



Optimisation des techniques de codage pour les transmissions radio avec voie de retour

Moustapha El Aoun

► To cite this version:

Moustapha El Aoun. Optimisation des techniques de codage pour les transmissions radio avec voie de retour. Traitement du signal et de l'image [eess.SP]. Télécom Bretagne, Université de Bretagne Occidentale, 2012. Français. NNT: . tel-00733300

HAL Id: tel-00733300

<https://theses.hal.science/tel-00733300>

Submitted on 18 Sep 2012

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.

Sous le sceau de l'Université européenne de Bretagne

Télécom Bretagne

En habilitation conjointe avec l'Université de Bretagne Occidentale

Ecole Doctorale – sicma

Optimisation des techniques de codage et de retransmission pour les systèmes radio avec voie de retour

Thèse de Doctorat

Mention : Sciences et technologies de l'information et de la communication

Présentée par **Moustapha EL AOUN**

Département : Signal et communications

Laboratoire : Lab-STICC Pôle : CACS/COM

Directeur de thèse : Xavier Lagrange

Soutenue le 12 juillet 2012

Jury :

M. Gilles Burel, Professeur, Université de Bretagne Occidentale (Président)
M. Charly Poulliat, Professeur, INP – ENSEEIHT Toulouse (Rapporteur)
M. Tijani Chahed, Professeur, Telecom SudParis (Rapporteur)
M. Jean-François Héland, Professeur, INSA Rennes (Examineur)
M. Xavier Lagrange, Professeur, Telecom Bretagne (Directeur de thèse)
M. Raphaël Le Bidan, Maître de conférences, Telecom Bretagne (Examineur)
M. Ramesh Pyndiah, Professeur, Telecom Bretagne (Invité)

Remerciements

J'adresse tout d'abord mes remerciements et toute ma reconnaissance à mes encadrants M. Raphaël Le Bidan et M. Xavier Lagrange qui m'ont guidé pour réaliser ce travail de thèse dans les meilleures conditions et qui m'ont énormément appris. J'en suis très reconnaissant. Je les remercie pour leurs patience, conseils, et disponibilité.

Je tiens également à remercier M. Ramesh Pyndiah chef du département signal et communications de Telecom Bretagne qui a suivi ce travail de très près. J'ai pleinement profité de ses orientations et conseils qui m'ont éclairé le chemin pour mener ce travail.

Je remercie le président de mon jury M. Gilles Burel et les rapporteurs M. Charly Poulliat et M. Tijani Chahed d'avoir pris le temps de lire mon manuscrit, de l'évaluer, et de me faire partager leurs remarques pertinentes. Je tiens également à remercier M. Jean-François Héliard d'avoir accepté d'être examinateur à ma soutenance de thèse.

Je tiens à remercier les permanents, les doctorants et les postdoctorants du département Signal et Communications pour l'ambiance agréable de travail qu'ils ont su m'offrir durant ces trois années de thèse.

Je remercie également ma famille, mes parents, frères et sœurs, pour le soutien et les encouragements qu'ils m'ont apporté pendant toutes ces années.

Table des matières

Acronymes	xvi
Introduction	1
1 Modèle de la transmission	5
1.1 Introduction	5
1.2 Architecture générale du système	5
1.3 Principales fonctions de la couche physique	6
1.3.1 Codage et Modulation	7
1.3.2 Modèles de canaux	8
1.3.3 Démodulation et décodage	9
1.4 Modèle de codes correcteurs d'erreurs	9
1.4.1 Codes gaussiens de longueur infinie	9
1.4.2 Codes optimaux de longueur finie	10
1.4.3 Codes convolutifs et modulation binaire	10
1.5 Critères des performances	11
1.5.1 Débit utile moyen	11
1.5.2 Efficacité énergétique	12
1.6 Conclusion	12
2 Protocoles de retransmission élémentaires	15
2.1 Introduction	15
2.2 Présentation des principaux protocoles	15
2.2.1 <i>Send and wait</i>	16
2.2.2 <i>Go-Back-N</i>	17
2.2.3 <i>Selective Reject</i>	19

2.2.4	<i>Send and wait</i> parallèle	19
2.2.5	Métriques de performances	20
2.3	Calcul et optimisation du débit utile moyen	21
2.3.1	Hypothèses et modèle temporel de la transmission	21
2.3.2	Débit utile moyen en SaW	23
2.3.3	Débit utile moyen en GBN	25
2.3.4	Débit utile moyen en SR	26
2.3.5	Débit utile moyen en SWP	26
2.3.6	Comparaison des différents protocoles	27
2.3.7	Optimisation du débit utile moyen	28
2.4	Étude et optimisation de l'efficacité énergétique	31
2.4.1	Définition et calcul de l'efficacité énergétique	31
2.4.2	Optimisation de l'efficacité énergétique	33
2.5	Extension à une voie de retour imparfaite	36
2.5.1	Modèle et impact des erreurs sur la voie de retour	36
2.5.2	Répartition optimale de l'énergie entre la voie directe et la voie de retour	37
2.6	Conclusion	39
3	Combinaison du codage et de la retransmission (HARQ)	43
3.1	Introduction	43
3.2	Principe de l'ARQ hybride	44
3.2.1	Mise en œuvre de l'HARQ type I	44
3.2.2	Mise en œuvre du schéma HARQ type II	45
3.3	Décodage optimal en HARQ-II	46
3.3.1	Transmissions multiples d'un même mot de code	46
3.3.2	Transmission de paquets de redondance distincts	48
3.4	Calcul et optimisation du débit utile moyen en HARQ	49
3.4.1	Calcul du débit utile moyen en HARQ-I	50
3.4.2	Calcul du débit utile moyen en HARQ-II	50
3.4.3	Performances limites en HARQ	52
3.4.4	Performances à longueur finie	57
3.4.5	Optimisation du débit utile moyen	61
3.5	Efficacité énergétique	64
3.5.1	Expression générale de l'efficacité énergétique	64

3.5.2	Efficacité énergétique en HARQ-I	64
3.5.3	Efficacité énergétique en HARQ-II	65
3.6	Extension à une voie de retour imparfaite	66
3.6.1	Prise en compte des pertes d'acquittements dans l'analyse	68
3.6.2	Influence de l'imperfection de la voie de retour sur le débit en HARQ	69
3.6.3	Répartition optimale de la puissance entre voie directe et voie de retour	73
3.7	Conclusion	74
4	Schéma HARQ avec redondance multipaquets	77
4.1	Introduction	77
4.2	Principe de la stratégie multi-paquets	78
4.2.1	Rappel du principe des schémas HARQ-SP	78
4.2.2	Principe des schémas HARQ multi-paquets	78
4.3	Étude des performances limites de l'approche multi-paquets	80
4.3.1	Expression générale du débit utile moyen en HARQ-MP	80
4.3.2	Application aux codes gaussiens asymptotiquement longs	81
4.3.3	Étude comparative	84
4.4	Schémas de codage pratiques pour le MP	85
4.4.1	Principe général de la construction proposée	85
4.4.2	Exemples de construction	87
4.4.3	Performances des schémas pratiques	90
4.5	Implémentation pratique et optimisation du protocole MP en présence d'erreurs sur la voie de retour	96
4.5.1	Protocole MP avec un acquittement pour L paquets au premier round	96
4.5.2	Protocole MP avec un acquittement par paquet au premier round	97
4.5.3	Performance en présence de pertes sur la voie de retour	97
4.6	Conclusion	102
5	Schéma HARQ hybride SP et MP	105
5.1	Introduction	105
5.2	Principe du protocole Hybride SP et MP	105
5.2.1	Fonctionnement de l'émetteur	107
5.2.2	Fonctionnement du récepteur	108

5.3	Analyse des performances du protocole hybride SP/MP	109
5.3.1	Modèle d'analyse	110
5.3.2	Application au calcul du débit utile	112
5.3.3	Performances du protocole hybride H-SP-MP	117
5.4	Exploitation de la connaissance du canal en réception	121
5.4.1	Présence d'une connaissance sur le RSB moyen	122
5.4.2	Présence d'une connaissance sur le RSB instantané	123
5.4.3	Performances de ces deux approches	123
5.5	Efficacité énergétique	125
5.6	Conclusion	126
Conclusion		130
A Schéma HARQ hybride SP et MP		131
A.1	Expressions des probabilités $P^{xx}(m)$ et $Q_l^{xx}(m)$	131
A.2	Probabilités de transition	132
A.2.1	Transitions à partir de l'état initial	132
A.2.2	Transitions en mode SP	132
A.2.3	Transitions en mode MP	133
A.2.4	Transitions en mode indéterminé	134
A.3	Performances du H-SP-MP avec connaissance du canal	135
Bibliographie		137

Liste des figures

1.1	Architecture générale du modèle de transmission	6
1.2	Traitements mis en œuvre par la couche PHY. Noter le codage séparé de l'entête, puis le multiplexage avec les données codées, avant l'opération de modulation proprement dite	7
1.3	Taux d'erreur paquet (borne et simulation) en fonction du RSB $\gamma = E_s/N_o$ pour le code convolutif de rendement 1/2 et de polynôme générateur (23,35), dans le cas d'une transmission sur le canal BABG réel, avec une taille de paquet de données $D = 256$ bits. Sont également présentées les performances du code optimal de même taille et même rendement, ainsi que les performances des codes aléatoires gaussiens de rendement 1/2.	12
2.1	Schéma illustrant le modèle de transmission entre les stations A et B (données de A vers B et acquittements de B vers A)	16
2.2	Protocole <i>Send and Wait</i> ($V(S)$: numéro de séquence du prochain paquet à transmettre. $V(R)$: numéro de séquence du paquet attendu). . .	18
2.3	Protocole à fenêtre d'anticipation de taille $w = 3$ paquets en émission et 1 paquet en réception.	18
2.4	Protocole Go-Back-N avec une fenêtre d'anticipation de taille 3	18
2.5	Protocole SR avec une fenêtre d'anticipation de taille 3	19
2.6	Protocole ARQ <i>Send and Wait</i> parallèle avec un nombre de processus SaW parallèles $w = 3$	20
2.7	Format de paquet transmis composé des données utiles de taille D bits et de l'entête de taille H bits.	21
2.8	Schéma explicatif montrant les durées d'émission de paquets et d'acquittements, les durées de traitement et le délai de propagation.	23
2.9	Débit utile normalisé en fonction du RSB en SaW, GBN, SR et PSW pour $D = 1000$ bits, $H = 40$ bits et $\tau = 1/2$	28
2.10	Débit utile normalisé en fonction du RSB en SaW, GBN, SR et PSW pour $D = 1000$ bits, $H = 40$ bits et $\tau = 3/2$	29

2.11	Schéma explicatif du cas de transmission avec un délai de propagation $T_p = \frac{T_e}{2}$ ($\tau = 0,5$).	29
2.12	Schéma explicatif du cas de transmission avec un délai de propagation $T_p = \frac{3}{2}T_e$ ($\tau = 1,5$).	29
2.13	Débit utile normalisé en fonction de la taille D des données utiles en SaW, GBN, SR et PSW pour $\gamma = 7$ dB, $H = 40$ bits et $\tau_o = 520$	30
2.14	Débit utile normalisé en fonction de la taille D des données utiles en SaW, GBN, SR et PSW pour $\gamma = 7$ dB et $H = 40$ bits et $\tau_o = 1560$. . .	31
2.15	Énergie ζ nécessaire par bit utile en fonction du rapport ($\gamma = E_s/N_o$) en dB pour une taille de données $D = 1000$ bits $H = 40$ bits pour différentes valeurs du temps de propagation ($T_p = 100T_s$, $T_p = 500T_s$, $T_p = 1000T_s$)	34
2.16	Energie ζ nécessaire par bit utile en fonction de la taille D d'un paquet de données, pour $\gamma = 7$ dB, et pour différentes valeurs du temps de propagation $\tau_o = 100$ et $\tau_o = 500$)	34
2.17	Allure de la fonction $\exp(-\zeta_{sw})$ en fonction du rapport $\gamma = E_s/N_o$ et de la taille D des données utiles pour $H = 40$ bits.	35
2.18	Débit utile normalisé en fonction du rapport signal à bruit γ en ARQ pour $D = 1000$ bits, $H = 40$ bits, $\tau = 0,5$ et $\epsilon = 0,1$ et $0,3$	38
2.19	Débit utile normalisé en fonction de α en ARQ pour $D = 1000$ bits, $H = 40$ bits, $\tau = 0,5$ et $E_g/N_o = 7$ dB.	40
3.1	Construction de paquets retransmis en IR par poinçonnage d'un code mère à faible rendement	46
3.2	Débit utile moyen simulé en HARQ-II avec codes convolutifs de polynôme générateur (23 35) pour $D = 256$, $H = 40$, $R_o = 0,46$ et $M_1 = 3$	52
3.3	Débit utile en HARQ-I, HARQ-II-CC et HARQ-II-IR pour une transmission sur canal BABG complexe avec codes gaussiens de longueur infinie et rendement $R_o = 0,96$, et avec $M_1 = 3$	55
3.4	Débit utile en HARQ-I, HARQ-II-CC et HARQ-II-IR pour une transmission sur canal de Rayleigh à évanouissements par blocs avec codes gaussiens de longueur infinie et rendement $R_o = 0,96$, et avec $M_1 = 3$	56
3.5	Débit utile en HARQ-I, HARQ-II-CC et HARQ-II-IR pour une transmission sur canal BABG complexe avec codes optimaux de longueur finie, pour $D = 512$, $H = 40$, $R_o = 0,96$ et $M_1 = 3$	59
3.6	Débit utile en HARQ-I et HARQ-II pour une transmission binaire avec codage convolutif de rendement $1/2$, $D = 512$ bits, $H = 40$ bits ($R_o = 0,48$ bits/symbole), et deux rounds au maximum ($M_1 = 1$). Les performances en ARQ-SR et HARQ-II + codes optimaux de mêmes paramètres sont également présentées. Courbes en traits pleins : débit théorique, courbes avec marqueurs : débit simulé.	62

3.7	Débit utile en HARQ-II-CC avec codes convolutifs de polynômes générateurs (23,35) en fonction de la taille de données utiles pour une transmission sur canal BABG réel avec $H = 40$ bits, $M_1 = 1$	63
3.8	Débit utile HARQ-II-CC avec codes optimaux en fonction du rendement R_o pour $D = 512$ bits, $H = 40$ bits et $M_1 = 3$	63
3.9	Energie moyenne nécessaire par bit utile en fonction du RSB $\gamma = E_s/N_o$ pour une transmission binaire sur canal BABG réel, en ARQ-SR et HARQ-I non tronqués avec codes convolutifs de polynômes générateurs (23 35) $D = 512$ bits et $H = 40$ bits.	66
3.10	Energie nécessaire ζ par bit utile en fonction du rapport signal à bruit $\gamma = E_s/N_o$ en HARQ-I, HARQ-II-CC et HARQ-II-IR avec des codes convolutifs et codes optimaux de longueur finie, pour une transmission binaire sur canal BABG réel avec $R_o = 0,48$ bits/symbole, $M_1 = 1$, $D = 512$ bits et $H = 40$ bits	67
3.11	Energie nécessaire par bit utile ζ en fonction du rapport signal à bruit $\gamma = E_s/N_o$ en HARQ-II-CC et HARQ-II-IR avec des codes en blocs optimaux de longueur finie et de rendement $R_o = 0,48$ bits/symbole, avec $D = 512$ bits utiles, $H = 40$ bits, et pour différentes valeurs du nombre maximum de rounds	67
3.12	Débit utile en HARQ-I, HARQ-II-CC et HARQ-II-IR pour une transmission sur canal BABG complexe avec des codes gaussiens aléatoires de longueur infinie et rendement $R_o = 0,96$ bits/symbole, $M_1 = 3$ et une probabilité de perte $\epsilon = 0,1$ sur la voie retour	70
3.13	Débit utile en HARQ-II-IR pour une transmission sur canal BABG complexe avec des codes gaussiens aléatoires de longueur infinie et rendement $R_o = 0,96$ bits/symbole, $M_1 = 3$ et pour différentes probabilités de perte sur la voie retour ($\epsilon = 0,01$, $\epsilon = 0,1$ et $\epsilon = 0,3$)	70
3.14	Débit utile en HARQ-I, HARQ-II-CC et HARQ-II-IR pour une transmission sur canal BABG complexe avec des codes optimaux de longueur finie et rendement $R_o = 0,96$ bits/symbole, $M_1 = 3$, $D = 512$ bits, $H = 40$ bits et une probabilité de perte $\epsilon = 0,1$ sur la voie retour . . .	71
3.15	Débit utile en HARQ-II-IR pour une transmission sur canal BABG complexe avec des codes optimaux de longueur finie et rendement $R_o = 0,96$ bits/symbole, $M_1 = 3$, $D = 512$ bits, $H = 40$ bits, et pour différentes probabilités de perte sur la voie retour ($\epsilon = 0,01$, $\epsilon = 0,1$ et $\epsilon = 0,3$) .	72
3.16	Débit utile en HARQ-I, HARQ-II-CC et HARQ-II-IR pour une transmission sur canal de Rayleigh à évanouissements par blocs avec des codes gaussiens aléatoires de longueur infinie et rendement $R_o = 0,96$ bits/symbole, $M_1 = 3$ et une probabilité de perte $\epsilon = 0,1$ sur la voie retour	72

3.17	Débit utile en HARQ-II-IR pour une transmission sur canal de Rayleigh à évanouissements par blocs avec des codes gaussiens aléatoires de longueur infinie et rendement $R_o = 0,96$ bits/symbole, $M_1 = 3$ et pour différentes probabilités de perte sur la voie retour ($\epsilon = 0,01$, $\epsilon = 0,1$ et $\epsilon = 0,3$)	73
3.18	Débit utile pour une transmission HARQ-II-IR avec codes optimaux de longueur finie sur canal BABG réel en fonction de α , pour $D = 512$ bits, $H = 40$ bits, $R_o = 0,96$ bits/symbole, $M_1 = 3$, et $E_g/N_o = 1$ dB	75
4.1	Principe d'une transmission utilisant le protocole HARQ-MP	79
4.2	Débit utile en HARQ-SP et HARQ-MP pour une transmission sur canal BABG complexe avec codes gaussiens de longueur infinie et rendement $R_o = 1$ bit/symbole, avec $M_1 = 2$, $M_2 = 2$ et $L = 4$	85
4.3	Débit utile en HARQ-SP et HARQ-MP pour une transmission sur canal de Rayleigh à évanouissements par blocs avec codes gaussiens de longueur infinie et rendement $R_o = 1$ bit/symbole, avec $M_1 = 2$, $M_2 = 2$ et pour différentes valeurs du nombre L de paquets de données considérés en MP.	86
4.4	Schéma de principe de la construction du code MP	86
4.5	Schéma HARQ-SP basé sur une concaténation parallèle de codes convolutifs (SP-PCCC)	88
4.6	Schéma HARQ-MP basé sur une concaténation parallèle de codes convolutifs (MP-PCCC)	89
4.7	Schéma HARQ-MP basé sur une concaténation série d'un code convolutif et d'un accumulateur (MP-SCCA)	90
4.8	Illustration du processus de retransmission dans le cas du schéma HARQ-MP-TPP (concaténation série d'un code convolutif et d'un code de parité, S_i/R_i : partie systématique/parité du paquet i)	91
4.9	Débit utile normalisé en fonction du SNR pour une transmission QPSK et un protocole HARQ-SP ou HARQ-MP sur un canal BABG complexe avec $L = 4$, $M_1 = 2$, $M_2 = 2$, $D = 1024$ bits, et $H = 40$ bits ($R_o = 0,98$ bits/symbole).	92
4.10	Taux d'erreur paquet pour différents schémas HARQ-SP, dans le cas d'une transmission QPSK sur un canal BABG complexe, avec $M_1 = 2$, $D = 1024$ bits, et $H = 40$ bits.	93
4.11	Taux d'erreur paquet pour différents schémas HARQ-MP, dans le cas d'une transmission QPSK sur un canal BABG complexe, avec $M_1 = 2$, $D = 1024$ bits, et $H = 40$ bits.	93
4.12	Débit utile normalisé pour différents schémas HARQ-SP et HARQ-MP dans le cas d'une transmission QPSK sur canal de Rayleigh à évanouissements par blocs, avec $L = 4$, $M_1 = 2$, $M_2 = 2$, $D = 1024$ bits, et $H = 40$ bits ($R_o = 0,98$ bits/symbole)	94

4.13	Taux d'erreur paquet pour différents schémas HARQ-SP, dans le cas d'une transmission QPSK sur canal de Rayleigh à évanouissements par blocs, avec $M_1 = 2$, $D = 1024$ bits, et $H = 40$ bits	94
4.14	Taux d'erreur paquet pour différents schémas HARQ-MP, dans le cas d'une transmission QPSK sur canal de Rayleigh à évanouissements par blocs, avec $M_1 = 2$, $D = 1024$ bits, et $H = 40$ bits	95
4.15	Automate à l'émission de la variante <i>MP-Acq. par paquet</i> ($V(A)$ variable interne contenant le numéro de séquence du dernier paquet positivement acquitté, $N(S)$ numéro de séquence du paquet en cours, $V(S)$ numéro de séquence du prochain paquet à envoyer, $N(R)$ numéro de séquence porté par l'acquittement reçu)	98
4.16	Automate en réception de la variante <i>MP-Acq. par paquet</i> ($V(R)$ variable interne contenant le numéro de séquence porté par l'acquittement transmis, $N(S)$ numéro de séquence du paquet reçu, $N(R)$ numéro de séquence porté par l'acquittement envoyé à la station A)	98
4.17	Débit utile moyen en HARQ-SP et HARQ-MP (deux variantes), pour une transmission sur canal BABG complexe avec des codes gaussiens aléatoires de longueur infinie et rendement $R_o = 1$, $L = 4$, $M_1 = 2$, $M_2 = 2$, $R_o = 1$ et $\epsilon = 0,1$	100
4.18	Débit utile moyen en HARQ SP et HARQ-MP (deux variantes), pour une transmission sur canal BABG complexe avec des codes gaussiens aléatoires de longueur infinie et rendement $R_o = 1$, $L = 4$, $M_1 = 2$, $M_2 = 2$, $R_o = 1$ et $\epsilon = 0,3$	101
4.19	Débit utile moyen en HARQ-SP et HARQ-MP (deux variantes), pour une transmission sur canal de Rayleigh à évanouissements par blocs avec des codes gaussiens aléatoires de longueur infinie et rendement $R_o = 1$, $L = 4$, $M_1 = 2$, $M_2 = 2$ et $\epsilon = 0,1$	101
4.20	Débit utile moyen en HARQ-SP et HARQ-MP (deux variantes), pour une transmission sur canal de Rayleigh à évanouissements par blocs avec des codes gaussiens aléatoires de longueur infinie et rendement $R_o = 1$, $L = 4$, $M_1 = 2$, $M_2 = 2$ et $\epsilon = 0,3$	102
5.1	Automate à l'émission du protocole H-SP-MP	107
5.2	Automate en réception du protocole H-SP-MP	109
5.3	Structure générale de la chaîne de Markov modélisant le protocole hybride H-SP-MP avec un nombre maximum de retransmissions fixé à M_1 en mode SP, à M_2 en mode MP, et un mode MP opérant sur des groupes de L paquets de données.	113
5.4	Figure complémentaire de la structure générale de chaîne de Markov de la figure 5.3 représentant les états du mode indéterminé du protocole hybride SP et MP avec paramètres M_1 , M_2 , L . L'état $(0, j, 0, j)$ vérifie $0 < j < L$	114

5.5	Chaîne de Markov représentant un schéma H-SP-MP pour $L = 3$, $M_1 = 2$ et $M_2 = 2$	115
5.6	Représentation d'une section de la chaîne de Markov en H-SP-MP. Les valeurs entourées des cercles pointillés représentent le nombre de paquets bien décodés relatifs à une transition particulière.	116
5.7	Nombre de paquets de données d_{i1} bien décodés associé aux transitions allant d'un état $s_i = (a, b, c, d)$ vers l'état initial $s_1 = (0, 0, 0, 0)$	116
5.8	Nombre de paquets bien décodés associé à la transition entre un état s_i et un état s_j tous deux différents de l'état initial	116
5.9	Débit utile en SP, MP et H-SP-MP pour une transmission sur canal BABG complexe avec codes gaussiens de longueur infinie et rendement $R_o = 1$ bit/symbole, avec $M_1 = 2$, $M_2 = 2$, $L = 3$ et $\epsilon = 0, 1$	117
5.10	Débit utile en SP, MP et H-SP-MP pour une transmission sur canal BABG complexe avec codes gaussiens de longueur infinie et rendement $R_o = 1$ bit/symbole, $M_1 = 2$, $M_2 = 2$, $L = 3$ et $\epsilon = 0, 3$	119
5.11	Débit utile en SP, MP et H-SP-MP pour une transmission sur canal BABG complexe avec codes gaussiens de longueur infinie et rendement $R_o = 1$ bit/symbole, $M_1 = 2$, $M_2 = 2$, $L = 3$ et $\epsilon = 0, 5$	120
5.12	Débit utile en SP, MP et H-SP-MP pour une transmission sur canal de Rayleigh à évanouissements par blocs avec codes gaussiens de longueur infinie et rendement $R_o = 1$ bit/symbole, $M_1 = 2$, $M_2 = 2$, $L = 3$ et $\epsilon = 0, 1$	120
5.13	Débit utile en SP, MP et H-SP-MP pour une transmission sur canal de Rayleigh à évanouissements par blocs avec codes gaussiens de longueur infinie et rendement $R_o = 1$ bit/symbole, $M_1 = 2$, $M_2 = 2$, $L = 3$ et $\epsilon = 0, 3$	121
5.14	Débit utile en SP, MP et H-SP-MP pour une transmission sur canal de Rayleigh à évanouissements par blocs avec codes gaussiens de longueur infinie et rendement $R_o = 1$ bit/symbole, $M_1 = 2$, $M_2 = 2$, $L = 3$ et $\epsilon = 0, 5$	122
5.15	Débit utile en SP, MP et H-SP-MP (aveugle/avec connaissance du RSB moyen) pour une transmission sur canal de Rayleigh à évanouissements par blocs avec codes gaussiens de longueur infinie et rendement $R_o = 1$ bit/symbole, $M_1 = 2$, $M_2 = 2$, $L = 3$, $\gamma_{th} = -1$ dB et $\epsilon = 0, 1$	123
5.16	Débit utile en SP, MP et H-SP-MP (aveugle/avec connaissance du RSB instantané) pour une transmission sur canal de Rayleigh à évanouissements par blocs avec codes gaussiens de longueur infinie et rendement $R_o = 1$ bit/symbole, $M_1 = 2$, $M_2 = 2$, $L = 3$, $\gamma_{th} = -1$ dB et $\epsilon = 0, 1$	124

5.17	Débit utile en HARQ-SP et H-SP-MP avec des codes gaussiens en fonction du coefficient de répartition α , pour une transmission sur un canal de Rayleigh à évanouissements par blocs sans codage d'acquittements et avec un RSB global moyen constant ($E_g/N_o = 5$ dB) et un rapport $1 - \rho = 0,001$ entre la taille d'un acquittement et la taille d'un paquet codé	126
R	Schéma coopératif avec présence d'un relais.	130

Acronymes

ACK	Positive Acknowledgement
ARQ	Automatic Repeat reQuest
BABG	Bruit Additif Blanc Gaussien
BPSK	Binary Phase Shift Keying
BSC	binary symmetric channel
CRC	Cyclic Redundancy Check
CSI	Channel State Information
FEC	Forward Error Correction
GBN	Go-back-N
HARQ	Hybrid Automatic Repeat reQuest
HARQ-I	HARQ type I
HARQ-II	HARQ type II
HARQ-II-CC	HARQ-II-Chase Combining
HARQ-II-IR	HARQ-II-Incremental Redundancy
HARQ-SP	HARQ simple paquet
HARQ-MP	HARQ multi-paquets
HARQ-H-SP-MP	HARQ hybride SP et MP
i.i.d.	indépendants et identiquement distribués
MAP	Maximum <i>a posteriori</i>
MCS	Modulation and Coding Scheme
MP-PCCC	MP avec concaténation parallèle de codes convolutifs
MP-SCCA	MP avec concaténation série d'un code convolutif et d'un accumulateur
MP-TPP	MP turbo produit paquet
MRC	Maximum Ratio Combining
MV	Maximum de Vraisemblance
NAK	Negative Acknowledgement
PCCC	Parallel Concatenated Convolutional Codes
PSW	Parallel Send and wait
RLC	Radio Link Control
RSB	Rapport Signal à Bruit
SaW	Send and Wait
SCCA	Serial Concatenation of Convolutional and Accumulator codes
SP-CC	SP avec codes convolutifs
SP-PCCC	SP avec concaténation parallèle de codes convolutifs

SR	Selective Reject
TEP	Taux d'Erreur Paquet
TPP	Turbo Produit Paquet

Introduction

Les systèmes de transmission radio avec voie de retour sont largement déployés actuellement. Dans ces systèmes, la transmission de données en mode paquet se fait généralement en combinant un protocole de retransmission (*Automatic Repeat reQuest* ou ARQ) avec un code correcteur d'erreur FEC (*Forward Error Correction*). Cette combinaison est appelée HARQ (*Hybrid ARQ*). Elle est implémentée dans de nombreuses normes de communication, en particulier pour les systèmes cellulaires de troisième et quatrième génération tel que l'UMTS (*Universal Mobile Telecommunication System*) et le LTE (*Long Term Evolution*). Dans ces systèmes, la voie de retour est utilisée pour informer l'émetteur de la bonne ou mauvaise réception des données transmises. Elle peut également être exploitée pour informer l'émetteur de l'état du canal afin de s'adapter à l'évolution des conditions de propagation.

Les performances de la combinaison ARQ/FEC sont directement liées au protocole de retransmission et au code FEC choisis. Dans la plupart des études présentées dans la littérature, l'optimisation des schémas HARQ porte essentiellement sur l'optimisation du schéma de codage de la couche physique (règles de poinçonnage). L'un des objectifs de cette thèse est de proposer une optimisation conjointe inter-couches des stratégies ARQ et FEC (*cross layer*) en considérant à la fois des aspects de communication numérique et des aspects protocolaires.

Ce mémoire de thèse est organisé de la manière suivante.

Le chapitre 1 présente l'architecture générale du système de transmission considéré tout au long de ce travail, les principales hypothèses que nous avons faites, ainsi que les différentes familles de codes et modèles de canaux considérés.

Alors que le chapitre 1 se concentre sur la couche physique, le chapitre 2 s'intéresse plus particulièrement aux fonctions de contrôle d'erreurs mises en œuvre par la couche liaison de données. Nous revisitons ainsi les protocoles de retransmission ARQ classiques, et comparons leurs performances aussi bien sous l'angle du débit utile que sous l'angle de l'efficacité énergétique.

Le chapitre 3 s'intéresse aux protocoles hybride HARQ qui combinent un code correcteur d'erreurs et un protocole de retransmission. Nous présentons et comparons les trois principales familles de protocoles hybrides (type-I, type-II avec combinaison de paquet, type-II avec redondance incrémentale) selon le critère du débit utile mais aussi de l'efficacité énergétique. L'une des contributions originales de cette thèse est d'avoir exploité des résultats récents de la théorie de l'information pour obtenir les

performances limites *non asymptotiques* des protocoles HARQ, lorsque l'on considère des codes optimaux de longueur finie.

Le chapitre 4 introduit la contribution centrale de la thèse. Dans les schémas HARQ classiques, dits *simple paquet* (HARQ-SP), les paquets de redondance transmis aux différents rounds sont tous relatifs au même paquet de données. Dans le cas où la taille des paquets retransmis est identique à la taille du paquet initial, chaque retransmission engendre une réduction importante du débit utile. D'autre part, lorsque le canal de retour n'est pas parfait, le débit utile est également pénalisé par les retransmissions inutiles effectuées par l'émetteur lorsque le paquet est bien reçu mais que l'acquittement positif est perdu. Pour remédier à ces problèmes, nous proposons donc une nouvelle approche dite *multi-paquets* (MP), qui vise à réduire le nombre moyen de retransmissions par paquet grâce à la construction de paquets de redondance pouvant aider au décodage simultané de plusieurs paquets de données. Les performances de cette approche sont présentées et comparées à celles des HARQ classiques. Les résultats présentés dans ce chapitre ont fait l'objet d'une publication en conférence [1].

L'étude menée au chapitre 4 montre que l'approche MP offre un gain significatif en débit par rapport aux protocoles HARQ classiques lorsque le rapport signal sur bruit est suffisamment élevé. Ces derniers demeurent toutefois plus performants à faible RSB. Ce constat nous conduit à proposer dans le chapitre 5 un protocole hybride qui combine les schémas classiques et l'approche multi-paquets dans le but de concevoir un schéma de retransmission plus performant sur toute la plage de RSB. Dans ce même chapitre, nous présentons également une méthode générale d'analyse des performances des schémas HARQ. Cette méthode se base sur les chaînes de Markov à temps discret et états finis. Les travaux présentés dans ce chapitre ont conduit à trois publications en conférence [2] [3] [4].

Enfin, le chapitre de conclusion résume les principaux résultats obtenus dans cette thèse, et présente quelques perspectives à ce travail.

Publications

- M. El Aoun, R. Le Bidan, X. Lagrange and R. Pyndiah, “Multiple-packet versus single-packet incremental redundancy strategies for type-II hybrid ARQ”, *International symposium on turbo codes and iterative information processing*, Sep. 2010.
- M. El Aoun, X. Lagrange, R. Le Bidan and R. Pyndiah, “Analysis and optimization of hybrid single packet and multiple-packets incremental redundancy in the presence of channel state information”, *14th International Symposium on Wireless Personal Multimedia Communications*, Oct. 2011, pp : 1–5.
- M. El Aoun, X. Lagrange, R. Le Bidan and R. Pyndiah, “Analyse des schémas HARQ classiques et évolués (schémas multi-paquets) en présence d’une voie de retour imparfaite”, *XXIIIe colloque GRETSI : traitement du signal et des images*, Sep. 2011
- M. El Aoun, X. Lagrange, R. Le Bidan and R. Pyndiah, “Throughput analysis of hybrid single-packet and multiple-packet truncated type-II HARQ strategies with unreliable feedback channel”, *IEEE Wireless Communications and Networking Conference*, Apr. 2012

CHAPITRE 1

Modèle de la transmission

1.1 Introduction

Ce chapitre détaille le modèle de transmission considéré tout au long de la thèse. Il s'agit d'une transmission point à point entre une station émettrice A et une station réceptrice B reliées par un canal radio duplex. On s'intéresse ici aux deux premières couches du modèle OSI, à savoir la couche physique (PHY) et la couche liaison de données (*Radio Link Control*, ou RLC).

Ce chapitre est organisé de la manière suivante. Dans un premier temps nous introduisons l'architecture générale du système de transmission. Nous nous détaillons ensuite les principales fonctions de la couche PHY, ainsi que les familles de codes et les modèles de canaux considérés dans ce travail. Les fonctions de la couche RLC relatives au contrôle d'erreurs seront abordées plus en détail dans le chapitre suivant. Nous présentons finalement les différents critères retenus pour analyser les performances du système.

1.2 Architecture générale du système

L'architecture générale du système de transmission considéré est présentée sur la figure 1.1. On considère une transmission point-à-point entre un émetteur (station A) et un récepteur (station B) sur un canal radio duplex bruité. La station A se contente de transmettre des paquets de données de taille constante à la station B. Cette dernière répond par des acquittements indiquant la bonne ou mauvaise réception des paquets. Les acquittements transmis sur le canal de retour (de B vers A) peuvent également être confrontés à des erreurs de transmission. On parlera alors de voie de retour imparfaite.

Au niveau de la couche RLC, les données issues des couches hautes sont découpées en blocs de taille identique, égale à D bits. Outre la segmentation, le rôle de la couche RLC est également de garantir l'intégrité des données vis-à-vis des couches supérieures, par un mécanisme de détection d'erreurs couplé à un protocole de retransmission. Pour cela, la couche RLC rajoute une somme de contrôle (*Cyclic Redundancy Check*, ou CRC) pour permettre la détection des erreurs au niveau de la station B. En règle générale, la taille du CRC est choisie suffisamment grande (typiquement 16 ou 32 bits) pour que

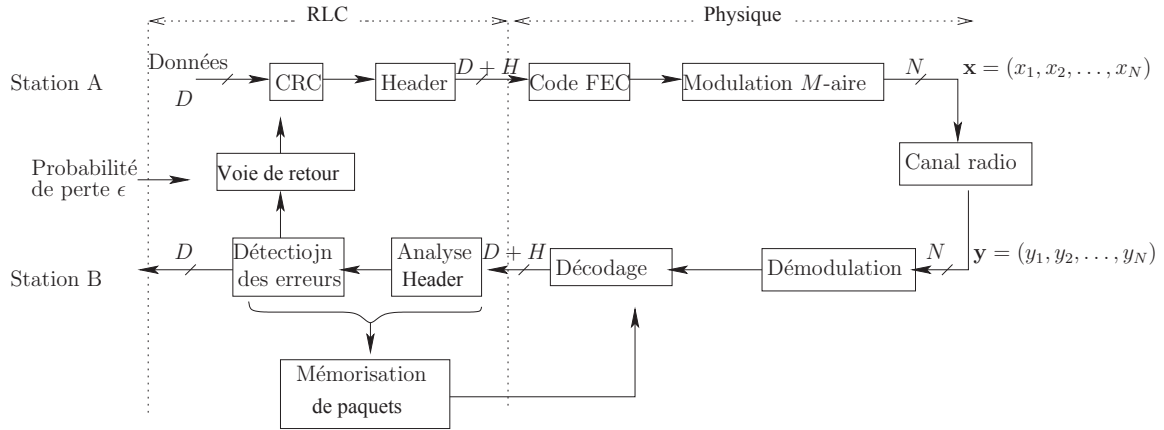


Figure 1.1 — Architecture générale du modèle de transmission

la probabilité de non détection des erreurs puisse être considérée comme négligeable. En plus de la somme de contrôle, la couche RLC rajoute également au bloc de données un entête permettant l'identification des paquets émis. Soit H la taille (en bits) de l'ensemble des champs rajoutés par la couche RLC (entête + CRC). Comme indiqué sur la figure 1.1, les paquets délivrés de la couche RLC vers la couche physique ont alors une taille de $D + H$ bits. La couche RLC est également responsable du mécanisme de retransmission. Elle doit pouvoir mémoriser chaque paquet de données transmis de A vers B (canal direct) jusqu'à sa bonne réception. La station B notifie la station A de la bonne ou mauvaise réception du paquet par l'envoi d'un acquittement sur la voie de retour. Dans cette thèse, pour simplifier, on considère que les acquittements ont la même taille que les entêtes (paquet de H bits). Suivant le protocole de retransmission considéré, la couche RLC de la station B doit pouvoir mémoriser, si nécessaire, tout ou partie des paquets reçus. Les protocoles de retransmission (règles de dialogue et format de trame) seront présentés en détail dans le chapitre suivant.

La couche physique est responsable de la transmission des trames RLC sur le milieu de propagation (canal radio). Pour cela, elle utilise un format de transmission approprié au contrainte du canal. Les principales fonctions et hypothèses sur la couche PHY sont détaillées dans la section suivante.

1.3 Principales fonctions de la couche physique

Nous présentons ici les principaux traitements mis en œuvre par la couche PHY en rapport avec la problématique de la thèse. Comme indiqué sur la figure 1.2, on se limite volontairement pour simplifier, aux fonctions de codage et modulation côté émetteur et aux fonctions duales côté récepteur.

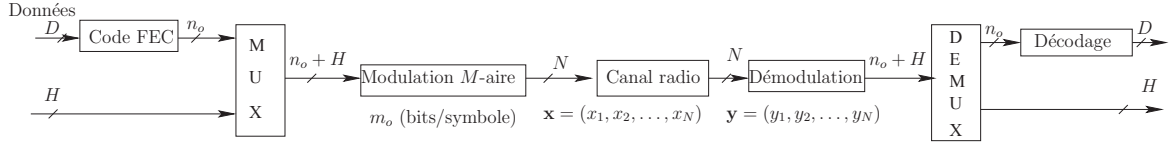


Figure 1.2 — Traitements mis en œuvre par la couche PHY. Noter le codage séparé de l'entête, puis le multiplexage avec les données codées, avant l'opération de modulation proprement dite

1.3.1 Codage et Modulation

Les données issues de la couche RLC sont encodées et modulées puis envoyées sur le canal. Le schéma de codage et modulation (*Modulation and Coding Scheme*, MCS) $\mathcal{C}$ est vu ici comme une fonction de codage qui associe un mot de code de N symboles modulés $\mathbf{x} = (x_1, x_2, \dots, x_N)$ à un bloc d'information de taille D bits utiles. Le taux de codage du MCS est défini comme le rapport :

$$R_o = \frac{D}{N} \text{ bits/symbole} \quad (1.1)$$

Bien que nous présentions l'opération de codage et modulation comme un seul traitement, en pratique comme indiqué sur la figure 1.2, ces deux fonctions sont habituellement réalisées de manière disjointe.

Dans cette thèse, on suppose que la station A démarre toujours la transmission sur le même MCS. Autrement dit, l'adaptation de lien (adaptation du MCS à l'évolution des conditions radio) n'a pas été prise en considération dans ce travail. Par ailleurs, on ne s'intéresse pas non plus à la façon dont l'entête est encodé. Comme indiqué sur la figure 1.2, on suppose que l'entête est encodé à part, et multiplexé avec le mot de code binaire produit par l'encodage FEC du paquet de données. Cette hypothèse nous permet de considérer que l'entête est toujours correctement décodé par la station B. Celle-ci est donc toujours capable d'identifier le numéro de la trame reçue, indépendamment du résultat du décodage du paquet de données.

On désigne par ρ le rapport entre la taille n_o du message binaire après encodage FEC et la taille $n_o + H$ du paquet global avant modulation (voir figure 1.2) :

$$\rho = \frac{n_o}{n_o + H} \quad (1.2)$$

En l'absence du codage FEC, ρ vaut tout simplement $D/(D + H)$.

Les symboles modulés $\{x_i\}$ prennent leurs valeurs dans un alphabet à M éléments réels ou complexes, M pouvant éventuellement être infini (codes gaussiens), et vérifiant $E[x_i] = 0$ et $E[|x_i|^2] = E_s$, où $E[\cdot]$ désigne l'espérance mathématique. Dans le cas le plus simple d'une modulation BPSK (*Binary Phase Shift Keying*) à $M = 2$ états, les symboles modulés $\{x_i\}$ appartiennent à l'ensemble fini $\{-\sqrt{E_s}, +\sqrt{E_s}\}$.

Le taux global de codage R_o du MCS peut s'exprimer en fonction de ρ , du rendement $r = D/n_o$ du code FEC binaire, et du nombre $m_o = \log_2(M)$ de bits par symbole dans

l'alphabet de modulation comme suit :

$$R_o = \frac{D}{N} = \frac{D}{n_o} \left(\frac{n_o}{n_o + H} \right) \left(\frac{n_o + H}{N} \right) = r \rho m_o \quad (1.3)$$

1.3.2 Modèles de canaux

Deux modèles de canaux ont été utilisés pour modéliser le lien direct (A vers B) dans nos études : le canal à bruit additif blanc gaussien (BABG) et le canal de Rayleigh à évanouissement par blocs. Le modèle de canal sur la voie de retour sera présenté dans les chapitres suivants.

Canal gaussien

Le canal à bruit additif blanc gaussien (BABG) est un modèle courant et très simple. En notant par $\mathbf{x} = (x_1, x_2, \dots, x_N)$ le vecteur du signal transmis et par $\mathbf{w} = (w_1, w_2, \dots, w_N)$ le vecteur du bruit rajouté par le canal, le vecteur d'observation $\mathbf{y} = (y_1, y_2, \dots, y_N)$ reçu par B s'écrit :

$$\mathbf{y} = \mathbf{x} + \mathbf{w} \quad (1.4)$$

où les échantillons de bruit w_i suivent une loi normale réelle $\mathcal{N}(0, N_o/2)$ de moyenne nulle et variance $N_o/2$ lorsque les symboles modulés $\{x_i\}$ sont réels, ou bien une loi gaussienne complexe circulaire symétrique $\mathcal{CN}(0, N_o)$ de moyenne nulle et variance N_o lorsque les symboles modulés $\{x_i\}$ sont complexes. Dans tous les cas, la probabilité conditionnelle est de la forme (à une constante près) :

$$\Pr\{\mathbf{y}|\mathbf{x}\} \propto \exp\left(-\frac{\|\mathbf{y} - \mathbf{x}\|^2}{N_o}\right) \quad (1.5)$$

Que ce soit pour un canal réel ou complexe, le rapport signal sur bruit sera défini comme $\gamma = E_s/N_o$. Le canal BABG est alors caractérisé par un RSB constant.

Canal de Rayleigh à évanouissements par blocs

Sur le canal de Rayleigh à évanouissements par blocs, chaque paquet transmis observe une réalisation différente du canal. Le modèle mathématique de ce canal est donné par :

$$\mathbf{y} = h\mathbf{x} + \mathbf{w} \quad (1.6)$$

où h est le coefficient d'évanouissement affectant le paquet transmis, et $\mathbf{w}$ est le vecteur du bruit additif blanc gaussien complexe. Les coefficients du canal h sont supposés indépendants, et identiquement distribués d'une transmission à l'autre, suivant une loi gaussienne complexe de moyenne nulle et de variance unitaire ($h \sim \mathcal{CN}(0, 1)$). Le RSB instantané (par paquet) $\gamma_{inst} = |h|^2\gamma$ est une variable aléatoire qui suit alors une loi exponentielle de moyenne γ . La probabilité conditionnelle $\Pr\{\mathbf{y}|\mathbf{x}, h\}$ est de la forme :

$$\Pr\{\mathbf{y}|\mathbf{x}, h\} \propto \exp\left(-\frac{\|\mathbf{y} - h\mathbf{x}\|^2}{N_o}\right) \quad (1.7)$$

1.3.3 Démodulation et décodage

Le rôle de la démodulation et du décodage est de prendre une décision sur le mot de code transmis $\mathbf{x}$, à partir de l'observation $\mathbf{y}$ reçue. La règle de décodage optimale est la règle du maximum de vraisemblance, qui consiste à rechercher le mot de code $\hat{\mathbf{x}}$ qui maximise la probabilité conditionnelle $\Pr\{\mathbf{y}|\mathbf{x}, h\}$:

$$\hat{\mathbf{x}} = \arg \max_{\mathbf{x} \in \mathcal{C}} \Pr\{\mathbf{y}|\mathbf{x}, h\} \quad (1.8)$$

$$= \arg \min_{\mathbf{x} \in \mathcal{C}} \|\mathbf{y} - h\mathbf{x}\|^2 \quad (1.9)$$

$$= \arg \max_{\mathbf{x} \in \mathcal{C}} \Re\{h^* \mathbf{y}, \mathbf{x}\} - \frac{|h|^2}{2} \|\mathbf{x}\|^2 \quad (1.10)$$

où $(.)^*$ est le complexe conjugué. On voit ici que la quantité $h^* \mathbf{y}$ constitue une statistique suffisante pour le décodage optimal.

1.4 Modèle de codes correcteurs d'erreurs

Dans cette section, nous introduisons les trois grandes familles de schémas de modulation et codage considérés dans ce travail, et nous donnons l'expression générale de leurs performances sur les différents canaux, mesurées par la probabilité d'erreur paquet après décodage $p_o = \Pr\{\hat{\mathbf{x}} \neq \mathbf{x}\}$.

1.4.1 Codes gaussiens de longueur infinie

Les codes gaussiens sont des codes essentiellement d'intérêt théorique. Ils sont utilisés pour établir les performances limites que l'on peut espérer atteindre asymptotiquement avec des codes pratiques. Les symboles $\{x_i\}$ de chaque mot de code $\mathbf{x} = (x_1, x_2, \dots, x_N)$ sont tirés aléatoirement suivant une loi gaussienne complexe $\mathcal{CN}(0, E_s)$. Dans la limite de très grande taille de blocs ($N \rightarrow \infty$), on montre que cette famille de codes atteint la capacité du canal BABG [5]. La probabilité d'échec de décodage sur le canal BABG (avec $R_o = \log_2(|\mathcal{C}|)/N$ constant pour $N \rightarrow \infty$) s'écrit :

$$p_o = \begin{cases} 1 & R_o \geq I(\gamma) \\ 0 & R_o < I(\gamma) \end{cases} \quad (1.11)$$

où $I(\gamma) = \log_2(1 + \gamma)$ est l'information mutuelle moyenne en sortie du canal BABG (en bits par symbole), et correspond ici à la capacité de ce canal.

Sur le canal de Rayleigh, la probabilité p_o correspond à la probabilité de coupure lorsque N tend vers l'infini. Elle est donnée par :

$$p_o = \Pr\{I(\gamma) \leq R_o\} \quad (1.12)$$

1.4.2 Codes optimaux de longueur finie

Le modèle des codes aléatoires gaussiens permet d'établir les performances optimales d'un système, dans la limite de très grande taille de blocs. Or en pratique, les codes utilisés sont de longueur finie et présentent donc des performances nécessairement en retrait vis-à-vis des performances asymptotiques des codes aléatoires gaussiens. Polyanskiy *et al* ont récemment introduit dans [6] et [7] une très bonne estimation de la probabilité d'erreur minimale atteignable par le meilleur code (dit code *optimal*) de longueur finie. Leur approche ne fait aucune hypothèse quant à la structure (alphabet de modulation notamment) du code optimal. Plus précisément, ils ont obtenu une expression qui relie le taux de codage maximal R_o^* que l'on peut atteindre avec un code de longueur finie N et une probabilité d'erreur inférieure ou égale à p_o . Elle est donnée par [6] :

$$NR_o^* = NC - \sqrt{NV}Q^{-1}(p_o) + \alpha_o \log_2 N + O(1) \quad (1.13)$$

où $Q(x) = \frac{1}{\sqrt{2\pi}} \int_x^{+\infty} e^{-\frac{t^2}{2}} dt$, α_o est une constante réelle positive ou nulle. C est la capacité du canal et V est un paramètre qui dépend uniquement des caractéristiques du canal appelé dispersion du canal. On peut alors estimer la probabilité d'erreur minimale $p_o^*(N, R_o)$ que l'on peut atteindre avec un code de longueur N et de rendement R_o :

$$p_o^*(N, R_o) \gtrsim Q \left(\sqrt{\frac{N}{V}} \left(C - R_o + \alpha_o \frac{\log_2 N}{N} \right) \right) \quad (1.14)$$

Pour un canal de transmission BABG réel à entrée et sortie continues, les paramètres V et C sont définis comme suit [6] :

$$C(\gamma) = \frac{1}{2} \log_2(1 + 2\gamma) \quad (1.15)$$

$$V(\gamma) = \frac{\gamma}{2} \frac{\gamma + 1}{(\gamma + \frac{1}{2})^2} (\log_2(e))^2 \quad (1.16)$$

Notons qu'il n'existe pas à ce jour d'expression disponible pour le canal de Rayleigh à évanouissements par blocs. Des premiers résultats ont toutefois été récemment établis pour le canal de Rayleigh ergodique dans présentés en [8].

1.4.3 Codes convolutifs et modulation binaire

La dernière famille de codes considérés associe un code convolutif binaire et une modulation binaire de type BPSK. Le taux de codage de ce MCS vaut alors $R_o = r\rho$. Sur le canal BABG réel, à fort RSB, la probabilité d'erreur par paquet est bien approchée par la borne de l'union [9][10] :

$$p_o \leq D \sum_{d_o \geq d_{free}} \frac{A_{d_o}}{2} \operatorname{erfc} \sqrt{d_o \gamma} \quad (1.17)$$

où d_{free} est la distance libre du code, $\gamma = \frac{E_s}{N_0}$ est le rapport signal à bruit ($E_s = rE_b$, E_b est l'énergie d'un bit avant codage) et les couples (d_o, A_{d_o}) désignent respectivement le

m	Séquences génératrice	d_{free}	$A_{d_{free}+h}$ ($h = 0, 1, \dots, 6$)
1	(1 3)	3	1, 1, 1, 1, 1, 1, 1
2	(5 7)	5	1, 2, 4, 8, 16, 32, 64
3	(15 17)	6	1, 3, 5, 11, 25, 55, 121
4	(23 35)	7	2, 3, 4, 16, 37, 68, 176
5	(53 75)	8	1, 8, 7, 12, 48, 95, 281
6	(133 171)	10	11, 0, 38, 0, 193, 0, 1331
7	(247 371)	10	1, 6, 12, 26, 52, 132, 317
8	(561 753)	12	11, 0, 50, 0, 286, 0, 1630
9	(1131 1537)	12	1, 7, 19, 28, 69, 185, 411
10	(2473 3217)	14	14, 0, 92, 0, 426, 0, 2595
11	(4325 6747)	15	14, 21, 34, 101, 249, 597, 1373

Tableau 1.1 — Premières valeurs du paramètre A_{d_o} pour différents codes convolutifs de rendement $r = 1/2$

poids de Hamming et la multiplicité des chemins erronés dans le treillis, et sont donnés dans la table 1.1. Ce tableau présente les premiers paramètres A_{d_o} pour différents codes convolutifs optimaux en terme de distance et de rendement $r = 1/2$ [11] [12].

La figure 1.3 compare la probabilité d'erreur paquet donnée par (1.17), avec le taux d'erreur paquet résultant de la simulation pour le code convolutif de rendement $r = 1/2$ et de polynôme générateurs (23,35) en octal, pour une transmission sur un canal BABG réel avec des paquet de données de taille $D = 256$ bits. On voit que l'approximation (1.17) est très précise au delà de 0 dB. A titre d'information, nous avons également tracé les performances du code optimal de même rendement et de même dimension $D = 256$ bits, ainsi que les performances asymptotiques des codes gaussiens de longueur infinie et rendement $1/2$.

L'évaluation des performances du code convolutif sur le canal de Rayleigh à évanouissements par blocs est également possible, mais plus compliquée à mener car elle dépend notamment des réalisations du canal vues par chaque mot de code [13][14]. Ce point n'a pas été donc considéré dans cette thèse.

1.5 Critères des performances

Il existe de nombreux critères d'évaluation des performances d'un système de transmission. Dans cette thèse, nous nous limitons à deux critères, à savoir le débit utile moyen et l'efficacité énergétique

1.5.1 Débit utile moyen

Le débit utile moyen η est défini comme le rapport du nombre moyen de bits d'information délivrés avec succès au destinataire pendant un certain intervalle de temps, sur le nombre moyen total de symboles transmis durant ce même intervalle de temps.

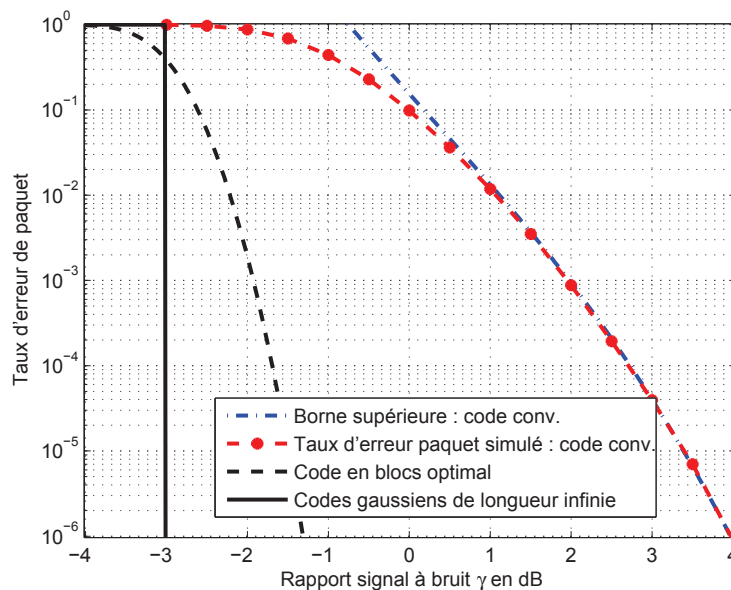


Figure 1.3 — Taux d'erreur paquet (borne et simulation) en fonction du RSB $\gamma = E_s/N_o$ pour le code convolutif de rendement $1/2$ et de polynôme générateur $(23,35)$, dans le cas d'une transmission sur le canal BABG réel, avec une taille de paquet de données $D = 256$ bits. Sont également présentées les performances du code optimal de même taille et même rendement, ainsi que les performances des codes aléatoires gaussiens de rendement $1/2$.

Défini ainsi, le débit utile moyen est homogène à un taux de codage.

1.5.2 Efficacité énergétique

Le second critère considéré, l'efficacité énergétique, mesure l'énergie moyenne ζ nécessaire pour transmettre correctement un bit de données au destinataire. En d'autres termes, c'est le rapport entre l'énergie consommée par un paquet correctement reçu (en tenant compte de l'énergie consommée par les retransmissions et les acquittements), et le nombre de bits utiles du paquet. Ce critère est très important pour les systèmes mettant en œuvre un protocole de retransmission des données erronées. En effet, la retransmission améliore la fiabilité du système en permettant la correction de configuration d'erreurs problématiques pour le code de canal, mais en même temps, augmente significativement l'énergie consommée. D'où l'intérêt de regarder le coût énergétique requis par bit utile correctement transmis.

1.6 Conclusion

Ce chapitre a présenté l'architecture générale du modèle de transmission qui constitue le point de départ de cette thèse. L'accent a été mis ici sur les principales fonctions de la couche physique. Le chapitre suivant développe les protocoles de retransmission

mis en œuvre au niveau de la couche liaison de données.

CHAPITRE 2

Protocoles de retransmission élémentaires

2.1 Introduction

Ce chapitre traite un cadre de transmission sans codage canal FEC. Il s'intéresse aux mécanismes de retransmission mis en œuvre par la couche liaison de données (couche RLC). Ces mécanismes sont appelés ARQ (*Automatic Repeat reQuest*). Ils sont largement utilisés pour améliorer la fiabilité des transferts de données sur la couche physique. Le principe consiste à rajouter de la redondance (CRC) à chaque paquet transmis de manière à ce que le destinataire puisse vérifier l'intégrité des paquets reçus. La détection d'un paquet erroné déclenche l'envoi d'une demande de retransmission à l'émetteur.

Ce chapitre est organisé comme suit. Dans un premier temps, nous introduisons les principaux protocoles ARQ en section 2.2. Nous développons ensuite une analyse unifiée de ces différents protocoles en terme de débit utile moyen et d'efficacité énergétique, en section 2.3 et 2.4. Les analyses précédentes sont finalement étendues au cas plus réaliste d'une transmission avec voie de retour imparfaite en section 2.5.

2.2 Présentation des principaux protocoles

Dans tout le document, nous considérons une transmission de données d'une station appelée A vers une station appelée B. La station A émet des données vers B et la station B émet des acquittements (positifs ou négatifs) vers A (cf. figure 2.1). Nous ne considérons pas le cas de transmissions simultanées de données de A vers B et de B vers A. En effet, la plupart des applications sont de type transactionnel avec des requêtes et des réponses à ces requêtes. Par exemple, un utilisateur consulte une page web en cliquant sur un lien puis la page est chargée sur son terminal. Les échanges entièrement duplex sont utiles lorsqu'on considère de nombreuses échanges différents multiplexés sur une même liaison de données (cas par exemple de liaisons entre routeurs). Dans le contexte de cette thèse, nous considérons des liaisons radios, soit dans un contexte cellulaire entre un terminal particulier et une station de base, soit dans un contexte

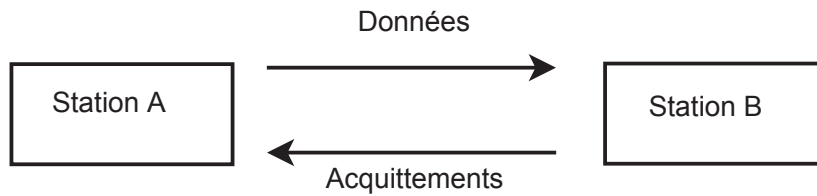


Figure 2.1 — Schéma illustrant le modèle de transmission entre les stations A et B (données de A vers B et acquittements de B vers A)

de réseau de capteurs entre deux nœuds. Dans ces 2 contextes, considérer un seul sens de transmission des données à un instant donné se justifie pleinement (la transmission ultérieure de données de B vers A est similaire à la transmission de données de A vers B).

Un protocole ARQ est défini par un format de trame ainsi que par des règles de dialogue [15]. Les règles de dialogue assurent une transmission fiable entre la station émettrice A et la station réceptrice B. En particulier, ces règles visent à garantir que le protocole fonctionne correctement (ne se bloque jamais). Les trames envoyées sont composées de blocs de données utiles et d'autres blocs nécessaires au bon fonctionnement du protocole. Ces derniers sont généralement appelés entête ou *Header*.

En particulier, dans un protocole ARQ, la couche RLC de la station A rajoute un champ appelé CRC (*Cyclic Redundancy Check*). Ce champ permet à la station B de détecter les erreurs et ensuite d'informer A de la bonne ou mauvaise réception de paquets transmis. Lorsque le paquet est bien reçu, B envoie un acquittement positif (ACK) pour dire à la station A de passer à la transmission du paquet suivant, ou bien un acquittement négatif (NAK) pour demander la retransmission du paquet mal reçu. La station A continue la retransmission d'un paquet jusqu'à sa bonne réception, ou bien jusqu'à ce que l'on atteigne le nombre maximum M_1 de retransmissions autorisées.

Il existe différentes variantes de protocoles ARQ dans la littérature. Dans cette thèse, on s'intéresse aux trois principales familles : les protocoles *Send and Wait* (SaW), *Go-back-N* (GBN) et *Selective Reject* (SR), ainsi qu'à une variante du SaW appelée *Send and wait* parallèle (PSW). Le principe ainsi que les avantages et les inconvénients de chacun de ces protocoles sont décrits dans les sous sections suivantes.

2.2.1 *Send and wait*

Le protocole *Send and Wait* (SaW) consiste à transmettre un paquet par la station A et attendre son acquittement avant de passer à la transmission du paquet suivant [16, 17]. Initialement, la station A transmet son premier paquet et attend la réponse de B. La station B vérifie la bonne réception du paquet et renvoie un acquittement. Si l'acquittement est positif, A passe au paquet suivant. Si l'acquittement est négatif, A, retransmet le paquet erroné et lorsqu'aucun acquittement n'est arrivé, par exemple suite à une perte d'acquittement sur la voie de retour, la station A retransmet le même paquet après expiration d'un temporisateur armé après émission de chaque paquet.

Dans ce protocole, la station A doit garder en mémoire le paquet transmis jusqu'à sa bonne réception par B. La station B rejette tout paquet erroné.

En ARQ, les paquets peuvent être retransmis plusieurs fois suite à un délai d'attente plus long que prévu ou à des pertes d'acquittements, ce qui peut causer en retour des problèmes de duplication de paquets. Pour remédier à ces problèmes, les protocoles ARQ numérotent les paquets transmis [15, 18]. On montre qu'en SaW, une numérotation sur un bit suffit. À chaque réception d'un paquet, B envoie un acquittement portant le numéro (0 ou 1) du prochain paquet attendu. La figure 2.2 illustre le principe du protocole SaW.

Le principal avantage du protocole SaW est la simplicité des règles d'émission et de réception, ainsi que son très faible coût mémoire (un seul paquet à mémoriser côté émetteur, B mémorise le paquet attendu lorsqu'il est correctement reçu). Son principal inconvénient est que l'émetteur ne fait rien tant qu'il n'a pas reçu l'acquittement de la station B. Durant ce temps, d'autant plus long que le temps de propagation est grand devant la durée d'un paquet, le canal est inutilisé. Ce protocole n'est donc pas optimum en terme de taux d'utilisation du lien.

2.2.2 *Go-Back-N*

Le protocole *Go-Back-N* (GBN) consiste à transmettre successivement des paquets sans attendre la réponse pour chacun. Il vise à remédier au délai d'attente en SaW par une transmission en continue sur une fenêtre appelée "fenêtre d'anticipation" dont le principe est présenté sur la figure 2.3. Cette stratégie d'envoi de paquets permet d'améliorer l'efficacité des transmissions (meilleur taux d'utilisation du lien). Le nombre maximum de paquets autorisés à être émis sans attendre un acquittement est appelé taille de la fenêtre d'anticipation, et noté w .

A chaque instant, la station A détient la liste des numéros de séquences des paquets qu'elle peut envoyer. De son côté, la station B détient le numéro de séquence du prochain paquet attendu. La station A transmet en continu les w paquets de sa fenêtre d'anticipation. Lorsque A reçoit un acquittement positif de B, elle déplace sa fenêtre d'une position et transmet les paquets suivants. Lorsque A reçoit un acquittement négatif, elle interrompt la transmission en cours et retransmet tous les paquets de la fenêtre à partir du paquet erroné. A doit donc stocker en mémoire les paquets non acquittés à l'intérieur de la fenêtre courante, pour pouvoir les retransmettre en cas de perte ou de mauvaise réception. De son côté, tout comme en SaW, la station B rejette tout paquet autre que le prochain paquet attendu. La taille de la fenêtre d'anticipation est calculée en fonction des paramètres du système (délai de propagation, taille de paquets, etc). Le nombre de bits b réservés à la numérotation de paquets est fonction de w . On montre qu'il est donné par $b = \lceil \log_2(w + 1) \rceil$ [18]. La figure 2.4 illustre un exemple de transmission s'appuyant sur le protocole GBN.

Grâce à l'utilisation de la fenêtre d'anticipation, le protocole GBN améliore le taux d'utilisation du lien comparativement au protocole SaW, en contrepartie, d'une complexité mémoire (on stocke w paquets) et d'une complexité de traitement supérieure côté émetteur. Côté récepteur, le fonctionnement du protocole est similaire au SaW.

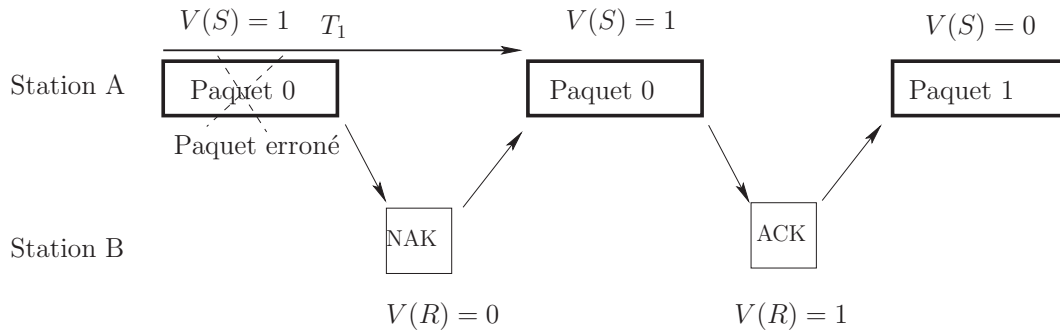


Figure 2.2 — Protocole *Send and Wait* ($V(S)$: numéro de séquence du prochain paquet à transmettre. $V(R)$: numéro de séquence du paquet attendu).

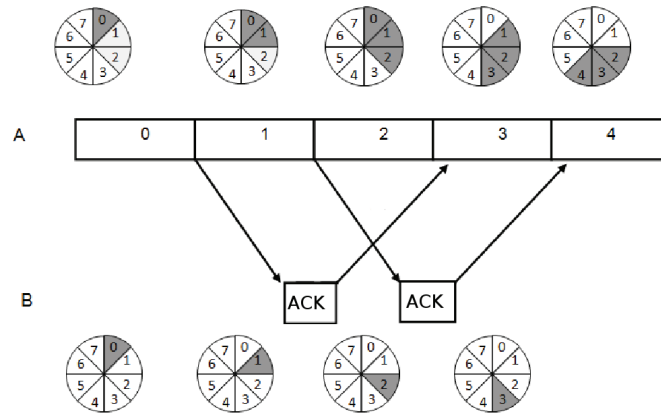


Figure 2.3 — Protocole à fenêtre d'anticipation de taille $w = 3$ paquets en émission et 1 paquet en réception.

Son principal inconvénient est qu'en cas de réception d'un acquittement négatif, on doit retransmettre tous les paquets déjà envoyés à partir du paquet signalé en erreur, ce qui réduit d'autant le débit utile du protocole.

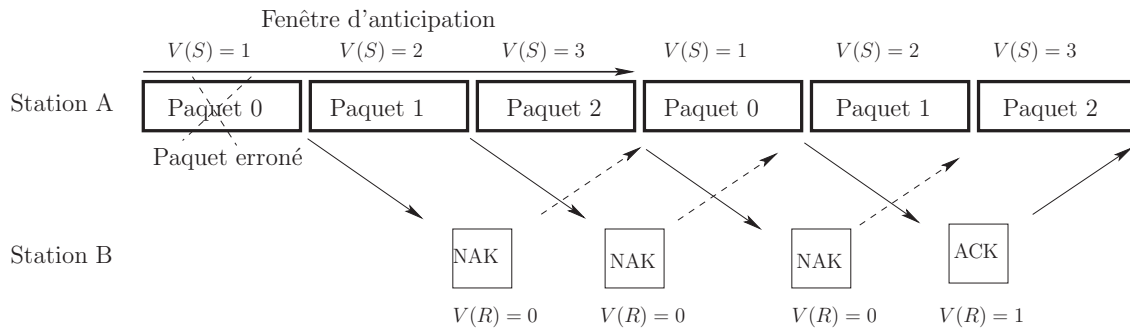


Figure 2.4 — Protocole Go-Back-N avec une fenêtre d'anticipation de taille 3

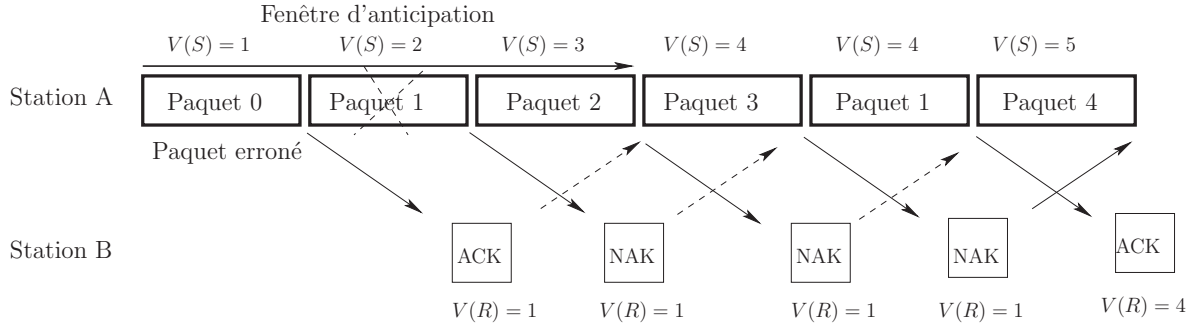


Figure 2.5 — Protocole SR avec une fenêtre d'anticipation de taille 3

2.2.3 Selective Reject

Le protocole *Selective Reject* (SR) cherche à améliorer l'efficacité du GBN en minimisant le nombre de retransmissions grâce à l'utilisation d'une fenêtre supplémentaire en réception. La station A dispose d'une fenêtre d'anticipation comme dans le protocole GBN. Initialement, A transmet les paquets de sa fenêtre sans attendre la réponse pour chaque paquet. La station B dispose de son côté d'une fenêtre de même taille pour mémoriser les prochains paquets attendus, dont elle détient et met à jour la liste des numéros de séquence. Lorsque A reçoit un acquittement positif, elle déplace sa fenêtre d'une position et continue sa transmission. Lorsqu'elle reçoit un acquittement négatif, elle retransmet uniquement le paquet rejeté. Par cette retransmission sélective, elle évite ainsi le problème de retransmission multiple qui affecte le protocole GBN. La figure 2.5 illustre un exemple de retransmission en SR.

Le protocole SR est le protocole le plus performant si les mémoires en émission et en réception sont suffisamment grandes, car il optimise au mieux l'utilisation du lien, et minimise par ailleurs le nombre de retransmissions relatives à chaque paquet. En contrepartie, les règles de traitement sont plus complexes, et les besoins mémoire supérieurs à ceux des autres protocoles. En SR, tout comme en GBN, la station A doit garder en mémoire les w paquets de sa fenêtre d'anticipation jusqu'à leur bonne réception. A la différence du SaW et du GBN, la station B doit ici disposer à son tour d'une mémoire suffisante pour stocker les paquets bien reçus et les délivrer dans le bon ordre aux couches supérieures. On montre que le nombre de bits b réservés à la numérotation de paquets pour une fenêtre d'anticipation de taille w vaut $b = \lceil \log_2(2w) \rceil$ [18].

2.2.4 Send and wait parallèle

Plusieurs généralisations du protocole SaW ont été proposées dans la littérature afin d'améliorer le taux d'occupation du lien tout en conservant au maximum la simplicité de ce protocole [19, 20, 21]. En particulier, la variante dite *Parallel Send and wait* (PSW) consiste à faire tourner plusieurs processus SaW en parallèle. Chaque processus fonctionne comme un protocole classique SaW. Le nombre maximal de processus w

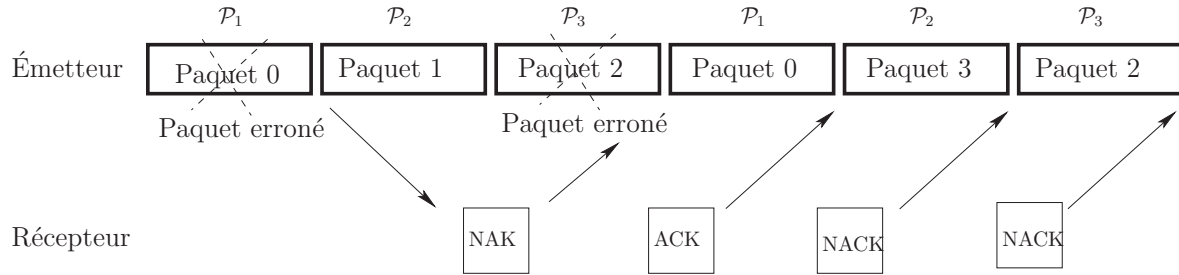


Figure 2.6 — Protocole ARQ *Send and Wait* parallèle avec un nombre de processus SaW parallèles $w = 3$.

doit être égal au nombre maximal de paquets que l'on peut transmettre pendant la durée d'aller-retour (taille d'une fenêtre d'anticipation). L'augmentation du nombre de processus SaW en parallèle au-delà du nombre maximum de paquets autorisés introduit des retards sur les retransmissions et n'est donc pas souhaitable [22]. La figure 2.6 illustre le principe du protocole PSW sur un exemple avec trois processus en parallèle ($\mathcal{P}_1$, $\mathcal{P}_2$ et $\mathcal{P}_3$).

L'avantage du protocole PSW est qu'il permet d'améliorer le débit du protocole SaW grâce à une transmission en continu, tout en conservant pour l'essentiel, la simplicité des règles du SaW classique. Par l'utilisation de w processus SaW en parallèle, le protocole PSW peut être vu comme un protocole à fenêtre d'anticipation de taille w . Lorsque le nombre de processus est ajusté au délai d'aller-retour, le protocole PSW permet d'occuper au maximum les délais d'attente. Grâce à un taux d'utilisation du lien de 100%, il aura un débit comparable au protocole SR, et par conséquent un débit supérieur au protocole GBN. Comme ce protocole est composé de plusieurs processus (protocoles) SaW, la complexité des règles de dialogue est réduite par rapport au SR. Les demandes en mémoire sont également réduites. Notons toutefois que les entêtes des paquets doivent inclure un champ supplémentaire pour la numérotation des processus SaW parallèles. Par ailleurs, tout comme en SR, le ré-ordonnancement des paquets en réception est nécessaire.

2.2.5 Métriques de performances

L'analyse des performances des protocoles ARQ peut se faire selon plusieurs critères. Les critères le plus souvent considérés sont le débit utile moyen, qui mesure le nombre moyen de bits transmis avec succès par intervalle de temps, et la fiabilité, qui mesure le taux d'échec de transmission à l'issue du protocole ARQ. D'autres métriques peuvent également être utilisées, telles que le délai de transfert (latence) ou bien les ressources mémoire en émission et en réception (dimensionnement des buffers) [23]. Une métrique complémentaire particulièrement pertinente dans certains contextes de transmission avec ressources limitées tels que les réseaux de capteurs est la notion d'efficacité énergétique.

Par la suite, nous avons choisi d'analyser et comparer les différents protocoles ARQ sur la base du débit utile moyen et d'efficacité énergétique.



Figure 2.7 — Format de paquet transmis composé des données utiles de taille D bits et de l'entête de taille H bits.

2.3 Calcul et optimisation du débit utile moyen

Dans cette section, on établit l'expression générique du débit utile moyen pour chaque protocole ARQ en fonction de différents paramètres du modèle de transmission et du protocole ARQ. On discute ensuite de l'optimisation du débit utile en jouant sur les paramètres du système.

2.3.1 Hypothèses et modèle temporel de la transmission

On considère un modèle de transmission entre un émetteur (station A) et un récepteur (station B) sur un canal duplex. Ce canal introduit un délai de propagation identique sur le canal direct et sur le canal de retour. On suppose que l'émetteur envoie uniquement des données au récepteur et que le récepteur répond par des acquittements.

Hypothèses clé

On suppose que les entêtes sont protégés séparément des données utiles. La protection séparée de l'entête permet d'identifier un paquet erroné lorsque son entête est correctement reçu. Dans ce cas, le récepteur envoie un acquittement négatif à l'émetteur. Dans cette étude on suppose que les entêtes et les acquittements sont fortement protégés et qu'ils ne sont pas sujets aux erreurs (probabilité d'erreur sur l'entête négligeable devant la probabilité d'erreur de paquets) [24, 25, 26]. On note par H la taille de l'ensemble des champs supplémentaires nécessaires au bon fonctionnement du protocole, notamment les numéros de paquets ainsi que la redondance (CRC pour la détection d'erreurs). Les données utiles sont découpées en blocs de taille identique notée D . La taille totale des paquets émis est donc $H + D$ (bits). La figure 2.7 illustre le format de paquet considéré. Les informations précisant le type de paquet sont généralement placées au début du paquet et le CRC est généralement placé à la fin. Nous nous intéressons seulement à la taille totale des champs supplémentaires et représentons ceux-ci à la fin du paquet pour simplifier les illustrations. Tous les calculs sont valides quelle soit la place effective des ces champs supplémentaires.

Le modèle du canal direct considéré introduit des pertes de paquet indépendantes et identiquement distribuées (i.i.d.) avec une probabilité p_o . Dans un premier temps, le canal de retour est supposé parfait (acquittements reçus sans erreurs). L'extension à une voie de retour imparfaite sera discutée dans la section 2.5.

Le protocole est organisé en cycles

On définit un *round ARQ* par la transmission d'un paquet par la couche physique suivi de la transmission de son acquittement. La durée d'un round comprend la durée de transmission du paquet et la durée de transmission de l'acquittement (on néglige ici les temps de traitement). On désigne par *cycle ARQ* la succession de plusieurs rounds successifs (au minimum un round, au maximum $M_1 + 1$ rounds). Un cycle ARQ est toujours relatif à un unique paquet de données. Il se termine lorsque le paquet a été acquitté positivement, ou bien lorsque l'on atteint le nombre maximal M_1 de rounds autorisés.

Modèle temporel considéré

Lorsqu'un paquet est transmis, la station A recevra l'acquittement de B après une durée T_{round} . Cette durée est la durée d'un round ARQ. Elle comporte la durée d'émission de paquet T_e , la durée d'émission d'acquittements T_a , les durées de traitement de paquet et d'acquittement T_t et les délais de propagation T_p . Elle est donnée par (cf. figure 2.8) :

$$T_{\text{round}} = T_e + 2T_p + 2T_t + T_a \quad (2.1)$$

Les durées de traitement de paquets et d'acquittements sont négligeables devant les durées d'émission et les délais de propagation. La durée élémentaire d'un round ARQ se réduit à :

$$T_{\text{round}} = T_e + T_a + 2T_p \quad (2.2)$$

Dans l'analyse à suivre, il est commode de raisonner en durées normalisées, relativement à la durée T_e d'un paquet de la couche physique. Notons par $\rho = \frac{D}{D+H}$ le rapport entre le nombre de bits utiles d'un paquet et le nombre total de bits par paquet, et par $\tau = \frac{T_p}{T_e}$ le délai de propagation normalisé. En considérant que les acquittements ont la même taille que les entêtes ($T_a/T_e = H/(D+H) = 1 - \rho$), la durée normalisée $T_1 = \frac{T_{\text{round}}}{T_e}$ d'un round ARQ s'écrit alors :

$$T_1 = 2 - \rho + 2\tau \quad (2.3)$$

Pour un délai de propagation T_p constant, le temps de propagation normalisé τ dépend du nombre N de symboles modulés d'un paquet. Pour cela, on introduit une grandeur normalisée notée τ_o définie par le rapport entre le délai de propagation T_p et la durée T_s d'un symbole. On a :

$$\tau_o = \frac{T_p}{T_s} \quad (2.4)$$

$$\tau = \frac{T_p}{T_e} = \frac{T_p}{NT_s} = \frac{1}{N}\tau_o \quad (2.5)$$

Dans la suite, on utilisera les paramètres D , H , ρ , τ et τ_o pour évaluer les performances des protocoles ARQ.

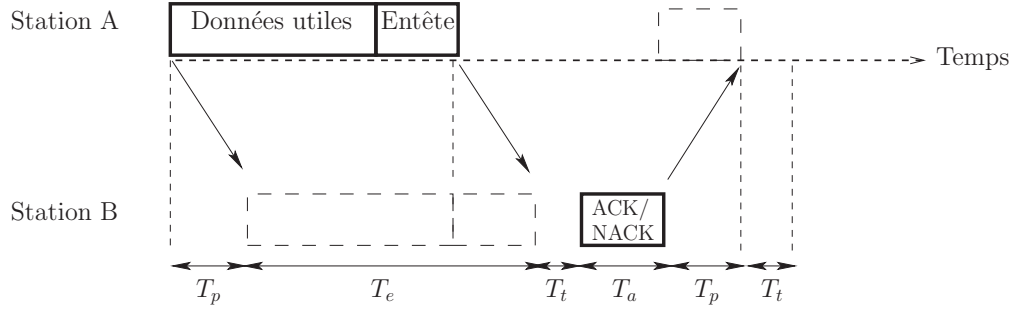


Figure 2.8 — Schéma explicatif montrant les durées d'émission de paquets et d'acquittements, les durées de traitement et le délai de propagation.

2.3.2 Débit utile moyen en SaW

Le débit utile moyen η est défini comme le rapport du nombre moyen de bits d'information correctement reçus par B au cours d'un cycle, $E[\mathcal{B}]$, sur la durée moyenne requise pour délivrer un paquet (durée moyenne d'un cycle) $E[\mathcal{T}]$, exprimée ici en nombre de périodes symbole de manière à avoir un débit homogène à un taux de codage :

$$\eta = \frac{E[\mathcal{B}]}{E[\mathcal{T}]} \text{ (bits/symbole)}. \quad (2.6)$$

Pour calculer ce débit, on introduit à cet effet différentes probabilités très importantes qui sont utilisées tout au long du document. On note par e_m l'événement “échec de décodage (perte de paquet) au round m ”, et par E_m l'événement “échec de transmission au round m ”. Notons que d'une manière générale, un échec de transmission peut être dû soit à la mauvaise réception du paquet, soit à la perte de l'acquittement (positif ou négatif). **Dans le cas particulier retenu ici d'une transmission avec voie de retour parfaite**, on a la simplification :

$$\Pr\{E_m\} = \Pr\{e_m\} \quad (2.7)$$

Soit $p(m) = \Pr\{e_0, e_1, \dots, e_{m-1}, e_m\}$ la probabilité conjointe d'avoir $m + 1$ échecs de décodage consécutifs aux rounds 0 à m compris. On définit d'une manière similaire la probabilité conjointe $P(m) = \Pr\{E_0, E_1, \dots, E_{m-1}, E_m\}$ d'avoir $m + 1$ échecs de transmission consécutifs aux rounds 0 à m compris. Enfin, on note par $Q(m) = \Pr\{E_0, E_1, \dots, E_{m-1}, \bar{E}_m\}$ la probabilité conjointe d'avoir m échecs de transmission consécutifs aux rounds 0 à $m - 1$ suivi d'un succès au round m . On remarque que $P(m) + Q(m) = P(m - 1)$. Par conséquent :

$$Q(m) = P(m - 1) - P(m) \quad (2.8)$$

Par ailleurs, **dans le cas d'une voie de retour parfaite**, on a l'égalité :

$$P(m) = p(m) \quad (2.9)$$

Pour calculer le débit utile, développons tout d'abord l'expression du nombre moyen $E[B]$ de bits d'information correctement reçus. On considère pour cela que la perte d'un

acquittement positif au dernier round ne constitue pas un échec de transmission. En effet, dans ce cas les données ont été bien reçues par la station B, et la station A ne retransmet pas ces données car le nombre maximal de retransmissions autorisées est déjà atteint. Les données transmises sont soit perdues avec une probabilité $p(M_1)$ (0 bits correctement reçu), soit bien reçues avec une probabilité $1 - p(M_1)$ (D bits correctement délivrés à B). Le nombre moyen de bits d'information correctement reçus par cycle ARQ vaut donc :

$$E[\mathcal{B}] = D(1 - p(M_1)) \quad (\text{bits}) \quad (2.10)$$

Calculons maintenant la durée moyenne d'un cycle ARQ en SaW. Celle ci dépend à la fois de la durée d'un round ARQ (exprimée ici en nombre de périodes symboles) et du nombre moyen $\bar{N}_R$ de rounds ARQ :

$$E[\mathcal{T}] = \bar{N}_R(2 - \rho + 2\tau)N \quad (\text{durées symbole}) \quad (2.11)$$

On peut développer le nombre moyen $\bar{N}_R$ de rounds par cycle ARQ comme suit :

$$\bar{N}_R = \sum_{m=0}^{M_1} (m+1)Q(m) + (M_1+1)P(M_1) \quad (2.12)$$

En utilisant le résultat (2.8), ce dernier peut encore s'écrire :

$$\bar{N}_R = 1 + \sum_{m=0}^{M_1-1} P(m). \quad (2.13)$$

On aboutit donc à l'expression suivante du débit utile moyen en SaW :

$$\eta_{sw} = \frac{D}{N} \frac{1}{2 - \rho + 2\tau} \frac{1 - p(M_1)}{1 + \sum_{m=0}^{M_1-1} P(m)} \quad (2.14)$$

En rappelant que le rendement d'une transmission vaut $R_o = \frac{D}{N}$ (bits/symbole), le débit utile moyen en SaW peut finalement s'exprimer comme :

$$\eta_{sw} = R_o \frac{1}{2 - \rho + 2\tau} \frac{1 - p(M_1)}{1 + \sum_{m=0}^{M_1-1} P(m)} \quad (2.15)$$

Cette expression est très générale. Dans le contexte considéré ici d'une transmission avec voie de retour parfaite et pertes de paquets i.i.d. de round en round avec une probabilité p_o , la probabilité conjointe $P(m)$ se simplifie en :

$$P(m) = p(m) = p_o^{m+1} \quad (2.16)$$

et le débit utile en SaW devient :

$$\eta_{sw} = R_o \frac{1}{2 - \rho + 2\tau} \frac{1 - p_o^{M_1+1}}{1 + \sum_{m=1}^{M_1} p_o^m} \quad (2.17)$$

Dans la littérature, le débit utile moyen des protocoles ARQ est habituellement donné pour un nombre infini M_1 de retransmissions (protocole *non tronqué*). En utilisant le résultat :

$$\sum_{m=0}^{\infty} p_o^m = \frac{1}{1 - p_o}, \quad (2.18)$$

on peut déduire de (2.17) le débit utile du protocole SaW non tronqué :

$$\lim_{M_1 \rightarrow \infty} \eta_{sw} = R_o \frac{1 - p_o}{2 - \rho + 2\tau} \quad (2.19)$$

Notons que les expressions habituellement données dans la littérature négligent souvent le temps normalisé $1 - \rho$ d'émission des acquittements [17, 27, 28].

On voit ici que le débit utile moyen en SaW dépend du délai de propagation normalisé τ . En particulier plus le délai de propagation est important comparé à la durée du paquet (τ grand devant 1), plus le débit utile est faible car l'émetteur passe beaucoup de temps à attendre.

2.3.3 Débit utile moyen en GBN

Le protocole GBN, utilise une fenêtre d'anticipation pour améliorer le taux d'utilisation du lien. Comme dans le protocole SaW, le nombre moyen de bits d'information correctement reçus est $D(1 - p(M_1))$. Pour calculer le débit utile moyen en GBN, il est nécessaire de calculer la durée moyenne consommée par cycle. Dans ce protocole, on suppose que la fenêtre d'anticipation est ajustée à la durée d'aller-retour $T_1 = 2 - \rho + 2\tau$. Ceci signifie que la durée consommée par l'émission de tous les paquets de la fenêtre d'anticipation vaut $(2 - \rho + 2\tau)N$ périodes symboles. Chaque round se concluant par un échec consomme précisément cette durée. Le dernier round consomme quant à lui N périodes symboles. La durée moyenne d'un cycle en GBN vaut donc :

$$E[\mathcal{T}] = (\bar{N}_R - 1)(2 - \rho + 2\tau)N + N \quad (2.20)$$

où $\bar{N}_R$ est donné par (2.12), et le débit utile moyen en GBN s'écrit :

$$\eta_{gbn} = R_o \frac{1 - p(M_1)}{1 + (2 - \rho + 2\tau) \sum_{m=0}^{M_1-1} P(m)} \quad (2.21)$$

En spécialisant l'expression précédente au scénario considéré dans cette section, on obtient finalement :

$$\eta_{gbn} = R_o \frac{1 - p_o^{M_1+1}}{1 + (2 - \rho + 2\tau) \sum_{m=1}^{M_1} p_o^m} \quad (2.22)$$

Le débit utile du protocole GBN non tronqué s'en déduit comme :

$$\lim_{M_1 \rightarrow \infty} \eta_{gbn} = R_o \frac{1 - p_o}{1 - p_o + (2 - \rho + 2\tau)p_o} \quad (2.23)$$

Cette expression est identique à l'expression donnée dans la littérature où $w = 2 - \rho + 2\tau$ est la taille de la fenêtre d'anticipation [17]. On remarque que, tout comme en SaW, le

débit utile moyen en GBN dépend du temps de propagation normalisé. Il est d'autant plus faible que τ est grand. Toutefois, lorsque le canal est relativement bon ($p_o \ll 1$), chacun des paquets de la fenêtre d'anticipation est correctement reçu avec une probabilité élevée. Le protocole GBN réalise alors une meilleure utilisation du lien que le SaW et son débit est beaucoup moins impacté par τ , quelque soit la valeur de ce dernier (dénominateur proche de 1 dans l'expression (2.23))

2.3.4 Débit utile moyen en SR

Comme dans les protocoles précédents, le nombre moyen de bits d'information correctement reçus par cycle est $D(1 - p(M_1))$. Contrairement au protocole GBN, l'échec de transmission d'un paquet engendre uniquement la retransmission sélective de ce paquet. La durée perdue pour chaque round correspond donc à la durée d'émission des N symboles du paquet, soit :

$$E[\mathcal{T}] = \bar{N}_R \times N \quad (2.24)$$

Par conséquent, le débit utile en SR est :

$$\eta_{sr} = R_o \frac{1 - p(M_1)}{1 + \sum_{m=0}^{M_1-1} P(m)} \quad (2.25)$$

Avec les hypothèses de transmission considérées dans cette section, l'expression précédente se simplifie en :

$$\eta_{sr} = R_o \frac{1 - p_o^{M_1+1}}{1 + \sum_{m=1}^{M_1} p_o^m} \quad (2.26)$$

On en déduit le débit utile moyen du protocole SR non tronqué :

$$\lim_{M_1 \rightarrow \infty} \eta_{sr} = R_o(1 - p_o) \quad (2.27)$$

On reconnaît ici l'expression habituellement donnée dans la littérature [17, 27, 28]. On peut noter en particulier que le débit est indépendant du délai de propagation normalisé τ (sous réserve d'avoir des mémoires de tailles suffisantes). Le débit en SR correspond à ce que l'on peut faire de mieux avec un protocole ARQ dans ce contexte de transmission.

2.3.5 Débit utile moyen en SWP

Le débit utile en PSW dépend du nombre w de processus SaW placés en parallèle, relativement à la durée normalisée d'aller-retour T_1 . D'une manière générale, le fait d'utiliser w processus en parallèle améliore le taux d'utilisation du lien d'un facteur w dans l'expression (2.17) du débit utile moyen en SaW. En particulier, lorsque $w \geq T_1$, le lien est utilisé à 100% et on atteint alors le débit idéal du protocole SR. Le débit

utile moyen en PSW peut donc s'écrire :

$$\begin{aligned}\eta_{psw} &= \min \left(w \frac{R_o}{2 - \rho + 2\tau} \frac{1 - p(M_1)}{1 + \sum_{m=0}^{M_1-1} P(m)}, R_o \frac{1 - p(M_1)}{1 + \sum_{m=0}^{M_1-1} P(m)} \right) \\ &= \min \left(\frac{w}{T_1}, 1 \right) \eta_{sr}\end{aligned}\quad (2.28)$$

Le nombre minimal de processus à mettre en parallèle pour avoir le débit maximal correspond donc au nombre maximal de paquet que l'on puisse transmettre pendant la durée d'un round ARQ, soit :

$$w = \lceil T_1 \rceil \quad (2.29)$$

Dans ce cas de figure, on a exactement le même débit utile que le protocole SR :

$$\eta_{psw} = \eta_{sr} \quad (2.30)$$

2.3.6 Comparaison des différents protocoles

Dans cette section, nous comparons les protocoles ARQ sur la base du débit utile moyen. Nous considérons pour cela une transmission de paquets de données de taille fixe $D = 1000$ bits avec des entêtes de $H = 40$ bits. Les symboles émis sont issus d'une modulation de phase à deux états (BPSK) et transmis sur un canal gaussien réel à entrées binaires de densité spectrale de puissance $N_o/2$. Les symboles transmis ont une énergie moyenne E_s . Le rapport signal à bruit γ est défini comme $\gamma = \frac{E_s}{N_o} = \frac{E_b}{N_o}$ ici (modulation BPSK, où E_b est l'énergie bit, cf. chapitre I). La taille d'un paquet est donc $N = 1040$ symboles BPSK, et le rendement de transmission vaut $R_o = 1000/1040 = 0,96$. La probabilité d'erreur binaire en sortie du démodulateur est donnée par :

$$p_b = \frac{1}{2} \operatorname{erfc}(\sqrt{\gamma}) \quad (2.31)$$

La probabilité d'échec de décodage d'un paquet s'écrit alors :

$$p_o = 1 - (1 - p_b)^N \quad (2.32)$$

On considère des protocoles non tronqués et on compare les débits moyens en SaW, GBN, SR et PSW, donnés respectivement par (2.19), (2.23), (2.27) et (2.28). Les figures 2.9 et 2.10 présentent l'évolution de η en fonction du rapport signal à bruit γ , pour un délai de propagation normalisé τ respectivement fixé à $1/2$ et $3/2$. Le cas $\tau = 1/2$ correspond à une situation où l'émetteur ne peut transmettre qu'un seul paquet entre l'instant de fin d'émission d'un paquet et l'instant de réception de son acquittement. Le second cas $\tau = 3/2$ correspond à un délai d'aller-retour suffisant pour transmettre 3 paquets (fenêtre d'anticipation de taille $w = 4$) (cf. figures 2.11 et 2.12).

On constate tout d'abord que le protocole SR offre le débit le plus important, et que ce débit est indépendant du temps de propagation normalisé τ . Le protocole SaW présente les moins bonnes performances. Son débit utile est inférieur à celui des protocoles à fenêtre d'anticipation, et très sensible à l'augmentation de τ . Les performances

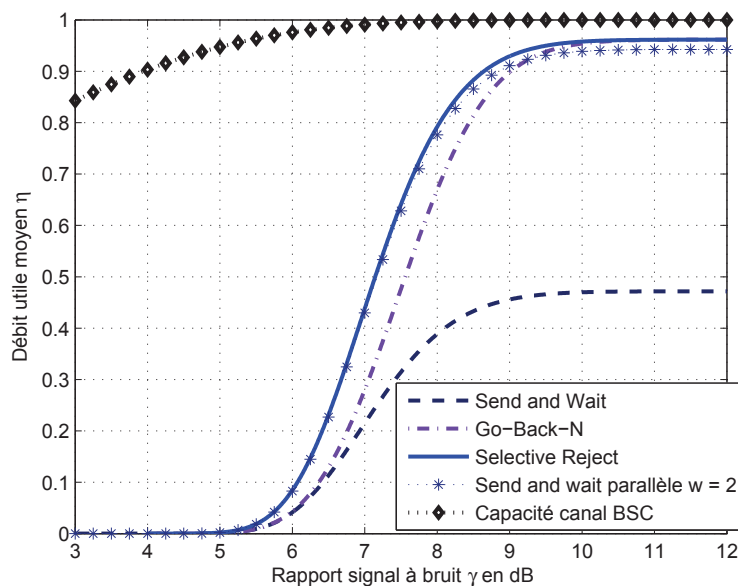


Figure 2.9 — Débit utile normalisé en fonction du RSB en SaW, GBN, SR et PSW pour $D = 1000$ bits, $H = 40$ bits et $\tau = 1/2$.

du GBN dépendent également de la valeur de τ , et sont d'autant meilleures que τ est petit. On peut noter que le GBN surpasse le SaW dans tous les cas, et qu'il atteint les performances du protocole SR à fort RSB (quelque soit la valeur de τ). Enfin, on observe que le protocole PSW présente des performances équivalentes à celles du SR. La différence asymptotique (fort RSB) observée entre ces deux protocoles sur les figures 2.9 et 2.10 est simplement due à un léger sous dimensionnement du nombre w de processus placés en parallèle. En effet, on a $T_1 = 2 - \rho + 2\tau$ et $\rho = 1000/1040$ soit $T_1 = 2,04$ pour $\tau = 1/2$, et $T_1 = 4,03$ pour $\tau = 3/2$. En toute rigueur, il aurait donc fallu choisir $w = 3$ et $w = 5$ au lieu de 2 et 4 pour atteindre exactement les performances asymptotiques du SR.

A titre d'information, nous avons également tracé sur les figures 2.9 et 2.10 la courbe de capacité du canal binaire symétrique (BSC) donnée par $C_{\text{BSC}} = 1 + p_b \log_2(p_b) + (1 - p_b) \log_2(1 - p_b)$ ¹. On remarque que pour le scénario de transmission considéré ici, les protocoles ARQ classiques sont nettement sous optimaux par rapport aux limites établies par la théorie de l'information.

2.3.7 Optimisation du débit utile moyen

Le débit utile en ARQ dépend d'un certain nombre de paramètres, et en particulier du rapport signal à bruit γ , des tailles respectives D et H du paquet de données et de son entête, et du nombre maximum M_1 de retransmissions autorisées. On peut noter que η est une fonction croissante du RSB et de M_1 . Reprenons le scénario étudié dans la section précédente, et fixons la taille H de l'entête ainsi que le RSB γ . Nous avons vu en

1. Sur un canal sans mémoire, la capacité avec voie de retour est égale à la capacité sans lien retour.

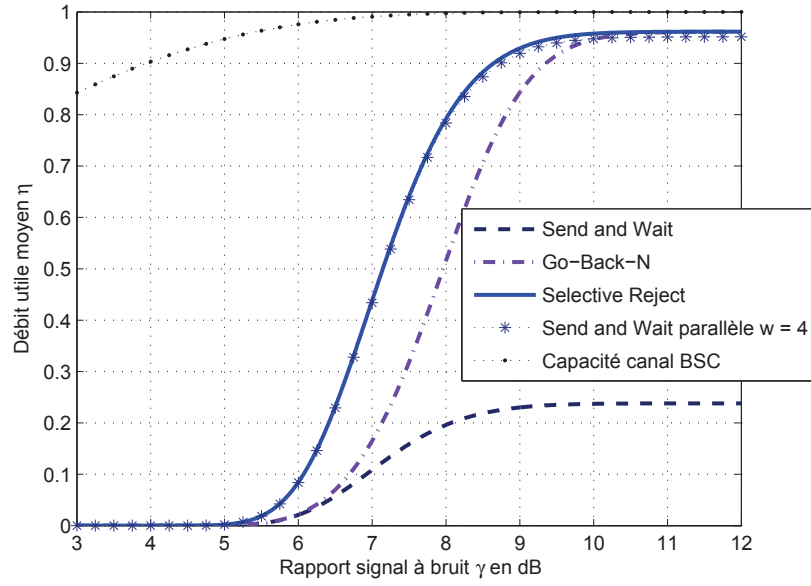


Figure 2.10 — Débit utile normalisé en fonction du RSB en SaW, GBN, SR et PSW pour $D = 1000$ bits, $H = 40$ bits et $\tau = 3/2$.

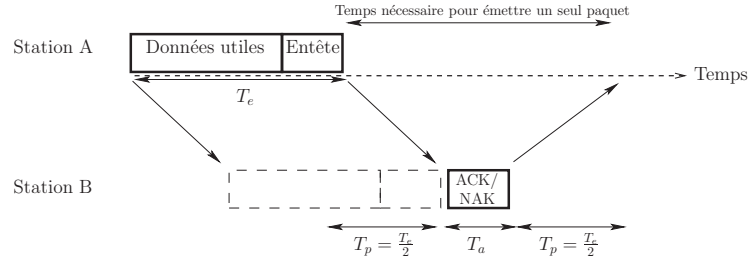


Figure 2.11 — Schéma explicatif du cas de transmission avec un délai de propagation $T_p = \frac{T_e}{2}$ ($\tau = 0,5$).

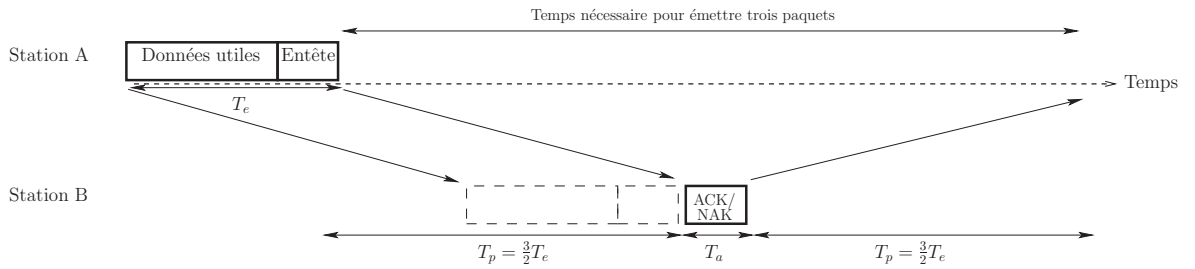


Figure 2.12 — Schéma explicatif du cas de transmission avec un délai de propagation $T_p = \frac{3}{2}T_e$ ($\tau = 1,5$).

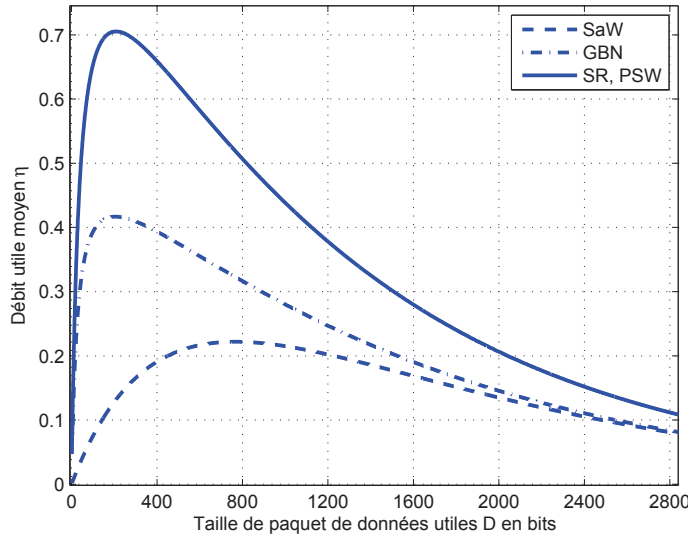


Figure 2.13 — Débit utile normalisé en fonction de la taille D des données utiles en SaW, GBN, SR et PSW pour $\gamma = 7$ dB, $H = 40$ bits et $\tau_o = 520$.

section 2.3.1 que la durée de propagation normalisée τ peut s'écrire comme $\tau = \tau_o/N$, où $\tau_o = \frac{T_p}{T_s}$ (durée de propagation normalisée par rapport à la période symbole) est une constante du scénario. Dans ce cas, à τ_o fixé, le débit utile moyen dépend uniquement de la taille D des paquets de données. Augmenter D accroît la probabilité de perte de paquet p_o donnée par (2.32) ($N = D + H$ pour la modulation BPSK), ce qui diminue le débit utile moyen. Diminuer D rend la proportion de l'entête moins négligeable devant D , et donc réduit également le débit utile moyen. Ceci suggère l'existence d'une taille optimale de paquet qui maximise le débit utile moyen dans ces conditions.

Pour illustrer ce point, nous présentons sur les figures 2.13 et 2.14 l'évolution du débit η en fonction de la taille D d'un paquet de données pour les différents protocoles ARQ, et pour deux valeurs distinctes du délai de propagation $\tau_o = 520$ et $\tau_o = 1560$. Ces deux choix correspondent à des délais $\tau = 1/2$ et $\tau = 3/2$ pour une taille de paquet $N = 1040$ symboles. Les figures ont été obtenues en fixant $H = 40$ bits et $\gamma = 7$ dB (zone de moyen RSB). On observe qu'il existe bien une taille optimale qui maximise le débit. Cette taille optimale dépend du protocole considéré. Elle est indépendante de τ_o en SR, et vaut $D = 208$ bits en SR dans l'exemple considéré. La taille optimale en GBN a le même ordre de grandeur et reste peu impactée par τ_o . En revanche, la taille optimale en SaW dépend fortement du délai de propagation τ_o . Cette taille est de 768 bits pour $\tau_o = 520$, et de 993 bits pour $\tau_o = 1560$. On remarque ici que l'augmentation du délai de propagation n'engendre pas systématiquement l'augmentation de la taille optimale de données utiles. En effet, plus D augmente, plus p_o augmente indépendamment du τ_o , ce qui réduit d'autant le débit utile moyen.

L'optimisation précédente a été réalisée à RSB et taille d'entête constants. Plus généralement, on pourrait chercher le jeu de paramètres (D, H) optimal ou de manière équivalente le couple (ρ, N) , qui maximise le débit pour chaque RSB [29, 30, 31, 32].

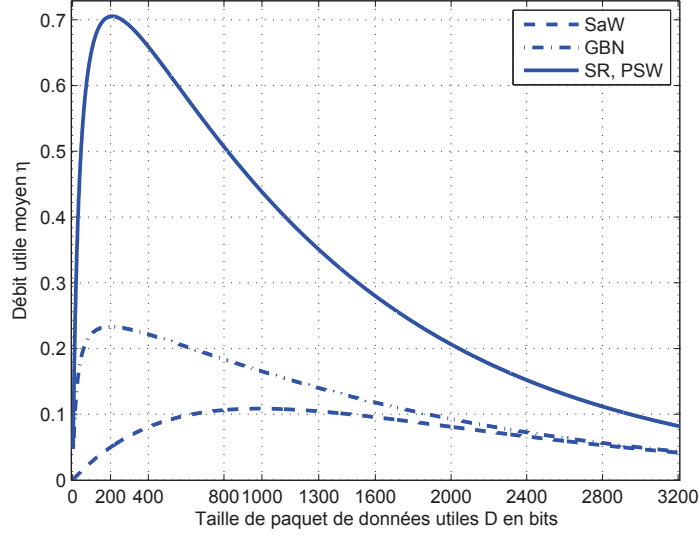


Figure 2.14 — Débit utile normalisé en fonction de la taille D des données utiles en SaW, GBN, SR et PSW pour $\gamma = 7$ dB et $H = 40$ bits et $\tau_o = 1560$.

2.4 Étude et optimisation de l'efficacité énergétique

Dans les sections précédentes, nous avons comparé les protocoles ARQ selon le critère du débit utile moyen. Ce critère n'est pas toujours le critère le plus pertinent. En effet, dans certains scénarios tels que les réseaux de capteurs, la bonne gestion de l'énergie consommée est une métrique clé [33]. Dans cette section, nous revisitions les protocoles ARQ sous l'angle de l'optimisation de la gestion de l'énergie. Plus précisément, nous étudions l'efficacité énergétique de chaque protocole, qui mesure l'énergie moyenne totale ζ requise pour transmettre avec succès un bit de donnée utile au destinataire.

2.4.1 Définition et calcul de l'efficacité énergétique

L'énergie nécessaire par bit utile ζ prend en compte l'énergie totale E_{tot} consommée par la transmission des paquets et des acquittements durant un cycle ARQ, normalisée par le nombre moyen de bits correctement reçus à l'issue du cycle. Elle s'écrit :

$$\zeta = \frac{E_{tot}}{D(1 - p(M_1))} \quad (2.33)$$

Il est à noter ici qu'avec cette définition, $\zeta \rightarrow \infty$ lorsque $p(M_1) \rightarrow 1$ (paquets jamais délivré au destinataire), et qu'elle est égale à $\frac{E_{tot}}{D}$ pour $p(M_1) \rightarrow 0$ (protocole très fiable), ce qui est intuitivement satisfaisant.

L'énergie nécessaire par bit utile ζ dépend de la taille D d'un paquet de données utiles, et de la taille H des entêtes et des acquittements (supposés ici, pour simplifier, de même taille que les entêtes). Elle dépend également de l'énergie d'un symbole modulé

sur la voie directe, et sur la voie de retour. Pour simplifier, on considère ici que l'on utilise la même énergie par symbole E_s pour les paquets et les acquittements, et l'on considère une modulation BPSK (énergie d'un symbole E_s = énergie d'un bit E_b). Dans la suite, nous dérivons l'expression de l'énergie nécessaire par bit utile ζ pour chacun des protocoles ARQ précédemment étudiés.

Protocoles SaW, SR et PSW

Un échec de transmission en SaW, SR ou PSW provoque la retransmission du paquet mal reçu. L'énergie consommée par round est la somme de l'énergie consommée par le paquet et son acquittement, soit $NE_s + HE_s = E_s N(2 - \rho)$. L'énergie moyenne totale consommée par cycle est l'énergie consommée par round multipliée par le nombre moyen $\bar{N}_R$ de rounds par cycle, donné par (2.13). Elle s'écrit :

$$E_{tot} = E_s N(2 - \rho) \left(1 + \sum_{m=0}^{M_1-1} P(m) \right) \quad (2.34)$$

On en déduit l'énergie moyenne ζ nécessaire par bit utile en SaW, SR et PSW en normalisant par le nombre moyen de bits correctement délivrés à l'issue du cycle :

$$\zeta = E_s \frac{(2 - \rho)}{R_o} \frac{1 + \sum_{m=0}^{M_1-1} P(m)}{1 - p(M_1)} \quad (2.35)$$

On en déduit l'expression de ζ dans le cas de protocoles SaW, SR ou PSW non tronqués ($M_1 \rightarrow \infty$) avec des pertes de paquets i.i.d. avec une probabilité p_o :

$$\zeta = E_s \frac{2 - \rho}{R_o} \frac{1}{1 - p_o} \quad (2.36)$$

Protocole GBN

Dans le cas du GBN, du fait de la fenêtre d'anticipation, on a une transmission continue sur la voie directe. A chaque round ARQ se soldant par un échec de transmission, on consomme donc une énergie proportionnelle à la durée d'aller-retour (taille de la fenêtre d'anticipation) sur la voie directe, à la quelle s'ajoute l'énergie d'un acquittement sur la voie de retour, soit $NE_s(2 - \rho + 2\tau) + HE_s$. L'énergie consommée au dernier round, terminant par succès ou échec, est la somme de l'énergie consommée par le paquet et son acquittement. L'énergie moyenne totale consommée par cycle en GBN vaut donc :

$$E_{tot} = E_s N(3 - 2\rho + 2\tau) \sum_{m=0}^{M_1-1} P(m) + NE_s(2 - \rho) \quad (2.37)$$

L'énergie moyenne ζ nécessaire par bit utile en GBN vaut alors :

$$\zeta = E_s \frac{1}{R_o} \frac{(3 - 2\rho + 2\tau) \sum_{m=0}^{M_1-1} P(m) + (2 - \rho)}{1 - p(M_1)} \quad (2.38)$$

Dans le cas d'un protocole GBN non tronqué ($M_1 \rightarrow \infty$) avec perte de paquets i.i.d. il vient :

$$\zeta = \frac{E_s}{R_o} \left((3 - 2\rho + 2\tau) \left(\frac{p_o}{1 - p_o} \right) + 2 - \rho \right) \quad (2.39)$$

2.4.2 Optimisation de l'efficacité énergétique

En observant par exemple les expressions (2.36) et (2.39), on constate que l'énergie moyenne ζ nécessaire par bit utile dépend de plusieurs paramètres : l'énergie transmise par symbole E_s (ou de manière équivalente le RSB γ , où N_o est fixé à 1), le rendement R_o de la transmission, le rapport $\rho = \frac{D}{D+H}$, la durée de propagation normalisée τ , et la probabilité de perte de paquet p_o . En rappelant que $R_o = \frac{D}{N}$ et que $\tau = \frac{\tau_o}{N}$, ζ est donc fonction de E_s , D , H et τ_o .

De la même façon que l'on a optimisé la taille des paquets pour maximiser le débit utile moyen en section 2.3.7, on cherche ici à optimiser les paramètres du système de manière à minimiser l'énergie ζ nécessaire par bit utile (maximiser l'efficacité énergétique). Plus précisément, on choisit ici de fixer les paramètres H et τ_o , et l'on s'intéresse à l'évolution de ζ en fonction de D et E_s .

Dans un premier temps, on fixe la taille d'un paquet de données à $D = 1000$ bits et on présente sur la figure 2.15 l'évolution de ζ en fonction du RSB $\gamma = E_s/N_o$, pour chaque protocole, et pour différentes valeurs du délai de propagation normalisé τ_o . Le paramètre H est fixé à 40 bits. On constate l'existence d'une énergie de transmission optimale dans chaque cas. Les protocoles SaW, SR et PSW ont le même optimum, indépendant de τ_o , et présentent la meilleure efficacité énergétique. L'énergie nécessaire par bit utile ζ est supérieure en GBN, et augmente avec τ_o . En effet, plus le délai de propagation est important, plus la fenêtre d'anticipation est grande, et plus le protocole GBN est pénalisé par la retransmission d'un nombre important de paquets en cas de réception d'un acquittement négatif. On peut noter également que lorsque l'énergie par symbole E_s est élevée (fort RSB), ζ est une fonction linéaire du RSB à D constant. En effet, chaque paquet transmis est correctement reçu avec une forte probabilité, et l'énergie totale consommée se réduit simplement à la somme de l'énergie consommée par le paquet et son acquittement, soit $E_s(D + 2H)$.

Sur la figure 2.16, nous présentons ensuite l'évolution de ζ en fonction de la taille D d'un paquet de données, pour un RSB fixé à $\gamma = 7$ dB, et pour différentes valeurs de τ_o . Comme précédemment il existe une taille optimale pour chaque protocole. En effet, augmenter D revient à augmenter la probabilité de perte p_o (à RSB constant), et donc par conséquent le nombre de retransmissions, ce qui accroît en retour ζ . Lorsque l'on diminue D l'énergie consommée par les entêtes et les acquittements devient non négligeable, ce qui contribue à augmenter ζ . De nouveau, les protocoles SaW, SR, et PSW surpassent le protocole GBN en terme d'efficacité énergétique.

Jusqu'à présent les paramètres E_s et D ont été optimisés séparément. A τ_o et H constants, on peut également chercher à les optimiser conjointement. Sur la figure 2.17, nous avons représenté l'évolution de $\exp(-\zeta)$ en fonction des deux variables (E_s , D) pour le protocole SaW, car il est plus facile de visualiser un maximum qu'un minimum

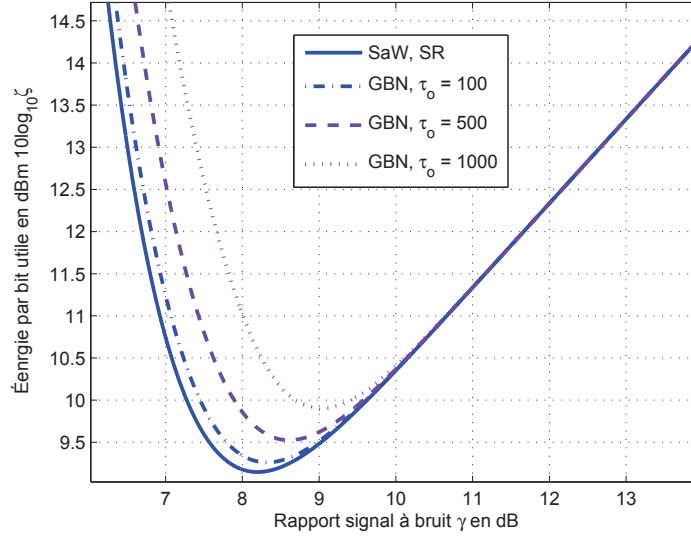


Figure 2.15 — Énergie ζ nécessaire par bit utile en fonction du rapport ($\gamma = E_s/N_o$) en dB pour une taille de données $D = 1000$ bits $H = 40$ bits pour différentes valeurs du temps de propagation ($T_p = 100T_s$, $T_p = 500T_s$, $T_p = 1000T_s$)

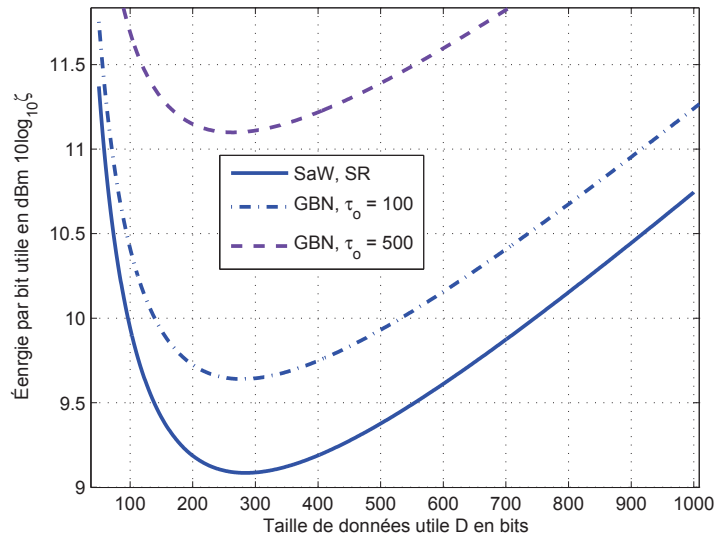


Figure 2.16 — Energie ζ nécessaire par bit utile en fonction de la taille D d'un paquet de données, pour $\gamma = 7$ dB, et pour différentes valeurs du temps de propagation $\tau_o = 100$ et $\tau_o = 500$)

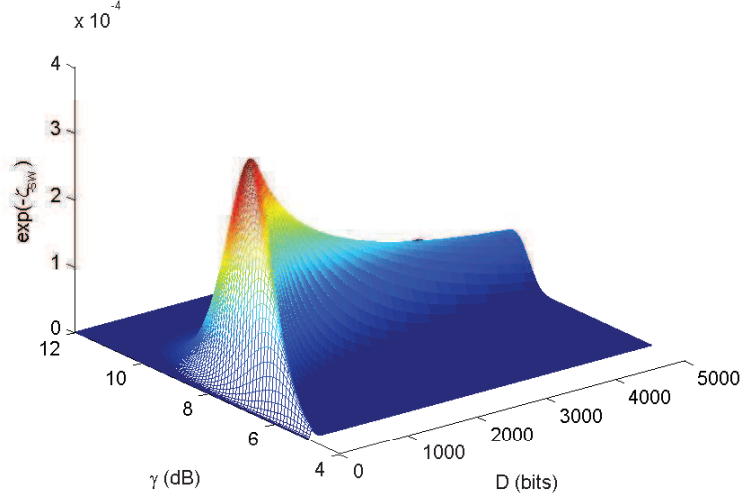


Figure 2.17 — Allure de la fonction $\exp(-\zeta_{sw})$ en fonction du rapport $\gamma = E_s/N_o$ et de la taille D des données utiles pour $H = 40$ bits.

sur une surface. L'énergie minimale correspond au pic visible sur la figure. On voit alors qu'il existe un minimum global unique atteint en une énergie E_s^{opt} et une taille de donnée D^{opt} optimales. Afin de comparer les différents protocoles entre eux, nous avons déterminé numériquement le couple optimal $(E_s^{\text{opt}}, D^{\text{opt}})$ qui minimise l'énergie ζ nécessaire par bit utile pour chaque protocole avec les hypothèses $H = 40$ bits et $\tau_o = 500$. Les résultats sont présentés dans le tableau 2.1. Outre l'énergie minimale $\zeta_{\min}$ atteinte, nous avons également indiqué le débit correspondant η . On voit que les protocoles SaW, SR et PSW présentent la meilleure efficacité énergétique, pour des paramètres optimaux identiques. En revanche, il n'ont pas le même débit à ce point de fonctionnement, l'avantage allant nettement aux protocoles SR/PSW. Le protocole GBN présente l'énergie minimale $\zeta_{\min}$ la plus élevée. En contrepartie, il offre également le meilleur débit à ce point de fonctionnement.

	SaW	GBN	SR	PSW
E_s^{opt}/N_o en dB	7,56	8,62	7,56	7,56
D^{opt} en bits	430	1035	430	430
$10 \log_{10} \zeta_{\min}$	7,96	8,96	7,96	7,96
η	0,24	0,84	0,78	0,78

Tableau 2.1 — Paramètres optimaux qui minimisent l'énergie moyenne nécessaire pour transmettre correctement un bit de données utiles pour un délai de propagation $\tau_o = 500$, et débit utile moyen.

2.5 Extension à une voie de retour imparfaite

Dans toutes les analyses précédentes, les calculs présentés supposaient une transmission avec voie de retour parfaite (les acquittements arrivent toujours correctement à la station A). En pratique, le canal de retour est bruité et peut introduire des erreurs sur les acquittements. Ces erreurs entraînent des confusions ou des pertes d'acquittements. On se focalise dans cette thèse sur un canal de retour avec perte d'acquittements. En effet, les acquittements sont supposés ici protégés de sorte que la station A puisse distinguer entre un acquittement correctement reçu et un acquittement erroné. Dans ce dernier cas l'acquittement est alors rejeté (perte d'acquittement).

Dans cette section, on présente tout d'abord le modèle à perte d'acquittements adopté pour modéliser les erreurs sur la voie de retour. On étudie ensuite l'impact de ces erreurs sur les principaux résultats développés dans les sections précédentes (analyse du débit et de l'efficacité énergétique).

Une manière de réduire le taux de perte sur la voie de retour consiste à augmenter la puissance d'émission sur ce lien. Cela nous amène à considérer le problème original de savoir comment répartir au mieux l'énergie de transmission entre la voie directe et la voie de retour sous la contrainte d'un budget total d'énergie constant.

2.5.1 Modèle et impact des erreurs sur la voie de retour

On modélise la voie de retour comme un canal à effacement par paquet où les pertes d'acquittements se produisent de manière i.i.d. avec une probabilité ϵ . Lorsqu'un acquittement est perdu, la station A le traite comme un acquittement négatif et effectue les retransmissions nécessaires. Ces retransmissions entraînent une perte en débit utile moyen dans le cas où les paquets retransmis avaient été correctement reçus par B.

Tous les résultats de calcul établis précédemment à partir de la probabilité conjointe d'un échec de transmission aux rounds 0 à m $P(m) = \Pr\{E_o, E_1, \dots, E_m\}$ restent valides pour le modèle de canal avec voie de retour imparfaite. Il suffit de préciser l'expression de $P(m)$, qui ne se réduit pas ici la probabilité conjointe d'un échec de décodage aux rounds 0 à m $p(m) = \Pr\{e_o, e_1, \dots, e_m\}$, car il faut également tenir compte des pertes d'acquittements. Plus précisément, on a un échec de transmission au round m lorsque le paquet est mal décodé par B et l'acquittement négatif bien reçu par A, ou bien lorsque l'acquittement (positif ou négatif) est perdu. D'où

$$\Pr\{E_m\} = \epsilon + (1 - \epsilon)p_o \quad (2.40)$$

Lorsque les pertes de paquets et les pertes d'acquittements sont supposées i.i.d. de round en round, les événements $\{E_m\}$ sont mutuellement indépendants, et la probabilité conjointe $P(m)$ s'écrit alors :

$$\begin{aligned} P(m) &= \prod_{i=0}^m \Pr\{E_i\} \\ &= (\epsilon + (1 - \epsilon)p_o)^{m+1} \end{aligned} \quad (2.41)$$

Il est important à noter qu'en pratique, les protocoles GBN et SR tolèrent des pertes d'acquittement. En effet, lorsque la taille de fenêtre d'anticipation est supérieure au temps d'aller-retour normalisé, il est possible qu'un acquittement soit perdu sans que le processus de retransmission ne se déclenche si l'acquittement du paquet suivant est bien reçu avant que la station A n'atteigne la fin de la fenêtre. La prise en compte de ce cas de figure, en toute rigueur, l'équation (2.40) n'est plus valide, un calcul exact de la probabilité $\Pr\{E_m\}$ devant alors faisant intervenir la probabilité de perte de plusieurs acquittements successifs. Par soucis de simplicité, ce cas de figure n'a pas été pris en compte dans la thèse. On suppose que la fenêtre d'anticipation est ajustée à la durée normalisée d'aller-retour, de sorte que lorsqu'un acquittement est mal décodé ou perdu, la station A commence à retransmettre immédiatement le ou les paquets non acquittés positivement. L'équation (2.40) s'applique alors.

En résumé, pour obtenir le débit utile ou l'efficacité énergétique des différents protocoles, il suffit de reprendre les expressions dérivées dans les sections 2.3 et 2.4 en remplaçant $P(m)$ par l'expression (2.41). En particulier, dans le cas des protocoles non tronqués, il suffit de remplacer p_o par $\epsilon + (1 - \epsilon)p_o$. Ainsi, à titre d'exemple, le débit utile moyen en SR non tronqué en présence de pertes sur la voie de retour s'écrit :

$$\eta_{sr} = R_o(1 - p_o)(1 - \epsilon) \quad (2.42)$$

et l'énergie nécessaire par bit utile devient :

$$\zeta_{sr} = E_s \frac{2 - \rho}{R_o} \frac{1}{(1 - p_o)(1 - \epsilon)} \quad (2.43)$$

La figure 2.18 montre l'impact des pertes d'acquittement sur le débit utile moyen des principaux protocoles ARQ non tronqués. Elle a été obtenue pour une transmission BPSK sur canal gaussien, pour une taille de paquet utile $D = 1000$ bits, une taille d'entête $H = 40$ bits, un délai de propagation normalisé $\tau = 0,5$, et pour deux valeurs de la probabilité de pertes d'acquittements $\epsilon = 0,1$ et $\epsilon = 0,3$. On constate que la présence de perte sur la voie de retour réduit le débit utile moyen des différents protocoles. Le protocole le plus affecté est le protocole GBN. Le débit utile pour ce protocole reste supérieur au débit en SaW mais la perte en débit causée par la mauvaise réception d'acquittement est plus élevée en GBN que pour les autres protocoles. En effet, une perte d'acquittement entraîne la retransmission de tous les paquets de la fenêtre d'anticipation, ce qui pénalise le débit utile moyen.

2.5.2 Répartition optimale de l'énergie entre la voie directe et la voie de retour

Nous venons de voir que les protocoles ARQ sont sensibles aux pertes sur la voie de retour. Une solution possible pour réduire ces pertes est d'augmenter l'énergie de transmission sur le lien de retour. Considérons un réseau de capteurs où chaque nœud dispose de ressources énergétiques propres et limitées. Dans ce contexte, on peut supposer que l'on a une énergie moyenne constante par round ARQ, et l'on peut s'interroger sur la meilleure manière de répartir l'énergie de transmission entre la voie directe et la

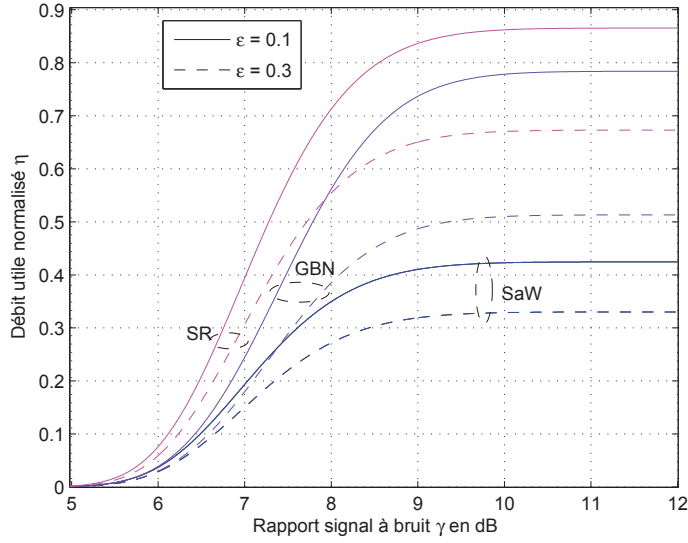


Figure 2.18 — Débit utile normalisé en fonction du rapport signal à bruit γ en ARQ pour $D = 1000$ bits, $H = 40$ bits, $\tau = 0,5$ et $\epsilon = 0,1$ et $0,3$.

voie de retour, au sens d'un critère à définir (maximisation du débit ou maximisation de l'efficacité énergétique par exemple).

Désignons par E_s l'énergie transmise par symbole sur le lien direct et par E_a l'énergie transmise par symbole sur le lien retour. L'énergie totale E_{round} est la somme entre l'énergie consommée par un paquet et l'énergie consommée par un acquittement :

$$E_{\text{round}} = NE_s + (1 - \rho)NE_a \quad (2.44)$$

Désignons par α la proportion de l'énergie totale E_{round} allouée à la transmission de l'acquittement :

$$\alpha = \frac{(1 - \rho)NE_a}{E_{\text{round}}} \quad (2.45)$$

$\alpha = 0$ correspond au cas limite où toute la puissance est portée sur le lien direct (acquittement jamais transmis). À l'inverse, lorsque $\alpha = 1$, toute la puissance est allouée au lien retour (ce qui n'a pas de sens pratique).

On définit par ailleurs l'énergie moyenne globale par symbole E_g comme le rapport entre l'énergie totale consommée par round et le nombre total de symboles transmis durant le round (paquet + acquittement).

$$E_g = \frac{E_{\text{round}}}{N(2 - \rho)} \quad (2.46)$$

On peut alors relier les énergies de transmission sur la voie directe (E_s) et sur la voie de retour (E_a) aux paramètres α et E_g par :

$$E_s = (1 - \alpha)(2 - \rho)E_g \quad (2.47)$$

$$E_a = \alpha \frac{2 - \rho}{1 - \rho} E_g \quad (2.48)$$

On se place dans le cadre d'une transmission BPSK sur canal BABG. La probabilité de perte d'un paquet sur le lien direct s'écrit alors :

$$p_o = 1 - \left(1 - \frac{1}{2} \operatorname{erfc} \sqrt{\frac{E_s}{N_o}} \right)^{(D+H)} \quad (2.49)$$

De la même façon la probabilité de perte d'acquittement sur le lien retour s'écrit :

$$\epsilon = 1 - \left(1 - \frac{1}{2} \operatorname{erfc} \sqrt{\frac{E_a}{N_o}} \right)^H \quad (2.50)$$

Avec ces hypothèses, le débit utile de la transmission, de même que l'énergie nécessaire ζ par bit utile correctement reçu, peuvent s'écrire comme des fonctions des paramètres E_g , α , D et H par l'intermédiaire des équations (2.47) et (2.48).

Dans un premier temps, on cherche à déterminer la répartition optimale α de l'énergie entre la voie directe et la voie de retour qui maximise le débit utile moyen. Pour cela, on fixe $D = 1000$ bits, $H = 40$ bits, $\tau = 0,5$ et $E_g/N_o = 7$ dB. La figure 2.19 présente l'évolution du débit utile moyen en fonction du paramètre α pour les différents protocoles. On constate qu'à τ constant, il existe une valeur optimale de α qui maximise le débit et que cette valeur est indépendante du protocole considéré. Dans l'exemple présenté, le maximum est obtenu pour $\alpha = 0,0372$, ce qui correspond à des énergies par symbole sur le lien direct $E_s/N_o = 6,99$ dB et $E_a/N_o = 7,02$ dB sur la voie de retour. D'une manière générale, on peut montrer que cet optimum correspond à la valeur qui annule la dérivée de p_o par rapport à α . En résumé, pour le critère de la maximisation du débit utile moyen, la répartition optimale de l'énergie est la même pour tous les protocoles ARQ.

Nous avons ensuite recherché la répartition de puissance entre voie directe et voie de retour optimale au sens de la maximisation de l'efficacité énergétique. Pour D , H , τ et E_g fixés, ceci revient à chercher la valeur de α qui minimise l'énergie ζ nécessaire par bit correctement reçu. Cette énergie est donnée par les équations (2.36) et (2.39) en remplaçant p_o par $\epsilon + (1 - \epsilon)p_o$. Le tableau 2.2 présente les valeurs optimales de α pour chaque protocole ainsi que l'énergie ζ_{min} correspondante. Ces valeurs ont été obtenues pour un RSB global $E_g/N_o = 7$ dB, une taille de paquet de données $D = 1000$ bits, une taille d'acquittement $H = 40$ bits et un délai de propagation $\tau = 0,5$. On remarque que la répartition optimale α dépend du protocole ARQ considéré. En particulier, l'énergie allouée à la voie de retour est supérieure en GBN, ce qui est cohérent avec le fait qu'une perte sur la voie de retour coûte cher en terme d'énergie pour ce protocole (retransmission de toute la fenêtre).

2.6 Conclusion

Dans ce chapitre, nous avons présenté les principaux protocoles ARQ, et revisité de manière unifiée leurs performances selon le critère du débit utile moyen. Nous avons vu en particulier que le protocole SR offre le meilleur débit utile, au prix d'une complexité

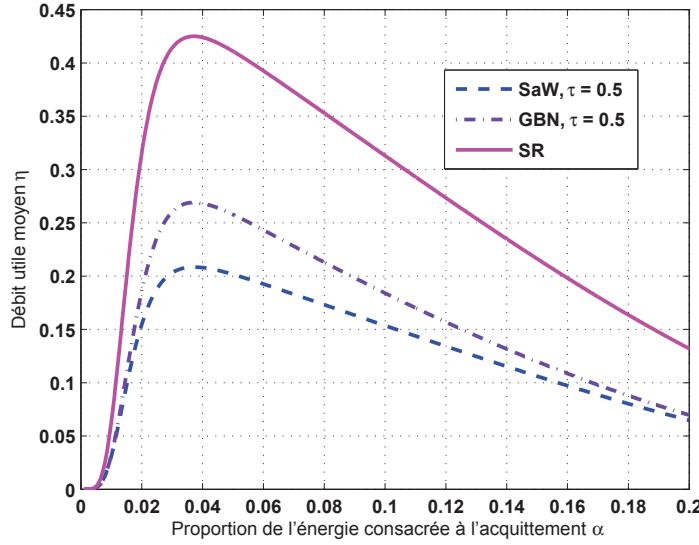


Figure 2.19 — Débit utile normalisé en fonction de α en ARQ pour $D = 1000$ bits, $H = 40$ bits, $\tau = 0,5$ et $E_g/N_o = 7$ dB.

	SaW, SR, PSW	GBN
α_{opt}	0,0389	0,0399
E_s^{opt}/N_o dB	6,992	6,897
E_a^{opt}/N_o dB	7,21	7,323
$10 \log_{10}(\zeta_{\min})$	10,87	12,45
ϵ	$2 \cdot 10^{-2}$	$2,3 \cdot 10^{-2}$

Tableau 2.2 — Paramètres optimaux qui minimisent l'énergie moyenne nécessaire pour transmettre correctement un bit de données utiles pour un délai de propagation $\tau = 0,5$ pour une transmission avec voie de retour imparfaite et pour $E_g/N_o = 7$

(mémoire et traitement) supérieure à celle des autres protocoles. Le protocole PSW offre un débit comparable à celui du protocole SR, mais avec des règles plus simples. Les protocoles GBN et SaW sont également moins complexes, mais sensibles au délai de propagation. A fort RSB, le débit utile en GBN rejoint celui du SR. Nous avons également discuté de l'optimisation des paramètres du système de manière à maximiser le débit utile pour chaque protocole.

Nous avons également comparé les différents protocoles en terme d'efficacité énergétique, qui mesure l'énergie minimale requise pour transmettre un bit de donnée utile avec succès. De cette étude, il ressort que le protocole GBN fait la moins bonne utilisation de l'énergie à sa disposition. Nous avons également discuté de l'optimisation des paramètres (taille de paquet et énergie d'émission) afin de maximiser l'efficacité énergétique.

Les études précédentes supposaient une voie de retour parfaite. Nous les avons ensuite étendues au cas plus réaliste d'une voie de retour avec pertes d'acquittements. Nous avons montré en particulier que le protocole GBN était le plus sensible aux erreurs

sur la voie de retour. Nous nous sommes également penchés sur le problème original de savoir comment répartir au mieux l'énergie à l'émission entre le lien direct et la voie de retour, sous l'hypothèse d'un budget total d'énergie constant par round ARQ. Les conclusions montrent que suivant le critère d'optimalité considéré, la répartition optimale n'est pas nécessairement dépendante du protocole.

Les protocoles ARQ génèrent un nombre important de retransmissions lorsque le bruit est élevé. Ceci diminue le débit utile de transmission et augmente les délais de transfert. Afin de réduire le nombre moyen de retransmissions, on peut envisager de protéger les paquets de données par un code correcteur d'erreurs, de manière à créer un canal équivalent plus fiable en sortie du décodeur. Les différentes techniques de combinaison ARQ/FEC font l'objet du chapitre suivant.

CHAPITRE 3

Combinaison du codage et de la retransmission (HARQ)

3.1 Introduction

Les protocoles ARQ offrent une solution simple et flexible pour garantir l'intégrité des données au niveau des couches hautes, en adaptant à la demande la redondance (nombre de retransmissions) à l'état du canal. C'est une solution particulièrement performante lorsque le canal est bon la majorité du temps. En revanche, lorsque le taux de perte est important sur la couche physique, le nombre de retransmissions augmente considérablement et le débit utile moyen diminue. À l'inverse, le codage canal FEC est une stratégie puissante pour la couche physique. Il peut garantir des transmissions fiables même sur un canal relativement mauvais. Par contre, il présente un risque de sur-protection des données (baisse du débit utile moyen) lorsque le canal est bon, car il est généralement dimensionné au pire cas.

La combinaison codage et retransmission (*Hybrid ARQ*, ou HARQ) réalise un compromis naturel entre les deux extrêmes précédents. Elle vise à bénéficier à la fois de la flexibilité des protocoles ARQ et du pouvoir de correction d'erreurs du codage FEC. L'optimisation de cette combinaison nécessite de considérer à la fois des aspects protocolaires (couche RLC) et des fonctions de communication numérique (couche physique), et donc une optimisation multi-couches (*cross-layer*). Des études récentes multi-couches de l'analyse des performances en HARQ intégrant les couches liaison de données et réseaux ont été présentées en [34][35].

Dans ce chapitre, on présente le principe des schémas HARQ. On développe ensuite l'expression des performances des principaux schémas HARQ, et on les compare entre eux en terme de débit utile moyen et d'efficacité énergétique, pour différents modèles de codes et de canaux. Dans un premier temps, comme dans le chapitre précédent, le canal de retour est supposé sans erreurs. Les résultats et les comparaisons sont ensuite étendus à une voie de retour avec perte d'acquitements.

3.2 Principe de l'ARQ hybride

Le codage ARQ hybride (HARQ) consiste à combiner un code correcteur d'erreurs (code de la couche physique) avec un protocole ARQ. L'idée est d'utiliser le code pour corriger les motifs d'erreur les plus fréquents, et de laisser la correction des motifs plus rares (motifs imprévus) à la charge du protocole de retransmission. En pratique, les données utiles sont tout d'abord protégées par un code détecteur d'erreurs, qui rajoute une somme de contrôle (CRC) à la trame d'information. Celle-ci est ensuite encodée par un deuxième code, le code correcteur d'erreurs (FEC). Le paquet doublement encodé est ensuite transmis sur le canal bruité. Le paquet reçu est tout d'abord décodé par le décodeur FEC. L'intégrité du message décodé est ensuite vérifiée par contrôle du CRC. En cas d'erreurs détectées, le récepteur rejette ou conserve le paquet reçu, suivant la stratégie HARQ considérée, puis demande une retransmission à la station A. Dans le cas contraire (message décodé valide), les données utiles sont délivrées à la couche supérieure.

On distingue deux grandes familles de protocoles ARQ hybrides. Dans la première appelée ARQ *hybride type-I* (HARQ-I), les paquets erronés sont rejetés [17] [28]. Dans la seconde appelée ARQ *hybride type-II* (HARQ-II), les paquets erronés sont mémorisés à chaque round, puis réutilisés par le décodeur FEC pour augmenter les chances de succès de décodage aux rounds suivants [28] [27].

3.2.1 Mise en œuvre de l'HARQ type I

Le fonctionnement du schéma HARQ type-I (HARQ-I) est similaire au fonctionnement des protocoles ARQ. La station B décode le paquet reçu. Si le décodage est correct, elle envoie un acquittement positif pour informer l'émetteur (station A) de la bonne réception du paquet. Si le décodage échoue, elle rejette le paquet erroné et renvoie une demande de retransmission (acquittement négatif). Au round suivant, la retransmission est une copie à l'identique du paquet erroné. La station B tente à nouveau de décoder l'observation reçue. Cette procédure de retransmission se répète jusqu'à ce que le paquet soit correctement décodé, ou que le nombre maximal M_1 de retransmissions autorisées soit atteint.

L'HARQ type-I est habituellement mis en œuvre en cascasant (concaténation série) un bon code détecteur d'erreurs de type CRC (code externe), avec un bon schéma de codage et modulation (code interne), comme indiqué sur le schéma général du chapitre 1 (cf. figure 1.1). L'optimisation du schéma HARQ-I s'en trouve ainsi simplifiée, puisqu'il suffit d'optimiser séparément chacun des deux codes (détecteur/correcteur). De même, les traitements en réception s'en trouvent simplifiés, puisque la détection d'erreurs et le décodage sont réalisés de manière disjointe, l'une après l'autre.

En théorie, on pourrait n'utiliser qu'un seul code pour réaliser ces deux fonctions. Considérons le cas particulier d'un code en blocs linéaire de pouvoir de correction t , muni d'un décodeur algébrique (code BCH ou RS par exemple). On peut utiliser un calcul de syndrome pour détecter la présence d'erreurs et évaluer leur nombre. Si ce dernier est inférieur ou égal au pouvoir de correction du code, on procède au décodage

du paquet reçu. Dans le cas contraire, on demande une retransmission. L'utilisation d'un code unique présente l'avantage de nécessiter moins de redondance qu'une approche disjointe. Cela requiert en contrepartie la conception de codes qui soient à la fois très performants en terme de correction et de détection, un problème d'autant plus délicat que l'on cherche à opérer au plus proche des limites de correction établies par la théorie de l'information (limite de Shannon). Les codes LDPC offrent une solution plus attractive dans ce contexte.

3.2.2 Mise en œuvre du schéma HARQ type II

Les protocoles HARQ type-II fonctionnent selon le même principe général que les protocoles HARQ-I, excepté que les paquets décodés en erreur ne sont plus rejetés mais au contraire mémorisés par le récepteur (station B) pour aider au décodage lors des rounds suivants. On distingue deux variantes principales en HARQ-II suivant la stratégie de retransmission adoptée.

HARQ-II *Chase Combining*

Dans la première variante appelée *Chase Combining* (HARQ-II-CC), les paquets retransmis sont identiques au paquet initial, comme en HARQ-I. Mais à la différence du HARQ-I, à chaque round, le récepteur combine toutes les copies reçues jusqu'à ce round pour procéder à une nouvelle tentative de décodage. La règle de combinaison optimale a été dérivée dans un cadre très général par Chase dans [36]. Dans le cas particulier d'un canal sans mémoire avec bruit additif blanc gaussien, on montre qu'elle se réduit à une combinaison à gain maximal (MRC *Maximum Ratio Combining*).

HARQ-II *Incremental Redundancy*

Dans la seconde approche appelée *Incremental Redundancy* (HARQ-II-IR), les retransmissions sont constituées de paquets distincts du paquet initial, chaque nouveau paquet apportant de la redondance supplémentaire (redondance incrémentale). Il se construit donc ainsi en réception un code correcteur de plus en plus puissant au fur et à mesure des retransmissions. En pratique, les paquets de redondance sont habituellement construits par poinçonnage compatible en rendement [37] [38] d'un code ou schéma de modulation et codage (MCS) à faible rendement, appelé code mère $\mathcal{C}$. À chaque round, la station A transmet un fragment différent du mot de code initial produit par $\mathcal{C}$. Dans l'exemple de la figure 3.1, le mot de code $\mathbf{x}$ délivré par $\mathcal{C}$ est ainsi divisé en $M_1 + 1$ mots de code partiels $(\mathbf{x}_0, \mathbf{x}_1, \dots, \mathbf{x}_{M_1})$, supposés ici de même taille. Au round m , la station A transmet le mot partiel $\mathbf{x}_m$. Le code équivalent vu par la station B est alors noté $\mathcal{C}_m$, et se compose de l'ensemble des mots de la forme $(\mathbf{x}_0, \mathbf{x}_1, \dots, \mathbf{x}_m)$. Son rendement s'écrit $R_m = \frac{R_o}{m+1}$, où R_o désigne le rendement de la transmission côté émetteur. Notons que l'on retrouve le code mère $\mathcal{C}_{M_1} = \mathcal{C}$ au dernier round.

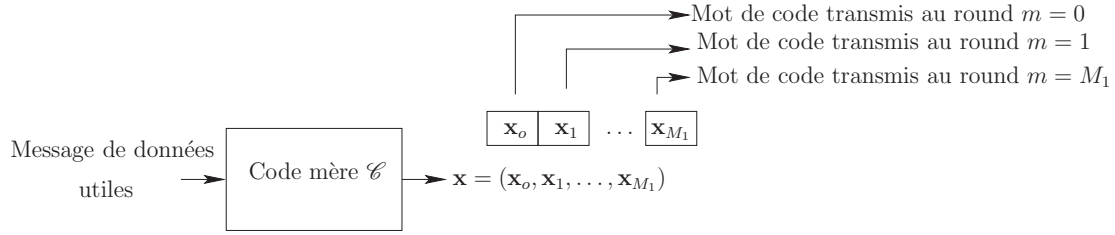


Figure 3.1 — Construction de paquets retransmis en IR par poinçonnage d'un code mère à faible rendement

3.3 Décodage optimal en HARQ-II

En HARQ-I, les paquets erronés sont rejetés et la station B procède au décodage du paquet reçu à chaque round sans tenir compte des observations précédentes. En HARQ-II, en revanche, les paquets erronés sont mémorisés et réutilisés dans le processus de décodage. L'objectif de cette section est d'établir la règle de décodage optimal en HARQ-II. Il s'agit d'un prérequis nécessaire pour pouvoir ensuite évaluer les performances des systèmes HARQ-II. Deux cas de figure sont à distinguer selon que l'on retransmette le même mot de code (HARQ-II-CC) ou bien des mots de code différents (HARQ-II-IR) à chaque round.

3.3.1 Transmissions multiples d'un même mot de code

On suppose ici que la station A renvoie le même mot de code $\mathbf{x}$ à chaque round. Notons $\mathbf{y}_m$ le signal reçu au round m et h_m le coefficient du canal affectant ce signal. Les signaux observés aux rounds 0 à m sont modélisés par les vecteurs :

$$\begin{aligned} \mathbf{y}_0 &= h_0 \mathbf{x} + \mathbf{w}_0 \\ \mathbf{y}_1 &= h_1 \mathbf{x} + \mathbf{w}_1 \\ &\vdots \\ \mathbf{y}_m &= h_m \mathbf{x} + \mathbf{w}_m \end{aligned} \tag{3.1}$$

Les réalisations $\{h_m\}$ du canal sont supposées indépendantes entre elles et indépendantes des réalisations $\{\mathbf{w}_m\}$ du bruit. Les vecteurs de bruit sont des vecteurs gaussiens complexes centrés et circulaires, de matrice de covariance $\mathbf{K}_w = N_o \mathbf{I}_N$, mutuellement indépendants de round en round.

La règle de décodage optimal qui minimise la probabilité d'erreur après décodage est la règle du maximum de vraisemblance (MV).

$$\hat{\mathbf{x}} = \arg \max_{\mathbf{x} \in \mathcal{C}} Pr\{\mathbf{y}_0, \dots, \mathbf{y}_m | \mathbf{x}, h_0, \dots, h_m\} \tag{3.2}$$

En utilisant la règle de Bayes, il vient :

$$\begin{aligned} Pr\{\mathbf{y}_0, \dots, \mathbf{y}_m | \mathbf{x}, h_0, \dots, h_m\} &= Pr\{\mathbf{y}_0 | \mathbf{x}, h_0, \dots, h_m\} \times Pr\{\mathbf{y}_1 | \mathbf{x}, \mathbf{y}_0, h_0, \dots, h_m\} \times \dots \\ &\quad \times Pr\{\mathbf{y}_m | \mathbf{x}, \mathbf{y}_0, \dots, \mathbf{y}_{m-1}, h_0, \dots, h_m\} \end{aligned} \tag{3.3}$$

Par l'indépendance des réalisations du bruit et du canal entre elles ainsi que de round en round, on a :

$$\Pr\{\mathbf{y}_j|\mathbf{x}, \mathbf{y}_o, \mathbf{y}_1, \dots, \mathbf{y}_{j-1}, h_o, h_1, \dots, h_m\} = \Pr\{\mathbf{y}_j|\mathbf{x}, h_j\} \quad (3.4)$$

L'équation (3.3) devient :

$$\Pr\{\mathbf{y}_o, \mathbf{y}_1, \dots, \mathbf{y}_m|\mathbf{x}, h_o, h_1, \dots, h_m\} = \prod_{j=0}^m \Pr\{\mathbf{y}_j|\mathbf{x}, h_j\} \quad (3.5)$$

Or par hypothèse, $\mathbf{w}_m \sim \mathcal{CN}(\mathbf{0}, N_o \mathbf{I}_N)$, d'où :

$$\Pr\{\mathbf{y}_j|\mathbf{x}, h_j\} = \frac{1}{(\pi N_o)^{N/2}} \exp\left(-\frac{\|\mathbf{y}_j - h_j \mathbf{x}\|^2}{N_o}\right) \quad (3.6)$$

Dans ce contexte, la règle de décodage devient alors :

$$\hat{\mathbf{x}} = \arg \min_{\mathbf{x} \in \mathcal{C}} \sum_{j=0}^m \|\mathbf{y}_j - h_j \mathbf{x}\|^2 \quad (3.7)$$

La métrique précédente peut encore s'écrire :

$$\sum_{j=0}^m \|\mathbf{y}_j - h_j \mathbf{x}\|^2 = \sum_{j=0}^m \|\mathbf{y}_j\|^2 + \|\mathbf{x}\|^2 \sum_{j=0}^m |h_j|^2 - 2 \sum_{j=0}^m \Re\{\langle \mathbf{y}_j, h_j \mathbf{x} \rangle\} \quad (3.8)$$

où $\Re(\cdot)$ désigne la partie réelle, et l'opérateur $\langle \mathbf{x}, \mathbf{y} \rangle = \sum_{j=1}^N x_j y_j^*$ désigne le produit scalaire entre les vecteurs complexes $\mathbf{x}$ et $\mathbf{y}$, et où $(\cdot)^*$ désigne le conjugué. Le facteur $\sum_{j=1}^m \|\mathbf{y}_j\|^2$ est commun à tous les mots de code et peut donc être omis. On en déduit une forme équivalente de la règle de décodage optimal :

$$\hat{\mathbf{x}} = \arg \max_{\mathbf{x} \in \mathcal{C}} \sum_{j=0}^m \Re\{\langle \mathbf{y}_j, h_j \mathbf{x} \rangle\} - \frac{\|\mathbf{x}\|^2}{2} \sum_{j=0}^m |h_j|^2 \quad (3.9)$$

$$= \arg \max_{\mathbf{x} \in \mathcal{C}} \sum_{j=0}^m \Re\{\langle h_j^* \mathbf{y}_j, \mathbf{x} \rangle\} - \frac{\|\mathbf{x}\|^2}{2} \sum_{j=0}^m |h_j|^2 \quad (3.10)$$

On remarque que la combinaison $\mathbf{y} = \sum_{j=0}^m h_j^* \mathbf{y}_j$ constitue une statistique suffisante pour le décodage optimal de $\mathbf{x}$ au round m . On reconnaît ici une combinaison à gain maximal (MRC) corrigée d'un terme correspondant à l'énergie du mot de code. Pour des mots de codes à énergie constante (modulation de phase par exemple), la règle optimale se simplifie en :

$$\hat{\mathbf{x}} = \arg \max_{\mathbf{x} \in \mathcal{C}} \Re\left\{\left\langle \sum_{j=0}^m h_j^* \mathbf{y}_j, \mathbf{x} \right\rangle\right\} \quad (3.11)$$

Après combinaison des copies du signal transmis, tout se passe comme si on n'avait réalisé qu'une seule transmission sur un canal équivalent de gain $h = \sum_{j=0}^m |h_j|^2$:

$$\mathbf{y} = (|h_o|^2 + |h_1|^2 + \dots + |h_m|^2) \mathbf{x} + \sum_{i=0}^m h_i^* \mathbf{w}_i \quad (3.12)$$

Le rapport signal à bruit sur le canal équivalent vu au round m s'écrit :

$$\begin{aligned}\gamma_m &= \frac{\left(\sum_{j=0}^m |h_j|^2\right)^2 E_s}{\left(\sum_{j=0}^m |h_j|^2\right) N_o} \\ &= \sum_{j=0}^m \frac{|h_j|^2 E_s}{N_o}\end{aligned}\quad (3.13)$$

où $\frac{|h_j|^2 E_s}{N_o}$ est le rapport signal à bruit vu lors de la transmission j . On voit donc que le rapport signal à bruit à l'entrée du décodeur augmente de manière monotone au fur et à mesure des retransmissions. La combinaison de paquets en HARQ-II-CC apporte un gain en RSB. Notons par ailleurs que le décodage optimal utilise ici le même décodeur qu'en HARQ-I (même complexité), mais alimenté ici par la combinaison $\mathbf{y}$ des copies, et tenant compte du gain global h .

3.3.2 Transmission de paquets de redondance distincts

En HARQ-II-IR, les paquets transmis à chaque round sont différents (redondance incrémentale). Soient $\mathbf{x}_o, \mathbf{x}_1, \dots, \mathbf{x}_m$ les signaux transmis aux rounds 0 à m et $\mathbf{y}_o, \mathbf{y}_1, \dots, \mathbf{y}_m$ les signaux observés par le récepteur :

$$\begin{aligned}\mathbf{y}_0 &= h_0 \mathbf{x}_0 + \mathbf{w}_0 \\ \mathbf{y}_1 &= h_1 \mathbf{x}_1 + \mathbf{w}_1 \\ &\vdots \\ \mathbf{y}_m &= h_m \mathbf{x}_m + \mathbf{w}_m\end{aligned}\quad (3.14)$$

La règle de décodage optimal au sens de la minimisation de la probabilité d'erreur de décodage (règle MV) au round m s'écrit :

$$\hat{\mathbf{x}} = \arg \max_{\mathbf{x}=(\mathbf{x}_o, \mathbf{x}_1, \dots, \mathbf{x}_m) \in \mathcal{C}_m} \Pr\{\mathbf{y}_o, \mathbf{y}_1, \dots, \mathbf{y}_m | \mathbf{x}, h_o, h_1, \dots, h_m\} \quad (3.15)$$

où $\mathcal{C}_m$ désigne le code équivalent vu par la station B au round m . Autrement dit, on procède à un décodage conjoint des paquets $\mathbf{x}_o, \mathbf{x}_1, \dots, \mathbf{x}_m$. En supposant que les réalisations du canal et du bruit sont indépendantes entre elles et de round en round, et en supposant par ailleurs que $\mathbf{w}_m \sim \mathcal{CN}(\mathbf{0}, N_o \mathbf{I}_N)$, la règle de décodage optimal devient :

$$\hat{\mathbf{x}} = \arg \min_{\mathbf{x}=(\mathbf{x}_o, \mathbf{x}_1, \dots, \mathbf{x}_m) \in \mathcal{C}_m} \sum_{i=0}^m \|\mathbf{y}_i - h_i \mathbf{x}_i\|^2 \quad (3.16)$$

$$= \arg \max_{\mathbf{x}=(\mathbf{x}_o, \mathbf{x}_1, \dots, \mathbf{x}_m) \in \mathcal{C}_m} \sum_{i=0}^m \left(\Re\{< h_i^* \mathbf{y}_i, \mathbf{x}_i >\} - \frac{|h_i|^2}{2} \|\mathbf{x}_i\|^2 \right) \quad (3.17)$$

On voit à nouveau que les quantités $h_i^* \mathbf{y}_i$ forment une statistique suffisante pour la détection. Mais contrairement au *Chase Combining*, la règle de décodage ne se simplifie pas davantage dans le cas général.

3.4 Calcul et optimisation du débit utile moyen en HARQ

L'objectif de cette section est d'établir l'expression générique du débit utile en HARQ type-I et type II, afin de comparer par la suite les performances de ces différents schémas. Dans le but d'estimer les performances limites de chaque système, les performances sont tout d'abord données pour le cas idéal de codes gaussiens de longueur infinie. Pour s'approcher d'un scénario plus réaliste, on présente ensuite les performances pour des codes optimaux de longueur finie. C'est l'une des contributions originales de cette thèse. On termine en présentant les performances de quelques schémas de codage utilisés en pratique. On discute également de l'optimisation du débit utile moyen.

Hypothèses clés

On s'intéresse à un schéma HARQ basé sur un protocole ARQ de type SR ou PSW (avec un nombre de processus SaW parallèles ajusté à la durée d'aller-retour). En effet, ces protocoles présentent les meilleures performances en terme de débit utile et d'efficacité énergétique (cf. chapitre II). L'extension à des schémas HARQ basés sur des protocoles ARQ de type SaW ou GBN est immédiate en partant des formules adéquates présentées dans le chapitre précédent.

Dans ce cas, l'expression générale du débit utile moyen est identique à celle établie pour le protocole SR, dans l'hypothèse où les paquets transmis sont de même taille à chaque round (toujours vrai en HARQ-I, HARQ-II-CC, mais pas nécessairement en HARQ-II-IR), hypothèse retenue ici. Désignons par E_m l'événement "échec de transmission au round m ", et par $P(m) = \Pr\{E_o, E_1, \dots, E_m\}$ la probabilité conjointe d'avoir $m+1$ échecs de transmissions successifs aux rounds 0 à m . De la même façon, désignons par e_m l'événement "échec de décodage au round m ", et par $p(m) = \Pr\{e_o, e_1, \dots, e_m\}$ la probabilité conjointe d'avoir $m+1$ échecs successifs de décodage aux rounds 0 à m . Le débit utile moyen s'écrit alors :

$$\eta_{HARQ} = R_o \frac{1 - p(M_1)}{1 + \sum_{m=0}^{M_1-1} P(m)} \quad (3.18)$$

Dans un premier temps, on se place ici dans l'hypothèse d'une voie de retour parfaite. L'extension à une transmission avec perte d'acquittements sera traitée en section 3.6. Dans ce cas, un échec de transmission est systématiquement dû à un échec de décodage ($E_m = e_m$), et l'on a la simplification importante :

$$P(m) = p(m) = \Pr\{e_o, e_1, \dots, e_m\} \quad (3.19)$$

Pour aller plus loin, il nous faut maintenant spécifier le calcul de $p(m)$, qui dépend du protocole HARQ considéré.

3.4.1 Calcul du débit utile moyen en HARQ-I

Les schémas HARQ-I ne mémorisent pas les paquets erronés. On peut donc les analyser de la même manière que les protocoles ARQ sur lesquels ils se basent. En particulier, du fait du rejet des paquets erronés les événements $\{e_m\}$ sont mutuellement indépendants, et la probabilité conjointe $p(m)$ se décompose en produit des probabilités marginales :

$$p(m) = \Pr\{e_o, e_1, \dots, e_m\} \quad (3.20)$$

$$= \prod_{i=0}^m \Pr\{e_i\} \quad (3.21)$$

$$= p_o^{m+1} \quad (3.22)$$

où p_o désigne la probabilité d'erreur en sortie du décodeur FEC (identique à tous les rounds). Cette dernière dépend du modèle de canal et de code considéré (cf. chapitre I). Le débit utile moyen en HARQ-I s'en déduit à partir de (3.18) et (3.19).

3.4.2 Calcul du débit utile moyen en HARQ-II

Le calcul du débit utile en HARQ-II est plus compliqué car, du fait de la mémorisation et de la réutilisation de round en round des paquets erronés, les événements $\{e_m\}$ ne sont pas mutuellement indépendants. Le calcul de la probabilité conjointe $p(m) = \Pr\{e_o, e_1, \dots, e_m\}$ d'un échec de décodage successif aux round 0 à m doit prendre en compte ce point, et s'avère alors très compliqué à mener d'une manière exacte. En revanche, on peut calculer relativement facilement la probabilité marginale $\Pr\{e_m\}$ d'un échec de décodage au round m , et s'en servir pour obtenir un encadrement relativement fin du débit utile moyen.

Expression de la probabilité d'erreur de décodage au round m

D'une manière générale, la probabilité $\Pr\{e_m\}$ dépend du modèle du canal, du code $\mathcal{C}$ considéré, et de la stratégie de retransmission HARQ-II adoptée (CC ou IR).

En CC, nous avons montré en section 3.3.1 que le paquet résultant de la combinaison des $m+1$ copies au round m est équivalent à la transmission du paquet initial sur un canal avec un RSB plus important d'un facteur $m+1$. Autrement dit, si $P_{\mathcal{C}}(\gamma)$ désigne la probabilité d'erreur pour le code $\mathcal{C}$ sur un canal BABG de RSB $\gamma = \frac{E_s}{N_o}$, on a :

$$\Pr\{e_m\} = P_{\mathcal{C}}(\gamma_m) = P_{\mathcal{C}}\left(\gamma \sum_{j=0}^m |h_j|^2\right) \quad (3.23)$$

En IR, d'une manière générale, le code équivalent $\mathcal{C}_m$ vu par la station B change à chaque round. La probabilité $\Pr\{e_m\}$, prend donc une expression différente à chaque round, en fonction du code mère et du masque de poinçonnage considéré.

Encadrement et approximation du débit utile en HARQ-II

On présente dans cette section un encadrement du débit utile moyen. Comme le débit est fonction des probabilités d'échecs conjoints $p(m)$, on commence par encadrer ces probabilités.

Remarquons d'abord que $p(m) \leq \Pr\{e_j\} \forall j = 0, 1, \dots, m$. En effet :

$$p(m) = \Pr\{e_o, e_1, \dots, e_j\} \quad (3.24)$$

$$= \Pr\{e_o, e_1, \dots, e_{j-1}, e_{j+1}, \dots, e_m | e_j\} \Pr\{e_j\} \quad (3.25)$$

$$\leq \Pr\{e_j\} \quad (3.26)$$

Par conséquent,

$$p(m) \leq \min_{0 \leq j \leq m} \Pr\{e_j\} \quad (3.27)$$

D'une manière générale, en HARQ-II, la probabilité d'erreur décroît à chaque round, on peut donc majorer $p(m)$ par :

$$p(m) \leq \Pr\{e_m\} \quad (3.28)$$

D'autre part, il est raisonnable de supposer qu'un échec de décodage au round m est d'autant plus probable que l'on sait qu'il y a eu un échec de décodage aux rounds précédents [39].

$$\Pr\{e_m | e_{m-1}, \dots, e_o\} \geq \Pr\{e_m\} \quad (3.29)$$

Partant de cette hypothèse et en utilisant la règle de Bayes, il s'ensuit que :

$$\Pr\{e_m, e_{m-1}, \dots, e_o\} \geq \Pr\{e_m\} \Pr\{e_{m-1}, \dots, e_o\} \quad (3.30)$$

On aboutit à la minoration suivante de $p(m)$:

$$p(m) = \Pr\{e_m, e_{m-1}, \dots, e_o\} \geq \prod_{i=0}^m \Pr\{e_i\}, \quad (3.31)$$

et donc à l'encadrement final,

$$\prod_{i=0}^m \Pr\{e_i\} \leq p(m) \leq \Pr\{e_m\} \quad (3.32)$$

On en déduit l'encadrement suivant du débit utile :

$$\eta_{inf} \leq \eta \leq \eta_{sup} \quad (3.33)$$

où η_{inf} est calculé en considérant que $p(m) = \Pr\{e_m\}$, et η_{sup} est calculé en considérant que $p(m) = \prod_{i=0}^m \Pr\{e_i\}$.

La figure 3.2 compare les bornes inférieure et supérieure sur η avec la valeur exacte du débit, obtenue ici par simulation, dans le cas d'une transmission HARQ-II-CC sur canal BABG, avec modulation BPSK et codage convolutif de rendement 1/2 et

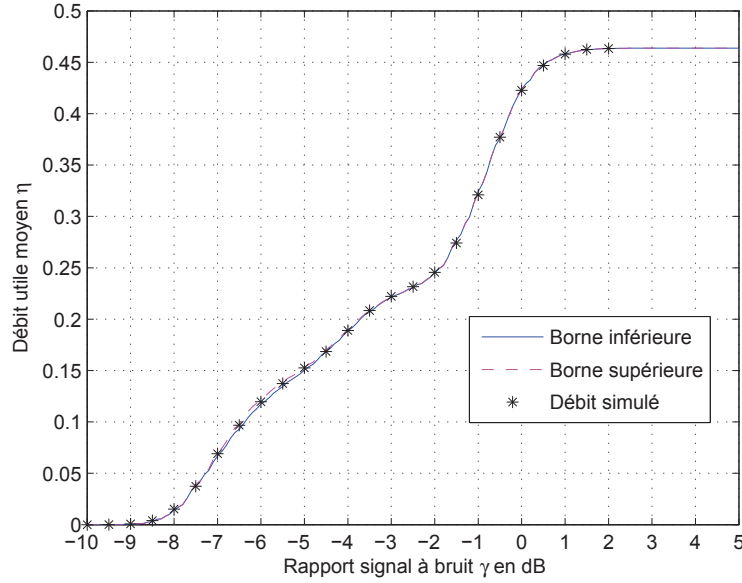


Figure 3.2 — Débit utile moyen simulé en HARQ-II avec codes convolutifs de polynôme générateur (23 35) pour $D = 256$, $H = 40$, $R_o = 0,46$ et $M_1 = 3$

polynôme générateur (23,35). Chaque paquet comprend $D = 256$ bits utiles et $H = 40$ bits d’entête, soit un rendement global $R_o = 0,46$ bits/symbole. Le nombre maximum de retransmissions est fixé à $M_1 = 3$. On voit sur cet exemple que l’encadrement précédent est très précis.

Dans la suite de la thèse, on utilisera l’approximation $p(m) \approx \Pr\{e_m\}$ car elle est simple à évaluer, et nous avons constaté dans nos simulation qu’elle donne des résultats très proche de la valeur exacte.

3.4.3 Performances limites en HARQ

On souhaite tout d’abord établir les performances limites en HARQ-I et II afin de pouvoir comparer ces stratégies de retransmission entre elles sous leur jour le plus favorable. Il s’agit également d’établir des bornes de performances auxquelles on pourra ensuite venir comparer des codes pratiques. Pour obtenir l’expression des performances limites, on suppose l’utilisation de codes gaussiens aléatoires arbitrairement longs (cf. chapitre 1 section 1.3.2), qui atteignent la capacité du canal BABG. Signalons que l’étude des performances limites a pour autre avantage d’aboutir à une expression analytique exacte du débit, relativement facile à évaluer numériquement.

L’étude des performances limites est tout d’abord menée sur le canal BABG, puis étendue ensuite au canal de Rayleigh à évanouissements par blocs.

Performances limites sur le canal BABG

Rappelons que l'information mutuelle moyenne (IMM) réalisée par la transmission de symboles gaussiens complexes circulaires i.i.d., centrés, de variance E_s , sur un canal BABG complexe de variance N_o s'écrit

$$I(\gamma) = \log_2(1 + \gamma) \quad (\text{bits/symbole}) \quad (3.34)$$

avec $\gamma = \frac{E_s}{N_o}$. L'information mutuelle moyenne coïncide avec la capacité du canal $C(\gamma)$ car elle est maximale pour une distribution gaussienne à l'entrée.

Considérons un code aléatoire gaussien composé de M mots de code de longueur N symboles, dont les composantes sont tirées aléatoirement suivant une loi $\mathcal{CN}(0, E_s)$. Le rendement du code vaut donc $R_o = \log_2(M)/N$ bits/symbole. Lorsque $N \rightarrow \infty$, la probabilité d'erreur après décodage p_o s'écrit alors :

$$p_o = \begin{cases} 1 & R_o \geq I(\gamma) \\ 0 & R_o < I(\gamma) \end{cases} \quad (3.35)$$

En HARQ-I, les paquets erronés sont rejetés, et la probabilité d'échec de décodage conjoint $p(m)$ aux rounds 0 à m vaut :

$$p(m) = \Pr\{e_o, e_1, \dots, e_m\} = p_o^{m+1} = \begin{cases} 1 & R_o \geq I(\gamma) \\ 0 & R_o < I(\gamma) \end{cases} \quad (3.36)$$

A R_o fixé, suivant que le RSB soit supérieur ou égal ou bien inférieur à la capacité, soit le paquet passe toujours correctement à la première transmission, soit il ne passe jamais (fonctionnement en tout ou rien).

Évaluons ensuite l'expression de la probabilité d'échec de décodage conjoint $p(m)$ aux rounds 0 à m en HARQ-II-CC. Nous avons montré en section 3.3.1 qu'au round m , après combinaison des paquets, tout se passe comme si le paquet initial avait été transmis sur un canal BABG équivalent caractérisé par un RSB $m + 1$ fois supérieur au RSB initial (le RSB s'accumule de round en round).

$$\Pr\{e_m\} = \Pr\{R_o \geq I((m+1)\gamma)\} = \begin{cases} 1 & R_o \geq I((m+1)\gamma) \\ 0 & R_o < I((m+1)\gamma) \end{cases} \quad (3.37)$$

La probabilité d'échec de décodage conjointe $p(m)$ s'écrit alors :

$$p(m) = \Pr\{e_o, e_1, \dots, e_m\} = \Pr\{I(\gamma) \leq R_o, I(2\gamma) \leq R_o, \dots, I((m+1)\gamma) \leq R_o\} \quad (3.38)$$

Or l'information mutuelle moyenne $I(\gamma)$ est une fonction strictement croissante du RSB γ . Il s'ensuit que dans le cas particulier de codes gaussiens asymptotiquement longs, les événements $\{e_m\}$ sont inclus les uns dans les autres : $e_o \in e_1 \in \dots \in e_m$. On a alors la simplification :

$$p(m) = \Pr\{e_o, e_1, \dots, e_m\} = \Pr\{e_m\} = \begin{cases} 1 & R_o \geq I((m+1)\gamma) \\ 0 & R_o < I((m+1)\gamma) \end{cases} \quad (3.39)$$

Autrement dit, la transmission échoue systématiquement, jusqu'à ce qu'au fil des retransmissions, le RSB équivalent cumulé devienne suffisamment grand pour que l'information mutuelle vue par la station B soit supérieure au rendement R_o d'une transmission. Le paquet passe alors avec certitude.

Terminons par l'expression de $p(m)$ dans le cas de la stratégie HARQ-II-IR. La transmission de redondance incrémentale crée un code gaussien équivalent $\mathcal{C}_m$ de rendement $R_m = R_o/(m+1)$ au round m , et la probabilité d'échec de décodage à ce round s'écrit alors :

$$\Pr\{e_m\} = \Pr\{I(\gamma) \leq R_m\} = \Pr\{(m+1)I(\gamma) \leq R_o\} \begin{cases} 1 & R_o \geq (m+1)I(\gamma) \\ 0 & R_o < (m+1)I(\gamma) \end{cases} \quad (3.40)$$

On voit bien ici qu'en IR, l'information mutuelle moyenne s'accumule de round en round. La probabilité d'échec de décodage conjoint $p(m)$ s'écrit :

$$p(m) = \Pr\{e_o, e_1, \dots, e_m\} = \Pr\{I(\gamma) \leq R_o, 2I(\gamma) \leq R_o, \dots, (m+1)I(\gamma) \leq R_o\} \quad (3.41)$$

Or l'information mutuelle moyenne $I(\gamma)$ est toujours positive ou nulle. Dans le cas du HARQ-II-IR avec codes gaussiens asymptotiquement longs, les événements $\{e_m\}$ sont à nouveau inclus les uns dans les autres : $e_o \in e_1 \in \dots \in e_m$, et on conserve la simplification :

$$p(m) = \Pr\{e_o, e_1, \dots, e_m\} = \Pr\{e_m\} = \begin{cases} 1 & R_o \geq (m+1)I(\gamma) \\ 0 & R_o < (m+1)I(\gamma) \end{cases} \quad (3.42)$$

La transmission échoue systématiquement jusqu'à ce que l'on ait accumulé suffisamment d'information mutuelle au fur et à mesure des rounds, relativement au rendement R_o fixé. Le paquet passe alors avec certitude.

La figure 3.3 compare le débit utile moyen en HARQ-I, HARQ-II-CC et HARQ-II-IR dans le cas d'une transmission avec codes gaussiens arbitrairement longs, de rendement $R_o = 0,96$ bits/symboles, sur canal BABG complexe, pour un nombre de retransmissions maximal $M_1 = 3$. Une première constatation est que le débit utile en HARQ-II est bien meilleur qu'en HARQ-I. On voit donc que la mémorisation et la réutilisation des paquets erronés en HARQ-II apporte un gain considérable par rapport au rejet de paquets erronés en HARQ-I. Le phénomène est particulièrement marqué à faible et moyen RSB. En effet, comme expliqué précédemment, le débit utile moyen en HARQ-I n'augmente pas avec le nombre de retransmissions. Il est soit nul, soit maximal, suivant que le RSB de la transmission soit inférieur ou supérieur ou égal à la limite de Shannon $\gamma_{lim} = 2^{R_o-1}$ correspondant au rendement R_o considéré. La comparaison entre le débit utile en HARQ-II-CC et le débit utile en HARQ-II-IR montre par ailleurs que l'envoi de redondance incrémentale est plus performant que la combinaison des paquets, par l'apport d'un gain de codage supplémentaire (contre un gain en RSB uniquement en CC). Notons finalement que les courbes de débit utile moyen présentent des paliers correspondants au rendement de codage global $R_o/(m+1)$ vus à chaque round par la station B. Plus le nombre de retransmissions est élevé, plus le nombre de paliers augmente.

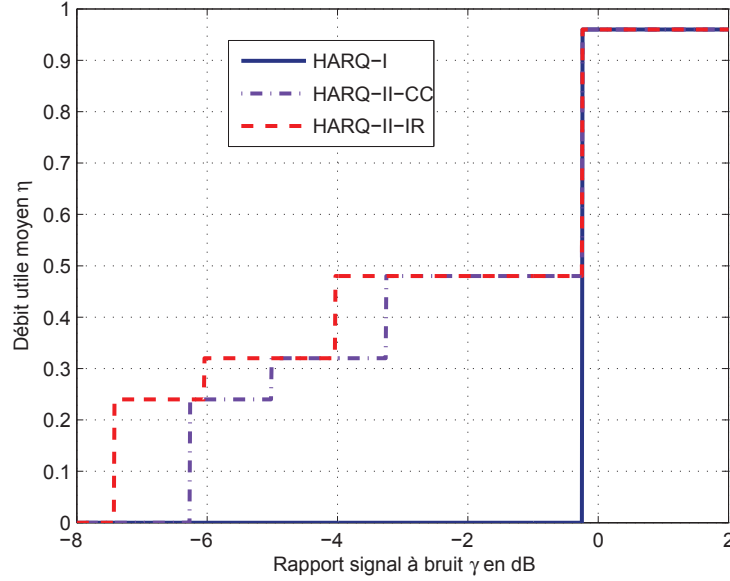


Figure 3.3 — Débit utile en HARQ-I, HARQ-II-CC et HARQ-II-IR pour une transmission sur canal BABG complexe avec codes gaussiens de longueur infinie et rendement $R_o = 0,96$, et avec $M_1 = 3$

Performances limites sur le canal de Rayleigh à évanouissements par blocs

Sur le canal de Rayleigh à évanouissements par blocs, chaque transmission sur la voie directe expérimente un RSB $|h_m|^2\gamma$ différent, où $\gamma = \frac{E_s}{N_o}$ désigne le RSB moyen, et h_m est un coefficient d'atténuation $\sim \mathcal{CN}(0, 1)$.

Considérons tout d'abord une transmission HARQ-I. Dans le cas de codes gaussiens aléatoires asymptotiquement longs, on montre que la probabilité d'échec de décodage à chaque round coïncide avec la probabilité de coupure, définie comme :

$$\Pr\{e_m\} = p_o = \Pr\{I(|h_m|^2\gamma) \leq R_o\} \quad (3.43)$$

Celle-ci dépend de R_o et de la distribution du RSB $|h_m|^2\gamma$. Dans le scénario considéré ici, c'est une variable aléatoire i.i.d. de round en round, et distribuée selon une loi exponentielle d'espérance γ . D'où :

$$p_o = 1 - \exp\left(-\frac{2^{R_o} - 1}{\gamma}\right) \quad (3.44)$$

En tenant compte du rejet des paquets erronés, la probabilité d'échec de décodage conjointe $p(m)$ en HARQ-I s'écrit :

$$p(m) = \Pr\{e_o, e_1, \dots, e_m\} = p_o^{m+1} = \left(1 - \exp\left(-\frac{2^{R_o} - 1}{\gamma}\right)\right)^{m+1} \quad (3.45)$$

En HARQ-II-CC, comme vu précédemment sur le canal BABG, le RSB se cumule au fur et à mesure des retransmissions, et la probabilité d'échec de décodage après

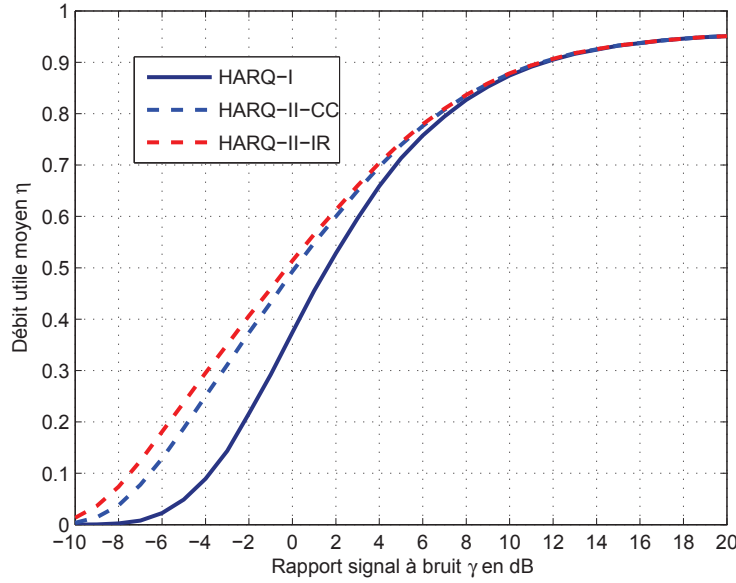


Figure 3.4 — Débit utile en HARQ-I, HARQ-II-CC et HARQ-II-IR pour une transmission sur canal de Rayleigh à évanouissements par blocs avec codes gaussiens de longueur infinie et rendement $R_o = 0,96$, et avec $M_1 = 3$

combinaison, pour un code gaussien aléatoire arbitrairement long, s'écrit :

$$\Pr(e_m) = \Pr \left\{ I \left(\gamma \sum_{i=0}^m |h_i|^2 \right) \leq R_o \right\} \quad (3.46)$$

Les événements $\{e_m\}$ sont inclus les uns dans les autres $e_o \in e_1 \in \dots \in e_m$, et la probabilité d'un échec de décodage conjoint $p(m)$ aux rounds 0 à m s'écrit :

$$p(m) = \Pr(e_o, e_1, \dots, e_m) = \Pr\{e_m\} = \Pr \left\{ I \left(\gamma \sum_{i=0}^m |h_i|^2 \right) \leq R_o \right\} \quad (3.47)$$

De la même façon, en HARQ-II-IR, l'information mutuelle moyenne s'accumule à chaque round [40], et la probabilité d'échec de décodage pour un code gaussien aléatoire s'écrit :

$$\Pr(e_m) = \Pr \left\{ \sum_{i=0}^m I(|h_i|^2 \gamma) \leq R_o \right\} \quad (3.48)$$

Dans ce cas également, les événements $\{e_m\}$ sont inclus les uns dans les autres $e_o \in e_1 \in \dots \in e_m$, et la probabilité d'un échec de décodage conjoint $p(m)$ aux rounds 0 à m s'écrit :

$$p(m) = \Pr(e_o, e_1, \dots, e_m) = \Pr\{e_m\} = \Pr \left\{ \sum_{i=0}^m I(|h_i|^2 \gamma) \leq R_o \right\} \quad (3.49)$$

La figure 3.4 compare le débit utile moyen en HARQ-I, HARQ-II-CC et HARQ-II-IR dans le cas d'une transmission avec codes gaussiens aléatoires de rendement $R_o = 0,96$

bits/symboles, sur canal de Rayleigh à évanouissements par blocs, pour un nombre de retransmissions maximal $M_1 = 3$. On retrouve la même hiérarchie de performances que sur le canal BABG, à savoir que le HARQ-II surpasse le HARQ-I à tous les RSB. Par ailleurs, l'IR apporte un gain supplémentaire par rapport à la combinaison des paquets, gain d'autant plus marqué que le RSB est faible.

Le débit utile en HARQ est calculé en utilisant l'expression donnée par (3.18). Les probabilités d'échec sont données par (3.45), (3.47) et (3.49). La figure 3.4 montre le débit utile moyen en HARQ-I, HARQ-II-CC et HARQ-II-IR pour une transmission sur un canal de Rayleigh à évanouissements par blocs. Le nombre maximal de retransmissions choisi est $M_1 = 3$. Le rendement de codage et modulation est $R_o = 0,96$. Ce rendement est choisi pour avoir les mêmes paramètres choisis précédemment. On constate en premier temps que le HARQ-II est meilleur en débit sur toute la plage du RSB que le HARQ-I. Ce résultat est déjà confirmé pour les transmissions sur canal BABG. Ensuite, on remarque que le HARQ-II-IR apporte un débit légèrement supérieur au débit en HARQ-II-CC pour les zones de faibles RSB. Aux forts RSB, les deux stratégies sont équivalentes.

3.4.4 Performances à longueur finie

Les résultats précédents représentent ce que l'on peut faire de mieux en terme de correction avec les protocoles HARQ-I/II classiques. Ils ne sont atteignables qu'asymptotiquement, dans la limite de grande taille de blocs. Pour se rapprocher d'un scénario plus réaliste, on propose ici de reprendre l'évaluation des performances pour des codes de longueur finie. Dans un premier temps, on établit les performances limites pour des codes optimaux de longueur finie. C'est l'une des contributions originales de cette thèse. Elle s'appuie sur des résultats établis récemment dans [41, 6], où les auteurs ont donné une approximation fine et simple à calculer de la probabilité d'erreur atteignable par le meilleur code de longueur N et rendement R_o . Dans un deuxième temps, on présente les performances obtenues avec des codes convolutifs car c'est un schéma de codage pratique utilisé dans de nombreuses normes radio.

Codes optimaux de longueur finie

On se limite ici à une étude sur le canal BABG, car le calcul des performances des codes optimaux de longueur finie sur le canal à évanouissements par blocs reste un problème ouvert à ce jour¹.

Pour une transmission sur canal BABG complexe de RSB $\gamma = \frac{E_s}{N_o}$, la probabilité d'erreur minimale atteignable avec le meilleur code de rendement R_o bits/symbole et de longueur N symboles, est approchée finement par [7] :

$$p(\gamma) \approx \frac{1}{2} \operatorname{erfc} \left(\sqrt{\frac{N}{2V(\gamma)}} (C(\gamma) - R_o + \alpha_o \frac{\log_2(N)}{N}) \right) \quad (3.50)$$

1. Signalons que des résultats ont été récemment obtenus sur le canal de Rayleigh ergodique à évanouissements symbole par symbole [8]

où α_o est une constante que l'on prend égale à $1/2$ ici (dépend de la contrainte imposée sur l'énergie des mots dans la composition du code. Pour une transmission sur canal gaussien à entrée et sortie continue et lorsque tous les mots de code ont la même énergie (modulation de phase), ou bien lorsque l'on fixe une limite d'énergie maximale par mot, on a $\alpha_o = 1/2$ [41, 6]), $C(\gamma)$ est la capacité du canal, et $V(\gamma)$ est un paramètre appelé dispersion du canal

$$C(\gamma) = \log_2(1 + \gamma) \quad (3.51)$$

$$V(\gamma) = \frac{\gamma}{2} \frac{\gamma + 2}{(\gamma + 1)^2} (\log_2(e))^2 \quad (3.52)$$

Dans le cas d'une transmission sur un canal BABG réel, en conservant la définition précédente de γ , seules changent la capacité et la dispersion, qui sont alors données par :

$$C(\gamma) = \frac{1}{2} \log_2(1 + 2\gamma) \quad (3.53)$$

$$V(\gamma) = \frac{\gamma}{2} \frac{\gamma + 1}{(\gamma + \frac{1}{2})^2} (\log_2(e))^2 \quad (3.54)$$

Pour calculer le débit utile moyen avec des codes optimaux de longueur finie en HARQ-I, on identifie la probabilité de perte de paquet p_o à l'expression (3.50) de $p(\gamma)$. La probabilité d'échec de décodage conjoint aux rounds 0 à m se calcule alors comme $p(m) = p_o^{m+1}$.

En HARQ-II-CC, on utilise le fait que le RSB se cumule à chaque round se sorte que la probabilité d'échec de décodage au round m est donnée par :

$$\Pr\{e_m\} = p(\gamma_m) \quad \text{avec } \gamma_m = (m + 1)\gamma \quad (3.55)$$

Pour calculer la probabilité d'échec de décodage conjointe $p(m)$, on utilise l'approximation $p(m) \approx \Pr\{e_m\}$.

Dans le cas de l'HARQ-II-IR, on fait l'hypothèse que le poinçonnage compatible en rendement d'un code optimal de longueur finie produit un code optimal. Cette hypothèse est vraisemblable dans la mesure où l'on sait construire aujourd'hui des familles très performantes de codes turbos ou LDPC compatibles en rendement, capables d'opérer très près de la limite de Shannon pour une grande variété de rendements. Ainsi, la probabilité d'échec de décodage au round m $\Pr\{e_m\}$ se calcule à partir de l'expression (3.50) de $p(\gamma)$, en notant que la station B voit alors un code équivalent de longueur $N_m = (m + 1)N$ et rendement $R_m = R_o/(m + 1)$, sur un canal de RSB γ . Pour calculer la probabilité d'échec de décodage conjointe $p(m)$, on utilise enfin l'approximation $p(m) \approx \Pr\{e_m\}$.

La figure 3.5 compare le débit utile moyen en HARQ-I, HARQ-II-CC et HARQ-II-IR pour une transmission sur canal gaussien complexe avec des codes optimaux de longueur $N = 512 + 40$ symboles et un rendement $R_o = 512/(512 + 40) = 0,96$ bits/symbole. En comparant avec les résultats de la figure 3.3, on observe un comportement similaire à celui obtenu avec des codes de longueur infinie. En particulier, l'HARQ-II-IR offre le

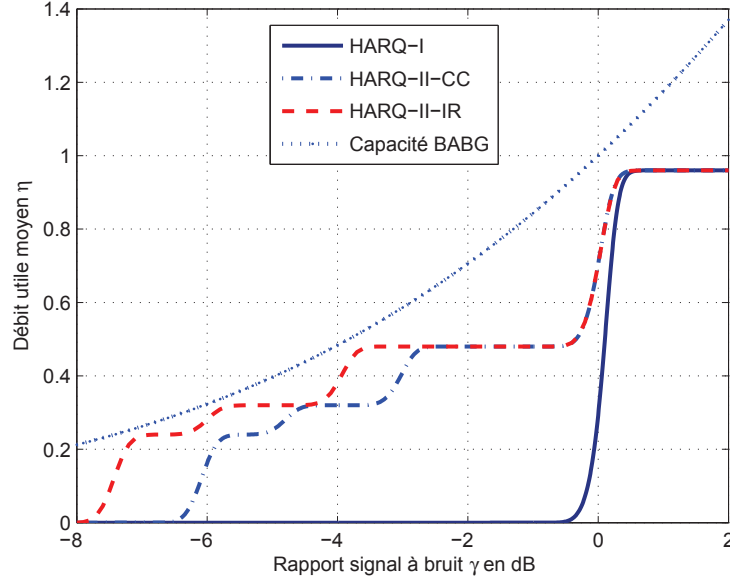


Figure 3.5 — Débit utile en HARQ-I, HARQ-II-CC et HARQ-II-IR pour une transmission sur canal BABG complexe avec codes optimaux de longueur finie, pour $D = 512$, $H = 40$, $R_o = 0,96$ et $M_1 = 3$

meilleur débit, et les deux protocoles HARQ-II surpassent le HARQ-I. Notons que les performances à longueur finie convergent vers les performances asymptotiques lorsque l'on augmente la taille N du code.

La figure 3.5 montre également la courbe de capacité $C(\gamma)$ du canal BABG complexe. Rappelons que cette limite est atteinte par un codage FEC basé sur des codes gaussiens asymptotiquement longs, dont le rendement est adapté à la capacité du canal à chaque valeur du RSB. On voit ici qu'en IR, aux RSBs coïncidants avec la limite de Shannon correspondant à chacun des rendements discrets $R_m = R_o/(m + 1)$ supportés par le protocole de retransmission, les performances obtenues sont proches de la capacité. Elles s'en éloignent ensuite rapidement lorsque le RSB augmente, du fait de la plage limitée des rendements supportés. Il est possible de coller plus finement à la courbe de capacité sur une grande plage de RSB en augmentant le nombre M_1 de retransmissions maximal et en faisant varier le nombre de symboles retransmis à chaque round (paquets de taille variable). Ce cas de figure n'a pas été considéré dans cette thèse.

Codes convolutifs

Pour terminer cette analyse de performances, on se propose ici d'évaluer le débit utile moyen lorsque l'on combine un protocole de retransmission avec un code convolutif, scénario fréquent en pratique dans les systèmes radio [42].

Nous considérons ici une transmission BPSK sur un canal BABG réel, avec $M_1 = 1$ retransmission au maximum. En HARQ-I, chaque paquet est protégé par le code convolutif à 16 états de rendement $r_o = 1/2$ et polynômes générateurs (23, 35). On

m	Séquences génératrice	d_{free}	A_{d_o+d} ($d = 0, 1, \dots, 6$)
1	(1 3)	3	1, 1, 1, 1, 1, 1, 1
2	(5 7)	5	1, 2, 4, 8, 16, 32, 64
3	(15 17)	6	1, 3, 5, 11, 25, 55, 121
4	(23 35)	7	2, 3, 4, 16, 37, 68, 176
5	(53 75)	8	1, 8, 7, 12, 48, 95, 281
6	(133 171)	10	11, 0, 38, 0, 193, 0, 1331

Tableau 3.1 — Premières valeurs des paramètres d_o (poids de Hamming) et A_{d_o} (multiplicité) pour différents codes convolutifs de rendement 1/2

utilise les premiers termes de la borne de l'union pour calculer la probabilité d'erreur après décodage p_o à chaque round (cf. chapitre I) :

$$p_o \leq D \sum_{d_o \geq d_{free}} \frac{A_{d_o}}{2} \operatorname{erfc} \sqrt{d_o \gamma} \quad (3.56)$$

Les couples (d_o, A_{d_o}) désignent respectivement le poids de Hamming et la multiplicité des chemins erronés dans le treillis, et sont donnés dans la table 3.1. Dans les simulations, nous avons considéré des paquets de $D = 512$ bits utiles avec un entête de $H = 40$ bits, soit un MCS de rendement $R_o = r_o D / (D + H) = 0,48$ bits/symbole. La probabilité d'échec de décodage conjoint est calculée comme $p(m) = p_o^{m+1}$

On procède de manière similaire pour évaluer les performances en HARQ-II-CC. Plus précisément, on utilise l'approximation $p(m) \approx \Pr\{e_m\}$, en tenant compte du fait que le RSB augmente à chaque transmission pour calculer $\Pr\{e_m\}$:

$$p(m) = \Pr\{e_m\} \leq D \sum_{d_o \geq d_{free}} \frac{A_{d_o}}{2} \operatorname{erfc} \sqrt{d_o(m+1)\gamma} \quad (3.57)$$

En HARQ-II-IR, le paquet initial est tout d'abord encodé par le code mère de rendement 1/4 et polynômes générateurs (23 35 27 33). A la première transmission, seule la redondance correspondant aux bits codés produits par les polynômes (23 35) est transmise. La redondance produite par les polynômes (27 33) est ensuite envoyée si besoin lors de la deuxième transmission. En utilisant à nouveau l'approximation $p(m) \approx \Pr\{e_m\}$, la probabilité d'échec de décodage conjoint aux rounds 0 à m est évaluée comme suit [37] [43] :

$$p(m) = \Pr\{e_m\} \leq \frac{D}{u} \sum_{d_m \geq d_{free}} \frac{A_{d_m}}{2} \operatorname{erfc} \sqrt{d_m \gamma} \quad (3.58)$$

Les couples (d_m, A_{d_m}) sont donnés dans la table 3.2 pour chaque round $m = 0, 1$. Le paramètre u désigne la période de poinçonnage et vaut ici $u = 8$ [37].

La figure 3.6 présente les courbes de débit utile moyen pour ce scénario. A titre de comparaison, nous avons également présenté les performances limites obtenues avec des codes optimaux de même rendement. On voit ici que la hiérarchie HARQ-II-IR >

m	0	1
d_m	7, 8, ..., 12	15, 16, ..., 20
$A_{d_m}^m$	16, 24, 32, 80, 296, 544	16, 8, 0, 8, 16, 32

Tableau 3.2 — Premières valeurs des paramètres d_m (poids de Hamming) et $A_{d_m}^m$ (multiplicité) à chaque round m pour le poinçonnage compatible en rendement du code convolutif mère de rendement 1/4 et de polynômes (23 35 27 33)

HARQ-II-CC > HARQ-I est conservée avec des codes pratiques. On constate également un écart de performance important avec les codes optimaux de longueur finie. Ceci est dû à l'utilisation de codes convolutifs simples à pouvoir de correction modérée. L'utilisation de codes LDPC ou turbocodes permettrait de se rapprocher davantage des performances des codes optimaux.

Nous avons également tracé sur la figure 3.6 les performances obtenues en l'absence de codage convolutif, lorsque la gestion des erreurs est entièrement reportée sur le protocole ARQ-SR de la couche RLC. On constate ici la nette supériorité de l'approche hybride à faible et moyen RSB. A fort RSB en revanche, l'HARQ est pénalisé par l'introduction de la redondance FEC, qui devient alors superflue, et une approche purement basée sur la retransmission devient préférable.

Nous avons également indiqué par des marqueurs sur la figure 3.6 les courbes de débit utile moyen simulé. On voit sur cet exemple que les prédictions analytiques basées sur l'approximation $p(m) \approx \Pr\{e_m\}$ coïncident avec les résultats de simulation, ce qui valide la démarche de calcul proposée.

3.4.5 Optimisation du débit utile moyen

De manière similaire à l'étude menée en section 2.3, on cherche ici à optimiser les paramètres du système de manière à maximiser le débit utile moyen. Les résultats de simulation précédents montrent que le débit utile en HARQ augmente avec le RSB, tous les autres paramètres étant fixés par ailleurs. A RSB constant, le débit utile moyen en HARQ dépend non seulement de la taille de données utiles D comme en ARQ, mais également du taux de codage R_o . Il existe alors deux degrés de liberté sur lesquels on peut jouer pour maximiser le débit utile moyen. Dans cette section, l'optimisation du débit sur la base des codes optimaux est un résultat original.

Influence de la taille des données

Dans cette sous section, on étudie la variation du débit utile moyen en fonction de la taille de données utiles D . On se limite ici à l'étude des schémas HARQ basés sur des codes convolutifs. En effet, pour les codes optimaux de longueur finie, la probabilité d'erreur diminue avec la taille des données utiles. Le débit maximal est alors atteint pour des tailles infinies.

La figure 3.7 présente l'évolution du débit utile moyen en fonction de la taille de

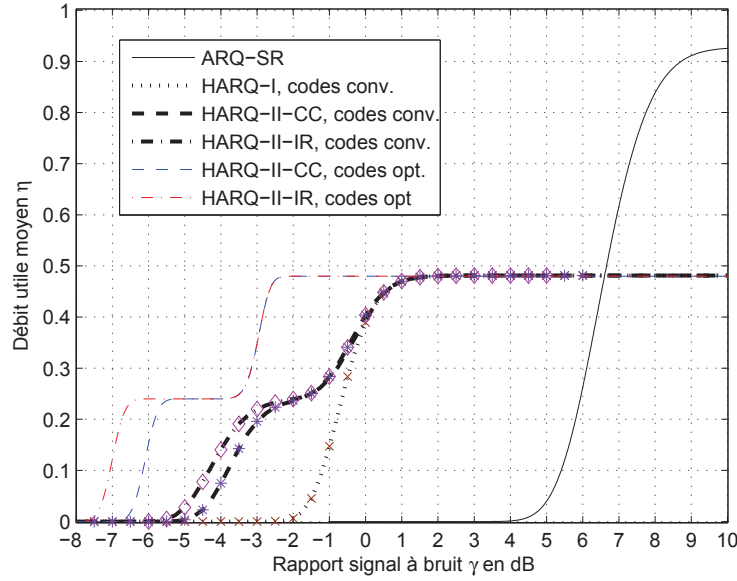


Figure 3.6 — Débit utile en HARQ-I et HARQ-II pour une transmission binaire avec codage convolutif de rendement $1/2$, $D = 512$ bits, $H = 40$ bits ($R_o = 0,48$ bits/symbole), et deux rounds au maximum ($M_1 = 1$). Les performances en ARQ-SR et HARQ-II + codes optimaux de mêmes paramètres sont également présentées. Courbes en traits pleins : débit théorique, courbes avec marqueurs : débit simulé.

données pour un schéma HARQ-II dans le cas d'une transmission binaire sur canal BABG, avec codage convolutif de rendement $1/2$ et polynôme générateur (23 35), pour une taille d'acquittements $H = 40$ bits et pour différentes valeurs de $\gamma = E_s/N_o$. Le nombre maximum de rounds est fixé ici à deux. On remarque que le débit utile admet un maximum pour une taille optimale de données utiles, d'autant plus grand que le RSB est élevé. En effet, pour un code convolutif, la probabilité d'erreur de décodage augmente avec D , ce qui diminue en retour η . Lorsque l'on réduit D , la proportion de données utiles devient petite devant la taille des entêtes, ce qui réduit également η . D'où l'existence d'un optimum.

Influence du taux de codage

Nous venons d'étudier la variation du débit utile moyen en fonction du RSB et de la taille des données utiles. Il est intéressant de traiter également la variation du débit en fonction du rendement R_o . A cet effet, on fixe une taille de données $D = 512$ bits et un nombre de retransmissions $M_1 = 3$. On se limite ici aux codes optimaux de longueur finie car l'étude avec des codes convolutifs nécessite de définir autant de masques de poinçonnage que de rendements souhaités. Par ailleurs, les rendements supportés sont nécessairement de la forme $k/(k+1)$. Il est donc plus simple de considérer des codes optimaux pour lesquels on peut faire varier le rendement R_o de manière continue. De plus les conclusions obtenues ont une portée plus générale.

Il s'agit donc ici de déterminer le rendement optimal R_o qui maximise le débit utile moyen, à RSB γ et taille de données D fixés. On présente sur la figure 3.8 l'évolution de

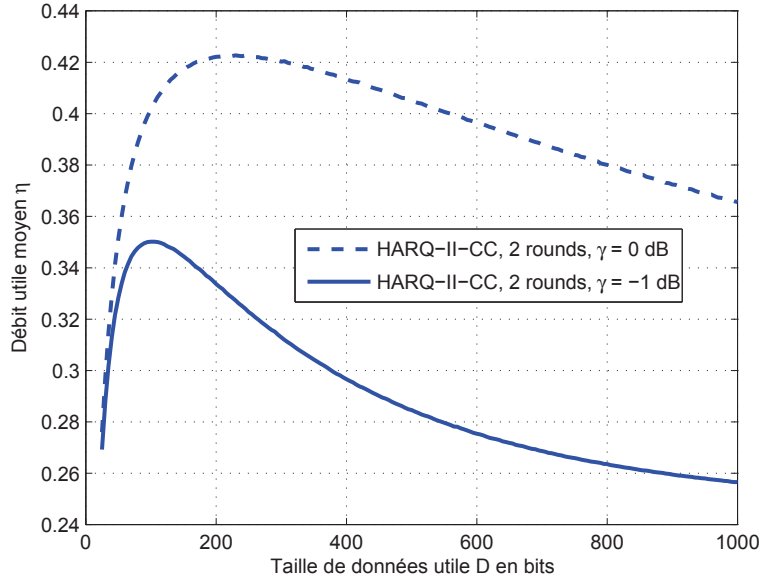


Figure 3.7 — Débit utile en HARQ-II-CC avec codes convolutifs de polynômes générateurs (23,35) en fonction de la taille de données utiles pour une transmission sur canal BABG réel avec $H = 40$ bits, $M_1 = 1$.

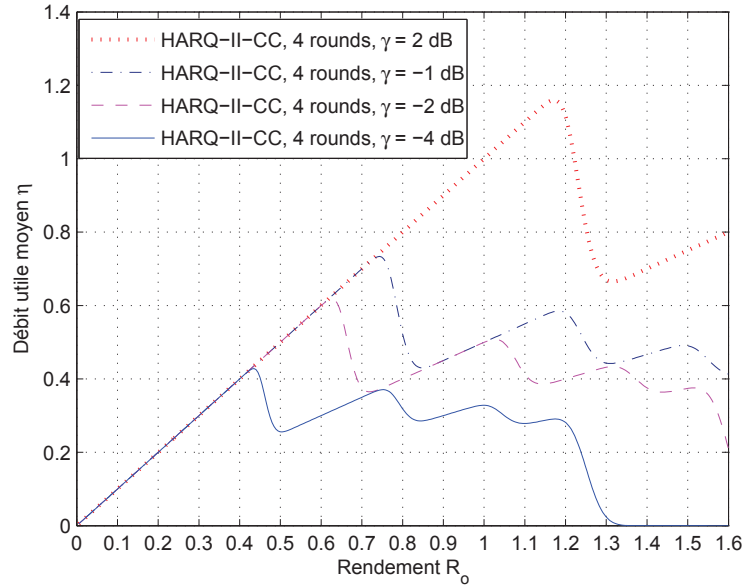


Figure 3.8 — Débit utile HARQ-II-CC avec codes optimaux en fonction du rendement R_o pour $D = 512$ bits, $H = 40$ bits et $M_1 = 3$.

η en fonction du R_o pour différents RSB, avec un maximum de 4 rounds. On constate qu'il existe un rendement optimal unique qui maximise le débit utile moyen, et que cet optimum dépend de γ pour une taille D de données constante.

A RSB constant, aux faibles rendements, le débit utile est croissant jusqu'à sa valeur maximale (maximum global). Cette zone de rendement correspond à un succès

de décodage au premier round. Il décroît ensuite puisque le récepteur n'arrive à décoder les données qu'au second round (peu de symboles sont erronés). On observe alors des maximums locaux selon que le récepteur décode le paquet au second round ou bien aux rounds suivants.

A rendement constant, on constate que différentes valeurs de l'énergie à l'émission (ou de manière équivalente du RSB γ) peuvent conduire au même débit utile moyen. Dans ce cas, il est plus judicieux d'émettre avec l'énergie minimale. Dans la section suivante, on étudie l'impact de la variation de l'énergie transmise par symbole sur le débit utile moyen.

3.5 Efficacité énergétique

Le codage canal permet de corriger les erreurs pour des faibles RSB, mais rajoute en contrepartie de la redondance aux paquets de données transmis. Ceci rend la transmission d'un paquet plus coûteuse en énergie. Dans cette section, on souhaite vérifier dans un premier temps si le codage apporte un gain du point de vue de l'efficacité énergétique. Il s'agit donc de comparer les schémas HARQ-I avec les protocoles ARQ vis-à-vis de cette métrique. On étudie ensuite l'apport en terme d'efficacité énergétique de la mémorisation des paquets erronés en réception, en comparant les schémas HARQ-I avec les schémas HARQ-II. On se penche enfin sur l'optimisation des paramètres du système de manière à minimiser l'énergie totale nécessaire ζ pour transmettre avec succès un bit utile à la station B. Nous proposons dans cette section une optimisation de l'efficacité énergétique en HARQ avec des codes optimaux de longueur finie. C'est la seconde grande contribution de ce chapitre.

3.5.1 Expression générale de l'efficacité énergétique

Tout comme en ARQ (section 2.4), l'énergie ζ est définie comme l'énergie moyenne totale consommée par cycle de transmission, normalisée par le nombre moyen de bits de données correctement reçus par cycle. En reprenant le résultat établi au chapitre 2 pour un protocole de retransmission de type *Selective Reject*, l'énergie moyenne nécessaire par bit utile s'écrit :

$$\zeta = E_s \frac{2 - \rho}{R_o} \frac{1 + \sum_{m=0}^{M_1-1} p(m)}{1 - p(M_1)} \quad (3.59)$$

où $p(m)$ ici prend la valeur adéquate selon le schéma HARQ (I/II-CC/II-IR) et le modèle de code considéré (codes optimaux de longueur finie ou codes convolutifs).

3.5.2 Efficacité énergétique en HARQ-I

L'expression de l'énergie nécessaire par bit utile montre que ζ dépend de l'énergie transmise par symbole E_s , du rendement de codage à la première transmission R_o , du paramètre ρ et de la probabilité d'échec $p(m)$. Le paramètre ρ et la probabilité $p(m)$ dépendent en retour de la taille de données utiles D , de la taille H des acquittements,

et de R_o . Par conséquent, l'énergie nécessaire par bit utile est donc fonction de R_o , D , H et E_s (ou γ).

On se propose ici d'étudier l'influence de E_s sur ζ , tous les autres paramètres étant fixés par ailleurs. A H et D constants, réduire E_s augmente la probabilité d'échec de décodage et accroît ζ du fait de l'augmentation du nombre moyen de retransmissions. A l'inverse augmenter E_s accroît l'énergie d'un paquet et son acquittement et donc ζ . Ceci suggère l'existence d'une valeur optimale de l'énergie de transmission.

Nous comparons sur la figure 3.9 l'évolution de ζ en fonction du RSB pour une transmission binaire sur canal BABG réel, en ARQ-SR et en HARQ-I avec codes convolutifs de rendement $1/2$ et polynômes générateurs (23,35). Les tailles de données utiles et d'acquittements sont $D = 512$ bits et $H = 40$ bits, soit un rendement de transmission $R_o = 0,48$ bits/symbole en HARQ et $R_o = 0,96$ bits/symbole en ARQ-SR, et un nombre maximum infini de rounds (protocole non tronqué). On remarque que les schémas HARQ font une meilleure utilisation de l'énergie que les protocoles ARQ et ce malgré l'introduction de redondance supplémentaire. En effet, l'énergie minimale ζ en HARQ est inférieure à l'énergie minimale en ARQ. Il en va de même pour le RSB optimal de fonctionnement. A fort RSB, l'énergie consommée par un bit utile est équivalente à l'énergie consommée par une seule transmission (forte probabilité de succès au premier round), normalisée par D . On retrouve donc une croissance asymptotique linéaire en E_s ($\zeta \approx \frac{2-\rho}{R_o} E_s$). Notons qu'à fort RSB, ζ est plus élevée en HARQ qu'en ARQ du fait de la transmission de symboles de redondance supplémentaires. A faible RSB, l'énergie moyenne consommée par un bit utile est élevée en ARQ et HARQ car le nombre moyen de bits décodés est faible comparé au nombre total de bits transmis.

3.5.3 Efficacité énergétique en HARQ-II

Dans cette section, on compare l'efficacité énergétique en HARQ-I et en HARQ-II. Il s'agit principalement d'évaluer l'apport de la mémorisation des paquets vis-à-vis de cette métrique.

La figure 3.10 présente l'évolution de ζ en fonction du RSB en HARQ-I et HARQ-II-CC/IR, pour une transmission binaire sur canal BABG réel avec codage convolutif. Nous avons repris ici le scénario considéré dans la section 3.4.4 avec le code convolutif (23 35) en HARQ-I et HARQ-II-CC, le code convolutif mère (23 35 27 33) en HARQ-II-IR, pour un maximum de 2 rounds ($M_1 = 1$), avec $D = 512$ bits et $H = 40$ bits ($R_o = 0,48$ bits/symbole). A titre de comparaison, nous avons également présenté sur la figure 3.10 l'efficacité énergétique en HARQ-I et HARQ-II-CC/IR pour des codes optimaux de longueur finie de même paramètres. On remarque tout d'abord que les schémas HARQ-II sont plus performants que les schémas HARQ-I en terme d'efficacité énergétique. En HARQ-I, ζ admet un minimum global pour une énergie symbole E_s optimale. Par ailleurs, l'énergie nécessaire par bit utile est très sensible à la variation de l'énergie symbole autour de sa valeur optimale. En HARQ-II, on remarque la présence de minima locaux et suivant les cas, on n'observe pas toujours un minimum global unique de ζ . Sans surprise, les schémas HARQ avec codes optimaux présentent

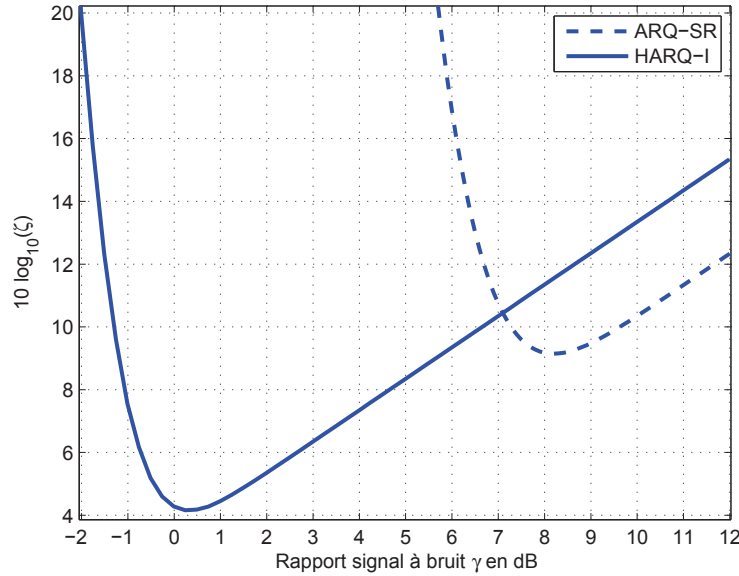


Figure 3.9 — Energie moyenne nécessaire par bit utile en fonction du RSB $\gamma = E_s/N_o$ pour une transmission binaire sur canal BABG réel, en ARQ-SR et HARQ-I non tronqués avec codes convolutifs de polynômes générateurs (23 35) $D = 512$ bits et $H = 40$ bits.

une meilleure efficacité énergétique que les schémas HARQ avec codes convolutifs. Les résultats obtenus avec codes optimaux peuvent être vus comme une limite théorique à laquelle on peut comparer des schémas de codage pratiques, car ces derniers ont une probabilité d'échec de décodage supérieure ou égale à celle des codes optimaux, et donc une consommation énergétique (nombre de retransmissions) plus importante. Notons que tous les schémas HARQ-I/II ont la même efficacité énergétique asymptotique (fort RSB) car la première transmission passe avec une forte probabilité dans tous les cas, et on retrouve la croissance linéaire de ζ en γ .

La figure 3.11 compare plus spécifiquement l'efficacité énergétique en HARQ-II-CC et HARQ-II-IR pour des codes optimaux de longueur finie avec respectivement $M_1 = 1$ et $M_1 = 3$. Comme précédemment, on observe que l'énergie minimale consommée par un bit utile ζ est inférieure en HARQ-II-IR, et qu'il en va de même pour l'énergie symbole E_s correspondante. On en conclut que l'envoi de redondance incrémentale est plus avantageux que la combinaison de paquets en terme d'efficacité énergétique, et la différence s'accroît avec le nombre de rounds. On constate par ailleurs que le nombre de minima locaux augmente avec M_1 .

3.6 Extension à une voie de retour imparfaite

Toutes les analyses menées jusqu'à présent dans ce chapitre faisaient l'hypothèse d'une voie de retour parfaite (sans erreurs). On se propose d'étudier dans cette dernière section la robustesse des schémas HARQ vis-à-vis des erreurs sur la voie de retour, et

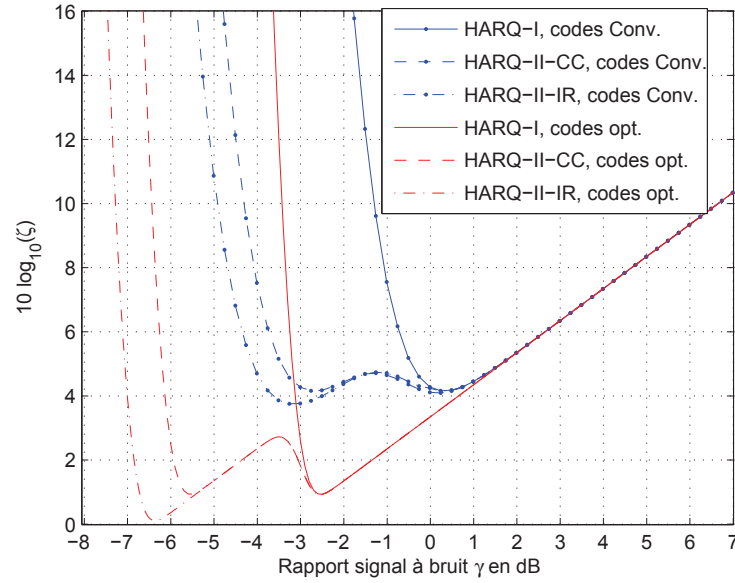


Figure 3.10 — Energie nécessaire ζ par bit utile en fonction du rapport signal à bruit $\gamma = E_s/N_o$ en HARQ-I, HARQ-II-CC et HARQ-II-IR avec des codes convolutifs et codes optimaux de longueur finie, pour une transmission binaire sur canal BABG réel avec $R_o = 0,48$ bits/symbole, $M_1 = 1$, $D = 512$ bits et $H = 40$ bits

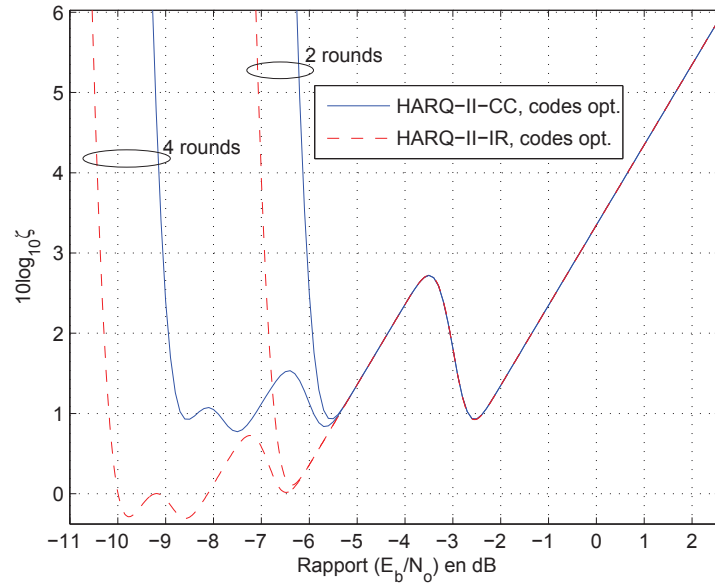


Figure 3.11 — Energie nécessaire par bit utile ζ en fonction du rapport signal à bruit $\gamma = E_s/N_o$ en HARQ-II-CC et HARQ-II-IR avec des codes en blocs optimaux de longueur finie et de rendement $R_o = 0,48$ bits/symbole, avec $D = 512$ bits utiles, $H = 40$ bits, et pour différentes valeurs du nombre maximum de rounds

de revisiter le problème de la répartition de puissance optimale entre la voie directe et la voie de retour en présence de codage.

3.6.1 Prise en compte des pertes d'acquittements dans l'analyse

Le modèle du lien retour est identique à celui considéré dans le chapitre précédent, à savoir un modèle avec perte d'acquittement i.i.d. de round en round avec une probabilité ϵ .

L'expression générale (3.18) du débit utile moyen pour un schéma HARQ-I/II basé sur un protocole SR s'applique aussi bien à une voie de retour parfaite qu'imparfaite. Dans ce dernier cas, il faut alors tenir compte du fait qu'un échec de transmission au round m (événement E_m) n'est plus nécessairement causé par un échec de décodage (événement e_m) mais peut être dû également à une perte d'acquittement. Ainsi :

$$\Pr\{E_m\} = \epsilon + (1 - \epsilon)\Pr\{e_m\} \quad (3.60)$$

En HARQ-I tout comme en ARQ, la probabilité conjointe d'observer $m + 1$ échecs de transmission successifs aux rounds 0 à m $P(m) = \Pr\{E_0, E_1, \dots, E_m\}$ se calcule d'une façon simple grâce à l'indépendance entre les événements $\{E_m\}$:

$$P(m) = \prod_{i=0}^m \Pr\{E_i\} \quad (3.61)$$

L'indépendance des événements $\{E_m\}$ ne s'applique plus en HARQ-II du fait de la mémorisation et de la réutilisation des paquets d'un round sur l'autre. Nous introduisons donc ici une expression approchée de la probabilité conjointe $P(m)$ pour un schéma HARQ-II avec voie de retour imparfaite.

Désignons par θ l'événement "perte d'acquittement". Ainsi, $\Pr\{\theta\} = \epsilon$. Au premier round ($m = 0$), on a :

$$P(0) = \Pr\{E_0\} = \epsilon + (1 - \epsilon) \Pr\{e_0\} \quad (3.62)$$

Au second round, la probabilité d'échec de transmission conjointe $P(1)$ s'écrit :

$$P(1) = \Pr\{E_0, E_1\} \quad (3.63)$$

$$= \Pr\{E_0, \theta\} + \Pr\{E_0, e_1, \bar{\theta}\} \quad (3.64)$$

$$= \epsilon P(0) + (1 - \epsilon) \Pr\{E_0, e_1\} \quad (3.65)$$

$$= \epsilon P(0) + (1 - \epsilon) (\Pr\{\theta, e_1\} + \Pr\{e_0, e_1, \bar{\theta}\}) \quad (3.66)$$

$$= \epsilon P(0) + (1 - \epsilon) (\epsilon \Pr\{e_1\} + (1 - \epsilon) \Pr\{e_0, e_1\}) \quad (3.67)$$

En utilisant l'approximation $\Pr\{e_0, e_1\} \approx \Pr\{e_1\}$ (cf. section 3.4.2), on aboutit à :

$$P(1) = \Pr\{E_0, E_1\} \quad (3.68)$$

$$\approx \epsilon P(0) + (1 - \epsilon) \Pr\{e_1\} \quad (3.69)$$

Au round 3,

$$P(2) = \Pr\{E_o, E_1, E_2\} \quad (3.70)$$

$$= \Pr\{E_o, E_1, \theta\} + \Pr\{E_o, E_1, e_2, \bar{\theta}\} \quad (3.71)$$

$$= \epsilon P(1) + (1 - \epsilon) \Pr\{E_o, E_1, e_2\} \quad (3.72)$$

En décomposant les événements E_o et E_1 ($E_m = \{\theta\} \cup \{e_m, \bar{\theta}\}$), on aboutit à :

$$P(2) = \epsilon P(1) + (1 - \epsilon) [\epsilon (\Pr\{e_2\} + (1 - \epsilon) \Pr\{e_o, e_2\}) + (1 - \epsilon) (\Pr\{e_1, e_2\} + (1 - \epsilon) \Pr\{e_o, e_1, e_2\})] \quad (3.73)$$

En approximant les probabilités conjointes $\Pr\{e_o, e_2\}$, $\Pr\{e_1, e_2\}$ et $\Pr\{e_o, e_1, e_2\}$ par $\Pr\{e_2\}$, on obtient :

$$P(2) \approx \epsilon P(1) + (1 - \epsilon) \Pr\{e_2\} \quad (3.74)$$

Plus généralement, en utilisant l'approximation : $\Pr\{e_{i_1}, e_{i_2}, \dots, e_{i_k}, e_m\} \approx \Pr\{e_m\}$, $\forall i_1 < i_2 < \dots < i_k < m$, on peut montrer qu'au round m :

$$P(m) \approx \epsilon P(m-1) + (1 - \epsilon) \Pr\{e_m\} \quad (3.75)$$

De manière équivalente, en développant la récursion la probabilité conjointe $P(m)$ peut s'écrire :

$$P(m) \approx \epsilon^{m+1} + \sum_{i=0}^m \Pr\{e_m\} \epsilon^{m-i} (1 - \epsilon) \quad (3.76)$$

On peut ainsi évaluer le débit utile des schémas HARQ-II dans le cas d'une voie de retour imparfaite en remplaçant $P(m)$ par (3.76) dans l'expression du débit (3.18).

3.6.2 Influence de l'imperfection de la voie de retour sur le débit en HARQ

Sur la base du résultat établi précédemment, on étudie ici l'impact des pertes d'acquittements sur les performances limites du système, sur des canaux BABG et à évanouissements par blocs.

Transmission sur canal BABG

On considère dans un premier temps des codes gaussiens aléatoires de longueur infinie. Pour cette famille de codes, l'approximation $\Pr\{e_{i_1}, e_{i_2}, \dots, e_{i_k}, e_m\} \approx \Pr\{e_m\}$ est une égalité stricte. La figure 3.12 présente l'évolution du débit utile limite moyen en HARQ-I et HARQ-II pour une transmission de rendement $R_o = 0,96$ bits/symbole sur un canal BABG complexe, avec $M_1 = 3$ et $\epsilon = 0,1$. On constate tout d'abord que la hiérarchie de performances (HARQ-II-IR > HARQ-II-CC > HARQ-I) observé dans le cas d'une voie de retour parfaite est conservée. En comparant avec la figure

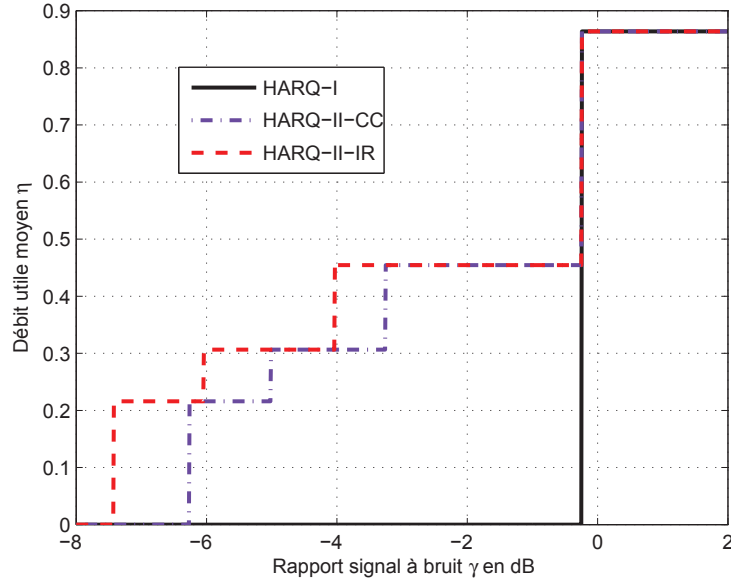


Figure 3.12 — Débit utile en HARQ-I, HARQ-II-CC et HARQ-II-IR pour une transmission sur canal BABG complexe avec des codes gaussiens aléatoires de longueur infinie et rendement $R_o = 0,96$ bits/symbole, $M_1 = 3$ et une probabilité de perte $\epsilon = 0,1$ sur la voie retour

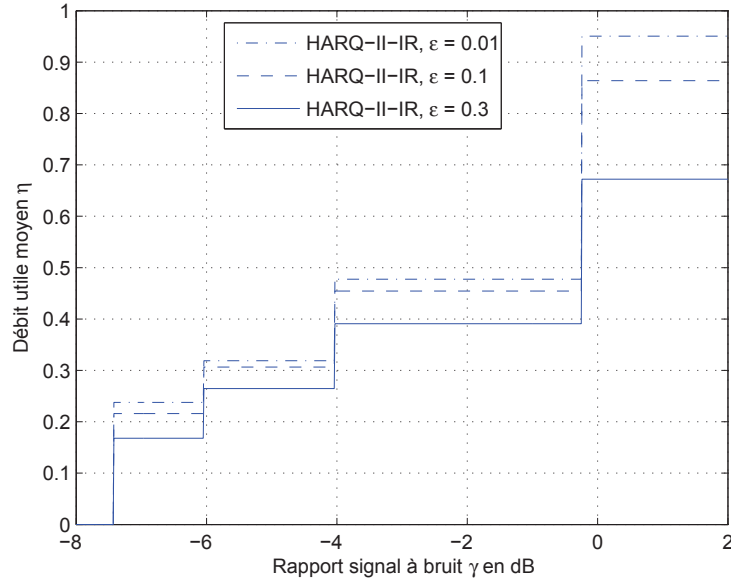


Figure 3.13 — Débit utile en HARQ-II-IR pour une transmission sur canal BABG complexe avec des codes gaussiens aléatoires de longueur infinie et rendement $R_o = 0,96$ bits/symbole, $M_1 = 3$ et pour différentes probabilités de perte sur la voie retour ($\epsilon = 0,01$, $\epsilon = 0,1$ et $\epsilon = 0,3$)

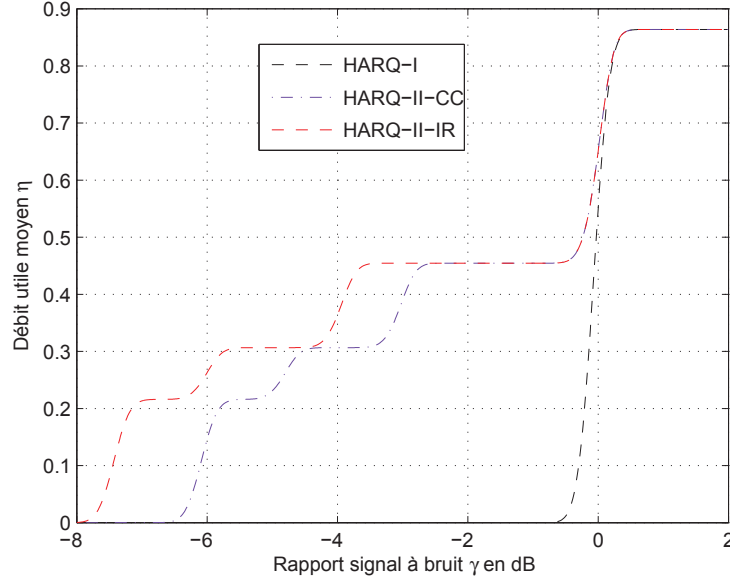


Figure 3.14 — Débit utile en HARQ-I, HARQ-II-CC et HARQ-II-IR pour une transmission sur canal BABG complexe avec des codes optimaux de longueur finie et rendement $R_o = 0,96$ bits/symbole, $M_1 = 3$, $D = 512$ bits, $H = 40$ bits et une probabilité de perte $\epsilon = 0,1$ sur la voie retour

3.3, on constate par ailleurs que le débit utile maximum à chaque round diminue. Les résultats présentés sur la figure 3.13 qui étudie l'impact de ϵ sur le débit utile moyen en HARQ-II-IR confirment que l'augmentation de la probabilité d'échec sur la voie de retour entraîne une perte en débit utile moyen.

Les figures 3.14 et 3.15 montrent que l'on observe un comportement tout à fait similaire avec des codes optimaux de longueur finie, tous les autres paramètres de la transmission étant identiques par ailleurs.

Transmission sur canal à évanouissements par blocs

Sur le canal de Rayleigh à évanouissements par blocs, on se limite à l'étude de codes gaussiens de longueur infinie. Les probabilités d'échec de décodage au round m en HARQ-I, HARQ-II-CC et HARQ-II-IR sont données par les équations (3.45), (3.47) et (3.49) respectivement. Les probabilités $P(m)$ sont obtenues grâce à la formule (3.76), et le débit utile est calculé à l'aide de son expression donnée par (3.18).

La figure 3.16 compare le débit utile limite moyen en HARQ-I et HARQ-II pour une transmission de rendement $R_o = 0,96$ bits/symbole sur le canal de Rayleigh à évanouissements par blocs, pour $M_1 = 3$ et $\epsilon = 0,1$. La figure 3.17 présente quant à elle l'impact de la probabilité de perte ϵ sur le débit utile moyen en HARQ-II-IR sur le canal de Rayleigh à évanouissements par blocs. Les conclusions sont les mêmes que dans le cas du canal BABG.

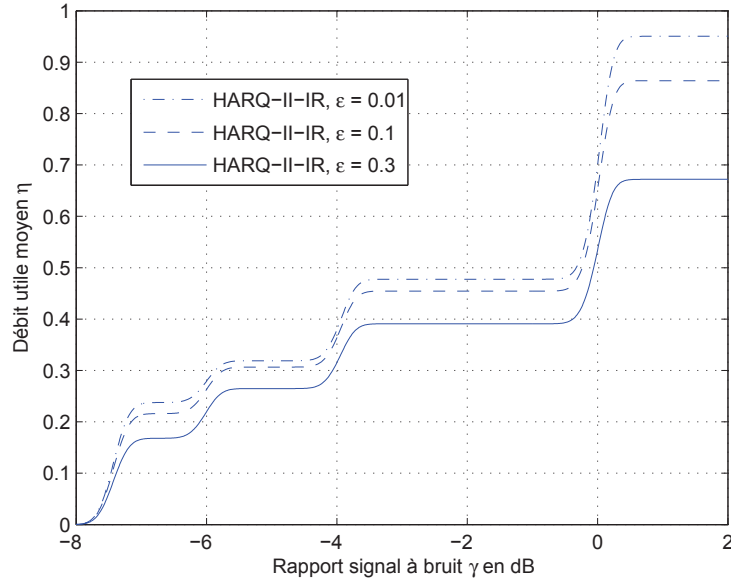


Figure 3.15 — Débit utile en HARQ-II-IR pour une transmission sur canal BABG complexe avec des codes optimaux de longueur finie et rendement $R_o = 0,96$ bits/symbole, $M_1 = 3$, $D = 512$ bits, $H = 40$ bits, et pour différentes probabilités de perte sur la voie retour ($\epsilon = 0,01$, $\epsilon = 0,1$ et $\epsilon = 0,3$)

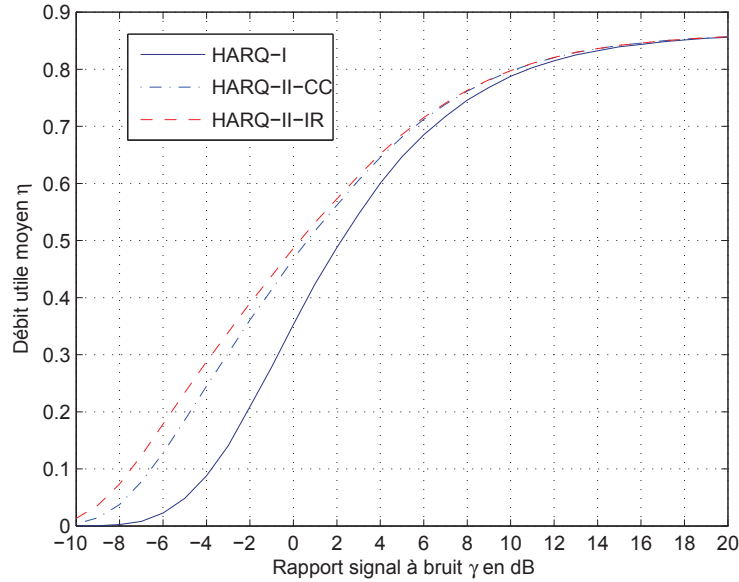


Figure 3.16 — Débit utile en HARQ-I, HARQ-II-CC et HARQ-II-IR pour une transmission sur canal de Rayleigh à évanouissements par blocs avec des codes gaussiens aléatoires de longueur infinie et rendement $R_o = 0,96$ bits/symbole, $M_1 = 3$ et une probabilité de perte $\epsilon = 0,1$ sur la voie retour

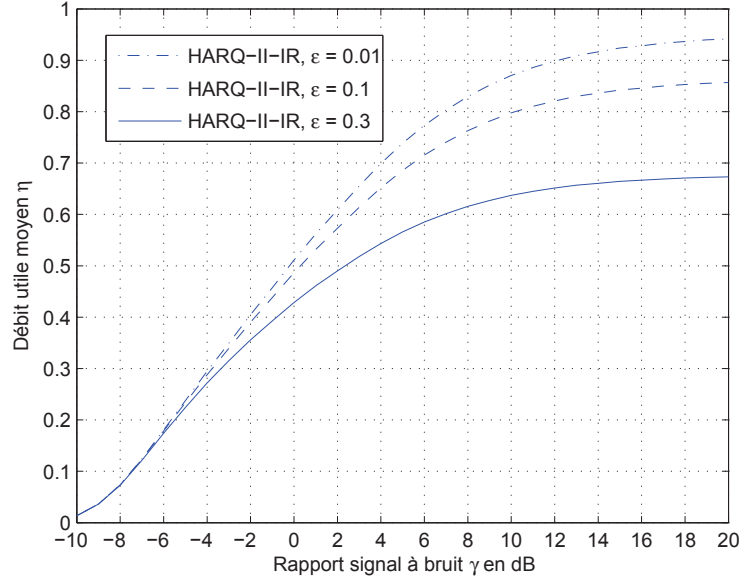


Figure 3.17 — Débit utile en HARQ-II-IR pour une transmission sur canal de Rayleigh à évanouissements par blocs avec des codes gaussiens aléatoires de longueur infinie et rendement $R_o = 0,96$ bits/symbole, $M_1 = 3$ et pour différentes probabilités de perte sur la voie retour ($\epsilon = 0,01$, $\epsilon = 0,1$ et $\epsilon = 0,3$)

3.6.3 Répartition optimale de la puissance entre voie directe et voie de retour

Dans cette section, à l'image de l'étude menée en section 2.5.2 au chapitre précédent, on se place par exemple dans un réseau de capteurs où les nœuds communiquent entre eux pair à pair, selon un protocole HARQ, et avec une autonomie limitée. On peut donc supposer que chaque nœud dispose d'un budget d'énergie global constant E_{round} par round, à dépenser à la fois sur le lien direct et la voie de retour. On peut alors s'interroger sur la meilleure façon de répartir l'énergie totale entre les deux liens, de manière à maximiser le débit utile moyen par exemple, ou bien l'efficacité énergétique. Rappelons que le fait d'augmenter l'énergie sur la voie de retour permet de réduire la probabilité de perte d'acquittement, et donc contribue à réduire le nombre de retransmissions. Par soucis de concision, on se limitera ici à la maximisation du débit utile.

Le modèle d'analyse est identique à celui utilisé pour l'ARQ en section 2.5.2. On rappelle que $0 \leq \alpha \leq 1$ est la variable qui définit la proportion de l'énergie totale d'un round allouée à la transmission de l'acquittement sur la voie de retour, et E_g désigne l'énergie globale par symbole (énergie totale par round E_{round} normalisée par le nombre total de symboles émis sur chacun des deux liens au cours du round). Dans le cas d'une modulation binaire, nous avons montré en section 2.5.2 que les énergies transmises par symbole sur la voie directe E_s et sur la voie de retour E_a peuvent s'exprimer en fonction

de α , ρ et E_g comme suit :

$$E_s = (1 - \alpha)(2 - \rho)E_g \quad (3.77)$$

$$E_a = \alpha \frac{2 - \rho}{1 - \rho} E_g \quad (3.78)$$

On se place dans le cadre d'une transmission sur canal gaussien réel, et l'on considère des schémas HARQ basés sur des codes optimaux de longueur finie. Le modèle d'erreur sur la voie de retour est un canal à effacement, de probabilité de perte ϵ . En supposant que les paquets d'acquittements sont transmis avec une modulation binaire, sans codage², la probabilité ϵ est alors donnée par :

$$\epsilon = 1 - \left(1 - \frac{1}{2} \operatorname{erfc} \sqrt{\frac{E_a}{N_o}} \right)^H \quad (3.79)$$

Le débit utile moyen η se calcule à partir de (3.18), où les probabilités conjointes $P(m)$ dépendent de E_s par l'intermédiaire de $\Pr\{e_m\}$, et de E_a par l'intermédiaire de ϵ . À énergie globale E_g constante, le débit utile est donc fonction de α et l'on peut rechercher l'existence d'un point de fonctionnement optimum.

La figure 3.18 montre l'évolution de η en fonction de α , pour une transmission HARQ-I, HARQ-II-CC et HARQ-II-IR sur canal gaussien réel avec codes optimaux de longueur finie, en prenant $D = 512$ bits, $H = 40$ bits, $R_o = 0,96$ bits/symbole, $M_1 = 3$ et $E_g/N_o = 1$ dB. Aux faibles valeurs de α (faible énergie allouée à l'émission des acquittements), le débit utile est faible à cause des nombreuses retransmissions engendrées par la perte d'acquittements. Les paquets sont correctement reçus sur la voie directe, mais l'émetteur continue à retransmettre tant qu'il ne reçoit pas l'acquittement positif (ou que le nombre maximal de retransmissions n'est pas atteint). Le débit utile moyen vaut alors $R_o/(M_1 + 1)$. Au fur et à mesure que α augmente, le débit utile s'améliore puisque l'émetteur reçoit certains acquittements. Au bout d'une certaine valeur de α , l'énergie allouée à la transmission de paquet devient insuffisante pour effectuer un décodage correct. Le récepteur ne décode alors pas le paquet à sa première transmission, et le débit se réduit progressivement, au fur et à mesure que le nombre moyen de retransmissions augmente. Il existe donc une valeur optimale de α qui maximise le débit utile, ce que l'on observe effectivement sur la figure 3.18. Les comportements des courbes du débit dépendent du schéma HARQ. En HARQ-I, la courbe du débit est croissante puis décroissante. En HARQ-II, on observe des plages de α sur lesquelles le débit est constant. À l'intérieur de ces plages, il faut un nombre minimum moyen requis de retransmissions pour réussir à décoder, et le saut d'un palier à l'autre traduit l'apparition d'une retransmission supplémentaire en moyenne.

3.7 Conclusion

Dans ce chapitre, nous avons tout d'abord décrit le principe des trois grandes familles de protocoles ARQ hybride qui combinent un code correcteur d'erreurs avec un

2. On pourrait également appliquer un code optimal de longueur finie sur la voie de retour, ϵ serait alors donné par (3.50).

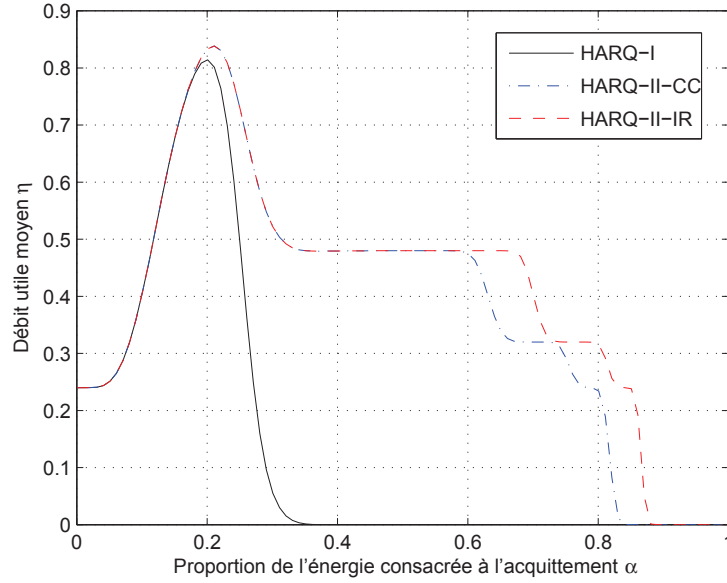


Figure 3.18 — Débit utile pour une transmission HARQ-II-IR avec codes optimaux de longueur finie sur canal BABG réel en fonction de α , pour $D = 512$ bits, $H = 40$ bits, $R_o = 0,96$ bits/symbole, $M_1 = 3$, et $E_g/N_o = 1$ dB

protocole de retransmission ARQ. Nous avons développé une analyse du débit utile moyen des protocoles HARQ-I, HARQ-II-CC et HARQ-II-IR, en considérant aussi bien des modèles théoriques des codes (afin d'établir les performances limites de ces schémas), que des codes utilisés en pratique tels que les codes convolutifs. De cette analyse, il ressort notamment que l'envoi de redondance incrémentale (HARQ-II-IR) est la solution la plus performante en terme de débit. La combinaison de paquets (HARQ-II-CC) est une bonne alternative, particulièrement attractive à fort RSB, compte tenu de la simplicité du récepteur. L'HARQ-I est nettement en retrait du fait du rejet des paquets erronés.

Nous nous sommes également intéressés à l'étude de l'efficacité énergétique des différents protocoles hybrides. Nous avons ainsi montré que le schéma HARQ-II-IR fait également la meilleure utilisation de l'énergie à sa disposition.

La plupart du temps, l'étude des schémas HARQ est menée en supposant une voie de retour parfaite. Nous avons donc étendu notre modèle d'analyse des performances au cas d'une voie de retour avec perte d'acquittements, et vérifié ainsi en particulier que, pour tous les schémas HARQ, le débit utile moyen diminue lorsque la probabilité de perte d'acquittements augmente. Notons également que le schéma HARQ-II-IR reste la solution la plus performante en terme de débit dans ce contexte.

Pour corriger les erreurs de transmission, les schémas HARQ classiques s'appuient sur la retransmission de paquets de redondance tous relatifs à un même paquet d'information. Il peut arriver des situations où seul un petit nombre de symboles de redondance peut suffire pour réussir à décoder. Dans ce cas, la retransmission d'un paquet entier de redondance induit une perte en débit. Dans le chapitre suivant, nous introduisons des schémas HARQ plus évolués, qui construisent des paquets de redondance

pouvant aider au décodage simultané de plusieurs paquets d'information.

CHAPITRE

4

Schéma HARQ avec
redondance
multipaquets

4.1 Introduction

Les schémas HARQ-II sont très efficaces pour réaliser des transmissions fiables sur des canaux inconnus ou variants dans le temps. La plupart des travaux relatifs aux schémas HARQ-II dans la littérature s'intéressent à l'optimisation du code correcteur d'erreurs (règle de poinçonnage) [37, 38, 44]. En revanche, l'étude et l'optimisation de la stratégie de retransmission en elle-même ont reçu très peu d'attention.

L'inconvénient des schémas HARQ classiques (dits *simple paquet* ici, ou SP), est que les paquets retransmis (paquets de redondance) ne sont relatifs qu'au paquet de donnée courant. La station A poursuit la retransmission de redondance jusqu'à bonne réception du paquet de donnée, ou bien jusqu'à ce que le nombre maximal M_1 de retransmissions soit atteint. La retransmission de paquets de redondance de même taille que le paquet initial engendre une réduction importante du débit utile à chaque nouvelle retransmission, alors que parfois, suivant le rapport signal sur bruit de fonctionnement, la retransmission d'un petit nombre de symboles erronés aurait suffi (avec une pénalité moindre sur le débit). Dans ce chapitre, on introduit une nouvelle stratégie de retransmission HARQ permettant de remédier à cet inconvénient. C'est l'une des contributions principales de cette thèse. Cette stratégie, appelée HARQ *multi-paquets* (HARQ-MP), procède à un encodage conjoint de plusieurs paquets de données afin de produire des paquets de redondance pouvant aider au décodage simultané de plusieurs paquets de données erronés. Ceci augmente en retour le débit utile moyen de la transmission.

Le chapitre est organisé comme suit. Nous commençons par introduire le principe général de la stratégie HARQ multi-paquets (HARQ-MP). Nous étudions ensuite les performances limites de cette approche, en terme de débit utile moyen, dans l'hypothèse d'une voie de retour parfaite. Cela nous conduit à proposer différents schémas de codage permettant de mettre en œuvre la stratégie MP en pratique, et à évaluer leurs performances respectives. Nous nous penchons finalement sur la meilleure manière de réaliser le protocole de retransmission multi-paquets en présence d'erreurs sur la voie de retour.

4.2 Principe de la stratégie multi-paquets

Dans cette section, nous rappelons tout d'abord le principe de la stratégie SP présentée dans le chapitre précédent. Nous présentons ensuite en détail le principe de l'approche MP proposée.

4.2.1 Rappel du principe des schémas HARQ-SP

Dans les schémas HARQ-II classiques, dits *simple paquet* (SP), étudiés au chapitre précédent, le paquet initial ainsi que les paquets de redondance envoyés au cours d'un cycle sont tous relatifs à un même paquet de donnée en provenance de la couche RLC. Le processus HARQ ne passe au paquet de données suivant que lorsque le cycle HARQ se termine, soit lorsque la station A a été notifiée de la bonne réception du paquet, ou bien lorsque le nombre maximum M_1 de retransmissions autorisées est atteint. Les paquets de redondance sont construits soit par répétition du paquet envoyé au premier round (variante *Chase Combining*), soit par poinçonnage compatible en rendement d'un code mère à faible rendement (variante *Incremental Redundancy*). En notant R_o le rendement d'une transmission de A vers B, le débit utile maximal réalisable au round m vaut alors $R_o/(m+1)$ bits/symbole ($m = 0, \dots, M_1$).

4.2.2 Principe des schémas HARQ multi-paquets

L'idée principale des schémas HARQ multi-paquets (HARQ-MP) consiste à construire des paquets de redondance portant, non pas sur un seul paquet de données comme en SP, mais sur $L > 1$ paquets de données consécutifs. Les paquets de redondance sont alors construits par encodage conjoint de ces L paquets de données. La mise en œuvre détaillée de la stratégie multi-paquets, illustrée sur la figure 4.1, est la suivante.

Au premier round ($m = 0$), L paquets de données successifs sont encodés séparément par un même code $\mathcal{C}_o$, de rendement R_o , puis envoyés les uns à la suite des autres à la station B. Celle-ci les décode séparément, comme en HARQ-SP. Si le décodage échoue pour l'un ou plusieurs de ces L paquets, la station B renvoie un acquittement négatif à la station A, ce qui enclenche alors le processus de transmission de redondance MP.

Au cours des rounds suivants ($m > 0$), des paquets de redondance, de taille N symboles chacun, sont construits par encodage conjoint des L paquets de données à l'aide d'un second code, noté $\mathcal{C}_{is}$, puis envoyés à la station B, à raison d'un paquet de redondance par round. Les paquets de redondance sont produits de manière incrémentale. La stratégie HARQ-MP appartient donc à la famille des protocoles de type HARQ-II-IR. En réception, la station B procède au décodage conjoint des L paquets de donnée, sur la base de l'ensemble des $L + m$ paquets reçus depuis le premier round. Le cycle ARQ s'arrête lorsque la station A a été notifiée de la bonne réception des L paquets de données, ou bien lorsque le nombre maximal de retransmissions autorisées M_2 est atteint.

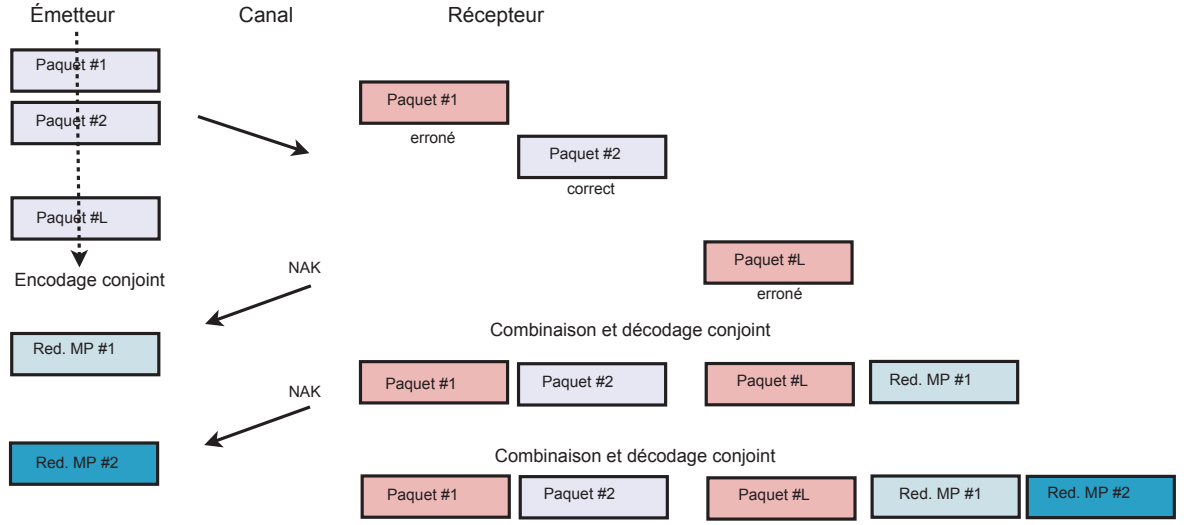


Figure 4.1 — Principe d'une transmission utilisant le protocole HARQ-MP

Au premier round ($m = 0$), du fait de l'encodage séparé des L paquets, la station B voit un code de rendement R_o bits/symbole. Aux rounds suivants ($m > 0$), au fur et à mesure de la réception de nouveaux paquets de redondance MP, la station B voit un code équivalent $\mathcal{C}_m$ de rendement $R_m = R_o L / (L + m)$ bits/symbole.

Dans le principe de construction des paquets de redondance MP, on reconnaît une forme de concaténation série avec un code externe $\mathcal{C}_o$, de rendement R_o , et un code interne systématique $\mathcal{C}_{is}$, de rendement $R_{is} = \frac{L}{L+M_2}$, dont la sortie est poinçonnée de manière incrémentale, sous contrainte de compatibilité en rendement. Cette analogie sera exploitée avantageusement en section 4.4.2 pour proposer des schémas de codage permettant d'implémenter le MP en pratique.

Notons qu'il existe des schémas de superposition de transmission de nouveaux paquets et des retransmissions relatives à des anciens paquets mal reçus. Dans le schéma introduit en [45], l'émetteur transmet un seul paquet au premier round. Au second round, il construit un paquet par un encodage conjoint du paquet mal reçu et d'un nouveau paquet de données. Ceci constitue un multiplexage de deux paquets de données différents. Il peut être vu comme un cas particulier de la stratégie MP proposé dans cette thèse. La référence [46] propose également une transmission simultanée d'un nouveau paquet de données avec les retransmissions relatives aux paquets précédents.

Signalons pour terminer qu'il existe plusieurs façons de signaler à la station A l'échec de décodage au premier round. Dans la description ci-dessus, nous avons ainsi proposé que la station B renvoie un acquittement collectif, relatif à l'ensemble des L paquets transmis au round 0. On peut également imaginer l'envoi d'un acquittement à chaque paquet reçu au cours du premier round, à la manière des protocoles SP classiques. Ces deux options seront discutées et comparées plus en détail en section 4.5, dans le contexte d'une voie de retour imparfaite.

4.3 Étude des performances limites de l'approche multi-paquets

L'objectif de cette section consiste à calculer le débit utile moyen que l'on peut atteindre avec l'approche MP, dans l'hypothèse où l'on utilise des codes gaussiens aléatoires de très grande longueur (performances limites), et valider ainsi que cette approche peut s'avérer avantageuse par rapport à l'approche SP classique. Dans toute la section, la voie de retour est supposée sans pertes ni erreurs.

4.3.1 Expression générale du débit utile moyen en HARQ-MP

De la même façon qu'en HARQ-SP, le débit utile moyen η en HARQ-MP est défini comme le rapport entre le nombre moyen $E[\mathcal{B}]$ de bits d'information correctement reçus par cycle ARQ, et le nombre moyen $E[\mathcal{T}]$ de symboles transmis par cycle.

$$\eta = \frac{E[\mathcal{B}]}{E[\mathcal{T}]} \quad \text{bits/symbole} \quad (4.1)$$

La démarche de calcul est identique à celle adoptée dans les chapitres précédents, à deux différences près. Tout d'abord, le nombre de symboles transmis n'est pas le même à tous les rounds. Plus précisément, on transmet L paquets de N symboles chacun au round 0, puis 1 paquet de N symboles à chaque round $m = 1, \dots, M_2$. D'autre part, un échec de décodage au premier round ne signifie pas nécessairement que les L paquets sont tous perdus. On peut avoir réussi à en décoder une partie, ce dont il faut tenir compte dans le calcul de $E[\mathcal{B}]$.

Commençons par calculer le nombre moyen $E[\mathcal{T}]$ de symboles transmis au cours d'un cycle HARQ. On se place pour cela dans le cas d'un protocole HARQ-II-IR générique qui transmet N_m symboles au round m . Reprenant les notations introduites au chapitre II, on désigne par $P(m) = \Pr(E_0, \dots, E_m)$ la probabilité d'observer $m + 1$ échecs de transmission successifs aux rounds 0 à m , et par $Q(m) = \Pr(E_0, \dots, E_{m-1}, \overline{E_m})$ la probabilité d'observer m échecs successifs aux rounds 0 à $m - 1$, suivi d'un succès de transmission au round m . Le nombre moyen de symboles transmis par cycle se calcule alors comme suit :

$$\begin{aligned} E[\mathcal{T}] &= N_0 Q(0) + (N_0 + N_1) Q(1) + \dots + \sum_{m=0}^{M_2} N_m Q(M_2) + \sum_{m=0}^{M_2} N_m P(M_2) \\ &= N_0 + \sum_{m=0}^{M_2-1} N_{m+1} P(m) \end{aligned} \quad (4.2)$$

où l'on a utilisé la relation $Q(m) = P(m-1) - P(m)$. Dans l'approche MP, le premier round est caractérisé par la transmission de $N_0 = LN$ symboles. Aux rounds $m > 0$, l'émetteur transmet des paquets de $N_m = N$ symboles chacun. D'où :

$$E[\mathcal{T}] = N \left(L + \sum_{m=0}^{M_2-1} P(m) \right) \quad (4.3)$$

Afin de calculer le nombre moyen $E[\mathcal{B}]$ de bits d'information correctement délivrés à la station B par cycle HARQ, il nous faut introduire de nouveaux événements. Rappelons que l'événement "échec de décodage au round m " avait été noté e_m , et que $p(m) = \Pr(e_0, \dots, e_m)$ désigne la probabilité conjointe d'observer $m + 1$ échecs de décodage consécutifs aux rounds 0 à m inclus. Pour la suite, on définit de nouveaux événements $\mathcal{A}_l =$ "échec de décodage au premier round, avec l paquets bien décodés parmi L ", pour $l = 0, \dots, L - 1$. On désigne alors par $p_l(m) = \Pr(\mathcal{A}_l, e_1, \dots, e_m)$ la probabilité d'observer $m + 1$ échecs de décodage consécutifs aux rounds 0 à m , avec l paquets parmi L bien décodés à l'issue du premier round. Remarquons qu'avec cette définition,

$$p(m) = \sum_{l=0}^{L-1} p_l(m) \quad (4.4)$$

La variable aléatoire discrète $\mathcal{B}$ peut alors prendre les valeurs suivantes :

$$\mathcal{B} = \begin{cases} 0 & \text{avec une probabilité } p_0(M_2) \\ D & \text{avec une probabilité } p_1(M_2) \\ \vdots & \\ lD & \text{avec une probabilité } p_l(M_2) \\ \vdots & \\ (L-1)D & \text{avec une probabilité } p_{L-1}(M_2) \\ LD & \text{avec une probabilité } 1 - p(M_2) \end{cases} \quad (4.5)$$

Le nombre moyen de paquets correctement décodés s'écrit alors :

$$E[\mathcal{B}] = D \left(L(1 - p(M_2)) + \sum_{l=1}^{L-1} l p_l(M_2) \right) \quad (4.6)$$

Au final, on obtient l'expression suivante du débit utile moyen en HARQ-II avec redondance multi-paquet :

$$\eta = R_o \frac{L(1 - p(M_2)) + \sum_{l=1}^{L-1} l p_l(M_2)}{L + \sum_{m=0}^{M_2-1} P(m)} \quad (4.7)$$

L'expression précédente est très générale. Dans l'hypothèse d'une voie de retour parfaite, un échec de transmission se réduit nécessairement à un échec de décodage, et l'on a la simplification $P(m) = p(m)$. Pour aller plus loin dans le calcul de η , il nous faut maintenant connaître l'expression des probabilités conjointes $p(m)$ et $p_l(m)$, qui dépendent du code FEC utilisé.

4.3.2 Application aux codes gaussiens asymptotiquement longs

Afin d'établir les performances limites de la stratégie HARQ-MP en terme de débit utile, nous développons ici le calcul des probabilités conjointes $p(m)$ et $p_l(m)$ dans le cas d'une transmission sur canal BABG et canal de Rayleigh à évanouissements par blocs, utilisant des codes aléatoires gaussiens asymptotiquement longs.

Transmission sur canal BABG

Considérons tout d'abord la transmissions de paquets encodés par un code gaussien aléatoire de rendement R_o bits/symbole, sur un canal BABG de rapport signal sur bruit $\gamma = \frac{E_s}{N_o}$. Au premier round, dans la limite de grandes tailles de blocs ($N \rightarrow \infty$), chaque paquet est correctement décodé avec une probabilité 1 lorsque $\log_2(1 + \gamma) > R_o$, et perdu avec la même probabilité dans le cas contraire. Ainsi,

$$p(0) = \Pr(e_0) = \begin{cases} 1 & \text{si } \log_2(1 + \gamma) \leq R_o \\ 0 & \text{si } \log_2(1 + \gamma) > R_o \end{cases} \quad (4.8)$$

Notons que suivant les valeurs de R_o et γ , les L paquets envoyés au premier round sont soit tous correctement reçus, soit tous perdus. D'où $p_0(0) = \Pr(e_0)$ et $p_l(0) = 0$ pour $l = 1, \dots, L - 1$.

Aux rounds suivants, le code équivalent vu par la station B au fur et à mesure de la réception de nouveaux paquets de redondance relatifs aux L paquets de données est un code à redondance incrémentale. Le décodage conjoint des L paquets réussit donc dès que l'on a accumulé suffisamment d'information mutuelle, soit à partir de la plus petite valeur de m vérifiant $(L + m) \log_2(1 + \gamma) > LR_o$. Il échoue dans tous les autres cas. La probabilité d'échec de décodage au round m s'écrit donc :

$$\Pr(e_m) = \begin{cases} 1 & \text{si } \log_2(1 + \gamma) \leq \frac{L}{L+m} R_o \\ 0 & \text{si } \log_2(1 + \gamma) > \frac{L}{L+m} R_o \end{cases} \quad (4.9)$$

La probabilité conjointe d'échec de décodage $p(m)$ s'écrit alors

$$\begin{aligned} p(m) &= \Pr(e_0, \dots, e_m) \\ &= \Pr(\log_2(1 + \gamma) \leq R_o, (L + 1) \log_2(1 + \gamma) \leq LR_o, \dots, (L + m) \log_2(1 + \gamma) \leq LR_o) \\ &= \Pr((L + m) \log_2(1 + \gamma) \leq LR_o) \\ &= \Pr(e_m) \end{aligned} \quad (4.10)$$

On voit donc que pour des codes aléatoires gaussiens, les événements $\{e_m\}$ sont inclus les uns dans les autres, ce qui implique que la probabilité conjointe d'échec de décodage $p(m)$ se réduit à la probabilité marginale $\Pr(e_m)$. Notons enfin que de la même manière qu'au premier round, les L paquets sont soit tous décodés correctement à l'issue du round m , soit tous en échec. D'où $p_0(m) = p(m)$ et $p_l(m) = 0$ pour $l = 1, \dots, L - 1$.

Transmission sur canal de Rayleigh

Considérons maintenant une transmission sur canal de Rayleigh à évanouissements par blocs, de RSB moyen $\gamma = E_s/N_o$. Chaque paquet transmis voit une réalisation différente $|h|^2\gamma$ du RSB. Plus précisément, le RSB instantané $|h|^2\gamma$ suit une distribution exponentielle de paramètre $1/\gamma$. Dans le cas de codes aléatoires gaussiens asymptotiquement longs de rendement R_o bits/symbole, la probabilité d'échec de décodage d'un paquet est égale à la probabilité de coupure et vaut :

$$p_o = \Pr(\log_2(1 + |h|^2\gamma) \leq R_o) = 1 - \exp\left(-\frac{2^{R_o} - 1}{\gamma}\right) \quad (4.11)$$

Le premier round se soldera par un échec de décodage si au moins un des paquets est perdu. D'où

$$p(0) = \Pr(e_0) = 1 - (1 - p_o)^L \quad (4.12)$$

De la même façon, la probabilité $p_l(0)$ d'aboutir à un échec de décodage au premier round tout en ayant décodé correctement $l < L$ paquets s'écrit :

$$p_l(0) = \Pr\{\mathcal{A}_l\} = \binom{L}{l} (1 - p_o)^l p_o^{L-l} \quad l = 0, \dots, L-1 \quad (4.13)$$

Au round $m > 0$, la station B effectue un décodage conjoint des L paquets sur la base d'un code gaussien équivalent de rendement $R_m = LR_o/(L + m)$. Si l'on désigne par $|h_i|^2\gamma$ le RSB instantané vu par le i -ème paquet transmis depuis le début du cycle, la probabilité d'échec de décodage au round m s'écrit alors :

$$\Pr(e_m) = \Pr\left(\sum_{i=1}^{L+m} \log_2(1 + |h_i|^2\gamma) \leq LR_o\right) \quad (4.14)$$

Cette probabilité n'admet pas de forme analytique simple, mais peut être évaluée numériquement par tirage de Monte-Carlo. La probabilité conjointe d'échec de décodage aux rounds 0 à m s'écrit alors :

$$\begin{aligned} p(m) &= \Pr(e_0, \dots, e_m) \\ &= \Pr(\log_2(1 + |h_1|^2\gamma) \leq R_o, \dots, \log_2(1 + |h_L|^2\gamma) \leq R_o, \\ &\quad \sum_{i=1}^{L+1} \log_2(1 + |h_i|^2\gamma) \leq LR_o, \dots, \sum_{i=1}^{L+m} \log_2(1 + |h_i|^2\gamma) \leq LR_o) \\ &= \Pr\left(\sum_{i=1}^{L+m} \log_2(1 + |h_i|^2\gamma) \leq LR_o\right) \\ &= \Pr(e_m) \end{aligned} \quad (4.15)$$

A nouveau, dans le cas des codes aléatoires gaussiens, les événements $\{e_m\}$ sont inclus les uns dans les autres et la probabilité conjointe d'échec de décodage $p(m)$ se réduit à la probabilité marginale $\Pr(e_m)$. Cette propriété nous permet par ailleurs d'exprimer la probabilité $p_l(m)$ sous la forme simplifiée suivante :

$$\begin{aligned} p_l(m) &= \Pr\{\mathcal{A}_l, e_1, e_2, \dots, e_m\} \\ &= \Pr\{\mathcal{A}_l, e_m\} \\ &= \Pr\{\log_2(1 + |h_1|^2\gamma) > R_o, \log_2(1 + |h_2|^2\gamma) > R_o, \dots, \log_2(1 + |h_l|^2\gamma) > R_o, \\ &\quad \log_2(1 + |h_{l+1}|^2\gamma) \leq R_o, \dots, \log_2(1 + |h_L|^2\gamma) \leq R_o, \\ &\quad \sum_{i=1}^{L+m} \log_2(1 + |h_i|^2\gamma) \leq LR_o\} \end{aligned} \quad (4.16)$$

Ne disposant de forme analytique simple pour les probabilités $p(m)$ et $p_l(m)$ au round $m > 0$, ces dernières seront évaluées par tirage de Monte-Carlo dans toute la suite de la thèse.

4.3.3 Étude comparative

Sur la base des calculs précédents, nous évaluons ici les performances limites des schémas HARQ-MP pour différents scénarios de transmission, et nous les comparons à celles des schémas HARQ-SP étudiés au chapitre précédent. Seuls les schémas à redondance incrémentale sont considérés ici, car il a été montré au chapitre précédent que le protocole HARQ-II-IR offre un débit au moins aussi bon et souvent meilleur que le protocole HARQ-II-CC. D'autre part, la stratégie HARQ-MP introduite dans ce chapitre se classe également dans la famille des protocoles IR.

Transmission sur canal gaussien

La figure 4.2 compare le débit utile moyen en HARQ-II-IR (schéma SP) et HARQ-MP dans le cas d'une transmission avec codes gaussiens aléatoires arbitrairement longs, de rendement $R_o = 1$ bits/symboles, sur un canal BABG complexe de RSB $\gamma = E_s/N_o$, en considérant un encodage conjoint sur $L = 4$ paquets en MP, et pour un nombre maximal de retransmissions fixé à $M_1 = M_2 = 2$ dans les deux cas. On observe qu'à moyen et fort RSB, le protocole MP offre un débit supérieur ou égal au SP. Asymptotiquement (fort RSB), les deux stratégies ont des performances équivalentes car la première transmission passe alors avec succès de manière quasi-certaine, et elle est identique en terme de rendement (R_o) dans les deux cas. En revanche, on constate que l'approche SP s'avère nettement plus robuste à faible RSB. En HARQ-SP, le rendement vu au m^e round est égal $R_o/(m+1)$, contre $R_o \frac{L}{L+m}$ bits/symbole en HARQ-MP. Ceci explique la différence de paliers de débits entre les deux stratégies. Notons que nous avons choisi ici de fixer le même nombre de retransmissions dans les deux cas ($M_1 = M_2$). On aurait également pu choisir de fixer le paramètre M_2 à la valeur LM_1 , de manière à avoir le même rendement final que le protocole SP au dernier round. Dans ce cas, nos analyses ont montré que le débit utile en HARQ-MP est alors supérieur à celui du SP sur toute la plage de RSB.

Transmission sur canal de Rayleigh à évanouissements par blocs

La figure 4.3 compare le débit utile moyen limite en HARQ-SP et en HARQ-MP avec des codes gaussiens de longueur infinie et rendement $R_o = 1$ bit/symbole, pour une transmission sur un canal de Rayleigh à évanouissements par blocs. Le nombre maximal de retransmissions est fixé à 2 pour les deux stratégies. Dans le cas du protocole MP, nous avons fait varier par ailleurs le nombre L de paquets traités conjointement par le protocole multi-paquets. Comme sur le canal gaussien, on voit que le protocole MP surpasse le SP à moyen et fort RSB. Le gain est d'autant plus important que la valeur de L est élevée. A nouveau, l'approche IR simple-paquet classique s'avère plus robuste à faible RSB. Ceci s'explique par le fait qu'à RSB fixé, un succès de transmission en MP au premier round est d'autant moins probable que L est élevé (le cas $L = 1$ correspond au protocole HARQ-SP). Lorsque le RSB est faible, il est préférable de cumuler toute l'énergie apportée par la redondance sur un seul paquet de données, afin de maximiser la probabilité de réussir à décoder. A moyen et fort RSB, en revanche, l'énergie apportée

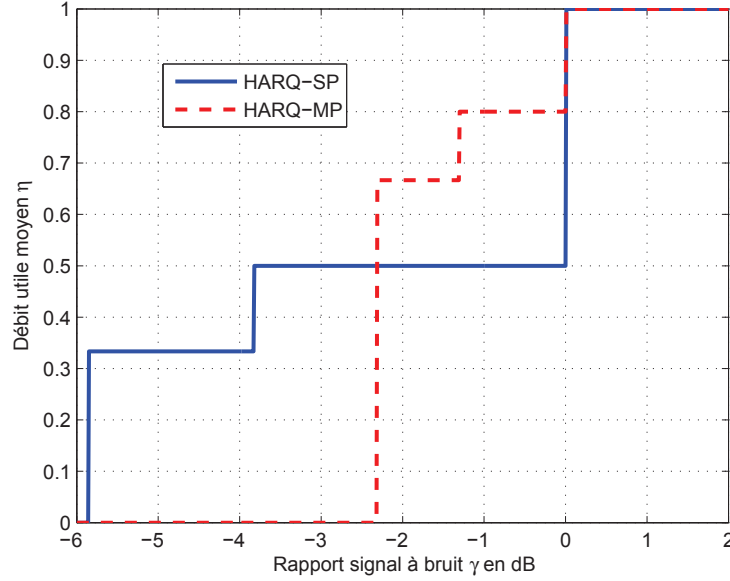


Figure 4.2 — Débit utile en HARQ-SP et HARQ-MP pour une transmission sur canal BABG complexe avec codes gaussiens de longueur infinie et rendement $R_o = 1$ bit/symbole, avec $M_1 = 2$, $M_2 = 2$ et $L = 4$.

par un paquet de redondance peut suffire à décoder plusieurs paquets de données, d'où l'avantage asymptotique (gain en débit) du MP. Le paramètre L sera fixé à la valeur 4 dans la suite de ce chapitre.

4.4 Schémas de codage pratiques pour le MP

L'analyse théorique précédente a démontré que, suivant la zone de RSB dans laquelle on se situe, l'approche HARQ-MP peut effectivement apporter un gain en terme du débit utile moyen par rapport à la stratégie SP classique. Cette analyse avait toutefois pour but d'évaluer les performances limites de la stratégie multi-paquets, et se basait sur des codes aléatoires, faciles à analyser, mais pas assez structurés pour être utilisés en pratique. L'objectif de cette section consiste donc à proposer différentes solutions de codage permettant de mettre en œuvre la stratégie MP en pratique. Les performances des schémas de codage proposés sont évaluées par simulation, et comparées à celles de schémas HARQ-SP pratiques, ainsi qu'aux performances limites de la section précédente.

4.4.1 Principe général de la construction proposée

Le schéma de codage que nous proposons est présenté sur la figure 4.4 et s'inspire directement du schéma concaténé générique introduit dans [47]. Il est formé de la concaténation d'un code externe $\mathcal{C}_o$ de rendement R_o et d'un code interne $\mathcal{C}_{is}$, de

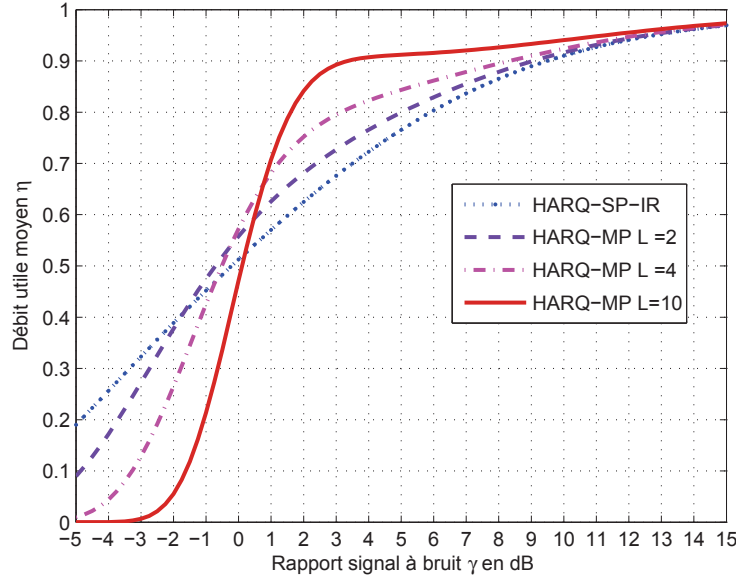


Figure 4.3 — Débit utile en HARQ-SP et HARQ-MP pour une transmission sur canal de Rayleigh à évanouissements par blocs avec codes gaussiens de longueur infinie et rendement $R_o = 1$ bit/symbole, avec $M_1 = 2$, $M_2 = 2$ et pour différentes valeurs du nombre L de paquets de données considérés en MP.

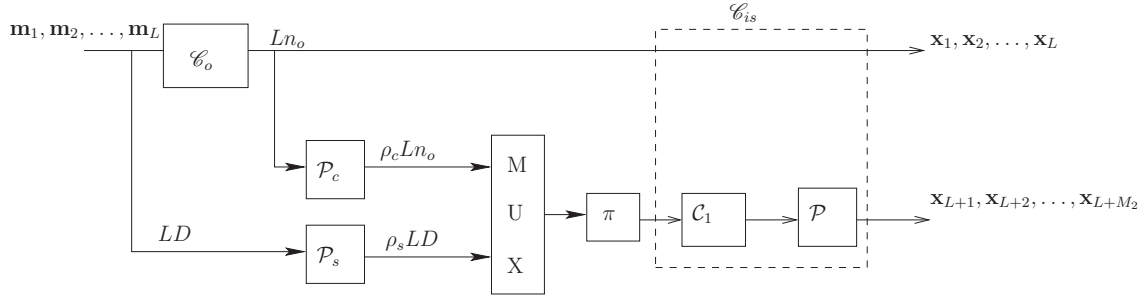


Figure 4.4 — Schéma de principe de la construction du code MP

rendement $R_i = \frac{L}{L+M_2}$. On n'impose aucune contrainte particulière sur le code externe $\mathcal{C}_o$. En revanche, le code interne $\mathcal{C}_{is}$ est supposé systématique, et produit sa redondance de manière incrémentale, par poinçonnage compatible en rendement.

La procédure de codage/décodage est la suivante. Au premier round, L messages de données successifs $(\mathbf{m}_1, \dots, \mathbf{m}_L)$ sont encodés un par un par le code $\mathcal{C}_o$ externe, pour donner L mots de code $(\mathbf{x}_1, \dots, \mathbf{x}_L)$. Ces L mots de code sont modulés, puis transmis successivement à la station B. Celle-ci décode séparément les L paquets reçus. En cas d'échec de décodage sur l'un ou plusieurs des paquets, la station B renvoie un acquittement négatif pour en informer la station A. Aux rounds suivants, les L messages de données ainsi que les L mots de codes produits au premier round sont poinçonnés séparément, puis multiplexés (on suppose ici que les messages et les mots de code sont binaires, la modulation vient après l'encodage FEC et n'apparaît pas sur la figure 4.4).

Les masques de poinçonnage appliqués aux données et aux mots de code sont notés respectivement $\mathcal{P}_s$ et $\mathcal{P}_c$. Les paramètres ρ_s et ρ_c désignent quant à eux la proportion de bits d'information et bits codés qui survit au poinçonnage. La sortie du multiplexeur est entrelacée, puis encodée par un second code noté $\mathcal{C}_1$. Comme illustré sur la figure 4.4, ce dernier produit uniquement la redondance du code interne $\mathcal{C}_{is}$, la partie systématique étant composée des L mots de code $(\mathbf{x}_1, \dots, \mathbf{x}_L)$ transmis au premier round. La séquence codée en sortie du code $\mathcal{C}_1$ est finalement poinçonnée de manière incrémentale, sous contrainte de compatibilité de rendement, par le masque $\mathcal{P}$, pour créer jusqu'à M_2 paquets de redondance distincts $(\mathbf{x}_{L+1}, \dots, \mathbf{x}_{L+M_2})$, à transmettre si nécessaire aux rounds $m = 1, \dots, M_2$.

La structure de la figure 4.4 est très générale. Notons que l'on retrouve une concaténation série classique en posant $\rho_s = 0$ et $\rho_c = 1$. A l'inverse, on obtient un schéma de concaténation parallèle en prenant $\rho_s = 1$ et $\rho_c = 0$.

4.4.2 Exemples de construction

Sur la base de la construction générale de la figure 4.4, nous proposons ici trois schémas de codage pratiques implémentant l'approche MP. Le premier s'appuie sur une concaténation parallèle de codes convolutifs (*Parallel Concatenated Convolutional Codes*, ou PCCC). Les deux autres schémas reposent une concaténation série : concaténation série d'un code convolutif et d'un accumulateur (*Serial Concatenation of Convolutional and Accumulator codes* ou SCCA), et concaténation série ligne-colonne (code produit) d'un code convolutif et d'un code de parité, appelée ici *Turbo produit paquet* (TPP). Dans les trois cas, on met en œuvre un décodage itératif de type turbo à partir du second round, afin d'exploiter pleinement la capacité de correction des codes proposés.

A des fins de comparaison, nous décrivons au préalable deux schémas HARQ-IR simple-paquet basés respectivement sur un code convolutif poinçonné, et sur une concaténation parallèle de codes convolutifs. Tous les schémas de codage présentés se basent sur une transmission QPSK, utilisant un code convolutif de rendement 1/2 à 8 états construit à partir des polynômes 15 et 17 (en octal). Le nombre maximal de rounds est fixé à 3, que ce soit en SP ou bien en MP. Enfin, les schémas MP opèrent sur des groupes de $L = 4$ paquets de données consécutifs.

Schéma HARQ-SP avec codes convolutifs (SP-CC)

La mise en œuvre de schémas HARQ-II-IR sur la base de codes convolutifs poinçonnés a été discuté dans le chapitre 3. Dans le cadre de la comparaison SP/MP, le code mère retenu est le code convolutif de rendement 1/6, à 8 états, de polynômes générateurs (15,17,13,15,17,13). La redondance incrémentale est générée par le processus de poinçonnage régulier suivant. Au premier round, on ne transmet que la redondance correspondant aux bits codés produits par les polynômes (15, 17). On envoie ensuite successivement aux rounds 1 et 2 les bits de redondance générés par les polynômes (13,15), puis (15,17). A chaque round, le récepteur effectue un décodage de

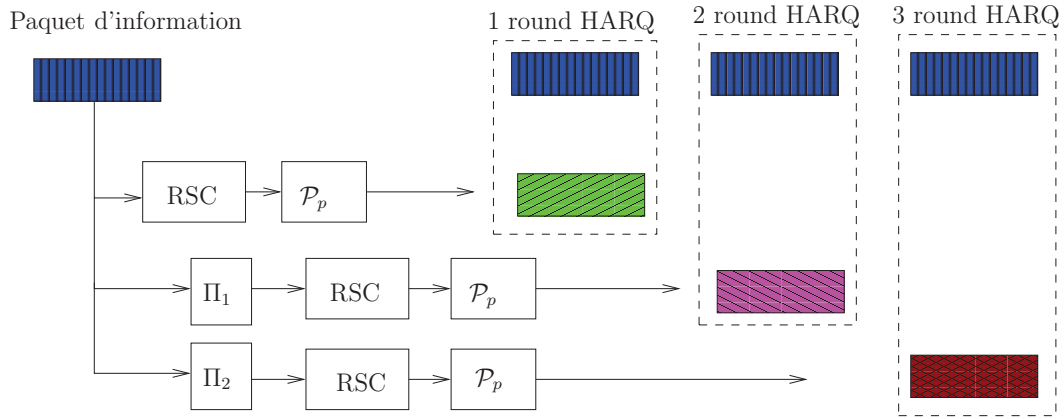


Figure 4.5 — Schéma HARQ-SP basé sur une concaténation parallèle de codes convolutifs (SP-PCCC)

Viterbi sur le treillis du code mère de rendement 1/6.

Schéma HARQ-SP avec concaténation parallèle de codes convolutifs (SP-PCCC)

Le second schéma SP considéré dans cette comparaison vise à créer un code plus puissant qu'un simple code convolutif à redondance incrémentale. Il est basé sur une concaténation parallèle de codes convolutifs (SP-PCCC) à l'émission, complété par un décodage itératif (turbo) en réception à partir du second round. Au premier round, le paquet d'information est encodé par le code convolutif récurrent systématique (RSC) de polynômes générateurs (1,15/17) et de rendement 1/2. Le code vu par la station B au premier round est donc équivalent à celui vu dans le scénario précédent (SP-CC). La station B utilise ici un décodeur Maximum *a posteriori* (MAP) pour décoder le paquet reçu. En cas d'erreurs détectées, un nouveau paquet de redondance est construit par entrelacement du message d'information (permutation Π_1), puis encodage du message entrelacé et poinçonnage des bits systématiques. Comme indiqué sur la figure 4.5, le nouveau paquet est formé de la répétition du message d'information, complété par la nouvelle redondance. Le code équivalent vu par la station B est alors identique à un turbo code parallèle [48], avec deux observations différentes relatives au message d'information, et deux séquences de redondance distinctes. La station B combine les bits systématiques (*Chase Combining*) avant de procéder au décodage itératif de l'ensemble. Si les erreurs persistent, on construit une nouvelle retransmission (round 3) de manière similaire, à l'aide d'une seconde fonction de permutation Π_2 . La station B combine les 3 observations relatives aux bits systématique, puis effectue ensuite un décodage turbo multiple de l'ensemble [49]. Dans nos simulations, les entrelaceurs Π_1 et Π_2 ont été tirés aléatoirement à chaque nouveau cycle HARQ.

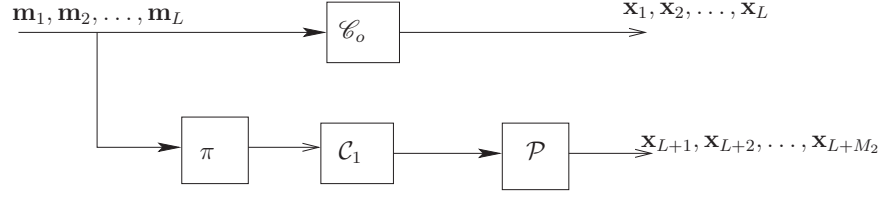


Figure 4.6 — Schéma HARQ-MP basé sur une concaténation parallèle de codes convolutifs (MP-PCCC)

Schéma HARQ-MP avec concaténation parallèle de codes convolutifs (MP-PCCC)

Le premier schéma HARQ-MP considéré est une concaténation parallèle de codes convolutifs (MP-PCCC). Il est présenté sur la figure 4.6, et se base sur la structure générale de la figure 4.4 en prenant $\rho_s = 1$ et $\rho_c = 0$. De la même manière que dans les deux schémas SP précédents, le code externe $\mathcal{C}_o$ est le code RSC à 8 états de rendement 1/2 et polynôme générateur (1,17/15). Au premier round, les L paquets de données sont encodés séparément par $\mathcal{C}_o$, puis modulés et transmis successivement sur le canal. La station B décode chaque paquet séparément, à l'aide d'un décodeur MAP. En cas d'échec sur l'un ou plusieurs de ces paquets, les messages de données sont entrelacés puis encodés une seconde fois par le code $\mathcal{C}_1$, constitué ici du code RSC à 8 états de polynôme générateur (17/15) (code générant uniquement la séquence de parité du code RSC (1,17/15)). On obtient ainsi une séquence de redondance de longueur L fois supérieure à la taille des paquets codés transmis au premier round. On vient alors appliquer un masque de poinçonnage $\mathcal{P}$ de manière à en extraire M_2 nouveaux paquets de redondance. Ces derniers sont transmis à B à raison d'un nouveau paquet par round. A partir du second round, le code équivalent vu par la station B est donc un turbo code parallèle [48]. Toutefois, à la différence du schéma SP-PCCC présenté auparavant, chaque paquet de redondance s'obtient ici par encodage *conjoint* des L paquets de données.

Schéma HARQ-MP avec concaténation série d'un code convolutif et d'un accumulateur (MP-SCCA)

Le second schéma HARQ-MP proposé est présenté sur la figure 4.7 et repose sur la concaténation série d'un code convolutif avec un accumulateur (MP-SCCA, *Serial Concatenation of Convolutional and Accumulator codes*). Il se déduit de la structure générale de la figure 4.4 en prenant $\rho_s = 0$ et $\rho_c = 1$. Le code externe est à nouveau le code RSC de rendement 1/2 à 8 états. Les L paquets transmis au premier round sont produits de la même façon que pour le schéma MP-PCCC, et décodés un par un par la station B (décodeur MAP). Seule change la génération des paquets de redondance à partir du second round. Ceux-ci sont obtenus par entrelacement aléatoire des L paquets codés en sortie de $\mathcal{C}_o$, puis encodage conjoint par le code interne $\mathcal{C}_1$ constitué ici d'un code accumulateur de rendement unitaire, à deux états. On applique finalement un masque de poinçonnage $\mathcal{P}$ à la sortie de l'accumulateur pour en extraire M_2 paquets

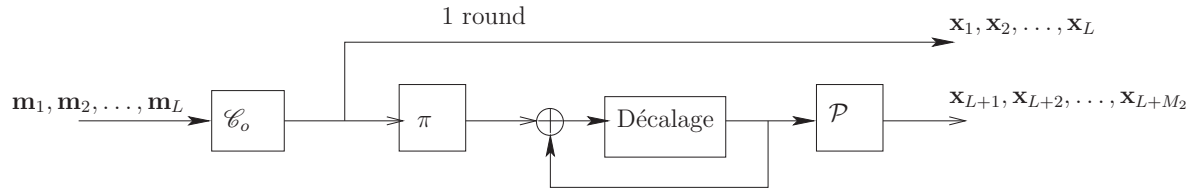


Figure 4.7 — Schéma HARQ-MP basé sur une concaténation série d'un code convolutif et d'un accumulateur (MP-SCCA)

de redondance. A partir du second round, la station B effectue un décodage itératif de l'ensemble des paquets reçus pour tenter de décoder conjointement les L paquets de données. Dans nos simulations, l'entrelaceur Π a été tiré aléatoirement à chaque nouveau cycle HARQ.

Schéma HARQ-MP turbo produit paquet (MP-TPP)

Le dernier schéma HARQ-MP proposé est une forme très particulière de concaténation série appelée code produit, associant ici le code convolutif RSC $\mathcal{C}_o$ à 8 états en ligne, avec un code de parité de dimension L en colonne. Durant le premier round, chaque paquet de donnée est encodé par $\mathcal{C}_o$, modulé, et transmis à B. Celle-ci décode séparément les différents paquets. En cas d'échec, les L paquets codés générés au premier round sont disposés en ligne dans une matrice. On vient alors créer un paquet de redondance supplémentaire de même taille que les paquets précédents, en appliquant le code de parité (somme modulo-2) sur chaque colonne, comme indiqué sur la figure 4.8. A réception de ce paquet de redondance, la station B effectue un décodage itératif du code produit (turbo décodage) pour tenter de décoder conjointement les L paquets de données [50] [51]. Si le décodage échoue de nouveau, on peut alors imaginer plusieurs solutions pour générer des paquets de redondance supplémentaires. L'une des plus simples consiste à répéter l'un des paquets erronés (la station B combine alors les observations relatives à chaque paquet, avant de procéder au turbo décodage de la matrice produit). La solution retenue dans nos simulations exploite quant à elle le fait que le code convolutif initial est de rendement $1/2$, et construit des paquets de redondance par concaténation des parties systématique de deux paquets erronés. Dans tous les cas, ceci suppose que la station B renvoie des acquittements identifiant les paquets mal décodés.

4.4.3 Performances des schémas pratiques

Nous présentons ici les performances simulées des différents schémas de codage MP proposés, et nous les comparons à celles des deux schémas SP décrit précédemment, ainsi qu'aux performances limites obtenues en section 4.3.3. Les résultats sont tout d'abord donnés pour un canal BABG, puis pour un canal de Rayleigh à évanouissements par blocs. Ces résultats ont été publiés en [1].

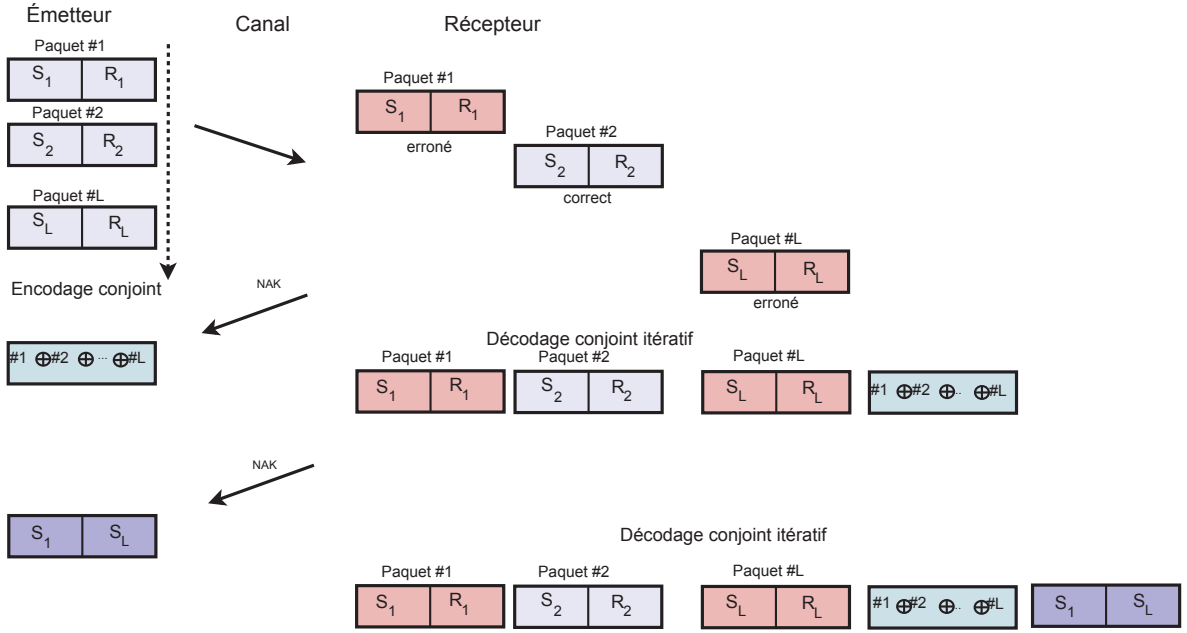


Figure 4.8 — Illustration du processus de retransmission dans le cas du schéma HARQ-MP-TPP (concaténation série d'un code convolutif et d'un code de parité, S_i/R_i : partie systématique/parité du paquet i)

Canal BABG

La figure 4.9 présente le débit utile moyen obtenu avec chacun des 5 schémas de codage précédent (2 SP + 3 MP). On s'est placé ici dans le cas d'une transmission QPSK sur canal gaussien complexe, avec un maximum de 3 rounds en SP et MP, une taille de données utiles $D = 1024$ bits, et une taille d'entête $H = 40$ bits (rendement $R_o = 0,98$ bits/symbole). A titre de référence, nous avons également tracé en pointillés les performances limites obtenues avec des codes aléatoires gaussiens asymptotiquement longs. En comparant les débits offerts par les schémas SP et MP, on constate effectivement qu'en accord avec les prédictions théoriques de la section 4.3.3, les schémas MP surpassent les schémas SP à moyen et fort RSB. Que ce soit en SP ou bien en MP, on voit également que la concaténation parallèle (PCCC) présente ici les meilleures performances. Parmi les trois schémas MP, l'approche MP-TPP a le débit le moins bon, mais reste toutefois meilleur que les schémas SP à moyen et fort RSB. En revanche, ces derniers conservent l'avantage à faible RSB.

Afin de mieux comprendre les résultats de la figure 4.9, nous présentons sur les figures 4.10 et 4.11 le taux d'erreur paquet (TEP) résiduel en sortie de chacun des schémas SP et MP, à chaque round. Les schémas présentant les probabilités d'erreur paquet les plus faibles sont les schémas qui auront les débits les plus élevés. Au premier round, le TEP résiduel correspond au taux d'erreur paquet du code utilisé. Il est identique pour les cinq schémas SP/MP car ils sont tous basés sur le code convolutif à 8 états et de rendement $1/2$. Aux rounds $m \geq 1$, on définit le TEP résiduel comme le rapport du nombre moyen de paquets erronés au round concerné, sur le nombre total de paquets d'information effectivement transmis. Dans le cas des schémas SP, on voit

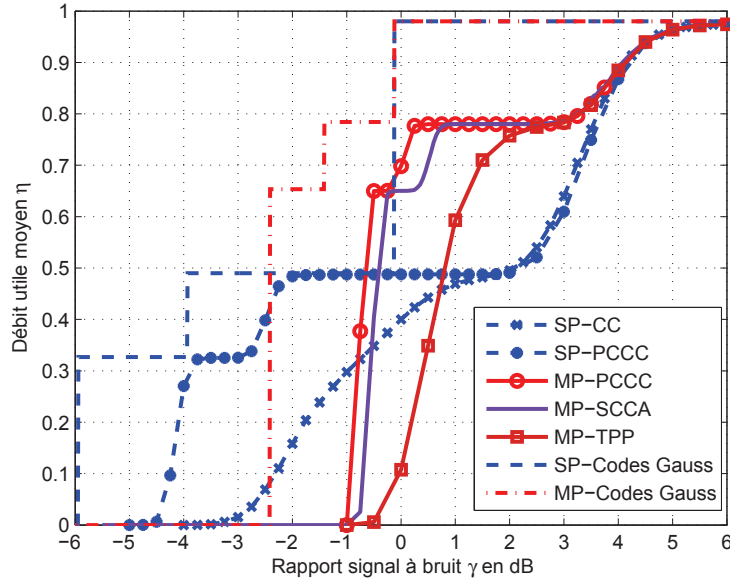


Figure 4.9 — Débit utile normalisé en fonction du SNR pour une transmission QPSK et un protocole HARQ-SP ou HARQ-MP sur un canal BABG complexe avec $L = 4$, $M_1 = 2$, $M_2 = 2$, $D = 1024$ bits, et $H = 40$ bits ($R_o = 0,98$ bits/symbole).

sur la figure 4.10 que le schéma SP-PCCC présente effectivement le TEP le plus faible. On observe un résultat similaire dans le cas des schémas MP, sur la figure 4.11. On peut noter dans le cas que la concaténation série MP-SCCA surpasse la concaténation parallèle MP-PCCC à fort RSB (apparition d'un plancher d'erreur pour ce dernier), mais l'avantage apparaît à des taux d'erreurs trop faibles pour que cela ait un impact notable sur la courbe du débit utile moyen.

Canal de Rayleigh à évanouissements par blocs

La figure 4.12 compare le débit utile moyen des schémas de codage SP et MP pour une transmission QPSK sur canal de Rayleigh à évanouissement par blocs, en conservant les mêmes paramètres de transmission que précédemment. On constate à nouveau que les schémas MP sont meilleurs pour les zones de moyens et forts RSB, conformément aux résultats établis en section 4.3.3. Pour ces schémas pratiques, la concaténation parallèle offre le meilleur débit, que ce soit en SP ou en MP. On note également que le débit du schéma MP-TPP est inférieur à celui du schéma SP-PCCC.

En complément, les figures 4.13 et 4.14 présentent le TEP résiduel à chaque round, pour chacun des schémas SP et MP. Il est intéressant de noter ici que tous les schémas bénéficient d'un gain de diversité à chaque round, chaque nouveau paquet transmis voyant une réalisation différente du canal. On peut noter également les mauvaises performances du schéma MP-TPP sur ce canal.

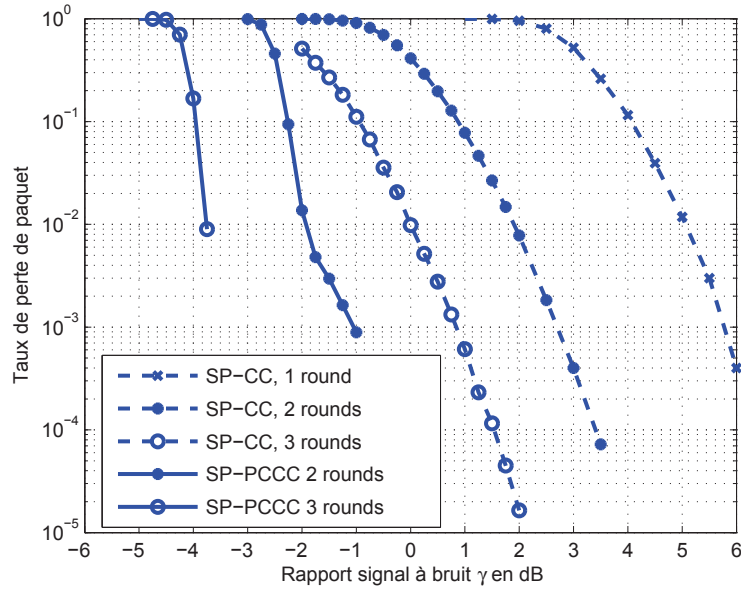


Figure 4.10 — Taux d'erreur paquet pour différents schémas HARQ-SP, dans le cas d'une transmission QPSK sur un canal BABG complexe, avec $M_1 = 2$, $D = 1024$ bits, et $H = 40$ bits.

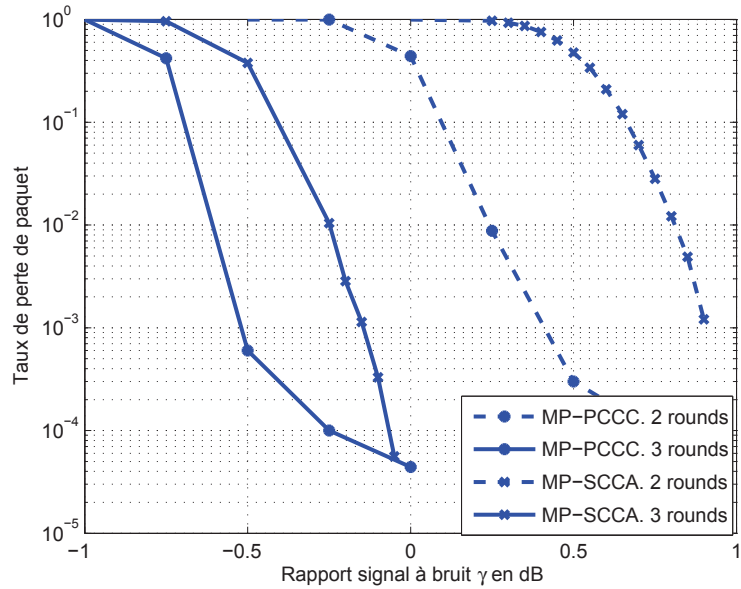


Figure 4.11 — Taux d'erreur paquet pour différents schémas HARQ-MP, dans le cas d'une transmission QPSK sur un canal BABG complexe, avec $M_1 = 2$, $D = 1024$ bits, et $H = 40$ bits.

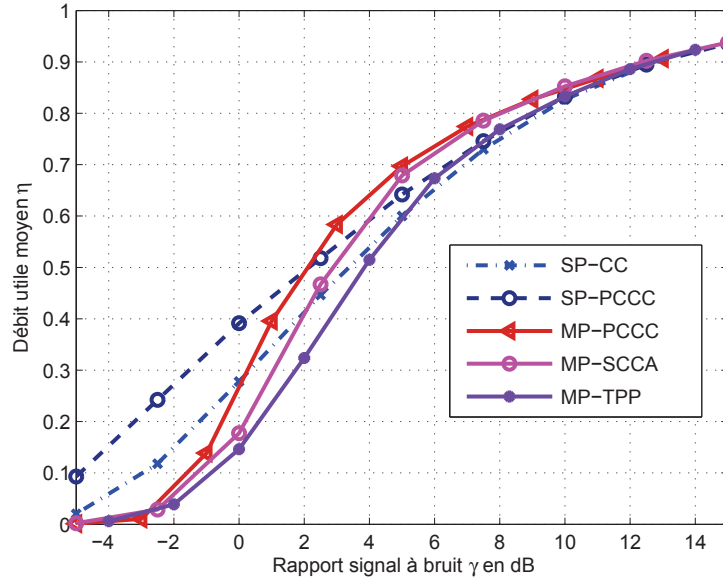


Figure 4.12 — Débit utile normalisé pour différents schémas HARQ-SP et HARQ-MP dans le cas d'une transmission QPSK sur canal de Rayleigh à évanouissements par blocs, avec $L = 4$, $M_1 = 2$, $M_2 = 2$, $D = 1024$ bits, et $H = 40$ bits ($R_o = 0,98$ bits/symbole)

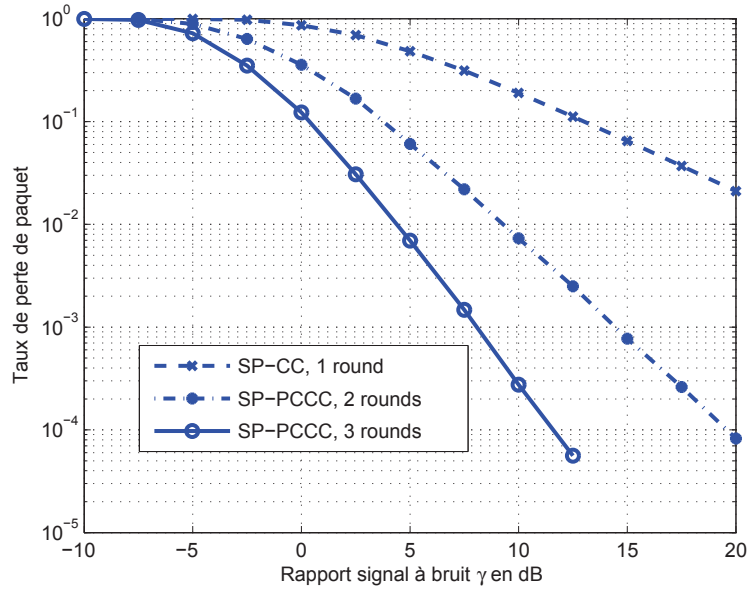


Figure 4.13 — Taux d'erreur paquet pour différents schémas HARQ-SP, dans le cas d'une transmission QPSK sur canal de Rayleigh à évanouissements par blocs, avec $M_1 = 2$, $D = 1024$ bits, et $H = 40$ bits

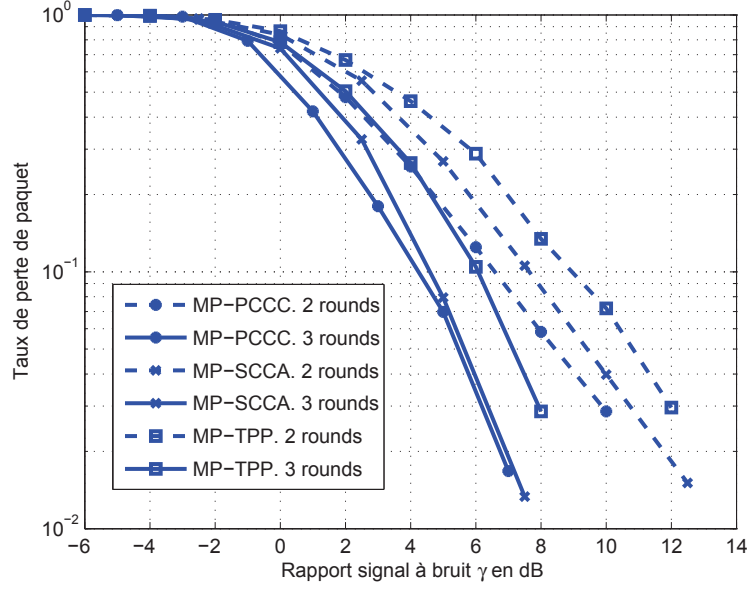


Figure 4.14 — Taux d'erreur paquet pour différents schémas HARQ-MP, dans le cas d'une transmission QPSK sur canal de Rayleigh à évanouissements par blocs, avec $M_1 = 2$, $D = 1024$ bits, et $H = 40$ bits

Remarques

Les simulations précédentes appellent quelques remarques complémentaires. On constate en effet un écart parfois relativement important entre les performances des meilleurs schémas SP/MP, et les performances limites correspondantes. Ceci est dû à plusieurs causes. Tout d'abord, les codes simulés ici sont relativement courts (basés sur des paquets de 1024 bits d'information, hors entête), et ne peuvent donc pas être comparés équitablement avec des codes asymptotiquement longs. D'autre part, tous les schémas proposés (SP ou MP) se réduisent au même code convolutif à 8 états au premier round, à la capacité de correction limitée. De bien meilleures performances pourraient être obtenues en remplaçant ce code convolutif par un code plus puissant, tel qu'un turbo code ou bien un code LDPC. Enfin, les codes proposés (concaténation série ou parallèle) n'ont absolument pas été optimisés par rapport aux canaux de transmission considérés. En particulier, les entrelaceurs et les masques de poinçonnage ont été tirés aléatoirement à chaque nouveau cycle. Le choix du code convolutif externe n'a pas non plus reçu d'attention spécifique. Les résultats présentés restent donc nettement perfectibles.

4.5 Implémentation pratique et optimisation du protocole MP en présence d'erreurs sur la voie de retour

Nous avons étudié les performances limites des schémas MP, validé ainsi son intérêt, et proposé des schémas de codage permettant d'implémenter cette stratégie en pratique. L'étude précédente supposait toutefois une transmission avec voie de retour sans erreurs. Autrement dit, les acquittements arrivent toujours correctement à la station A. Dans cette section, nous nous plaçons dans le cas plus réaliste d'une transmission avec voie de retour imparfaite, et nous étudions l'impact des erreurs sur la voie de retour sur les performances du protocole MP. La manière d'implémenter le protocole (détail des règles de dialogue entre A et B) devient alors très importante.

Nous avons signalé en section 4.2.2 qu'il y avait principalement deux solutions possibles pour implémenter le protocole MP. Ces deux solutions diffèrent essentiellement sur la notification d'un échec de décodage au premier round. La première solution consistant à renvoyer un seul acquittement pour l'ensemble des L paquets a été décrite en section 4.2.2. Dans la deuxième solution, la station B renvoie un acquittement pour chaque paquet. Le point en commun entre ces deux variantes est la retransmission de redondance portant sur plusieurs paquets de données utiles aux rounds $m > 0$. Ces deux variantes donnent des débits comparables lorsque la voie de retour est parfaite. Ce n'est plus le cas en présence d'erreurs sur la voie de retour.

Pour illustrer ce point, nous présentons tout d'abord en détail chacune des deux variantes du protocole MP. Nous les comparons ensuite en terme de débit utile moyen, afin d'identifier la variante la plus performante.

4.5.1 Protocole MP avec un acquittement pour L paquets au premier round

Cette première variante, initialement décrite en section 4.2.2, suppose que la station B renvoie un seul acquittement relatif à l'ensemble des L paquets transmis au premier round HARQ. Cet acquittement informe la station A de la bonne ou mauvaise réception des L paquets transmis au premier round. Lorsque cet acquittement est perdu, la station A fait l'hypothèse qu'il s'agit d'un acquittement négatif, et enclenche l'envoi de redondance MP. S'il s'agit effectivement d'un NAK, à réception du paquet de redondance, la station B effectue le décodage conjoint et renvoie un acquittement informant A du résultat. Dans le cas contraire (perte d'un acquittement positif), la station B ignore le paquet de redondance et renvoie à nouveau l'acquittement positif. De son côté, la station A continue à retransmettre des paquets de redondance tant qu'elle n'a pas reçu d'acquittement positif ou bien que le nombre maximal M_2 de retransmissions autorisées n'est pas atteint. Le fonctionnement de cette variante du MP en émission et en réception, en présence d'erreurs sur la voie de retour, est donc tout à fait similaire au fonctionnement de la stratégie SP, à la seule différence qu'en MP, la station A transmet L paquets au premier round au lieu d'un seul paquet en SP. Les

automates d'émission/réception du protocole sont donc identiques aux automatiques ARQ classiques. Par la suite, cette variante sera désignée par l'abréviation “*MP-Acq par L paquets*”.

4.5.2 Protocole MP avec un acquittement par paquet au premier round

Dans la seconde variante du protocole MP, appelée “*MP-Acq par paquet*”, la station A transmet son premier paquet de données. Si ce paquet est correctement décodé et si l'acquittement positif est bien reçu, la station A passe à la transmission du paquet suivant. Si le paquet est mal décodé et que l'acquittement négatif arrive correctement, la station A passe à la transmission des $L-1$ paquets de données suivants, puis retransmet nécessairement au moins un paquet de redondance MP, relatif aux L paquets de données déjà transmis. Lorsqu'aucun acquittement n'arrive, la station A suppose que le paquet a été correctement décodé, et passe à l'émission du paquet suivant, dans l'attente d'un prochain acquittement dont la valeur déterminera si les paquets émis précédemment ont bien été décodés ou non. Plus précisément, si le nouvel acquittement valide le deuxième paquet, la station A passe donc à la transmission du troisième paquet, comme si aucun acquittement n'avait été perdu. Lorsqu'aucun acquittement n'est reçu pour L paquets de données successifs, la station A enclenche alors la transmission de redondance MP. La figure 4.15 synthétise le fonctionnement à l'émission du protocole *HARQ-MP-Acq par paquet*.

Le fonctionnement en réception est simple, et similaire au protocole SP classique. Au départ, la station B décode séparément les paquets de données reçus, et renvoie un acquittement pour chacun des paquets. A réception de paquets de redondance en provenance de A, la station B décode conjointement les paquets de données et de redondance, et renvoie un acquittement relatif à l'ensemble des L paquets. La figure 4.16 résume le fonctionnement du protocole *HARQ-MP-Acq par paquet* en réception.

4.5.3 Performance en présence de pertes sur la voie de retour

En présence d'erreurs sur la voie de retour, les performances des schémas HARQ se dégradent, puisque dans certains cas, l'émetteur effectue des retransmissions inutiles. Afin de comparer les performances des schémas SP et MP on présente ici une méthode de calcul du débit utile pour des transmissions avec perte sur la voie de retour. L'objectif est double. Il s'agit d'une part de comparer les deux variantes du protocole MP entre elles, et d'autre part, de comparer la robustesse relative des schémas MP à celles des schémas SP dans ce contexte.

Rappelons tout d'abord que, pour les schémas HARQ-SP, comme nous l'avons vu dans le chapitre précédent, le débit utile moyen en présence de pertes sur la voie de retour peut être calculé à l'aide de la formule 3.18, rappelée ci-dessous :

$$\eta = R_o \frac{1 - p(M_1)}{1 + \sum_{m=0}^{M_1-1} P(m)}$$

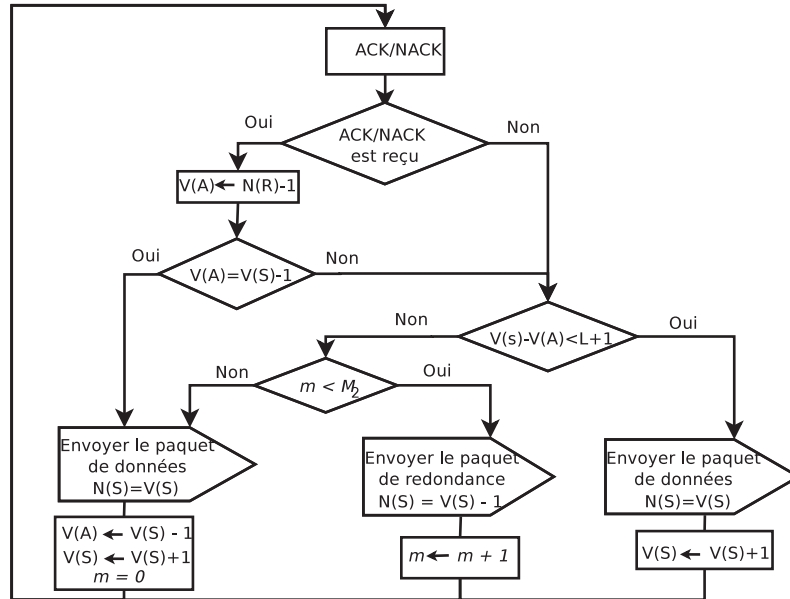


Figure 4.15 — Automate à l'émission de la variante *MP-Acq. par paquet* ($V(A)$ variable interne contenant le numéro de séquence du dernier paquet positivement acquitté, $N(S)$ numéro de séquence du paquet en cours, $V(S)$ numéro de séquence du prochain paquet à envoyer, $N(R)$ numéro de séquence porté par l'acquittement reçu)

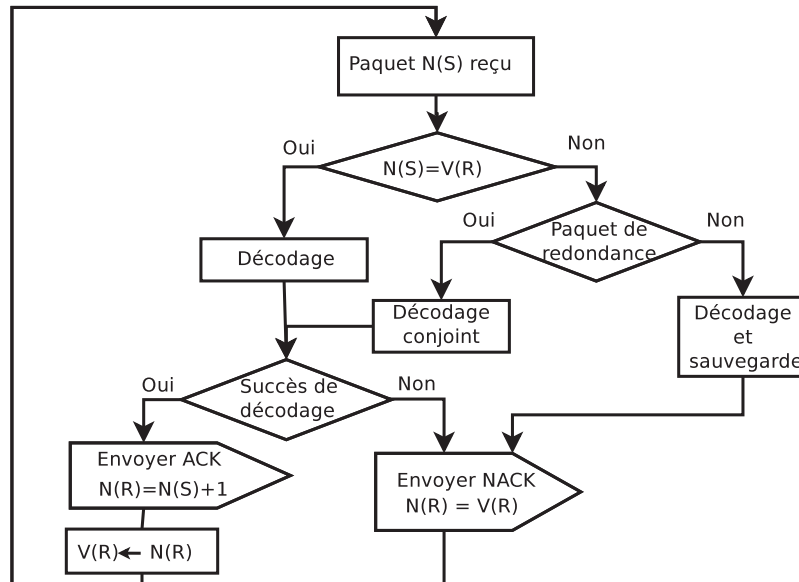


Figure 4.16 — Automate en réception de la variante *MP-Acq; par paquet* ($V(R)$ variable interne contenant le numéro de séquence porté par l'acquittement transmis, $N(S)$ numéro de séquence du paquet reçu, $N(R)$ numéro de séquence porté par l'acquittement envoyé à la station A)

Un échec de transmission aura lieu si l’acquittement est perdu, ou bien si le paquet est mal décodé et l’acquittement négatif correctement reçu. La probabilité d’échec de transmission au premier round $P(0)$ est donc donnée par :

$$P(0) = \epsilon + \Pr(e_0)(1 - \epsilon) \quad (4.17)$$

où $\Pr(e_0)$ désigne la probabilité d’échec de décodage au premier round, et ϵ est la probabilité de perte sur la voie retour. Plus généralement, nous avons proposé dans le chapitre précédent de calculer la probabilité conjointe d’un échec de transmission aux rounds 0 à m inclus par l’expression :

$$P(m) \approx \epsilon^{m+1} + \sum_{i=0}^m \Pr\{e_i\} \epsilon^{m-i} (1 - \epsilon) \quad (4.18)$$

où $\Pr(e_m)$ est la probabilité d’échec de décodage au round m .

Dans le cas du protocole MP avec un acquittement pour L paquets au premier round, la formule précédente s’applique à l’identique, en remplaçant M_1 par M_2 . Seule change la probabilité d’échec de décodage $\Pr(e_i)$. On s’intéresse ici aux performances limites des deux variantes du protocole HARQ-MP. On suppose donc l’utilisation de codes gaussiens asymptotiquement longs. Dans ce cas, l’approximation précédente devient une égalité stricte, et la probabilité d’échec de décodage est donnée par les expressions (4.8) et (4.10) sur le canal BABG, ou bien par (4.12) et (4.15) sur le canal de Rayleigh à évanouissements par blocs.

En revanche, le calcul exact du débit utile moyen pour le protocole HARQ-MP avec un acquittement par paquet au premier round est beaucoup plus difficile à mener analytiquement, car il est difficile de déterminer à quel instant la station A décide de transmettre les paquets de redondance MP. Cela reste donc un problème ouvert à ce jour. Dans le cadre de ce chapitre, les performances de cette variante seront donc évaluées uniquement par simulation. Nous présenterons dans le chapitre suivant une méthode semi-analytique permettant de mener à bien ce calcul.

Transmission sur canal BABG

La figure 4.17 compare le débit utile moyen en SP et MP pour une transmission sur canal BABG complexe avec des codes gaussiens asymptotiquement longs, avec $M_1 = M_2 = 2$, $R_o = 1$ bit/symbole et $\epsilon = 0, 1$. On observe que les schémas MP sont asymptotiquement plus robustes aux erreurs sur la voie de retour que les schémas SP, au sens où le débit utile maximal atteint à fort RSB est plus élevé en MP. D’autre part, des deux variantes considérées, le schéma “*MP-Acq par paquet*” s’avère être le schéma le plus robuste aux pertes d’acquittements. Afin de mieux visualiser la supériorité de cette approche, on présente sur la figure 4.18 les résultats obtenus avec une probabilité de perte d’acquittements plus élevée ($\epsilon = 0, 3$). Les performances en retrait de la première variante s’expliquent par le fait que pour ce dernier, la perte de l’acquittement au premier round est particulièrement critique au sens où elle conduit systématiquement à l’émission de redondance MP, ce qui baisse en retour le débit utile moyen, phénomène d’autant plus visible à moyen et fort RSB. Dans le cas de la seconde variante, à fort

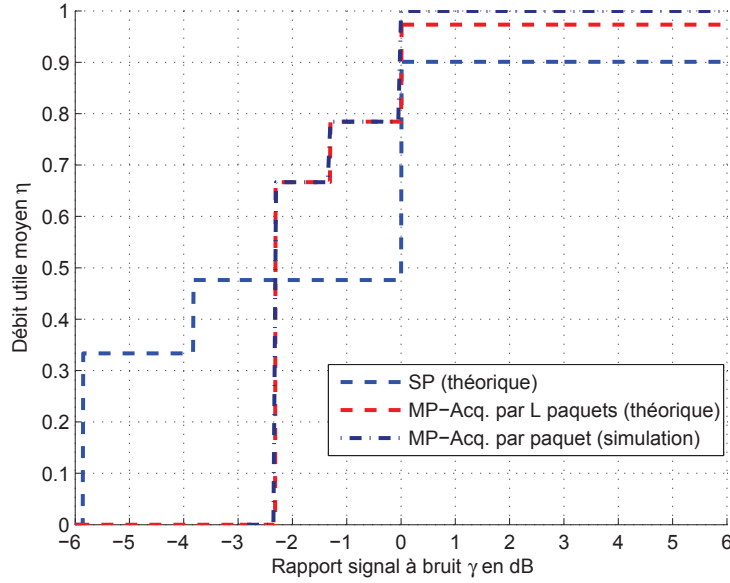


Figure 4.17 — Débit utile moyen en HARQ-SP et HARQ-MP (deux variantes), pour une transmission sur canal BABG complexe avec des codes gaussiens aléatoires de longueur infinie et rendement $R_o = 1$, $L = 4$, $M_1 = 2$, $M_2 = 2$, $R_o = 1$ et $\epsilon = 0,1$

RSB, il faut que L acquittements successifs soient perdus pour enclencher la transmission de redondance, ce qui est beaucoup plus improbable. Autrement dit, les pertes d'acquittements sont alors transparentes pour la grande majorité, et l'émetteur ne transmet aucune redondance inutile.

Transmission sur canal de Rayleigh à évanouissements par blocs

Les performances obtenues en HARQ-SP et HARQ-MP (deux variantes) dans le cas du canal de Rayleigh à évanouissements par blocs sont présentées sur les figures 4.19 et 4.20, en considérant respectivement une probabilité de perte $\epsilon = 0,1$ et $\epsilon = 0,3$. Tous les autres paramètres sont identiques à ceux du scénario précédent. Dans le cas de la variante MP avec acquittement pour L paquets, nous avons présenté à la fois les prédictions théoriques du débit, calculées à partir (4.7), ainsi que les performances simulées. Pour la seconde variante MP, avec acquittement pour chaque paquet, seuls les résultats de simulation sont donnés. On voit ici de nouveau l'avantage asymptotique (fort RSB) du protocole MP, et plus particulièrement de la variante avec acquittement pour chaque paquet. Le SP reste toutefois plus intéressant à faible RSB. En effet, dans cette zone, il est plus probable que l'acquittement perdu soit un acquittement négatif. La retransmission d'un paquet de redondance relatif à un seul paquet de donnée ne réduit pas le débit, mais au contraire maximise les chances de bien décoder. Dans le cas des schémas MP, à faible RSB, il est très probable que plusieurs paquets soient affectés par le canal. La retransmission d'un seul paquet de redondance MP ne suffit pas nécessairement pour décoder l'ensemble des paquets perdus. Ce n'est plus vrai à fort RSB.

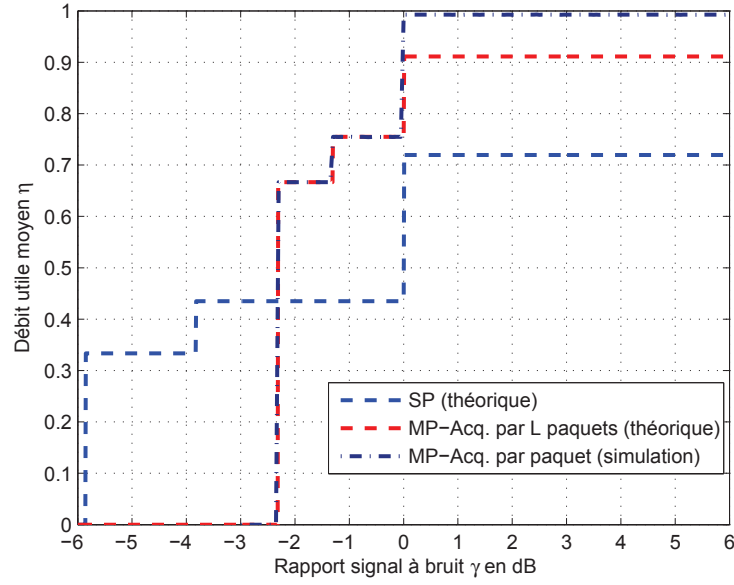


Figure 4.18 — Débit utile moyen en HARQ SP et HARQ-MP (deux variantes), pour une transmission sur canal BABG complexe avec des codes gaussiens aléatoires de longueur infinie et rendement $R_o = 1$, $L = 4$, $M_1 = 2$, $M_2 = 2$, $R_o = 1$ et $\epsilon = 0,3$

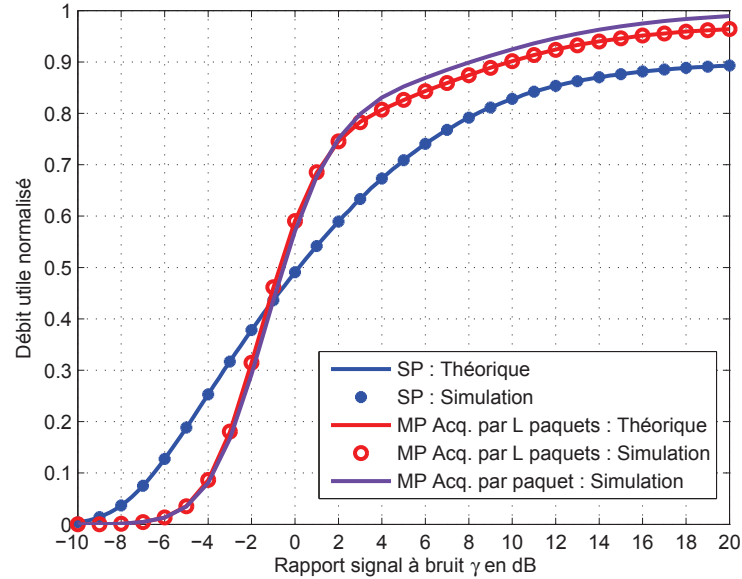


Figure 4.19 — Débit utile moyen en HARQ-SP et HARQ-MP (deux variantes), pour une transmission sur canal de Rayleigh à évanouissements par blocs avec des codes gaussiens aléatoires de longueur infinie et rendement $R_o = 1$, $L = 4$, $M_1 = 2$, $M_2 = 2$ et $\epsilon = 0,1$

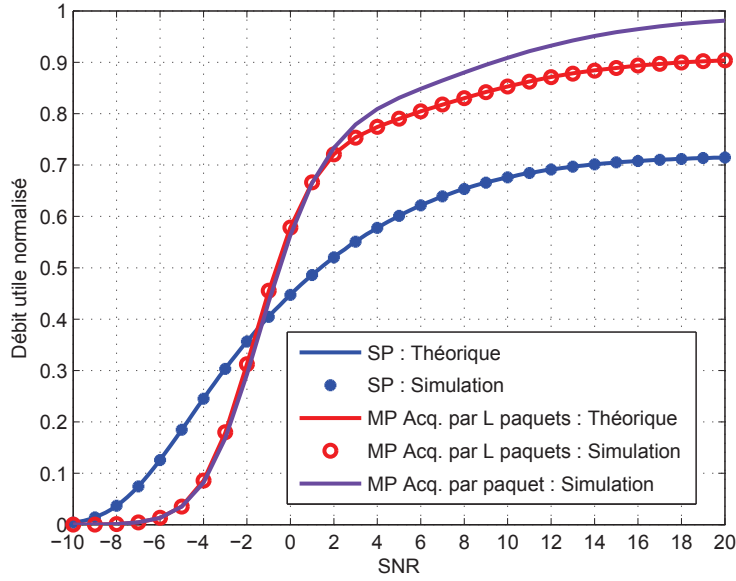


Figure 4.20 — Débit utile moyen en HARQ-SP et HARQ-MP (deux variantes), pour une transmission sur canal de Rayleigh à évanouissements par blocs avec des codes gaussiens aléatoires de longueur infinie et rendement $R_o = 1$ $L = 4$, $M_1 = 2$, $M_2 = 2$ et $\epsilon = 0,3$

4.6 Conclusion

Dans ce chapitre, nous avons proposé une stratégie HARQ-II multi-paquets (MP) à redondance incrémentale, permettant d'améliorer le débit utile moyen par rapport aux schémas HARQ-II-IR classiques (SP), par la retransmission de paquets de redondance pouvant aider au décodage simultané de plusieurs paquets de données erronés. Nous avons présenté une comparaison du débit utile moyen en MP et en SP en considérant à la fois des schémas HARQ théoriques et des schémas pratiques. De cette comparaison, il ressort que la stratégie HARQ-MP offre le débit le plus élevé à moyen et fort RSB. En revanche, l'approche SP reste plus performante à faible RSB.

Dans un deuxième temps, nous avons analysé les performances des stratégies MP et SP en présence d'erreurs sur la voie de retour. Ceci nous a conduit à comparer deux variantes du protocole MP, pour conclure que la variante consistant à renvoyer un acquittement pour chaque paquet reçu au premier round s'avère être la plus robuste aux erreurs sur la voie de retour. De plus, cette variante du protocole MP surpasse nettement le protocole classique SP à fort RSB.

En conclusion, les schémas MP restent toujours moins bons que l'approche classique à faible RSB. Il semble donc intéressant de chercher à combiner ces deux approches afin d'avoir le meilleur débit possible sur toute la plage de RSB. D'autre part, du fait de la complexité de l'analyse de la variante MP avec acquittement par paquet, l'évaluation des performances de cette approche n'a pu se faire que par simulation. Ces deux points sont donc l'objet du chapitre suivant, où nous proposons une combinaison SP/MP

ainsi qu'une méthode semi-analytique très générale pour évaluer les performances des schémas HARQ.

CHAPITRE

5

Schéma HARQ hybride
SP et MP

5.1 Introduction

Dans le chapitre précédent, nous avons proposé une stratégie de construction de paquets de redondance *multi-paquets* (MP). Nous avons montré qu'à moyen et fort RSB, l'approche multi-paquet surpasse l'approche HARQ-II classique dite *simple paquet* (SP) en terme de débit utile. D'autre part, nous avons montré que les schémas MP sont plus robustes que les schémas SP aux erreurs sur la voie de retour. Ces derniers offrent toutefois un débit plus élevé que les schémas MP à faible RSB. Dans ce chapitre, nous proposons une stratégie hybride SP et MP, notée par la suite H-SP-MP, qui vise à tirer profit des avantages propres à chacune de ces deux approches, suivant la zone de RSB dans laquelle on se situe. Le protocole hybride ne fait aucune hypothèse quant à la fiabilité sur la voie de retour, et s'applique donc naturellement à une voie de retour imparfaite. Nous introduisons également une méthode analytique permettant de prédire les performances théoriques de ce protocole.

Ce chapitre est organisé comme suit. Nous introduisons tout d'abord le principe général du protocole hybride SP et MP. Nous présentons ensuite le modèle d'analyse des performances, et nous l'utilisons pour prédire les performances limites du protocole hybride. Dans sa version basique, le protocole hybride n'utilise pas la connaissance du canal en réception. Or cette information est disponible dans la plupart des systèmes radio actuels. Ceci nous conduit à présenter et comparer deux variantes du protocole hybride qui exploitent la présence d'une information sur le RSB pour adapter la stratégie de retransmission. Enfin, nous regardons, suivant le critère de maximisation du débit utile, la meilleure façon de répartir la puissance d'émission entre le lien direct et la voie de retour, à budget d'énergie constant.

5.2 Principe du protocole Hybride SP et MP

Les schémas HARQ-SP se caractérisent par la transmission de paquets de redondance tous relatifs à un même paquet d'information. Les paquets de redondance peuvent être identiques au paquet initial (*Chase Combining*) ou bien distincts à chaque retransmission (*Incremental Redundancy*). A l'inverse, les schémas HARQ-MP

construisent des paquets de redondance relatifs à plusieurs paquets d'information. Nous avons vu au chapitre 4 qu'ils offrent ainsi un gain en débit à moyen et fort RSB puisqu'un seul paquet de redondance peut aider au décodage simultané de plusieurs paquets d'information. Mais l'approche classique simple paquet reste plus performante à faible RSB. Notons par ailleurs qu'en présence d'erreurs sur la voie de retour, le débit utile en HARQ-SP se dégrade fortement, alors que le débit utile en HARQ-MP est beaucoup moins affecté par cette perturbation. Il semble donc judicieux de chercher à combiner au mieux ces deux approches, somme toute complémentaires, afin de concevoir un protocole efficace sur toute la plage des RSB. C'est le principe du protocole hybride proposé ici.

Le protocole hybride SP et MP est conçu pour des systèmes de transmission avec voie de retour imparfaite. Il vise à limiter l'impact des pertes d'acquittements sur les performances du système, et à bénéficier des performances du SP à faible RSB. Pour atteindre cet objectif, le protocole H-SP-MP combine astucieusement les stratégies SP et MP (avec acquittement par paquet) au travers des règles de transmission suivantes. Lorsqu'un acquittement négatif est correctement reçu et porte sur un paquet précédemment transmis, l'émetteur transmet un paquet redondance relatif à ce paquet mal reçu, comme dans un protocole SP classique. Lorsque l'acquittement n'est pas correctement décodé ou perdu, l'émetteur considère que le paquet émis précédemment est correctement reçu, et passe au paquet de données suivant. Si c'est bien le cas, on évite ainsi d'envoyer un paquet de redondance inutilement. S'il s'avère, par la suite, que le paquet de données est mal reçu par la station B, la station A continue sa transmission pendant L paquets de données successifs à partir du plus ancien paquet non reçu et/ou non acquitté, puis émet une redondance multi-paquet calculée sur les L paquets de données. Cette redondance est construite de la même manière qu'en HARQ-MP.

Dans la suite, on présente le fonctionnement détaillé du protocole, et plus particulièrement les automates d'émission/réception, afin de clarifier les règles de transmission. Pour décrire le fonctionnement du protocole H-SP-MP, nous faisons référence aux variables $V(S)$ et $V(A)$ internes à l'émetteur, et à la variable $V(R)$ interne au récepteur, présentés dans le chapitre précédent. On rappelle que $N(S)$ désigne le numéro de séquence du paquet en cours, $V(S)$ indique le numéro de séquence du prochain paquet à transmettre et $V(A)$ indique le numéro de séquence du dernier paquet positivement acquitté dont l'émetteur a connaissance. Les paquets $V(A)+1$ à $N(S)$ sont sauvegardés en mémoire de l'émetteur jusqu'à leur bonne réception. En réception, $N(R)$ désigne le numéro porté par l'acquittement et $V(R)$ est la variable mémorisant le numéro du paquet attendu (tous les paquets jusqu'à $V(R) - 1$ ont été correctement reçus). Dans ce protocole, on introduit un nouveau champ dans l'entête des paquets appelé *RED* (redondance) encodé sur 2 bits qui a la signification suivante :

- $RED = 00$: paquet de données transmis à la première fois ;
- $RED = 01$: paquet de redondance relatif au paquet numéroté $N(S)$;
- $RED = 10$: paquet de redondance relatif aux paquets $N(S) - L + 1$ à $N(S)$ (multipaquet).

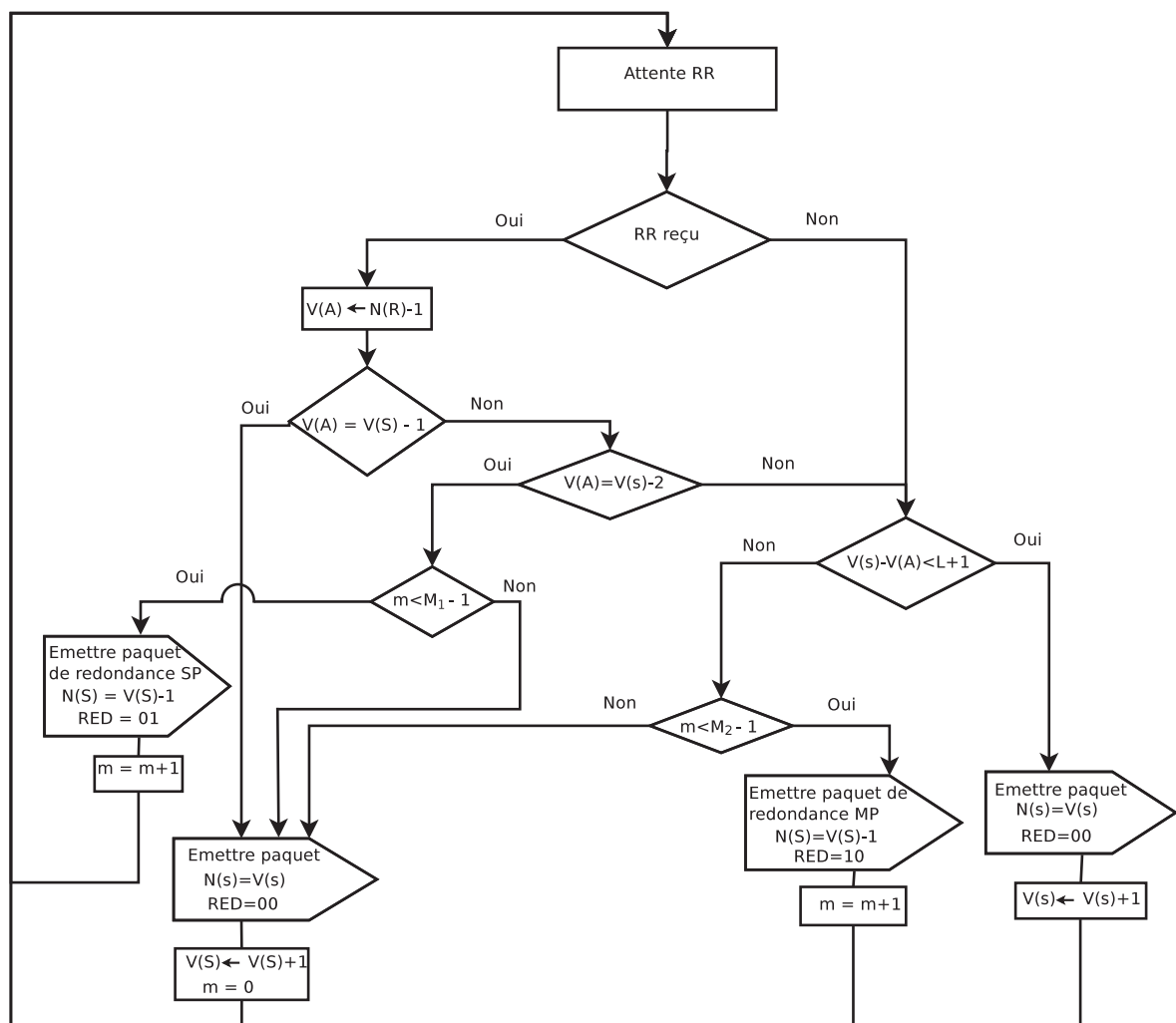


Figure 5.1 — Automate à l'émission du protocole H-SP-MP

5.2.1 Fonctionnement de l'émetteur

A réception d'un acquittement suite à l'envoi d'un paquet, trois scénarios sont possibles : bonne réception d'un acquittement positif indiquant la bonne réception du paquet ainsi que la bonne réception des paquets précédents, bonne réception d'un acquittement négatif indiquant la nécessité de retransmettre un paquet de redondance, ou enfin perte de l'acquittement. Rappelons que la station A fait la différence entre les deux premiers scénarios suivant les valeurs prises par le numéro de l'acquittement $N(R)$ et le numéro $V(S)$ du prochain paquet à transmettre. Le fonctionnement du protocole vis-à-vis de ces trois scénarios est alors le suivant :

Acquittement positif du paquet courant ($N(R) = V(S)$)

Lorsque l'acquittement reçu porte un numéro identique à $V(S)$, il est dit positif. Il permet à la station A de libérer sa mémoire et de passer à la transmission du paquet

de données suivant ($N(S) = V(S)$). Celle ci met à jour $V(A)$ et incrémente la valeur de $V(S)$.

Acquittement négatif du paquet courant ($N(R) < V(S)$)

Lorsque l'acquittement porte un numéro inférieur à $V(S)$, la station A est informée de la mauvaise réception du paquet de données portant le numéro $N(S) = N(R)$. Elle est également informée par la même occasion de la bonne réception des paquets de données numérotés de $V(A) + 1$ à $N(R) - 1$. Elle actualise $V(A)$ à la valeur $N(R) - 1$. Désormais $V(S) - 1 - V(A)$ indique le nombre de paquets dont la station A n'a pas la certitude qu'ils ont été bien reçus. Le protocole proposé adapte sa stratégie à la valeur de $V(S) - 1 - V(A)$. Si la valeur de $V(S) - 1 - V(A)$ est égale à 1 (acquittement négatif du paquet courant), la station A transmet une redondance relative au paquet erroné $N(S) = N(R)$. Si la valeur de $V(S) - 1 - V(A)$ est strictement supérieure à 1, la station A continue sa transmission jusqu'à ce que $V(S) - 1 - V(A)$ atteigne la valeur L . Elle retransmet ensuite un paquet de redondance MP portant sur l'ensemble des L paquets de données précédents. Ces L paquets contiennent nécessairement le paquet négativement acquitté. Après chaque transmission d'un paquet de redondance, la station A incrémente la valeur du compteur de nombre de rounds noté m . La figure 5.1 résume les règles de fonctionnement du protocole H-SP-MP en émission.

Acquittement perdu

Lorsque la station A détecte une perte d'acquittement, elle teste le nombre de paquets transmis et non acquittés. Si ce nombre est inférieur à L , elle transmet un nouveau paquet de données utiles et attend l'acquittement de la station B. Elle peut alors recevoir un acquittement positif indiquant la bonne réception de tous les paquets précédents. Elle peut aussi recevoir un acquittement portant un numéro $N(R)$ inférieur à $V(S)$ (cas précédents). Si aucun acquittement n'est reçu pour L paquets de données successifs, la station A enclenche la retransmission de redondance MP (cf. figure 5.1).

Au final, on voit donc que la retransmission en H-SP-MP suit un seul mode à la fois. En effet, lorsque la station A décide de retransmettre en mode SP ou MP, elle reste dans ce mode jusqu'à la bonne réception de données ou bien jusqu'à ce que le nombre maximal de retransmissions soit atteint. La séparation des modes SP/MP fait que le nombre rounds en SP (M_1) peut ne pas être identique à celui en MP (M_2). De la même manière qu'en SP ou MP, le cycle ARQ se termine lorsque la station A est notifiée de la bonne réception des paquets transmis ou que bien lorsque le nombre maximal des rounds est atteint.

5.2.2 Fonctionnement du récepteur

Lorsqu'un paquet est reçu, le récepteur vérifie si ce paquet est bien le paquet attendu. Si ceci est le cas, il envoie un acquittement positif ou négatif selon le résultat du décodage. Sinon, il sauvegarde le paquet reçu et renvoie un acquittement négatif

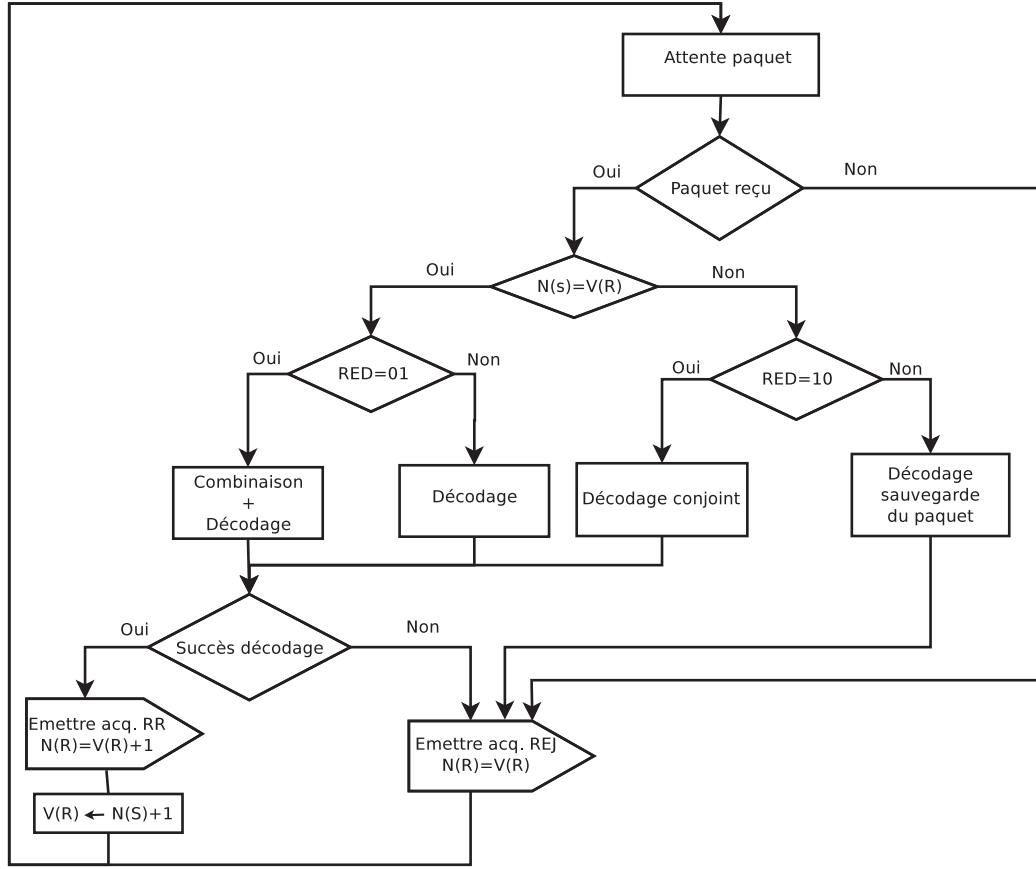


Figure 5.2 — Automate en réception du protocole H-SP-MP

pour demander la retransmission. L'automate en réception du protocole H-SP-MP est représenté à la figure 5.2.

5.3 Analyse des performances du protocole hybride SP/MP

L'analyse des performances du protocole hybride est rendue difficile par le fait qu'à l'inverse des protocoles SP ou MP considérés dans les chapitres précédents, du fait de la complexité du fonctionnement du protocole, il ne nous est pas possible ici de donner une expression analytique compacte du débit utile moyen. Pour résoudre ce problème, on propose ici une méthode analytique très générale de calcul du débit utile en HARQ. Cette méthode s'applique aussi bien aux schémas SP, MP, ainsi qu'à la combinaison des deux, et prend en compte la présence d'erreurs sur la voie de retour dans le modèle. Elle permet également de prédire les performances de schémas telle que la variante MP avec acquittement par paquet introduite dans la section 4.5.2. Elle se base sur la modélisation du protocole par une chaîne de Markov, et s'inspire en particulier des références [52] et [53], qui utilisent ce type de représentation pour calculer respectivement le débit en ARQ ainsi qu'en HARQ avec voie de retour parfaite. Les états constituant la chaîne sont définis à partir des paramètres du protocole.

Cette section est divisée en trois parties. Dans un premier temps, on présente la modélisation du protocole par une chaîne de Markov. Ensuite, on montre comment utiliser cette chaîne de Markov pour calculer le débit utile du protocole. On présente enfin les performances limites du protocole hybride H-SP-MP, calculées à l'aide de la méthode proposée, et on les compare à celles des protocoles SP et MP.

5.3.1 Modèle d'analyse

Afin de définir les états de la chaîne de Markov associée au protocole H-SP-MP, nous fixons les hypothèses suivantes. On suppose que la temporisation T_1 est fixée à sa valeur minimale (voir chapitre 2). Cette temporisation est la durée au bout de laquelle l'émetteur reçoit l'acquittement. Ensuite, on considère que les paquets de données et de redondance ont la même durée et il n'y a qu'un seul type d'acquittement. Cette modélisation permet de considérer une structure de slot. Il est alors possible de modéliser l'évolution du protocole sous la forme d'une chaîne de Markov à temps discret. Pour décrire complètement cette chaîne, il nous faut préciser la définition des états ainsi que les probabilités de transition.

Définition des états de chaîne de Markov

On définit un état de la chaîne de Markov par un quadruplet de variables liées au fonctionnement du protocole. La transition d'un état vers un autre état est caractérisée par la transmission d'un paquet par la station A et l'émission de son acquittement par la station B. Lorsque les paquets et les acquittements sont bien reçus, la station A libère de sa mémoire les paquets qui viennent d'être transmis. Dès qu'il y a une mauvaise réception de données ou d'acquittement, la station A doit mémoriser un ou plusieurs paquets. Chaque état est donc défini par la donnée des quatre variables suivantes :

- a nombre de paquets de données mal décodés par la station B.
- b nombre de paquets (données et redondance) transmis et non acquittés.
- c permet de distinguer les transmissions en mode SP ($c = 1$) de celles en mode MP ($c = 0$).
- $d = V(R) - 1 - V(A)$, est le nombre de paquets successifs bien décodés depuis l'instant initial mais dont l'acquittement n'a encore été reçu par la station A.

L'état initial de la chaîne se caractérise par une mémoire libre en émission et en réception, soit $(a = 0, b = 0, c = 0, d = 0)$. Les paramètres du protocole imposent que les variables a , b , c et d sont bornées. En effet :

- $0 \leq a \leq L$ et $a \leq b$ car au bout de L paquets de données successifs émis et non acquittés, la station A transmet une redondance et il ne peut jamais y avoir plus de paquets non reçus que de paquets envoyés.
- $0 \leq b \leq \max(M_1, L + M_2 - 1)$ car la station A libère tous les paquets de sa mémoire après M_1 paquets de redondance SP ou M_2 paquets de redondance MP.
- $0 \leq c \leq 1$ par définition de c .
- $0 \leq d \leq L - 1$ car on peut avoir au maximum L paquets de données bien reçus et non acquittés avant de transmettre de la redondance. Si le nombre de paquets bien reçus et non acquittés est égal à L , on force la valeur de d à zéro. En effet, le

prochain paquet à transmettre est un paquet de redondance MP. Par conséquent, la valeur maximale considérée de la variable d est $L - 1$.

La station A décide de transmettre de la redondance MP après avoir envoyé L paquets de données sans recevoir de réponse de la station B. Elle est incapable de savoir quels sont les paquets correctement reçus et dans quel ordre ils ont été reçus. Elle envoie donc une redondance portant sur l'ensemble des L paquets de données. Dans notre modèle d'état, lorsque $b = L$ et $0 \leq a \leq L$, on force alors la valeur de d à zéro ($d = 0$).

La forme générale de la chaîne de Markov à temps discret et à états finis modélisant le protocole H-SP-MP est présentée sur les figures 5.3 et 5.4. Sa spécialisation au cas particulier où $M_1 = M_2 = 2$ et $L = 3$ est illustrée sur la figure 5.5.

L'examen de la figure 5.5 montre que les états peuvent être partitionnés en trois classes. Une première classe regroupe les états relatifs au mode "SP". Ce mode correspond aux états caractérisés par des paramètres de la forme $(a, b, 1, 1)$ avec $0 \leq a \leq 1$ et $1 \leq b \leq M_1$. La deuxième classe contient les états relatifs au mode "MP". Ces états sont caractérisés par $(a = 0, b \geq L)$ ou $(1 \leq a \leq L, 1 \leq b \leq L + M_2 - 1)$ et $(c = 0, d = 0)$. Enfin, la troisième classe d'états définit un mode dit "indéterminé", caractérisé par la bonne réception d'au moins un paquet de données et un acquittement non reçu, à partir duquel on peut tout aussi bien revenir à l'état initial directement, que basculer en mode SP ou bien en mode MP. Les états de ce mode sont caractérisés par $a \geq 0$, $1 \leq b \leq L - 1$, $c = 0$ et $1 \leq d \leq L - 1$.

Définition des probabilités de transition

La station A transmet un paquet et attend son acquittement. La station B répond par un acquittement positif ou négatif selon le résultat du décodage. Cet acquittement peut être bien reçu par la station A ou bien perdu. On a donc quatre événements possibles : bonne réception du paquet (succès de décodage) et bonne réception de son ACK, bonne réception du paquet et mauvaise réception (perte) de l'ACK, mauvaise réception du paquet (échec de décodage) et bonne réception du NAK, ou mauvaise réception du paquet et mauvaise réception du NAK. On utilise une notation avec 2 indices hauts valant chacun soit b (pour bon), soit m (pour mauvais). Les quatre événements précédents sont donc référencés respectivement par bb, bm, mb et mm. La plupart des transitions sur la figure 5.3 font intervenir des probabilités notées soit $P^{xx}(m)$, soit $Q_l^{xx}(m)$, où m désigne le numéro du round en cours, et $xx = \{bb, bm, mb, mm\}$ désigne l'événement associé à la transition. Les probabilités $P^{xx}(m)$ sont utilisées lorsque l'on réalise un décodage relatif à un seul paquet de données. A titre d'exemple, $P^{bb}(m)$ désigne ainsi la probabilité de l'événement "succès de décodage d'un paquet de données et succès de transmission de l'acquittement au round m sachant qu'il y a eu échecs de décodage aux rounds 0 à $m - 1$ ". Les probabilités $Q_l^{xx}(m)$ sont utilisées lorsque l'on réalise un décodage conjoint relatif à plusieurs paquets de données. Plus précisément, $Q_l^{bb}(m)$ représente ainsi la probabilité de l'événement "succès de décodage conjoint des L paquets de données et bonne réception de l'acquittement au round m sachant qu'il y a eu des échecs de décodage aux rounds 0 à $m - 1$ et qu'il y a eu l paquets de données bien décodés au premier round ($m = 0$)".

Les probabilités $P^{xx}(m)$ sont calculées comme suit. On rappelle que $p(m) = \Pr\{e_o, e_1, \dots, e_m\}$ désigne la probabilité conjointe d'un échec de décodage aux rounds 0 à m . Les probabilités marginales $\Pr\{e_i\}$ sont fonction du schéma de modulation et codage et du modèle de canal considéré. En posant $p(-1) = 1$, on montre que pour tout $m \geq 0$ on a (voir annexe A) :

$$P^{bb}(m) = (1 - p(m)/p(m-1))(1 - \epsilon). \quad (5.1)$$

On a de même :

$$P^{bm}(m) = (1 - p(m)/p(m-1))\epsilon \quad (5.2)$$

$$P^{mb}(m) = (p(m)/p(m-1))(1 - \epsilon) \quad (5.3)$$

$$P^{mm}(m) = (p(m)/p(m-1))\epsilon \quad (5.4)$$

De la même façon, les probabilités $Q_l^{xx}(m)$ sont calculées à partir de la probabilité $p_l(m) = \Pr(\mathcal{A}_l, e_1, \dots, e_m)$ d'observer $m+1$ échecs de décodage consécutifs aux rounds 0 à m , avec l paquets parmi L bien décodés à l'issue du premier round en MP. Plus précisément, on montre dans l'annexe A que pour tout $m \geq 1$:

$$Q^{bb}(m) = (1 - p_l(m)/p_l(m-1))(1 - \epsilon) \quad (5.5)$$

De façon similaire, on a :

$$Q^{bm}(m) = (1 - p_l(m)/p_l(m-1))\epsilon \quad (5.6)$$

$$Q^{mb}(m) = p_l(m)/p_l(m-1)(1 - \epsilon) \quad (5.7)$$

$$Q^{mm}(m) = p_l(m)/p_l(m-1)\epsilon \quad (5.8)$$

Les différentes probabilités de transition entre les états de la même classe ou entre états de classes différentes sont présentées en détail en annexe A

Equation d'équilibre de la chaîne

Dans le cas général d'une chaîne de Markov à n états et de matrice de transition $\mathbf{T} = (t_{ij})$, il existe une loi stationnaire []. Lorsque chaque état de la chaîne est accessible à partir de n'importe quel autre état, la loi stationnaire est unique [54]. On note par $\pi = (\pi_1, \pi_2, \dots, \pi_n)$, le vecteur ligne des probabilités d'équilibre de la chaîne. Ce vecteur est la solution du système linéaire suivant :

$$\begin{cases} \pi = \pi \mathbf{T} \\ \sum_{i=1}^n \pi_i = 1. \end{cases} \quad (5.9)$$

5.3.2 Application au calcul du débit utile

Comme dans les chapitres précédents, le débit utile moyen est défini comme le rapport du nombre moyen de bits d'information correctement reçus sur le nombre moyen de symboles transmis. Il peut alors exprimer en fonction du rendement de transmission

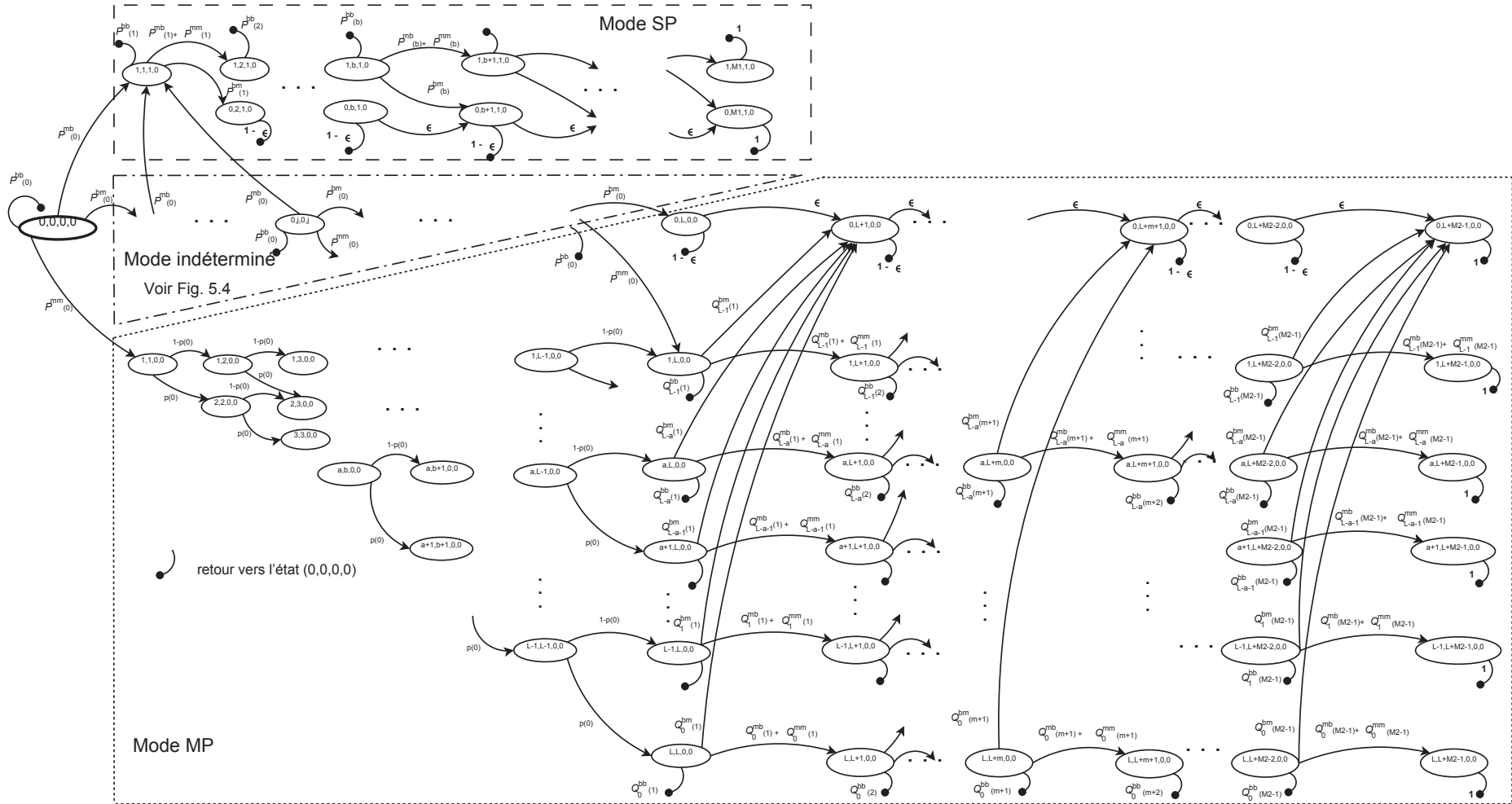


Figure 5.3 — Structure générale de la chaîne de Markov modélisant le protocole hybride H-SP-MP avec un nombre maximum de retransmissions fixé à M_1 en mode SP, à M_2 en mode MP, et un mode MP opérant sur des groupes de L paquets de données.

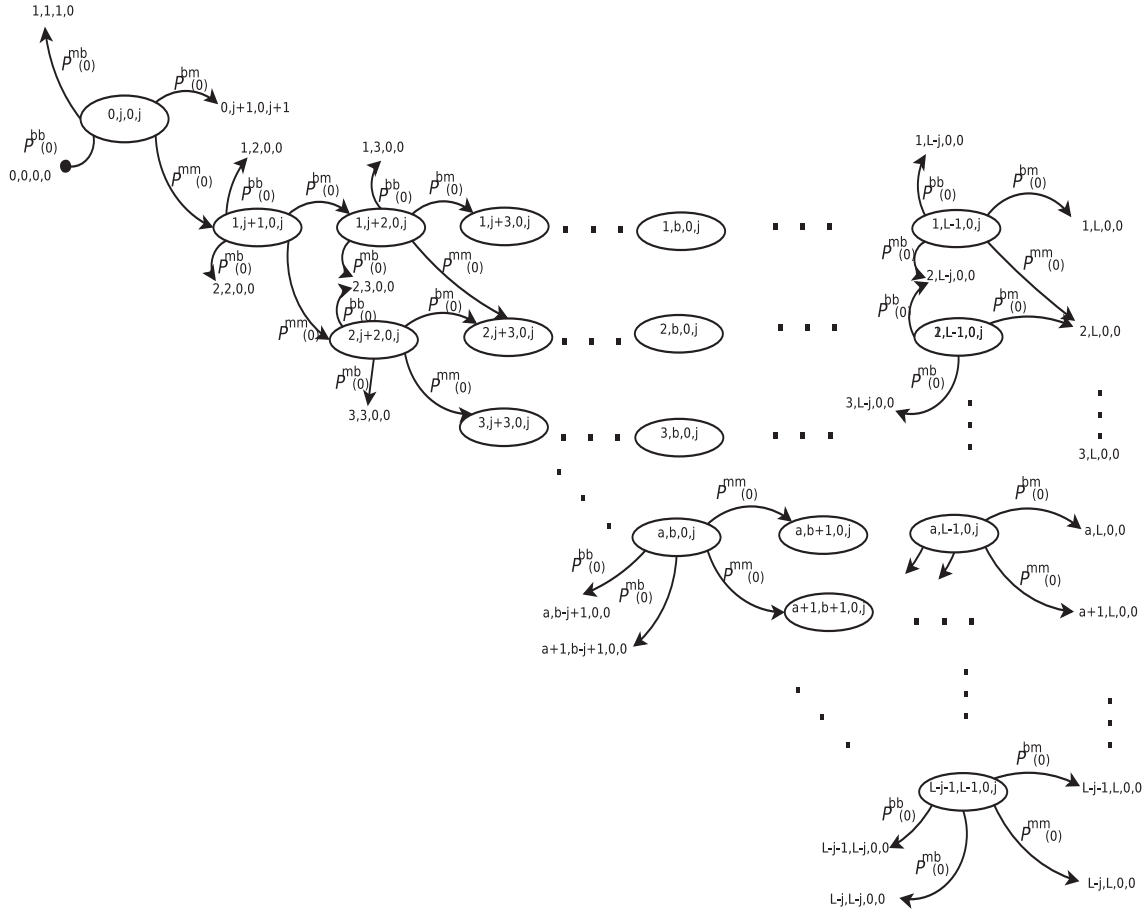


Figure 5.4 — Figure complémentaire de la structure générale de chaîne de Markov de la figure 5.3 représentant les états du mode indéterminé du protocole hybride SP et MP avec paramètres M_1 , M_2 , L . L'état $(0, j, 0, j)$ vérifie $0 < j < L$.

R_o , du nombre moyen de paquets correctement reçus $E[\mathcal{B}_p]$ et du nombre moyen de paquets transmis $E[\mathcal{T}_p]$ sur une durée d'un cycle de temps :

$$\eta = R_o \frac{E[\mathcal{B}_p]}{E[\mathcal{T}_p]} \text{ bits/symbole} \quad (5.10)$$

La modélisation du comportement du protocole par une chaîne de Markov permet de calculer les quantités $E[\mathcal{B}_p]$ et $E[\mathcal{T}_p]$, lorsque la chaîne est dans son état d'équilibre. Il suffit pour cela de garder trace du nombre moyen d_{ij} de paquets correctement décodés lors d'une transition entre deux états $s_i = (a, b, c, d)$ et $s_j = (a', b', c', d')$. On utilise pour cela une matrice $\mathbf{D} = (d_{ij})_{1 \leq i, j \leq n}$ de mêmes dimensions que la matrice de transition $\mathbf{T}$. La figure 5.6 illustre une section de la chaîne montrant des valeurs du nombre de paquets correctement décodés lors des transitions en états de la chaîne. A chaque transition, le nombre de paquets correctement décodés est compris entre 0 et L . d_{ij} prend sa valeur maximale L lorsque le récepteur décode correctement L paquets en mode MP après transmission de redondance. Le tableau 5.7 montre le nombre de paquets correctement décodés associé aux transitions d'un état s_i vers l'état initial. Ce tableau regroupe toutes les transitions des états du mode SP, MP et indéterminé

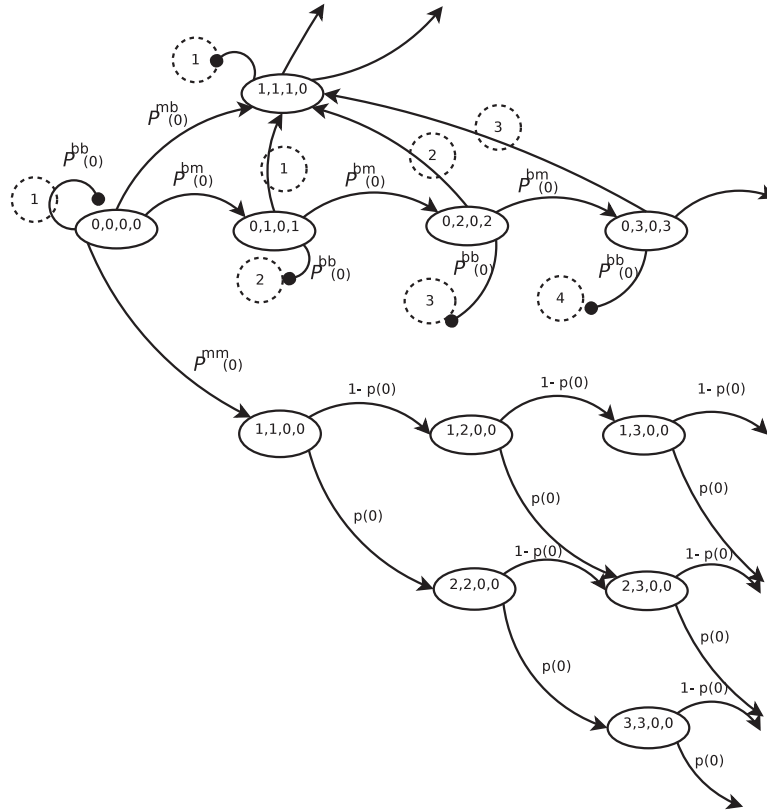


Figure 5.6 — Représentation d'une section de la chaîne de Markov en H-SP-MP. Les valeurs entourées des cercles pointillés représentent le nombre de paquets bien décodés relatifs à une transition particulière.

	a	b	c	d	Nombre de paquets décodés d_{i1}
s_i	0	$1 \leq b \leq M_1$	1	0	1
	1	$1 \leq b \leq M_1 - 1$	1	0	1
	1	$b = M_1$	1	0	$1 - p(M_1)/p(M_1 - 1)$
	0	$0 \leq b \leq L - 1$	0	b	$b + 1$
	$0 \leq a \leq L$	$L \leq b \leq L + M_2 - 2$	0	0	L
	$0 \leq a \leq L$	$b = L + M_2 - 1$	0	0	$L - ap_l(m)$

Figure 5.7 — Nombre de paquets de données d_{i1} bien décodés associé aux transitions allant d'un état $s_i = (a, b, c, d)$ vers l'état initial $s_1 = (0, 0, 0, 0)$

s_i	a	b	c	d	d_{ij}
	$0 \leq a \leq L - 1$	$1 \leq b \leq L - 2$	0	$1 \leq d \leq L - 2$	d
s_j	$a' \neq 0$	$b' \neq 0$	$0 \leq c' \leq 1$	$d' = 0$	

Figure 5.8 — Nombre de paquets bien décodés associé à la transition entre un état s_i et un état s_j tous deux différents de l'état initial

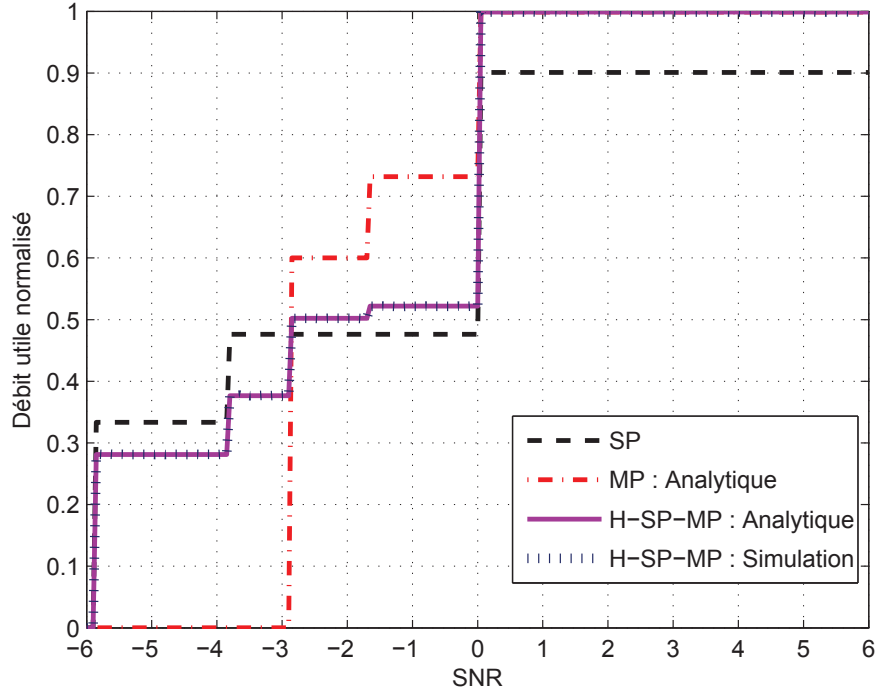


Figure 5.9 — Débit utile en SP, MP et H-SP-MP pour une transmission sur canal BABG complexe avec codes gaussiens de longueur infinie et rendement $R_o = 1$ bit/symbole, avec $M_1 = 2$, $M_2 = 2$, $L = 3$ et $\epsilon = 0,1$

En utilisant (5.10), (5.12) et (5.13), le débit utile moyen en H-SP-MP est finalement donné par :

$$\eta = R_o \pi(d_1, d_2, \dots, d_n)^T \quad (5.14)$$

où $(.)^T$ désigne la transposée. Le débit utile dépend du rendement de la transmission, du nombre moyen de paquets correctement décodés en chacun des états de la chaîne de Markov, et du vecteur des probabilités à l'équilibre.

5.3.3 Performances du protocole hybride H-SP-MP

Le modèle d'analyse précédent nous permet d'évaluer les performances théoriques du protocole H-SP-MP. Dans cette section, on s'intéresse plus particulièrement aux performances limites obtenues en considérant des codes aléatoires gaussiens de très grande taille. Les résultats sont présentés tout d'abord sur le canal BABG, puis sur le canal de Rayleigh à évanouissements par blocs, et comparés à chaque fois aux performances des protocoles HARQ-II-SP et HARQ-II-MP.

Transmission sur canal Gaussien

Pour utiliser la méthode d'analyse précédente, il suffit d'explicitier le calcul des probabilités de transition pour le modèle de transmission considéré (code + canal). La

transmission sur canal gaussien est caractérisée par un rapport signal à bruit constant. Dans le cas de codes gaussiens aléatoires asymptotiquement longs, nous avons vu dans les chapitres précédents qu'un échec de décodage au round m implique nécessairement un échec de décodage aux rounds 0 à $m - 1$. Dans ce cas, les probabilités de transition $P^{xx}(m)$ s'écrivent en fonction de $p(m)$ et ϵ comme suit :

$$P^{bb}(m) = (1 - p(m))(1 - \epsilon) \quad (5.15)$$

$$P^{bm}(m) = (1 - p(m))\epsilon \quad (5.16)$$

$$P^{mb}(m) = p(m)(1 - \epsilon) \quad (5.17)$$

$$P^{mm}(m) = p(m)\epsilon \quad (5.18)$$

$$(5.19)$$

Les probabilités $p(m)$ sont données par (3.42) en SP.

Pour ce canal les L paquets transmis au premier round en MP sont soit tous décodés soit tous erronés. D'où $Q_l^{xx}(m) = 0$ pour $l \neq 0$. D'autre part, lorsque il y a un échec de décodage conjoint au round $m \geq 1$, cela implique nécessairement un échec de décodage conjoint aux rounds 0 à $m - 1$. Les probabilités $Q_l^{xx}(m)$ sont exprimées ainsi :

$$Q_l^{bb}(m) = (1 - p(m))(1 - \epsilon) \quad (5.20)$$

$$Q_l^{bm}(m) = (1 - p(m))\epsilon \quad (5.21)$$

$$Q_l^{mb}(m) = p(m)(1 - \epsilon) \quad (5.22)$$

$$Q_l^{mm}(m) = p(m)\epsilon \quad (5.23)$$

où $p(m)$ est donnée par (4.10).

Le calcul des probabilités de transition permet de définir la matrice $\mathbf{T}$ des transitions de la chaîne de Markov. On résout ensuite le système linéaire (5.9) pour déterminer le vecteur de probabilités de transition à l'équilibre π . La matrice $\mathbf{D}$ est produite grâce aux tableaux 5.7 et 5.8, et le débit utile η est finalement calculé en utilisant (5.14).

Dans la suite, on considère un protocole H-SP-MP avec les paramètres $L = 3$, $M_1 = M_2 = 2$ dont la chaîne de Markov associée est illustrée sur la figure 5.5.

Les figures 5.9, 5.10 et 5.11 présentent les courbes du débit utile analytique et simulé en H-SP-MP pour une transmission avec codes gaussiens aléatoires arbitrairement longs, de rendement $R_o = 1$ bits/symboles, sur un canal BABG complexe de RSB $\gamma = E_s/N_o$, en considérant un encodage conjoint sur $L = 3$ paquets en MP, pour un nombre maximal de rounds fixé à $M_1 = M_2 = 2$ et pour différentes valeurs de la probabilité de perte sur la voie de retour ($\epsilon = 0, 1$, $\epsilon = 0, 3$ et $\epsilon = 0, 5$). On remarque dans premier temps que les courbes du débit analytique et simulé sont confondues. Ceci valide la méthode analytique de calcul des performances. Ensuite, on constate que le schéma H-SP-MP apporte une nette amélioration du débit utile moyen à faible RSB. Il conserve également la robustesse des schémas MP aux erreurs sur la voie de retour. A fort RSB, le H-SP-MP atteint le débit maximal du MP. A faible RSB, plus la probabilité de perte sur la voie de retour augmente, plus le protocole H-SP-MP se comporte comme un protocole MP et donc plus le débit utile moyen se dégrade, comparé à un protocole SP classique. Les protocoles SP et MP peuvent être vus comme des cas particuliers du protocole H-SP-MP. En effet, le protocole H-SP-MP se réduit au protocole

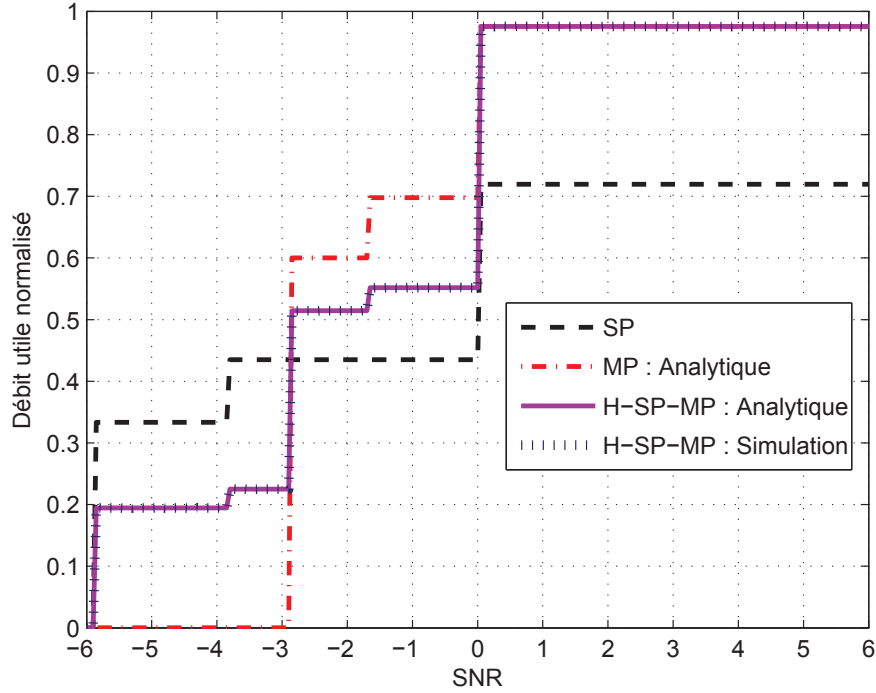


Figure 5.10 — Débit utile en SP, MP et H-SP-MP pour une transmission sur canal BABG complexe avec codes gaussiens de longueur infinie et rendement $R_o = 1$ bit/symbole, $M_1 = 2$, $M_2 = 2$, $L = 3$ et $\epsilon = 0, 3$

SP pour $\epsilon = 0$. À l'inverse, on retrouve le protocole MP lorsque $\epsilon = 1$. On peut donc modéliser les protocoles SP et MP par des chaînes de Markov et les analyser à l'aide de la méthode précédente. Nous avons ainsi validé que pour une valeur quelconque de ϵ , les performances obtenues étaient rigoureusement conformes aux résultats théoriques et simulés établis dans les chapitres précédents. Ces différents résultats ont fait l'objet des publications [2] et [4].

Transmission sur canal de Rayleigh à évanouissements par blocs

Dans le cas d'une transmission sur canal de Rayleigh à évanouissements par blocs avec des codes gaussiens aléatoires asymptotiquement longs, la probabilité d'échec de décodage conjoint $p(m)$ est donnée par (3.49) en SP et par (4.15) en MP. La probabilité $p_l(m)$ est donnée par (4.16). Les probabilités de transition entre les états de la chaîne s'en déduisent à l'aide des expressions (5.1), (5.2), (5.3), (5.4), (5.5), (5.6), (5.7) et (5.8) en fonction de $p(m)$ et $p_l(m)$ et ϵ .

Les figures 5.12, 5.13 et 5.14 comparent le débit analytique et simulé en SP, MP et H-SP-MP pour ce scénario. De nouveau, le protocole H-SP-MP montre une amélioration du débit à faible RSB par rapport au schéma MP. À fort RSB, il atteint le débit asymptotique du MP. D'autre part, en augmentant la probabilité de perte sur la voie de retour, le protocole H-SP-MP reste plus robuste que le SP. En effet, à fort RSB, il est plus probable qu'un paquet de données soit correctement décodé à sa première

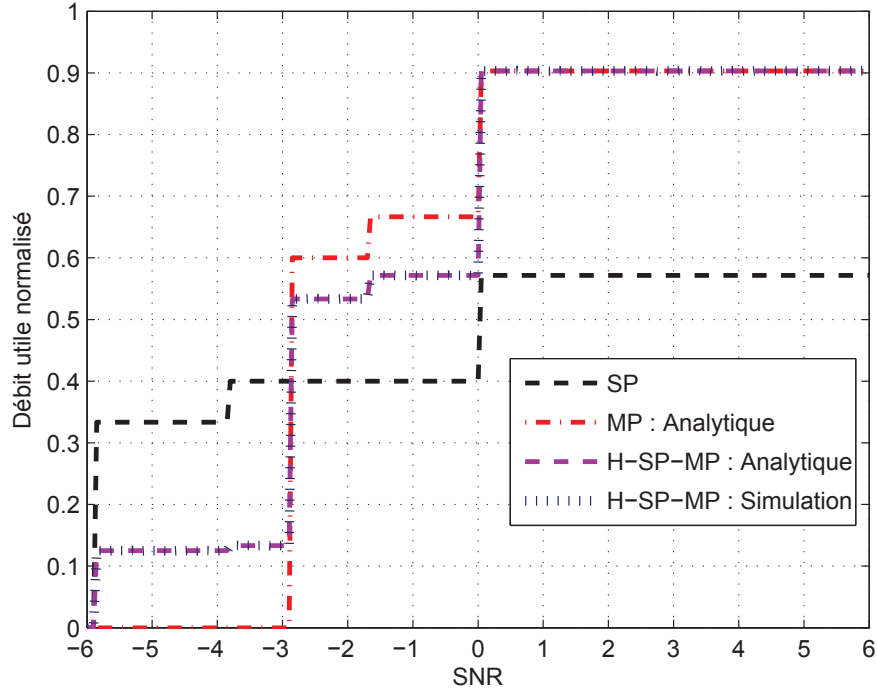


Figure 5.11 — Débit utile en SP, MP et H-SP-MP pour une transmission sur canal BABG complexe avec codes gaussiens de longueur infinie et rendement $R_o = 1$ bit/symbole, $M_1 = 2$, $M_2 = 2$, $L = 3$ et $\epsilon = 0,5$

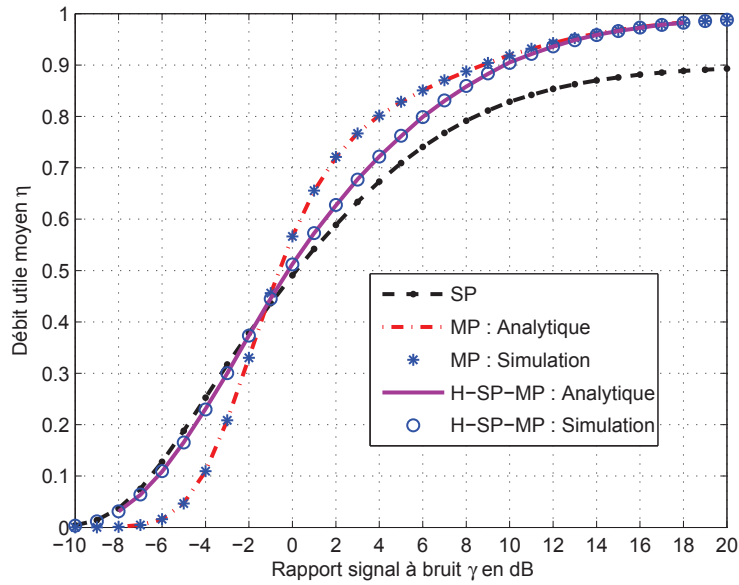


Figure 5.12 — Débit utile en SP, MP et H-SP-MP pour une transmission sur canal de Rayleigh à évanouissements par blocs avec codes gaussiens de longueur infinie et rendement $R_o = 1$ bit/symbole, $M_1 = 2$, $M_2 = 2$, $L = 3$ et $\epsilon = 0,1$

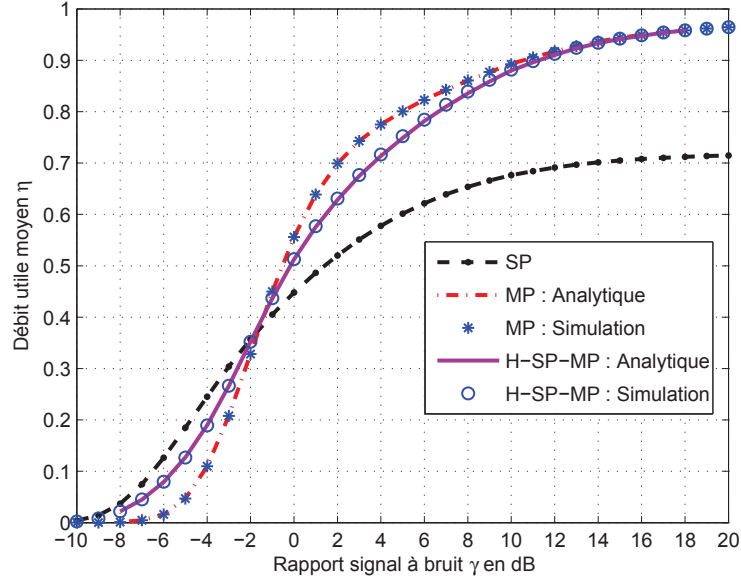


Figure 5.13 — Débit utile en SP, MP et H-SP-MP pour une transmission sur canal de Rayleigh à évanouissements par blocs avec codes gaussiens de longueur infinie et rendement $R_o = 1$ bit/symbole, $M_1 = 2$, $M_2 = 2$, $L = 3$ et $\epsilon = 0,3$

transmission. Lorsque l’acquittement positif est mal reçu par l’émetteur, le schéma SP gaspille des ressources à retransmettre le paquet déjà reçu tandis que le schéma H-SP-MP suppose que ce paquet a été bien décodé et passe à la transmission du paquet suivant. Il recevra par la suite un acquittement positif indiquant la bonne réception des paquets précédemment transmis. La perte de l’acquittement est alors transparente. D’où le gain en débit des schémas MP et H-SP-MP par rapport au schéma SP. Le protocole H-SP-MP constitue donc un bon compromis entre les stratégies SP et MP vu qu’il bénéficie de la robustesse de la stratégie SP à faible RSB, et du gain asymptotique (fort RSB) en débit offert par le MP.

5.4 Exploitation de la connaissance du canal en réception

Le schéma H-SP-MP proposé en section 5.2 peut être qualifié d’aveugle dans le sens où il n’exploite aucune information sur l’état du canal. Or cette information est présente dans la plupart des systèmes de transmission radio modernes, en particulier les systèmes cellulaires. Lorsque le canal est estimé et connu du réception, il est possible d’exploiter cette information pour adapter le type de redondance à transmettre aux conditions du canal afin de réduire la perte en débit en H-SP-MP par rapport au SP à faible RSB, voire par rapport au MP à moyen RSB. En effet, dans certaines plages de RSB, il est préférable de retransmettre en mode SP. Dans d’autres, la transmission de redondance MP est au contraire plus profitable.

Ce constat nous a amené à proposer une variante du protocole H-SP-MP qui exploite

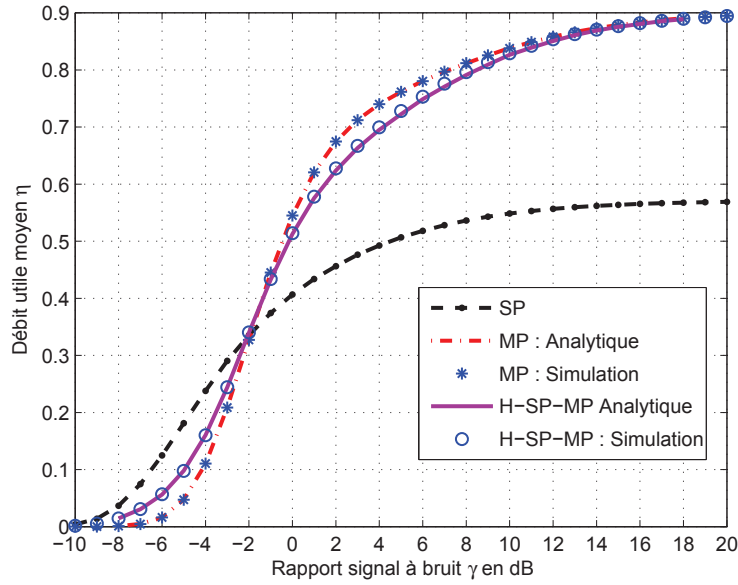


Figure 5.14 — Débit utile en SP, MP et H-SP-MP pour une transmission sur canal de Rayleigh à évanouissements par blocs avec codes gaussiens de longueur infinie et rendement $R_o = 1$ bit/symbole, $M_1 = 2$, $M_2 = 2$, $L = 3$ et $\epsilon = 0,5$

la connaissance de l'état du canal en réception pour améliorer le débit utile de la transmission. Plus précisément, l'idée consiste à favoriser la retransmission en mode MP dans certaines zones du RSB. Il suffit pour cela de forcer la station B à s'abstenir de renvoyer les acquittements négatifs. Ceci conduit nécessairement à transmettre en mode MP par définition du protocole hybride SP et MP. La station A considère que l'acquittement a été perdu et passe automatiquement à la transmission du paquet de données suivant. Cela ne nécessite pas de renvoyer l'état du canal à la station A. On ne rajoute donc pas de signalisation supplémentaire sur la voie de retour. La connaissance de l'état du canal est simplement utilisée pour identifier la zone du RSB dans laquelle on se situe. La station B peut connaître soit le RSB moyen de la transmission, soit le RSB instantané de la dernière transmission. Ces deux cas sont adressés séparément et comparés dans les sous-sections suivantes.

Notons que lorsque la station B renvoie l'état du canal à la station A, cette information peut être utilisée pour sélectionner le schéma de modulation et codage le mieux adapté aux conditions de propagation (optimisation du rendement R_o). Ce point n'a pas du tout été considéré dans cette thèse.

5.4.1 Présence d'une connaissance sur le RSB moyen

Dans certaines conditions (situation de faible mobilité), le canal de transmission varie relativement lentement par rapport à la durée d'un paquet. L'estimation du RSB moyen constitue alors une bonne approximation du RSB sur le canal. Lorsque cette information est disponible en réception, elle peut être utilisée pour adapter le renvoi des acquittements négatifs. Plus précisément, le récepteur renvoie les acquittements

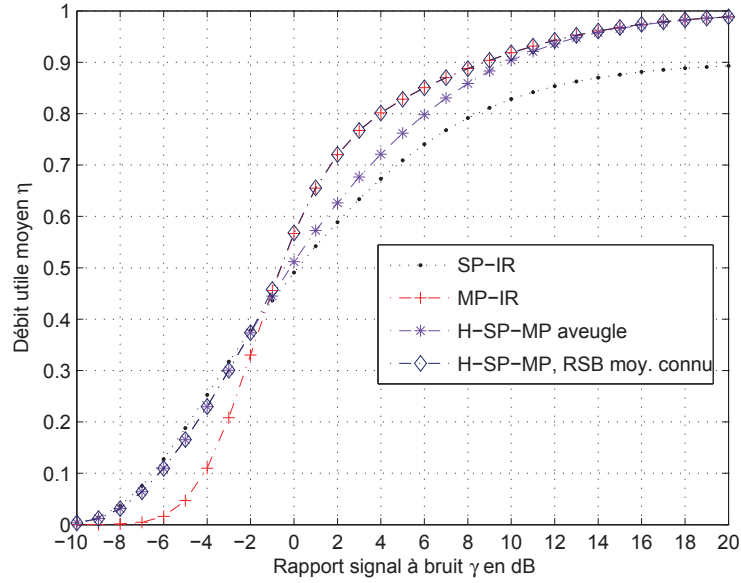


Figure 5.15 — Débit utile en SP, MP et H-SP-MP (aveugle/avec connaissance du RSB moyen) pour une transmission sur canal de Rayleigh à évanouissements par blocs avec codes gaussiens de longueur infinie et rendement $R_o = 1$ bit/symbole, $M_1 = 2$, $M_2 = 2$, $L = 3$, $\gamma_{th} = -1$ dB et $\epsilon = 0,1$

négatifs uniquement lorsque le RSB estimé est inférieur à un seuil γ_{th} . Lorsque le RSB est supérieur au seuil γ_{th} , le récepteur s'abstient de tout acquittement négatif, ce qui enclenche la retransmission en mode MP (redondance MP). Le choix du seuil γ_{th} est basé sur l'analyse des courbes du débit en SP et MP. Ce seuil correspond au RSB à partir duquel le débit utile en MP devient supérieur au débit SP.

5.4.2 Présence d'une connaissance sur le RSB instantané

Lorsque le canal de transmission varie rapidement, le RSB moyen ne constitue plus une bonne estimation du canal. Dans ce cas, le protocole H-SP-MP fonctionne de manière similaire à celui de la section précédente, à la différence qu'il se base sur la connaissance instantanée du canal. La station B ne renvoie pas les acquittements négatifs lorsque le RSB instantané est supérieur au seuil γ_{th} .

5.4.3 Performances de ces deux approches

Le débit utile moyen pour les schémas H-SP-MP avec connaissance sur l'état du canal est calculé d'une façon similaire à celui en H-SP-MP présenté au début de ce chapitre (sans connaissance du canal ou *aveugle*). Le calcul analytique est fait en utilisant un modèle de Markov dans lequel les probabilités de transition P^{mb} et P^{mm} doivent prendre en compte la décision sur le renvoi des acquittements négatif lorsque le RSB est supérieur au seuil γ_{th} (voir annexe A).

Pour les RSB inférieurs à γ_{th} , la variante H-SP-MP avec connaissance du RSB

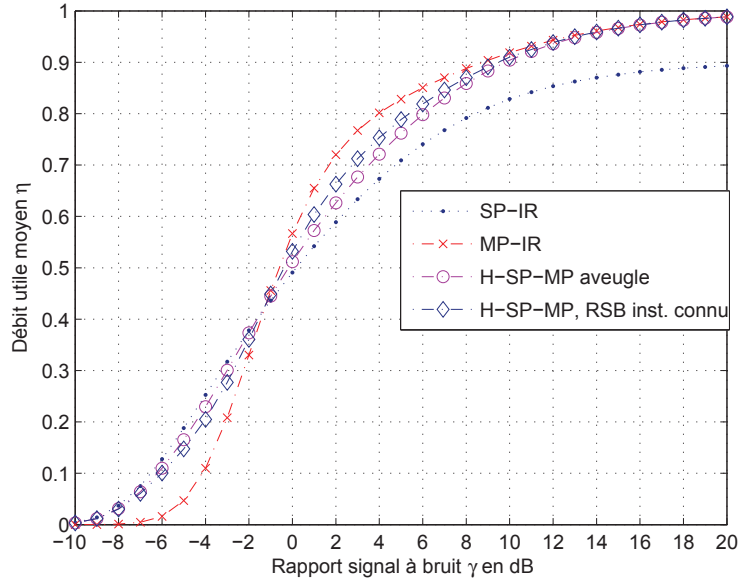


Figure 5.16 — Débit utile en SP, MP et H-SP-MP (aveugle/avec connaissance du RSB instantané) pour une transmission sur canal de Rayleigh à évanouissements par blocs avec codes gaussiens de longueur infinie et rendement $R_o = 1$ bit/symbole, $M_1 = 2$, $M_2 = 2$, $L = 3$, $\gamma_{th} = -1$ dB et $\epsilon = 0, 1$

moyen fonctionne de façon identique à l'H-SP-MP aveugle pour les RSB inférieurs à γ_{th} . Pour des RSB supérieurs à γ_{th} , l'H-SP-MP avec RSB moyen connu privilège la retransmission de redondance en mode MP. La figure 5.15 présente le débit utile moyen pour ces différents schémas dans le cas où $M_1 = 2$, $M_2 = 2$, $\gamma_{th} = -1$ dB et $\epsilon = 0, 1$. On remarque bien l'amélioration du débit apportée par le non renvoi des acquittements négatifs à moyen RSB. En particulier, au delà du seuil γ_{th} , le débit utile en H-SP-MP avec RSB moyen connu est identique au débit utile moyen en MP.

La figure 5.16 compare le débit utile offert par les protocoles SP, MP, H-SP-MP aveugle, et H-SP-MP avec connaissance du RSB instantané pour $\gamma_{th} = -1$ dB et pour une probabilité d'erreur sur la voie de retour $\epsilon = 0, 1$. Le débit utile en H-SP-MP avec connaissance sur le RSB instantané est inférieur au débit en H-SP-MP aveugle pour les RSBs inférieurs à γ_{th} . En effet, pour γ_{th} fixe, et pour des RSBs moyens sur le canal inférieurs à γ_{th} , la valeur instantanée du RSB ne constitue pas une bonne estimation du RSB moyen sur le canal. Il est possible que la réalisation du RSB soit supérieure à γ_{th} . Dans ce cas, la décision du récepteur de s'abstenir de renvoyer les acquittements négatifs conduisant à des retransmissions en mode MP n'est pas la bonne car le SP est meilleur que le MP pour les RSB moyens inférieurs à γ_{th} . Pour cette raison, le débit utile en H-SP-MP avec connaissance du RSB instantané est inférieur au débit en H-SP-MP aveugle pour les RSBs moyens inférieurs à γ_{th} car l'occurrence du basculement en mode MP devient plus importante en H-SP-MP avec connaissance du RSB instantané. En revanche, on remarque bien l'amélioration du débit pour les RSBs supérieurs à γ_{th} . Ces résultats ont fait l'objet de la publication [3].

5.5 Efficacité énergétique

L'analyse théorique du débit utile montre que le schéma hybride SP et MP est plus robuste aux erreurs sur la voie de retour. Il est alors possible de limiter la consommation en énergie allouée à la transmission des acquittements sans trop dégrader les performances du système. Dans ce cadre, il paraît intéressant de chercher les paramètres optimaux qui maximisent le débit utile pour une énergie globale constante. Cette énergie globale est répartie entre la transmission des symboles du paquet sur la voie directe et la transmission des symboles d'acquiescement sur la voie de retour. Elle est notée E_{round} en référence à l'analyse similaire menée dans les chapitres 2 et 3. Dans le contexte du protocole hybride, sa signification est toutefois différente car elle ne correspond pas nécessairement à l'énergie consommée au cours d'un round, mais plus simplement à l'énergie consommée par l'envoi d'un paquet et son acquiescement.

Tout au long de ce chapitre, nous avons considéré des codes gaussiens asymptotiquement longs. Dans cette section, nous pouvons garder cette hypothèse car la répartition optimale de l'énergie ne va pas dépendre directement de la taille du code mais du rapport entre la taille du paquet codé et la taille du paquet d'acquiescement. Comme dans le chapitre 3, on note par α la proportion de l'énergie allouée à la transmission des acquittements. Les énergies transmises par symbole de données E_s et par symbole d'acquiescement E_a dépendent de l'énergie globale par symbole E_g définie dans les chapitres 2 (section 2.5.2) et 3 (section 3.6.3). Ces énergies sont respectivement données par les équations (3.77) et (3.78) en fonction de α , ρ et E_g .

Les probabilités d'échec de transmissions sont calculées en fonction de E_g , α et ρ . Le débit utile de la transmission est calculé en utilisant le modèle de Markov décrit en section 5.3.2. Il dépend de E_g , α et ρ .

Dans cette section, on considère une transmission sur un canal de Rayleigh dans les deux sens de transmission. Sur la voie de retour, pour simplifier l'analyse, les acquittements sont supposés transmis à l'aide d'une modulation binaire et sans codage de canal. La probabilité d'erreur binaire sur la voie de retour s'exprime alors en fonction de E_a comme suit :

$$P_b = 1/2 \left(1 - \sqrt{\frac{E_a/N_o}{1 + E_a/N_o}} \right) \quad (5.24)$$

La probabilité de perte d'acquiescement vaut alors :

$$\epsilon = 1 - (1 - P_b)^H \quad (5.25)$$

La figure 5.17 présente l'évolution du débit utile en HARQ-SP et HARQ-hybride SP et MP avec codes gaussiens aléatoires asymptotiquement longs en fonction de α pour une transmission sur canal de Rayleigh à évanouissements par blocs avec une énergie globale par symbole $E_g = 5$ dB et un rapport $1 - \rho \approx 10^{-3}$ entre la taille d'un acquiescement et la taille d'un paquet codé. On remarque que pour ce choix de paramètre, le schéma hybride SP-MP offre un débit supérieur au SP pour toutes les valeurs de α . Pour les deux protocoles considérés, il existe une valeur optimale de α qui maximise le débit utile moyen. Cet optimum est plus petit dans le cas du schéma hybride. Ceci

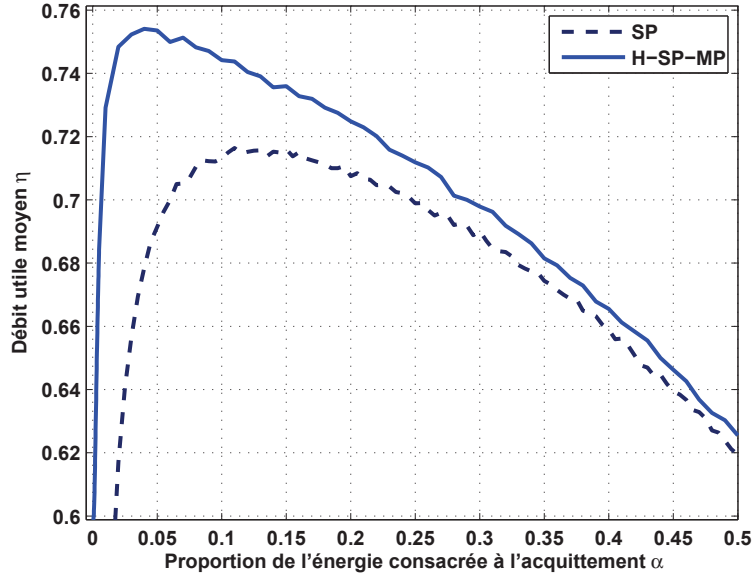


Figure 5.17 — Débit utile en HARQ-SP et H-SP-MP avec des codes gaussiens en fonction du coefficient de répartition α , pour une transmission sur un canal de Rayleigh à évanouissements par blocs sans codage d’acquitements et avec un RSB global moyen constant ($E_g/N_o = 5$ dB) et un rapport $1 - \rho = 0,001$ entre la taille d’un acquitement et la taille d’un paquet codé

montre que le schéma hybride SP et MP demande moins d’énergie transmise par acquitement. Cette conclusion était prévisible car nous avons vu que ce schéma était plus robuste que le SP aux erreurs sur la voie de retour.

5.6 Conclusion

Dans ce chapitre, nous avons proposé et analysé un protocole hybride combinant la stratégie SP classique avec la stratégie MP développée dans le chapitre précédent. Nous avons détaillé le fonctionnement protocolaire de cette stratégie hybride au niveau de l’émetteur et au niveau du récepteur. Les performances de cette combinaison ont été présentées et comparées avec celles des schémas SP et MP. Il a été démontré que la combinaison hybride SP/MP constitue une bonne alternative pour les transmissions radio avec voie de retour. En effet, elle a permis d’améliorer le débit utile en MP aux zones de faibles RSB et de conserver le gain en débit apporté par la stratégie MP à fort RSB. Elle garde également la robustesse de MP aux erreurs sur la voie de retour.

Pour calculer les performances des différents schémas HARQ, nous avons présenté une méthode analytique très générale basée sur la modélisation des protocoles HARQ par des chaînes de Markov à temps discret et à états finis. Les performances prédites par cette méthode ont été validées par des résultats de simulation.

Enfin, sous l’hypothèse d’un budget d’énergie limité pour la transmission d’un paquet et son acquitement, nous avons regardé, suivant le critère de maximisation du

débit utile, la répartition optimale de l'énergie entre le lien direct et la voie de retour. Les résultats obtenus montrent que le protocole hybride SP/MP demande moins de ressources énergétiques sur le lien retour que le SP, et donne un débit nettement supérieur au débit en SP au point optimal de fonctionnement.

Conclusion

Dans cette thèse, nous avons d'abord revisité les protocoles de retransmission élémentaires (ARQ) et évalué leurs performances en terme du débit offert et d'efficacité énergétique. Nous avons ensuite étendu l'évaluation des performances pour les schémas ARQ hybride ou HARQ classiques. Ces schémas, caractérisés par la retransmission de paquets relatifs à un seul paquet de données erroné (schémas simple-paquet ou HARQ-SP), souffrent d'une dégradation du débit utile après chaque retransmission. La comparaison entre les différents schémas HARQ classiques a montré que le schéma HARQ-II-IR est celui qui offre les meilleures performances.

Afin d'améliorer le débit utile en HARQ-SP, nous avons proposé une approche visant à réduire le nombre moyen de retransmissions par paquet de données. Cela a nécessité la création de paquets de redondance pouvant aider au décodage de plusieurs paquets de données utiles. Cette stratégie est nettement meilleure que la stratégie SP à moyen et à fort RSB. Elle est également plus robuste aux erreurs sur la voie de retour que le SP. A faible RSB, le débit offert par la stratégie SP est meilleur que celui offert par le MP. Pour améliorer le débit utile en MP à faible RSB et maintenir sa robustesse aux erreurs sur la voie de retour, nous avons proposé une nouvelle approche combinant les schémas SP et MP. La comparaison entre les différentes stratégies a été proposée dans cette thèse. Les performances de ces stratégies ont été évaluées théoriquement puis par simulation. Les résultats de simulation ont été en accord avec les modèles théoriques d'analyse des performances.

Les approches proposées se basent sur des hypothèses fortes. Nous avons supposé que les paquets transmis ont tous la même taille. D'autre part, nous avons considéré que le schéma de codage et modulation (MCS) est fixe durant tout le cycle ARQ. Or, dans certains scénarios, la retransmission de paquet de petite taille peut suffire pour effectuer un décodage correct. Nous évoquons ici deux points essentiels qui peuvent être pris en compte dans le modèle de transmission.

Dans tout le document, les acquittements ne transportent pas d'information supplémentaire autre que l'indication du bon ou mauvais décodage. Il est alors possible de renvoyer une information sur l'état du canal ou sur la fiabilité des symboles décodés :

Lorsque les acquittements transportent une information sur l'état du canal, la station A peut alors adapter sa transmission en choisissant le bon MCS. Le choix du rendement pourra se baser sur la maximisation du débit utile de transmission.

Les acquittements peuvent également transporter une information sur les symboles

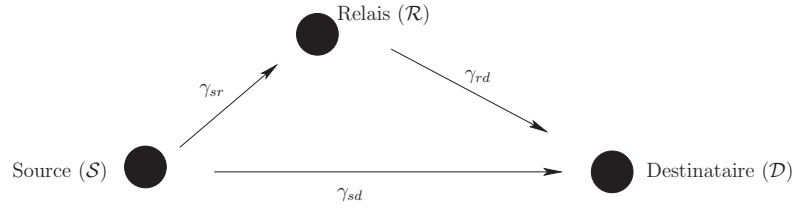


Figure R — Schéma coopératif avec présence d'un relais.

fortement affectés par le canal (symboles les moins fiables) ou bien sur la fiabilité des paquets erronés (en MP par exemple). Dans le premier cas, l'émetteur pourra choisir les symboles à retransmettre afin de réaliser un succès de décodage et donc une amélioration du débit utile. Dans le second cas, il est possible de concevoir un protocole privilégiant la retransmission de redondance combinée avec les symboles critiques. Les différents protocoles possibles dépendent du type d'information portée par l'acquiescement.

Dans cette thèse, avons considéré une transmission entre deux stations A (émettrice) et B (réceptrice). Dans ce cadre de transmission, nous avons comparé les performances des schémas de codage HARQ classiques (SP) avec des schémas HARQ proposés. Il est alors possible d'étendre cette étude pour des transmissions coopératives entre A (source $\mathcal{S}$) et B (destinataire $\mathcal{D}$) en présence de relais $\mathcal{R}$ (*Network Coding* ou codage réseau, cf. figure R), et de concevoir un protocole de retransmission plus performant que les schémas conventionnels. Dans ces derniers, les paquets erronés sont retransmis, dans la majorité des cas, par le relais lorsqu'il arrive à les décoder correctement. En effet, dans la plupart des scénarios, le relais est supposé plus proche du destinataire que la source. En spécifiant le scénario considéré (position de source et du relais par rapport au destinataire, gestion des transmissions, gestion des acquiescements, etc), il est possible de concevoir, pour le même scénario de transmission, un protocole hybride SP et MP qui retransmet, non seulement des paquets de redondance relatifs à un seul paquet de données erroné, mais aussi des paquets de redondance MP portant sur plusieurs paquets déjà transmis. On peut imaginer un scénario avec une source ou un destinataire mobiles. En effet, dans la plupart de temps, nous avons fait référence aux systèmes cellulaires. Le scénario de transmission coopérative pourra consister alors à un schéma avec mobilité du destinataire (lien descendant) ou de la source (lien montant) tout en gardant le relais fixe.

ANNEXE

A

Schéma HARQ hybride
SP et MPA.1 Expressions des probabilités $P^{xx}(m)$ et $Q_l^{xx}(m)$

Nous démontrons ici les expressions des probabilités $P^{xx}(m)$ et $Q_l^{xx}(m)$ présentées dans la section 5.3.1 du rapport. Revenons aux définitions de ces probabilités. Par exemple, la probabilité $P^{bb}(m)$ désigne la probabilité de l'événement "succès de décodage d'un paquet de données et succès de transmission de l'acquittement au round m sachant qu'il y a eu échecs de décodage aux rounds 0 à $m - 1$ ". De cette définition, il ressort :

$$\begin{aligned}
P^{bb}(m) &= \Pr\{\bar{e}_m, \bar{\theta} | e_o, \dots, e_{m-1}\} \\
&= \Pr\{\bar{e}_m, | e_o, \dots, e_{m-1}\} (1 - \epsilon) \\
&= (\Pr\{e_o, e_1, \dots, e_{m-1}, \bar{e}_m\} / \Pr\{e_o, \dots, e_{m-1}\}) (1 - \epsilon) \\
&= ((\Pr\{e_o, e_1, \dots, e_{m-1}\} - \Pr\{e_o, e_1, \dots, e_{m-1}, e_m\}) / \Pr\{e_o, \dots, e_{m-1}\}) (1 - \epsilon) \\
&= (1 - p(m)/p(m-1)) (1 - \epsilon).
\end{aligned} \tag{A.1}$$

D'une façon similaire, nous démontrons les expressions des probabilités $P^{xx}(m)$ pour $xx = \{bm, mb, mm\}$.

Considérons maintenant le cas des probabilités $Q_l^{xx}(m)$. Prenons le cas de la probabilité $Q_l^{bb}(m)$ d'avoir un "succès de décodage conjoint des L paquets de données et bonne réception de l'acquittement au round m sachant qu'il y a eu des échecs de décodage aux rounds 0 à $m - 1$ et qu'il y a eu l paquets de données bien décodés au premier round ($m = 0$)". Cette probabilité est exprimée ainsi :

$$\begin{aligned}
Q_l^{bb}(m) &= \Pr\{\bar{e}_m, \bar{\theta} | \mathcal{A}_l, e_1, \dots, e_{m-1}\} \\
&= \Pr\{\bar{e}_m, | \mathcal{A}_l, e_1, \dots, e_{m-1}\} (1 - \epsilon) \\
&= (\Pr\{\mathcal{A}_l, e_1, \dots, e_{m-1}, \bar{e}_m\} / \Pr\{\mathcal{A}_l, e_1, \dots, e_{m-1}\}) (1 - \epsilon) \\
&= ((\Pr\{\mathcal{A}_l, e_1, \dots, e_{m-1}\} - \Pr\{\mathcal{A}_l, e_1, \dots, e_{m-1}, e_m\}) / \Pr\{\mathcal{A}_l, e_1, \dots, e_{m-1}\}) (1 - \epsilon) \\
&= (1 - p_l(m)/p_l(m-1)) (1 - \epsilon).
\end{aligned} \tag{A.2}$$

De la même manière on exprime $Q_l^{xx}(m)$ pour $xx = \{bm, mb, mm\}$ en fonction de $p_l(m)$ et ϵ .

A.2 Probabilités de transition

Dans le chapitre 5, nous avons présenté le modèle de chaîne de Markov représentant le protocole H-SP-MP. Nous avons défini les probabilités $P^{xx}(m)$ et $Q_l^{xx}(m)$ nécessaires pour le calcul des probabilités de transition. Nous présentons ici en détail les probabilités de transition entre tous les états de la chaîne. Pour cela, on utilise la division de la chaîne en classes présentée dans le chapitre 5. Nous décrivons dans cette annexe les différentes probabilités de transition entre les états de la même classe et entre les états de classes différentes. Considérons d'abord les transitions à partir de l'état initial. Cet état peut appartenir à toutes les classes car on peut partir de cet état directement vers un état de la classe SP, "indéterminée", ou "MP".

A.2.1 Transitions à partir de l'état initial

A partir de l'état initial, on peut avoir quatre transitions possibles :

- Lorsque un paquet est correctement transmis et que son acquittement est bien reçu, l'émetteur et le récepteur libèrent leurs mémoires et initialisent leurs variables internes. On passe de l'état initial vers lui même. On a :

$$\Pr\{(0, 0, 0, 0) \rightarrow (0, 0, 0, 0)\} = P^{bb}(0) \quad (\text{A.3a})$$

- Lorsque un paquet est correctement transmis et que son acquittement est perdu, on passe de l'état initial vers l'état $(0, 1, 0, 1)$ du mode indéterminé. On a :

$$\Pr\{(0, 0, 0, 0) \rightarrow (0, 1, 0, 1)\} = P^{bm}(0) \quad (\text{A.3b})$$

- Lorsque un paquet est mal décodé et que son acquittement est bien reçu, on passe de l'état initial vers l'état $(1, 1, 1, 0)$ du mode SP. On a :

$$\Pr\{(0, 0, 0, 0) \rightarrow (1, 1, 1, 0)\} = P^{mb}(0) \quad (\text{A.3c})$$

- Lorsque un paquet est mal décodé et que son acquittement est perdu, on passe de l'état initial vers l'état $(1, 1, 0, 0)$ du mode MP. On a :

$$\Pr\{(0, 0, 0, 0) \rightarrow (1, 1, 0, 0)\} = P^{mm}(0) \quad (\text{A.3d})$$

A.2.2 Transitions en mode SP

Le mode SP correspond aux états $(a, b, 1, 1)$ avec $0 \leq a \leq 1$ et $1 \leq b \leq M_1$. Chaque round correspond à la transmission d'un paquet. Lorsque $a = 1$, cela signifie que le paquet de données n'a pas été correctement décodé. Par conséquent, pour $1 \leq b \leq M_1 - 1$ on a :

$$\Pr\{(1, b, 1, 0) \rightarrow (0, 0, 0, 0)\} = P^{bb}(b) \quad (\text{A.4a})$$

$$\Pr\{(1, b, 1, 0) \rightarrow (0, b+1, 1, 0)\} = P^{bm}(b) \quad (\text{A.4b})$$

$$\Pr\{(1, b, 1, 0) \rightarrow (1, b+1, 1, 0)\} = P^{mb}(b) + P^{mm}(b) \quad (\text{A.4c})$$

Les états $(0, b, 1, 0)$ sont atteints lorsque le paquet est bien décodé mais que l'émetteur n'en est pas informé (cf. équation A.4b). Il suffit alors de recevoir correctement l'acquittement pour revenir à l'état initial. On a pour $1 \leq b \leq M_1 - 1$:

$$\Pr\{(0, b, 1, 0) \rightarrow (0, b + 1, 1, 0)\} = \epsilon \quad (\text{A.5a})$$

$$\Pr\{(0, b, 1, 0) \rightarrow (0, 0, 0, 0)\} = 1 - \epsilon \quad (\text{A.5b})$$

Lorsque le nombre maximal de transmissions est atteint, on revient à l'état initial quelque soit le résultat de l'échange :

$$\Pr\{(a, M_1, 1, 0) \rightarrow (0, 0, 0, 0)\} = 1 \quad (\text{A.6})$$

A.2.3 Transitions en mode MP

Les états du mode MP sont caractérisés par $(a = 0, b \geq L)$ ou $(1 \leq a \leq L, 1 \leq b \leq L + M_2 - 1)$ et $(c = 0, d = 0)$. Au premier round, l'émetteur transmet L paquets de données numérotés de $V(A) + 1$ à $V(S) - 1$. Lorsque $b \leq L$, la variable b correspond au nombre de paquets en attente d'acquittement et $b = V(S) - V(A) + 1$. Lorsque $L < b \leq L + M_2 - 1$, la variable b permet d'indiquer le nombre de retransmissions et $b = L + m - 1$. À partir de l'état initial, on rentre en mode MP seulement si au moins un paquet n'a pas été reçu ($a \geq 1$) et que l'émetteur n'en a pas été informé immédiatement. La première phase du mode MP correspond aux transmissions de paquets de données. On a pour $1 \leq b \leq L - 1$ et $d = 0$:

$$\Pr\{(a, b, 0, 0) \rightarrow (a, b + 1, 0, 0)\} = 1 - p(0) \quad (\text{A.7a})$$

$$\Pr\{(a, b, 0, 0) \rightarrow (a + 1, b + 1, 0, 0)\} = p(0) \quad (\text{A.7b})$$

Lorsque L paquets ont été transmis, on passe à la phase de transmission de redondance MP. On a donc pour $L \leq b \leq L + M_2 - 2$ et $1 \leq a \leq L$:

$$\Pr\{(a, b, 0, 0) \rightarrow (0, b + 1, 0, 0)\} = Q_{L-a}^{bm}(b - L + 1) \quad (\text{A.8a})$$

$$\Pr\{(a, b, 0, 0) \rightarrow (a, b + 1, 0, 0)\} = Q_{L-a}^{mb}(b - L + 1) + Q_{L-a}^{mm}(b - L + 1) \quad (\text{A.8b})$$

$$\Pr\{(a, b, 0, 0) \rightarrow (0, 0, 0, 0)\} = Q_{L-a}^{bb}(b - L + 1) \quad (\text{A.8c})$$

L'équation (A.8a) correspond à un décodage correct des paquets (la variable a est mise à 0) mais comme l'émetteur ne reçoit pas l'acquittement, il continue à transmettre la redondance. Pour revenir à l'état initial, il suffit de recevoir correctement l'acquittement. On a pour $L \leq b \leq L + M_2 - 2$:

$$\Pr\{(0, b, 0, 0) \rightarrow (0, b + 1, 0, 0)\} = \epsilon \quad (\text{A.9a})$$

$$\Pr\{(0, b, 0, 0) \rightarrow (0, 0, 0, 0)\} = 1 - \epsilon \quad (\text{A.9b})$$

Lorsque le nombre maximal de paquets de redondance est atteint, on revient à l'état initial quoiqu'il arrive :

$$\Pr\{(a, L + M_2 - 1, 0, 0) \rightarrow (0, 0, 0, 0)\} = 1 \quad (\text{A.10})$$

A.2.4 Transitions en mode indéterminé

On définit un mode indéterminé caractérisé par la bonne réception d'au moins un paquet de données et un acquittement non reçu mais on peut revenir à l'état initial directement ou bien passer en SP ou en MP. Les états de ce mode sont caractérisés par $a \geq 0$, $1 \leq b \leq L - 1$, $c = 0$ et $1 \leq d \leq L - 1$.

Pour $1 \leq b \leq L - 2$, on peut passer de l'état $(0, b, 0, d)$ vers :

– l'état initial avec

$$\Pr\{(0, b, 0, d) \rightarrow (0, 0, 0, 0)\} = P^{bb}(0) \quad (\text{A.11a})$$

– l'état $(1, 1, 1, 0)$ du mode SP

$$\Pr\{(0, b, 0, d) \rightarrow (1, 1, 1, 0)\} = P^{mb}(0) \quad (\text{A.11b})$$

– l'état $(0, b + 1, 0, d + 1)$ du mode indéterminé

$$\Pr\{(0, b, 0, d) \rightarrow (0, b + 1, 0, d + 1)\} = P^{bm}(0) \quad (\text{A.11c})$$

– l'état $(1, b + 1, 0, d)$ du mode indéterminé

$$\Pr\{(0, b, 0, d) \rightarrow (1, b + 1, 0, d)\} = P^{mm}(0) \quad (\text{A.11d})$$

Pour $b = L - 1$, on peut passer de l'état $(0, L - 1, 0, d)$ vers :

– l'état initial avec

$$\Pr\{(0, b, 0, d) \rightarrow (0, 0, 0, 0)\} = P^{bb}(0) \quad (\text{A.12a})$$

– l'état $(1, 1, 1, 0)$ du mode SP

$$\Pr\{(0, L - 1, 0, d) \rightarrow (1, 1, 1, 0)\} = P^{mb}(0) \quad (\text{A.12b})$$

– l'état $(0, L, 0, 0)$ du mode MP

$$\Pr\{(0, L - 1, 0, d) \rightarrow (0, L, 0, 0)\} = P^{bm}(0) \quad (\text{A.12c})$$

– l'état $(1, L, 0, 0)$ du mode MP

$$\Pr\{(0, b, 0, d) \rightarrow (1, L, 0, 0)\} = P^{mm}(0) \quad (\text{A.12d})$$

Lorsque le nombre de paquets successifs à partir de l'état initial bien reçus et non acquittés est non nul ($a \neq 0$) et pour $1 \leq b \leq L - 1$, on a :

– On passe vers des états du mode MP avec les probabilités :

$$\Pr\{(a, b, 0, d) \rightarrow (a, b - d + 1, 0, 0)\} = P^{bb}(0) \quad (\text{A.13a})$$

$$\Pr\{(a, b, 0, d) \rightarrow (a + 1, b - d + 1, 0, 0)\} = P^{mb}(0) \quad (\text{A.13b})$$

– On reste au mode indéterminé avec les probabilités :

$$\Pr\{(a, b, 0, d) \rightarrow (a, b + 1, 0, d)\} = P^{bm}(0) \quad (\text{A.13c})$$

$$\Pr\{(a, b, 0, d) \rightarrow (a + 1, b + 1, 0, d)\} = P^{mm}(0) \quad (\text{A.13d})$$

Lorsque b atteint la valeur L , la valeur de d est mise à zéro. L'état avec une valeur de b égale à L est toujours un état du mode MP. En effet, l'émetteur retransmet immédiatement un paquet de redondance MP lorsqu'il ne reçoit pas l'acquiescement.

A.3 Performances du H-SP-MP avec connaissance du canal

Les schémas H-SP-MP avec connaissance du canal imposent que la station B s'abstient de renvoyer les acquittements négatifs pour un RSB supérieur γ_{th} . La modification par rapport au schéma H-SP-MP aveugle est alors dans le cas d'échec de décodage d'un paquet transmis par A. Lorsque le RSB moyen γ est connu, la chaîne de Markov modélisant le H-SP-MP, pour les RSBs supérieurs à γ_{th} , est identique à la chaîne modélisant le protocole MP pur. Pour les RSB inférieurs à γ_{th} , le protocole H-SP-MP avec connaissance de l'état du canal est identique au protocole H-SP-MP aveugle.

Lorsque la station B dispose uniquement d'une connaissance sur le RSB instantané, les seules probabilités de transition qui changent sont $P^{mb}(0)$ et $P^{mm}(0)$. C'est à dire les probabilités de passage au mode SP et MP. En effet, lorsque $\gamma_{inst} > \gamma_{th}$, la station B s'abstient de renvoyer les acquittements négatifs. On note ici par $\gamma_o = 2^{R_o} - 1$ le RSB à partir duquel la station B décode correctement le paquet. Il est à signaler que $\gamma_{th} < \gamma_o$ car pour $\gamma > \gamma_o$, le paquet est bien décodé et il est inutile de s'abstenir de renvoyer les acquittements (positifs dans ce cas). En gardant les mêmes notations de probabilités sur les transitions de la chaîne (sans regarder ici la signification du mb et mm), les probabilités $P^{mb}(0)$ et $P^{mm}(0)$ sur les transitions de la chaîne doivent prendre en compte la modification du protocole ($P^{mm}(0)$ est alors la probabilité de transition vers un état du mode MP et $P^{mb}(0)$ est la probabilité de transition vers un état du mode SP, cf. figure 5.3). On passe vers le mode MP si on a un échec de décodage du paquet ($\gamma_{inst} < \gamma_o$) et ($\gamma_{th} < \gamma_{inst}$) ou bien un échec de décodage et ($\gamma_{inst} < \gamma_{th}$) et échec de transmission de l'acquittement. On a donc :

$$\begin{aligned} P^{mm}(0) &= \Pr\{\gamma_{th} < \gamma_{inst} < \gamma_o\} + \Pr\{\gamma_{inst} < \gamma_{th}\}\epsilon \\ &= \exp(-\gamma_{th}/\gamma) - \exp(-\gamma_o/\gamma) + (1 - \exp(-\gamma_{th}/\gamma))\epsilon \end{aligned} \quad (A.15)$$

D'autre part, on passe vers un état du mode SP si on a un échec de décodage et $\gamma_{inst} < \gamma_{th}$ et un succès de transmission de l'acquittement. On alors :

$$\begin{aligned} P^{mb}(0) &= \Pr\{\gamma_{inst} < \gamma_{th}\}(1 - \epsilon) \\ &= (1 - \exp(-\gamma_{th}/\gamma))(1 - \epsilon) \end{aligned} \quad (A.16)$$

Les probabilités $P^{bb}(0)$ et $P^{bm}(0)$ restent inchangées et sont données par :

$$\begin{aligned} P^{bb}(0) &= \Pr\{\gamma_{inst} > \gamma_o\}(1 - \epsilon) \\ &= (\exp(-\gamma_o/\gamma))(1 - \epsilon) \end{aligned} \quad (A.17)$$

et

$$\begin{aligned} P^{bm}(0) &= \Pr\{\gamma_{inst} > \gamma_o\}\epsilon \\ &= (\exp(-\gamma_o/\gamma))\epsilon \end{aligned} \quad (A.18)$$

On vérifie bien que $P^{bb}(0) + P^{bm}(0) + P^{mb}(0) + P^{mm}(0) = 1$.

Bibliographie

- [1] M. El Aoun, R. Le Bidan, X. Lagrange and R. Pyndiah. Multiple-packet versus single-packet incremental redundancy strategies for type-II hybrid ARQ. In *International symposium on turbo codes and iterative information processing*, Sep. 2010.
- [2] M. El Aoun, X. Lagrange, R. Le Bidan and R. Pyndiah. Analyse des schémas HARQ classiques et évolués (schémas multi-paquets) en présence d'une voie de retour imparfaite. In *XXIIIe colloque GRETSI : traitement du signal et des images*, Sep. 2011.
- [3] M. El Aoun, X. Lagrange, R. Le Bidan and R. Pyndiah. Analysis and optimization of hybrid single packet and multiple-packets incremental redundancy in the presence of channel state information. In *The 14th International Symposium on Wireless Personal Multimedia Communications*, pages 1 – 5, Oct. 2011.
- [4] M. El Aoun, X. Lagrange, R. Le Bidan and R. Pyndiah. Throughput analysis of hybrid single-packet and multiple-packet truncated type-II HARQ strategies with unreliable feedback channel. In *IEEE Wireless Communications and Networking Conference*, Apr. 2012.
- [5] T. M. Cover and J. A. Thomas. *Elements of Information Theory*. John Wiley & Sons, 2006.
- [6] Y. Polyanskiy, H. V. Poor and S. Verdú. Dispersion of gaussian channels. In *Proc IEEE Int. Symp. Inform. Theory ISIT*, pages 2204–2208, Seoul, Korea, June 28-July 3, 2009.
- [7] Y. Polyanskiy, H. V. Poor and S. Verdú. Channel Coding Rate in the Finite Blocklength Regime. *IEEE Trans. Inform. Theory*, vol. 56, no. 5 :2307–2359, May 2010.
- [8] Y. Polyanskiy, and S. Verdú. Scalar coherent fading channel : dispersion analysis. *IEEE Int. Symp. Inf. Theory (ISIT)*, page 2979–2982, Saint Petersburg, Russia, Aug. 2011.
- [9] Stephen G. Wilson. *Digital Modulation and Coding*. Prentice Hall, 1995.
- [10] W. E. Ryan and S. Lin. *Channel Codes : Classical and Modern*. Cambridge University Press, 2009.
- [11] J.-J. Chang, D.-J. Hwang and M.-C. Lin. Some Extended Results on the Search for Good Convolutional Codes. *IEEE Trans. Inf. Theory*, vol. 43, no. 5, pages : 1682–1697, 1997.

- [12] P. Frenger, P. Orten and T. Ottosson. Convolutional Codes with Optimum Distance Spectrum. *IEEE Commun. Letters*, vol. 3, no. 11, pages : 317–319, 1999.
- [13] E. Malkamäki and H. Leib. Performance of truncated type-II hybrid ARQ schemes with noisy feedback over block fading channels. *IEEE Trans. Commun.*, vol. 48, no. 9 :1477–1487, Sep. 2000.
- [14] R. Knopp, and P. A. Humblet. On coding for block fading channels. *IEEE Trans. Inform. Theory*, vol. 46, n 1 :189–205, Jan 2000.
- [15] Xavier Lagrange et Dominique Seret. *Introduction aux Réseaux*. Hermes, 1998.
- [16] H.O. Burton and D.D. Sullivan. Errors and Error Control. *Proceedings of the IEEE*, vol. 60, pp : 1293–1303, 1972.
- [17] S. Lin, D.J. Costello and M.J. Miller. Automatic Repeat Request Error-Control Schemes. *IEEE Commun. Mag*, vol. 22, pages : 5–17, 1984.
- [18] Andrew Tanenbaum. *Computer Networks*. Prentice Hall PTR, 2003.
- [19] J. F. Chang and T. H. Yang. Multichannel ARQ protocols. *IEEE Trans. Commun.*, vol. 41, pp : 592–598, 1993.
- [20] M. Moeneclaey, H. Nurneel, I. Bruyland, and D. Y. Chung. Throughput optimization for a generalized stop-and-wait ARQ scheme. *IEEE Trans. Commun.*, COM-34, pp : 205–207, 1986.
- [21] Z. Ding and M. Rice. Throughput analysis of ARQ protocols for parallel multi-channel communications. *In Proceedings of IEEE GLOBECOM*, pp : 1279–1283, 2005.
- [22] Erik Dahlman, Stefan Parkvall, Johan Skold et Per Beming. *3G Evolution : HSPA and LTE for Mobile Broadband*. Academic Press Inc, 2007.
- [23] Moustapha El Aoun, Xavier Lagrange, Raphaël Le Bidan et Ramesh Pyndiah. *Efficacité des protocoles de retransmission ARQ*. Rapports Pracom, 2009.
- [24] R. J. Benice and A. H. Frey. An Analysis of Retransmission System. *IEEE Trans. Commun*, COM-12, pp : 135–145, 1964.
- [25] Shu Lin et Daniel J. Costello. *Error Control Coding*. Pearson-Prentice Hall, 2004.
- [26] Andrew Tanenbaum. *Réseaux, Architectures, protocoles, application*. Interéditions, 1991.
- [27] S. Lin and P.S. Yu. A Hybrid ARQ Scheme with Parity Retransmission for Error Control of Satellite Channels. *IEEE Trans. Commun*, vol. 30, no. 7 pages : 1701–1719, 1982.
- [28] S. B. Wicker. *Error Control System for Digital Communication and Storage*. Prentice Hall, 1995.
- [29] J. M. Morris. Optimal Blocklengths for ARQ Error Control Schemes. *IEEE Trans. Commun.*, vol. 27, pp : 488–493, 1979.
- [30] A.C. Martins and J.C. Alves. ARQ Protocols with Adaptive Block Size Perform Better over a Wide Range of Bit Error Rates. *IEEE Trans. Commun.*, vol. 38, no. 6, pp : 737–739, 1990.

- [31] S. Hara, A. Ogino, M. Araki, M. Okada and N. Morinaga. Throughput Performance of SAW-ARQ Protocol with Adaptive Packet Length in Mobile Packet Data Transmission. *IEEE Trans. Vehic. Tech.*, vol. 45, no. 3, pp : 561–569, 1996.
- [32] J. Qi, S. Aissa, X. Zhao. Optimal Frame Length for Keeping Normalized Goodput with Lowest Requirement on BER. *In Proc. IEEE International Conference on Innovations in Information Technology (Innovations'07)*, pp : 715–719, 2007.
- [33] Y. Sankarasubramaniam, I. F. Akyildiz, and S. W. McLaughlin. Energy efficiency based packet size optimization in wireless sensor networks. *In Proceedings of the 1st IEEE International Workshop on Sensor Network Protocols and Applications (SNPA '03)*, pp : 1–8, Anchorage, Alaska, USA, 2003.
- [34] S. Marcille, P. Ciblat, C. Le Martret . Etude au niveau IP d'un protocole ARQ Hybride avec voie de retour imparfaite. *GRETSI, Bordeaux (France)*, Sep. 2011.
- [35] Le Martret, A. Le Duc, S. Marcille, and P. Ciblat. Analytical Performance derivation of Hybrid ARQ Schemes at IP layer. *IEEE Trans. Commun.*, vol. 60, no. 5, pp : 1305–1314, May 2012.
- [36] D. Chase. Code combining - A maximum-likelihood decoding approach for combining an arbitrary number of noisy packets. *IEEE Trans. Commun.*, vol. COM-33, no. 5, pages : 385–393, 1985.
- [37] J. Hagenauer. Rate-compatible punctured convolutional codes (RCPC codes) and their applications. *IEEE Trans. Commun.*, vol. 36 no. 4, pages : 389–400, 1988.
- [38] D. N. Rowitch and L. B. Milstein. On the performance of hybrid FEC/ARQ systems using rate compatible punctured turbo (RCPT) codes. *IEEE Trans. Commun.*, 2000.
- [39] S. Kallel and D. Haccoun. Generalized Type II Hybrid ARQ Scheme Using Punctured Convolutional Coding. *IEEE Trans. Commun.*, vol. 38 :1938–1946, Nov. 1990.
- [40] G. Caire and D. Tuninetti. The Throughput of Hybrid-ARQ protocols for the Gaussian Collision Channel. *IEEE Transactions on Information Theory, Volume 47 n 5 - July 2001*, 2001.
- [41] Y. Polyanskiy, H. V. Poor and S. Verdú. New channel coding achievability bounds. *In Proc. IEEE Int. Symp. Inform. Theory ISIT*, pages 1763–1767, July 6-8, 2008.
- [42] Timo Halonen, Javier Romero and Juan Melero. *GSM, GPRS Performance and EDGE : Evolution Towards 3G/UMTS*. John Wiley & Sons Ltd, 2003.
- [43] J. Proakis. *Digital Communications*. Mc Graw Hill, 4 dition 2001.
- [44] S. Kallel. Sequential decoding with an efficient incremental redundancy ARQ strategy. *IEEE Trans. Commun.*, vol. 40, no. 10 :1588–1593, Oct. 1992.
- [45] C. Hausl, and A. Chindapol. Hybrid ARQ with cross-packet channel coding. *IEEE Comm. Letters.*, vol. 11, no. 5 :434–436, May. 2007.
- [46] R.Zhang, L.Hanzo. Superposition-Coding Aided Multiplexed Hybrid ARQ Scheme for Improved End-to-End Transmission Efficiency. *IEEE Trans. Veh. Tech.*, vol. 58 :4681–4686, Oct. 2009.

- [47] A. Graell i Amat, G. Montorsi, and F. Vatta. Design and performance analysis of a new class of rate compatible serially concatenated convolutional codes. *IEEE Trans. Commun.*, vol. 57, no.8 :2280–2289, Aug. 2009.
- [48] Claude Berrou, Alain Glavieux, and Punya Thitimajshima. Near Shannon limit error-correcting coding and decoding : turbo-codes. In *IEEE ICC '93, Geneva*, pages 1064 – 1070, 1993.
- [49] Divsalar D. and Pollara F. Multiple turbo codes. In *IEEE Milcom-95*, volume 1, pages 279 – 285, Nov. 1995.
- [50] D. M. Rankin and T. A. Gulliver. Single parity-check product codes. *IEEE Trans. Commun.*, vol. 49, no.8 :1354–1362, Aug. 2001.
- [51] R. Pyndiah. Near-Optimum Decoding of Product Codes : Block Turbo Codes. *IEEE Trans. Commun.*, vol. 46, no.8 :1003–1010, Aug. 1998.
- [52] M. Zorzi and R. Rao. On the Use of Renewal Theory in the Analysis of ARQ protocols. *IEEE Trans. Commun.*, vol. 44 :1077–1081, Sep. 1996.
- [53] L. Badia, M. Levorato, M. Zorzi. Markov analysis of selective repeat type II hybrid ARQ using block codes. *IEEE Trans. Commun.*, vol. 56 :1434–1441, Sep. 2008.
- [54] R. W. Wolff. *Stochastic Modeling and the Theory of Queues*. Prentice-Hall, New York, NY, 1989.