a) ne représente que les routes *avec qualité de service* découvertes par BRuIT. Lorsque aucune route ne répond aux critères de qualité de service spécifiés, le protocole recherche une route au mieux et le nombre de routes découvertes est au total similaire à celui du protocole AODV.

2.4.2 Routage

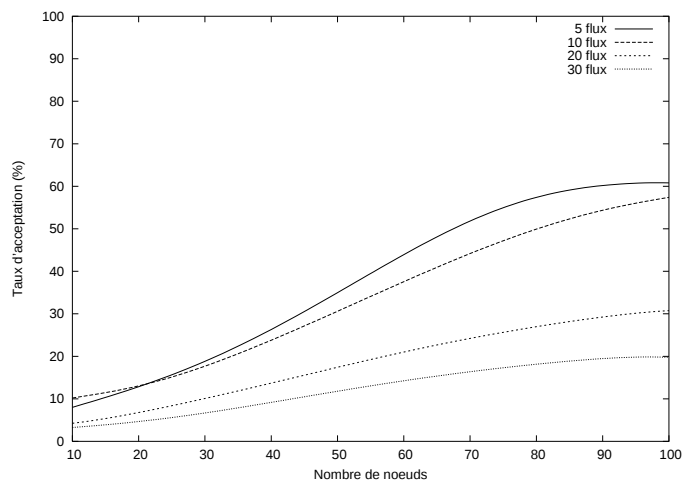
Le routage de BRuIT est effectué de manière réactive. Les routes sont recherchées à la demande et seules les routes admissibles en terme de bande passante disponible sont considérées. La qualité des routes déterminées par cette approche est fortement dépendante des conditions du réseau. Un routage réactif ne produira pas systématiquement des routes optimales en terme de nombre de sauts.

La figure 2.16(b) représente la moyenne réalisée sur 100 simulations de la longueur des routes générées par l'implantation du protocole de routage réactif AODV sous NS-2 sur des graphes géométriques aléatoires en fonction du nombre de nœuds et du nombre de flux en présence sur le réseau. Cette figure permet de déterminer la charge du réseau résultant de la présence des flux. La longueur des routes, croissant de 1 saut à 7 sauts, est peu dépendante du nombre de flux en présence, mais dépend essentiellement du nombre de nœuds en présence. Ceci peut s'expliquer par le fait que plus la densité du réseau augmente, plus la probabilité de collision entre deux messages de recherche de route sera élevée. Ces messages étant transmis en diffusion et non acquittés, leur perte n'est pas détectée et ils ne sont pas retransmis. En conséquence, de nombreux chemins ne seront pas explorés. La longueur moyenne des routes générées par BRuIT, représentée en figure 2.16(a) présente les mêmes caractéristiques, ce qui est peu surprenant compte tenu de la similitude des deux modes de recherche de route. La longueur de ces routes varie cette fois de 2 sauts à 7 sauts. On constate que BRuIT génère des routes plus longues qu'AODV sur des graphes peu denses. Ceci est encore une fois dû au fait que BRuIT ne recherche pas simplement une route mais une route répondant à un critère de qualité de service. En conséquence, une partie des chemins est ignorée lors de la recherche de route. Les performances des deux protocoles se rejoignent lorsque la densité du réseau augmente du fait de l'accroissement du nombre de routes existant entre deux mobiles avec le nombre de nœuds en présence.

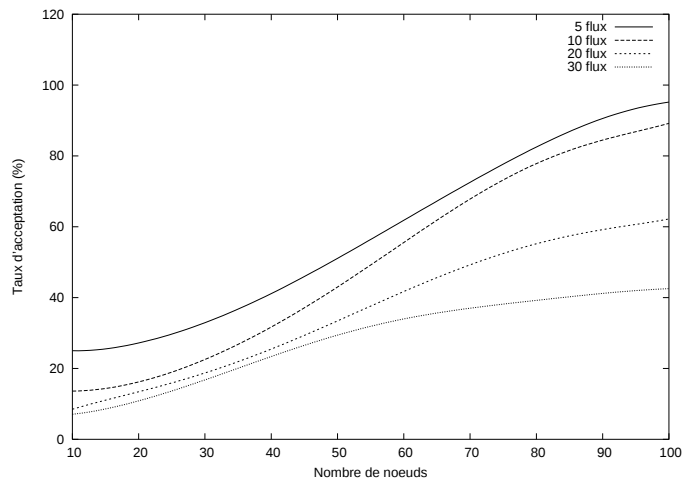
La figure 2.17(b) compare la longueur des routes obtenues par AODV à la longueur du plus court chemin entre les nœuds source et destination correspondants. La figure 2.17(a) représente la même métrique pour BRuIT. On peut remarquer qu'AODV génère des routes très proches des plus courts chemins. En revanche les routes découvertes par BRuIT sont jusqu'à deux fois plus longues en moyenne que les plus courts chemins. On peut remarquer que la longueur des routes découvertes par AODV croît avec la densité du réseau alors que dans le cas de BRuIT, le phénomène inverse se produit. BRuIT profite de l'accroissement de la densité du réseau pour explorer un plus grand nombre de chemin. Les routes déterminées par BRuIT se rapprochent le plus des plus courts chemins pour une densité moyenne (environ 60 à 70 nœuds). La longueur des routes déterminées par AODV peut être considérée comme une borne inférieure de la longueur des routes que BRuIT pourra découvrir puisque BRuIT ajoute une restriction supplémentaire.

Les figures 2.18(b) et 2.18(a) représentent le temps moyen d'établissement des routes pour AODV et pour BRuIT respectivement. Ce temps est mesuré comme la différence entre l'instant où l'initiateur d'une communication reçoit la confirmation de route et l'instant où il avait émis la requête correspondante. Si

⁴Sous NS-2, la portée radio est de 161,2 m et les simulations ont été réalisées dans un carré de 1000 m × 1000 m

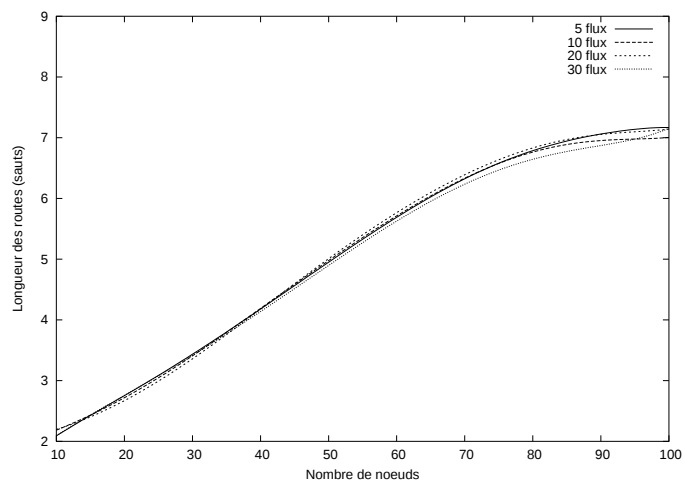


(a) BRuIT

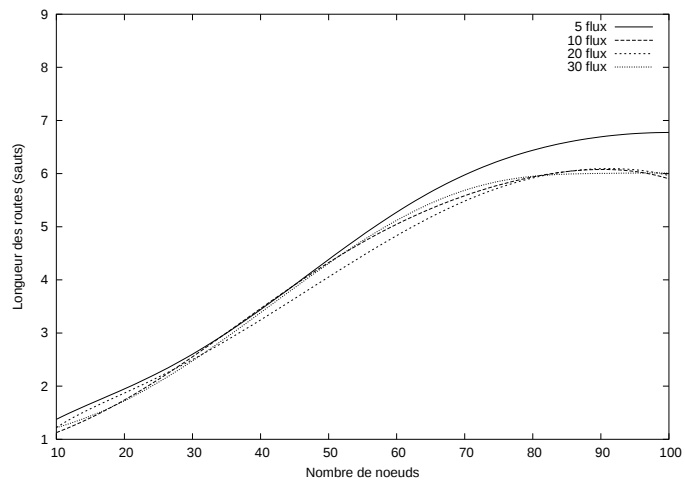


(b) AODV

FIG. 2.15 – Taux d'acceptation moyen des flux sur des graphes géométriques aléatoires (portée = 0,16) ; moyenne sur 100 simulations)

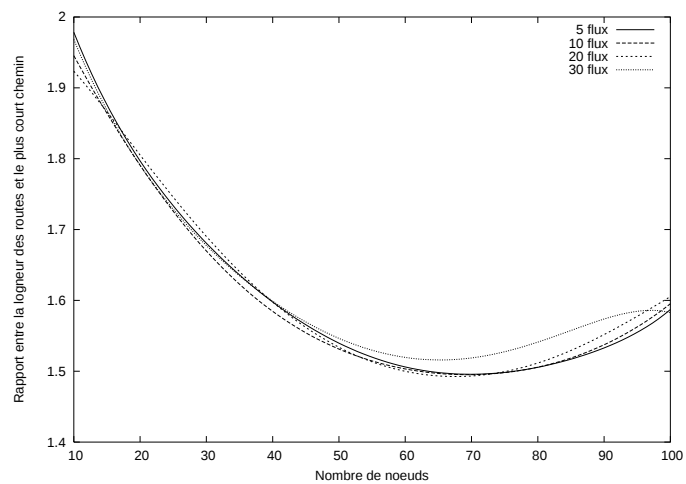


(a) BRuIT

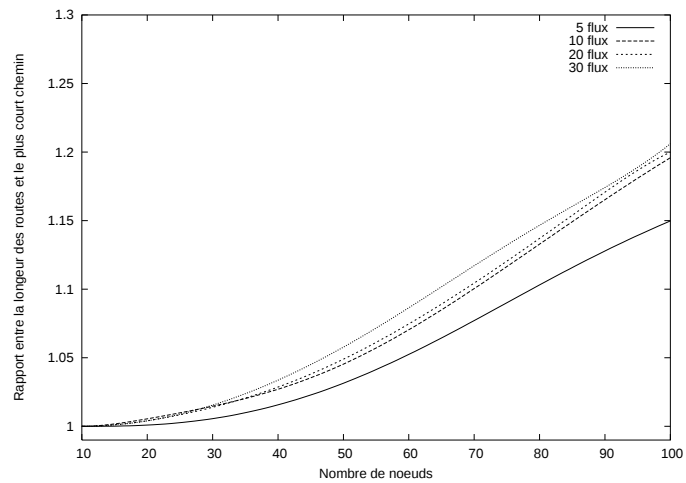


(b) AODV

FIG. 2.16 – Longueur moyenne des routes sur des graphes géométriques aléatoires (portée = 0,16) ; moyenne sur 100 simulations)



(a) BRuIT



(b) AODV

FIG. 2.17 – Rapport entre la longueur des routes déterminées et les plus courts chemins correspondants sur des graphes géométriques aléatoires (portée = 0,16) ; moyenne sur 100 simulations)

ces résultats sont similaires lorsque le réseau est peu dense, AODV prend l'avantage lorsque le réseau devient dense.

2.4.3 Évaluation du coût de ce mécanisme

Transmettre les informations nécessaires au contrôle d'admission engendre un surcoût en trafic de contrôle, les paquets *Hello* apportant aux mobiles la connaissance de leur voisinage étendu sont émis régulièrement. Nous avons choisi une fréquence d'émission de deux paquets par seconde et par nœud. Le protocole de routage proactif OLSR, se basant aussi sur une transmission régulière de paquets *Hello* utilise une fréquence d'émission d'un paquet par seconde. Toutefois, lorsqu'il s'agit de fournir des garanties de qualité de service, il est nécessaire de réagir à la fois à la mobilité des nœuds et à la dynamique des flux. La figure 2.19 compare les volumes de paquets de contrôle émis par l'ensemble des nœuds du réseau dans le scénario représenté en figure 2.12 en utilisant le protocole AODV et le protocole BRuIT. Si le trafic généré par BRuIT est permanent, son volume est régulier alors que le protocole AODV génère peu de paquets lorsque le réseau n'est pas surchargé mais le volume de trafic engendré par la disparition des routes due à un trafic élevé est nettement plus important. Ce scénario plaide encore une fois en faveur d'une limitation des débits des différents nœuds afin d'éviter autant que possible l'apparition de congestions.

Les figures 2.20(b) et 2.20(a) représentent le volume de signalisation moyen généré par l'ensemble des nœuds dans le cas de graphes géométriques aléatoires respectivement pour AODV et pour BRuIT. Le volume de trafic de contrôle croît dans les deux cas avec le nombre de nœuds et, dans une moindre mesure avec le nombre de flux en présence. L'ordre de grandeur reste le même et l'on peut remarquer que, si BRuIT génère de façon permanente un certain volume de trafic de contrôle alors qu'AODV n'émet des requêtes que lorsqu'une route est requise, la stratégie de BRuIT se révèle payante lorsque la charge du réseau devient importante. Dans ce cas, en effet, AODV souffrira de nombreuses cassures de routes dues à la présence de congestion dans le réseau et générera en conséquence un volume important de signalisation.

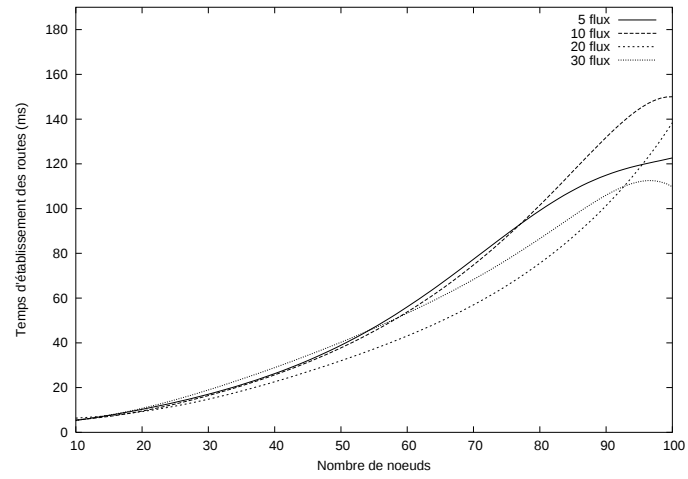
2.4.4 Dégradation des flux

Lorsque des dépassements de capacité surviennent, dus à une mauvaise identification des brouilleurs potentiels ou à la mobilité des nœuds, les réservations acceptées par les routeurs ne pourront être honorées. Ces mêmes routeurs constateront alors un accroissement du taux de pertes de paquets des flux privilégiés et prendront alors la décision de dégrader les flux en accord avec les profils de dégradation spécifiés par leurs émetteurs. Afin d'illustrer ce phénomène, examinons le scénario représenté en figure 2.21. Dans ce scénario, deux ensembles de trois nœuds, initialement indépendants, mettent en place des réservations pour deux flux. Le premier flux, de *A* à *B* a un débit de 600 kbit/s et le second, de *C* à *D*, un débit de 400 kbit/s. Les deux sous réseaux sont mobiles et se rapprochent jusqu'à arriver à distance de détection de porteuse après 50 s de simulation.

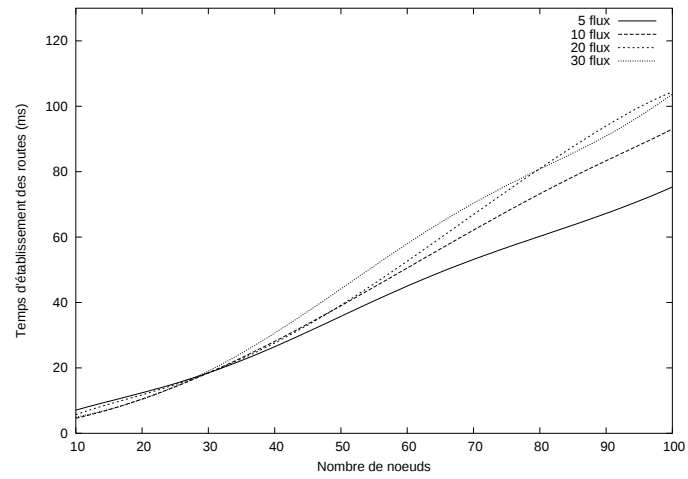
La figure 2.22 représente les débits atteints par les deux flux avec le protocole de routage AODV. Lorsque les différents émetteurs ont à partager la ressource radio, les deux flux sont pénalisés et le partage du médium est équitable et irrégulier. La figure 2.23 représente les débits atteints par les mêmes flux en utilisant le protocole BRuIT. Lorsque le débit des deux flux commence à chuter, le mécanisme de dégradation de BRuIT limite les débits de ces deux flux en accord avec leur profil initial. Les rapports entre les deux flux sont conservés et les débits sont stables même s'ils sont réduits.

2.5 Conclusion

Le protocole BRuIT présenté dans ce chapitre constitue une première solution au problème de la réservation de bande passante dans les réseaux *ad hoc*. Ce protocole adopte un fonctionnement de type réactif tout en tirant parti d'informations collectées de façon proactive. La transmission régulière de paquets *Hello* par chacun des nœuds du réseau permet d'apporter à ces mêmes nœuds une connaissance leur permettant d'affiner le contrôle d'admission. Toutefois, il est difficile de ne considérer ni trop ni pas assez d'informations car l'identification des interactions pouvant exister entre les différents nœuds est extrêmement dépendante de l'environnement et de la topologie précise du réseau, paramètres qui *a priori* peuvent évoluer rapidement. En conséquence, BRuIT intègre un mécanisme de dégradation et



(a) BRuIT



(b) AODV

FIG. 2.18 – Temps moyen d'établissement des routes sur des graphes géométriques aléatoires (portée = 0,16) ; moyenne sur 100 simulations)

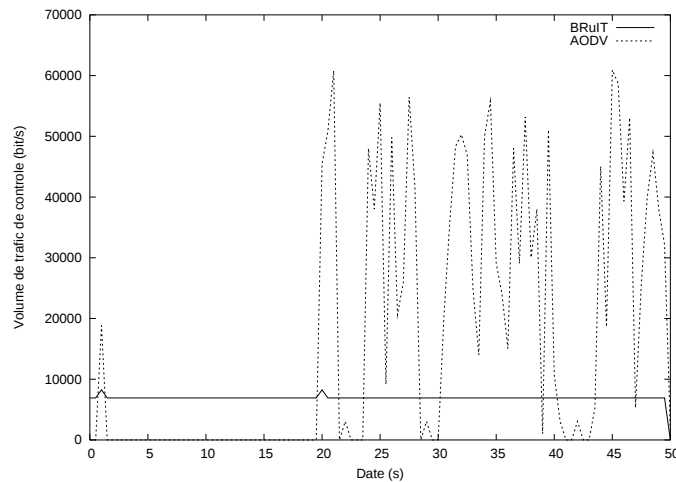


FIG. 2.19 – Volume de trafic de contrôle généré par BRuIT et par AODV

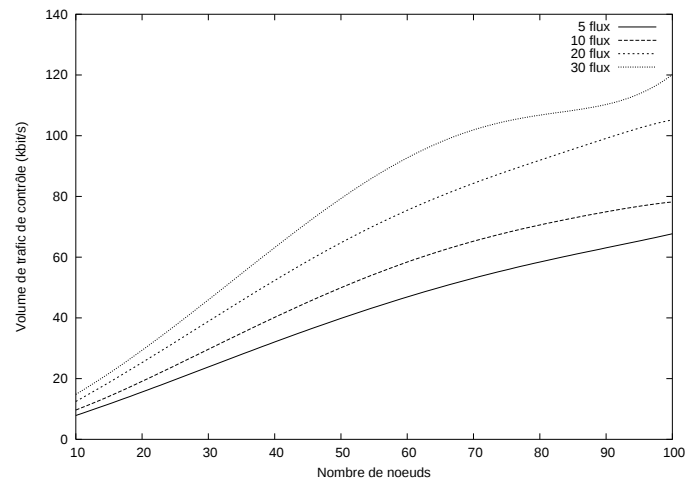
restauration des garanties permettant à chaque application de communiquer aux routeurs un profil de qualité de service lui correspondant.

Les simulations présentées ici ont permis de valider le fonctionnement du protocole BRuIT. Elles indiquent en particulier que s'il est possible d'améliorer la connaissance qu'ont les applications de l'état du réseau, ce processus a un coût en terme de performances. La sélectivité dans le choix des routes conduit à la génération de routes plus longues, mais tend à équilibrer la charge du réseau puisque le contrôle d'admission échouera dans les zones proches de la congestion. Le temps de mise en place des routes est accru, mais la prévention des congestions aura pour effet de limiter les procédures de reconstruction de routes.

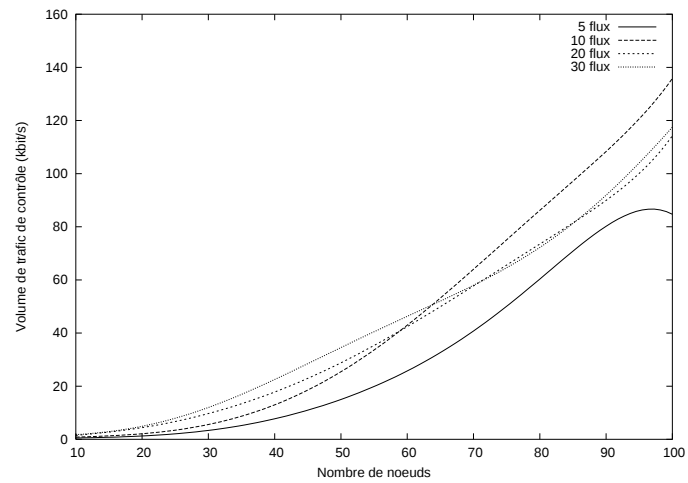
Toutefois, les performances réelles de ce protocole ne peuvent être déduites de simples simulations. Il est en effet difficile, en simulation, de concevoir les scénarios pertinents conduisant à l'évaluation du comportement du protocole. D'autre part, BRuIT n'a pas été conçu dans l'unique but de rajouter quelques lignes de code au simulateur NS. C'est pour répondre à ce besoin d'expérimentation que Sébastien Hinderer, dans [Hin02], a défini une architecture d'implantation sous le système Linux. L'implantation reste encore à effectuer, mais elle permettra de valider ou de modérer les résultats obtenus par simulation par des expérimentations sur des réseaux *ad hoc* réels tels que ceux étudiés par le projet WAND (*Wireless Ad hoc Networks for Dublin*) de l'équipe systèmes distribués de Trinity College, Dublin⁵.

Un problème n'a cependant pas été abordé lors de ce chapitre. En effet, fournir des garanties est impossible si l'on ne maîtrise pas le débit du trafic au mieux utilisant les mêmes ressources. Il est possible d'affecter une partie fixe de la bande passante du canal à ces trafics au mieux mais compte tenu des interactions existant entre les différents émetteurs, cette valeur conduira souvent à une sous-utilisation ou à une sur-utilisation de la capacité du canal.

⁵<http://wand.dsg.cs.tcd.ie/>



(a) BRuIT



(b) AODV

FIG. 2.20 – Volume moyen de trafic de contrôle généré dans des graphes géométriques aléatoires (portée = 0,16) ; moyenne sur 100 simulations)

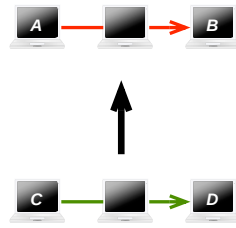


FIG. 2.21 – Deux sous-réseaux se rapprochent jusqu'à partager le médium

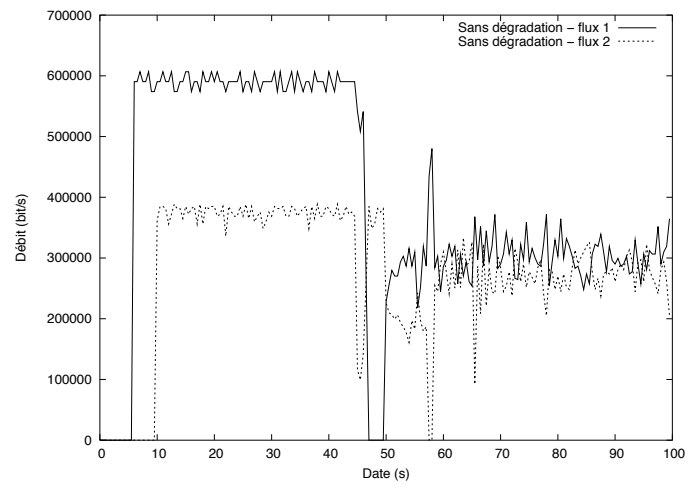


FIG. 2.22 – Diminution du débit des flux sans politique de gestion de dégradation

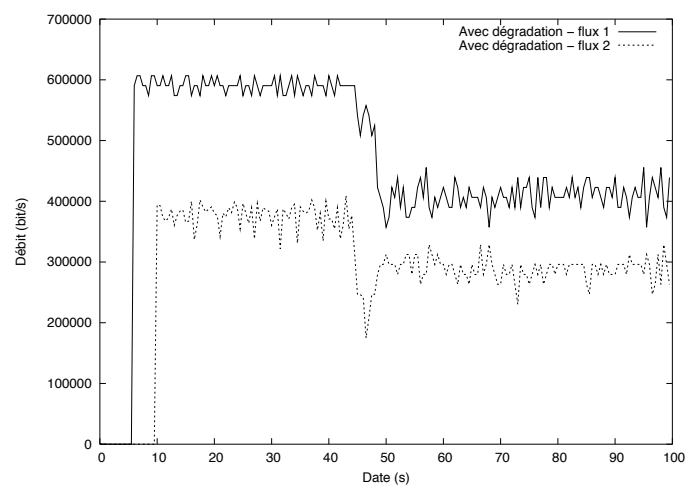


FIG. 2.23 – Dégradation du débit des flux avec BRuIT

Dans l'architecture proposée, les flux privilégiés doivent coexister avec les flux au mieux. En effet, s'il est possible de séparer ces différents trafics en utilisant deux canaux distincts, cette approche nécessite la mise en œuvre d'un mécanisme d'allocation de canaux capable de réagir lors de la fusion de deux réseaux distincts. En effet, si l'un des deux réseaux utilise pour les trafics au mieux le canal choisi par le second pour le routage des trafics privilégiés, les garanties deviendront rapidement caduques. Une autre solution pourrait résider dans la différenciation de services proposée par le protocole IEEE 802.11e.

Ce n'est toutefois pas l'approche que nous avons choisi d'étudier. Compte tenu du mode de fonctionnement de BRuIT, il semble possible d'utiliser les informations transmises régulièrement par les différents nœuds afin de réguler le débit du trafic au mieux. Tous les trafics coexisteront sur le même canal et le débit maximal du trafic au mieux sera limité afin de ne pas empiéter sur les trafics privilégiés. Cependant, comment limiter le trafic au mieux au plus juste, sans provoquer de dépassement de la capacité du médium radio et sans sur-limiter ce type de trafic ?

Dans cette section, nous présentons un algorithme itératif dont l'objectif est d'allouer à chaque mobile une quantité de bande passante qu'il pourra utiliser pour le routage des trafics au mieux. Cet algorithme réalise un partage de la bande passante disponible en tenant compte du mode de partage du médium actuellement en vigueur dans les réseaux *ad hoc*. La bande passante allouée doit être la plus importante possible sans pour autant provoquer de congestion. L'allocation réalisée se doit par ailleurs de maintenir un certain niveau d'équité entre les différents mobiles, et ce en réduisant les écarts entre les bandes passantes allouées.

Il n'existe, à notre connaissance, aucun travail similaire adapté aux réseaux *ad hoc*. En effet, les algorithmes d'allocation de bande passante fonctionnent généralement de manière totalement centralisée, faisant appel à des techniques d'optimisation ou ne garantissent qu'une équité en terme de nombre de trames transmises, tels qu'un processus de tourniquet distribué.

Dans [HB01], par exemple, Huang et Bensaou proposent un algorithme centralisé et un algorithme distribué permettant d'allouer à différents flux une proportion de la capacité du médium radio en fonction des interactions existant entre ces différents flux. La version centralisée de cet algorithme se base sur une décomposition en cliques du graphe de contention entre les trafics dans le réseau afin de déterminer les proportions de la capacité de chacune de ces cliques que peut utiliser un flux. La version distribuée se base sur une vision locale à chaque nœud de ce graphe de contention afin d'aboutir au même type d'allocation. Les auteurs montrent que cette allocation respecte les propriétés d'une allocation équitable en terme d'équité *max-min*. Cependant, cet algorithme opère sur une vision du réseau figée dans le sens où l'apparition ou la disparition d'un flux dans le réseau nécessite une nouvelle exécution de l'algorithme. Par ailleurs, ce type de stratégie est peu adaptée au problème de l'allocation équitable de bande passante afin de limiter les débits des trafics au mieux puisque les caractéristiques de ces trafics ne sont pas connues.

3.1 Description formelle du problème

L'objectif que nous recherchons est une allocation de la bande passante non utilisée par les trafics privilégiés entre les différents mobiles d'un réseau *ad hoc*. Soit $G = (V, E, b)$ le graphe pondéré représentant le réseau considéré. Un sommet dans ce graphe représente un mobile et un lien entre deux sommets signifie que les deux mobiles correspondants se partagent la bande passante du canal radio. En d'autres

termes, un lien existe entre deux mobiles s'ils ne peuvent émettre un paquet simultanément sur le canal. Dans la suite de cette section, nous utiliserons les notations suivantes :

- $G = (V, E, b)$ est le graphe non orienté représentant le réseau *ad hoc* considéré. V est l'ensemble des sommets de ce graphe, E l'ensemble de ses arêtes et b une fonction associant à chaque sommet $v \in V$ un poids $b(v) \in \mathbb{R}$ représentant la capacité du canal radio disponible pour le mobile correspondant. Il s'agit du débit maximal pouvant être émis dans un voisinage de v incluant v sans provoquer un dépassement de la capacité du médium ;
- Pour un sommet $v \in V$, on notera $N(v) = \{u \in V, (u, v) \in E\}$ son voisinage ouvert et $N[v] = N(v) \cup \{v\}$ son voisinage fermé.

Une allocation de bande passante $(x(v))_{v \in V}$ associe à chaque mobile une valeur représentant le débit maximal auquel il est autorisé à émettre. Une telle allocation sera qualifiée d'admissible dès lors qu'en aucun point du réseau, la capacité du médium n'est excédée. Le respect de cette contrainte est une condition *sine qua non* à laquelle doit se conformer toute allocation. Celle-ci peut se traduire formellement par l'ensemble d'inéquations suivant :

$$\forall v \in V, \sum_{u \in N[v]} x(u) \leq b(v). \quad (3.1)$$

Ces contraintes définissent un ensemble d'allocations admissibles formant un polyèdre convexe ou un ensemble polyédrique convexe non borné $\mathbb{R}^{|V|}$. Notre objectif est de déterminer la meilleure allocation possible satisfaisant ces contraintes. La meilleure allocation pourrait signifier l'allocation maximisant l'utilisation totale de la bande passante. Il pourrait aussi s'agir de l'allocation la plus équitable possible minimisant les écarts entre les bandes passantes allouées. Considérons par exemple le réseau représenté par le graphe de la figure 3.1. Dans ce réseau, tous les mobiles disposent de la même capacité disponible initiale et nous supposons que le partage de la bande passante ne s'effectue qu'entre voisins directs.

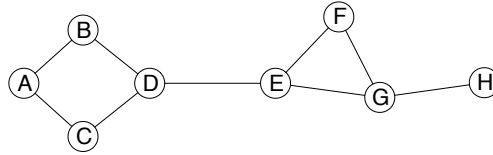


FIG. 3.1 – Exemple de graphe associé à un réseau *ad hoc*

Dans cette configuration, l'allocation la plus équitable associerait la même quantité de bande passante à tous les nœuds du réseau, c'est-à-dire 25 % de la capacité résiduelle. Allouer une proportion supérieure de bande passante à chaque nœud du réseau provoquerait un dépassement de capacité au niveau de tous les sommets de degré 3. Cependant, une telle allocation sous-utilise le réseau puisqu'il est possible d'allouer à A une plus grande passante sans provoquer de dépassement de capacité. En effet, tous ses voisins peuvent supporter une charge supplémentaire.

A contrario, une allocation optimale en terme d'utilisation du réseau consisterait à allouer $1/3$ de la capacité du médium à chaque mobile sauf au mobile E auquel serait affectée une bande passante nulle. Cette allocation, malgré son déséquilibre manifeste, est la solution la plus équitable conduisant à une allocation totale de $7/3$ fois la capacité du médium, dans le sens où s'agit de la solution minimisant la variance entre les débits alloués aux différents nœuds.

Le problème de la maximisation de l'utilisation globale du réseau tout en respectant l'ensemble de contraintes 3.1 s'apparente à un problème de *fractional packing*. Ce type de problème peut être résolu par des algorithmes de programmation linéaire et un certain nombre d'algorithmes d'approximation sont présentés dans [PST95]. Ces algorithmes sont séquentiels et difficiles à adapter dans un contexte distribué. Dans [BBR97], un algorithme distribué de programmation linéaire permet d'obtenir une approximation en $1 + \epsilon$ en un nombre poly-logarithmique d'échanges de messages. Cependant, ce type de solution n'intègre aucune notion d'équité et peut, comme dans l'exemple précédent, déconnecter le réseau en allouant une bande passante nulle à certains mobiles.

La solution que nous recherchons devra représenter un compromis entre utilisation du réseau et équité de l'allocation. Par ailleurs, afin d'être peu coûteuse dans un contexte totalement distribué, elle devra être basée uniquement sur des informations locales.

3.1.1 Une suite d'allocations admissibles

Afin d'introduire la procédure d'allocation que nous avons choisi d'étudier, considérons un nœud particulier du réseau. Ce nœud peut disposer, pour un coût relativement faible en terme d'occupation du médium, d'informations sur tous ces voisins immédiats dans le graphe G . Afin à la fois de s'assurer du respect des contraintes, d'essayer d'utiliser au mieux la capacité du médium et d'assurer un comportement le plus équitable possible, ce mobile ne pourra s'octroyer plus que le minimum des bandes passantes libres dans son voisinage fermé divisé par le degré maximum parmi ses voisins et lui-même augmenté d'une unité. Dans la suite de cette section, nous utiliserons les notations suivantes *en sus* de celles qui ont été introduites au paragraphe précédent :

- $d(v)$ représentera le degré du sommet $v \in V$ dans le graphe G ;
- $\Delta_1(v) = \max_{u \in N[v]} d(u)$ est le degré maximum dans le voisinage fermé du sommet v ;
- $b_1(v) = \min_{u \in N[v]} b(u)$ est la bande passante (*i.e.* le poids) minimal dans le voisinage fermé du sommet v ;
- $x(v)$ représente la bande passante allouée au mobile correspondant à $v \in V$.

Lemme 3.1.1. *Autoriser chaque mobile à utiliser une bande passante $x(v) = \frac{b_1(v)}{\Delta_1(v)+1}$ conduit à une allocation respectant les contraintes 3.1.*

Preuve Par définition de Δ_1 et b_1 , $u \in N[v]$ implique que $d(v) \leq \Delta_1(u)$ et que $b(v) \geq b_1(u)$. En conséquence,

$$\begin{aligned} \sum_{u \in N[v]} x(u) &\leq \sum_{u \in N[v]} \frac{b_1(u)}{\Delta_1(u)+1} \leq \sum_{u \in N[v]} \frac{b(v)}{d(v)+1} \leq \frac{d(v)+1}{d(v)+1} b(v) \\ &\leq b(v). \end{aligned}$$

■

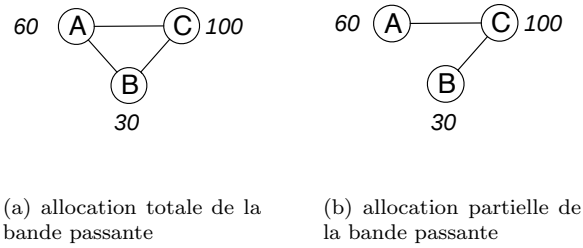


FIG. 3.2 – Deux scénarios similaires conduisant à des allocations de qualités sensiblement différentes

Considérons par exemple la situation représentée par la figure 3.2(a). Sur cette figure représentant un réseau comportant trois nœuds, les bandes passantes disponibles pour chaque nœud sont indiquées à coté de chacun des terminaux. L'allocation la plus équitable respectant les contraintes permet à chaque mobile d'utiliser 10 unités de bande passante puisque B ne dispose que de 30 unités de bande passante. Il n'y a, dans ce cas, pas d'allocation plus efficace présentant un écart aussi faible entre les différentes bandes passantes allouées.

Dans cette situation, cette allocation conduit à une utilisation totale de la bande passante disponible. En effet, aucun mobile ne peut émettre à un débit supérieur sans provoquer un dépassement de capacité au niveau du nœud B . Toutefois, dans la plupart des cas, cette allocation étant précautionneuse, elle ne conduit pas à une utilisation totale de la capacité du médium. Par exemple, dans la topologie représentée par la figure 3.2(b), les mobiles B et C obtiendront chacun 10 unités de bande passante et le mobile A 20 unités, compte tenu de leurs voisinages respectifs. Toutefois, à l'issue de cette affectation, au moins 10 unités de bande passante sont encore inutilisées dans le voisinage de chacun des nœuds. Cette topologie pourtant ne diffère de la précédente que d'un unique lien.

Ce type de situation pourrait bien sûr être résolu en apportant aux différents nœuds du réseau une meilleure connaissance de leur voisinage. Il est cependant souhaitable que l'algorithme d'allocation

soit le plus distribué possible et ne se base donc que sur des informations locales. Une autre approche permettant de résoudre ce problème de sous-utilisation consiste à itérer le processus précédent. Une allocation résultera en un nouvel ensemble de bandes passantes disponibles sur lequel il sera possible d'appliquer à nouveau le processus tout en respectant toujours les contraintes 3.1.

L'allocation proposée pourra donc s'exprimer pour chaque mobile $v \in V$ par une suite $(x^{(i)}(v))_{i \in \mathbb{N}}$. Parallèlement, nous définirons et utiliserons la suite $(e^{(i)}(v))_{i \in \mathbb{N}}$ dont les valeurs seront définies par les bandes passantes résiduelles successives au niveau du nœud v . Les deux suites sont définies comme suit :

$$\begin{cases} x^{(0)}(v) = \frac{b_1(v)}{\Delta_1(v) + 1}; \\ e^{(0)}(v) = b(v) - \sum_{u \in N[v]} x^{(0)}(u). \end{cases} \quad (3.2)$$

$$\forall i \in \mathbb{N}, \begin{cases} x^{(i+1)}(v) = x^{(i)}(v) + \frac{1}{\Delta_1(v) + 1} \cdot \min_{u \in N[v]} e^{(i)}(u); \\ e^{(i+1)}(v) = b(v) - \sum_{u \in N[v]} x^{(i+1)}(u). \end{cases}$$

Le calcul des termes de cette suite d'allocations n'est basé que sur des informations recueillies dans un voisinage à une distance 1 dans le graphe. Si le graphe considéré représente les exclusions mutuelles entre les différents nœuds, les principes de cet algorithme pourront aisément être couplés au protocole BRuIT. Étudions maintenant en détail ses différentes propriétés.

3.1.2 Propriétés de la suite

Lemme 3.1.2. *Tous les termes de cette suite d'allocations respectent les contraintes (3.1) :*

$$\forall v \in V, \forall i \in \mathbb{N}, \sum_{u \in N[v]} x^{(i)}(u) \leq b(v).$$

Preuve Par définition, $\forall v \in V, \forall i \in \mathbb{N}$,

$$\forall u \in N[v], \min_{w \in N[u]} e^{(i)}(w) \leq e^{(i)}(v).$$

$$\Rightarrow \sum_{u \in N[v]} \left(\min_{w \in N[u]} e^{(i)}(w) \right) \leq |N[v]| \times e^{(i)}(v).$$

Comme $N[v]$ contient au moins v , $|N[v]| > 0$, on peut alors écrire :

$$\frac{\sum_{u \in N[v]} \left(\min_{w \in N[u]} e^{(i)}(w) \right)}{|N[v]|} \leq e^{(i)}(v).$$

Comme $\forall u \in V, \forall v \in N[u], 1 + \Delta_1(u) \geq 1 + d(v)$ et, par définition, $1 + d(v) = |N[v]|$, alors

$$\sum_{u \in N[v]} \min_{w \in N[u]} e^{(i)}(w) \times \frac{1}{\Delta_1(u) + 1} \leq e^{(i)}(v).$$

En conséquence,

$$\begin{aligned} \sum_{u \in N[v]} (x^{(i+1)}(u) - x^{(i)}(u)) &\leq b(v) - \sum_{u \in N[v]} x^{(i)}(u). \\ \Rightarrow \sum_{u \in N[v]} x^{(i+1)}(u) &\leq b(v). \end{aligned}$$

On a donc $\forall v \in V, \forall i \in \mathbb{N}^*, \sum_{u \in N[v]} x^{(i)}(u) \leq b(v)$.

Grâce au lemme 3.1.1, on sait que

$$\forall v \in V, \sum_{u \in N[v]} x^{(0)}(u) \leq b(v).$$

Donc,

$$\forall v \in V, \forall i \in \mathbb{N}, \sum_{u \in N[v]} x^{(i)}(u) \leq b(v).$$

■

Tous les termes de la suite sont donc admissibles au regard des contraintes. Il s'agit maintenant de déterminer la qualité de la solution. Tout d'abord, au regard du lemme 3.1.2 et de la définition de $x^{(0)}(v)$, nous pouvons aisément affirmer que, si aucune bande passante disponible n'est nulle lors du calcul du terme initial de la suite, tout nœud obtient et conserve une valeur d'allocation non nulle. Formellement :

Lemme 3.1.3. $\forall v \in V, b(v) > 0 \Rightarrow \forall i \in \mathbb{N}, x^{(i)}(v) > 0$.

Quel que soit le nœud considéré, cette suite est donc positive, et pour peu qu'aucune bande passante initiale ne soit nulle, strictement positive. Elle respecte les contraintes et détermine donc une allocation admissible. Chaque suite $(x^{(i)}(v))_{i \in \mathbb{N}}$ est donc bornée.

Lemme 3.1.4. La suite $(x^{(i)}(v))_{i \in \mathbb{N}}$ est convergente.

Preuve Considérant la définition de $e^{(i)}(v)$ et le lemme 3.1.2, on peut écrire que $\forall v \in V, \forall i \in \mathbb{N}, e^{(i)}(v) \geq 0$. En conséquence, la suite $(x^{(i)}(v))_{i \in \mathbb{N}}$ est monotone croissante. Étant majorée par $b(v)$, la suite $(x^{(i)}(v))_{i \in \mathbb{N}}$ est convergente. ■

Ce lemme montre donc qu'il existe une limite. Toutefois, il semble difficile de caractériser cette limite pour un graphe quelconque. Il est aussi difficile de donner une indication de la vitesse de convergence de cette suite. En effet, si l'on considère un graphe complet dans lequel tous les poids des nœuds sont identiques, le premier terme de la suite sera sa limite. En revanche, si l'on considère la topologie simple représentée par la figure 3.3, à l'issue de la première étape, le rapport entre les différentes bandes passantes est conservé. Cette situation va alors perdurer et la limite de la suite ne sera pas atteinte en un nombre fini d'étapes. Si le rapport des bandes passantes initiales avait été $b(A) = b(C) = 2 \cdot b(B)$, la limite de la suite aurait été atteinte dès son premier terme.

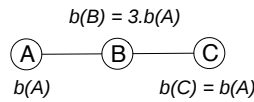


FIG. 3.3 – Une chaîne de trois nœuds

S'il est impossible de borner la vitesse de convergence, il est cependant possible de donner quelques indications sur les conditions permettant une convergence.

Lemme 3.1.5. Un nœud v ne pourra disposer d'une bande passante résiduelle nulle, à l'étape i que s'il présente à la fois la bande passante minimale et le degré maximal parmi son voisinage fermé à l'étape $i - 1$. Sa bande passante résiduelle sera effectivement nulle s'il possède la bande passante minimale et le degré maximal parmi son voisinage fermé à deux sauts.

$$\forall i \in \mathbb{N}^*; \forall v \in V, e^{(i)}(v) = 0 \Rightarrow d(v) = \max_{u \in N[v]} d(u) \text{ et } e^{(i-1)}(v) = \min_{u \in N[v]} e^{(i-1)}(u).$$

$$\forall i \in \mathbb{N}^*; \forall v \in V, d(v) = \max_{u \in N_2[v]} d(u) \text{ et } e^{(i-1)}(v) = \min_{u \in N_2[v]} e^{(i-1)}(u) \Rightarrow e^{(i)}(v) = 0.$$

Preuve Supposons que le nœud $v \in V$ a une bande passante résiduelle non nulle à l'étape $i - 1$ et n'a pas le degré maximum parmi ses voisins. Alors,

$$e^{(i-1)}(v) > 0 \text{ et } \exists z \in N(v) \text{ tel que } d(z) = \Delta_1(v) > d(v).$$

Par définition :

$$\begin{aligned} e^{(i)}(v) &= e^{(i-1)}(v) - \sum_{u \in N[v]} \frac{\min_{w \in N[u]} e^{(i-1)}(w)}{\Delta_1(u) + 1} \\ &= e^{(i-1)}(v) - \left(\sum_{u \in N[v] \setminus \{z\}} \frac{\min_{w \in N[u]} e^{(i-1)}(w)}{\Delta_1(u) + 1} \right) - \frac{\min_{w \in N[z]} e^{(i-1)}(w)}{\Delta_1(z) + 1}. \end{aligned}$$

Comme $\forall u \in N[v], d(v) \leq \Delta_1(u)$ et comme $\forall w \in N[u], \min_{w \in N[u]} e^{(i-1)}(w) \leq e^{(i-1)}(v)$, nous pouvons écrire :

$$\frac{\min_{w \in N[u]} e^{(i-1)}(w)}{\Delta_1(u) + 1} \leq \frac{e^{(i-1)}(v)}{d(v) + 1}.$$

Et étant donné que $\Delta_1(z) \geq d(z) > d(v)$:

$$\frac{\min_{w \in N[z]} e^{(i-1)}(w)}{\Delta_1(z) + 1} < \frac{e^{(i-1)}(v)}{d(v) + 1}.$$

En injectant ces deux expressions dans la précédente, on obtient :

$$e^{(i)}(v) > e^{(i-1)}(v) - d(v) \cdot \frac{e^{(i-1)}(v)}{d(v) + 1} - \frac{e^{(i-1)}(v)}{d(v) + 1} > 0.$$

En conséquence, si, à une certaine étape, un nœud possède une bande passante non-nulle et ne possède pas le degré maximal dans son voisinage fermé, il disposera d'une bande passante résiduelle non nulle lors de l'étape suivante. Il est par ailleurs immédiat de montrer qu'un nœud ne possédant pas la bande passante résiduelle minimale parmi son voisinage à une certaine étape ne présente pas non plus une bande passante résiduelle nulle à l'étape suivante.

Un nœud n'atteint donc une bande passante résiduelle nulle que s'il possède à la fois la bande passante minimale et le degré maximal parmi ses voisins.

Supposons maintenant que $d(v) = \max_{u \in N_2[v]} d(u)$ et que $e^{(i-1)}(v) = \min_{u \in N_2[v]} e^{(i-1)}(u)$. Dans ce cas :

$$\forall u \in N[v], \min_{w \in N[z]} e^{(i-1)}(w) = e^{(i-1)}(v) \text{ et } \Delta_1(u) = d(v)$$

En conséquence,

$$\forall u \in N[v], x^{(i)}(u) = x^{(i-1)}(u) + \frac{e^{(i-1)}(v)}{d(v) + 1}$$

et donc,

$$e^{(i)}(v) = e^{(i-1)}(v) - (d(v) + 1) \cdot \frac{e^{(i-1)}(v)}{d(v) + 1} = 0$$

Lorsqu'un nœud atteint cette limite, il force l'ensemble de ses voisins à ne plus progresser. En effet, aucun voisin ne pourra allouer de bande passante puisqu'il possède un voisin dont la bande passante est nulle. Il existe donc au moins deux catégories de nœuds. Certains dont la bande passante résiduelle convergera vers 0 et les voisins de ceux-ci dont la progression de l'allocation sera bloquée par les précédents. Ces deux catégories ne sont, bien sûr pas disjointes. ■

Lemme 3.1.6. *Tous nœud possède au moins un voisin dont la bande passante résiduelle converge vers 0 :*

$$\forall v \in V, \exists u \in N[v] \text{ tel que } \lim_{i \rightarrow +\infty} e^{(i)}(u) = 0.$$

Preuve Étant donné que la suite $(x^{(i)}(v))_{i \in \mathbb{N}}$ est convergente, en utilisant la définition des termes de la suite, $x^{(i+1)}(v) = x^{(i)}(v) + \frac{1}{\Delta_1(v)+1} \min_{u \in N[v]} e^{(i)}(u)$, on peut constater que la grandeur $x^{(i+1)}(v) - x^{(i)}(v) = \frac{1}{\Delta_1(v)+1} \min_{u \in N[v]} e^{(i)}(u)$ converge vers 0. En conséquence, $\min_{u \in N[v]} e^{(i)}(u)$ converge aussi vers 0. ■

Les deux classes de nœuds décrites précédemment sont donc les deux seules classes de nœuds possibles. Dans notre effort pour caractériser la convergence de cette suite, il est encore possible de dégager quelques propriétés.

Lemme 3.1.7. *La suite des écarts $(x^{(i+1)}(v) - x^{(i)}(v))_{i \in \mathbb{N}}$ est décroissante.*

Preuve La suite $(x^{(i)}(v))_{i \in \mathbb{N}}$ étant croissante, la suite $(e^{(i)}(v))_{i \in \mathbb{N}}$ est décroissante.

$$\forall u \in N[v], e^{(i)}(u) \leq e^{(i-1)}(u).$$

En conséquence,

$$\exists z \in N[v] \text{ tel que } \min_{u \in N[v]} e^{(i-1)}(u) = e^{(i-1)}(z) \geq e^{(i)}(z) \geq \min_{u \in N[v]} e^{(i)}(u).$$

et

$$(x^{(i+1)}(v) - x^{(i)}(v)) \leq (x^{(i)}(v) - x^{(i-1)}(v)).$$

■

Lemme 3.1.8. *Le vecteur X constitué par les limites des suites $(x^{(i)}(v))_{i \in \mathbb{N}}$ est un optimum de Pareto.*

Preuve Considérons la relation d'ordre partiel $\leq$, ordre naturel sur $\mathbb{R}^N$, c'est-à-dire $x \leq y$ si et seulement si $x_i \leq y_i, \forall i \in [1, N]$. Un vecteur solution est un optimum de Pareto s'il est un maximum au sens de $\leq$ sur l'ensemble des solutions admissibles.

Considérons qu'il existe un vecteur S respectant l'ensemble de contraintes (3.1) et tel que $S > X$ au sens de l'ordre précédemment défini. Ceci signifie qu'il existe un nœud particulier $v \in V$ tel que $S(v) > X(v)$, $S(v)$ représentant la composante de S correspondant à un nœud v . Le lemme 3.1.6 indique qu'il existe un nœud $u \in N[v]$ disposant d'une bande passante résiduelle nulle. Ceci signifie que $\sum_{w \in N[u]} X(w) = b(u)$. Or, $S(v) > X(v)$ et S est conforme aux contraintes (3.1). Il existe donc nécessairement un nœud $z \in N[u] \setminus \{v\}$ tel que $S(z) < X(z)$. La propriété $S > X$ n'est donc pas vérifiée.

■

Nous avons maintenant identifié et caractérisé la suite d'allocations de bande passante que nous étudierons par la suite. Sa limite semble montrer des propriétés intéressantes au regard du problème que nous cherchons à résoudre. La vitesse de convergence reste le point faible de cette suite. Il existe un certain nombre de cas pathologiques, en particulier les topologies que l'on pourrait qualifier de périodiques, pour lesquelles les rapports entre les bandes passantes résiduelles sont identiques lors d'une étape à ce qu'il était une ou plusieurs étapes auparavant. Toutefois, nous verrons que cette suite converge rapidement vers une solution proche de sa limite.

3.2 Algorithme distribué dérivé de la suite d'allocations

L'algorithme 1 présenté ici est basé sur la suite définie en section 3.1.1. Il détermine itérativement les éléments successifs de la suite définie par le système d'équations 3.2. La bande passante allouée est stockée dans la variable X sur chaque nœud et la bande passante résiduelle dans la variable E . À l'issue de l'exécution de l'algorithme, X contient la limite de la suite précédemment définie. L'algorithme se base sur des informations transmises à chaque étape par les voisins à une distance d'un saut dans le graphe G . Les informations à collecter sont, à chaque étape $i \in \mathbb{N}$ et pour chaque voisin u , les valeurs $x^{(i)}(u)$ et $e^{(i)}(u)$. La transmission de ces informations requiert deux étapes distinctes de communication, étant donné que les deux suites sont interdépendantes.

Algorithme 1 – **Allocation de bande passante** (pour le nœud $v \in V$).

Données : la liste des voisins de v et $b(v)$ la bande passante à allouer dans le voisinage de v .

Résultat : X , la bande passante allouée au nœud v .

$E := b(v)$;

$X := 0$;

tant que X *n'est pas constant* **faire**

 transmettre E et d , degré du nœud, à tous ses voisins ;

 recevoir $E(u)$ et $d(u)$ de chaque voisin u ;

 Calculer $\Delta_1(v)$;

$X := X + \frac{1}{\Delta_1(v)+1} \min_{u \in N[v]} E(u)$;

 transmettre X à tous ses voisins ;

 recevoir $X(u)$ de chaque voisin u ;

$E := b(v) - \sum_{u \in N[v]} X(u)$;

Le lemme 3.1.4 montre que cet algorithme converge vers une solution X respectant l'ensemble de contraintes (3.1). À chaque étape, l'allocation calculée respecte, elle aussi, ces contraintes, comme le montre le lemme 3.1.2. L'allocation ainsi déterminée est équitable dans le sens où chaque mobile obtient une proportion de la bande passante non nulle (lemme 3.1.3) et dépendante de la densité et de la capacité de son voisinage.

Le lemme 3.1.6 montre que le réseau est pleinement utilisé, dans le sens où l'émission par un nœud quelconque d'un volume de données supérieur à la bande passante qui lui a été allouée conduira à un dépassement de la capacité du médium dans son voisinage. Il ne s'agit pas de l'allocation optimale en terme d'utilisation globale du réseau, comme nous le verrons par la suite.

Le lemme 3.1.7 montre que la différence entre deux allocations successives décroît. Si la limite de l'allocation peut être très longue à atteindre, une solution acceptable peut, en revanche, être déterminée au bout d'un nombre faible d'étapes. La procédure d'allocation pourra être considérée comme achevée lorsque la différence entre deux valeurs successives n'excède pas un seuil déterminé en fonction de la qualité de la solution recherchée.

Enfin, il est possible d'augmenter la vitesse de convergence si chaque mobile transmet à ses voisins la bande passante minimale restante dans son voisinage. Une valeur nulle indique que le nœud concerné ne prendra pas part aux prochains tours d'allocation. Il peut donc ne pas être pris en compte lors de la détermination du degré maximal présent dans le voisinage. Cependant, il convient de le considérer pour la détermination de la bande passante minimale dans le voisinage.

Si un certain nombre de propriétés de cet algorithme restent difficilement quantifiables, les divers lemmes présentés précédemment laissent entrevoir un comportement satisfaisant. Dans la suite de cette section, nous nous intéresserons aux performances de cet algorithme en terme de qualité de l'allocation mais aussi en terme de vitesse de convergence.

3.2.1 Simulation de l'algorithme

Afin d'obtenir une évaluation statistique et quantitative des performances de l'algorithme 1, un certain nombre de simulations ont été réalisées au moyen d'une implantation autonome réalisée en langage C++. Afin de déterminer son comportement exact, différentes formes de topologies ont été considérées.

Afin de déterminer la qualité de la solution obtenue, l'allocation résultante de l'application de cet algorithme sera comparée à la solution maximisant la bande passante allouée globalement dans le réseau la plus équitable ainsi qu'à la solution la plus équitable au sens de l'équité max-min [BM01, BG87].

Déterminer une solution optimale en terme de bande passante globale allouée revient à résoudre le programme linéaire suivant, connu sous le nom de *fractional packing* :

$$\max \sum_{v \in V} x(v), \text{ sous contraintes } \forall v \in V, \sum_{u \in N[v]} x(u) \leq b(v).$$

Une solution à ce problème peut être obtenue au moyen d'outils de programmation linéaire tels que l'algorithme du Simplex [Dan49]. Toutefois, la solution obtenue par ce biais n'est assujettie à aucune

contrainte d'équité. Certains nœuds pourront se voir affecter une bande passante nulle alors que d'autres obtiendront une allocation correspondant à la totalité de la capacité du médium. L'ensemble des solutions d'un tel programme linéaire n'est, en général, pas réduit à un seul élément, mais forme un convexe de $\mathbb{R}^n$. Afin de permettre une comparaison en terme d'équité entre la solution fournie par l'algorithme 1 et cette solution que nous qualifierons d'optimale par la suite, il convient de sélectionner dans cet ensemble la solution la plus équitable. Nous avons choisi comme critère d'équité l'écart-type des bandes passantes allouées à chaque nœud. L'algorithme du simplex nous permet de déterminer la valeur maximale de l'allocation totale de bande passante dans le réseau. Si l'on note x_{max} cette valeur, rechercher la solution minimisant l'écart-type entre les composantes parmi les solutions maximisant la bande passante totale allouée peut s'écrire :

$$\min \sqrt{\frac{\sum_{v \in V} (x(v) - \bar{x})^2}{|V|}}, \text{ sous contraintes } \begin{cases} \forall v \in V, \sum_{u \in N[v]} x(u) \leq b(v); \\ \sum_{v \in V} x(v) = x_{max}. \end{cases}$$

La fonction racine carrée étant croissante et le nombre de nœuds dans le réseau étant constant et positif, ce problème est équivalent au programme quadratique :

$$\min \sum_{v \in V} (x(v) - \bar{x})^2, \text{ sous contraintes } \begin{cases} \forall v \in V, \sum_{u \in N[v]} x(u) \leq b(v); \\ \sum_{v \in V} x(v) = x_{max}. \end{cases}$$

La résolution de ce programme quadratique a été confiée à la bibliothèque OOQP (*Object-Oriented Software for Quadratic Programming* [GW01]) développée par le département d'informatique de l'université du Wisconsin. La solution obtenue de cette manière pourra être qualifiée de solution optimale la plus équitable.

Nous comparerons aussi l'allocation déterminée par l'algorithme 1 à une solution considérée comme équitable mais non optimale *a priori* en terme de bande passante totale. La solution minimisant simplement l'écart-type conduirait à une allocation égale au minimum du rapport entre la bande passante disponible et le degré des nœuds. Cependant, appliquer une telle politique conduit à une allocation laissant des ressources inoccupées dans certaines zones. Nous comparerons donc notre solution à la solution la plus équitable au sens de l'équité max-min [BM01, BG87]. Dans cette allocation, aucune bande passante allouée ne peut être augmentée sans diminuer la bande passante allouée à un autre nœud à qui il a été octroyé une bande passante plus faible qu'au premier. Afin de déterminer les paramètres affectant les performances de l'algorithme 1, son fonctionnement a été simulé sur différentes topologies.

Topologies régulières

Les topologies régulières, telles que tout nœud possède exactement le même nombre de voisins, présentent, pour un degré donné, un diamètre croissant en fonction du nombre de nœuds uniquement. Par exemple, la figure 3.4 représente deux graphes réguliers comportant 8 nœuds et de degrés différents. L'exemple le plus simple en est les topologies en anneau dans lesquelles chaque nœud possède exactement deux voisins. Dans ce cas, le diamètre du graphe est égal à la moitié du nombre de nœuds, arrondi à l'inférieur. Considérons un graphe de n nœuds de degré d . Le lemme des poignées de main impose que le produit $d \cdot n$ soit pair. Ce type de graphe permet d'étudier en partie l'influence du degré et du nombre de nœuds sur le résultat de l'allocation.

Lorsque les bandes passantes disponibles des nœuds sont égales, l'algorithme converge en une étape. Dans ce cas, la solution optimale est identique à la solution max-min et à la solution obtenue par notre algorithme. Quel que soit le nombre de nœuds, la bande passante allouée à chaque nœud est identique et est égale à une proportion $1/(d+1)$ de la capacité du médium (d étant le degré des nœuds du graphe).

Considérons maintenant des graphes dans lesquels les bandes passantes initiales sont réparties aléatoirement entre des valeurs de 50 unités et 150 unités.

La figure 3.5 représente l'allocation moyenne en fonction du nombre de nœuds et du degré de ces nœuds. Cette figure représente la moyenne sur 1000 simulations différentes. L'allocation moyenne est encore une fois essentiellement fonction du degré des nœuds et non du nombre de nœuds ou du diamètre du graphe. La figure 3.6 représente le rapport entre la bande passante moyenne allouée dans le cas de la

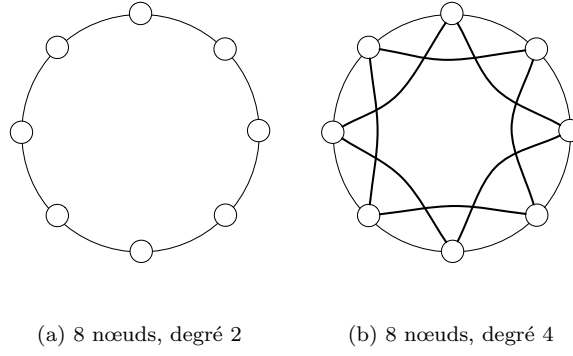


FIG. 3.4 – Exemples de topologies régulières

solution optimale et la bande passante allouée par notre algorithme. Le facteur entre les deux allocations possibles est compris entre 1 et 1,34 et est en moyenne de 1,18. La figure 3.7 représente le rapport entre l'allocation max-min et le résultat de l'algorithme. Le rapport varie peu est cette fois en moyenne de 1,01.

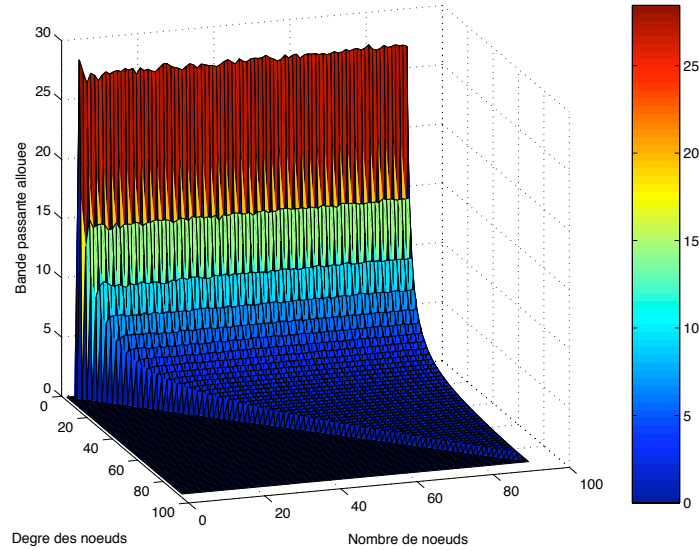


FIG. 3.5 – Allocation moyenne – graphes réguliers, bande passante aléatoire

La figure 3.8 représente l'écart-type de l'allocation réalisée par l'algorithme. Cet écart-type est cette fois uniquement fonction du nombre de nœuds ou du diamètre du graphe et non plus du degré des nœuds. La figure 3.9 représente le rapport entre l'écart-type de la solution optimale et celui de l'allocation réalisée par l'algorithme. Ce rapport croît en fonction du nombre de nœuds et du degré des nœuds et vaut en moyenne 33,9. De même, le rapport entre l'allocation max-min et l'allocation réalisée par l'algorithme est présenté en figure 3.10. Ce rapport vaut, en moyenne 2,6 et croît en fonction du nombre de nœuds et du degré des nœuds.

Les simulations menées sur ce type de graphes montrent que si le degré de nœuds a une influence sur la valeur moyenne de l'allocation de bande passante, il n'en est pas de même du nombre total de nœuds dans le réseau, ni du diamètre du graphe. En effet, les graphes possédant 100 nœuds de degré 2 ont un diamètre de 50 et l'algorithme y obtient les mêmes valeurs qu'appliqué sur un graphe de 10 nœuds de degré 2 présentant un diamètre de 5. Ceci confirme le caractère local de l'algorithme. Au niveau global, l'écart entre les bandes passantes initiales a une influence sur l'écart-type entre les différentes allocations.

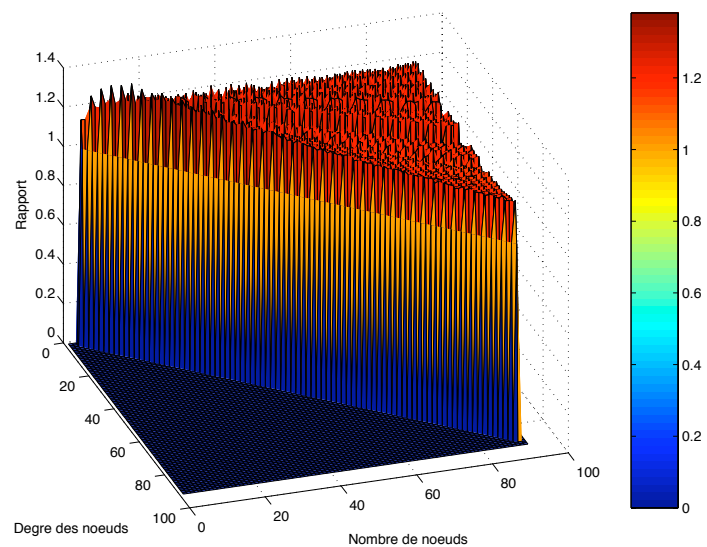


FIG. 3.6 – Rapport de l'allocation moyenne optimale et de l'allocation réalisée par l'algorithme – graphes réguliers, bande passante aléatoire

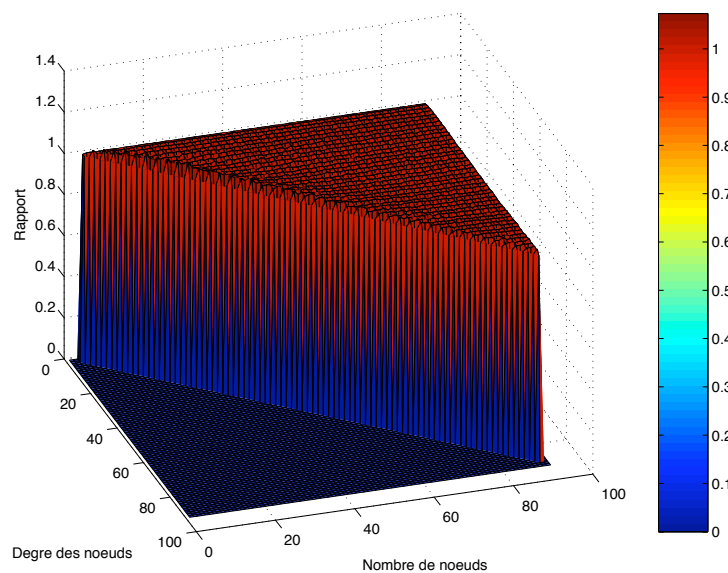


FIG. 3.7 – Rapport de l'allocation moyenne max-min et de l'allocation réalisée par l'algorithme – graphes réguliers, bande passante aléatoire

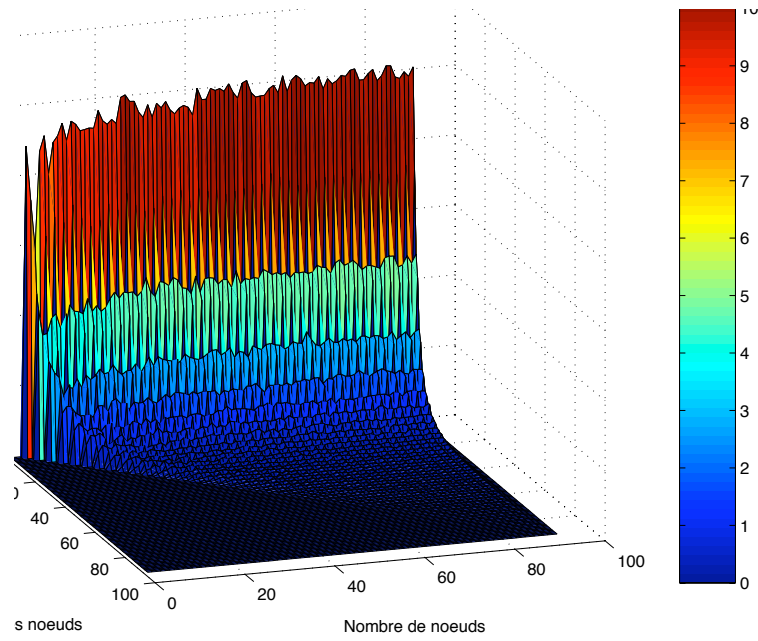


FIG. 3.8 – Écart-type – graphes réguliers, bande passante aléatoire

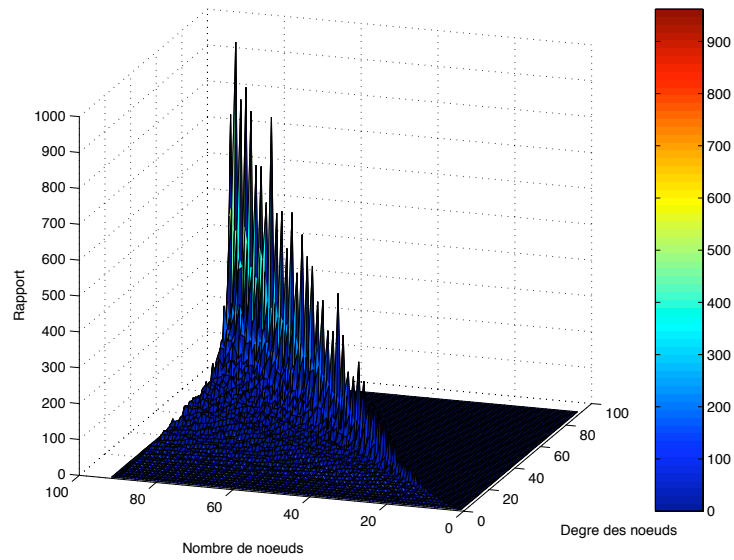


FIG. 3.9 – Rapport de l'écart-type de la solution optimale et de l'écart-type de l'allocation réalisée par l'algorithme – graphes réguliers, bande passante aléatoire

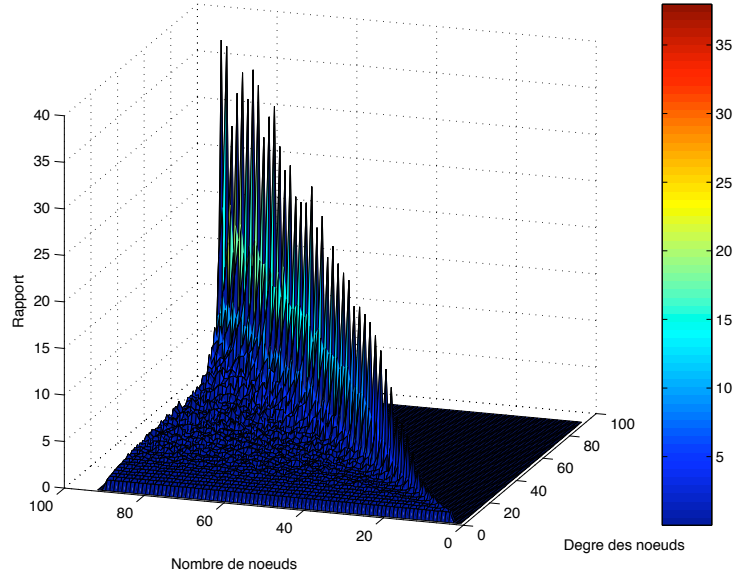


FIG. 3.10 – Rapport de l'écart-type de la solution max-min et de l'écart-type de l'allocation réalisée par l'algorithme – graphes réguliers, bande passante aléatoire

Graphes géométriques aléatoires

Les graphes géométriques aléatoires, dont une étude complète est disponible dans [Pen03], permettent de représenter fidèlement des réseaux *ad hoc* quelconques. Un tel graphe est généré en disposant de façon aléatoire un nombre déterminé de nœuds sur un plan et en leur associant une portée de communication. Lorsque deux mobiles sont à une distance inférieure à cette portée, un lien existe entre eux dans le graphe G . La portée de communication choisie va déterminer un certain nombre de paramètres du graphe associés tels que le degré moyen ou encore le nombre de composantes connexes du graphe. La figure 3.2.1 présente quelques propriétés des graphes géométriques aléatoires générés pour évaluer les performances de l'algorithme 1. Ces courbes représentent en fonction de la portée des nœuds dans un carré unité le degré moyen des nœuds, l'écart-type moyen entre les degrés des nœuds, le nombre de composantes connexes du graphe et la moyenne du maximum des degrés dans le voisinage fermé d'un nœud. Ces courbes représentent la moyenne de ces valeurs obtenues sur 10000 simulations pour des graphes de 10 à 100 nœuds.

Ces figures montrent que le degré moyen des nœuds augmente avec la portée de communication et que le degré maximum dans le voisinage d'un nœud augmente plus rapidement et atteint la limite dès que la portée atteint $2/3$ de la taille du carré contenant les nœuds. L'écart-type entre les degrés reste faible et atteint son niveau maximal de 1,6 pour une portée avoisinant les $2/3$ de la taille de la zone géographique contenant les nœuds. Enfin, dès que la portée atteint $1/2$, le réseau est connecté quel que soit le nombre de nœuds. Dans cette première analyse, nous ne nous intéresserons qu'à des réseaux connectés, même pour des faibles portées de transmission.

La figure 3.12 représente l'allocation moyenne réalisée par l'algorithme sur de tels graphes en fonction du nombre de nœuds et de la portée des nœuds. La figure 3.13 représente, quant à elle, la bande passante totale allouée au niveau du réseau. La bande passante initiale de chaque nœud est tirée aléatoirement entre 50 et 150 unités et les résultats présentés sont les résultats moyens obtenus sur 1000 simulations. La quantité de bande passante allouée semble être essentiellement fonction de la portée de transmission, donc du degré des nœuds. L'écart-type de cette allocation est représenté en figure 3.14 et présente les mêmes caractéristiques, il est essentiellement fonction du degré des nœuds.

Comparons maintenant le résultat obtenu à l'allocation optimale en terme de bande passante allouée présentant l'écart-type le plus faible et la solution la plus équitable au sens de l'équité max-min. La figure 3.15 représente le rapport entre l'allocation moyenne optimale et l'allocation réalisée par l'algorithme. Ce rapport est en moyenne de 1,3 et n'excède pas 1,7. La figure 3.15 représente le rapport entre

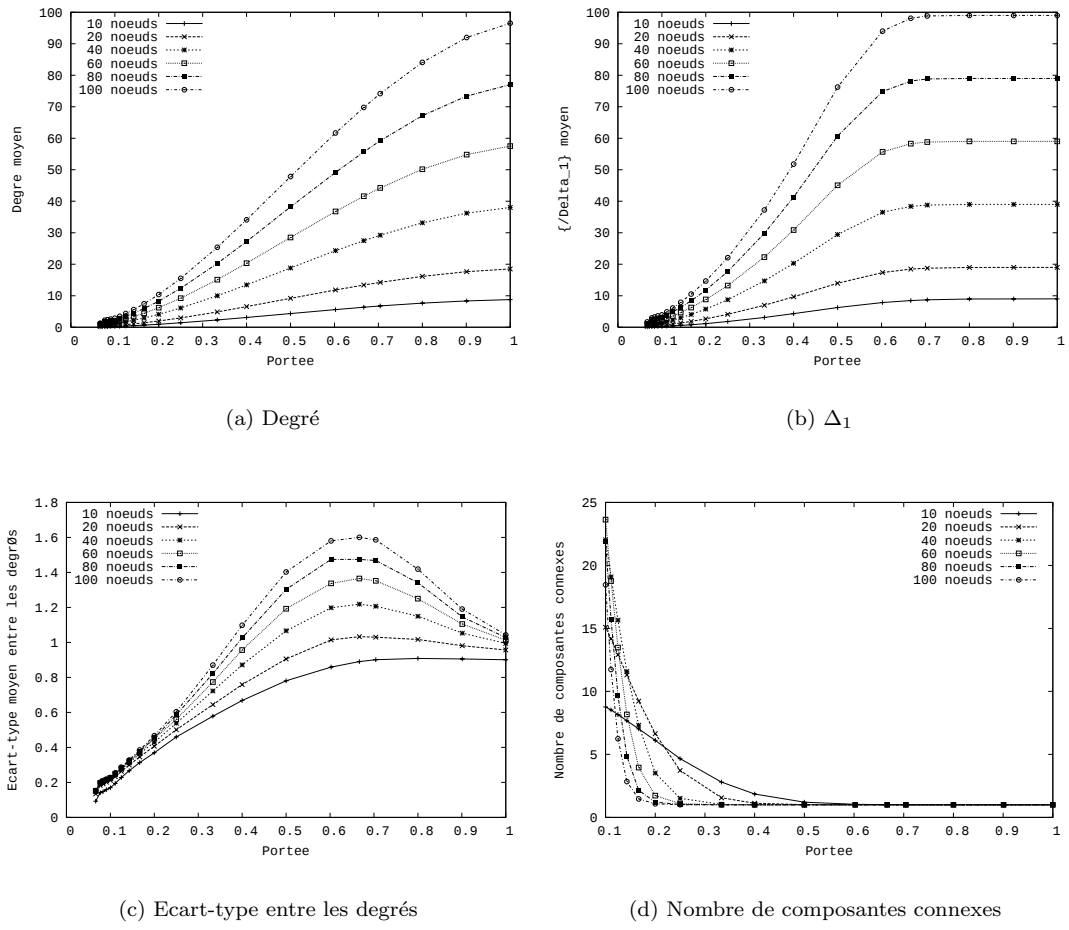


FIG. 3.11 – Propriétés des graphes géométriques aléatoires utilisés (Moyenne sur 10000 graphes)

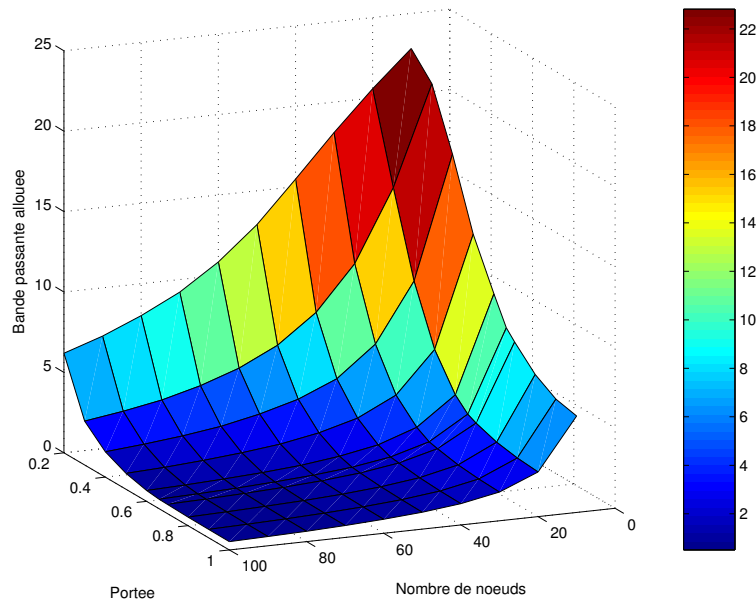


FIG. 3.12 – Allocation moyenne – graphes géométriques aléatoires, bande passante aléatoire

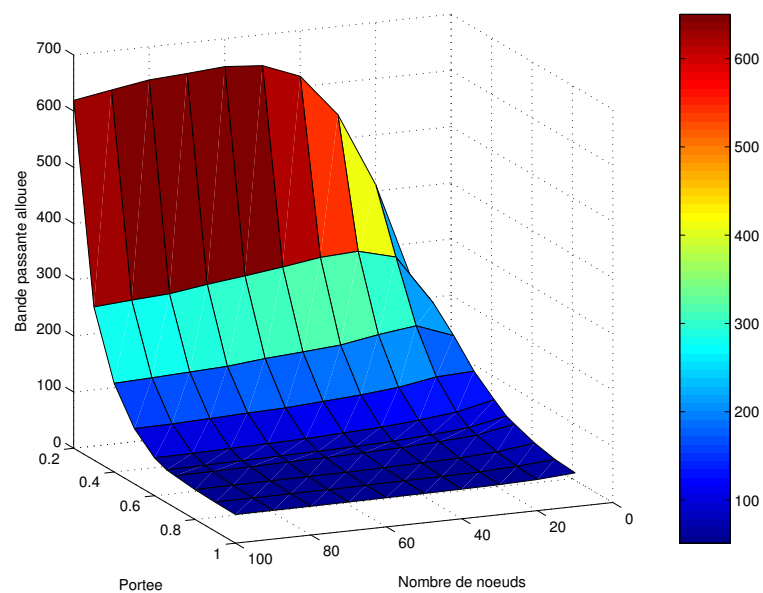


FIG. 3.13 – Allocation globale – graphes géométriques aléatoires, bande passante aléatoire

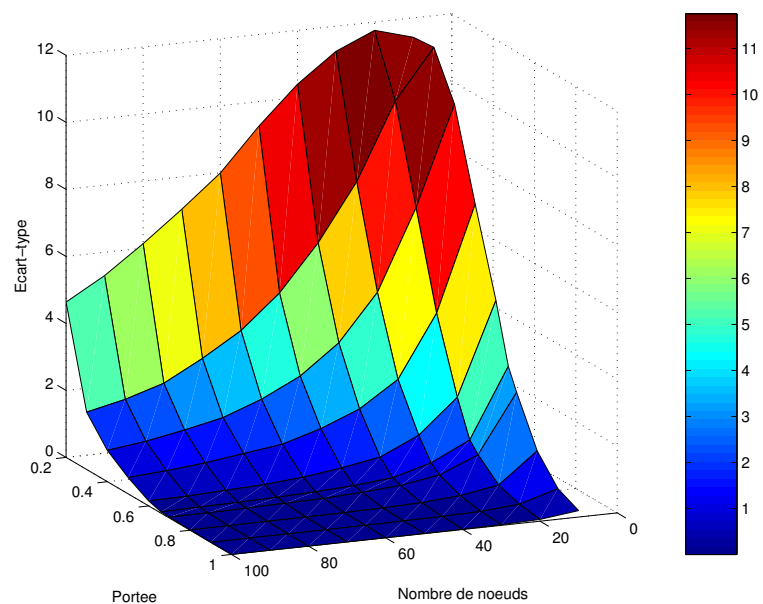


FIG. 3.14 – Écart-type – graphes géométriques aléatoires, bande passante aléatoire

l'allocation moyenne au sens de l'équité max-min et l'allocation réalisée par l'algorithme. Ce rapport est en moyenne de 1,05 et ne dépasse pas 1,15. L'allocation réalisée par l'algorithme est donc proche de l'allocation max-min. L'écart avec l'optimal est plus élevé, mais l'algorithme proposé conduit à une allocation sensiblement plus équitable, comme le montre la figure 3.17 représentant le rapport entre l'écart-type de l'allocation optimale et l'écart-type de l'allocation réalisée par notre algorithme. Ce rapport est en moyenne de 13 et peut atteindre 73. Enfin, la figure 3.18 représente le rapport entre l'écart-type de l'allocation max-min et l'allocation réalisée par l'algorithme. Ce rapport est en moyenne de l'ordre de 4 et peut atteindre 13,3. En rapprochant ces résultats des caractéristiques des graphes considérés, les rapports entre allocations moyennes semblent d'autant plus élevés que l'écart-type entre les degrés dans le graphe est élevé alors que les rapports entre écart-types semblent varier avec le degré moyen dans le graphe.

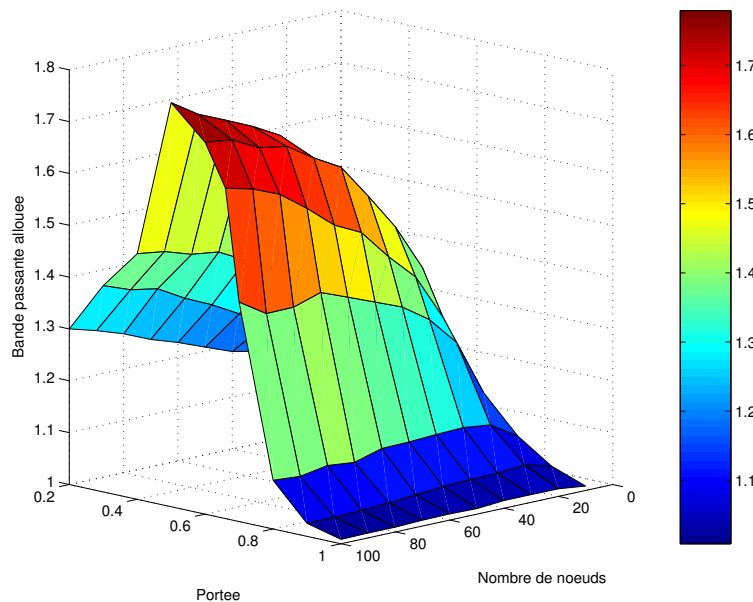


FIG. 3.15 – Rapport de l'allocation moyenne optimale et de l'allocation réalisée par l'algorithme – graphes géométriques aléatoires, bande passante aléatoire

L'allocation réalisée par notre algorithme est donc relativement proche, en terme de bande passante, de l'allocation la plus équitable au sens de l'équité max-min. L'écart-type entre les bandes passantes allouées à chaque nœud est faible. Le volume de bande passante allouée semble être fonction à la fois du degré moyen dans le graphe et de l'écart-type entre ces degrés. En revanche, l'écart-type des différentes allocations varie uniquement en fonction du degré du graphe. Cet algorithme présente en outre l'avantage d'être distribué, son exécution n'est conditionnée qu'à la transmission d'informations locales dans le voisinage de chaque nœud.

3.3 Vitesse de convergence

La vitesse de convergence de l'algorithme dépend essentiellement de la précision recherchée. Atteindre la limite de l'algorithme n'est pas toujours possible en un nombre fini d'étapes, comme le montre le scénario simple de la figure 3.19. Dans cette configuration, trois mobiles disposent des bandes passantes $2 \cdot b$, $3 \cdot b$ et $2 \cdot b$ comme indiqué sur la figure. Dans ce cas, l'algorithme, à la première étape, alloue $2 \cdot b/3$ à chaque mobile et les bandes passantes résiduelles sont alors respectivement $2 \cdot b/3$, b et $2 \cdot b/3$. Si $b \neq 0$, aucune de ces bandes passantes n'est nulle et il est possible de poursuivre l'allocation. Les nouvelles bandes passantes à allouer ont conservé le même rapport entre elles et l'algorithme convergera en un nombre infini d'étapes.

Toutefois, la bande passante restante sera, à chaque étape, divisée par 3 et deviendra vite quantité négligeable. La suite étant croissante, à chaque étape un volume de bande passante inférieur à celui qui avait été alloué à l'étape précédente est distribué. En conséquence, il est possible de considérer

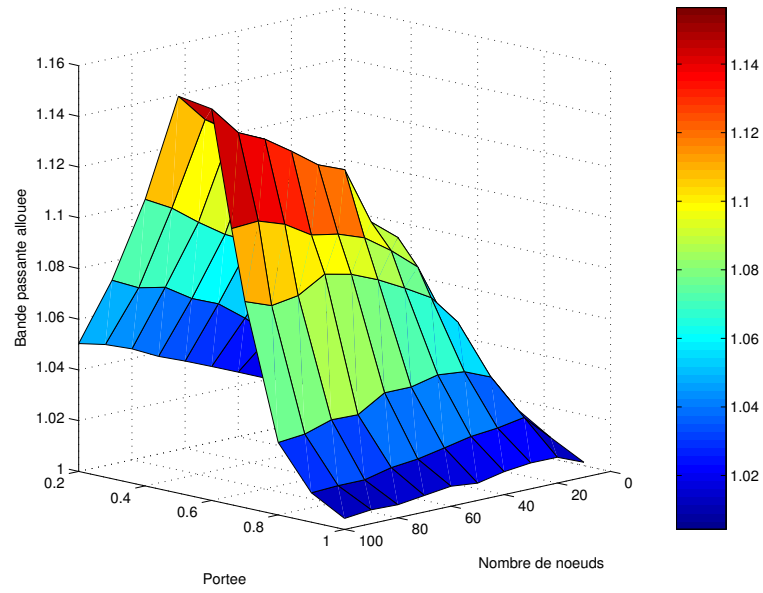


FIG. 3.16 – Rapport de l'allocation moyenne max-min et de l'allocation réalisée par l'algorithme – graphes géométriques aléatoires, bande passante aléatoire

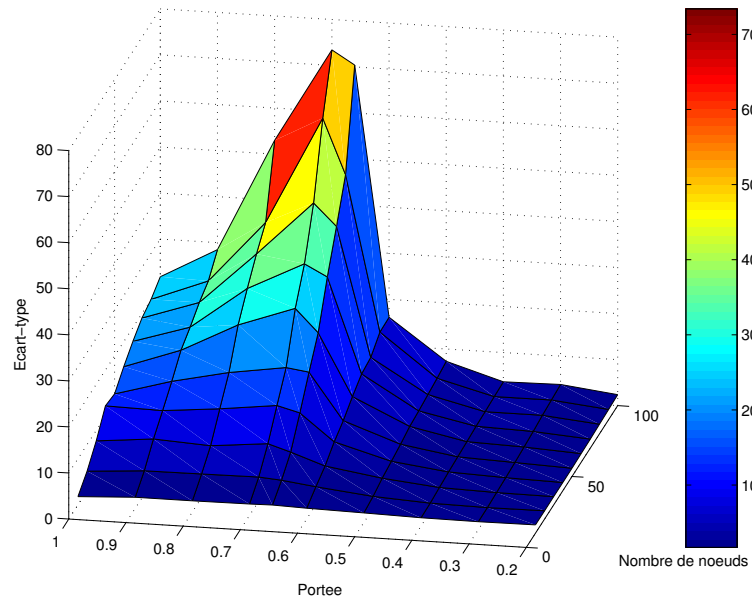


FIG. 3.17 – Rapport de l'écart-type de la solution optimale et de l'écart-type de l'allocation réalisée par l'algorithme – graphes géométriques aléatoires, bande passante aléatoire

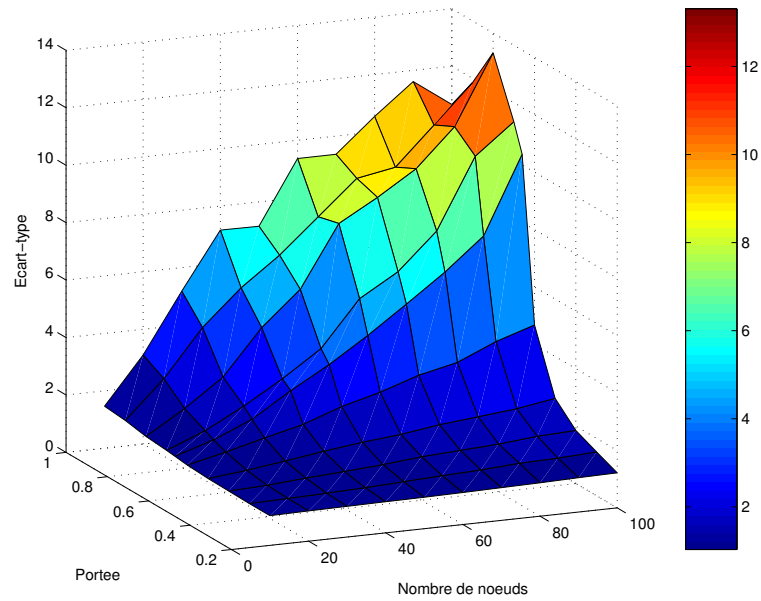


FIG. 3.18 – Rapport de l'écart-type de la solution max-min et de l'écart-type de l'allocation réalisée par l'algorithme – graphes géométriques aléatoires, bande passante aléatoire

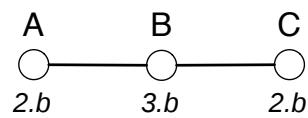


FIG. 3.19 – Scénario dans lequel l'algorithme converge en un nombre infini d'étapes

que l'exécution de l'algorithme est terminée lorsque l'allocation entre deux étapes diffère de moins d'un certain seuil.

La figure 3.20 représente le nombre d'étapes nécessaires à la convergence de l'algorithme pour les graphes réguliers présentés précédemment. Ce nombre d'étapes, valant en moyenne 2,8, semble être directement fonction du diamètre du graphe puisqu'il augmente lorsque le degré des nœuds décroît.

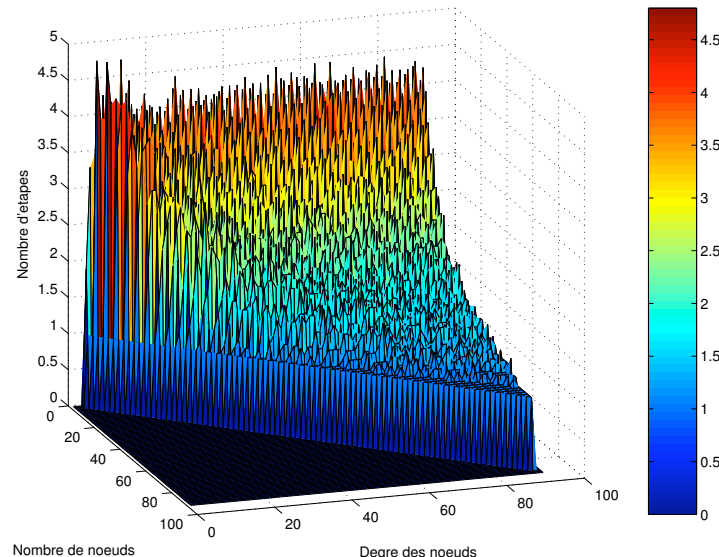


FIG. 3.20 – Nombre d'étapes avant convergence de l'algorithme – graphes réguliers, bande passante aléatoire

La figure 3.21 représente ce nombre d'étapes pour les graphes géométriques aléatoires précédents et vaut, en moyenne 5,38. L'allure de la courbe est semblable à l'allure de la figure 3.11(c) représentant l'écart-type entre les degrés dans un graphe géométrique aléatoire.

Afin de déterminer l'influence de la valeur du seuil d'arrêt sur la vitesse de convergence et sur la qualité de la solution, plusieurs simulations ont été réalisées en considérant trois valeurs de ce seuil : 0,1 %, 1 % et 10 %. Pour plus de clarté, nous ne présenterons dans ce cas que les résultats concernant des graphes comportant 50 nœuds, les résultats étant similaires pour les autres configurations. La figure 3.22 représente la bande passante allouée moyenne dans de tels graphes. La différence entre les allocations réalisées avec des seuils de 0,1 % et 1 % ne dépasse pas les 7 % et la différence entre les allocations réalisées avec des seuils de 0,1 % et 10 % atteint les 20 %. La figure 3.23 représente l'écart-type moyen entre les bandes passantes allouées. Ici, la différence est plus sensible. Cependant, on peut constater que poursuivre l'exécution de l'algorithme plus longtemps conduit à un accroissement de cet écart-type. La suite étant croissante, on peut en déduire que les dernières étapes ne concernent essentiellement que les nœuds disposant déjà d'une allocation élevée. Enfin, la figure 3.24 représente le nombre d'étapes nécessaires pour atteindre la limite de l'algorithme pour chacune des précisions considérées. Dans ce cas, la différence est flagrante. Sur les graphes considérés ici, l'algorithme atteint en à peine plus d'une étape en moyenne une valeur satisfaisant la condition d'arrêt si deux allocations successives diffèrent de moins de 10 %. Lorsque le seuil est fixé à 1 %, l'algorithme converge en environ 5 étapes pour une portée de communication telle que Δ_1 est sensiblement inférieur au nombre de nœuds du graphe, selon la figure 3.11(b). Enfin, pour un seuil de 0,1 %, le nombre d'étapes peut atteindre une valeur de 35. Ces résultats montrent que l'algorithme peut converger rapidement vers une solution acceptable et relativement proche de la limite. Les dernières étapes de l'algorithme ne servent qu'à raffiner l'allocation.

3.4 Gestion de la mobilité

L'algorithme 1 n'est utilisable que dans le cadre de réseaux statiques. En effet, la présence de nœuds mobiles dans un tel réseau conduirait à des situations dans lesquelles les contraintes (3.1) seraient violées.

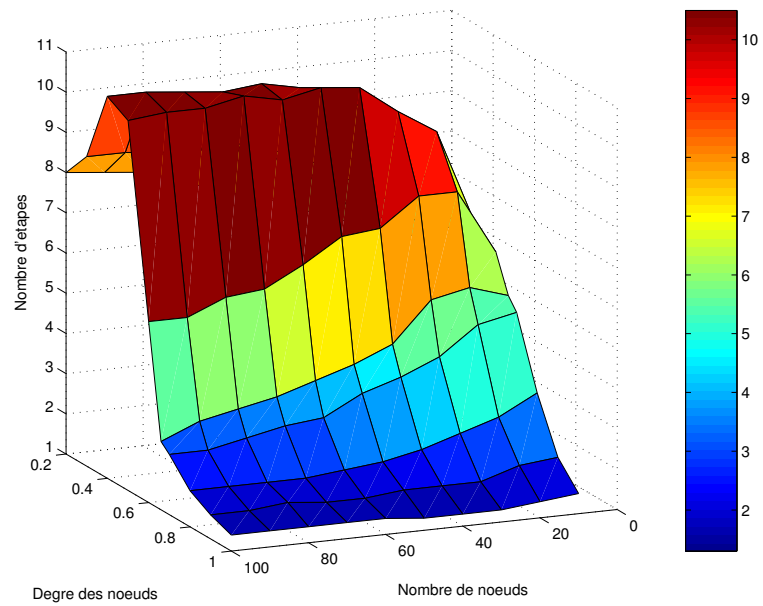


FIG. 3.21 – Nombre d'étapes avant convergence de l'algorithme – graphes géométriques aléatoires, bande passante aléatoire

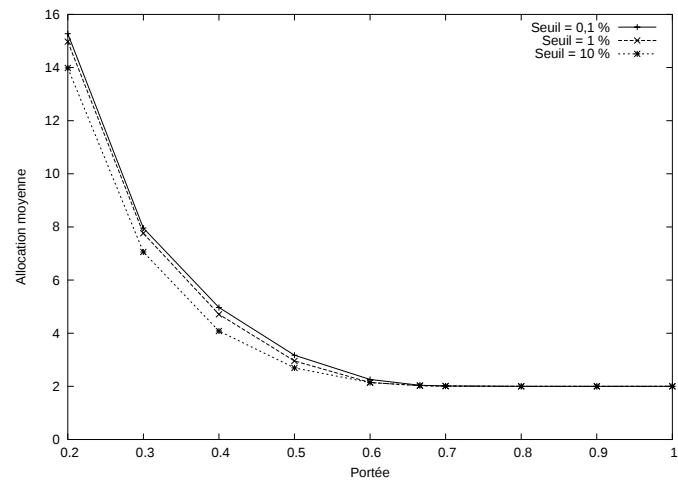


FIG. 3.22 – Allocation moyenne pour différentes valeurs du seuil d'arrêt – graphes géométriques aléatoires, 50 nœuds, bande passante aléatoire

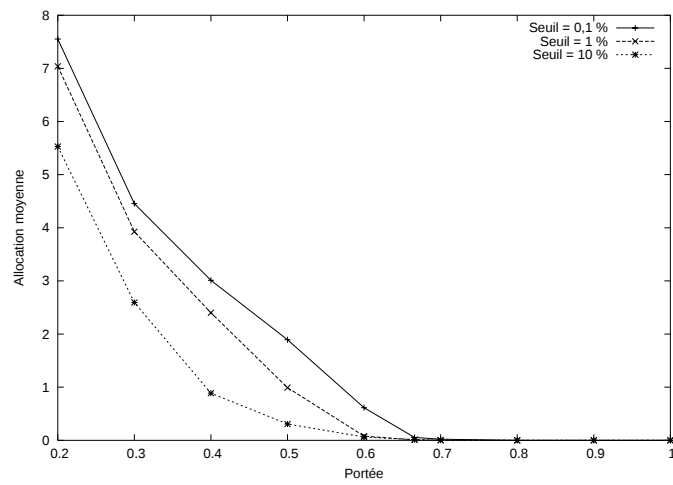


FIG. 3.23 – Écart-type moyen pour différentes valeurs du seuil d'arrêt – graphes géométriques aléatoires, 50 nœuds, bande passante aléatoire

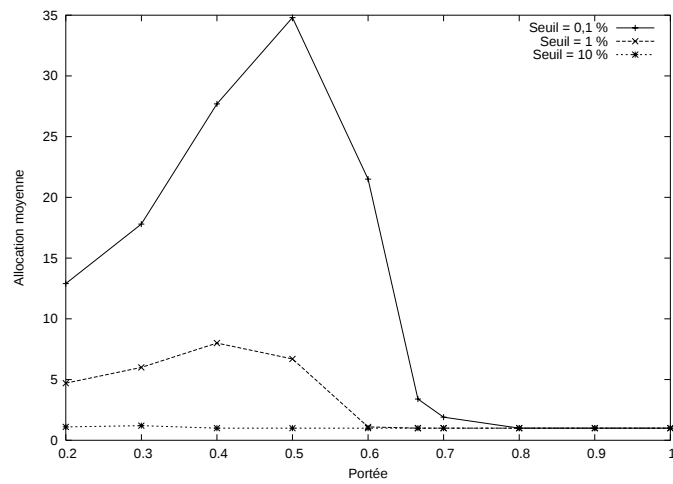


FIG. 3.24 – Nombre d'étapes moyen pour différentes valeurs du seuil d'arrêt – graphes géométriques aléatoires, 50 nœuds, bande passante aléatoire

Prenons, par exemple, le cas de deux mobiles isolés possédant une bande passante disponible identique. Chaque mobile s'alloue naturellement la totalité de la bande passante disponible. Si les deux mobiles se rapprochent, la capacité du canal sera dépassée dès qu'ils arriveront à portée de communication l'un de l'autre. Dans ce cas simple, chaque mobile percevra une bande passante disponible égale à -100% de la capacité initiale. L'étape suivante de l'algorithme les conduira alors à modifier leur allocation pour atteindre 50% de la capacité initiale. Cependant, si les bandes passantes initiales avaient été déséquilibrées — par exemple si l'un des mobiles ne disposait que d'une bande passante égale à 20% de l'autre —, appliquer l'algorithme 1 conduirait à une situation dans laquelle un mobile obtiendrait une allocation négative — le mobile le plus faible obtiendrait une allocation égale à -30% de la bande passante dans l'exemple.

Précisons dans un premier temps de quelle façon sera modélisée la mobilité des nœuds afin d'évaluer les performances des différentes adaptations. Les mobiles, durant une simulation, pourront se déplacer et auront à calculer une nouvelle allocation ainsi qu'à transmettre périodiquement des informations. Nous négligerons le temps de transmission d'un message ainsi que les problèmes pouvant survenir à la suite de collisions ou de pertes de paquets. En d'autres termes, les transmissions de messages seront considérées comme fiables, simultanées et instantanées. Le temps de calcul étant considéré, lui aussi, comme négligeable face à la vitesse de mobilité des nœuds, il est possible d'exécuter un nombre quelconque d'itérations de notre algorithme entre deux variations de la topologie.

Il s'agit, bien évidemment, d'une approximation extrêmement forte dont il faudra tenir compte lors de la traduction de cet algorithme en protocole. La capacité de communication des différents mobiles étant limitée, il ne sera possible d'échanger qu'un nombre limité de messages avant que les relations de voisinage dans le graphe n'évoluent. Ces considérations nous permettront de déterminer la mobilité maximale supportée par notre algorithme. Dans un premier temps, cependant nous considérerons que les mobiles se déplacent, formant éventuellement une nouvelle topologie, puis un nombre défini d'étapes de l'algorithme seront effectuées, conduisant à une nouvelle allocation et le processus se répètera.

Les résultats de simulations de scénarios impliquant des nœuds mobiles montrent que l'application simple de l'algorithme précédent conduit presque toujours à une allocation négative. Il est donc nécessaire de modifier cet algorithme afin de le rendre réactif aux changements de topologie locale aussi bien qu'aux variations dans la bande passante restant à distribuer. L'augmentation du nombre de nœuds dans un voisinage ou la diminution de la bande passante disponible devra provoquer au plus vite une libération de ressources par tous les nœuds présents dans la zone congestionnée. De même, une diminution dans le nombre de voisins en contention ou une augmentation de la bande passante disponible devra conduire à un partage des ressources libres.

Il n'est pas envisageable de recommencer globalement le processus d'allocation lorsque la mobilité des nœuds conduit à une modification de la topologie du réseau. En effet, un processus global est extrêmement coûteux en terme d'échanges de messages et l'influence des mobiles distants sur une allocation locale est faible.

Plusieurs stratégies locales sont envisageables permettant de réagir à une variation des caractéristiques du voisinage d'un nœud. Une première stratégie pourrait consister à remplacer le passage de l'étape i d'une allocation à l'étape $i + 1$ par l'expression suivante :

$$x^{(i+1)}(v) = \frac{d^{(i+1)}(v)}{d^{(i)}(v)} \cdot \left(x^{(i)}(v) + \frac{1}{\Delta_1(v) + 1} \cdot \min_{u \in N[v]} e^{(i)}(u) \right).$$

Une seconde stratégie envisageable consiste à tester toute allocation avant de l'appliquer. Lorsque l'application de l'algorithme d'allocation conduit à un dépassement de la capacité du médium, tous les mobiles saturés ainsi que tous leurs voisins recommenceront l'algorithme d'allocation depuis la première étape. Il est nécessaire d'impliquer les mobiles voisins des zones saturées dans ce processus, faute de quoi la nouvelle allocation ne serait pas conforme aux contraintes (3.1). Dans ce but, nous introduisons deux nouvelles suites $(y^{(i)}(v))_{i \in \mathbb{N}}$ et $(f^{(i)}(v))_{i \in \mathbb{N}}$ et modifions le calcul de $(x^{(i)}(v))_{i \in \mathbb{N}}$ et $(e^{(i)}(v))_{i \in \mathbb{N}}$ comme suit :

$$\text{left} \begin{cases} y^{(i+1)}(v) = x^{(i)}(v) + \frac{1}{\Delta_1(v)+1} \cdot \min_{u \in N[v]} e^{(i)}(u); \\ f^{(i+1)}(v) = b(v) - \sum_{u \in N[v]} y^{(i+1)}(u). \end{cases}$$

$$\text{Si } \min_{u \in N[v]} f^{(i+1)}(u) \geq 0, \text{ alors } \begin{cases} x^{(i+1)}(v) = y^{(i+1)}(v); \\ e^{(i+1)}(v) = f^{(i+1)}(v). \end{cases}$$

$$\text{Et si } \min_{u \in N[v]} f^{(i+1)}(u) < 0, \text{ alors } \begin{cases} x^{(i+1)}(v) = 0; \\ e^{(i+1)}(v) = b(v) - \sum_{u \in N[v]} x^{(i+1)}(u). \end{cases}$$

Lemme 3.4.1. *L'allocation décrite par les équations précédentes demeure positive pour chaque nœud et respecte les contraintes (3.1). $\forall i \in \mathbb{N}, \forall v \in V, x^{(i)}(v) \geq 0$ and $e^{(i)}(v) \geq 0$*

Preuve Cette propriété peut être démontrée par récurrence. Dans un premier temps, montrons que si, à l'étape $i \in \mathbb{N}$, $\forall v \in V, x^{(i)}(v) \geq 0$ et $e^{(i)}(v) \geq 0$, alors ces propriétés sont toujours valides à l'étape $i + 1$.

Supposons que, dans ces conditions, $\exists v \in V, e^{(i+1)}(v) < 0$. Nécessairement, $\exists u \in N[v], x^{(i+1)}(u) > 0$. En se référant à la définition des suites, cela signifie que $\min_{w \in N[u]} f^{(i)}(w) \geq 0$.

Étant donné que $x^{(i)}(u) \geq 0$ et comme $v \in N[u]$, on peut vérifier la relation suivante, ce qui contredit l'hypothèse de départ :

$$e^{(i+1)}(v) \geq f^{(i+1)}(v) \geq \min_{w \in N[u]} f^{(i+1)}(w) \geq 0.$$

En ce qui concerne $x^{(i+1)}(v)$, cette valeur est soit nulle soit égale à $x^{(i)}(v) + \frac{1}{\Delta_1(v)+1} \cdot \min_{u \in N[v]} e^{(i)}(u)$, somme de termes positifs.

En conséquence, $\forall v \in V, x^{(i)}(v) \geq 0$ et $e^{(i)}(v) \geq 0 \Rightarrow \forall v \in V, x^{(i+1)}(v) \geq 0$ and $e^{(i+1)}(v) \geq 0$. Compte tenu du fait que tous les mobiles ont au départ une allocation nulle et une bande passante disponible positive, la propriété $\forall i \in \mathbb{N}, \forall v \in V, x^{(i)}(v) \geq 0$ and $e^{(i)}(v) \geq 0$ est vraie. ■

Le lemme 3.1.1 est donc toujours valide pour cette suite d'allocations et à chaque étape, $i \in \mathbb{N}$, $(x^{(i)}(v), v \in V)$ respecte les contraintes (3.1). Une version plus faible du lemme 3.1.3 est aussi vérifiée. Compte tenu de la mobilité des nœuds, il est possible d'imaginer un scénario dans lequel un nœud se déplace de zone surchargée en zone surchargée, indéfiniment et ne parvient pas à obtenir une allocation strictement positive. Cependant ce type de situation est peu probable si l'on considère des réseaux présentant une mobilité raisonnable.

Le lemme 3.1.4 concernant la convergence de la suite d'allocations n'est cependant plus valide. En effet, la suite des allocations d'un nœud n'est plus monotone puisqu'elle peut souffrir de retours à zéro occasionnels. Cependant, si la topologie se stabilise, l'allocation converge à nouveau. La limite de l'allocation ne sera pas identique à celle qui aurait été atteinte sur une topologie identique stable puisque le point de départ de cette allocation est différent.

L'algorithme 2 dérive de l'algorithme 1 et permet le calcul itératif et distribué des suites précédentes. Cet algorithme requiert trois phases de communication par étape et un soin particulier devra être apporté à sa traduction en protocole.

Algorithme 2 – **Allocation de bande passante dans un contexte mobile** (pour le nœud $v \in V$).

Données : la liste des voisins de v et $b(v)$ la bande passante à allouer dans le voisinage de v .

Résultat : X , la bande passante allouée au nœud v .

$E := b(v)$;

$X := 0$;

$d := d(v)$;

tant que *le réseau existe*, **faire**

 transmettre E et d à tous ses voisins ;

 recevoir $E(u)$ et $d(u)$ de chaque voisin u ;

$Y := X + \frac{1}{\Delta_1(v)+1} \min_{u \in N[v]} E(u)$;

 transmettre Y à tous ses voisins ;

 recevoir $Y(u)$ de chaque voisin u ;

$F := b(v) - \sum_{u \in N[v]} Y(u)$;

si $F \geq 0$, **alors**

$X := Y$;

sinon

$X := 0$;

 transmettre X à tous ses voisins ;

 recevoir $X(u)$ de chaque voisin u ;

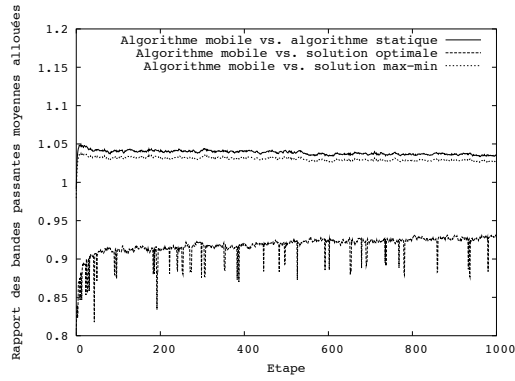
$E := b(v) - \sum X(u)$;

3.4.1 Simulation de l'algorithme mobile

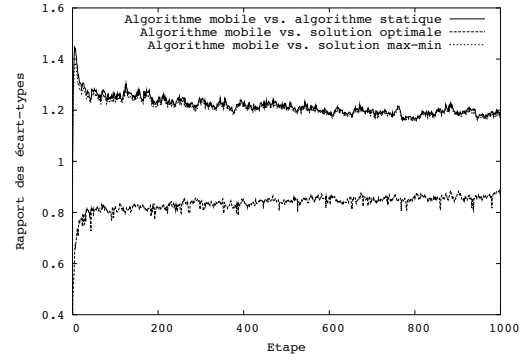
L'algorithme 2 a été simulé de la même manière que le précédent. Les simulations ont été réalisées sur des graphes géométriques aléatoires en utilisant un modèle de mobilité de type direction aléatoire (*Random direction* [RMSM01]). En effet, ce modèle de mobilité présente de bonnes propriétés stochastiques dans le sens où les mobiles couvrent tout l'espace et ne provoquent pas de variations par vagues de la densité locale de nœuds. Un nœud suivant ce modèle de mobilité tire une direction aléatoire ainsi qu'une vitesse aléatoire. Il suit cette direction jusqu'à rencontrer le bord de la zone de simulation. Il tire alors une nouvelle direction et une nouvelle vitesse et réitère le processus.

Pour évaluer les performances de cet algorithme, les mobiles exécuteront un nombre défini d'étapes de l'algorithme, se déplaceront d'une distance correspondant à une unité de temps, puis recommenceront ce processus. Chacun des résultats de simulation présenté ci-après représente la valeur moyenne de la grandeur considérée sur un total de 1000 simulations. Chaque simulation est composée de 1000 étapes de mouvement puis de calcul. La figure 3.4.1 représente le rapport entre les résultats obtenus en appliquant l'algorithme mobile 2 et le résultat obtenu par l'une des trois méthodes précédentes appliquées sur la topologie courante considérée comme statique : le résultat obtenu par l'algorithme 1, la solution optimale en terme de bande passante totale présentant l'écart-type le plus faible et la solution la plus équitable au sens de l'équité max-min. Les simulations présentées ici ont été réalisées sur un réseau comportant 20 nœuds disposés dans un carré unité et dont la portée est de 0,13. Après une phase d'initialisation, les différents rapports deviennent presque constants. L'algorithme mobile conduit à une allocation moyenne plus élevée que l'algorithme statique et que la solution max-min, mais il conduit aussi à un écart-type plus élevé que ces deux dernières. En comparaison de la solution optimale en terme de bande passante allouée, l'algorithme mobile conduit à une allocation plus faible d'environ 10 % mais présentant un écart-type plus faible. Ce type de simulation a été réalisé pour différentes valeurs des paramètres du réseau tels que le nombre de nœuds ou la portée des nœuds. Les conclusions sont similaires même si les valeurs des rapports diffèrent quelque peu.

La figure 3.26 représente l'influence du nombre d'étapes de l'algorithme mobile exécutées entre deux mouvements des nœuds du réseau. Les courbes, obtenues pour un réseau de 20 nœuds ayant une portée de 0,13, représentent le rapport entre le résultat de l'application de l'algorithme mobile et le résultat de l'algorithme initial appliqué sur la topologie considérée comme statique après un mouvement. Accroître le nombre d'étapes entre deux changements de topologie résulte en une augmentation à la fois de l'allocation



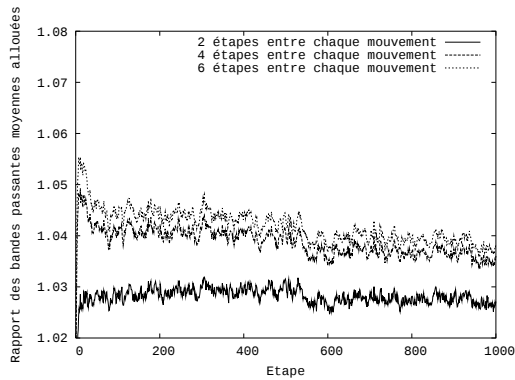
(a) Allocation moyenne



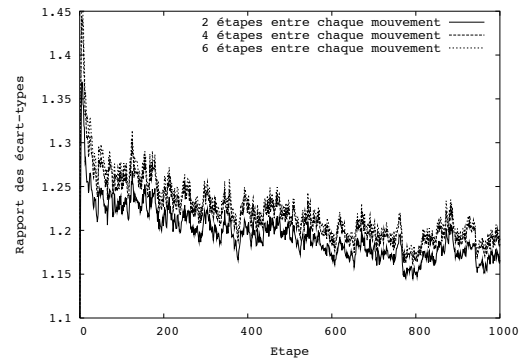
(b) Écart-type

FIG. 3.25 – Comparaison entre l'algorithme mobile et les différentes allocations statiques

moyenne et de l'écart-type, à l'identique de ce qui avait été constaté sur l'algorithme statique en faisant évoluer la condition d'arrêt. Les dernières étapes de l'algorithme servent à raffiner l'allocation. Ces résultats restent cependant proches et il est intéressant de noter qu'un faible nombre d'étapes suffit pour obtenir une solution convenable. L'algorithme pourra donc *a priori* supporter une mobilité relativement élevée.



(a) Allocation moyenne



(b) Écart-type

FIG. 3.26 – Influence du nombre d'étapes entre deux mouvements sur la qualité de l'allocation

Enfin, la figure 3.27 représente l'influence des caractéristiques du réseau sur la qualité de l'allocation réalisée par l'algorithme mobile. Les différentes courbes, lissées pour plus de lisibilité, représentent le rapport entre les résultats de l'algorithme mobile et les résultats de l'algorithme statique appliqué sur la topologie après chaque mouvement sur un réseau de 20 mobiles et pour différentes valeurs de portée de transmission. Lorsque la portée est élevée, le graphe résultant est presque toujours complet et l'allocation est similaire à l'allocation réalisée par l'algorithme statique. Diminuer la portée des nœuds accroît tout d'abord ce rapport. Les graphes utilisés ici ne sont plus obligatoirement connexes, à l'inverse de lors de l'étude de l'algorithme statique puisqu'il est impossible de garantir la connexité d'un réseau mobile. Aussi, lorsque le réseau devient partitionné en plusieurs composantes connexes, ce qui est le cas pour

un rayon de transmission inférieur à 0,3 pour un réseau comportant 20 nœuds, le rapport entre les allocations décroît à nouveau.

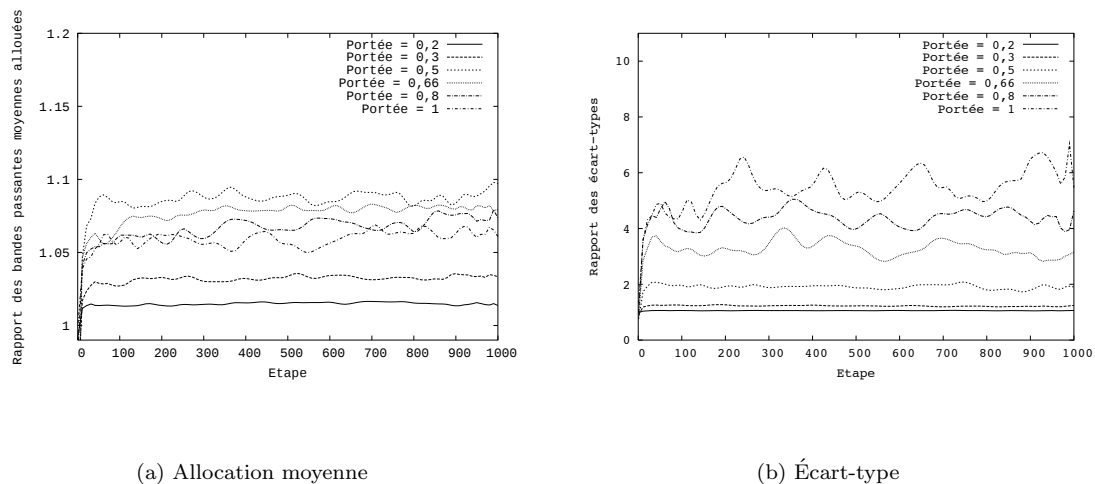


FIG. 3.27 – Influence des caractéristiques de la topologie sur l'allocation

3.5 Conclusion

Dans ce chapitre, nous avons présenté deux algorithmes distribués permettant de réaliser une allocation équitable de bande passante dans les réseaux *ad hoc*, approchant ainsi la solution d'un problème d'optimisation de façon distribuée. Les caractéristiques de ces algorithmes déterminées par une étude théorique et par des simulations montrent que l'allocation obtenue représente un bon compromis entre équité et utilisation globale du réseau. Les performances obtenues par l'algorithme mobile, résolvant le problème de la mobilité en réinitialisant le processus d'allocation au niveau des nœuds constatant une violation des contraintes imposées, restent proches de celles de son homologue statique. D'autres stratégies peuvent bien évidemment être imaginées afin de gérer la mobilité. Les résultats obtenus par cet algorithme sont cependant encourageants et il est maintenant nécessaire d'en dériver un protocole d'allocation de bande passante.

Cependant, la conception de ce protocole devra prendre en compte un certain nombre de contraintes supplémentaires. Tout d'abord, les communications ne s'effectuent pas en temps nul et le délai nécessaire à la collecte d'un volume d'informations suffisant pour itérer l'algorithme limitera la mobilité maximale du réseau. De plus, les transmissions de messages ont lieu de manière asynchrone et un mécanisme doit être mis en œuvre afin de s'assurer avant de déclencher le calcul d'une étape de l'algorithme, que tous les mobiles ont bien réalisé le calcul de l'étape précédente. Ceci peut être réalisé en choisissant un délai suffisamment grand entre deux étapes ou de façon explicite par le biais de messages de contrôle. Enfin, il est nécessaire de prendre en compte le mode de transmission des informations imposé par les protocoles sous-jacents. Les paquets de contrôle permettant aux nœuds d'itérer cet algorithme peuvent être soumis à des collisions.

Tous ces problèmes peuvent être résolus en accroissant le volume de trafic de contrôle généré. Cependant, comme le montrera le chapitre suivant, il n'est pas toujours pertinent de chercher à obtenir les informations les plus précises possible, si cette précision engendre une perte de performance globale trop importante. Toutefois, l'algorithme présenté ici conserve un comportement local et il est possible d'imaginer un mécanisme adaptatif adaptant la fréquence d'émission des paquets de contrôle en fonction de la densité du réseau, de la mobilité, etc.

4.1 Introduction

Il est difficile à l'heure actuelle de prévoir ce que sera le cadre d'utilisation exact des réseaux *ad hoc*. Sans doute ces réseaux seront-ils utilisés de manière autonome, formant des réseaux locaux. Sans doute seront-ils aussi utilisés comme extensions sans fil de réseaux filaires. Dans ce cadre, l'architecture Mobile IP [Per96] définit les modalités de la gestion de la mobilité des utilisateurs entre domaines. À l'intérieur d'un même domaine, le suivi de la mobilité des utilisateurs peut être confié à un protocole de micromobilité. Dans ce chapitre, nous nous intéresserons au fonctionnement de l'un de ces protocoles de micromobilité dans des réseaux hybrides constitués par une infrastructure fixe mais interconnectée au moyen de liens sans fil et par un réseau *ad hoc* sous-jacent.

Les protocoles de micromobilité utilisent certains messages de contrôle afin de garder une trace de la localisation des différents mobiles et de maintenir l'arbre de routage. Ces messages sont transmis en mode point à point, limitant la transmission des informations au récepteur ou en mode diffusion, rendant impossible la détection de collision ou de perte du paquet. Or, l'infrastructure de ces réseaux est organisée en arbre dont la racine est un routeur particulier jouant le rôle de passerelle et pouvant ainsi fournir accès à Internet et à différents services aux mobiles. Il existe donc une relation de hiérarchie entre les différents routeurs d'infrastructure et il est possible d'utiliser cette hiérarchie pour mettre en place et évaluer un mécanisme de transmission des informations intermédiaire, c'est-à-dire une diffusion acquittée.

L'étude de la transmission des notifications de mobilité dans un tel réseau constitue une première étude de la diffusion de paquets de contrôle de manière fiable. Ce chapitre présente tout d'abord en détail l'architecture considérée, les mécanismes et le but de la notification de mobilité. Différentes optimisations sont alors étudiées afin de déterminer le bon compromis entre volume de trafic de contrôle et précision des informations véhiculées par ces mêmes paquets de contrôle.

4.2 L'architecture Mobile IP et la micromobilité

Dans un contexte fortement mobile, il est aisé d'imaginer un utilisateur mobile, muni d'un terminal portable tel un assistant personnel, se déplaçant dans une zone géographique offrant une couverture réseau. Cet utilisateur sera sans doute amené, au gré de ses mouvements, à changer de réseau d'accès, profitant ainsi d'une couverture étendue par l'agrégation de plusieurs réseaux. Il est légitime, dans ce cas, de souhaiter conserver une continuité dans les communications établies, à l'image de la téléphonie mobile.

Cependant, le routage dans les réseaux IP est basé sur l'hypothèse qu'une adresse IP est attribuée à un terminal fixe faisant partie d'un réseau particulier. Un mobile dont l'adresse IP ne correspondrait pas à l'adresse du réseau dont il fait partie ne serait pas accessible de l'extérieur. Un changement de localisation s'accompagne donc d'un changement d'adresse IP. Toutefois, la préservation des connexions TCP, par exemple, requiert l'utilisation d'une adresse IP constante. Comment, alors, permettre une mobilité inter-domaines des terminaux transparente pour les utilisateurs ?

4.2.1 Mobile IP

L'architecture Mobile IP [Per96], conçue par l'IETF, résout ce problème en associant à chaque terminal deux adresses IP distinctes. La première adresse (appelée *home address*) est sa véritable adresse, associée à son réseau d'origine. La seconde adresse (appelée *care-of-address*) est fonction du domaine visité

et permet au mobile d'être joint depuis l'extérieur. Un terminal entrant dans un nouveau domaine devra, dans un premier temps, déterminer cette *care-of-address* identifiant un membre particulier du réseau visité appelé agent étranger (*Foreign Agent*). L'agent étranger peut correspondre, selon la politique en œuvre dans le réseau visité, à la passerelle d'accès au réseau, à un nœud distinct, ou au terminal visiteur lui-même. Dans ce dernier cas, la *care-of-address* pourra être obtenue dynamiquement (par exemple par le biais de DHCP) et sera nommée, selon la terminologie de Mobile IP, *colocated care-of-address*.

La deuxième étape du processus d'attachement à un réseau consiste à transmettre l'adresse locale (*care-of-address*) à un routeur particulier appartenant au réseau d'origine du terminal appelé agent mère (*Home agent*). Ce dernier établit alors un tunnel IP entre lui-même et l'agent étranger dont l'adresse lui a été communiquée. L'architecture résultante est représentée en figure 4.1. Dès lors, lorsqu'un paquet à destination du terminal mobile sera dirigé vers l'agent mère par le routage IP en vigueur dans Internet, ce dernier le transmettra *via* le tunnel à l'agent étranger. Ce dernier se chargera alors de le transmettre au terminal destinataire. Les paquets originaux du terminal mobile, quant à eux, prendront un chemin direct vers leur destination.

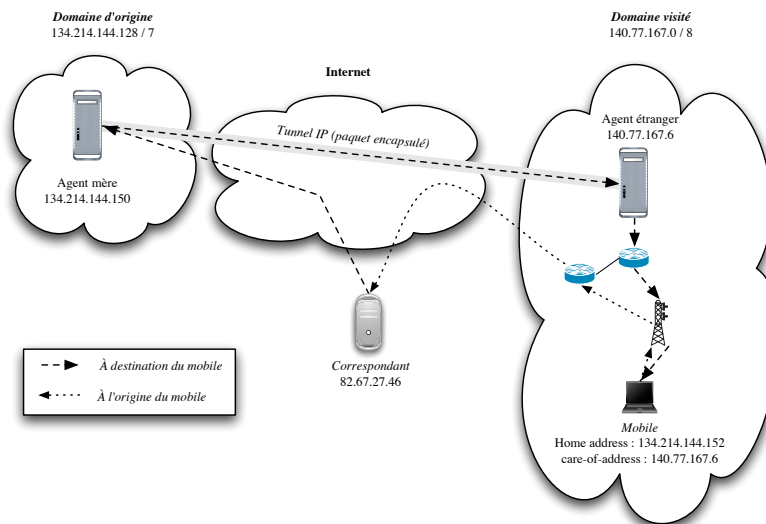


FIG. 4.1 – Architecture d'une communication Mobile IP

Mobile IP est aujourd'hui largement adopté. La plupart des routeurs intègrent ce protocole. Si le routage triangulaire introduit un délai supplémentaire dans l'acheminement des paquets à destination du terminal mobile, un certain nombre d'optimisations ont été proposées afin de permettre une communication directe entre terminaux. Certains problèmes liés à la sécurité des réseaux ont par ailleurs été mis en lumière. Par exemple, l'utilisation d'un pare-feu bloquant les paquets sortants du domaine visité qui possèdent une adresse source extérieure au domaine pourrait empêcher l'émission de paquets par le mobile visiteur. L'utilisation d'un pare-feu empêchant les paquets ayant une adresse source appartenant au réseau mère d'entrer dans le réseau mère pourrait compliquer la communication entre un mobile et son agent mère. Résoudre ce type de situations peut difficilement incomber aux pare-feux, car modifier les règles de filtrage pour permettre l'utilisation de Mobile IP créerait des failles dans la sécurité des réseaux concernés. Aussi, certaines modifications peuvent être apportées à Mobile IP pour permettre son utilisation dans des environnements sécurisés.

4.2.2 Micromobilité

La mobilité, telle qu'elle est définie par l'architecture Mobile IP, est appelée macromobilité. Les protocoles de micromobilité cherchent quant à eux, à améliorer les performances de Mobile IP lorsque les terminaux mobiles changent régulièrement de point d'attache au sein d'un même domaine (cette opération est appelée *handoff* ou *handover*). Par exemple, dans le cas d'un réseau comportant plusieurs stations de base sans fil offrant une couverture étendue, la communication de l'adresse du nouveau

point d'attache à l'agent mère nécessite l'échange de messages transitant par Internet, engendrant un certain volume de trafic de contrôle et allongeant la durée du *handoff*. Dans [PJ01], Perkins et Johnson proposent d'optimiser la procédure de changement d'agent étranger. Un mobile sélectionnant un nouvel agent étranger notifie l'ancien au moment du *handoff* afin que ce dernier transmette les paquets qui lui parviendront encore au nouvel agent étranger. Cette approche permet de réduire le délai nécessaire au changement de point d'attache, mais engendre toujours un volume de trafic de contrôle transitant par Internet.

Les protocoles de micromobilité proposent de gérer localement la mobilité des terminaux tant qu'ils demeurent dans le même domaine. La gestion de la macromobilité, c'est-à-dire des déplacements des mobiles entre domaines est laissée aux soins de Mobile IP. Ces protocoles permettent un changement rapide de station de base peu coûteux en terme d'utilisation des ressources du réseau. Ils sont étudiés au sein du groupe de travail *Mobile IP* de l'IETF. La plupart des protocoles proposés permettent par ailleurs aux terminaux d'entrer en état de veille, le réseau conservant une trace de leur position approximative ; ce procédé est appelé localisation ou *paging* et est l'objet d'étude du groupe de travail *Seamoby* de l'IETF.

La plupart des protocoles de micromobilité proposés à ce jour organisent les réseaux d'accès hiérarchiquement comme préconisé par [CP96]. Les stations de base offrant la connectivité sans fil aux terminaux sont reliées entre elles et à la passerelle faisant l'interface entre le réseau local et Internet par un ensemble de routeurs organisés en une structure arborescente. Une base de données de localisation est distribuée au sein du réseau d'accès afin de garder une trace de l'emplacement des mobiles. La mobilité intra-domaine ne se traduit alors qu'en une modification des tables de routage internes au domaine. Plusieurs études comparatives des principales propositions en matière de micromobilité ont été réalisées [CGK⁺02, CGC00] et une plate-forme d'évaluation a été déployée à l'université de Columbia (New York), permettant une comparaison des performances de ces divers protocoles.

Les auteurs de *Hierarchical Mobile IP* (HMIP) [GJP03] proposent d'établir une architecture arborescente de routeurs constituant un réseau de tunnels interne au domaine. La passerelle décapsule les paquets à destination d'un mobile visiteur puis les réencapsule avant de les transmettre dans le domaine local par le biais du tunnel approprié. Les terminaux mobiles se comportent comme des terminaux Mobile IP classiques et les notifications de changement de point d'attache destinées à l'agent mère sont interceptées par les routeurs du domaine visité et utilisées afin de mettre à jour le réseau de tunnels. Cette approche ne nécessite aucun changement dans le comportement des mobiles mais le processus d'encapsulations et de décapsulations successives peut être coûteux en termes de temps de calcul. C'est pourquoi l'architecture se limite très souvent à un seul niveau de hiérarchie.

Handoff-Aware Wireless Access Internet Infrastructure (Hawaii) [RLPT⁺02, RLPTV99] alloue à chaque mobile visiteur une adresse IP propre qu'il conservera durant son séjour dans le domaine. Le mobile fait alors partie du réseau local comme n'importe quel autre nœud. Les routeurs, organisés encore une fois en arborescence, utilisent un routage IP usuel et les changements de point d'attache, explicitement déclenchés par le mobile, ne provoquent qu'une mise à jour des tables de routage internes au domaine. Plusieurs politiques sont envisageables pour la redirection des paquets durant un *handoff*. Une limitation de cette architecture réside dans la capacité des routeurs actuels ne proposant qu'un faible nombre de règles permettant de définir des routes vers des hôtes particuliers plutôt que vers des sous-réseaux. Par ailleurs, le mécanisme de *handoff* proposé, ne mettant à jour les tables de routage que dans une partie du réseau local, peut conduire à l'établissement de routes sous-optimales au sein du réseau visité.

Cellular IP [Val99], fournit, quant-à-lui, une architecture se voulant la plus légère possible pour la gestion de la micromobilité, en particulier en minimisant le nombre de paquets de contrôle nécessaires. Le réseau hôte est, encore une fois, organisé hiérarchiquement en aval de la passerelle d'accès à Internet jouant le rôle d'agent étranger pour Mobile IP, comme représenté en figure 4.2. Les mobiles conservent leur adresse d'origine au sein du réseau et le routage à l'intérieur du domaine visité n'est pas un routage IP usuel mais un routage hiérarchique. Chaque nœud d'infrastructure conserve l'adresse d'un voisin lui permettant de joindre la passerelle (lien montant) et les adresses de plusieurs voisins lui permettant de joindre les mobiles (liens descendants). La mise à jour des tables de routage internes au domaine est réalisée en interceptant les paquets de données émis par les mobiles et ne nécessite aucun paquet de contrôle particulier. Les mobiles émettant peu de trafic peuvent émettre régulièrement des paquets de données vides (ces paquets sont alors appelés *Route Update*) destinés uniquement à la mise à jour de ces tables. La maintenance des routes montantes à destination de la passerelle est réalisée de la même manière, la passerelle inondant régulièrement le réseau local de messages particuliers appelés *Gateway*

advertisement. Enfin, les stations de base envoient des messages **BS advertisement** réguliers afin de signaler leur présence et les nœuds mobiles choisissent dynamiquement leur point d'ancrage sur la base de ces informations. Cellular IP propose deux techniques permettant à un mobile de changer de point d'attache. La première (appelée *semi-soft handoff*) consiste à transmettre les paquets de données par l'intermédiaire de l'ancienne et de la nouvelle station de base durant un certain intervalle de temps. La seconde approche (appelée *hard handoff*) cherche à minimiser le trafic dans le réseau en supprimant les paquets restant à transmettre dans l'ancienne station de base d'attache du mobile lorsque celui-ci déclare son nouveau point d'ancrage. Les routes inactives sont supprimées lors de l'expiration d'une temporisation.

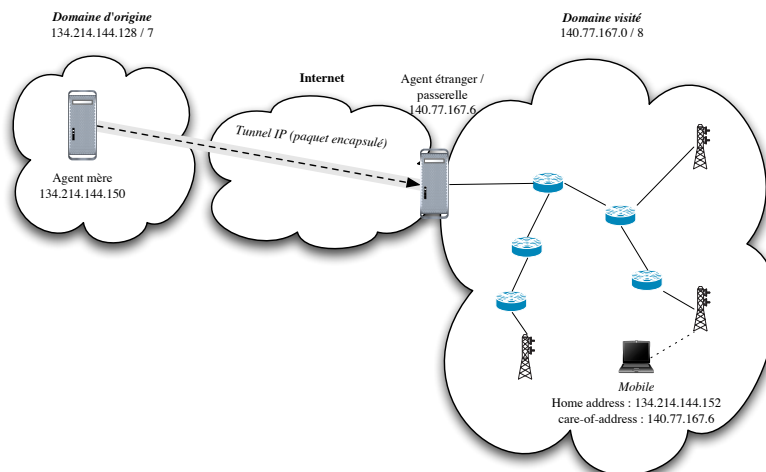


FIG. 4.2 – Architecture typique d'un réseau Cellular IP

4.2.3 Architecture considérée

Déployer un réseau d'accès filaire demeure cependant coûteux et l'utilisation d'un réseau d'accès sans fil peut présenter de nombreux avantages dans une telle architecture. La mise en place du réseau, ainsi que son extension, sont aisés et rapides, quelle que soit la configuration géographique de la zone à couvrir. L'utilisation d'un réseau d'accès sans fil s'accompagne cependant de l'ensemble des problèmes usuels de la transmission radio. Les ressources doivent *a priori* être partagées entre les trafics des mobiles et les trafics internes au réseau d'infrastructure, bien qu'il soit possible d'imaginer utiliser conjointement plusieurs canaux non recouvrants. Le transit des paquets dans l'infrastructure ne bénéficie plus ni de la fiabilité de transmission ni des performances des réseaux filaires. Toutefois, la souplesse d'une architecture telle que celle qui est représentée en figure 4.3, dans laquelle un réseau d'accès sans fil offre des services et une connectivité à Internet à un réseau *ad hoc* sous-jacent, reste très attractive. L'abondance de travaux publiés dans ce domaine [Typ01, WP00, CF03] le prouve. Dans l'architecture que nous considérerons par la suite, chaque nœud d'infrastructure joue à la fois le rôle de routeur d'accès et de station de base pour les terminaux mobiles.

Les différences introduites par un réseau d'infrastructure sans fil auront un impact sur les performances, et éventuellement sur le bon fonctionnement des protocoles de micromobilité. De même l'interaction du réseau d'infrastructure utilisant un routage défini par le protocole de micromobilité et du réseau *ad hoc* utilisant un routage dédié nécessite une définition de l'interface entre les deux types de routage. Plusieurs stratégies sont envisageables.

Il est possible d'utiliser un routage *ad hoc* dans l'ensemble du réseau, les stations de base de l'infrastructure étant considérées comme des nœuds *ad hoc*. Les travaux présentés dans [HKWHC01, BMAA02] utilisent cette stratégie représentée en figure 4.4. Ce type d'architecture est équivalent à un réseau *ad hoc* possédant une passerelle vers Internet et ne tirant aucun parti du caractère statique des nœuds d'infrastructure. Malgré sa simplicité, ce type d'approche souffre de problèmes relatifs au passage à l'échelle et ne permet pas une gestion rapide de la micromobilité.

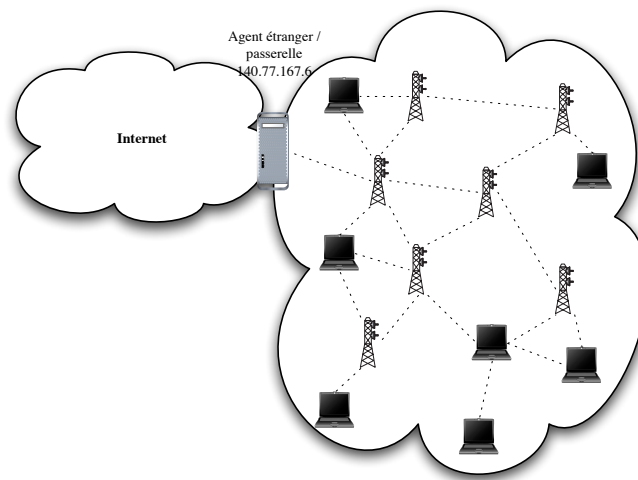


FIG. 4.3 – Architecture considérée : un réseau d’infrastructure sans fil sert un réseau *ad hoc*

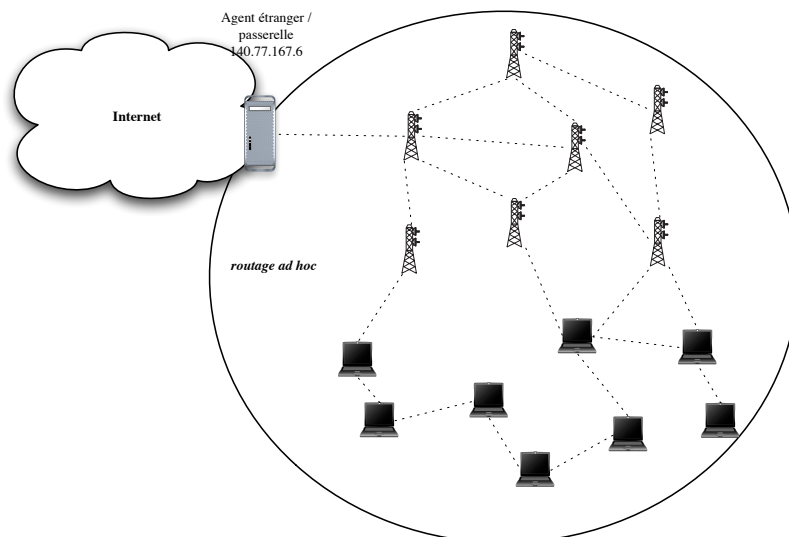


FIG. 4.4 – Stratégie de routage uni-polaire

Une seconde stratégie, représentée en figure 4.5 consisterait à distinguer les nœuds d'infrastructure et les nœuds *ad hoc*. Les stations de base seraient alors chargées d'effectuer la traduction du routage d'infrastructure en routage *ad hoc*. Ce type de stratégie est utilisé dans [Typ01, WP00].

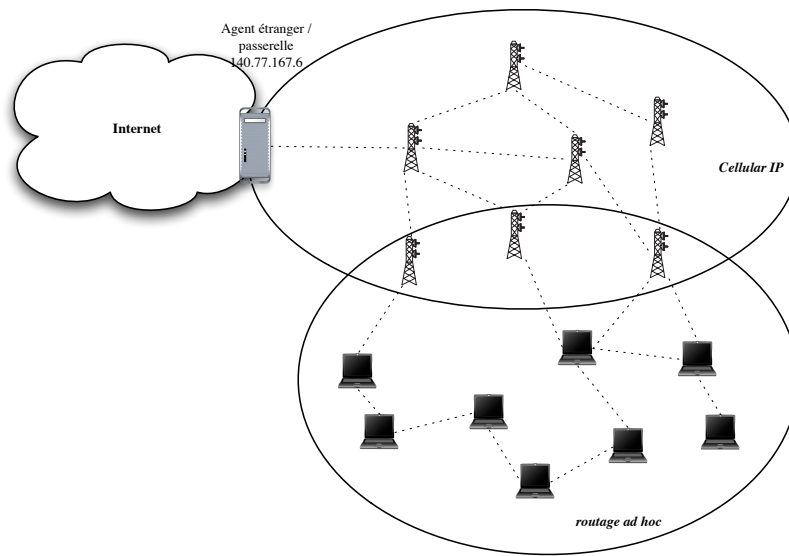


FIG. 4.5 – Stratégie de routage bipolaire

Enfin, la proposition d'architecture présentée dans [CF03] permet de partitionner le réseau hybride en sous-réseaux logiques, aboutissant à la situation représentée en figure 4.6. Un sous-réseau sera associé aux nœuds d'infrastructure et le réseau *ad hoc* peut se subdiviser en plusieurs réseaux logiques. Les stations de base participent à la fois au routage *ad hoc* et au routage d'infrastructure.

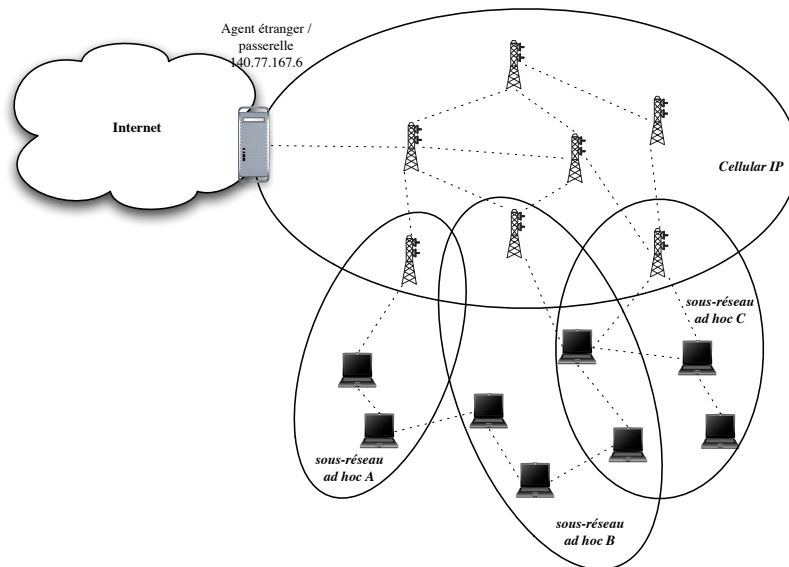


FIG. 4.6 – Stratégie de routage multipolaire

Dans ce type de réseau, le routage *ad hoc* impose que tous les mobiles communiquent en utilisant le même canal. En conséquence, le passage d'une station de base à une autre ne sera pas physique, dans le sens où les mobiles n'auront pas à changer leur canal d'émission mais plutôt logique dans le sens où un mobile sélectionnera la station de base la plus proche de lui afin de relayer ses paquets efficacement.

dans le réseau d'infrastructure. Les *handoffs* seront donc forcément de type *semi-soft* puisque le mobile sera toujours capable de recevoir des paquets de son ancien point d'attache alors qu'il est associé à une nouvelle station de base, pour peu qu'il se trouve toujours à portée de celle-ci.

4.2.4 Interactions entre le routage d'infrastructure et le routage *ad hoc*

Dans l'architecture que nous avons choisie, les mobiles communiquent en utilisant un protocole de routage *ad hoc* et les stations de base communiquent entre elles au moyen du protocole Cellular IP et avec les mobiles au moyen du protocole de routage *ad hoc*. Il est nécessaire de concevoir une interface aussi légère que possible afin de ne pas surcharger ces routeurs par lesquels passera l'essentiel du trafic réseau.

La plupart des protocoles de routage *ad hoc* font usage de paquets de contrôle destinés à découvrir le voisinage (paquets *Hello* présents dans OLSR et AODV par exemple). L'émission régulière de ces paquets est redondante avec l'émission de certains paquets par les nœuds mettant en œuvre Cellular IP. Les paquets **BS advertisement** émis périodiquement par les stations de base ainsi que les paquets **Route update** émis régulièrement par les mobiles ne transportant aucune information particulière en dehors de l'adresse de leur émetteur et de la station de base qu'ils ont choisie, ils peuvent avantageusement être remplacés par les paquets de contrôle du protocole de routage *ad hoc*. Il suffit d'intégrer aux paquets de contrôle *ad hoc* émis par les stations de base l'information qu'ils proviennent d'un routeur d'infrastructure. Quant aux paquets émis par les terminaux, il suffit pour les stations de base de les transmettre à travers l'infrastructure en modifiant simplement leur type afin d'éviter une mauvaise interprétation de ces paquets.

Cette fusion de paquets de contrôle présente un double avantage. Dans un premier temps, il s'agit de réduire le volume de trafic de contrôle transitant dans le réseau. Une telle modification permet par ailleurs d'accorder le protocole de routage *ad hoc* et le protocole de micromobilité afin qu'ils présentent la même réactivité face aux changements de topologie.

4.3 Stratégies de transmission des notifications de mobilité

L'utilisation d'un médium sans fil dans le réseau d'infrastructure présente un certain nombre d'inconvénients, principalement en terme de performances. La latence d'un lien Ethernet avoisine les $0,2\text{ ms}$ lorsque la latence d'un lien IEEE 802.11 approche 1 ms . De plus, la nature diffusante du médium radio peut rendre caducs certains choix réalisés lors de la conception de Cellular IP pour des réseaux d'infrastructure filaires.

Tout d'abord, la topologie du réseau d'infrastructure ne sera probablement pas arborescente, à moins de faire usage d'antennes directionnelles ou de disposer les stations de base de telle manière à empêcher les cycles dans la topologie. Compte tenu du caractère dynamique du routage d'infrastructure, les routes étant mises à jour au fur et à mesure que les paquets transitent à travers ce réseau, il est probable que la structure du réseau d'accès soit régulièrement modifiée. En conséquence, il peut être profitable de ne plus transmettre les paquets **Route update** en mode point à point (*unicast*) au sein du réseau d'accès, mais plutôt en mode diffusion locale (*broadcast*). Ce type d'approche peut, se révéler profitable puisque le mode diffusion du protocole IEEE 802.11 n'est pas acquitté, il présente de plus faibles délais de transmission. L'absence d'acquittements peut se révéler problématique sur un médium peu fiable, mais il est possible pour les routeurs non-destinataires de ce paquet de contrôle de tirer parti des informations qu'il contient, mettant à jour ses tables de routage plus rapidement qu'il ne l'aurait fait en attendant l'expiration d'une temporisation.

Une approche intermédiaire est par ailleurs envisageable. Compte tenu de la structure logique arborescente prévue par Cellular IP, chaque nœud d'infrastructure peut identifier précisément son père dans l'arbre de routage. Il est donc possible de concevoir, pour la transmission des paquets **Route Update**, un mode hybride entre le mode point à point et le mode diffusion. Ce mode de transmission, que nous appellerons diffusion acquittée (*Acknowledged Broadcast*), aurait le comportement d'une diffusion dans le sens où tous les nœuds à portée de transmission pourraient tirer parti des informations contenues dans le paquet de signalisation. Parallèlement, la hiérarchie de routage permet au père du nœud émetteur dans l'arbre, principal destinataire du message, d'émettre un acquittement en réponse. L'architecture hiérarchique résout dans ce cas le problème de la sélection du nœud chargé d'acquitter les messages diffusés par un autre nœud sans fil, sans risque de collision entre de multiples acquittements. Par ailleurs, l'absence de

mobilité des nœuds d'infrastructure élimine la confusion usuelle entre collision et sortie du destinataire de la zone de couverture. En pratique, ce type d'approche, peut être réalisé au moyen d'émissions en mode point à point alors que tous les nœuds d'infrastructure sont en mode écoute (*promiscuous*).

Le mode de transmission point à point représente la solution la plus fiable et la plus coûteuse en terme de ressources du réseau. Chaque trame est précédée par un échange RTS-CTS, optionnel mais utilisé ici afin d'en évaluer le coût, et doit être acquittée. Le mode diffusion n'est, à l'inverse, ni protégé, ni acquitté. La diffusion acquittée, en termes de performances, se comportera comme le mode point à point sans échange RTS-CTS préalable. Si l'on considère un canal à 2 Mbit/s , envoyer une trame **Route Update** de 36 octets en mode diffusion prend, en moyenne, $832\text{ }\mu\text{s}$. L'envoi de la même trame en mode diffusion acquittée durera $1146\text{ }\mu\text{s}$ et $1646\text{ }\mu\text{s}$ en mode point à point.

Nous avons étudié les trois modes de transmission des paquets **Route Update** possibles par le biais de simulations au moyen du simulateur NS-2¹ en version 2.27. Afin d'étudier les performances brutes de ces trois modes de transmission, nous avons choisi, dans un premier temps, de ne pas tirer parti des informations pouvant être obtenues par les nœuds d'infrastructure lors de l'utilisation des modes diffusion et diffusion acquittée.

La topologie considérée pour ces simulations et pour les suivantes est constituée de neuf nœuds d'infrastructure autour desquels s'articule un réseau *ad hoc* de deux à soixante-quatre nœuds. Les mobiles se déplacent selon le modèle de mobilité *Random Waypoint* [CBD02], c'est-à-dire que chaque mobile tire une destination aléatoire, s'y rend avec une vitesse aléatoire, reste sur place pendant un temps aléatoire et recommence le processus. La vitesse maximale des mobiles a été fixée à 50 m/s . Le trafic dans le réseau est constitué de un à trente-deux flux à un débit constant de cinq paquets de 500 octets par seconde, c'est-à-dire 20 kbit/s . Toutes les données sont émises à un débit de 2 Mbit/s afin de ne pas introduire de différence de débit entre les différents modes de transmission. Il faut cependant noter que les modes point à point et diffusion acquittée pourraient bénéficier de taux de transfert plus élevés. Les résultats présentés ici représentent les résultats moyens sur 50 simulations.

Le protocole étudié est une version modifiée de Cellular IP afin d'interagir avec un réseau *ad hoc* sous-jacent et utilisant un réseau d'infrastructure sans fil. Les mobiles ainsi que les stations de base participent au routage *ad hoc*. Tous émettent donc des paquets de contrôle *ad hoc* contenant la liste de leurs voisins ainsi que l'identité de la station de base qu'ils ont choisie pour les mobiles. Les stations de base remplacent cette dernière information par l'indication du fait qu'elles sont des stations de base. Ces messages sont émis à une période de 200 ms . Ces paquets jouent à la fois le rôle de paquets *Hello* pour le protocole de routage *ad hoc* et de paquets **Route Update** ou **BS advertisement** pour Cellular IP. Les notifications de mobilité sont implicitement initiées par les mobiles. Lorsqu'une station de base reçoit un paquet de routage *ad hoc* de la part d'un mobile, elle convertit ce dernier en paquet **Route Update** et le transmet à destination de la passerelle par le biais de l'infrastructure. Une route dans l'infrastructure expire au bout de 500 ms , soit un temps équivalent à la perte de deux paquets **Route Update** et demi. Si aucune route vers le destinataire d'un paquet de données n'est présente dans l'infrastructure, ce paquet est transmis à la passerelle qui le re-dirige ou le détruit si elle non plus ne possède aucune information sur la localisation du destinataire. Les mobiles transmettent toujours leurs paquets de données à la station de base à laquelle ils sont attachés.

Ces paquets de contrôle peuvent représenter un volume de données conséquent. Cependant, ils transportent des informations importantes pour la gestion du réseau et ne devraient ni être perdus, ni être retardés. La perte de tels paquets entraînera une inconsistance des tables de routage ne comportant pas suffisamment d'informations ou voyant des informations valides expirer. Un trop grand délai dans la transmission de ces paquets, résultera, à l'inverse, dans la persistance d'informations périmées. Il n'est donc pas aisé de déterminer intuitivement s'il est préférable d'utiliser une transmission protégée introduisant un délai dans l'accès au médium ou une transmission rapide et fragile.

La figure 4.7 présente le nombre de paquets de données perdus à cause d'une absence de route vers la destination. Ce type de situation survient lorsque les entrées dans les tables de routage correspondant au destinataire ont expiré et que les nouvelles informations ne sont pas parvenues à la passerelle. Les paquets de données concernés sont transmis à la passerelle qui, ne possédant aucune information sur la localisation du mobile, détruit le paquet. On peut remarquer que la transmission des paquets **Route Update** en mode point à point conduit au plus grand volume de pertes lorsque le médium est saturé. Le délai introduit par la protection des trames provoque un retard dans la propagation de ces trames et une expiration des routes dans les tables.

¹<http://www.isi.edu/nsnam/ns/>

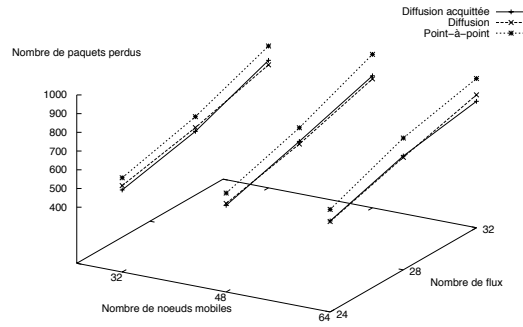


FIG. 4.7 – Pertes de paquets de données dues à une absence de route

La figure 4.8 présente le nombre de paquets de données perdus après avoir dépassé le nombre de retransmissions au niveau MAC autorisé. En général, ce type de pertes survient lorsqu'un mobile n'est plus accessible à partir de cette station de base mais que la route correspondante dans les tables de routage fait encore référence son ancienne station de base d'attache. Dans ce type de situation, la transmission des messages **Route Update** en mode diffusion conduit à la plus grande perte de paquets de données en raison de la faible fiabilité de ce mode de transmission.

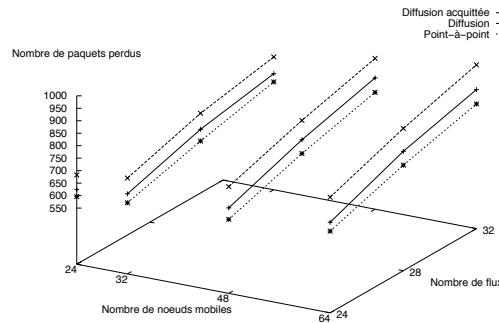


FIG. 4.8 – Pertes de paquets de données dues à une disparition du destinataire

Chacune des deux stratégies présente, comme nous venons de le voir, des avantages et des inconvénients. Il est toujours difficile de prévoir laquelle conduira au meilleur fonctionnement du réseau. La diffusion acquittée représentant une approche médiane, elle pourrait aussi bien conduire au meilleur résultat qu'au pire. La figure 4.9 représente le total de paquets de données correctement transmis pour chacune des situations simulées. Dès que le médium radio est saturé, le mode diffusion représente la solution la plus performante et son taux de réussite dépasse celui du mode point à point de 25 % dans le meilleur cas. La diffusion acquittée reste une approche médiane.

Le nombre de paquets transmis avec succès est donc plus sensible à la charge du médium qu'à la pertinence des tables de routage. Tant que la capacité du médium n'est pas saturée, les performances des trois approches sont équivalentes. Puis, le surcoût introduit par le mode de transmission des paquets de contrôle provoque une dégradation des performances du protocole. Enfin, lorsque le médium devient saturé quel que soit le mode de transmission des paquets de contrôle, les performances des trois approches redeviennent équivalentes.

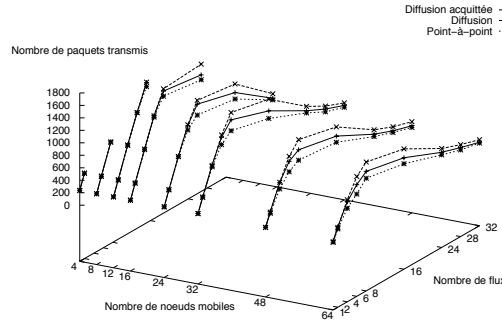


FIG. 4.9 – Total de paquets de données correctement transmis

4.4 Optimisations diverses

Parmi toutes les causes possibles de pertes de paquets, les collisions dues à la charge du médium radio ou les pertes dues à la mobilité rendent compte d'un phénomène physique et sont difficilement réductibles. En revanche, certaines pertes sont dues à des problèmes protocolaires et peuvent être résolues aisément, par exemple, les pertes de paquets dues aux échecs du protocole de résolution d'adresse (ARP – *Address Resolution Protocol*) ou à la longueur des files d'attente internes aux mobiles.

4.4.1 Remplissage gratuit de cache ARP

Lorsqu'un mobile entre dans le réseau, tout paquet destiné à ce mobile sera détruit par les routeurs tant que la correspondance entre son adresse IP et son adresse MAC n'aura pas été faite. Cette correspondance est déterminée au moyen d'un échange requête-réponse spécifique déclenché par la tentative d'émission d'un paquet en mode point à point à destination du nouvel arrivant. La requête est envoyée en mode diffusion et lorsque la résolution échoue, le paquet de données est détruit. Dans le contexte sans fil et mobile que nous étudions, deux causes peuvent être à l'origine d'un échec dans la résolution. Tout d'abord les requêtes peuvent être victimes de collisions, comme tout paquet. D'autre part, la mobilité des terminaux peut provoquer un tel échec si, par exemple, un mobile sort de la zone de couverture de la station de base à laquelle il s'était attaché avant que celle-ci n'ait effectué cette correspondance. Lorsque la station de base tentera de joindre le mobile, celui-ci ne pourra répondre à la requête.

Le protocole ARP pourrait tirer parti des émissions régulières par les mobiles de messages de routage *ad hoc*. Lors de la réception d'un tel paquet, une station de base pourrait extraire à la fois l'adresse IP et l'adresse MAC de l'émetteur et compléter sa table de correspondance de cette façon. Cette optimisation, que nous appellerons *remplissage gratuit de cache ARP*, ne peut toutefois se baser que sur certains paquets de contrôle n'ayant pas été retransmis et dont l'adresse IP de l'émetteur correspond bien à l'adresse MAC du transmetteur. La figure 4.10 compare une comparaison du nombre de paquets perdus à cause de l'absence dans les tables ARP de correspondance entre adresse IP et adresse MAC. Les résultats présentés correspondent à une simulation sur un réseau comportant 64 nœuds mobiles, l'infrastructure fonctionnant en mode diffusion acquittée avec et sans utilisation du mécanisme de remplissage gratuit de cache ARP. Les résultats obtenus sur d'autres topologies sont similaires. L'utilisation de ce mécanisme permet de s'affranchir en totalité des pertes dues à l'échec de la résolution ARP.

La figure 4.11 compare le total de paquets de données perdus à cause d'un dépassement du nombre maximal de retransmissions de paquets avec et sans cette optimisation. L'utilisation du remplissage gratuit des caches ARP conduit à une augmentation de ces pertes. En effet, le remplissage gratuit des caches ARP ne permet pas d'éviter les pertes de paquets dues à la mobilité des terminaux. Auparavant, les pertes de paquets destinés à un mobile absent auquel on n'avait pas encore transmis de paquet étaient imputées à ARP, maintenant ces pertes sont considérées comme des pertes dues à la mobilité.

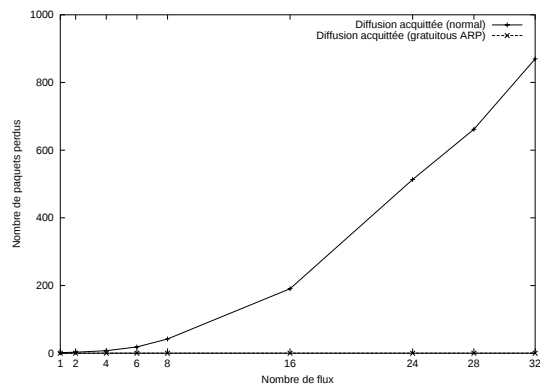


FIG. 4.10 – Pertes de paquets dues à une absence de correspondance ARP

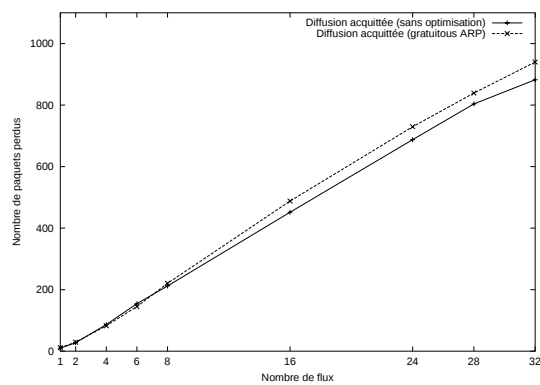


FIG. 4.11 – Pertes de paquets de données dues à des retransmissions excessives en utilisant le mécanisme de remplissage gratuit de cache ARP

4.4.2 Longueur des files d'attente

Les résultats présentés en section 4.3 indiquaient que le délai de transmission des paquets de contrôle avait un impact sur la performance du réseau. Ce problème ne peut simplement être résolu en accroissant le délai d'expiration des routes dans l'infrastructure. En effet, une telle opération provoquerait un accroissement des pertes dues à la persistance des routes invalides. Les paquets de contrôle transportant des informations périmées ne présentant aucun intérêt pour la gestion du réseau, ils ne devraient pas être transmis. Diminuer la taille de la file d'interface située entre la couche routage et la couche accès au médium dans la pile protocolaire des mobiles pourrait avoir un impact sur la validité des informations transmises.

La figure 4.12 compare le volume de paquets perdus dans cette file d'attente pour des longueurs de file de 10 paquets et de 100 paquets. On peut tout d'abord remarquer que le nombre de paquets détruits par ces files est très important. Il représente jusqu'à 60 % du volume total de pertes dans le pire cas. Augmenter la taille de ces files peut conduire à un accroissement du taux de transmission des paquets, mais aura un impact négatif sur les délais. La figure 4.13 compare le nombre de paquets perdus à cause de l'expiration des routes valides dans l'infrastructure et la figure 4.14 compare les pertes de paquets dues à la présence d'informations périmées dans les tables de routage. Au vu de ces résultats, si accroître la taille des files d'attente conduit à une réduction du nombre de paquets de données détruits, le délai introduit dans la transmission des paquets de contrôle provoque un accroissement des pertes dues à des erreurs de routage, principalement lorsque la charge du réseau augmente mais en nombre beaucoup plus faible que les pertes précédemment dues à une erreur ARP.

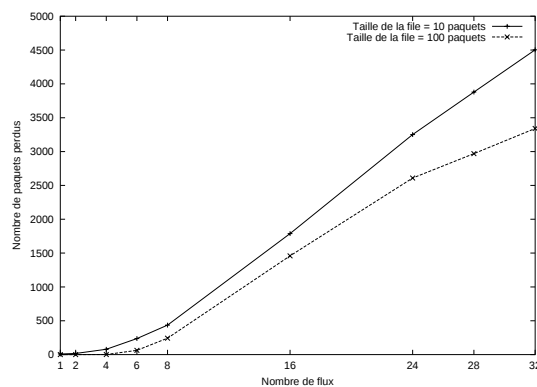


FIG. 4.12 – Volume de paquets détruits dans les files d'interface

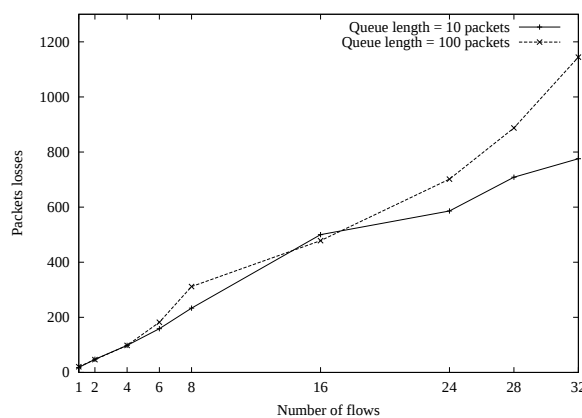


FIG. 4.13 – Volume de paquets perdus à cause de l'absence de route vers le destinataire

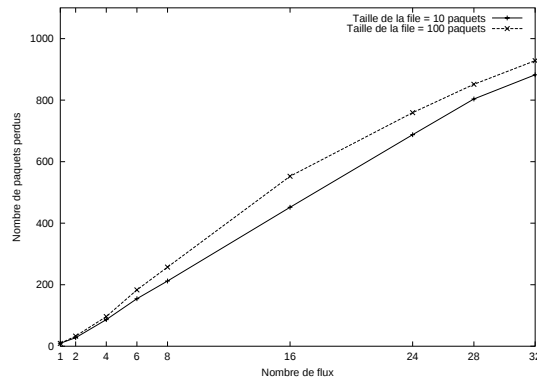


FIG. 4.14 – Volume de paquets perdus à cause de l’existence de route périmée vers le destinataire

La figure 4.15 présente le nombre de paquets correctement transmis pour ces deux longueurs de file dans le cas d’un réseau comportant 64 mobiles. Si utiliser une taille de file importante est profitable lorsque le médium est peu chargé, le délai de transmission des paquets de contrôle devient un paramètre primordial lorsque la charge du réseau augmente.

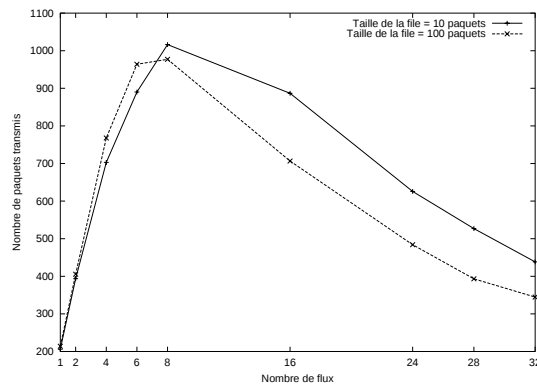


FIG. 4.15 – Volume de paquets transmis avec succès

4.5 Optimisation de la transmission de notifications de mobilité

Les résultats précédents indiquent que les pertes de paquets de données sont principalement dues à la charge du médium, et dans une moindre mesure à l’inconsistance des tables de routage. Dans un réseau dense, le volume de trafic de contrôle deviendra rapidement important. En effet, chaque mobile transmet périodiquement des paquets de contrôle qui seront relayés par les nœuds d’infrastructure. Le but des optimisations présentées dans ce paragraphe est de réduire le volume de trafic de contrôle échangé tout en essayant de conserver une réactivité identique face aux changements de topologie.

4.5.1 *Differential Route Update*

Lorsqu’un mobile vient de changer de point d’attache, il est important de s’assurer de la bonne notification du réseau d’infrastructure et donc de transmettre fréquemment des notifications de mobilité afin de garantir leur bonne réception. En revanche, lorsqu’un nœud reste attaché à la même station

de base, il n'a pas besoin de l'indiquer régulièrement. Pour réduire le volume de trafic de contrôle, il est possible d'accroître l'intervalle entre deux transmissions successives de notifications de mobilité. Cependant, il ne s'agit pas ici de réduire le nombre de paquets destinés au routage *ad hoc* car la mobilité d'un nœud proche d'une station de base n'est pas *a priori* corrélée à la mobilité des autres nœuds du réseau *ad hoc*. Cette optimisation réduira donc la fréquence de retransmission des paquets **Route Update** par les nœuds d'infrastructure et sera nommée *Differential Route Update* par la suite.

Les mobiles étant les seuls à décider de la station de base à laquelle ils s'associent, ils seront à l'origine de la décision de transmettre ou non les notifications de mobilité au sein de l'infrastructure. Le mobile indiquera explicitement dans chaque paquet de routage *ad hoc* si celui-ci doit être retransmis par l'infrastructure ou non. Le tableau 4.1 présente les intervalles séparant l'émission de deux paquets devant être convertis en paquets **Route Update** ainsi que le délai d'expiration des routes correspondantes en fonction du temps depuis lequel le mobile est associé à la même station de base. Lorsque le mobile change de point d'attache, cette valeur est réinitialisée à la fréquence la plus haute.

k ^e Route Update	1	2	3	4	5	6	7	8	10	11	12
Intervalle	0, 2 s	0, 3 s	0, 4 s	0, 6 s	0, 8 s	1, 0 s	1, 2 s	1, 4 s	1, 6 s	1, 8 s	2, 0 s
Expiration	0, 5 s	0, 75 s	1, 0 s	1, 5 s	2, 0 s	2, 5 s	3, 0 s	3, 5 s	4, 0 s	4, 5 s	5, 0 s

TAB. 4.1 – Intervalle de temps séparant l'émission de deux **Route Update** successifs.

4.5.2 *Nack Route*

Toutefois, l'utilisation de *Differential Route Update* peut provoquer des inconsistances dans les tables de routage. En effet, émettre un nombre réduit de notifications de mobilité signifie que le délai d'expiration des routes dans le réseau d'infrastructure doit être augmenté. Dans cette situation, une station de base détectant le départ d'un nœud de sa zone d'influence détruira l'entrée correspondante dans sa table de routage. Si le réseau d'infrastructure n'en est pas notifié, il est aisé de créer une boucle de routage temporaire, la station de base ne possédant aucune route vers le mobile transmettant les paquets vers son père qui possède une route passant par cette dernière pour joindre le mobile. Aussi, lorsqu'une station de base détectera le départ d'un mobile, elle devra notifier explicitement tout le réseau d'infrastructure au moyen d'un paquet **Route Delete** provoquant la suppression de la route dans le réseau d'infrastructure à mesure qu'il est retransmis. Cette optimisation représentant des émissions de paquets supplémentaires, sera étudiée séparément, sous le nom *Nack Route*.

4.5.3 *Nack Only*

En poussant le raisonnement à l'extrême, il est possible d'adopter une stratégie optimiste consistant à ne transmettre qu'une seule fois un message **Route Update** lors d'un *handoff*, de supposer que ce message sera bien reçu et provoquera la création d'une nouvelle route ainsi que la destruction explicite de l'ancienne. Les routes n'auraient alors plus de délai d'expiration. Si cette stratégie, que nous nommerons *Nack Only*, semble très fragile, car la perte d'un paquet **Route Update** aura des conséquences importantes, elle permettrait de réduire de façon drastique le volume de trafic de contrôle.

4.5.4 Résultats de simulation

Toutes ces optimisations présentent des avantages ainsi que des inconvénients évidents. L'évaluation par simulation de leurs performances a été réalisée sur des topologies variées et nous ne présenterons ici que les résultats obtenus sur des réseaux comportant 64 nœuds mobiles et transmettant les notifications de mobilité en mode diffusion acquittée, les résultats obtenus sur des topologies différentes ou avec des modes de transmission différents étant similaires.

L'objectif de chacune de ces optimisations est de réduire la proportion de bande passante dédiée au trafic de contrôle en limitant la transmission d'informations redondantes. Il s'agit de déterminer le niveau de redondance exact à utiliser afin d'assurer une transmission correcte et économique des informations. La figure 4.16 compare le nombre de paquets **Route Update** émis par les nœuds d'infrastructure pour

chacune des quatre stratégies présentées. L'activation des optimisations successives permet, comme prévu, de réduire grandement le volume du trafic de contrôle émis.

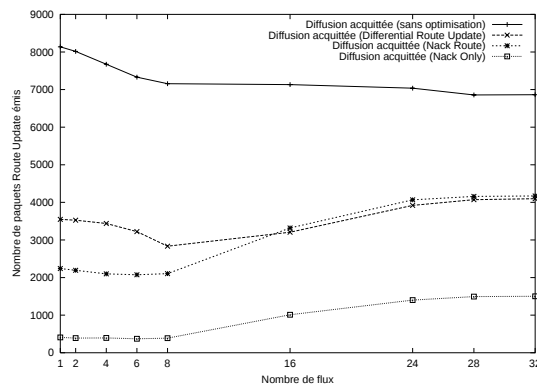


FIG. 4.16 – Nombre de paquets **Route Update** émis

Afin d'étudier l'impact de ces optimisations sur les performances du réseau, intéressons-nous maintenant à leur influence sur la validité des tables de routage. La figure 4.17 présente, pour chacune des quatre stratégies étudiées, le nombre de paquets perdus à cause de la disparition de toute route à destination du mobile. Ce type de perte survient lorsque les routes à destination d'un mobile ont expiré alors que l'information sur sa localisation ne s'est pas encore propagée à travers le réseau d'infrastructure ou a été perdue. Lorsque le canal radio est peu chargé, ces optimisations semblent avoir un impact négatif sur la performance du protocole de routage. Cependant, dès que la charge augmente, l'optimisation *Nack Only* se révèle très efficace. Le gain en terme de bande passante provoqué par cette optimisation repousse l'apparition du point de saturation. Cependant, l'impact de la perte d'un paquet de contrôle devient beaucoup plus dommageable, aussi lorsque le point de saturation du réseau est à nouveau atteint, l'intérêt de cette optimisation décroît.

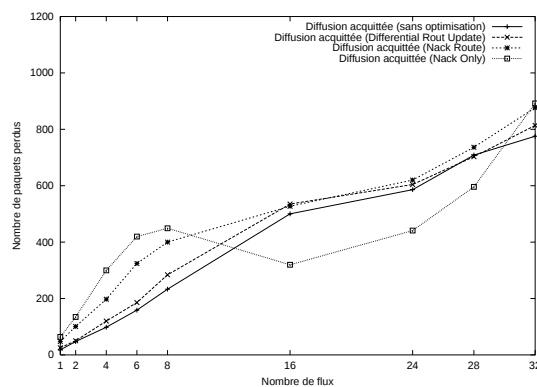


FIG. 4.17 – Nombre de paquets de données perdus à cause de l'absence de route vers la destination

Le figure 4.18 représente le nombre de paquets de données perdus à cause de la présence de routes périmées dans les tables de routage. Chacune des trois optimisations présentées a un impact négatif sur ce total et ce, quelle que soit la charge du médium. En effet, chacune de ces optimisations accroissant le délai d'expiration des routes, les routes non explicitement supprimées persisteront plus longtemps.

Les optimisations présentées ici soulèvent les mêmes questions que la comparaison des différentes stratégies pour la transmission des messages **Route Update**. En effet, réduire le coût du trafic de contrôle

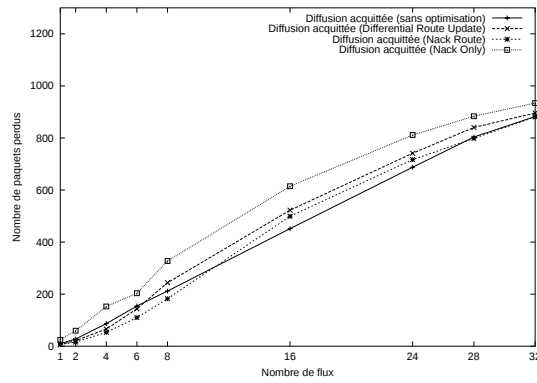


FIG. 4.18 – Nombre de paquets de données perdus à cause de la présence de routes incorrectes dans les tables

se solde par une réactivité réduite face à la mobilité des terminaux et une certaine inertie du mécanisme de routage. Au contraire, accroître ce volume de trafic provoque une occupation du médium qui se traduira par un accroissement des délais de propagation des informations. La figure 4.19 compare le nombre de paquets de données correctement transmis avec et sans optimisations. La stratégie *Differential Route Update* représente toujours un gain de performance, quelle que soit la charge du médium. L'optimisation *Nack Only* se révèle utile lorsque le médium est surchargé mais représente une perte de performance lorsque la nécessité d'économiser de la bande passante ne se fait pas sentir.

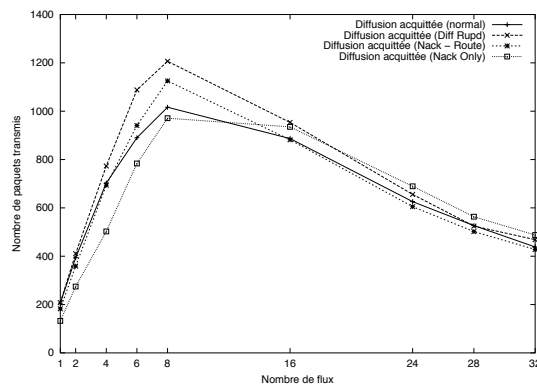


FIG. 4.19 – Nombre de paquets de données correctement transmis avec et sans optimisations

En guise de conclusion, la figure 4.20 compare le nombre de paquets correctement transmis sans aucune optimisation au nombre de paquets correctement transmis en activant l'optimisation *Differential Route Update* et le remplissage gratuit de cache ARP. Le gain de performance résultant de l'application des techniques présentées ici atteint 40 % dans le meilleur cas.

4.6 Conclusion

Dans ce chapitre, nous avons étudié l'impact de différentes optimisations, agissant à différents niveaux, sur la transmission des notifications de mobilité dans des réseaux hybrides sans fil. Dans de tels réseaux, les causes de pertes de paquets sont nombreuses. La congestion joue un rôle important, comme dans tout type de réseau radio et le seul moyen d'action sur ce phénomène est de limiter les émissions des

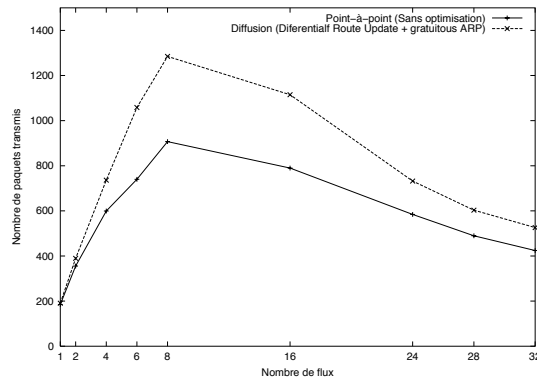


FIG. 4.20 – Nombre de paquets de données correctement transmis

différents nœuds du réseau. Toutefois, des problèmes liés à la validité des tables de routage rentrent ici en ligne de compte et sont directement liés à la qualité de la transmission des notifications de mobilité. Les informations contenues dans ces tables peuvent en effet être périmées, un mobile ayant changé de station de base sans que l'information ne se soit propagée dans l'infrastructure. À l'inverse, un mobile peut être considéré comme injoignable, toute route vers lui ayant expiré dans l'infrastructure.

Différentes stratégies ont été étudiées dans ce chapitre. Toutes les études menées ici montrent qu'il est primordial de ne pas surcharger le réseau par un volume de signalisation trop important. Accroître la fréquence de transmission des notifications de mobilité ne fait qu'aggraver les situations de congestion, ces mêmes notifications de mobilité étant retardées dans les files d'attente et transportant des informations périmées. Il est en général préférable d'adapter cette fréquence de transmission aux changements survenant dans la topologie du réseau.

Le même type de conclusion peut s'appliquer à BRuIT. La transmission de paquets transportant des informations de voisinage ne doit pas nécessairement être périodique. Déclencher l'émission d'un tel paquet lorsqu'un changement dans le voisinage est détecté permettrait de diminuer le volume de signalisation tout en conservant une réactivité importante.

Le travail présenté dans ce chapitre est loin d'être terminé. Dans un premier temps, il sera nécessaire d'étudier l'impact de l'utilisation des informations transmises en mode diffusion acquittée sur la maintenance des tables de routage. Par la suite, l'étude de l'intégration d'un tel mécanisme dans un protocole purement *ad hoc* tel que BRuIT devra être réalisée afin d'évaluer le gain en fiabilité induit par ce mode de fonctionnement. En effet, l'utilisation d'une diffusion acquittée peut se révéler profitable afin de s'assurer de la bonne transmission des informations qui auparavant étaient simplement diffusées. Toutefois, l'utilisation d'un tel mécanisme dans un cadre moins structuré tel qu'un réseau *ad hoc* nécessite l'élection dynamique pour chacun des nœuds d'un voisin chargé d'émettre les acquittements, ce qui représente un surcoût supplémentaire en terme de trafic de contrôle.

Conclusion et perspectives

Les réseaux *ad hoc* sont actuellement un sujet de recherche très actif même si les caractéristiques de ces réseaux tels que leur topologie ou leur utilisation ne sont pas clairement définies. L'utilisation d'un médium radio, fondamentalement différent du médium filaire usuel, ainsi que l'absence d'infrastructure fixe introduisent de nombreuses contraintes dont il faut tenir compte lors de la conception de protocoles. Les problématiques de qualité de service dans ces réseaux ont été le sujet de nombreuses publications. Toutefois, il est peu aisé de concevoir une solution générique répondant aux exigences de tous les types d'applications dans tout type de topologie. Dans cette thèse, nous avons choisi d'étudier la problématique de la réservation de bande passante dans les réseaux *ad hoc* basés sur le protocole IEEE 802.11. Ce travail peut par ailleurs aisément être adapté à tout protocole d'accès au médium ayant un fonctionnement similaire.

Le premier chapitre de ce document présente une modélisation d'un scénario mettant en jeu trois couples émetteurs-récepteurs et présentant un fort déséquilibre dans l'accès au médium. L'un des trois émetteurs ne peut, en effet, disposer que d'une proportion de la bande passante du canal avoisinant les 5 %. Ce problème de performances est en grande partie dû à l'absence de synchronisation des deux autres émetteurs lors de l'accès au médium. Ce phénomène est aggravé par l'utilisation d'une temporisation longue (*EIFS*) provoquée par la position relative des émetteurs voisins. Le comportement de ce scénario a été modélisé sous forme d'une chaîne de Markov en temps discret et validé par simulation. Les résultats sont par ailleurs proches des résultats d'expérimentation correspondants. Cette étude permet de mettre en lumière l'existence dans des réseaux *ad hoc* multi-sauts de configurations pathologiques présentant de réels problèmes de performances. Ce type de problèmes de performances peut parfois être résolu ou du moins détecté en effectuant un contrôle d'admission basé sur une estimation de la disponibilité des ressources.

Le deuxième chapitre présente BRuIT, un protocole de réservation de bande passante pour réseaux *ad hoc*. Ce protocole adopte un fonctionnement réactif, recherchant des routes admissibles au regard des critères de qualité de service par inondation. Le contrôle d'admission réalisé par les différents routeurs potentiels est basé sur la transmission régulière par chaque nœud d'informations concernant le volume de trafic émis et routé dans un voisinage d'une taille approchant la taille de la zone de détection de porteuse des nœuds utilisant le protocole IEEE 802.11. Le but de cette propagation d'informations est d'identifier au mieux l'ensemble des mobiles avec lesquels chaque nœud aura à concourir pour l'accès au médium afin d'évaluer la quantité de ressources disponibles. Les limites d'un tel mécanisme sont identifiées et une étude par simulation a été réalisée afin d'évaluer l'apport et le coût d'un tel mécanisme par rapport à un protocole de routage sans qualité de service. Les résultats obtenus montrent qu'effectuer un contrôle du débit de chaque nœud permet de limiter l'apparition de congestions, améliorant ainsi le fonctionnement du réseau. La mise en œuvre de ce mécanisme représente toutefois un surcoût en terme de volume de trafic de contrôle. La mise en place d'un mécanisme de dégradation des garanties est par ailleurs nécessaire afin de réagir aux estimations erronées pouvant survenir ainsi qu'aux variations des ressources disponibles dues à la mobilité des nœuds.

Un mécanisme de réservation de bande passante ne peut cependant offrir de garanties sans séparer ou contrôler le débit du trafic au mieux. Dans des réseaux *ad hoc*, il est difficile de séparer deux types de trafics car ce processus nécessite la propagation d'informations à travers tout le réseau et devient caduque lorsque plusieurs réseaux disjoints se rejoignent. Aussi, dans ce but, le troisième chapitre présente et étudie un algorithme distribué d'allocation de bande passante de façon équitable dans un réseau *ad hoc*. Cet algorithme se base sur des informations locales afin de permettre à chaque nœud de calculer itérativement le volume de bande passante qu'il est en mesure d'utiliser. L'étude des performances de cet algorithme sur des graphes réguliers et des graphes géométriques aléatoires montre qu'il s'agit d'un bon

compromis entre utilisation du réseau et équité dans l'accès au médium. L'algorithme initial est destiné à des réseaux statiques tels que des réseaux de capteurs. Une modification de cet algorithme permet de gérer la mobilité des nœuds. L'évaluation des performances de ce second algorithme montre que la perte de performances liées à la mobilité est minime. Cet algorithme est destiné à être utilisé en conjonction avec BRuIT afin de limiter le débit du trafic au mieux. Toutefois, il peut aussi être utilisé seul afin d'assurer une certaine équité dans l'accès au médium et palier ainsi les problèmes liés aux déséquilibres de la topologie de certains réseaux. Il est intéressant de noter que procédé utilisé ici permet de calculer une approximation de la solution d'un problème d'optimisation de façon distribuée. Généraliser cette étude à tout problème de maximisation ou de minimisation d'une fonction sous contrainte et déterminer la qualité de l'approximation ainsi obtenue est une perspective connexe importante de ce travail.

Enfin, le dernier chapitre étudie différentes optimisations relatives à la transmission des notifications de mobilité dans des réseaux hybrides possédant une infrastructure sans fil. Le but de cette étude était de déterminer l'impact et le coût relatifs à l'utilisation d'une diffusion acquittée. La structure hiérarchique fournie par ces réseaux hybrides permet de s'affranchir momentanément des problèmes d'élection de voisins privilégiés chargés d'acquitter les diffusions dans des réseaux totalement *ad hoc*. Il se dégage de cette étude que le volume de trafic de contrôle émis sur le médium a un impact important sur le fonctionnement du réseau, et ce même lorsque les informations transportées sont supposées améliorer le fonctionnement global du réseau.

La réservation de bande passante dans les réseaux *ad hoc* est une problématique difficile. De nombreuses contraintes liées aux aléas de la propagation radio, à la mobilité et à l'absence de modèle de topologie ou de scénario d'utilisation imposent le développement de solutions génériques. Au regard des résultats obtenus tout au long de cette thèse, il semble impossible, à un coût raisonnable, de fournir une solution assurant la conservation des garanties dans toute situation. Un tel niveau de garanties nécessiterait en effet une connaissance parfaite de la topologie afin de déterminer précisément les interactions entre mobiles et des différents trafics dans le réseau. L'utilisation de technologies de localisations telles qu'un GPS n'est pas d'un grand secours dans ce type de situation puisque l'atténuation sur un lien entre deux nœuds n'est pas uniquement fonction de leurs positions relatives. Il est cependant possible, sans parler de garanties, d'offrir des indications sur l'état des ressources dans le réseau aux différentes applications. Ces indications peuvent être utilisées afin d'évaluer la possibilité pour un flux d'être transmis correctement. Par ailleurs, mettre en place un tel système de réservations permet d'optimiser le routage dans ces réseaux en effectuant une répartition de la charge. Les mobiles appartenant aux zones congestionnées du réseau refuseront de transmettre de nouvelles requêtes et des routes distinctes des plus courts chemins seront explorées.

La combinaison des mécanismes présentés dans ce document constitue une première solution au problème de la réservation de bande passante dans les réseaux *ad hoc*. De nombreuses améliorations restent cependant à apporter à cette architecture. Premièrement, une évaluation plus précise de la bande passante disponible doit être mise en place afin de prendre en compte des situations comme le scénario des trois paires lorsque aucun mobile ne peut transmettre des informations d'un émetteur à l'autre. Une solution basée sur une communication entre les couches MAC et routage utilisée en conjonction de BRuIT permettrait d'améliorer l'évaluation effectuée. BRuIT effectue une évaluation *a priori* des ressources disponibles et des indications provenant des couches inférieures sur la disponibilité des ressources pourraient être utilisées afin d'effectuer une surveillance *a posteriori* et de détecter les évaluations erronées.

Un deuxième axe de recherche réside en une gestion plus fine de la mobilité des nœuds. Actuellement, BRuIT gère la mobilité des nœuds par une reconstruction à la source des routes. S'il semble difficile d'offrir des garanties durables dans un contexte mobile ou même de conserver des routes de secours risquant d'être invalidées du fait de la dynamique de la topologie et de l'occupation du médium, une solution de reconstruction de routes au niveau du point de cassure basée sur les informations de voisinage disponibles serait à étudier.

Les études présentées ici n'ont été pour la plupart validées que par simulation. Ces mesures ont permis de valider le fonctionnement des différents mécanismes et d'en évaluer les performances. Cependant, les simulateurs, pour des raisons de complexité et de temps de calcul, effectuent un certain nombre d'approximations vis-à-vis de la propagation radio. Il est donc indispensable d'implanter les différentes solutions présentées dans ce document et d'en réaliser une évaluation réelle. De nombreux paramètres tels que la taille de la zone de détection de porteuse, l'influence des différents brouilleurs, le mode exact de partage du médium seront sans doute différents. Cependant, les différences constatées n'invalident pas les mécanismes, mais conduiront certainement à un réglage des paramètres de BRuIT.

Spécification des paquets de contrôle de BRuIT

A.1 Paquets *Hello*

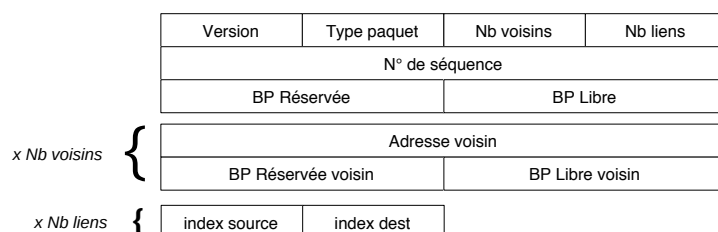


FIG. A.1 – Format d'un paquet *Hello*

Les paquets *Hello* sont transmis régulièrement en mode diffusion par tous les mobiles du réseau. Ils sont utilisés afin de propager des informations sur la topologie et l'état du canal radio dans un rayon de deux sauts. Leurs en-têtes sont composés des champs suivants :

- **Version** : ce champ indique la version du protocole BRuIT utilisée, il permet de faire évoluer le format des paquets en fonction de la version du protocole ;
- **Type paquet** : ce champ indique le type du paquet de contrôle. Dans ce cas, il indiquera qu'il s'agit d'un paquet *Hello* ;
- **Nb voisins** : ce champ indique le nombre de nœuds sur lesquels le corps du paquet contient des informations, le nœud émetteur exclu.
- **Nb liens** : ce champ indique le nombre de liens décrits dans le corps du paquet ;
- **N° de séquence** : ce champ, couplé à l'adresse de l'émetteur du paquet permet d'identifier le paquet de façon unique ;
- **BP réservée** : ce champ indique le volume de bande passante utilisé par le nœud émetteur pour l'émission et le routage des paquets de données ;
- **BP libre** : ce champ indique le volume de bande disponible au niveau de l'émetteur du paquet, compte tenu des trafics présents dans son voisinage.

Pour chaque voisin, le corps du paquet indique par ailleurs :

- **Adresse voisin** : ce champ indique l'adresse IP du voisin concerné ;
- **BP réservée voisin** : ce champ indique le volume de bande passante utilisé par le voisin concerné pour l'émission et le routage des paquets de données ;
- **BP libre** : ce champ indique le volume de bande disponible au niveau du voisin concerné, compte tenu des trafics présents dans son voisinage.

Enfin, pour chaque lien, le couple (index source ; index dest) indique la présence d'un lien entre deux voisins. Ces informations sur les liens entre voisins permettent de compléter la vision de la topologie à deux sauts dont disposent les mobiles.

A.2 Paquets de recherche de route

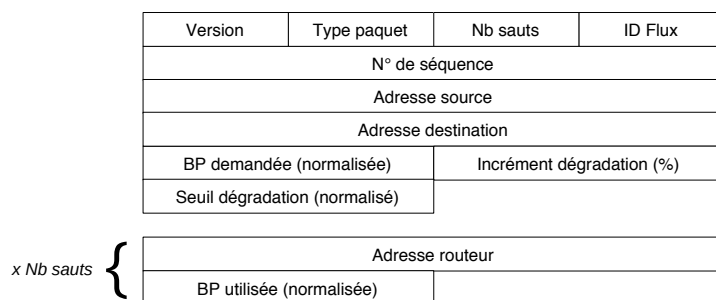


FIG. A.2 – Format d'un paquet *Route Request*

Les paquets *Route Request* sont émis lorsqu'une application désire rechercher une route répondant à certains critères de qualité de service vers une destination du réseau. Ces paquets sont transmis par inondation et sont composés des champs suivants :

- **Version** : ce champ indique la version du protocole BRuIT utilisée, il permet de faire évoluer le format des paquets en fonction de la version du protocole ;
- **Type paquet** : ce champ indique le type du paquet de contrôle. Dans ce cas, il indiquera qu'il s'agit d'un paquet *Route Request* ;
- **Nb sauts** : ce champ indique le nombre de sauts parcourus par la requête jusque là ;
- **ID Flux** : ce champ indique l'identifiant du flux correspondant à la requête. Les paquets de données de ce flux pourront être identifiés par la valeur du champ TOS de l'adresse IP ;
- **N° de séquence** : ce champ, couplé à l'adresse de l'émetteur du paquet, à l'adresse du récepteur et à l'identifiant de flux permet d'identifier la requête de façon unique ;
- **Adresse source** : adresse IP de la source du flux ;
- **Adresse destination** : adresse IP de la destination du flux ;
- **BP demandée** : ce champ indique le volume de bande passante normalisé demandé par l'application ;
- **Incrément dégradation** : ce champ indique le pourcentage de bande passante à retirer au débit accordé au flux lorsqu'un problème de performances est constaté ;
- **Seuil dégradation** : ce champ indique le seuil en deçà duquel il est préférable de reconstruire la route par rapport à effectuer une nouvelle dégradation du débit.

Pour chaque routeur traversé, le corps du paquet contient en outre :

- **Adresse routeur** : ce champ indique l'adresse IP du routeur concerné ;
- **BP utilisée** : ce champ indique le volume de bande passante normalisé utilisé par le routeur concerné pour la retransmission du flux.

A.3 Paquets de confirmation de réservation

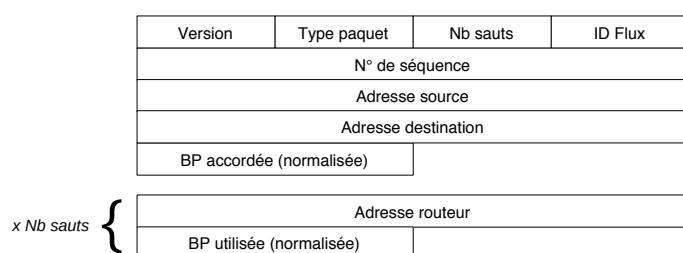


FIG. A.3 – Format d'un paquet *Resources Reservation*

Les paquets de réservation sont émis par la destination d'un *Route request* lorsque celle-ci accepte la réservation de bande passante demandée. Ces paquets empruntent le chemin inverse de la demande de route correspondante, provoquant la réservation effective des ressources sur son chemin. Ils sont émis en mode point-à-point et sont composés des champs suivants :

- **Version** : ce champ indique la version du protocole BRuIT utilisée, il permet de faire évoluer le format des paquets en fonction de la version du protocole ;
- **Type paquet** : ce champ indique le type du paquet de contrôle. Dans ce cas, il indiquera qu'il s'agit d'un paquet *Route Reply* ;
- **Nb sauts** : ce champ indique le nombre de sauts parcourus par la requête jusque là ;
- **ID Flux** : ce champ indique l'identifiant du flux correspondant à la requête. Les paquets de données de ce flux pourront être identifiés par la valeur du champ TOS de l'adresse IP ;
- **N° de séquence** : ce champ, couplé à l'adresse de l'émetteur du paquet, à l'adresse du récepteur et à l'identifiant de flux permet d'identifier la réponse de façon unique ;
- **Adresse source** : adresse IP de la source du flux ;
- **Adresse destination** : adresse IP de la destination du flux ;
- **BP accordée** : ce champ indique le volume de bande passante normalisé accordé à l'application ;
- **Incrément dégradation** : ce champ indique le pourcentage de bande passante à retirer au débit accordé au flux lorsqu'un problème de performances est constaté ;
- **Seuil dégradation** : ce champ indique le seuil en deçà duquel il est préférable de reconstruire la route par rapport à effectuer une nouvelle dégradation du débit.

Pour chaque routeur traversé, le corps du paquet contient en outre :

- **Adresse routeur** : ce champ indique l'adresse IP du routeur concerné ;
- **BP utilisée** : ce champ indique le volume de bande passante normalisé utilisé par le routeur concerné pour la retransmission du flux.

A.4 Paquets de libération et d'erreurs

Version	Type paquet	ID Flux
Adresse source		
Adresse destination		

FIG. A.4 – Format des paquets *Route Teardown* et *Route Broken*

Les paquets de libération de route et de notification d'erreur sont similaires, seul le champ *Type* diffère. Ces paquets ont pour vocation de provoquer la libération des ressources réservées de façon explicite dans le réseau. Les paquets de notification d'erreur permettent par ailleurs à la source du flux d'être notifié rapidement d'une cassure de route, lui permettant ainsi de ré-effectuer une demande rapidement. Ces paquets sont composés des champs suivants :

- **Version** : ce champ indique la version du protocole BRuIT utilisée, il permet de faire évoluer le format des paquets en fonction de la version du protocole ;
- **Type paquet** : ce champ indique le type du paquet de contrôle. Dans ce cas, il indiquera qu'il s'agit d'un paquet *Route Error* ou *Route Teardown* selon les cas ;
- **ID Flux** : ce champ indique l'identifiant du flux correspondant à la requête. Les paquets de données de ce flux pourront être identifiés par la valeur du champ TOS de l'adresse IP ;
- **Adresse source** : adresse IP de la source du flux ;
- **Adresse destination** : adresse IP de la destination du flux.

B.1 Journaux nationaux

- Claude Chaudet, Isabelle Guérin-Lassous - *Routage QoS et réseaux ad-hoc : de l'état de lien à l'état de noeud* – Accepté pour publication dans Techniques et Science Informatiques (TSI) - Numéro spécial réseaux et protocoles (Rapport de recherche INRIA n° 4700 – janvier 2003)

B.2 Conférences internationales avec comité de lecture

- Guillaume Chelius, Claude Chaudet – *Handoff Notification in Wireless Hybrid Networks* – Proceedings of the sixth IFIP/IEEE International Conference on Mobile and Wireless Communication Networks (MWCN) 2004 – Octobre 2004 – Paris, France.
- Claude Chaudet, Isabelle Guérin Lassous, Eric Thierry, Bruno Gaujal – *Study of the Impact of Asymmetry and Carrier Sense Mechanism in IEEE 802.11 Multi-hops Networks through a Basic Case* – Proceedings of The ACM Workshop on Performance Evaluation of Wireless Ad Hoc, Sensor, and Ubiquitous Networks (PE-WASUN '04) – Octobre 2004 – Venise, Italie.
- Claude Chaudet, Olivier Festor, Isabelle Guérin Lassous, Radu State – *A Managed Bandwidth Reservation Protocol for Ad Hoc Networks* – Proceedings of The First International Workshop on Service Assurance with Partial and Intermittent Resources (SAPIR 2004) held in conjunction with ICT 2004 – Août 2004 – Fortaleza, Brésil.
- Claude Chaudet, Isabelle Guérin Lassous, Janez Žeronvik – *A Distributed Algorithm for Bandwidth Allocation in Stable Ad Hoc Networks* – First Working Conference on Wireless On-demand Network Systems (WONS 2004) – Janvier 2004 – Madonna di Campiglio, Italie.
- Claude Chaudet, Isabelle Guérin-Lassous – *BRuIT - Bandwidth Reservation under InTerferences influence* – European Wireless 2002 (EW2002) – Février 2002 – Florence, Italie.
- Karell Bertet, Claude Chaudet, Isabelle Guérin Lassous, Laurent Viennot – *Impact of Interferences on Bandwidth Reservation for Ad Hoc Networks : a First Theoretical Study* – IEEE Symposium on Ad-Hoc Wireless Networks (GLOBECOM SAWN'2001) – Novembre 2001 – San Antonio, USA.

B.3 Conférences nationales avec comité de lecture

- Guillaume Chelius, Claude Chaudet, Natalie Whitlock – *Notification de mobilité dans les réseaux hybrides avec infrastructure sans-fil* – Algotel 2004 – Mai 2004 – Batz-sur-Mer, France.
- Claude Chaudet, Isabelle Guérin-Lassous – *Influence de l'asymétrie et des interférences sur l'équité de l'accès au médium dans les réseaux 802.11b* – Algotel 2003 – Mai 2003 – Banyuls-sur-Mer, France.
- Claude Chaudet, Isabelle Guérin-Lassous – *Prise en compte des interférences dans la réservation de bande passante : le protocole BRuIT* – Algotel 2002 - Mai 2002 - Mèze, France.
- Claude Chaudet - *Qualité de service et réseaux ad-hoc - un état de l'art* – MS3G 2001 – Services liés à la mobilité et réseaux mobiles de 3ème génération – Décembre 2001 – Lyon, France.

Bibliographie

- [Abr70] Norman Abramson. The ALOHA System – Another Alternative for Computer Communications. Dans *Proceedings of the 1970 Fall Joint Computer Conference*, volume 36, pages 177–186, Montvale, New Jersey, USA, 1970. AFIPS Press.
- [AC01] Imad Aad et Claude Castelluccia. Differentiation mechanisms for IEEE 802.11. Dans *Proceedings of the Twentieth Annual Joint Conference of the IEEE Computer and Communications Societies (IEEE Infocom)*, Anchorage, Alaska, USA, Avril 2001.
- [AC03a] Imad Aad et Claude Castelluccia. Enhancing IEEE 802.11 performance in congested environments. *Annales des télécommunications – Réseaux locaux : perspectives de l'accès radio*, 3-4 :397–416, Avril 2003.
- [AC03b] Imad Aad et Claude Castelluccia. Priorities in WLANs. *Computer Networks*, 41(4) :505–526, Mars 2003.
- [ACVS02] Gahng-Seop Ahn, Andrew T. Campbell, Andras Veres, et Li-Hsiang Sun. Supporting Service Differentiation for Real-Time and Best-Effort Traffic in Stateless Wireless Ad Hoc Networks (SWAN). *IEEE Transactions on Mobile Computing*, 1(3) :192–207, Septembre 2002.
- [ACVS03] Gahng-Seop Ahn, Andrew T. Campbell, Andras Veres, et Li-Hsiang Sun. SWAN. Internet Draft – draft-ahn-swan-manet-00.txt, Août 2003.
- [ADLK01] Patrick R. Amestoy, Iain S. Duff, Jean-Yves L'Excellent, et Jacko Koster. MUMPS : A General Purpose Distributed Memory Sparse Solver. *Lecture Notes in Computer Science*, 1947, 2001.
- [BBC⁺98] Steven Blake, David Black, Mark Carlson, Elwyn Davies, Zheng Wang, et Walter Weiss. An Architecture for Differentiated Services. Internet Request For Comments RFC 2475, Internet Engineering Task Force, Décembre 1998.
- [BBR97] Yair Bartal, Johna W. Byers, et Danny Raz. Global optimization using local information with applications to flow control. Dans *Proceedings of the thirty-eighth IEEE Symposium on Foundations of Computer Science (FOCS'97)*, pages 303–312, Miami Beach, Floride, USA, Octobre 1997.
- [BCG02] Raffaele Bruno, Marco Conti, et Enrico Gregori. Bluetooth : Architecture, Protocols and Scheduling Algorithms. *Cluster Computing*, 5(2) :117–131, Avril 2002.
- [BCGLV01] Karell Bertet, Claude Chaudet, Isabelle Guérin Lassous, et Laurent Viennot. Impact of Interferences on Bandwidth Reservation for Ad Hoc Networks : a First Theoretical Study. Dans *Proceedings of the IEEE Global Telecommunications Conference (Globecom 2001)*, San Antonio, Texas, USA, Novembre 2001.
- [BCS94] Robert Braden, David Clark, et Scott Shenker. Integrated Services in the Internet Architecture : an Overview. Internet Request For Comments RFC 1633, Internet Engineering Task Force, Juin 1994.
- [BCV01] Michael Barry, Andrew T. Campbell, et Andras Veres. Distributed Control Algorithms for Service Differentiation in Wireless Packet Networks. Dans *Proceedings of the Twentieth Annual Joint Conference of the IEEE Computer and Communications Societies (IEEE Infocom)*, Anchorage, Alaska, USA, Avril 2001.
- [BDSZ94] Vaduvur Bharghavan, Alan Demers, Scott Shenker, et Lixia Zhang. MACAW : a media access protocol for wireless LAN's. Dans *Proceedings of the conference on Communications*

- architectures, protocols and applications (ACM Sigcomm '94)*, pages 212–225, Londres, Royaume-Uni, Août 1994.
- [BG87] Dimitri P. Bertsekas et Robert Gallager. *Data Networks*. Prentice Hall, Englewood Cliffs, New Jersey, 1987.
- [BGLV00] Karell Bertet, Isabelle Guérin Lassous, et Laurent Viennot. Un premier pas vers la réservation de bande passante dans les réseaux radio. Dans *Acte des Deuxièmes Rencontres Francophones sur les aspects Algorithmiques des Télécommunications (Algotel 2000)*, pages 25–30, La Rochelle, France, Mai 2000.
- [Bia00] Giuseppe Bianchi. Performance Analysis of the IEEE 802.11 Distributed Coordination Function. *IEEE Journal on Selected Areas in Communications*, 18(3) :535–547, Mars 2000.
- [Blu99] Bluetooth Special Interest Group. *Specification of the Bluetooth System 1.0b, Volume 1 : Core*. Gästebuch, Foren, Décembre 1999.
- [BM01] Thomas Bonald et Laurent Massoulié. Impact of Fairness on Internet Performance. Dans *Proceedings of the Joint International Conference on Measurement and Modeling of Computer Systems (SIGMETRICS)*, Cambridge, Massachusetts, USA, Juin 2001.
- [BMAA02] Mounir Benzaid, Pascale Minet, et Khaldoun Al Agha. Integrating fast mobility in the OLSR routing protocol. Dans *Proceedings of the Fourth IEEE Conference on Mobile and Wireless Communications Networks (MWCN 2002)*, Stockholm, Suède, Septembre 2002.
- [BS92] Piotr Berman et Georg Schnitger. On the complexity of approximating the independent set problem. *Information and Computation*, 96(1) :77–94, Janvier 1992.
- [BWK00] Brahim Bensaou, Yu Wang, et Chi Chung Ko. Fair Medium Access in 802.11 Based Wireless Ad-Hoc Networks. Dans *Proceedings of the First International Workshop on Mobile Ad Hoc Networking and Computing (MobiHoc 2000)*, Boston, Massachusetts, USA, Août 2000.
- [BZB⁺97] Robert Braden, Lixia Zhang, Steven Berson, Shai Herzog, et Sigih Jamin. Resource ReSerVation Protocol (RSVP) – Version 1 Functional Specification. Internet Request For Comments RFC 2205, Internet Engineering Task Force, Septembre 1997.
- [CBD02] Tracy Camp, Jeff Boleng, et Vanessa Davies. A Survey of Mobility Models for Ad Hoc Network Research. *Wireless Communication and Mobile Computing WCMC Special issue on Mobile Ad Hoc Networking Research, Trends and Applications*, 2(5) :483–502, Septembre 2002.
- [CCG98] Federico Cali, Marco Conti, et Enrico Gregori. IEEE 802.11 Wireless LAN : Capacity Analysis and Protocol Enhancement. Dans *Proceedings of the Seventeenth Annual Joint Conference of the IEEE Computer and Communications Societies (IEEE Infocom)*, San Francisco, Californie, USA, Mars 1998.
- [CF03] Guillaume Chelius et Éric Fleury. Design of a hybrid routing architecture. Dans *Proceedings of the Fifth IFIP TC6 International Conference on Mobile Wireless Communications Networks (MWCN 2003)*, Singapour, Octobre 2003.
- [CG97] Harshal S. Chhaya et Sanjay Gupta. Performance modeling of asynchronous data transfer methods of IEEE 802.11 MAC protocol. *Wireless Networks*, 3(3) :217–234, 1997.
- [CGC00] A. Campbell et J. Gomez-Castellanos. IP Micromobility Protocols. *ACM SIGMOBILE Mobile Computing and Communications Review*, 4(4) :45–53, Octobre 2000.
- [CGK⁺02] A. Campbell, J. Gomez, S. Kim, Z. Turanyi, C-Y. Wan, et A. Valkó. Comparison of IP Micromobility Protocols. *IEEE Wireless Communications*, 9(1) :72–82, Février 2002.
- [CGT97] Tsu-Wei Chen, Mario Gerla, et Jack Tzu-Chieh Tsai. QoS routing performance in a multi-hop, wireless networks. Dans *Proceedings of the sixth International Conference on Universal Personal Communications (ICUPC 97)*, San Diego, Californie, USA, Octobre 1997.
- [CJ03] Thomas Clausen et Philippe Jacquet. Optimized Link State Routing Protocol (OLSR). Internet Request For Comments RFC 3626, Internet Engineering Task Force, Octobre 2003.

- [CM99] Scott Corson et Joseph Macker. Mobile Ad hoc Networking (MANET) : Routing Protocol Performance Issues and Evaluation Considerations. Internet Request For Comments RFC 2501, Internet Engineering Task Force, Janvier 1999.
- [CN99] Shigang Chen et Klara Nahrstedt. Distributed Quality of Service Routing in Ad-Hoc Networks. *IEEE Journal on Selected Areas in Communications*, 17(8) :1488–1505, Août 1999.
- [CP96] Ramón Cáceres et Venkata N. Padmanabhan. Fast and scalable handoffs for wireless internetwork. Dans *Proceedings of the second annual international conference on Mobile computing and networking (MobiCom '96)*, pages 56–66, Rye, New York, USA, Novembre 1996.
- [Dan49] George B. Dantzig. Programming of interdependant activities : II. mathematical model. *Econometrica*, 17 :200–211, 1949.
- [DC99] Jiunn Deng et Ruay-Shiung Chang. A Priority Scheme for IEEE 802.11 DCF Access Method. *IEICE Transactions on Communications*, ES82-B(1) :96–102, Janvier 1999.
- [Dea96a] Deane, John. WATM MAC requirements,” ATM Forum/96-0786, Juin 1996.
- [Dea96b] Deane, John. WATM PHY requirements,” ATM Forum/96-0785, Juin 1996.
- [DGL02] Dominique Dhoutaut et Isabelle Guérin Lassous. Impact of Heavy Traffic Beyond Communication Range in Multi-Hops Ad Hoc Networks. Dans *Proceedings of the Third International Network Conference (INC 2002)*, Plymouth, Royaume-Uni, Juillet 2002.
- [DGL03] Dominique Dhoutaut et Isabelle Guérin Lassous. Experiments with 802.11b in ad hoc configurations. Dans *Proceedings of the fourteenth IEEE International Symposium on Personal, Indoor and Mobile Radio Communications (PIMRC 2003)*, pages 1618–1622, Beijing, Chine, Septembre 2003. IEEE Press.
- [Dho02] Dominique Dhoutaut. *Analyse de 802.11 dans un contexte ad hoc : de l'analyse à l'expérimentation*. PhD thesis, Institut National des Sciences Appliquées de Lyon, Décembre 2002.
- [Dij59] Edsger Wybe Dijkstra. A note on two problems in connexion with graphs. *Numerische Mathematik*, 1 :269–271, 1959.
- [ETS98] ETSI BRAN. High Performance Radio Local Area Network (HIPERLAN) Type 1 – Functional specification , Juillet 1998.
- [ETS00a] ETSI BRAN. High Performance Radio Local Area Network (HIPERLAN) Type 2 Technical Specification TS 101 475 – Physical Layer, Avril 2000.
- [ETS00b] ETSI BRAN. High Performance Radio Local Area Network (HIPERLAN) Type 2 Technical Specification TS 101 493-1 – Packet based Convergence Layer ; Part 1 : Common Part, Avril 2000.
- [ETS00c] ETSI BRAN. High Performance Radio Local Area Network (HIPERLAN) Type 2 Technical Specification TS 101 493-2 – Packet based Convergence Layer ; Part 2 : Ethernet Service Specific Convergence Sublayer (SSCS), Avril 2000.
- [ETS00d] ETSI BRAN. High Performance Radio Local Area Network (HIPERLAN) Type 2 Technical Specification TS 101 761-1 – Data Link Control (DLC) Layer ; Part 1 : Basic Data Transport Functions, Avril 2000.
- [ETS00e] ETSI BRAN. High Performance Radio Local Area Network (HIPERLAN) Type 2 Technical Specification TS 101 761-2 – Data Link Control (DLC) Layer ; Part 2 : Radio Link Control (RLC) Sublayer, Avril 2000.
- [ETS00f] ETSI BRAN. High Performance Radio Local Area Network (HIPERLAN) Type 2 Technical Specification TS 101 763-1 – Cell based Convergence Layer ; Part 1 : Common Part, Avril 2000.
- [ETS00g] ETSI BRAN. High Performance Radio Local Area Network (HIPERLAN) Type 2 Technical Specification TS 101 763-2 – Cell based Convergence Layer ; Part 2 : UNI Service Specific Part, Avril 2000.

- [FBW02] Zuyuan Fang, Brahim Bensaou, et Yu Wang. Performance evaluation of a fair backoff algorithm for IEEE 802.11 DFWMAC. Dans *Proceedings of the Third ACM International Symposium on Mobile Ad Hoc Networking and Computing (MobiHoc 2002)*, pages 48–57, Lausanne, Suisse, Juin 2002.
- [GGK01] Piyush Gupta, Robert Gray, et P.R. Kumar. An Experimental Scaling Law for Ad Hoc Networks. Preprint, en ligne : http://black1.csl.uiuc.edu/~prkumar/html_files/postscript_files.html, Mai 2001.
- [GJM03] Leonidas Georgiadis, Philippe Jacquet, et Bernard Mans. Bandwidth Reservation in Multihop Wireless Networks : Complexity and Mechanisms. Technical Report 4876, INRIA, en ligne et disponible à <http://www.inria.fr/rrrt/rr-4876.html>, Juillet 2003.
- [GJM04] Leonidas Georgiadis, Philippe Jacquet, et Bernard Mans. Bandwidth Reservation in Multihop Wireless Networks : Complexity and Mechanisms. Dans *Proceedings of the International Workshop on Wireless Ad Hoc Networking (WWAN 2004)*, Tokyo, Japon, Mars 2004.
- [GJP03] Eva Gustafsson, Annika Jonsson, et Charles E. Perkins. Mobile IPv4 Regional Registration. Internet Draft – draft-ietf-mobileip-reg-tunnel-08.txt, Novembre 2003.
- [GK00] Piyush Gupta et P.R. Kumar. The capacity of wireless networks. *IEEE Transactions on Information Theory*, IT-46(2) :388–404, Mars 2000.
- [GN02] Antonio Grilo et Mario Nunes. Performance Evaluation of IEEE 802.11e. Dans *Proceedings of the Thirteenth IEEE International Symposium on Personal, Indoor and Radio Communications (PIMRC 2002)*, Lisbonne, Portugal, Septembre 2002.
- [GS02] Irina Gerasimov et Robert Simon. A Bandwidth-Reservation Mechanism for On-Demand Ad hoc Path Finding. Dans *Proceedings thirty-fifth Annual Simulation Symposium (ANSS-35 2002)*, San Diego, Californie, USA, Avril 2002.
- [GW01] Mike Gertz et Steve Wright. Object-oriented Software for Quadratic Programming. Technical Report ANL/MCS-P891-1000, Argonne National Laboratory, Mathematics and Computer Science Division, 2001.
- [HB01] Xiao Long Huang et Brahim Bensaou. On Max-min Fairness and Scheduling in Wireless Ad-Hoc Networks : Analytical Framework and Implementation. Dans *Proceedings of the 2001 ACM Symposium on Mobile Ad Hoc Networking and Computing (MobiHoc 2001)*, Long Beach, Californie, USA, Octobre 2001.
- [HBBW04] Marc Heissenbüttel, Torsten Braun, Thomas Bernoulli, et Markus Wächli. BLR : Beacon-Less Routing Algorithm for Mobile Ad-Hoc Networks. *Computer Communications Journal*, 27(11) :1076–1086, Juillet 2004.
- [HBWW99] Juha Heinanen, Fred Baker, Walter Weiss, et John Wroclawski. Assured Forwarding PHB Group. Internet Request For Comments RFC 2597, Internet Engineering Task Force, Juin 1999.
- [HG01] Armin Heindl et Reinhard German. Performance modeling of IEEE 802.11 wireless LANs with stochastic Petri nets. *Performance Evaluation*, 44(1-4) :139–164, Avril 2001.
- [Hin02] Sébastien Hinderer. Mise en œuvre d’un protocole de routage avec qualité de service pour réseaux *ad hoc*. Rapport de stage, ENS Lyon, en ligne et disponible à http://ens-lyon.free.fr/rapports/info/Sebastien.Hinderer_1.pdf, Juillet 2002.
- [HKWHC01] Eric Hsiao-Kuang Wu, Yi-Zhan Huang, et Jui-Hao Chiang. Dynamic Adaptive Routing for Heterogeneous Wireless Network. Dans *Proceedings of the IEEE Global Telecommunications Conference (Globecom 2001)*, San Antonio, Texas, USA, Novembre 2001.
- [HL00] Ying-Kwei Ho et Ru-Sheng Liu. On-Demand QoS-Based Routing Protocol for Ad Hoc Mobile Wireless Networks. Dans *Proceedings of the Fifth IEEE Symposium on Computers and Communications (ISCC 2000)*, Antibes, France, Juillet 2000.
- [HPS02] Zygmunt J. Haas, Marc R. Pearlman, et Prince Samar. The Zone Routing Protocol (ZRP) for Ad Hoc Networks. Internet Draft – draft-ietf-manet-zone-zrp-04.txt, Juillet 2002.

- [HRBSD03] Martin Heusse, Franck Rousseau, Gilles Berger-Sabbatel, et Andrzej Duda. Performance Anomaly of 802.11b. Dans *Proceedings of the Twenty-Second Annual Joint Conference of the IEEE Computer and Communications Societies (Infocom 2003)*, pages 836–843, San Francisco, Californie, USA, Avril 2003.
- [IEE85] IEEE Computer Society Std 802.3-1985. *Carrier Sense Multiple Access with Collision Detection (CSMA/CD) Access Method and Physical Layer Specifications*. The Institute of Electrical and Electronics Engineers, 1985.
- [IEE97] IEEE Standard for Information Technology – Telecommunications and Information Exchange between Systems. Local and Metropolitan Area Network – Specific Requirements – Part 11 : Wireless LAN Medium Access Control (MAC) and Physical Layer (PHY) Specifications, 1997.
- [IEE99a] IEEE Standard for Information Technology – Telecommunications and Information Exchange between Systems. Local and Metropolitan Area Network – Specific Requirements – Part 11 : Wireless LAN Medium Access Control (MAC) and Physical Layer (PHY) Specifications — Higher-speed physical layer extension in the 2.4 GHz band, 1999.
- [IEE99b] IEEE Standard for Information Technology – Telecommunications and Information Exchange between Systems. Local and Metropolitan Area Network – Specific Requirements – Part 11 : Wireless LAN Medium Access Control (MAC) and Physical Layer (PHY) Specifications — High-speed physical layer in the 5 GHz band, 1999.
- [IEE03] IEEE Standard for Information Technology – Telecommunications and Information Exchange between Systems. Local and Metropolitan Area Network – Specific Requirements – Part 11 : Wireless LAN Medium Access Control (MAC) and Physical Layer (PHY) Specifications — Further Higher-Speed Physical Layer Extension in the 2.4 GHz Band, 2003.
- [JMH03] David B. Johnson, David A. Maltz, et Yih-Chun Hu. The Dynamic Source Routing Protocol for Mobile Ad Hoc Networks (DSR). Internet Draft – draft-ietf-manet-dsr-09.txt, Avril 2003.
- [JNP99] V. Jacobson, Kathleen Nichols, et Kedarnath Poduri. An Expedited Forwarding PHB. Internet Request For Comments RFC 2598, Internet Engineering Task Force, Juin 1999.
- [JS03] Jangeun Jun et Mihail L. Sichitiu. The nominal capacity of wireless mesh networks. *IEEE Wireless Communications*, 10(5) :8–14, Octobre 2003.
- [JT87] John Jubin et Janet D. Tornow. The DARPA packet radio network protocols. *Proceedings of the IEEE, Special Issue on Packet Radio Networks*, 75(1) :21–32, Janvier 1987.
- [Kar90] Phil Karn. MACA - a new channel access method for packet radio. Dans *Proceedings of the ARRL/CRRL Amateur Radio Ninth Computer Networking Conference*, pages 134–140, Ontario, Canada, Septembre 1990.
- [KG02] Manthos Kazantzidis et Mario Gerla. End-to-end versus Explicit Feedback Measurement in 802.11 Networks. Dans *Proceedings of the Seventh IEEE Symposium on Computers and Communications (ISCC 2002)*, Taormina, Italie, Juillet 2002.
- [KGBK78] Robert E. Kahn, Steven A. Gronemeyer, Jerry Burchfiel, et Ronald C. Kunzelman. Advances in packet radio technology. *Proceedings of the IEEE*, 66(11) :1468–1496, Novembre 1978.
- [KGL01] Manthos Kazantzidis, Mario Gerla, et Sung-Ju Lee. Permissible Throughput Network Feedback in AODV MANETs. Dans *Proceedings of the thirty-seventh IEEE International Conference on Communications (ICC)*, Helsinki, Finlande, Juin 2001.
- [KMT98] Frank P. Kelly, Aman K. Maulloo, et David H. K. Tan. Rate control in communication networks : shadow prices, proportional fairness and stability. *Journal of the Operational Research Society*, 49(3) :237–252, Mars 1998.
- [KT75] Leonard Kleinrock et Fouad A. Tobagi. Random access techniques for data transmission over packet-switched radio channels. Dans *Proceedings of National Computer Conference*, pages 187–201, Anaheim, Californie, USA, Mai 1975.

- [LAC01] Seoung Bum Lee, Gahng Seop Ahn, et Andrew T. Campbell. Improving UDP and TCP Performance in Mobile Ad Hoc Networks with INSIGNIA. *IEEE Communication Magazine*, 36(6) :156–165, Juin 2001.
- [LAS01a] Anders Lindgren, Andreas Almquist, et Olov Schelén. Evaluation of Quality of Service Schemes for IEEE 802.11 Wireless LANs. Dans *Proceedings of The twenty-sixth Annual IEEE Conference on Local Computer Networks (LCN)*, pages 348–351, Tampa, Floride, USA, Novembre 2001.
- [LAS01b] Anders Lindgren, Andreas Almquist, et Olov Schelén. Quality of Service Schemes for IEEE 802.11 : A Simulation Study. Dans *Proceedings of the ninth International Workshop on Quality of Service (IWQoS)*, pages 281–287, Karlsruhe, Allemagne, Juin 2001.
- [LAS03] Anders Lindgren, Andreas Almquist, et Olov Schelén. Quality of Service Schemes for IEEE 802.11 Wireless LANs - An Evaluation. *Mobile Networks and Applications*, 8(3) :223–235, Juin 2003.
- [LAZC99] Seoung Bum Lee, Gahng Seop Ahn, Xiaowei Zhang, et Andrew T. Campbell. INSIGNIA. Internet Draft – draft-ietf-manet-insignia-01.txt, Novembre 1999.
- [LAZC00] Seoung Bum Lee, Gahng Seop Ahn, Xiaowei Zhang, et Andrew T. Campbell. INSIGNIA : An IP-Based Quality of Service Framework for Mobile ad Hoc Networks. *Journal on Parallel and Distributed Computing*, 60(4) :374–406, Avril 2000.
- [LBDC⁺01] Jinyang Li, Charles Blake, Douglas S.J. De Couto, Hu Imm Lee, et Robert Morris. Capacity of Ad Hoc wireless networks. Dans *Proceedings of the Seventh annual international conference on Mobile computing and networking (MobiCom 2001)*, pages 61–69, Rome, Italie, Juillet 2001.
- [LG97] Chunhung Richard Lin et Mario Gerla. Asynchronous and Multimedia Multihop Wireless Networks. Dans *Proceedings of the Sixteenth Annual Joint Conference of the IEEE Computer and Communications Societies (IEEE Infocom)*, pages 118–125, Kobe, Japon, Avril 1997.
- [Lin01] Chunhung Richard Lin. An On-demand QoS Routing Protocol for Mobile Ad Hoc Networks. Dans *Proceedings of the Twentieth Annual Joint Conference of the IEEE Computer and Communications Societies (IEEE Infocom)*, Anchorage, Alaska, USA, Avril 2001.
- [LL99] Chunhung Richard Lin et Jain-Shing Liu. QoS Routing in Ad Hoc Wireless Networks. *IEEE Journal on Selected Areas in Communications*, 17(8) :1426–1438, Août 1999.
- [LL00] Chunhung Richard Lin et Chung-Ching Liu. An On-Demand QoS Routing Protocol for Mobile Ad Hoc Networks. Dans *Proceedings of the IEEE Global Telecommunications Conference (Globecom 2000)*, volume 3, pages 1783–1787, San Francisco, Californie, USA, Novembre 2000.
- [LTWS01] Wen-Hwa Liao, Yu-Chee Tseng, Shu-Ling Wang, et Jang-Ping Sheu. A Multi-path QoS Routing Protocol in a Wireless Mobile ad Hoc Network. Dans *Proceedings of the First International Conference on Networking (ICN 2001)*, volume 2094 of *Lecture Notes in Computer Science*, pages 158–167, Colmar, France, Juillet 2001. Springer.
- [LTWS02] Wen-Hwa Liao, Yu-Chee Tseng, Shu-Ling Wang, et Jang-Ping Sheu. A Multi-path QoS Routing Protocol in a Wireless Mobile ad Hoc Network. *Telecommunication Systems*, 19(3-4) :329–347, Avril 2002.
- [LÊG04a] Zhifei Li, SukumarÊ ÊNandi, et AnilÊK. Gupta. Improving MAC Performance in Wireless Ad Hoc Networks Using Enhanced Carrier Sensing (ECS). Dans *Proceedings of the Third IFIP-TC6 Networking Conference (Networking 2004)*, pages 600–612, Athènes, Grèce, Mai 2004.
- [LÊG04b] Zhifei Li, SukumarÊ ÊNandi, et AnilÊK. Gupta. Modeling the Short-Term Unfairness of IEEE 802.11 in Presence of Hidden Terminals. Dans *Proceedings of the Third IFIP-TC6 Networking Conference (Networking 2004)*, pages 613–625, Athènes, Grèce, Mai 2004.
- [MB76] Robert Metcalfe et David Boggs. Ethernet : Distributed Packet Switching for Local Computer Networks. *Communications of the ACM*, 19(7) :395–404, Juillet 1976.

- [MBAAP02] Anelise Munaretto, Hakim Badis, Khaldoun Al Agha, et Guy Pujolle. A Link-state QoS Routing Protocol for Ad Hoc Networks. Dans *Proceedings of the Fourth IEEE Conference on Mobile and Wireless Communications Networks (MWCN 2002)*, Stockholm, Suède, Septembre 2002.
- [MCM⁺02] Stefan Mangold, Sunghyun Choi, Peter May, Ole Klein, Guido Hiertz, et Lothar Stibor. IEEE 802.11e Wireless LAN for Quality of Service. Dans *Proceedings of European Wireless 2002 (EW2002)*, pages 32–39, Florence, Italie, Février 2002.
- [MLGTR03] Rabah Meraihi, Gwendal Le Grand, Samir Tohmé, et Michel Riguidel. Gestion multi-couches de la qualité de service dans un réseau ad hoc à cœur stable. Dans *Actes du Colloque Francophone sur l'Ingenierie des Protocoles (CFIP)*, Evry, France, Octobre 2003.
- [Moy89] John T. Moy. The OSPF Specification. Internet Request For Comments RFC 1131, Internet Engineering Task Force, Octobre 1989.
- [Moy91] John T. Moy. OSPF Version 2. Internet Request For Comments RFC 1247, Internet Engineering Task Force, Juillet 1991.
- [NABT03] Qiang Ni, Imad Aad, Chadi Barakat, et Thierry Turetletti. Modeling and analysis of slow CW decrease for IEEE 802.11 WLAN. Dans *Proceedings of the fourteenth IEEE International Symposium on Personal, Indoor and Mobile Radio Communications (PIMRC 2003)*, pages 1717–1721, Beijing, Chine, Septembre 2003. IEEE Press.
- [NBBB98] Kathleen Nichols, Steven Blake, Fred Baker, et David Black. Definition of the Differentiated Services Field (DS Field) in the IPv4 and IPv6 Headers. Internet Request For Comments RFC 2474, Internet Engineering Task Force, Décembre 1998.
- [NBMR02] Navid Nikaein, Christian Bonnet, Yan Moret, et Idris A. Rai. 2LQoS- Two-Layered Quality of Service Model for Reactive Routing Protocols for Mobile Ad Hoc Networks. Dans *Proceedings of the Sixth World Multiconference on Systemics, Cybernetics and Informatics (SCI 2002)*, Orlando, Floride, USA, Juillet 2002.
- [NKGB00] Thyagarajan Nandagopal, Tae-Eun Kim, Xia Gao, et Vaduvur Bharghavan. Achieving MAC Layer Fairness in Wireless Packet Networks. Dans *Proceedings of the Sixth Annual International Conference on Mobile Computing and Networking (MobiCom)*, pages 87–98, Boston, Massachusetts, USA, Août 2000.
- [OTG04] Richard G. Ogier, Fred L. Templin, et Lewis Mark G. Topology Dissemination Based on Reverse-Path Forwarding (TBRPF). Internet Request For Comments RFC 3684, Internet Engineering Task Force, Février 2004.
- [PBRD03] Charles E. Perkins, Elizabeth M. Belding-Royer, et Samir Das. Ad hoc On-Demand Distance Vector (AODV) Routing. Internet Request For Comments RFC 3561, Internet Engineering Task Force, Juillet 2003.
- [Pen03] Mathew Penrose. *Random Geometric Graphs*. Oxford University Press, 2003.
- [Per96] Charles E. Perkins. IP Mobility Support. Internet Request For Comments RFC 2002, Internet Engineering Task Force, Octobre 1996.
- [PJ01] Charles E. Perkins et David B. Johnson. Route optimization in mobile ip. Internet Draft – draft-ietf-mobileip-optim-11.txt, Septembre 2001.
- [PM01] Xavier Pallot et Leonard E. Miller. Implementing Message Priority Policies over an 802.11 Based Mobile Ad Hoc Network. Dans *Proceedings of the Twentieth Military Communications Conference (MILCOM 2001)*, pages 855–859, Washington, USA, Octobre 2001.
- [PST95] Serge A. Plotkin, David B. Shmoys, et Éva Tardos. Fast approximation algorithms for fractional packing and covering problems. *Mathematics Of Operations Research*, 20(2) :257–301, Mai 1995.
- [RHZ00] G.V.S. Raju, G. Hernandez, et Q. Zou. Quality of Service Routing in Ad Hoc Networks. Dans *Proceedings of the IEEE Wireless Communications and Networking Conference 2000 (WCNC 2000)*, pages 263–265, Chicago, Illinois, USA, Septembre 2000.
- [RL94] Yakov Rekhter et Tony Li. A Border Gateway Protocol 4 (BGP-4). Internet Request For Comments RFC 1654, Internet Engineering Task Force, Juillet 1994.

- [RLPT⁺02] Ramachandran Ramjee, Thomas La Porta, Sandy Thuel, Kannan Varadhan, et Shie-Yuan Wang. HAWAII : A Domain-based Approach for Supporting Mobility in Wide-Area Wireless Networks. *IEEE/ACM Transactions on Networking*, 6(2), Juin 2002.
- [RLPTV99] Ramachandran Ramjee, Thomas La Porta, Sandy Thuel, et Kannan Varadhan. IP micro-mobility support using HAWAII. Internet Draft – draft-ramjee-micro-mobility-hawaii-00.txt, Février 1999.
- [RMSM01] Elizabeth M. Royer, P. Michael Melliar-Smith, et Louise E. Moser. An Analysis of the Optimum Node Density for Ad hoc Mobile Networks. Dans *Proceedings of the thirty-sixth IEEE International Conference on Communications (ICC 2001)*, Helsinki, Finlande, Juin 2001.
- [RNT03] Lamia Romdhani, Qiang Ni, et Thierry Turletti. AEDCF : Enhanced Service Differentiation for IEEE 802.11 Wireless Ad-Hoc Networks. Dans *Proceedings of the IEEE Wireless Communications and Networking Conference (WCNC 2003)*, La Nouvelle Orleans, Louisiane, USA, Mars 2003.
- [Rob75] Lawrence Roberts. Aloha packet system with and without slots and capture. *Computer Communications Review*, 5(2) :28–42, Avril 1975.
- [RS98] Ram Ramanathan et Martha Steenstrup. Hierarchically-organized, multihop mobile wireless networks for quality-of-service support. *Mobile Networks and Applications, Special issue on mobile multimedia communications*, 3(1) :101–119, Juin 1998.
- [SCN03] Samarth H. Shah, Kai Chen, et Klara Nahrstedt. Available Bandwidth Estimation in IEEE 802.11-based Wireless Networks. Dans *Proceedings of The first ISMA/CAIDA Bandwidth Estimation Workshop (BEst 2003)*, San Diego, Californie, USA, Décembre 2003.
- [SK96] João L. Sobrinho et A. S. Krishnakumar. Real-Time Traffic over the IEEE 802.11 Medium Access Control Layer. *Bell Labs Technical Journal*, 1(2) :172–187, 1996.
- [SPG97] Scott Shenker, Craig Partridge, et Roch Guerin. Specification of Guaranteed Quality of Service. Internet Request For Comments RFC 2212, Internet Engineering Task Force, Septembre 1997.
- [SSB99] Prasun Sinha, Raghupathy Sivakumar, et Vaduvur Bharghavan. CEDAR : a Core Extraction Distributed Ad hoc Routing algorithm. *IEEE Journal on Selected Areas in Communications*, 17(8) :1454–1465, Août 1999.
- [TC01] Y. C. Tay et K. C. Chua. A Capacity Analysis for the IEEE 802.11 MAC Protocol. *Wireless Networks*, 7(2) :159–171, Mars 2001.
- [TK75a] Fouad A. Tobagi et Leonard Kleinrock. Packet Switching in Radio Channels : Part I – Carrier Sense Multiple-Access Modes and Their Throughput-Delay Characteristics. *IEEE Transactions on Communications*, COM-23(12) :1400–1416, Décembre 1975.
- [TK75b] Fouad A. Tobagi et Leonard Kleinrock. Packet Switching in Radio Channels : Part II – The hidden terminal problem in carrier multiple-access and the busy-tone solution. *IEEE Transactions on Communications*, COM-23(12) :1417–1433, Décembre 1975.
- [Typ01] Ville Typpö. Mobility within Wireless Ad Hoc Networks : Towards Hybrid Wireless Multihop Networks. Master’s thesis, VTT Electronics and University of Oulu, Oulu, Finlande, 2001.
- [Val99] Andras G. Valkó. Cellular IP - A New Approach to Internet Host Mobility. *ACM Computer Communication Review*, 29(1) :50–65, Janvier 1999.
- [VBG00] Nitin H. Vaidya, Paramvir Bahl, et Seema Gupta. Distributed fair scheduling in a wireless LAN. Dans *Proceedings of the sixth annual international conference on Mobile computing and networking (MobiCom 2000)*, pages 167–178, Boston, Massachusetts, USA, Août 2000.
- [VL02a] Vladimir Vishnevsky et Andrey Lyakhov. 802.11 LANs : Saturation Throughput in the Presence of Noise. Dans E. et als Gregori, editor, *Proceedings of the Second International IFIP-TC6 Networking Conference (Networking)*, volume 2345 of *LNCS*, pages 1008–1019, Pise, Italie, 2002. Springer-Verlag.
- [VL02b] Vladimir Vishnevsky et Andrey Lyakhov. IEEE 802.11 Wireless LAN : Saturation Throughput Analysis with Seizing Effect Consideration. *Cluster Computing*, 5(2) :133–144, Avril 2002.

- [WB01] Yu Wang et Brahim Bensaou. Achieving fairness in IEEE 802.11 DFWMAC with Variable Packet Lengths. Dans *Proceedings of the IEEE Global Telecommunications Conference (GlobeCom 2001)*, San Antonio, Texas, USA, Novembre 2001.
- [WH01] Kui Wu et Janelle Harms. QoS Support in Mobile Ad Hoc Networks. *Crossing Boundaries – the GSA Journal of University of Alberta*, 1(1) :92–106, Novembre 2001.
- [WP00] Carl S. Wijting et Ramjee Prasad. Evaluation of mobile ad-hoc network techniques in a cellular network. Dans *Proceedings of the fifty-first IEEE Vehicular Technology Conference (VTC 2000-Spring)*, pages 1025–1029, Tokyo, Japon, Mai 2000.
- [Wro97] John Wroclawski. Specification of the Controlled-Load Network Element Service. Internet Request For Comments RFC 2211, Internet Engineering Task Force, Septembre 1997.
- [XG03] Qi Xue et Aura Ganz. Ad hoc QoS on-demand routing (AQOR) in mobile ad hoc networks. *Journal of Parallel and Distributed Computing – Special issue on Routing in mobile and wireless ad hoc networks*, 63(2) :154–165, Février 2003.
- [XSLC00] Hannan Xiao, Winston K.G. Seah, Anthony Lo, et Kee Chaing Chua. A Flexible Quality of Service Model for Mobile Ad Hoc Networks. Dans *Proceedings of the fifty-first IEEE Vehicular Technology Conference (VTC 2000-Spring)*, pages 445–449, Tokyo, Japon, Mai 2000.
- [XV02] Xue Xang et Nitin H. Vaidya. Priority Scheduling in Wireless Ad Hoc Networks. Dans *Proceedings of the Third ACM International Symposium on Mobile Ad Hoc Networking and Computing (MobiHoc 2002)*, Lausanne, Suisse, Juin 2002.
- [YCG⁺03] Lily L. Yang, S. Conner, Xingang Guo, M. Hazra, et J. Zhu. Common Wireless Ad Hoc Network Usage Scenarios. Internet Draft – draft-irtf-yang-ans-scenarios-00.txt, Octobre 2003.
- [ZC02] Chenxi Zhu et Scott Corson. QoS routing for mobile ad hoc networks. Dans *Proceedings of the Twenty-first Annual Joint Conference of the IEEE Computer and Communications Societies (IEEE Infocom)*, New York, USA, Juin 2002.
- [ZCYM04] Hao Zhu, Guohong Cao, Alyin Yener, et Allen D. Mathias. EDCF-DM : A Novel Enhanced Distributed Coordination Function for Wireless Ad Hoc Networks. Dans *Proceedings of the thirty-ninth IEEE International Conference on Communications (ICC)*, Paris, France, Juin 2004.