

THÈSE DE DOCTORAT DE

L'UNIVERSITÉ DE RENNES 1

ÉCOLE DOCTORALE N° 601

*Mathématiques et Sciences et Technologies
de l'Information et de la Communication*

Spécialité : Signal, Image, Vision

Par

Marie-Anne LACROIX

Détection d'événements sonores environnementaux dans les réseaux de capteurs sans fil à faible consommation

Thèse présentée et soutenue à Lannion, le 05/04/2022

Unité de recherche : IRISA, UMR CNRS 6074

Rapporteurs avant soutenance :

Jean-Marie GORCE Professeur des Universités, INSA Lyon
Romain SERIZEL Maître de Conférences, Université de Lorraine

Composition du Jury :

Examineurs :	Jean-Marie GORCE	Professeur des Universités, INSA Lyon - Rapporteur
	Yannis POUSSET	Professeur des Universités, Université de Poitiers
	Romain SERIZEL	Maître de Conférences, Université de Lorraine - Rapporteur
Dir. de thèse :	Pascal SCALART	Professeur des Universités, Université de Rennes 1
Encadrants :	Nancy BERTIN	Chargée de Recherche, Irisa (Rennes)
	Romuald ROCHER	Maître de Conférences, Université de Rennes 1

Invité(s) :

David Sodoyer Chargé de Recherche, IFSTTAR (Lille)

REMERCIEMENT

La thèse, à bien des égards, est un travail extrêmement solitaire. Néanmoins, j'ai eu la chance d'être entourée de nombreuses personnes qui, chacune à leur manière, m'ont permis d'arriver au bout de ce marathon qu'est le doctorat.

En premier lieu, je remercie bien évidemment mes encadrants pour leur soutien et leurs conseils tout au long de ces trois années : Pascal Scalart, Romuald Rocher et Nancy Bertin. Il n'est pas simple de se comprendre lorsque l'on se trouve dans des domaines de recherche et des villes différents, et c'est pourtant une tâche particulièrement réussie pour notre groupe. Chacun d'entre eux a été complémentaire de l'autre, et cela a été un plaisir de travailler toutes ces années avec eux.

Je tiens également à remercier l'ensemble du jury, MM Jean-Marie Gorce, Yannis Pousset, Romain Serizel et David Sodoyer, pour leur bienveillance et l'intérêt qu'ils ont porté à mes recherches.

Ma reconnaissance va aussi à toutes les personnes que j'ai pu rencontrer au cours de ce long trajet au sein de l'Irisa. Merci aux équipes Granit et Cairn pour m'avoir prouvé que la dégustation de galettes tout au long de janvier est possible et appris les autres significations de l'acronyme "CSI". Merci à Panama, pour leur accueil chaleureux malgré la brièveté de leur échange. Cette thèse n'aurait pas été la même sans Audrey et ses cours de sport, Marwa et sa bonne humeur, ou encore Robin et ses conseils avisés. Merci aussi à toutes les personnes rencontrées au cours de ces années pour leurs discussions, leur gentillesse et leurs conseils, souvent prodigués malgré le Covid et la distance : Joëlle, Peggy, César, Claire, et beaucoup d'autres que je ne saurais tous citer.

Je souhaite également remercier ma famille pour l'intérêt qu'ils ont porté à mon travail, même si n'était pas toujours facile ni à expliquer ni à comprendre. Un grand merci aux amis avec qui j'ai traversé ces dernières années, même quand le lien était difficile à conserver et malgré les différentes épreuves qui se sont présentées. Mes pensées vont également à tous les membres du serveur Discord PhD students avec qui j'ai partagé conseils, soutien, et de nombreux autres moments qui ont permis de briser la solitude et de mieux comprendre le travail de thèse en particulier et le monde en général.

Enfin, je souhaite exprimer mes remerciements les plus sincères à Erwan, qui est resté à mes côtés malgré les nombreuses difficultés qui se sont posées, et qui continue de m'offrir jour après jour son soutien et sa confiance indéfectibles.

TABLE DES MATIÈRES

1	Introduction	1
1.1	Introduction générale	1
1.2	Motivations et problématiques	2
1.3	Contributions de nos travaux	3
1.4	Structure de la thèse	5
2	Détection d'événements sonores	9
2.1	Introduction	9
2.2	Généralités sur la classification acoustique	10
2.2.1	Les types de la classification	10
2.2.2	Les étapes d'une détection d'événements acoustiques	12
2.2.3	Les propriétés des données audio	15
2.2.4	Le choix d'un classifieur	16
2.3	Les outils de la classification	18
2.3.1	Les bases de données	18
2.3.2	Descripteurs d'un signal audio	21
2.3.3	Critères de performance en classification	23
2.4	Les modèles de mélange gaussien	27
2.4.1	Fonctionnement théorique d'un modèle de mélange gaussien soustractif	28
2.4.2	Implémentation et vérification de l'algorithme	30
2.4.3	Etude du comportement	36
2.5	Le réseau neuronal convolutif et récurrent	41
2.5.1	Fonctionnement théorique d'un réseau neuronal convolutif et récurrent	41
2.5.2	Implémentation et vérification de l'algorithme	47
2.5.3	Etude du comportement	50
2.6	Conclusion	50
3	Consommation d'énergie d'un système communicant	53
3.1	Introduction	53
3.2	Transmission point-à-point au travers d'un canal de communication	53
3.2.1	Modèle d'une communication point-à-point	54
3.2.2	Le canal de propagation	55
3.2.3	Modulation du signal	59
3.3	Transmission coopérative à l'aide d'un noeud-relai	64
3.3.1	Modèle d'une communication avec relai	64
3.3.2	Protocoles de coopération	65
3.3.3	Combinaison des signaux à la réception	67
3.4	Mesurer les performances de la transmission	70
3.4.1	Probabilité d'erreur en réception	70
3.4.2	Efficacités du canal	74
3.4.3	Capacité de Shannon	75

3.4.4	Métriques du coopératif	78
3.5	Consommation d'énergie d'une communication	78
3.5.1	Modèle énergétique d'une transmission point-à-point	78
3.5.2	Modèle énergétique d'une transmission coopérative	84
3.6	Lien entre énergie et efficacité spectrale	93
3.6.1	Lien entre consommation d'énergie et capacité	93
3.6.2	Capacité pour un signal modulé	94
3.6.3	Minimisation de la consommation	98
3.7	Conclusion	100
4	Analyse d'un modèle de consommation réaliste	101
4.1	Introduction	101
4.2	Vers un modèle de consommation plus réaliste	101
4.2.1	Coefficient de l'amplificateur de puissance	102
4.2.2	Puissance des convertisseurs	110
4.2.3	Limitation de la puissance de transmission	113
4.3	Une meilleure estimation de l'efficacité spectrale b^*	115
4.3.1	Capacité du canal	116
4.3.2	Efficacité spectrale pour un modèle de consommation plus réaliste	120
4.3.3	Résultats expérimentaux	122
4.4	Influence de la position du relai	126
4.4.1	Variation du relai sur l'axe source-destinataire	127
4.4.2	Variation du relai dans un espace à deux dimensions	129
4.5	Conclusion	133
5	Détection d'événements sonores au sein d'un système embarqué multicap-	135
	teurs	
5.1	Introduction	135
5.2	Contexte d'une classification sonore embarquée	136
5.2.1	Introduction d'un système distribué multi-nœuds	136
5.2.2	Exemples de systèmes utilisables	139
5.3	Introduction d'éléments de spatialisation dans la classification	140
5.3.1	Descripteurs spatiaux	141
5.3.2	Utilisation des différents canaux et nœuds au sein du système	144
5.4	Amélioration des données dans un système contraint	146
5.4.1	Augmentation de données compatibles avec la polyphonie	146
5.4.2	Expériences sur le recouvrement et la technique du <i>mixup</i>	147
5.5	Encodage des descripteurs en virgule fixe	150
5.5.1	L'algorithme Linde-Buzo-Gray (LBG)	151
5.5.2	Quantification des descripteurs	152
5.5.3	Résultats expérimentaux	152
5.6	Influence de la transmission sur la classification d'événements sonores	159
5.6.1	Présentation du système	159
5.6.2	Résultats expérimentaux	160
5.6.3	Augmentation de données pour la transmission de descripteurs	165
5.7	Conclusion	168

6	Conclusions et perspectives	169
6.1	Conclusion	169
6.2	Perspectives	172
6.3	Publications	174
A	Liste des descripteurs dits "artisansaux"	175
B	Description de l'initialisation aléatoire d'un GMM	176
C	Preuves de convexité et calcul de b^*	177
C.1	Démonstration de la convexité de l'énergie totale pour un modèle simplifié issue de l'équation 3.101	177
C.2	Calcul de b^* pour un modèle simplifié de l'énergie	177
C.3	Validation de la convexité de l'énergie totale pour un modèle avancé sur un canal additive white Gaussian noise (AWGN) issue de l'équation 4.26	178
D	Détails sur le facteur de roll-off et l'efficacité du drain	180
D.1	Origine du facteur de roll-off	180
D.2	Calcul et propriétés de l'efficacité du drain	180
	Bibliographie	183

TABLE DES FIGURES

1.1	Visualisation des contributions apportées dans cette thèse	3
1.2	Schématisation de la structure de la thèse	6
2.1	Schéma récapitulatif des différentes méthodes de classification acoustique	12
2.2	Schéma récapitulatif des différentes étapes de la classification d'événements sonores	13
2.3	Etapes de calcul des coefficients <i>mel-frequency cepstral coefficients</i> (MFCC) . . .	22
2.4	Comptage des vrais et faux positifs, et vrais et faux négatifs, pour une classifi- cation monophonique	24
2.5	Schéma du comptage des substitutions, délétions et insertions	24
2.6	Schématisation des conséquences du micro et du macro-moyennage	27
2.7	Schéma de fonctionnement du <i>Gaussian mixture model</i> (GMM) soustractif . . .	28
2.8	Evolution des paramètres du GMM soustractif au cours du temps (en époque) pour la classe <i>people speaking</i>	32
2.9	Exemple théorique de distribution de la différence de log-vraisemblance pour l'ensemble des échantillons sur lesquels la classe est présente (dite distribution de présence) ou absente (dite distribution d'absence)	34
2.10	Courbe du taux d'erreur en fonction de la F-mesure pour différents seuils de détection	35
2.11	Courbe du taux d'erreur en fonction de la F-mesure pour différentes valeurs du nombre de composantes gaussiennes	37
2.12	Courbe du taux d'erreur en fonction de la F-mesure pour différents seuils de détection et une matrice de covariance pleine ou diagonale	38
2.13	Courbe du taux d'erreur en fonction de la F-mesure pour différents seuils de détection, sans ou avec normalisation pour une matrice pleine	40
2.14	Courbe du taux d'erreur en fonction de la F-mesure pour différents seuils de détection, sans ou avec normalisation pour une matrice diagonale	41
2.15	Graphique des différentes couches du modèle du CRNN (partie 1)	42
2.16	Graphique des différentes couches du modèle du CRNN (partie 2)	43
2.17	Fonctionnement d'une validation croisée à 4 plis	46
2.18	Courbe du taux d'erreur en fonction de la F-mesure pour différentes époques . .	49
2.19	Courbe du taux d'erreur en fonction de la F-mesure micro ou macro-moyenné pour différents seuils de décision	49
3.1	Schéma simplifié du canal physique de communication	54
3.2	Présentation des différents modules d'une communication sans fil	55
3.3	Schéma des différentes déperditions subies par le signal sur le canal	56
3.4	Distribution de la loi de Rayleigh pour différentes valeurs du paramètre σ	58
3.5	Exemple de décomposition du signal binaire en symboles et créneaux associés . .	60
3.6	Constellation de la modulation 16-QAM	62
3.7	Constellation de la modulation 8-PSK	63
3.8	Schéma simplifié du canal physique d'une communication coopérative	64

3.9	Probabilité d'erreur symbole en fonction du rapport signal sur bruit (RSB) symbole E_s/N_0 en dB, pour une communication point-à-point avec un signal modulé par 2 ⁶ -QAM	72
3.10	Probabilité d'erreur symbole en fonction du RSB symbole E_s/N_0 en dB, pour une communication coopérative avec protocole <i>amplify-and-forward</i> (AF) et un signal modulé par 2 ⁶ -QAM	73
3.11	Capacité de Shannon normalisée en fonction du RSB symbole	76
3.12	Efficacité énergétique η_{EE} normalisée en fonction de l'efficacité spectrale θ . . .	77
3.13	Blocs du circuit de l'ensemble transmetteur-récepteur	79
3.14	Energie consommée par bit pour une M-QAM ($M = 2^b$) et un canal AWGN ou de Rayleigh en fonction de l'ordre de modulation et pour différentes probabilité d'erreur binaire (PEB)	83
3.15	Energies totale, de transmission et du circuit consommées par bit pour une M-QAM($M = 2^b$), un canal AWGN et une distance $d_{sd} = 1\text{m}$ en fonction de la taille de constellation b de la modulation et pour PEB de 10^{-3}	83
3.16	Schéma simplifié du canal physique de communication coopérative	85
3.17	Energie consommée par bit lors d'une transmission coopérative par une M-QAM ($M = 2^b$) et un canal AWGN ou de Rayleigh en fonction de l'ordre de modulation et pour différentes PEB (nœud-relai à mi-distance source-destinataire)	87
3.18	Gain de coopération entre une transmission point-à-point et coopérative pour un canal AWGN ou de Rayleigh en fonction de l'ordre de modulation et pour différentes PEB (nœud-relai à mi-distance source-destinataire)	88
3.19	Gain de coopération entre une transmission point-à-point et coopérative pour un canal AWGN ou de Rayleigh en fonction de la localisation du relai et pour différentes PEB	89
3.20	Énergie consommée par bit transmis avec succès lors d'une communication coopérative sur des canaux de Rayleigh et avec un signal modulé par M-QAM . . .	91
3.21	Gain de coopération pour l'énergie de transmission avec succès lors d'une communication sur des canaux de Rayleigh et avec un signal modulé par M-QAM . . .	92
3.22	Efficacité spectrale en fonction du RSB symbole et PEB en fonction de la pénalité en RSB Γ_{AWN} pour une source modulée par M-QAM dans un canal AWGN . . .	94
3.23	Efficacité spectrale en fonction du RSB symbole pour une source gaussienne et PEB en fonction de la pénalité du RSB liée à l'utilisation d'une source M-QAM au lieu d'une source gaussienne pour un canal de Rayleigh à évanouissement plat . . .	97
3.24	Comparaison de l'énergie consommée par bit pour le modèle exact et le modèle calculé à partir de la pénalité	98
3.25	Évolution de l'efficacité spectrale b^*_d minimisant la consommation d'énergie par bit pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction de la distance d_{sd}	100
4.1	PAPR en fonction du facteur de roll-off ρ pour une 4-QAM et une 64-QAM . . .	104
4.2	Comparaison de l'énergie consommée par bit lors d'une communication point-à-point pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction de l'efficacité spectrale (ES) b pour un <i>peak to average power ratio</i> (PAPR) $\xi = 1$ ou dépendant du facteur de roll-off pour différentes valeurs de celui-ci	105

4.3	Comparaison de l'énergie consommée par bit lors d'une communication coopérative (protocole AF) pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction de l'ES b pour un PAPR $\xi = 1$ ou dépendant du facteur de roll-off pour différentes valeurs de celui-ci	105
4.4	Evolution de l'efficacité du drain normalisée en fonction de la puissance de transmission pour des émetteurs-récepteurs TI CC2420 et TI CC1000	107
4.5	Comparaison de l'énergie consommée par bit lors d'une communication point-à-point pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction de l'ES b pour une efficacité du drain $\eta = 0.35$ ou dépendante de la puissance de transmission pour différentes valeurs de l'exposant β	108
4.6	Comparaison de l'énergie consommée par bit lors d'une communication point-à-point pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction de l'ES b pour un modèle du coefficient de l'amplificateur de puissance (AP) constant ou dépendant du système pour différentes valeurs de l'exposant β	109
4.7	Gain de coopération sur un canal AWGN ou de Rayleigh en fonction de l'ES b pour un modèle du coefficient de l'AP constant ou dépendant du système pour différentes valeurs de l'exposant β	110
4.8	Schéma interne simplifié du convertisseur numérique-analogique (CNA)	111
4.9	Comparaison de l'énergie consommée par bit lors d'une communication point-à-point pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction de l'ES b pour différents modèles de P_c et différentes valeurs de la largeur de bande passante	113
4.10	Comparaison de l'énergie consommée par bit lors d'une communication point-à-point pour une M-QAM ($M = 2^b$) sur un canal AWGN en fonction de l'ES b pour différents modèles de P_c et différentes valeurs de la résolution n_2	114
4.11	Comparaison de l'énergie consommée par bit lors d'une communication point-à-point pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction de l'ES b et pour une puissance de transmission P_t limitée ou non	115
4.12	Probabilité d'erreur binaire d'une M-QAM sur un canal AWGN en fonction de la pénalité Γ_{AWGN} du RSB par symbole pour différents ordres de modulation et différentes approximations	117
4.13	Efficacité spectrale en fonction du RSB symbole pour différentes approximations	117
4.14	Probabilité d'erreur binaire d'une M-QAM sur un canal de Rayleigh en fonction de la pénalité $\Gamma_{Rayleigh}$ du RSB par symbole pour différents ordres de modulation et pour l'approximation 4.22	118
4.15	Évolution de l'efficacité spectrale (ES) b^*_d minimisant la consommation d'énergie par bit pour une M-QAM ($M = 2^b$) et un canal AWGN ou de Rayleigh en fonction de la d_{sd}	120
4.16	Efficacité spectrale b^* minimisant l'énergie consommée par bit transmis pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction de la distance d_{sd}	122
4.17	Efficacité spectrale b^* minimisant l'énergie consommée par bit transmis pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction de la largeur de bande B	123

4.18	Efficacité spectrale b^* minimisant l'énergie consommée par bit transmis pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction du coefficient de l'efficacité du drain β	124
4.19	Efficacité spectrale b^* minimisant l'énergie consommée par bit transmis pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction du facteur roll-off ρ	125
4.20	Efficacité spectrale b^* minimisant l'énergie consommée par bit transmis pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction de la PEB	125
4.21	Efficacité spectrale b^* minimisant l'énergie consommée par bit transmis pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction de la résolution binaire n_1 du convertisseur numérique-analogique (CNA)	126
4.22	Énergie consommée par bit lors d'une communication coopérative pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction de l'ES pour différentes positions du relai sur l'axe source-destinataire	128
4.23	Gain de coopération pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction de l'ES pour différentes positions du relai sur l'axe source-destinataire	128
4.24	Gain de coopération pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction de la position normalisée du relai pour différentes tailles de constellation	129
4.25	Gain de coopération pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction de la position normalisée du relai pour différentes PEB	130
4.26	Energie minimale consommée par bit lors d'une communication coopérative AF en fonction de la position du relai pour un canal AWGN, une ES $b = 2$ et une distance source-destinataire $d_{sd} = 10\text{m}$	130
4.27	Energie minimale consommée par bit lors d'une communication coopérative AF en fonction de la position du relai pour un canal de Rayleigh, une ES $b = 2$ et une distance source-destinataire $d_{sd} = 10\text{m}$	131
4.28	Efficacité spectrale minimisant l'énergie consommée par bit émis lors d'une transmission coopérative AF en fonction de la position du relai pour un canal de Rayleigh et une distance source-destinataire $d_{sd} = 10\text{m}$	132
4.29	Gain de coopération en fonction de la position du relai pour un canal AWGN, une ES $b = 2$ et une distance source-destinataire $d_{sd} = 10\text{m}$	132
4.30	Gain de coopération en fonction de la position du relai pour un canal de Rayleigh, une ES $b = 2$ et une distance source-destinataire $d_{sd} = 10\text{m}$	133
5.1	Étapes de la transmission des données d'un nœud de capteur vers le moniteur central	138
5.2	Vue schématique du format d'émission d'une trame selon la norme IEEE 802.15.4 [81]	140
5.3	Principe de l'algorithme LBG	151
5.4	Schéma simplifié des différents formes par lesquelles passent les descripteurs lors de leur transmission sur le canal de communication	159
5.5	Évolution de la F-mesure pour un GMM soustractif muni d'une matrice de covariance pleine et dont les paires de descripteurs sont codées sur n bits, en fonction de l'apprentissage et pour différentes PEB en réception	161

5.6	Évolution de la F-mesure pour un GMM soustractif muni d'une matrice de covariance diagonale et dont les paires de descripteurs sont codées sur n bits, en fonction de l'apprentissage et pour différentes PEB en réception	163
5.7	Évolution de la F-mesure pour un <i>convolutional and recurrent neural network</i> (CRNN) et dont les paires de descripteurs sont codées sur n bits, en fonction de l'apprentissage et pour différentes PEB en réception	164
5.8	Évolution de la F-mesure pour un CRNN et dont les paires de descripteurs sont codées sur n bits, pour différentes PEB en réception et un apprentissage réalisé avec ou sans augmentation de données	167
C.1	Dérivée seconde de $f_1 f_2$ en fonction de l'efficacité spectrale b	179
D.1	Exemple du rôle du facteur de roll-of ρ pour un filtre <i>square root raised cosine</i> (SRRC). Le dépassement de la bande passante est indiqué pour $\rho = 0.5$	180
D.2	Schéma d'un transistor et des différents puissances impliquées	181

LISTE DES TABLEAUX

2.1	Bases de données réelles, polyphoniques et bicanales possédant un étiquetage fort	20
2.2	Statistiques de la base de données TUT2017-Development en fonction des classes	20
2.3	Statistiques de la base de données TUT2017-Evaluation en fonction des classes	21
2.4	Paramètres utilisés pour le calcul du mel-spectrogramme et des MFCC	23
2.5	Description de l'algorithme d'espérance-maximisation	29
2.6	Paramètres utilisés dans le GMM soustractif	30
2.7	Comparaison entre les taux d'erreur (ER) et F-mesure de notre implémentation et de celle de référence	33
2.8	Comparaison entre les taux d'erreur (ER) et F-mesure pour l'initialisation aléatoire et par K-moyenne	35
2.9	Comparaison entre les taux d'erreur (ER) et F-mesure pour une matrice de covariance pleine et diagonale	39
2.10	Comparaison entre les taux d'erreur (ER) et F-mesure pour des descripteurs normalisés ou non	40
2.11	Paramètres utilisés dans chaque bloc convolutif	45
2.12	Paramètres utilisés dans le CRNN	46
2.13	Comparaison entre les taux d'erreur (ER) et F-mesure de notre implémentation et de celle de référence	48
2.14	Comparaison entre les taux d'erreur (ER) et F-mesure de notre implémentation avec ou sans normalisation	50
3.1	Paramètres du modèle	81
3.2	Valeurs de b^* en fonction du canal, de la distance source-destinataire et la PEB	84
4.1	Efficacité spectrale discrète b_d^* d'une communication point-à-point pour une M-QAM ($M = 2^b$) et un canal AWGN ou de Rayleigh en fonction du modèle de l'efficacité du drain	107
4.2	Paramètres du convertisseur analogique-numérique (CAN) et du CNA	112
5.1	Méthode d'extraction des <i>time difference of arrival</i> (TDOA)	143
5.2	Comparaison entre les taux d'erreur (ER) et F-mesure pour une classification par GMM soustractifs prenant en entrée les MFCC et leurs coefficients dérivés ou des TDOA (1)	145
5.3	Comparaison entre les taux d'erreur (ER) et F-mesure pour une classification par GMM soustractifs prenant en entrée les MFCC et leurs coefficients dérivés ou des TDOA (2)	145
5.4	Comparaison entre les taux d'erreur (ER) et F-mesure pour une classification par GMM soustractif pour un recouvrement entre les trames de 50% ou de 75%	148
5.5	Comparaison entre les taux d'erreur (ER) et F-mesure pour une classification par CRNN pour un recouvrement entre les trames de 50% ou de 75%	148
5.6	Description de la technique du <i>mixup</i>	149

5.7	Comparaison entre les taux d'erreur (ER) et F-mesure pour une classification par CRNN dont l'entraînement est réalisé avec la base de données (BDD) d'origine ou avec une augmentation de données (AD) par technique du <i>mixup</i>	150
5.8	Comparaison des taux d'erreur (ER) et F-mesure à l'issue d'un test sur différentes longueurs de quantification pour une classification par GMM soustractif paramétré par une matrice pleine et dont l'apprentissage a été réalisé pour des descripteurs codés en virgule flottante sur 64 bits	156
5.9	Comparaison des taux d'erreur (ER) et F-mesure à l'issue d'un test sur différentes longueurs de quantification pour une classification par GMM soustractif paramétré par une matrice pleine et dont l'apprentissage a été réalisé sur les données quantifiées correspondantes	156
5.10	Comparaison des taux d'erreur (ER) et F-mesure à l'issue d'un test sur différentes longueurs de quantification pour une classification par GMM soustractif paramétré par une matrice diagonale et dont l'apprentissage a été réalisé pour des descripteurs codés en virgule flottante sur 64 bits	157
5.11	Comparaison des taux d'erreur (ER) et F-mesure à l'issue d'un test sur différentes longueurs de quantification pour une classification par GMM soustractif paramétré par une matrice diagonale et dont l'apprentissage a été réalisé sur les données quantifiées correspondantes	157
5.12	Comparaison des taux d'erreur (ER) et F-mesure à l'issue d'un test sur différentes longueurs de quantification pour une classification par CRNN dont l'apprentissage a été réalisé pour des descripteurs codés en virgule flottante sur 64 bits	158
5.13	Comparaison des taux d'erreur (ER) et F-mesure à l'issue d'un test sur différentes longueurs de quantification pour une classification par CRNN dont l'apprentissage a été réalisé sur les données quantifiées correspondantes	158
5.14	Description de la méthode d'AD proposée	166

ACRONYMES

- AD** augmentation de données. 12, 156, 157, 159, 160, 169, 176, 177, 182–184
- AF** *amplify-and-forward*. 7–9, 16, 17, 75–78, 80, 82, 83, 87, 94, 95, 97, 115, 140–142, 181, 183
- AP** amplificateur de puissance. 8, 88, 111, 112, 116–120, 125, 133, 144
- ASK** *amplitude-shift keying*. 70
- AWGN** additive white Gaussian noise. 5, 7–9, 11, 15–17, 68, 80–84, 90–93, 96–98, 102–109, 111, 114–120, 122–127, 129–143, 181, 188, 189
- BDD** bases de données. 20, 27–29
- BDD** base de données. 12, 17, 27–30, 41, 49, 145, 146, 156, 159, 160, 182, 184
- CAG** contrôle automatique du gain. 64, 77, 78
- CAN** convertisseur analogique-numérique. 9, 11, 15, 87–91, 94, 110–112, 120–122, 124, 132, 135, 136, 144, 149, 180, 181
- CF** *compress-and-forward*. 76
- CNA** convertisseur numérique-analogique. 8, 9, 11, 15, 87, 89, 90, 94, 109, 111–113, 120–122, 124, 132, 135, 136, 144, 180, 181
- CNN** *convolutional neural network*. 27, 152
- CQCC** *constant-Q cepstral coefficients*. 186
- CRNN** *convolutional and recurrent neural network*. 10–12, 16, 17, 20, 28, 32, 33, 51, 56–58, 60, 61, 152, 157–160, 162–165, 168–170, 172, 174, 175, 177, 182, 183
- DCASE** *Detection and Classification of Acoustic Sound Events*. 13, 28, 30, 37, 43, 51, 57, 147
- DDF** *dynamic decode-and-forward*. 76
- DF** *decode-and-forward*. 75, 76, 87
- EE** efficacité énergétique. 84, 102, 111, 116, 118, 125
- EGC** *equal gain combiner*. 77
- EM** espérance-maximisation. 39, 41, 44, 187
- ES** efficacité spectrale. 7–9, 15–17, 84, 102–104, 107–109, 111, 114–120, 123–133, 135, 137, 138, 141–144, 181, 183
- FPGA** *field-programmable gate array*. 160, 178
- FSK** *frequency-shift keying*. 70
- GFCC** *gammatone-frequency cepstral coefficients*. 186
- GMM** *Gaussian mixture model*. 5, 6, 10–12, 16, 17, 20, 27, 28, 32, 33, 37–41, 43, 46, 47, 49, 51, 54, 57, 58, 61, 148, 152, 155, 157–159, 162–167, 169–173, 175, 182, 183, 187
- GRU** *gated recurrent unit*. 54, 55

- HMM** *hidden Markov model*. 27
- IFA** *intermediate frequency amplifier*. 88, 90
- IoT** *internet des objets*. 13, 14, 91
- LBG** *Linde-Buzo-Gray*. 4, 9, 161, 162, 170, 182
- LNA** *low noise amplifier*. 88, 90
- LPWAN** *low power wireless access network*. 180
- LSTM** *long short-term memory*. 27
- LTW DF** *Laneman-Tse-Wornell decode-and-forward*. 76
- MAQ** *modulation d'amplitude en quadrature*. 71, 80, 83
- MDA** *modulation par déplacement d'amplitude*. 70
- MDF** *modulation par déplacement de fréquence*. 70
- MDP** *modulation par déplacement de phase*. 72
- MEMS** *micro-electro-mechanical system*. 180
- MFCC** *mel-frequency cepstral coefficients*. 6, 11, 31–33, 37, 41, 49, 51, 148, 152, 154, 155, 162, 172, 182, 184, 186
- MRC** *maximal ratio combiner*. 77, 78, 80, 82, 95, 98, 99, 137, 183
- NAF** *non-orthogonal amplify-and-forward*. 76
- NBK DF** *Nabar-Bölcskei-Kneubühler decode-and-forward*. 76
- NMF** *non-negative matrix factorization*. 27, 32
- OAF** *orthogonal amplify-and-forward*. 76
- PAPR** *peak to average power ratio*. 7, 8, 15, 89, 111–115, 117, 118, 129, 133, 144, 180, 181
- PEB** *probabilité d'erreur binaire*. 7, 9–11, 16, 17, 78, 79, 91–93, 95–99, 103–108, 111, 125–130, 132, 135, 139, 140, 144, 169–171, 173–178, 181, 183
- PES** *probabilité d'erreur symbole*. 79–83, 91, 106–108, 128
- PSK** *phase-shift keying*. 72
- QAM** *quadrature amplitude modulation*. 71
- ReLU** *rectified linear unit*. 54
- RF** *radio fréquence*. 88, 116, 133, 149, 192
- RN** *réseaux de neurones*. 27, 28, 32, 178
- RNN** *recurrent neural network*. 27
- RSB** *rapport signal sur bruit*. 7, 8, 15, 16, 64, 69, 70, 74–82, 85, 86, 90, 91, 93, 95, 96, 98–100, 102–107, 111, 125–130, 135, 137, 157, 181, 183
- SAF** *slotted amplify-and-forward*. 76
- SRRC** *square root raised cosine*. 10, 113, 191

- SVM** séparateurs à vaste marge. 27
- TCD** transformée en cosinus discrète. 33, 148
- TDOA** *time difference of arrival*. 11, 151–155, 178
- TEB** taux d’erreur binaire. 79
- TES** taux d’erreur symbole. 79, 80
- TFCT** transformée de Fourier à court terme. 31, 33
- TFD** transformée de Fourier discrète. 31, 33
- TFR** transformée de Fourier rapide. 33, 148, 153

INTRODUCTION

1.1 Introduction générale

La classification d'événements sonores est un domaine en pleine expansion. En effet, ces méthodes sont favorisées par des avantages certains. D'une part, l'enregistrement de signaux audio est moins cher et perçu comme moins intrusif que l'enregistrement vidéo [94] ; le traitement de ce signal est souvent plus simple et plus efficace dans certaines situations, il permet donc de réagir rapidement. D'autre part, la baisse des coûts des microphones et leur présence dans de nombreux appareils qui ont intégré notre vie courante – comme les téléphones mobiles, les ordinateurs, ou encore les enceintes ou montres connectées – grâce à l'essor de l'internet des objets (IoT), permettent de mettre en place facilement et à faible coût des systèmes de détection d'événements sonores [190]. Ainsi, un tel système apparaît comme idéal pour de nombreuses applications allant du domaine du médical (détection d'appel à l'aide) [193] à la sécurité (détection d'intrusion) [194] en passant par des usages aussi variés que la détection de cris d'animaux [180], d'environnements sonores [200] ou de présences.

L'installation de tels systèmes requiert cependant de nombreuses connaissances à l'avancement desquelles la recherche participe activement. En premier plan, le monde de la recherche s'intéresse beaucoup à l'étude des algorithmes de détection d'événements sonores, que ce soit sur l'aspect des algorithmes de classification ou d'extraction des descripteurs. La popularisation des réseaux de neurones a notamment permis des avancées majeures en terme de performance de classification et d'étiquetage des données utilisées [196]. Le développement de la recherche reproductible et le partage des bases de données et algorithmes a également participé à ces progressions. Le challenge *Detection and Classification of Acoustic Sound Events* (DCASE), lancé pour la première fois en 2013 et qui a lieu tous les ans depuis 2016, est un acteur majeur de la recherche mondiale dans ce domaine. Il donne chaque année l'impulsion pour la popularisation de nouveaux axes à travers différentes tâches, qu'elles soient applicatives (détection de chant d'oiseaux, de bruits intérieurs ou extérieurs...) ou orientées vers de nouveaux systèmes (détection mono ou multi-capteurs, détection associée à la localisation des sources...) [47, 74].

D'autre part, l'IoT requiert de nombreuses transmissions sans fil, opérations parmi les plus coûteuses dans les systèmes à faible consommation d'énergie [72]. Or, pour des considérations écologiques et financières, il est nécessaire de minimiser la consommation d'énergie des communications sans fil afin de prolonger la durée de vie sans recharge des nœuds, tout en maintenant un débit binaire assez haut pour ne pas entacher leur performance. Dans le domaine des réseaux de capteurs sans fil, un tel compromis représente une importante problématique [108] ; en effet, l'énergie d'une transmission augmentant exponentiellement avec la distance, il semble intéressant d'utiliser des nœuds intermédiaires, appelés nœuds-relais, qui peuvent retransmettre l'information, ou parfois diffuser une information complémentaire [46]. De nombreux travaux s'intéressent alors à la localisation de ces nœuds afin de connaître les caractéristiques liées à la position du capteur, et éventuellement choisir le placement ou le nœud le plus avantageux

[222].

Si ces deux domaines de recherches sont prolifiques, peu de travaux s’y intéressent conjointement alors qu’ils sont intrinsèquement liés dans un système réaliste. En effet, si l’on souhaite conserver les bonnes performances obtenues en classification à l’aide de données en précision infinie, il faut mobiliser d’importantes ressources afin de transmettre le plus fidèlement possible ces données. En revanche, si l’on souhaite travailler dans un système contraint en énergie, avec des nœuds à faible consommation, il faudra réaliser un compromis entre la précision des données transmises et la consommation d’énergie.

1.2 Motivations et problématiques

La problématique générale de ces travaux consiste à comprendre comment réduire la consommation énergétique d’un dispositif de détection d’événements sonores distribué tout en maintenant les performances de cette détection à des niveaux satisfaisants. Le but est de pouvoir implémenter un tel dispositif sur des réseaux de capteurs sans fil à faible consommation tels qu’employés pour les applications de l’IoT. S’agissant d’une thématique récente et encore peu étudiée, nos travaux visent à explorer les contraintes et possibilités intrinsèques à ce type de systèmes embarqués. Nous nous placerons pour cela dans le contexte de nœuds de capteurs transmettant des descripteurs audio à un moniteur central réalisant la détection d’événements sonores via un canal de communication sans fil.

Nos travaux s’ouvrent ainsi sur deux axes de recherche.

La consommation d’une communication sans fil Afin de comprendre comment diminuer la consommation d’une transmission en limitant l’altération du signal et le débit de la communication, nous nous intéressons au compromis entre les efficacités spectrale et énergétique. Ces deux critères mesurant la qualité du service sont en effet intrinsèquement liés. Afin d’évaluer au mieux ce compromis, nous cherchons à modéliser la consommation de la communication de la manière la plus réaliste possible. L’utilisation d’un nœud-relai et son positionnement dans l’espace afin de réduire la consommation sont également des sujets que nous souhaitons interroger.

La détection d’événements sonores La détection d’événements sonores, comme tout traitement du signal, nous confronte à des obstacles de diverses natures que nous souhaitons étudier lors d’une implémentation matérielle. L’un d’entre eux est la difficulté, voire l’impossibilité de disposer en contexte réel de données encodées en virgule flottante comme cela est habituellement considéré lors des tests sur ordinateur. La quantification en virgule fixe qui s’impose entraîne alors des imprécisions sur le signal. Dans le cadre de nos travaux, nous aspirons à comprendre comment la quantification puis la transmission de descripteurs d’un nœud de capteurs à un moniteur central chargé de la classification peut influencer les performances de celle-ci et, le cas échéant, comment y remédier. Cependant, une implémentation matérielle n’apporte pas forcément que du négatif, et nous nous interrogeons également sur les apports que peut amener la combinaison des signaux de différents nœuds multi-capteurs.

1.3 Contributions de nos travaux

Afin de répondre aux problématiques énoncées dans la section précédente, nous avons mené des travaux qui nous ont conduit à apporter plusieurs contributions, tant dans le domaine de la consommation d'énergie des réseaux de capteurs sans fil que de la détection d'événements sonores. Ces contributions sont récapitulées sur la figure 1.1 pour chacun de ces champs de recherche, et détaillées ci-dessous.

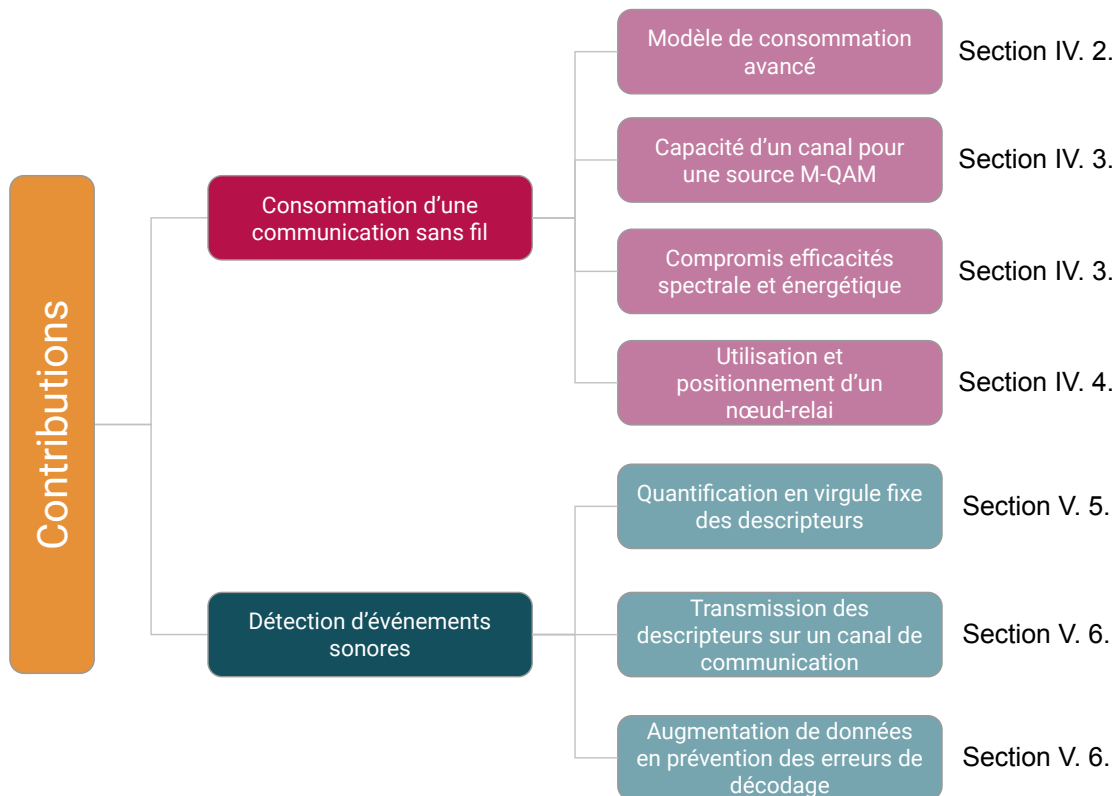


FIGURE 1.1 – Visualisation des contributions apportées dans cette thèse

La consommation d'une communication sans fil

1. Dans la plupart des études, la consommation de l'amplificateur de puissance (qui dépend de l'efficacité du drain et du PAPR) et celles du CNA et du CAN sont souvent simplifiées à l'aide de coefficients constants, alors que différentes études prouvent que ces consommations dépendent en réalité de paramètres du système. Si des modélisations avancées pour chacune de ces composantes existent, elles sont généralement étudiées indépendamment les unes des autres. Nous proposons dans cette thèse de joindre l'ensemble des connaissances avancées sur l'efficacité du drain, le PAPR et les puissances dissipées par les CNA et CAN pour réaliser une analyse inédite de la consommation de la transmission d'une source modulée par une modulation d'amplitude en quadrature (MAQ) (ou M-QAM) sur un canal AWGN ou de Rayleigh en fonction de l'ES. Cette contribution est détaillée dans la section 4.2.

2. L'efficacité spectrale (ES) est en lien direct avec la capacité d'un canal, qui elle-même est une fonction du RSB en réception. Dans nos travaux, nous proposons de compléter l'état de l'art en exprimant la capacité d'un canal d'une source gaussienne sur un canal de Rayleigh sous une forme davantage conforme aux simulations numériques. Nous exprimons ensuite le RSB en réception en fonction de cette capacité. Pour les sources modulées par M-QAM que nous considérons, des recherches indiquent qu'il est possible de déterminer une pénalité de ce RSB par rapport à des sources gaussiennes, et que cette pénalité peut être exprimée en fonction de la seule PEB du signal reçu. Nous proposons alors de nouvelles expressions de la pénalité en fonction de la PEB pour des canaux AWGN et de Rayleigh. En particulier, la pénalité que nous proposons pour les canaux de Rayleigh est, à notre connaissance, la première expression analytique exprimée en fonction de la PEB. Ces apports sont présentés dans la section 4.3.
3. Les modèles de consommation d'énergie exprimant le RSB à partir de la fonction réciproque de la PEB ne sont pas intéressants pour calculer la valeur de l'ES optimale minimisant cette énergie, car le calcul de cette réciproque mène à des expressions difficiles à manipuler. En travaillant avec la notion de capacité et de pénalité du RSB par rapport à une source gaussienne – comme décrit précédemment –, il est possible d'analyser simplement l'ES optimale. Nous présentons ainsi une analyse poussée de l'ES optimale à partir du modèle de consommation proposé. Les résultats de cette étude sont également à retrouver dans la section 4.3.
4. Enfin, nous nous intéressons à la question de l'utilisation du nœud-relai dans la section 4.4. Nous étudions la consommation d'une transmission coopérative avec le protocole AF en fonction de sa position dans un espace à une puis deux dimensions. Nous analysons notamment dans quels cas l'utilisation d'un relai est avantageuse énergétiquement par rapport à son absence, et comment évolue son ES en fonction de sa position.

Ces travaux ont mené aux publications suivantes :

- M. Lacroix, R. Rocher et P. Scalart, "Realistic power amplifier model for energy optimization in wireless networks", *Inter. Symp. on Wireless Comm. Systems (ISWCS)*, 2021.
- M. Lacroix, R. Rocher et P. Scalart, "Advanced energy model and spectral efficiency optimization in short-range wireless sensor networks", *Inter. Conf. on Green Computing and Comm. (GreenCom)*, 2021.
- M. Lacroix, R. Rocher et P. Scalart, "Realistic spectral efficiency analysis for energy-efficient wireless communications", à soumettre.
- M. Lacroix, R. Rocher et P. Scalart, "Spectral and energy efficiencies tradeoff analysis from an advanced energy model in wireless communication", à soumettre.

La détection d'événements sonores

5. Notre première contribution au champ du traitement audio est l'étude des possibilités de quantification en virgule fixe d'un vecteur de descripteurs. Nous proposons alors une analyse inédite de l'influence de ce codage et de la longueur de quantification choisie sur les performances d'un GMM soustractif et d'un CRNN. Les résultats de ces travaux, réalisés à partir d'apprentissages sur des données quantifiées ou non, sont exposés dans la section 5.5.3.

6. La quantification des descripteurs permet ensuite de les transmettre à faible coût d'un capteur-source vers un moniteur central via une communication sans fil. Cependant, une telle transmission peut être à l'origine d'erreurs qui amènent à décoder les mauvais descripteurs en réception. Nous proposons d'analyser l'impact de ces erreurs sur la détection d'événements sonores pour plusieurs longueurs de quantification et différentes PEB en réception. On présente nos conclusions dans la section 5.6.3.
7. Nos analyses montrent que la transmission des descripteurs sur un canal de communication sans fil sont à l'origine d'une chute des performances pour le CRNN. Les erreurs de décodage agissent en effet sur les descripteurs comme un bruit qui va perturber la détection des événements. En se basant sur plusieurs études qui insistent sur l'importance de prendre en compte le bruit et toute irrégularité des signaux réels dans un algorithme de classification, nous proposons une nouvelle méthode d'augmentation de données basée sur un ajout artificiel d'erreurs de décodage pour prévenir l'impact de celles-ci sur la détection d'événements sonores. Cette dernière contribution est détaillée dans la section 5.6.3.

1.4 Structure de la thèse

Cette thèse est structurée autour de quatre chapitres scientifiques qu'encadrent le présent chapitre d'introduction ainsi qu'un chapitre de conclusion générale. Le plan de la thèse et la structuration des chapitres par rapport à nos deux axes de recherche est récapitulé sur la figure 1.2.

- Le **chapitre 2** s'axe sur la découverte d'un système de détection d'événements sonores. Il introduit les notions essentielles pour comprendre les enjeux de la classification acoustique, et passe en revue les outils nécessaires à la mise en place d'un tel dispositif, tels que les BDD, les descripteurs ou les critères de performance. Ce chapitre est également l'occasion de présenter et d'étudier les classifieurs de l'état de l'art qui nous serviront dans la suite de nos travaux, à savoir un GMM soustractif et un CRNN.
- Le **chapitre 3** poursuit notre état de l'art, cette fois-ci dans l'optique de présenter le concept d'une transmission sans fil, qu'elle soit coopérative (avec un nœud-relai) ou non. Nous introduisons alors des notions telles que la modulation du canal, les atténuations du canal physique, les protocoles coopératifs ou encore la combinaison de signaux en réception. Nous poursuivons en nous intéressant à la qualité de service qui peut être fixée par la probabilité d'erreur en réception, la capacité ou encore les efficacités du canal. Ces différentes notions nous permettent de présenter le cœur de nos travaux en communication, à savoir la consommation d'une transmission point-à-point ou coopérative, et son lien avec l'ES.
- Ces études sont prolongées par le **chapitre 4**, qui présente nos différentes contributions dans le domaine des communications sans fil. Une analyse de plusieurs modèles de consommation est réalisée, suivie par la proposition d'une synthèse de ceux-ci en un unique modèle. La capacité d'un canal AWGN ou de Rayleigh est ensuite étudiée afin de proposer une analyse réaliste de l'ES minimisant la consommation d'énergie d'une transmission point-à-point. Notre propos est complété par une étude de la consommation d'une communication coopérative employant un protocole AF en fonction de la position du nœud-relai dans un espace à une ou deux dimensions. Nous discutons également de l'ES dans cet espace et de l'intérêt énergétique du relai dans différents contextes.

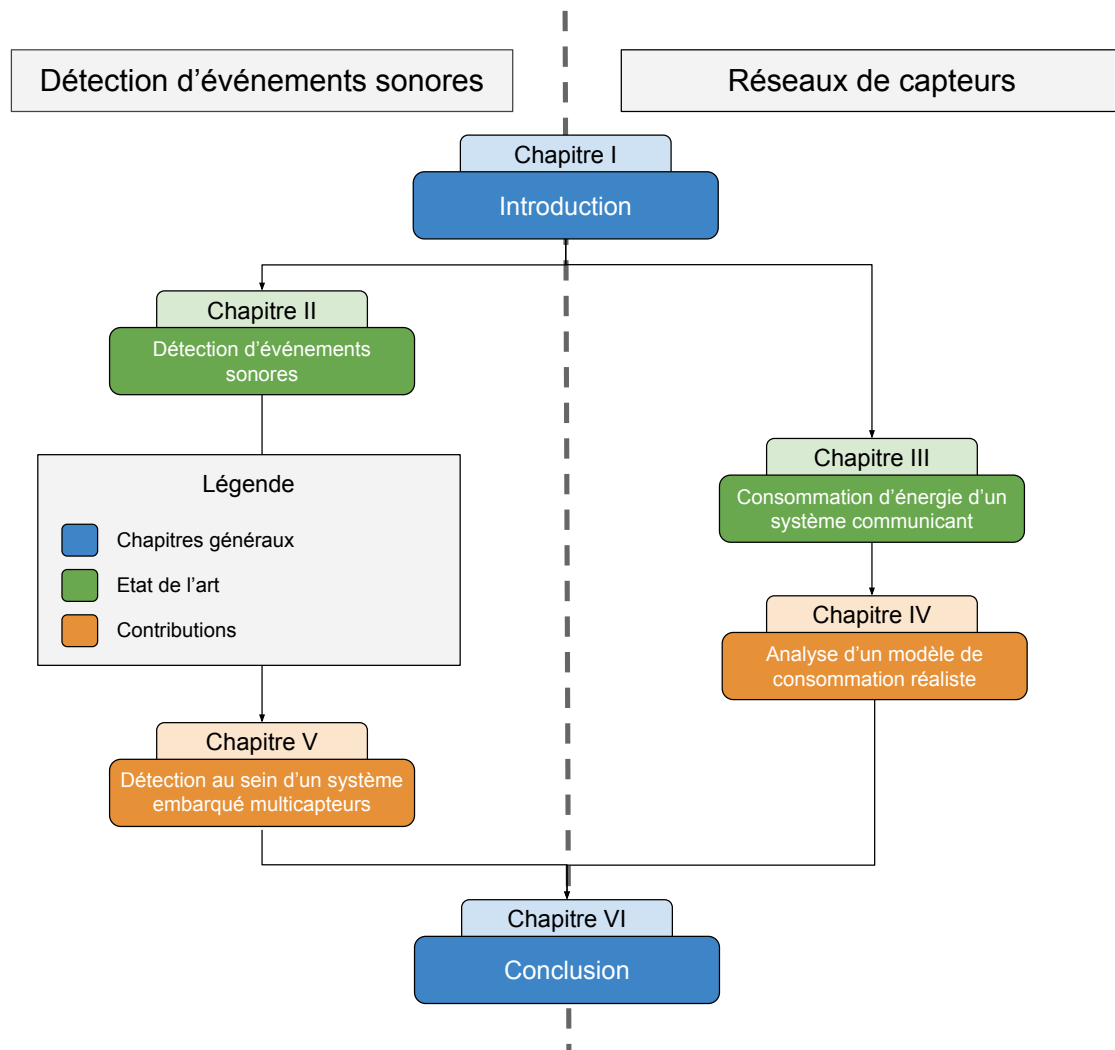


FIGURE 1.2 – Schématisation de la structure de la thèse

- Finalement, le **chapitre 5** propose une étude applicative de l'implémentation d'un dispositif de détection d'événements sonores sur un réseau de capteurs sans fil. Après avoir présenté le contexte général, nous étudions les possibilités de valoriser la présence de plusieurs capteurs via une spatialisation de la classification acoustique. Nous nous intéressons ensuite à la notion d'augmentation de données, souvent essentielle dans les systèmes appliqués pour lesquels il est difficile de récolter suffisamment de données, et expérimentons quelques méthodes. La dernière partie du chapitre porte sur l'analyse de l'influence sur la détection d'événements sonores de la quantification puis de l'émission des descripteurs sur un canal de communication sans fil. Des solutions à la dégradation des performances sont apportées via un apprentissage réalisé sur données quantifiées et une technique originale d'augmentation des données est proposée pour pallier ce problème.

DÉTECTION D'ÉVÉNEMENTS SONORES

2.1 Introduction

Le cerveau humain est capable d'interpréter une multitude de sons différents. Cette capacité n'est pas innée, elle se développe tout au long de la vie par apprentissage. La signification attachée à ces sons peut être abstraite (comme certains mots) ou plus tangible (comme le chant d'un oiseau, des verbes d'actions, etc). On dissocie généralement la nature des sons selon différentes catégories qui ont leurs caractéristiques et leur propre domaine de recherche : la musique, la parole, et les sons environnementaux. L'apprentissage de ces sons, quelle que soit leur catégorie, est généralement dit multimodal, c'est-à-dire qu'il ne s'appuie pas uniquement sur le signal sonore lui-même : le son est associé à un concept visuel ou une sensation (tondeuse à gazon utilisée par le voisin, vibration du moteur en fonction de l'appui sur la pédale de débrayage...) [107]. Une fois mémorisé, ce son n'a cependant pas besoin d'être associé à une autre source pour être reconnu : on peut détecter une tondeuse sans regarder par la fenêtre et connaître le comptage tour-minute de sa voiture à l'oreille. Le processus d'apprentissage permet en réalité d'associer un signal audio à ce qu'il représente, mais la reconnaissance de sa source ne nécessite pas d'avoir plus que celui-ci. Par exemple, beaucoup de gens savent reconnaître la sonnerie d'une cloche d'église sans avoir jamais vu la cloche se mouvoir. Ces éléments extérieurs participent alors activement à l'étiquetage des événements acoustiques par notre cerveau. La classification supervisée d'événements sonores tente de reproduire ce phénomène en déterminant un modèle statistique capable d'apprendre, à partir de signaux acoustiques et d'un ensemble d'étiquettes, à relier une portion de signal à un événement.

La classification d'événements sonores est un domaine de recherche en pleine expansion, notamment de par ses nombreux domaines d'application et le développement des applications de traitement du signal pour un large public. Elle est par exemple très répandue dans la surveillance acoustique [195]. Cela peut être installé dans un milieu médical, pour des systèmes d'aide au maintien de l'autonomie chez les personnes âgées [193], aux alentours de bâtiments pour des systèmes de surveillance (détection d'intrusion, de bris de glace, de coups de feu...) [194], ou encore dans le cadre d'application d'économie d'énergie pour piloter l'extinction des lumières ou des appareils de chauffage. Elle peut également être un soutien à certaines équipes de recherche, par exemple en ornithologie pour différencier les espèces d'oiseaux [180]. Avec le développement des objets connectés et la miniaturisation des microphones, des applications de classification peuvent également être développées pour l'assistance multimédia (extinction du microphone lors d'une téléconférence quand il n'y a pas d'orateur, sous-titrage automatique pour malentendants) ou sensibiliser les objets à leur environnement (téléphone dont la sonnerie se coupe automatiquement dans un restaurant, robotique) [200]. Chaque application nécessite une base de données appropriée, mais aussi le bon modèle de classification muni des critères de performance correspondants à la tâche visée.

Dans ce chapitre, nous commencerons par détailler le principe de la classification d'événe-

ments sonores, en expliquant à la fois les différents types et étapes de classification, les caractéristiques du signal sonore ou encore la nature des classifieurs. Par la suite, nous établirons un état de l'art des différents outils utilisés, tels que les bases de données et leurs caractéristiques, les critères de performance avec leurs spécificités et le compromis à réaliser entre eux, et les descripteurs du signal audio. Enfin, on s'intéressera au fonctionnement de deux classifieurs phares de l'état de l'art en détection d'événements sonores : le GMM et le CRNN.

2.2 Généralités sur la classification acoustique

La classification de signaux audio environnementaux est un domaine qui a connu de récentes avancées grâce à la popularisation des méthodes d'apprentissage automatique, et en particulier des réseaux de neurones. À ce titre, il s'agit d'un domaine qui regroupe un grand nombre de sous-domaines ayant chacun ses propres méthodes. Il s'apparente également à d'autres sujets de recherche tels que la recherche d'information musicale ou la reconnaissance de la parole. Néanmoins, la classification d'événements a ses propres caractéristiques sonores et ses propres méthodes de classification.

Dans cette section, nous allons présenter l'ensemble des éléments constitutifs de la classification d'événements sonores. Dans un premier temps, on introduira les différents types de classification existants, puis on expliquera les étapes d'un tel système. Dans un deuxième temps, on donnera les caractéristiques des données audio utilisées. Enfin, on listera les différents classifieurs utilisés dans la littérature et on justifiera nos choix quant aux classifieurs de référence.

2.2.1 Les types de la classification

La classification est le terme général utilisé pour définir les méthodes ayant pour but de regrouper des échantillons manifestant une proximité acoustique sous l'égide d'une même classe supposée homogène. Pour déterminer quelles méthodes utiliser, il est nécessaire de répondre à plusieurs questions qui permettent de dégager les objectifs du dispositif de classification.

Quel type d'apprentissage ?

Pour réaliser une classification acoustique, on entraîne un modèle statistique à reconnaître les caractéristiques qui différencient ou au contraire rassemblent ces échantillons selon s'ils appartiennent ou non à un même groupe. Ces échantillons peuvent être "annotés" – ou "étiquetés" –, c'est-à-dire que l'on sait à quelle classe chacun appartient. On connaît donc aussi le nombre total de classes. La classification supervisée consiste à comparer les annotations avec la sortie du modèle statistique afin d'évaluer ou de modifier ce dernier. Quand on ne dispose ni du nombre de classes, ni des étiquettes des échantillons, on effectue alors une classification dite non supervisée. Aussi appelée regroupement ou partitionnement, elle consiste à regrouper ensemble des échantillons de caractéristiques similaires et à les séparer au mieux des autres, sans information supplémentaire. En audio, en raison de la complexité de la tâche, la classification non supervisée n'est utilisée que pour certaines tâches spécifiques avec très peu de classes, comme la détection d'anomalies [10, 43, 95], ou en temps que pré-entraînement sur des bases de données (BDD) généralistes [76]. En général, on lui préférera la classification semi-supervisée. À la croisée entre supervisé et non-supervisé, celle-ci consiste à entraîner le modèle à la fois sur des données étiquetées et non étiquetées. Dans les méthodes les plus classiques, les échantillons

supervisés sont utilisés pour modéliser une première règle de classification, et les échantillons non supervisés pour affiner cette règle [38]. La classification semi-supervisée représente un réel avantage lorsqu'elle permet d'éviter un étiquetage long et coûteux des données [38], ce qui est justement le cas pour la classification acoustique, comme nous en reparlerons dans la section 5.4.1 [196]. Dans cette thèse, nous considérerons des bases de données munies d'annotations suffisamment fiables et quantitatives pour permettre une classification supervisée.

Quel type de sons ?

La classification sonore environnementale vise à classer deux types d'éléments : les événements d'une part, les scènes d'autre part. Un événement sonore renvoie à un son spécifique généralement produit par une seule source. Il peut s'agir du chant d'un oiseau, du cri d'un enfant ou encore du klaxon d'une voiture [196]. Par leur nombre et leur variabilité infinie, des événements ou des sources sont intrinsèquement des éléments subjectifs. Par exemple, doit-on considérer une machine à laver en fonctionnement comme un bruit non utile ou un événement ? Est-ce le moteur ou la voiture qui est la source du vrombissement ? Une chaîne hifi avec des enceintes en stéréo constitue-t-elle une ou plusieurs sources sonores ? Un événement doit être défini avec attention en fonction de l'application visée. Une scène sonore, quant à elle, est un concept plus large constitué d'une multitude d'événements individuels. Par exemple, une rue peut être caractérisée par des piétons marchant et parlant, des voitures circulant ou encore des oiseaux piaillant ; un environnement de bureau peut lui aussi contenir des personnes marchant et parlant, mais surtout des bruits de chaises, de clavier ou encore de photocopieuse. Certaines caractéristiques acoustiques pourraient également être révélées si la scène se déroule en intérieur ou en extérieur, comme le niveau sonore ou la réponse impulsionnelle. Dans nos travaux, nous nous intéresserons uniquement à la classification d'événements sonores, et non à des scènes.

Quel type d'événements ?

Finalement, en analyse d'événements sonores, plusieurs types d'applications se détachent en fonction des objectifs visés. Le signal audio est généralement découpé en trames de très courtes durées pour être traité numériquement. Ces trames sont censées être plus courtes que les événements étudiés pour permettre d'une part pour traiter en temps-réel le signal audio, d'autre part de moyenniser la sortie du classifieur sur plusieurs trames, comme on le détaillera dans la section suivante [196]. Les décisions en sortie du classifieur peuvent ainsi avoir plusieurs formes. Par exemple, la sortie peut être présentée sous la forme d'une séquence temporelle. On a alors un début et une fin approximative de chaque événement, déterminé à partir des probabilités de présence sur chacune des trames. Un tel procédé est appelé détection d'événements sonores, et peut traiter des événements simultanés ou non. On peut également choisir de traiter l'information à une échelle temporelle plus large, et de découper nos signaux sous la forme de segments, voire d'une piste complète. Chaque segment est alors traité indépendamment des autres, et la décision est effectuée sur l'ensemble des trames de ces segments. La présence d'un événement sur celui-ci nous indique son apparition sur une partie, mais pas forcément sur l'ensemble du segment temporel. Dans le cas où un événement au plus a lieu sur chaque segment, on parle de classification mono-label ou monophonique, parfois appelée simplement classification. Si plusieurs événements peuvent survenir, il s'agit de classification multi-label ou polyphonique, plus communément appelée étiquetage audio (*audio tagging*) [196]. L'ensemble de ces méthodes, qui sont souvent regroupées sous l'appellation de classification sonore, est récapitulé sur la figure

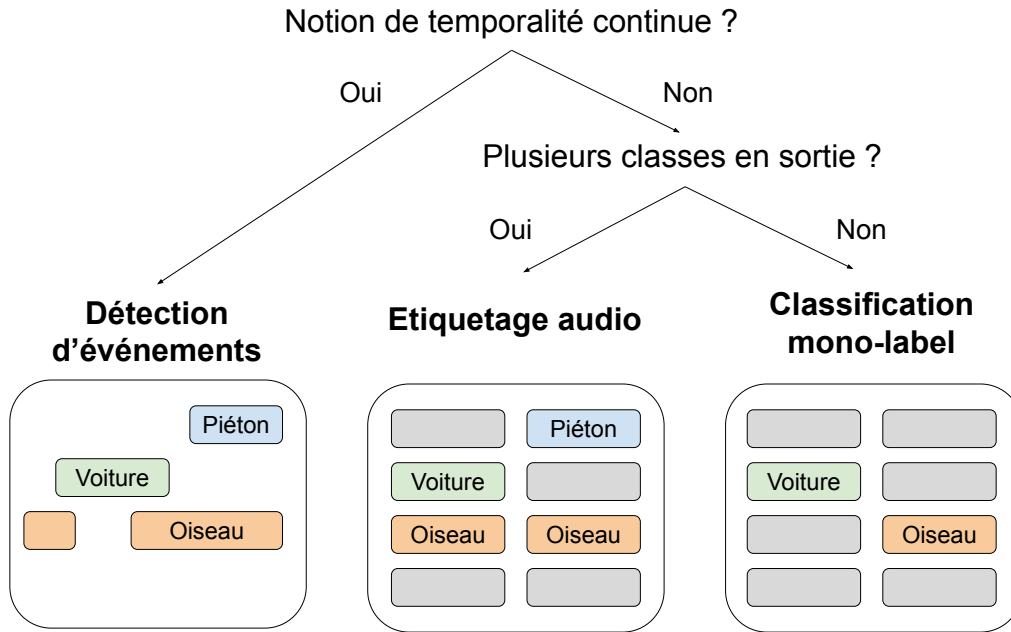


FIGURE 2.1 – Schéma récapitulatif des différentes méthodes de classification acoustique

2.1. Dans ces travaux, nous souhaitons étudier un dispositif de classification sonore qui puisse réagir en temps réel aux événements sonores, ce qui nous conduit à choisir de traiter le cas de la détection d'événements.

2.2.2 Les étapes d'une détection d'événements acoustiques

La détection d'événements acoustiques se déroule en plusieurs étapes indiquées sur la figure 2.2, que l'on va détailler.

Pré-traitement

Le signal reçu en entrée du système est souvent brut. On peut être amené à vouloir l'améliorer, l'homogénéiser par rapport à d'autres signaux ou encore à exacerber certaines de ses caractéristiques. L'un des traitements récurrent en traitement du signal est le débruitage, qui enlève les interférences et autres parties non utiles du signal. Néanmoins, ce traitement est assez dangereux car la notion bruit est subjective et, dans ce cas précis, il peut faire partie des événements sonores à reconnaître. Parmi les autres traitements, on peut également citer l'annulation d'écho. La séparation de sources permet quant à elle de traiter les sources séparément, ce qui peut par la suite faciliter la tâche du classifieur. Il est également possible de normaliser le niveau ou les fréquences du signal dans le cas d'enregistrements provenant de différents microphones, en particulier car les descripteurs peuvent dépendre de ces niveaux [196]. À noter que le pré-traitement n'a pas forcément à être identique pour les échantillons de test et d'apprentissage, par exemple si les deux n'ont pas été enregistrés dans les mêmes conditions.

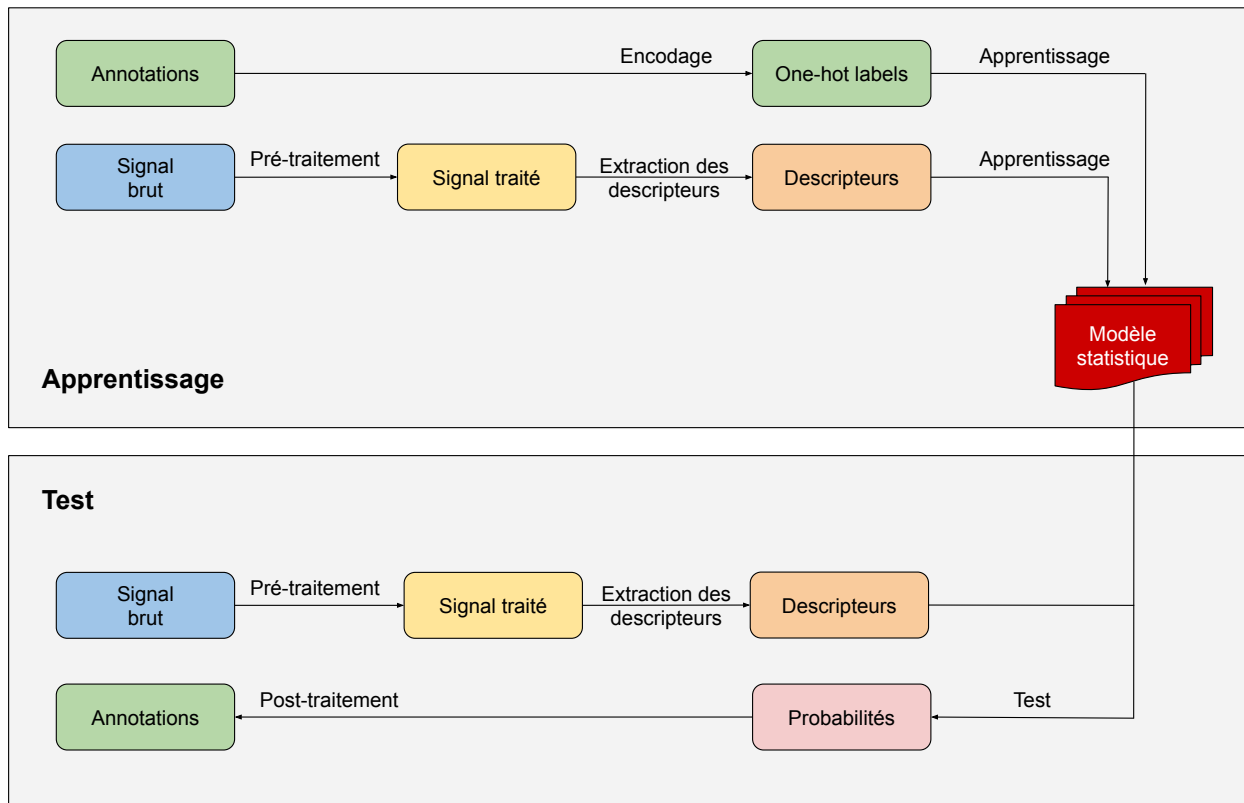


FIGURE 2.2 – Schéma récapitulatif des différentes étapes de la classification d'événements sonores

Extraction des descripteurs

Le signal en sortie du pré-traitement n'est pas donné sous une forme efficace pour la classification : les échantillons sont nombreux et l'information est redondante et peu efficace (deux échantillons peuvent avoir des caractéristiques temporelles similaires sans être de la même classe). On va donc en extraire des descripteurs regroupant le plus d'informations pertinentes pour réduire la complexité du traitement et améliorer les performances de la classification [196].

Le processus d'extraction consiste à découper le signal en une suite de trames fenêtrées de longueur T_t (en secondes), avec éventuellement un recouvrement entre les trames. Ces trames subissent alors une transformation, dont la représentation la plus courante est la représentation spectrale. Cette représentation n'est pas effectuée à l'échelle d'une piste audio entière, mais à très courte échelle sur des périodes de quasi-stationnarité du signal audio [5]. Afin de travailler sur des périodes de quasi-stationnarité, on choisit alors généralement une longueur de trame comprise entre 20 et 50ms [196].

La notion de descripteurs sera abordée plus amplement dans la section 2.3.2.

Encodage des annotations

Le fichier d'annotations peut se trouver sous diverses formes. En détection, il s'agit souvent d'un fichier de texte indiquant le temps de début et de fin de chaque événement [60, 129]. Ce genre de fichiers ne peut être traité directement par le classifieur, il faut donc un format

commun. La forme la plus répandue est l'encodage 1 parmi n . Cela consiste à remplir une matrice de taille $n \times L_T$ avec n le nombre des classes et L_T le nombre de trames du signal. L'encodage 1 parmi n a initialement été prévu pour la classification mono-label, dans laquelle un unique événement est présent à chaque trame. Celui-ci est indiqué par une valeur unitaire au sein du vecteur de taille n associé à la trame étudiée, tandis que les valeurs correspondant aux autres classes sont nulles. Par extension, pour la classification multi-label, on remplit chaque case (i, j) de la matrice en indiquant 1 si l'événement numéro i est présent à la trame j , et 0 sinon.

Apprentissage

L'apprentissage consiste à chercher le modèle statistique qui représente le mieux les données d'apprentissage. Lors de la construction d'un modèle, on sépare généralement les données en deux paquets : les données d'apprentissage d'une part, qui vont servir à entraîner le modèle et atteindre la convergence, et les données de validation, qui permettent de vérifier les performances. Le modèle se construit par évaluation statistique des données et en comparant les étiquettes d'origine avec la décision en sortie du classifieur. Afin de retirer un maximum d'information de l'ensemble des données, celles-ci traversent plusieurs fois le système ; chaque passage est appelé une époque.

Le fonctionnement précis de l'entraînement dépend du classifieur choisi. Dans certaines applications qui conduisent à une variabilité importante des paramètres statistiques du classifieur, il est possible de réaliser un apprentissage à plusieurs plis (appelés communément *fold*). Le principe est de réaliser l'apprentissage de plusieurs modèles à partir des mêmes paramètres initiaux mais en changeant la base de validation, et de sélectionner le meilleur modèle parmi tous. On évite ainsi d'entraîner le classifieur sur un unique pli pour lequel les échantillons les plus représentatifs feraient partie de la base de validation, donc l'apprentissage serait moins significatif [196].

Test

Une fois le modèle construit et validé, il est possible de l'utiliser avec les descripteurs de données de test, de préférence jamais vues par le modèle, afin d'obtenir en sortie un vecteur exprimant la probabilité de présence d'un événement de chacune des classes au sein de la trame d'analyse. Après un éventuel post-traitement, il est possible de binariser ces probabilités selon un seuil fixé pour connaître les étiquettes finales. Ces étiquettes pourront par la suite être comparées à la matrice des annotations pour connaître les performances du classifieur.

À noter qu'il s'agit là du même traitement que celui subi par les données de validation lors de la phase d'apprentissage.

Post-traitement

L'étape de reconnaissance détermine les probabilités de présence des événements à chaque trame. Chacune d'entre elle a une longueur très courte, bien plus que les événements analysés. En conséquence, un événement couvre plusieurs trames [196]. On peut se servir de cette information pour lisser les probabilités ou les décisions en sortie du classifieur à l'aide d'un algorithme de lissage. En effet, un événement présent sur une seule trame constitue probablement une erreur. On choisit alors une fenêtre de trames consécutives, de longueur cependant

inférieure à celle d'un événement, pour lisser les décisions [127, 196].

2.2.3 Les propriétés des données audio

Il n'existe pas une, mais des méthodes de classification. Chaque méthode répond à un contexte matériel et sonore, une application, un but spécifique. Malgré un bagage commun d'outils et d'algorithmes, il est important de définir en amont les caractéristiques du système, et en particulier ceux liés au signal sonore, qui constitue la base des systèmes de classification d'événements acoustiques.

Les propriétés du son

Les méthodes de classification sont généralement déterminées à partir des objectifs applicatifs que l'on se donne. Dans la classification d'événements acoustiques, ces objectifs sont eux-mêmes définis par l'environnement sonore que l'on cherche à comprendre. On se propose donc de lister les différentes caractéristiques particulières des signaux audio.

Dans un premier temps, il est important de connaître la nature du son qui va être enregistré. Par exemple, la durée du signal émis par les sources, selon qu'il soit long ou au contraire très bref, peut amener à adapter les étapes du post-traitement. C'est le cas également en cas d'événement rare ; une alarme incendie, par exemple, se déclenche rarement mais est importante à détecter [69].

Au-delà du son lui-même, certaines informations sur les classes utilisées sont nécessaires. En effet, un grand nombre de classes ou des classes très semblables (comme le passage d'un camion et d'une voiture) peuvent complexifier leur différenciation. Il s'agit donc de faire attention dans le choix des descripteurs, dans le pré-traitement éventuel, mais aussi dans le choix du classifieur [196]. Un élément crucial, comme nous l'avons vu dans la section 2.2.1, est la possibilité pour les événements de se produire en simultané ou pas. En effet, on peut entendre un chien aboyer et un chat miauler en même temps, mais on ne peut pas avoir l'ambiance sonore de l'intérieur d'un train et celui d'une forêt. Dans le premier cas, il s'agit d'une classification multi-label ; dans le second, d'une classification mono-label. Ceci a bien évidemment un impact sur le choix du classifieur, mais aussi sur le post-traitement [196].

Enfin, il est souvent utile de connaître les dispositions de l'enregistrement. Certains signaux peuvent être multi-canaux et, selon que les signaux des différents canaux proviennent ou non du même nœud de capteurs, des traitements particuliers peuvent être appliqués. Il peut s'agir de descripteurs spatiaux ou encore d'une classification de chaque canal en parallèle [4, 226]. On reviendra sur ces éléments dans la section 5.3.2. En outre, les signaux peuvent provenir d'un enregistrement réel, ou être le fruit d'un travail synthétique sur différentes pistes sonores. Cela peut jouer sur les seuils de détection en fin de classification, ou sur la classification elle-même [196].

Un point sur les annotations

L'annotation des signaux audio joue un rôle essentiel dans la détermination d'un système de classification, car elle permet de savoir à quelle classe appartient chaque échantillon audio. Elle est donc nécessaire en classification supervisée pour l'apprentissage, mais également en non supervisée pour la phase de test. Il arrive également de ne posséder qu'une partie des annotations ; on peut alors se tourner vers des techniques semi-supervisées. Par ailleurs, leur

précision peut différer. Sur un signal audio, on peut indiquer précisément le début et la fin de l'événement : il s'agit d'un étiquetage fort (*strong labelling*). On peut aussi se contenter d'indiquer la présence d'une classe sur tout ou partie d'un segment ou piste audio : il s'agit d'un étiquetage faible (*weak labelling*) [98, 196].

De telles annotations étant difficiles à obtenir automatiquement, on fait généralement appel à l'oreille humaine. Des annotateurs écoutent alors plusieurs fois les enregistrements audio et les étiquettent approximativement, librement [128] ou à partir d'une liste d'événements donnés [46]. Néanmoins, une telle annotation est très coûteuse (en temps, en argent...), d'autant plus que les classifieurs à base de réseaux de neurones, qui sont majoritaires aujourd'hui, requièrent un grand nombre de données. La fatigue et l'effort de concentration que demande l'exercice d'annotation peuvent également amener à des annotations temporelles peu précises ; en outre, l'étiquetage peut être subjectif – en particulier s'il est libre – ce qui rend le traitement difficile. Des méthodes d'annotation automatiques sont généralement ajoutées pour affiner les résultats, notamment la fin et le début pour les étiquettes fortes [46]. Il existe également d'autres méthodes, qui peuvent être associées aux précédentes. Tout d'abord, à l'écoute pure peuvent s'adjoindre des données multimodales. On peut s'aider par exemple des images vidéo enregistrées en même temps que l'audio, ou avoir quelqu'un qui indique en temps-réel ses activités ou les sons qu'il perçoit [23]. Enfin, il est possible de faire appel au *crowdsourcing*, c'est-à-dire qu'on sollicite un grand nombre d'utilisateurs, surtout sur Internet, pour leur faire étiqueter chacun quelques échantillons [166, 196].

2.2.4 Le choix d'un classifieur

Dans cette section, nous réalisons une vue d'ensemble des méthodes de classification pour la détection d'événements acoustiques. Cet état de l'art se présente sous le prisme d'une vue haut niveau : nous n'entrerons pas dans les détails de leur fonctionnement. En outre, nous nous intéresserons principalement aux méthodes valables pour la classification de sons polyphoniques.

Les catégories de classifieurs

On sépare généralement les classifieurs en deux catégories : les génératifs et les discriminatifs [196].

Les classifieurs génératifs, aussi appelés descriptifs [38], ont pour objectif de modéliser la distribution de chaque classe. Ces méthodes sont les plus intuitives : leur comportement est souvent très simple à comprendre, ainsi que celui des classes. Néanmoins, la modélisation et la discrimination peuvent être difficiles lorsqu'on assiste à une grande variabilité au sein des classes. Parmi ces classifieurs, on peut citer l'estimation par maximum de vraisemblance, le classifieur bayésien naïf, les modèles de Markov cachés, les modèles de mélanges gaussiens ou encore la factorisation de matrices non négatives.

Les classifieurs discriminatifs sont aussi appelés inductifs ou régressifs [38, 196]. Ils consistent cette fois-ci à modéliser non pas les classes elles-mêmes, mais la limite entre ces classes ; une fois le modèle obtenu, il est donc possible de recenser directement les échantillons qui appartiennent ou non à chaque classe. Néanmoins, cela complexifie souvent la classification multi-label, car il faut alors entraîner un classifieur par classe. En outre, le comportement des classifieurs et des classes est difficile à interpréter. Parmi ces classifieurs, on peut citer les séparateurs à vaste marge, les arbres, mais aussi les réseaux de neurones.

Au départ : des méthodes traditionnelles

Jusqu'en 2016, les principaux classifieurs utilisés pour la classification de sons environnementaux sont génératifs [64].

L'une des plus anciennes méthodes utilisées est le classifieur naïf bayésien. Issu de la classification de la parole [228], il consiste à calculer le maximum de vraisemblance *a posteriori* pour chaque classe [196]. S'il est robuste au sur-apprentissage, il donne cependant de mauvaises performances face à d'autres classifieurs [118], en particulier dans le cas d'une BDD déséquilibrée [25]. On lui préfère donc généralement d'autres méthodes.

Le modèle à mélanges gaussiens, ou *Gaussian mixture model* (GMM), est une méthode statistique généralement très appréciée, qui est longtemps restée une référence dans ce domaine [129] car elle est plus robuste au déséquilibre des BDD que le classifieur bayésien naïf. Sa difficulté d'utilisation réside dans le choix de ses paramètres, tels que le nombre de gaussiennes K , la méthode d'initialisation ou encore celle d'estimation des paramètres [196].

Le modèle de Markov cachés, ou *hidden Markov model* (HMM), également issu de la parole [31], modélise la dynamique temporelle des sons. Néanmoins, cette méthode est rarement utilisée pour elle-même, mais plutôt en post-traitement et généralement associée avec des GMM [154, 196].

Enfin, la factorisation de matrices non négatives, ou *non-negative matrix factorization* (NMF), est une méthode qui est apparue plus récemment. Initialement utilisée pour la séparation de sources, elle a été intégrée au domaine de la classification en tant que classifieur par dictionnaire [86]. Néanmoins, il est apparu qu'elle ne présentait pas de très bons résultats pour la classification polyphonique [31], elle est donc aujourd'hui plutôt utilisée comme post-traitement ou pour l'extraction de descripteurs [21].

Quelques méthodes régressives ont également fait leurs preuves. En classification multi-label, la principale est celle des séparateurs à vaste marge (SVM). Simples à mettre en place, ils discriminent efficacement les classes les unes des autres en modélisant leurs frontières [38]. Pour de la classification multi-label, il faut cependant entraîner un système SVM par classe [196].

Une autre méthode utilisée est celle des forêts aléatoires. Alors qu'une classification par arbre de décision est facilement sujette au sur-apprentissage, la combinaison de la sortie de plusieurs arbres permet d'éviter ce problème [99, 196]. Dans un tel système, chaque arbre est entraîné sur un sous-ensemble des données d'apprentissage et/ou des descripteurs ; chacun peut également avoir sa propre structure (*i.e.* les règles d'élégage ne sont pas les mêmes entre les arbres). Néanmoins, les décisions sont très longues à prédire et le système n'est pas facile à comprendre. En outre, les forêts sont très difficiles à utiliser en multi-label, et il faut soit un système pour chaque classe, soit un système très complexe.

L'avènement des réseaux neuronaux

L'année 2016 marque un tournant dans la classification par l'avènement des réseaux de neurones (RN) [129]. En effet, pour peu qu'ils disposent des bonnes BDD, les RN peuvent résoudre des tâches très compliquées et permettent une forte amélioration des résultats [38]. Initialement cantonnés aux perceptrons multi-couches, les réseaux se sont rapidement complexifiés avec l'utilisation de réseaux à mémoire temporelle (réseaux *long short-term memory* (LSTM), *recurrent neural network* (RNN)) ou traitement local (*convolutional neural network* (CNN)) issus du traitement d'image [129]. Les caractéristiques peuvent être combinées pour obtenir

des réseaux hybrides très prisés actuellement [31]. Les RN ne sont cependant pas une solution miracle : contrairement aux méthodes antérieures, ils nécessitent un très grand nombre de données d'entraînement, en particulier les plus profonds d'entre eux. En outre, leur complexité calculatoire les rend lents à entraîner et parfois à utiliser. La présence de multiples couches empêche également d'interpréter réellement leur fonctionnement.

Choix de références

Dans ces travaux, nous choisissons d'utiliser un classifieur discriminatif et un génératif issus de l'état de l'art : un CRNN inspiré des références des challenges DCASE 2019 et 2020 [191], et un système de GMM binaires donné comme référence lors du challenge de 2016 [129], tous deux sur des tâches de détection d'événements sonores issus de signaux réels. Ce challenge, organisé annuellement depuis 2016, propose aux chercheurs du monde entier de se pencher sur différentes tâches de classification acoustique (classification de scène, détection d'événements, détection non supervisée, etc.) afin de regrouper et partager les nouvelles connaissances du domaine [73].

Le choix de ces classifieurs, reconnus par la communauté de classification audio, nous permet de travailler à partir d'une base solide dont nous connaissons les performances initiales. Cela nous permet également de positionner nos travaux par rapport à un premier classifieur reconnu pour ces performances (le CRNN) et à un second plus fonctionnel dans un cadre pratique où il est difficile d'obtenir un nombre suffisant de données étiquetées (le GMM).

2.3 Les outils de la classification

Si le classifieur est souvent considéré comme le cœur de la détection d'événements sonores, il n'en reste pas moins un unique élément de la chaîne. En effet, un classifieur ne peut être performant si ses entrées ne lui sont pas adaptées. Pour répondre à la tâche qui lui est imposée, il est nécessaire de disposer d'une BDD adaptée, c'est-à-dire comprenant notamment les classes que l'on souhaite détecter et ayant les mêmes caractéristiques que les signaux que l'on va chercher à traiter. Cependant, on a vu dans la partie précédente qu'il n'est pas souhaitable que l'entrée du classifieur soit un signal audio complet en raison de la difficulté à le traiter et de la redondance des informations. Pour contrer ce problème, on cherche à extraire les descripteurs les plus convenables. Une fois ces éléments déterminés et la décision en sortie du classifieur obtenue, on se focalise sur l'évaluation des performances du système. Une telle analyse nécessite des outils particuliers, adaptés aux objectifs qu'on souhaite atteindre.

Dans cette section, on s'attachera donc d'abord à comprendre les critères qui permettent de choisir une BDD audio. On donnera à cette occasion l'ensemble des BDD qui correspondent à nos propres critères, et on justifiera nos choix. On dressera ensuite un état de l'art des descripteurs audio et des usages actuels. Enfin, on s'intéressera aux critères de performance en sortie du système, et comment les calculer.

2.3.1 Les bases de données

Le choix ou la construction d'une BDD est un élément à établir avec soin. En effet, elle est décisive dans la fiabilité du système de classification en fonctionnement réel [196].

Les propriétés des BDD

Afin de déterminer la BDD à utiliser, on doit connaître les paramètres qui la caractérisent. Dans un premier temps, il est important de s'intéresser à la façon dont le signal a été enregistré. Par exemple, la fréquence d'échantillonnage du signal est une donnée qui va préciser la bande passante de celui-ci ; pour un signal audio, on estime que la fréquence minimale acceptable est de 16kHz. Le type de microphones, mais surtout le nombre de canaux, voire de capteurs utilisés, peut également être un critère de choix si l'on souhaite travailler sur la spatialisation du son. La BDD peut par ailleurs être enregistrée dans un environnement réaliste, ou être issue de données collectées dans d'autres BDD, voire être créée synthétiquement en mélangeant plusieurs pistes audio les unes avec les autres pour pouvoir contrôler la densité d'événements présents [196]. Quelle que soit la façon dont est bâtie la BDD, un choix s'opère quant au recouvrement des différents événements : si ce recouvrement est possible (*i.e.* si plusieurs événements peuvent avoir lieu simultanément), alors l'enregistrement est dit polyphonique (ou multi-label) ; sinon, il est dit monophonique (ou mono-label) [179].

Dans un deuxième temps, l'un des critères particulièrement excluant pour une BDD va être la possibilité d'utilisation de ces données. En effet, selon si la licence permet un accès public ou propriétaire, l'intérêt ne va pas être le même. Néanmoins, avec la montée de l'intérêt pour la recherche reproductible, de plus en plus de BDD sont libres d'utilisation pour la recherche. Cela permet notamment de pouvoir comparer sa propre implémentation de classifieur à l'état de l'art [179].

Enfin, le contenu lui-même a son importance. L'un des éléments dont nous avons déjà parlé est la qualité de l'annotation : selon si celle-ci est forte (*i.e.* continue) ou faible (*i.e.* par segment), voire partiellement absente, les mêmes tâches de classification ne pourront être effectuées. Les signaux peuvent provenir de différents environnements (rue, bureau, habitation...) qui peuvent ou pas correspondre à l'application visée. Au-delà de l'aspect qualitatif, l'aspect quantitatif est stratégique : la durée de l'enregistrement, le nombre de classes, le nombre d'échantillons par classe, mais aussi l'équilibre ou le déséquilibre des classes (*i.e.* si certaines classes ont beaucoup d'échantillons et d'autres peu) sont des éléments absolument nécessaires à prendre en compte [179, 196].

Choix d'une BDD

Dans nos travaux, nous nous intéressons à des enregistrements réels (enregistrés pour la BDD mentionnée ou issue d'une BDD réaliste), polyphoniques et avec un étiquetage fort pour réaliser une détection d'événements. S'il ne nous est pas nécessaire de posséder des enregistrements multi-canaux, nous souhaitons tout de même que les signaux soient au minimum stéréophoniques (ou bicaux) afin de travailler sur la spatialisation du son. Enfin, nous cherchons à utiliser des BDD publiques pour pouvoir comparer nos résultats à l'état de l'art, et avec peu de classes (moins d'une dizaine) pour travailler avec un système d'une complexité abordable. L'ensemble des BDD qui, à notre connaissance, respectent ces critères, sont représentées sur le tableau 2.1.

Parmi ces BDD, DARE-G1 n'est plus disponible en téléchargement et TUT2016 comprend deux types de scènes avec chacun leurs événements, ce qui oblige à créer un système de classification pour chacun. Reste ainsi TUT2017 et Chime-Home. Nous choisissons finalement de retenir la base TUT2017. En effet, malgré un certain déséquilibre de la BDD – malheureusement inévitable dans ce genre d'enregistrements réels, et présent également pour Chime-Home

Nom	Contenu	Fréquence	Durée totale	Classes	Référence
TUT 2016	Ville + Habitation	44.1kHz	49 min	7+11	[128]
TUT 2017	Ville	44.1kHz	92 min	6	[129]
Chime-Home	Habitation	96kHz	6h45	9	[60]
DARE-G1	Mixte	44.1kHz	2h	761	[66]

TABLE 2.1 – Bases de données réelles, polyphoniques et bicanales possédant un étiquetage fort

–, elle représente un choix judicieux en cela qu'elle a été utilisée dans beaucoup de travaux à l'occasion de la tâche 3 du challenge DCASE 2017, et que nous possédons donc des résultats de référence reconnus par la communauté [131].

Description de la base TUT2017

La base TUT2017 est une sous-base de TUT2016 contenant uniquement des signaux issus de la scène extérieure, qui peuvent avoir été remixés par la suite. Elle inclut deux parties : une base d'apprentissage, nommée TUT2017-Development, et une base de test pour évaluer les modèles appris, appelée TUT2017-Evaluation. Chacune comprend six classes d'événements qui peuvent être considérées comme appartenant à deux catégories : celle des véhicules avec les classes *brakes squeaking* (frein grinçant), *car* (voiture qui passe) et *large vehicle* (véhicule important, comme un bus ou un camion), et celle des humains avec *children* (enfant), *people speaking* (personnes discutant) et *people walking* (personnes marchant). La distribution de ces événements au sein des deux sous-bases est indiquée dans la table 2.2 pour la base de développement et la table 2.3 pour la base de test. Nous remarquons dans la base d'apprentissage que la classe *car* y est largement majoritaire avec 2471s de présence totale, contre 97s pour la classe minoritaire *brakes squeaking*. Une grande longueur de présence totale indiquant une grande quantité de données pour entraîner le modèle du classifieur, une telle différence témoigne d'un fort déséquilibre entre les classes. Nous pouvons également observer la présence d'événements longs (en moyenne 15.14s pour la classe *large vehicle*) ou courts (en moyenne 1.87s pour *brakes squeaking*). Enfin, nous pouvons noter des différences parfois importantes entre la BDD d'apprentissage et d'évaluation. En effet, nous observons que les occurrences de *large vehicle* sont en moyenne deux fois moins longues dans la base de test que dans celle d'apprentissage ; au contraire, celles de *brakes squeaking* gagnent en moyenne 1s.

Classe	Nb de trames	Nb d'événements	Longueur totale	Longueur moyenne d'un événement
<i>brakes squeaking</i>	4946	52	97s	1.87s
<i>car</i>	112465	304	2471s	8.13s
<i>children</i>	15784	44	351s	7.98s
<i>large vehicle</i>	43358	61	924s	15.14s
<i>people speaking</i>	35956	89	716s	8.04s
<i>people walking</i>	62479	109	1247s	11.44s

TABLE 2.2 – Statistiques de la base de données TUT2017-Development en fonction des classes

Classe	Nb de trames	Nb d'événements	Longueur totale	Longueur moyenne d'un événement
<i>brakes squeaking</i>	3376	23	67s	2.89s
<i>car</i>	31971	106	663s	6.26s
<i>children</i>	1587	15	31s	2.08s
<i>large vehicle</i>	8196	24	172s	7.15s
<i>people speaking</i>	10364	37	206s	5.56s
<i>people walking</i>	16088	42	320s	7.62s

TABLE 2.3 – Statistiques de la base de données TUT2017-Evaluation en fonction des classes

2.3.2 Descripteurs d'un signal audio

Il existe un très grand nombre de descripteurs adaptés au signal audio. Il est possible de les calculer sur le signal entier, ce qu'on appelle descripteurs globaux, ou indépendamment sur chacune des trames, ce qu'on nomme descripteurs instantanés [150]. Dans le cadre d'un système fonctionnant en temps-réel – ce qui nous intéresse ici –, il est préférable de calculer les descripteurs à chaque trame : on se concentrera donc sur les descripteurs instantanés.

De nombreux descripteurs

Représentation fréquentielle Le choix d'une représentation est la première étape pour obtenir des descripteurs [190]. Le plus souvent, on ne souhaite pas travailler avec la représentation temporelle du signal, on passe alors à une représentation spectrale. Pour cela, on effectue une transformée de transformée de Fourier discrète (TFD) à court-terme sur chaque trame d'analyse (TFCT) pour obtenir, après des transformations détaillées plus loin, une représentation temps-fréquence du signal audio sur une échelle de fréquence linéaire, appelée spectrogramme [196]. Néanmoins, toutes les bandes de fréquence du spectrogramme n'ont pas la même pertinence en terme d'information acoustique. Ainsi, on transpose généralement ce spectrogramme sur une échelle dite perceptuelle, qui tend à mettre en exergue certaines bandes de fréquence pour se rapprocher au mieux de la sensibilité de l'oreille humaine. C'est le cas pour les filtres gammatones, l'échelle mel ou Bark, ou encore la transformée Q-constante [196].

Les descripteurs "artisanaux" La représentation ainsi obtenue est rarement utilisée telle quelle [150, 196], et on lui préfère généralement des descripteurs calculés de manière "artisanale". Ces derniers peuvent eux-mêmes être divisés en plusieurs sous-catégories, comme les descripteurs temporels, spatiaux ou encore perceptuels [150], qu'on retrouvera détaillés en annexe A. Parmi les descripteurs les plus utilisés, on retrouve les MFCC, calculés à partir du mel-spectrogramme [196] et qu'on détaillera plus loin. Ces descripteurs sont parfois dérivés temporellement sur plusieurs trames successives pour obtenir de nouveaux coefficients. En effet, ils sont calculés indépendamment sur chaque trame ; on peut vouloir introduire une dépendance pour modéliser l'évolution temporelle.

Des descripteurs avancés De nombreuses études ont cherché à améliorer ces descripteurs [65, 139, 214] ou à sélectionner les meilleurs ensembles [49, 148, 162, 183]. En outre, ces descripteurs proviennent pour la grande majorité de la recherche sur l'information musicale ou la

reconnaissance de la parole. Les signaux de ces derniers sont très caractéristiques, au contraire de ceux de la détection d'événements environnementaux qui peuvent être de nature très variable. Aussi, depuis quelques années, on s'intéresse à déterminer des descripteurs plus adaptés à la détection d'événements sonores à l'aide de techniques d'apprentissage automatique, tel que la NMF ou des réseaux de neurones [21, 105].

Choix des descripteurs

Dans nos travaux, nous ne souhaitons pas étudier le rôle des descripteurs. Nous nous contenterons donc d'utiliser les descripteurs traditionnels de l'état de l'art, toujours dans le but de comparer nos résultats avec ceux de méthodes reconnues. Pour cela, nous utiliserons un vecteur de descripteurs contenant des MFCC, ainsi que leurs première et deuxième dérivées – appelées Δ MFCC et Δ^2 MFCC – qui est couramment associé au GMM, et un mel-spectrogramme, généralement utilisé en entrée des RN, pour le CRNN.

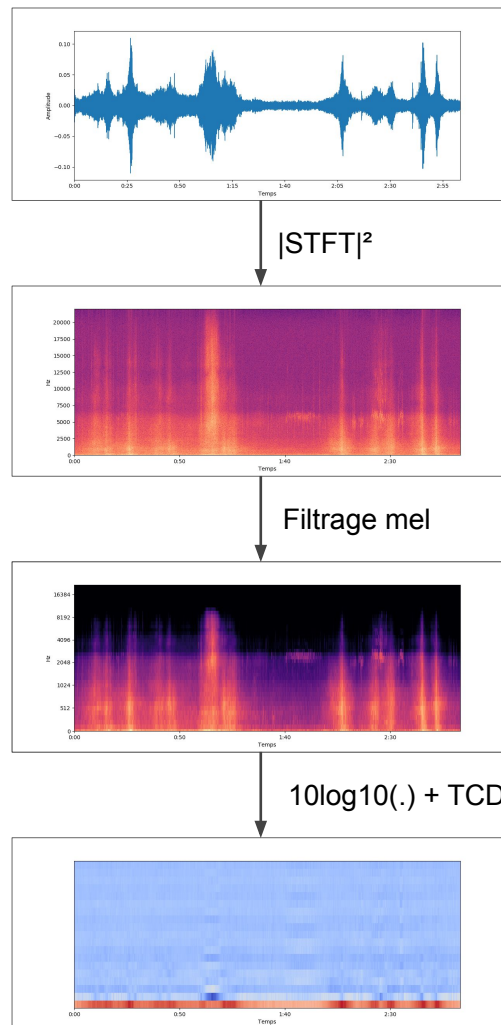


FIGURE 2.3 – Etapes de calcul des coefficients MFCC

Pour rappel, ces descripteurs sont calculés en suivant les étapes indiquées sur la figure 2.3. Dans un premier temps, on découpe le signal, échantillonné à la fréquence f_s , en trames de durée

T_w avec un recouvrement de T_h en appliquant une fenêtre de Hamming. On effectue ensuite une TFD sur chaque trame. En pratique, cela revient à calculer la transformée de Fourier à court terme (TFCT) X sur le signal x à l'aide de la transformée de Fourier rapide (TFR) (sur n_{fft} points). En prenant son module au carré $|X|^2$, on obtient son périodogramme. Ce périodogramme est filtré par un banc de n_{mel} filtres mels, des filtres triangulaires, pour obtenir le mel-spectrogramme M utilisé par le CRNN. Ce mel-spectrogramme est exprimé en décibels (c'est-à-dire qu'on calcule $10 \log(M)$), puis on effectue une transformée en cosinus discrète (TCD) dont on garde les n_{mfcc} ($n_{\text{mfcc}} \leq n_{\text{mel}}$) premiers MFCC pour l'ensemble du signal. Cette étape permet de compresser et décorréliser l'information [132]. On dérive ensuite localement (avec n_w le nombre de trames utilisées pour la dérivation) deux fois les MFCC à l'aide de l'algorithme de Savitzky-Golay pour calculer les ΔMFCC et $\Delta^2\text{MFCC}$. Par la suite, le premier coefficient des MFCC (appelé coefficient d'ordre 0), qui est relié à l'énergie du signal, est généralement supprimé. Pour obtenir le vecteur de descripteurs d'entrée du GMM, on concatène ainsi les $n_{\text{mfcc}} - 1$ derniers MFCC (*i.e.* tous sauf le premier), les n_{mfcc} ΔMFCC et les n_{mfcc} $\Delta^2\text{MFCC}$.

Paramètre	Description	Valeur
f_s	Fréquence d'échantillonnage	44.1 kHz
T_w	Durée d'une trame	40 ms
T_h	Durée de recouvrement	20 ms
n_{fft}	Taille de la TFR	2048
n_{mel}	Nombre de filtres mels	40
n_{mfcc}	Nombre de MFCC	20
n_w	Largeur de dérivation	9

TABLE 2.4 – Paramètres utilisés pour le calcul du mel-spectrogramme et des MFCC

Dans l'ensemble de ces travaux, les paramètres choisis sont ceux présentés dans la table 2.4

2.3.3 Critères de performance en classification

Comptage des échecs et succès

Lors de la phase de test, à l'issue du passage d'une trame dans un classifieur, on récupère un vecteur contenant l'ensemble des classes estimées. Pour calculer les performances de ce classifieur, on va pouvoir comparer ce vecteur à celui qui liste des événements réellement présents sur cette trame. On comptabilise alors plusieurs données qui forment les statistiques intermédiaires [127]. Pour chaque trame, on peut alors déterminer (voir également la figure 2.4) :

- **Les vrais positifs (*true positive*, TP)** : une classe estimée active est réellement présente sur la trame ;
- **Les faux positifs (*false positive*, FP)** : une classe estimée active ne l'est pas d'après le vecteur de référence ;
- **Les faux négatifs (*false negative*, FN)** : une classe estimée inactive est pourtant présente sur la trame ;
- **Les vrais négatifs (*true negative*, TN)** : une classe estimée inactive l'est réellement.

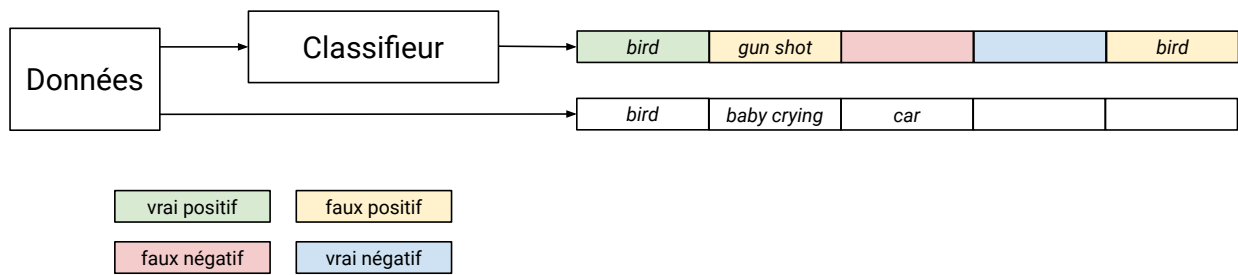


FIGURE 2.4 – Comptage des vrais et faux positifs, et vrais et faux négatifs, pour une classification monophonique

Si les événements doivent être détectés temporellement (c'est-à-dire que l'on cherche à déterminer leur début et leur fin), ces valeurs peuvent être redéfinies à l'échelle d'une occurrence entière de l'événement :

- **Les vrais positifs** : on considère qu'un événement est un vrai positif quand il est détecté et correctement étiqueté, et que le temps de début (et éventuellement de fin, en fonction de l'application) est acceptable. En effet, on peut se permettre de laisser une marge de tolérance pour la détection [127] ;
- **Les faux positifs** : de la même manière que pour un ensemble de trames, on considère un faux positif quand, sur la période de détection d'un événement, on ne trouve aucune correspondance sur la référence, en prenant en compte la marge de tolérance ;
- **Les faux négatifs** : dans le respect de la marge de tolérance, un événement présent dans les annotations de référence n'apparaît pas dans ceux détectés au moment où il aurait dû ;
- **Les vrais négatifs** : cette fois-ci, la classe n'est pas présente et pas détectée. Cependant, compte tenu de la tolérance, on ne comptabilise généralement pas cette statistique.

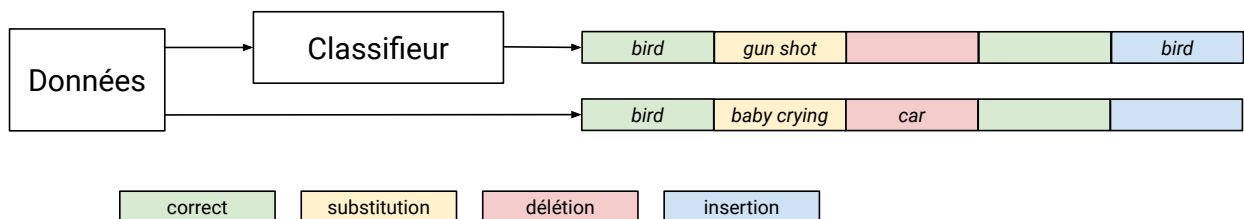


FIGURE 2.5 – Schéma du comptage des substitutions, délétions et insertions

Dans les deux cas, on peut également extraire de ces valeurs de nouvelles statistiques en terme d'erreur (voir également la figure 2.5) :

- **Les substitutions (S)** : on considère qu'il y a substitution quand une classe prend la place d'une autre. Bien entendu, les estimations ne sont pas ordonnées : une substitution correspond donc à un faux positif plus un faux négatif. On peut ainsi calculer : $S = \min(FP, FN)$;

- **Les délétions (D)** : on considère qu'il y a délétion d'une classe quand celle-ci n'est pas présente sur le vecteur des classes estimées, mais qu'elle n'a pas été substituée par une autre, c'est-à-dire que le nombre de faux négatifs est plus important que celui de faux positifs. On peut donc calculer : $D = \max(0, FN - FP)$;
- **Les insertions (I)** : on considère qu'il y a insertion d'une classe quand celle-ci est présente sur le vecteur des classes estimée, mais qu'elle ne se substitue pas à une autre, c'est-à-dire que le nombre de faux positifs est plus important que celui de faux négatifs. On peut donc calculer : $D = \max(0, FP - FN)$.

Les principaux critères de performance

À partir des décomptes précédemment effectués, il est possible de calculer différentes valeurs qui constitueront nos critères de performance. Il en existe un grand nombre dans le domaine de la classification, aussi n'allons-nous présenter que les principaux.

- **La précision (P)** : Elle correspond à la proportion d'événements correctement estimés parmi tous les événements estimés [125]. Elle est comprise entre 0 et 1, et meilleure quand elle approche de 1 :

$$P = \frac{TP}{TP + FP}. \quad (2.1)$$

- **Le rappel (R)** : Il représente la proportion d'événements détectés parmi tous les événements présents. Il est compris entre 0 et 1, et meilleur quand il tend vers 1 :

$$R = \frac{TP}{TP + FN}. \quad (2.2)$$

- **La F-mesure (F)** : Elle représente la proportion d'événements correctement estimés parmi les événements présents ou estimés (elle synthétise la précision et le rappel). Elle est comprise entre 0 et 1, et meilleure quand elle tend vers 1 :

$$F = \frac{2P.R}{P + R} = \frac{2TP}{2TP + FN + FP}. \quad (2.3)$$

Ces trois dernières mesures sont issues du domaine de la recherche d'information [127]. La F-mesure est un critère très apprécié dans la détection d'événements [191, 192] par sa facilité d'interprétation. Néanmoins, elle est très sensible au choix du moyennage expliqué ci-après [127].

- **Le taux d'erreurs (ER)** : Le taux d'erreur est un critère issu de la reconnaissance de la parole [196]. Soit N le nombre d'événements réellement présents, le taux d'erreur est la proportion d'erreurs par rapport aux nombres d'événements [152] :

$$ER = \frac{S+D+I}{N}. \quad (2.4)$$

On peut avoir un taux $ER > 1$ dans le cas où le nombre d'erreurs est supérieur à celui des événements présents. Cela indique normalement un grand nombre d'insertions. Plus un taux d'erreur est proche de 0, meilleur il est.

- **La spécificité** : Il s'agit de la proportion de silence correctement détectée parmi tous les silences présents [62]. Elle est comprise entre 0 et 1, et meilleure quand elle approche de 1 :

$$\text{spécificité} = \frac{\text{TN}}{\text{TN} + \text{FP}}. \quad (2.5)$$

- **L'exactitude** : Il s'agit cette fois-ci d'une mesure globale de la proportion des détections correctes (silence et événements) sur l'ensemble des détections effectuées. Elle est comprise entre 0 et 1, et meilleure quand elle tend vers 1 :

$$\text{exactitude} = \frac{\text{TP} + \text{TN}}{\text{TP} + \text{TN} + \text{FP} + \text{FN}}. \quad (2.6)$$

- **Le rapport de recouvrement moyen** : Ce critère n'est utilisé qu'en détection d'événements mesurant le début et la fin de ceux-ci [163]. Il est calculé pour un événement correctement détecté, avec *onsets* faisant référence aux indices temporels du début pour l'événement détecté et celui de référence, et *offsets* défini de la même manière pour la fin de l'événement. On peut ensuite le moyenner pour tous les événements. Il est compris entre 0 et 1, et indiquant de meilleures performances quand il tend vers 1 :

$$\text{OR} = \frac{\min \text{offsets} - \max \text{onsets}}{\max \text{offsets} - \min \text{onsets}}. \quad (2.7)$$

Moyennage des mesures

Ces critères sont généralement calculés sur chacune des classes indépendamment. A la suite de ces mesures, il est possible de les moyenner afin d'obtenir des indicateurs plus généraux sur les performances du système. Il existe deux manières de faire, chacune donnant dans certains cas des résultats bien différents les uns des autres, et possédant donc ses propres caractéristiques. Il est ainsi nécessaire de les connaître et de se fixer des objectifs afin de choisir la plus judicieuse.

Le moyennage par instance, ou micro-moyennage, consiste à moyenner un même critère de performance sur l'ensemble des échantillons de test. Comme illustré sur la figure 2.6, chaque instance, que ce soit une trame, un segment ou un événement, a alors le même poids dans la synthèse des performances. Cela signifie que les performances sont mesurées sur la détection réelle du système. Néanmoins, si une classe est surreprésentée, elle aura un grand poids dans les performances globales ; si ses propres résultats sont faibles, cela aura donc un fort impact sur ceux du classifieur. Par exemple, la F-mesure totale micro-moyennée sera calculée

$$F_{\text{micro}} = \frac{2 \sum_{k=1}^C \text{TP}_k}{2 \sum_{k=1}^C (\text{TP}_k + \text{FN}_k + \text{FP}_k)} \quad (2.8)$$

avec C le nombre total de classes et TP_k , FN_k et FP_k respectivement le total de vrais positifs, faux négatifs et faux positifs de la classe $n^\circ k$.

Pour éviter ce problème, il est possible de moyenner par classe, ou macro-moyenner. Ici, on procède en deux étapes comme indiqué sur la figure 2.6 : tout d'abord, on effectue un micro-moyennage pour chacune des classes. On calculera ensuite une moyenne classique sur les résultats obtenus. Dans ce cas, même les plus petites classes ont une importance similaire dans les résultats agrégés. Par exemple, la F-mesure totale macro-moyennée sera calculée

$$F_{\text{micro}} = \sum_{k=1}^C \left(\frac{2\text{TP}_k}{2\text{TP}_k + \text{FN}_k + \text{FP}_k} \right). \quad (2.9)$$

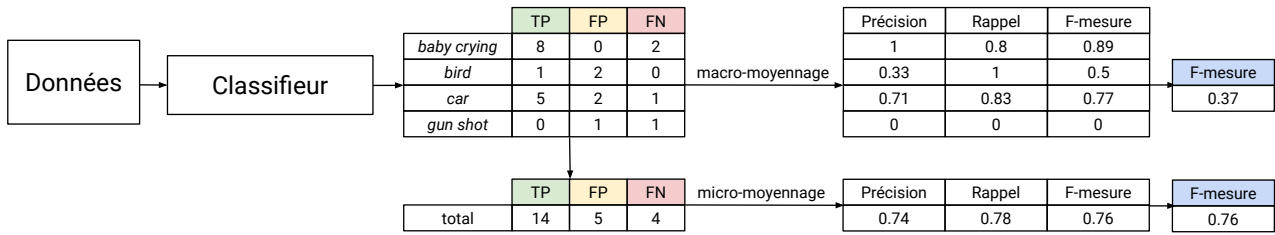


FIGURE 2.6 – Schématisation des conséquences du micro et du macro-moyennage

La différence entre les deux approches est illustrée sur la figure 2.6. Sur ce schéma, on constate les remarques précédentes : les classes *baby crying* et *car* sont les plus représentées, et ce sont également elles qui affichent les meilleures performances. Par micro-moyennage, on obtient donc une F-mesure élevée. Cependant, si on moyenne plutôt les performances de chaque classe, chacune d'entre elle obtient le même poids dans le score final. Ainsi, les deux classes minoritaires qui affichent une F-mesure plus faible, et en particulier la classe *gun shot* avec sa F-mesure nulle, font radicalement diminuer les performances mesurées du classifieur. C'est ce qu'on appelle le paradoxe de Simpson.

Le choix des critères de performance

Le choix d'un critère de performance doit être fait avec soin en fonction des objectifs visés [196]. Dans notre cas, on sélectionne le taux d'erreur et la F-mesure, calculés à partir des segments d'une seconde, de par leur popularité en détection d'événements [191, 192] ; nous pouvons ainsi comparer plus facilement nos résultats avec l'état de l'art. La précision et le rappel seront également donnés pour approfondir certaines analyses. Nous conservons à la fois le micro et le macro-moyennage dans nos analyses, avec une préférence pour le macro-moyennage en raison du déséquilibre des classes.

2.4 Les modèles de mélange gaussien

Dans cette section, nous présentons notre premier classifieur, à savoir le GMM soustractif. Issu de la référence de la tâche 3 du challenge DCASE 2016 [128], il consiste à apprendre un modèle pour détecter la présence et un autre pour détecter l'absence de chaque classe. Il s'agit par ailleurs d'une méthode très prisée avant l'avènement des réseaux de neurones [130]. Comme expliqué dans la section 2.3.2, on l'utilise avec un vecteur d'entrée comprenant 20 MFCC auxquels on enlève le premier, 20 Δ MFCC et 20 Δ^2 MFCC. Cette section vise à reproduire et étudier superficiellement ce classifieur afin de s'assurer de la justesse de son comportement par rapport à l'état de l'art, avant d'en analyser l'implémentation sur système embarqué dans le chapitre 5.

On commencera par exposer le fonctionnement précis de ce GMM soustractif. Ensuite, on l'implémentera et on vérifiera son bon fonctionnement. Enfin, on étudiera le comportement du GMM soustractif face à différents paramètres qui influent sur la complexité de la phase de test du système.

2.4.1 Fonctionnement théorique d'un modèle de mélange gaussien soustractif

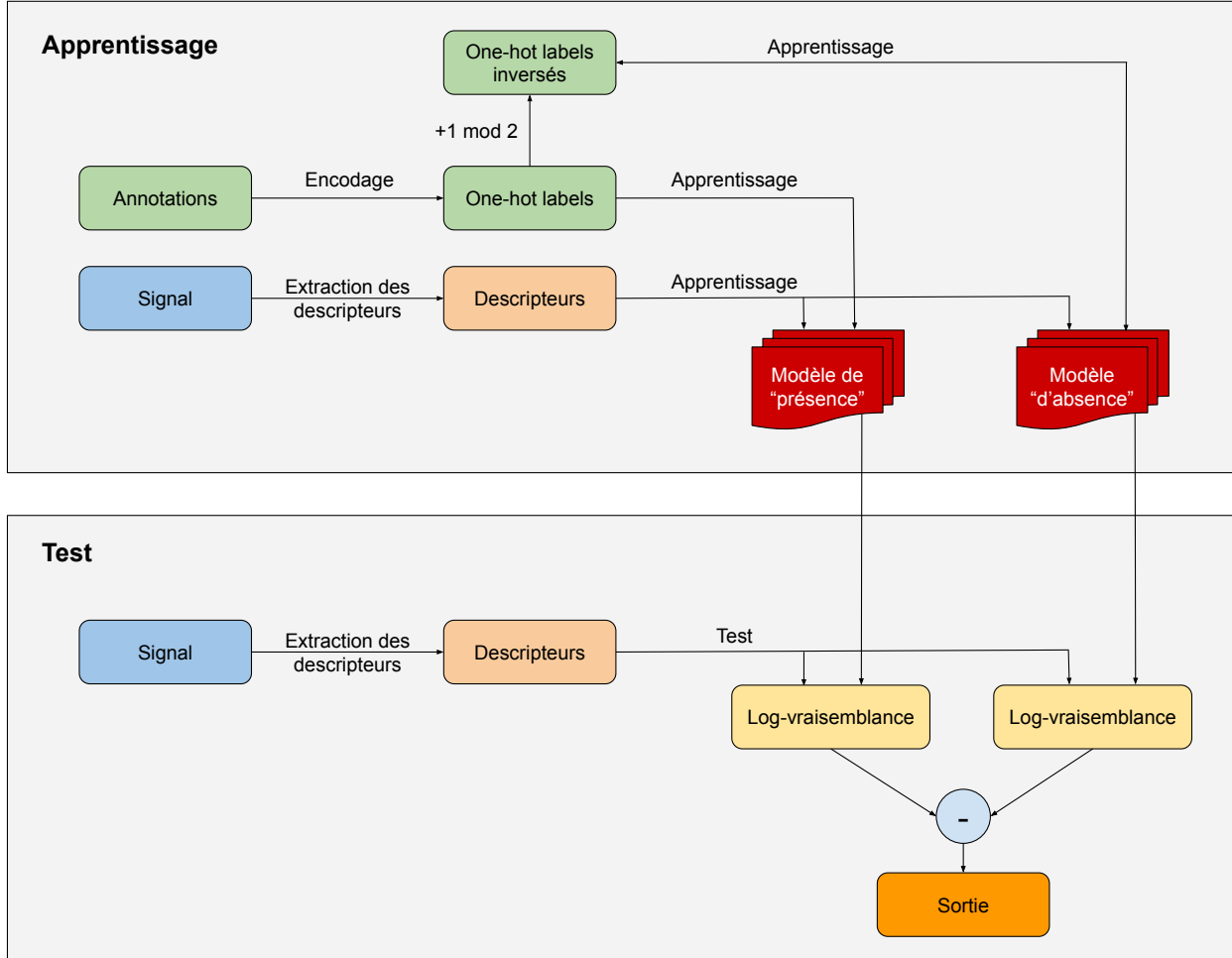


FIGURE 2.7 – Schéma de fonctionnement du GMM soustractif

Pour construire un GMM, on suppose que les descripteurs de chaque classe suivent un mélange d'une ou plusieurs lois gaussiennes. Ainsi, chacune des C classes va être représentée indépendamment par un mélange de variables suivant chacune une loi normale multi-dimensionnelle, c'est-à-dire ayant une densité de probabilité pour chacune des n_c gaussiennes :

$$f(x; \mu_k, \Sigma_k) = \frac{1}{(2\pi)^{N/2} \det(\Sigma_k)^{1/2}} \exp \left(-\frac{1}{2} (x - \mu_k)^\top \Sigma_k^{-1} (x - \mu_k) \right) \quad (2.10)$$

pour un vecteur $x \in \mathbb{R}^N$ avec N la dimension (ici la longueur du vecteur des descripteurs), $\mu_k \in \mathbb{R}^N$ le vecteur moyenne de la k e gaussienne et $\Sigma_k \in \mathcal{M}_N(\mathbb{R})$ sa matrice de covariance définie positive. Pour cela, comme on l'observe sur la figure 2.7 qui résume le fonctionnement du GMM soustractif, on extrait les descripteurs du signal audio et on encode les annotations sous la forme d'une matrice 1-parmi-n, ou *one-hot labels* (étendue à un signal polyphonique comme vu dans la section 2.2.2). A chaque trame, on vérifie dans la matrice la présence de chaque classe i : si elle est présente, on met à jour le modèle de mélange gaussien de classe i

Algorithme d'espérance-maximisation

Jusqu'à convergence du modèle, on alterne les deux étapes suivantes.

Espérance Soit π_k la probabilité d'un individu quelconque suivre la loi de la gaussienne k , et $\Phi_k = \{\pi_k, \mu_k, \Sigma_k\}$ le paramètre global de celle-ci. L'étape d'estimation consiste à calculer les probabilités *a posteriori* qu'un vecteur x_i suive la loi de la gaussienne k , soit :

$$t_{i,k} = \frac{\pi_k \times f(x_i; \mu_k, \Sigma_k)}{\sum_{j=1}^{n_c} \pi_j \times f(x_i; \mu_j, \Sigma_j)}. \quad (2.11)$$

Maximisation On maximise ensuite l'espérance. Pour cela, on met à jour ses paramètres à partir des probabilités *a posteriori* obtenues pour les N individus de la base d'apprentissage :

$$\pi_k \leftarrow \frac{1}{N} \sum_{i=1}^N t_{i,k} \quad (2.12)$$

$$\Sigma_k \leftarrow \frac{\sum_{i=1}^N t_{i,k} (x_i - \mu_k)(x_i - \mu_k)^\top}{\sum_{i=1}^N t_{i,k}} \quad (2.13)$$

$$\mu_k \leftarrow \frac{\sum_{i=1}^N x_i t_{i,k}}{\sum_{i=1}^N t_{i,k}}. \quad (2.14)$$

TABLE 2.5 – Description de l'algorithme d'espérance-maximisation

suivant la méthode d'espérance-maximisation (EM) [38] décrite dans l'encadré 2.5 ; sinon, on ne change rien.

Dans la phase de test d'un GMM simple, on calcule la log-vraisemblance de tous les événements pour chaque classe i par :

$$\ln(\mathcal{L}(x)) = \ln \left(\sum_{k=1}^{n_c} \pi_k \times f(x_i; \mu_k, \Sigma_k) \right). \quad (2.15)$$

Après un post-traitement éventuel (normalisation, intégration sur plusieurs trames...), on vérifie si la log-vraisemblance dépasse un seuil s fixé au préalable (ce qui signifie que l'événement est présent) ou non (l'événement n'est donc pas présent), soit

$$P(\text{classe } i \text{ présente}) = P(\ln(\mathcal{L}(x)) > s). \quad (2.16)$$

Pour un GMM soustractif, on combine, comme on le voit sur la figure 2.7, deux systèmes de GMM. D'une part, on construit le modèle classique que l'on vient de présenter et qui permet de calculer la vraisemblance de la présence de chaque classe. D'autre part, un second modèle est déterminé, non pas à partir de la matrice des annotations, mais de la matrice opposée

(tous les 1 deviennent des 0 et inversement), c'est-à-dire une matrice qui signale l'absence de chaque classe; nommé modèle d'absence, il calcule la vraisemblance que l'événement ne soit pas présent.

Une fois les deux log-vraisemblances obtenues, on soustrait celle d'absence à celle de présence, *i.e.*

$$D(x) = \ln(\mathcal{L}_{pos}(x)) - \ln(\mathcal{L}_{neg}(x)) \quad (2.17)$$

avec $\mathcal{L}_{pos}$ l'opération de vraisemblance du modèle de présence et $\mathcal{L}_{neg}$ celle du modèle d'absence. Le raisonnement à l'origine de cette opération est que, si un événement a lieu sur une trame, la log-vraisemblance de son modèle de présence doit être plus importante que celle du modèle d'absence, et leur différence est donc positive; dans le cas contraire, elle est négative. Le but de ce système est de prévenir certaines erreurs liées au mauvais positionnement du seuil, notamment s'il est identique pour toutes les classes et que celles-ci sont déséquilibrées. En effet, une classe majoritaire aura tendance à avoir une log-vraisemblance plus haute que la classe minoritaire. Or dans ce système, si la log-vraisemblance de la classe minoritaire est faible mais que la log-vraisemblance d'absence l'est encore plus, on sait que l'événement est présent. Malgré tout, on verra dans la section suivante que le seuil d'un tel GMM soustractif n'a pas toujours intérêt à être mis à 0.

Dans la suite de nos expériences, on fixe les paramètres du GMM aux valeurs de la table 2.6. Les GMM sont initialisés à l'aide d'une méthode *init*, ici fixée à "random", c'est-à-dire une initialisation aléatoire; celle-ci est réalisée un nombre n_{init} de fois, et celle ayant la plus faible inertie (*i.e.* la somme des distances au carré entre chaque individu et son centroïde le plus proche) est conservée. Chaque mélange contient un nombre de n_c gaussiennes, qui seront apprises par le passage successif de l'ensemble des données d'apprentissage pendant n_{epoch} époques. La matrice de covariance peut être calculée normalement (elle est alors de type *covtype* dite "*full*") ou être simplifiée comme matrice diagonale (alors dite "*diag*"). Enfin, les descripteurs en entrée du classifieur peuvent ou non être normalisés, ce qui sera décrit par le paramètre *norm*.

Paramètre	Description	Valeur
<i>init</i>	Type d'initialisation	random
n_{init}	Nombre d'initialisations	1
n_c	Nombre de gaussiennes	4
n_{epoch}	Nombres d'époques de l'entraînement	80
<i>covtype</i>	Nature de la matrice de covariance	full
<i>norm</i>	Normalisation	faux

TABLE 2.6 – Paramètres utilisés dans le GMM soustractif

2.4.2 Implémentation et vérification de l'algorithme

Dans cette section, on s'intéresse à vérifier le fonctionnement du GMM soustractif. Celui-ci est implémenté en Python 3 à l'aide des bibliothèques *dcase_util* pour la gestion de la chaîne de classification, *librosa* pour le traitement du signal audio [124], *scikit-learn* pour l'apprentissage du GMM [149] et *sed_eval* pour le calcul des critères de performance [127].

Bien que l'algorithme d'EM utilisé pour les GMM converge toujours vers un maximum local [38], il reste nécessaire de vérifier que notre classifieur converge au cours de l'apprentissage, ceci afin de s'assurer qu'aucun dysfonctionnement ne vienne perturber nos résultats.

Après avoir vérifié que ses performances correspondent à l'état de l'art, on examine le rôle de paramètres qui jouent sur l'apprentissage, à savoir la nature de l'initialisation et le seuil de décision.

Validation du fonctionnement du GMM soustractif

Pour chaque classe, l'apprentissage d'un GMM consiste à entraîner différentes gaussiennes en déterminant leurs statistiques (probabilité d'apparition, vecteur moyenne, matrice de covariance) sur une séquence de données. Ces données, initialisées à une certaine valeur selon le type d'initialisation choisi (K-moyenne ou aléatoire), doivent tendre vers un maximum local à une vitesse que l'on souhaite la plus haute possible, ces deux éléments dépendant notamment de l'initialisation. Dans le cas présent, on fixe une initialisation aléatoire, qui est détaillée dans l'annexe B.

On visualise différents éléments nous permettant de nous assurer de la convergence : la norme du vecteur moyenne, le déterminant et le conditionnement de la matrice de covariance, le tout pour chacune des quatre gaussiennes de chaque mélange, ainsi que la borne inférieure de la log-vraisemblance calculée sur les données d'entraînement, et ce pour les GMM de présence et d'absence de toutes les classes. On entraîne nos classifieurs sur cent époques, et on relève les mesures toutes les dix époques. Parmi les six classes de notre base de données, on illustre sur la figure 2.8 la convergence du système par la classe *people speaking* (personnes parlant), pour laquelle la convergence est la plus lente. On observe, comme on s'y attendait, que les valeurs de toutes ces statistiques convergent correctement vers une valeur constante. Le GMM d'absence, qui est entraîné par davantage de données (en effet, chaque classe est majoritairement absente sur l'ensemble de la BDD), converge très rapidement. Au contraire, le GMM de présence, qui est entraîné par moins de données, n'atteint sa convergence qu'au bout de 80 époques. Ce phénomène peut également être observé pour les autres classes, mêmes si celles-ci ont une convergence plus rapide, probablement en raison de la complexité du signal de parole et de sa relation étroite avec les MFCC. Par ailleurs, on note que, pour les deux GMM, les gaussiennes 1 et 3 ont des statistiques très proches les unes des autres ; cela n'est cependant pas le cas pour toutes les classes. Considérant à présent la norme des vecteurs moyenne et la borne inférieure de la log-vraisemblance, les valeurs de celles-ci sont tout à fait représentatives de ce qui peut être obtenu pour notre application. Il en est de même pour le déterminant et le conditionnement des matrices de covariance, que nous souhaitons tout de même commenter ici. En effet, le conditionnement d'une matrice représente le ratio entre sa valeur propre la plus élevée et celle la plus faible, ce qui nous donne ici un ratio très important pour les gaussiennes 2 et 4 dont les valeurs sont, par exemple, de 1063 et 3753 après 80 époques. Il est important d'y faire attention car des valeurs d'ordres de grandeur supérieur pourraient empêcher d'obtenir une vraisemblance convenable. De la même manière, le déterminant des matrices peut prendre des valeurs extrêmement faibles (respectivement 3.57×10^{-22} et 2.26×10^{-10} pour les deux gaussiennes étudiées précédemment) ou au contraire extrêmement fortes (9.42×10^6 pour la gaussienne 3 du GMM de présence). Or dans le calcul de la vraisemblance, la racine carrée du déterminant est au dénominateur ; dans certains cas, même si la racine carrée a tendance à écraser les valeurs extrêmes, la log-vraisemblance va se révéler très forte ou au contraire très faible quel que soit le vecteur de descripteurs de test, ceci à cause de ce déterminant. Cependant,

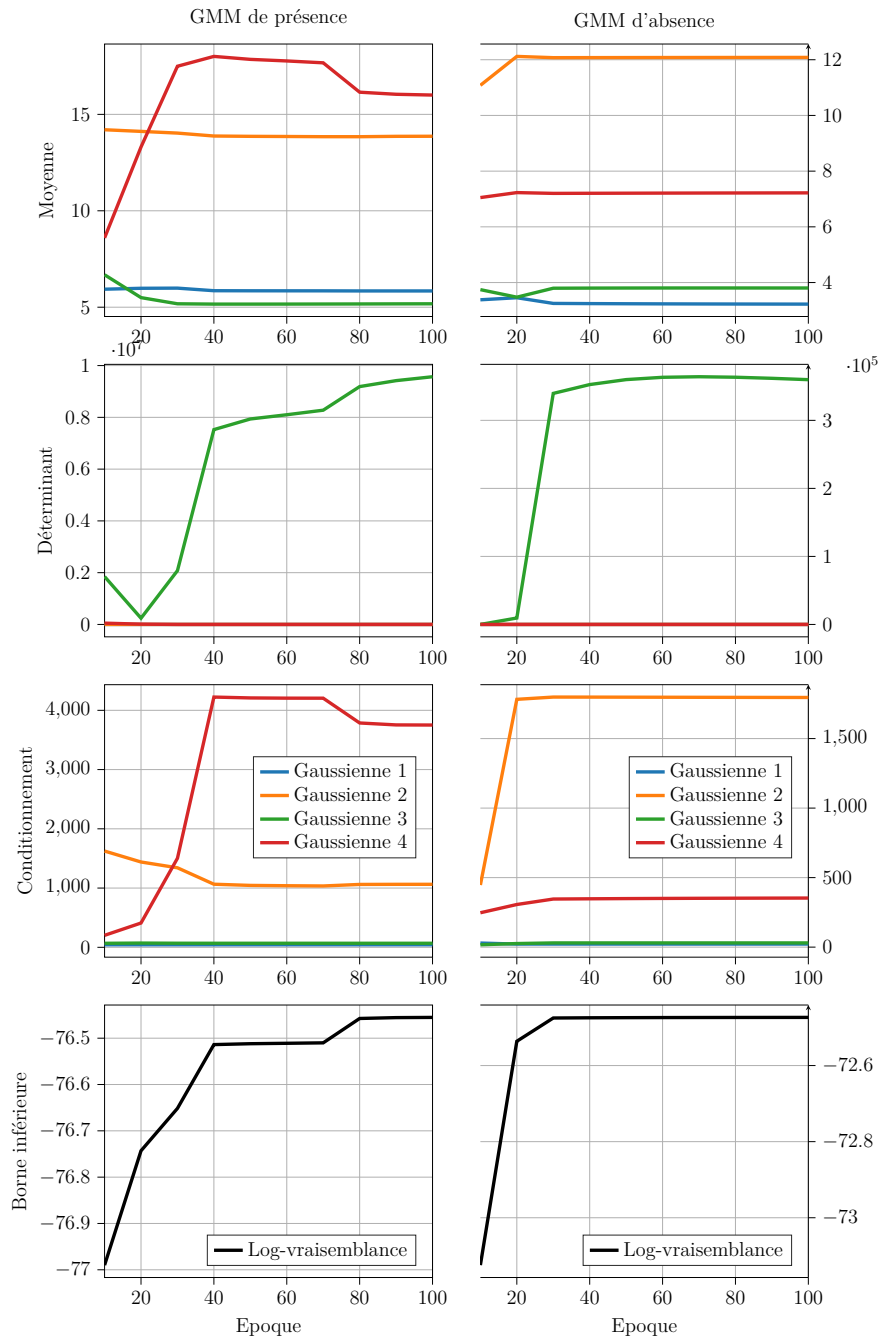


FIGURE 2.8 – Evolution de la moyenne, du déterminant et du conditionnement de la matrice de covariance pour chacune des quatre gaussiennes, et borne inférieure de la log-vraisemblance au cours du temps (en époque), pour les classifieurs de présence (à gauche) et d'absence (à droite) et la classe *people speaking*

les résultats de ce classifieur vont nous prouver que ce n'est pas le cas.

Finalement, on se fixe un nombre d'époques pour l'apprentissage de $n_{\text{epoch}} = 80$ pour s'assurer de la convergence de tous les GMM.

Pour valider définitivement notre implémentation, on compare les résultats de notre code avec celui proposé lors du challenge DCASE 2016 [128] sur la base de données TUT 2017. En fixant le seuil de décision sur la différence de log-vraisemblance à 0, on obtient les performances en terme de taux d'erreur et de F-mesure présentées sur la table 2.7. On remarque que par micro-moyennage, les valeurs sont toujours pratiquement identiques. Pour le macro-moyennage en revanche, on observe une certaine différence qui penche en faveur de la référence, avec une F-mesure de 37.80% contre 34.10% pour notre implémentation, et un taux d'erreur de 2.19 contre 2.50. Ce changement est majoritairement lié à la classe *brakes squeaking*, qui perd pratiquement 29.8% de F-mesure, et gagne 0.74 de taux d'erreur, alors que les autres classes montrent relativement peu de différence entre notre implémentation et celle de référence. Ce changement est lié au fait que la classe *brakes squeaking* est l'une de celles avec le moins d'échantillons tant pour la base d'apprentissage que pour la base de test, et la moindre variation peut conduire les performances à croître ou décroître fortement. Néanmoins, nous validons l'implémentation de notre algorithme.

	Implémentation		Référence	
	ER	F-mesure	ER	F-mesure
Total par micro-moyennage	1.42	43.79%	1.37	43.37%
Total par macro-moyennage	2.50	34.10%	2.19	37.80%
<i>brakes squeaking</i>	1.80	19.4%	1.06	49.2%
<i>car</i>	0.76	63.3%	0.74	60.8%
<i>children</i>	4.93	9.8%	3.84	9.4%
<i>large vehicle</i>	3.50	33.0%	3.29	34.9%
<i>people speaking</i>	2.43	30.2%	2.56	24.9%
<i>people walking</i>	1.59	48.8%	1.64	47.7%

TABLE 2.7 – Comparaison entre les taux d'erreur (ER) et F-mesure de notre implémentation (à gauche) et de l'implémentation de référence (à droite) pour un seuil de décision de 0

Validation des paramètres

Comme indiqué sur la table 2.6, nous avons jusqu'à présent considéré que les GMM sont entraînés à partir d'une initialisation aléatoire. On va donc étudier l'effet de cette initialisation. En effet, pour les GMM, celle-ci peut être réalisée selon plusieurs méthodes, dont les principales sont les initialisations aléatoires et par K-moyenne.

La première méthode consiste à attribuer des poids aléatoires à l'ensemble des gaussiennes n_c et des N individus de la base d'apprentissage, puis d'en déduire les paramètres du GMM. La méthode est détaillée en annexe B.

La seconde méthode, celle de la K-moyenne, consiste à minimiser la moyenne de la distance au carré entre les individus et un ensemble de n_c centroïdes. A l'issue de la convergence de ces centroïdes, la proportion d'individus appartenant à chaque centroïde est assimilée au poids π_k de chaque gaussienne, les centroïdes aux moyennes μ_k , puis les matrices de covariance Σ_k sont définies selon l'équation 2.13.

L'initialisation peut être répétée plusieurs fois. Dans ce cas, après chaque initialisation, on calcule la log-vraisemblance minimale sur la base d'apprentissage, et on sélectionne les résultats de l'itération ayant donné la plus haute. Dans notre cas, on effectue une unique initialisation car l'algorithme EM ajuste les paramètres par la suite.

En parallèle, on va étudier l'influence du seuil de décision. Fixé jusqu'à présent à 0, il est défini comme la valeur de la différence entre log-vraisemblances de présence et d'absence qui implique qu'un événement est considéré comme présent. En effet, en théorie, la différence de log-vraisemblance négative lorsque l'événement est absent, et elle est positive quand l'événement est présent. En observant la distribution de cette différence pour chaque classe, on obtient dans l'idéal le résultat de la figure 2.9a. Dans ce cas, fixer un seuil à 0 permet en effet de séparer efficacement les échantillons sur lesquels l'événement est présent ou non. En pratique, les deux distributions ne sont pas aussi disjointes et la séparation optimale ne se fait pas à 0, comme on peut le remarquer sur la figure 2.9b qui simule une distribution proche de ce qui peut être observé dans la réalité. Comme il est difficile de placer un unique seuil pour toutes les classes juste en observant la distribution de la différence de log-vraisemblance, on calcule plutôt les performances pour différents seuils allant de -100 à 100 par pas de 5.

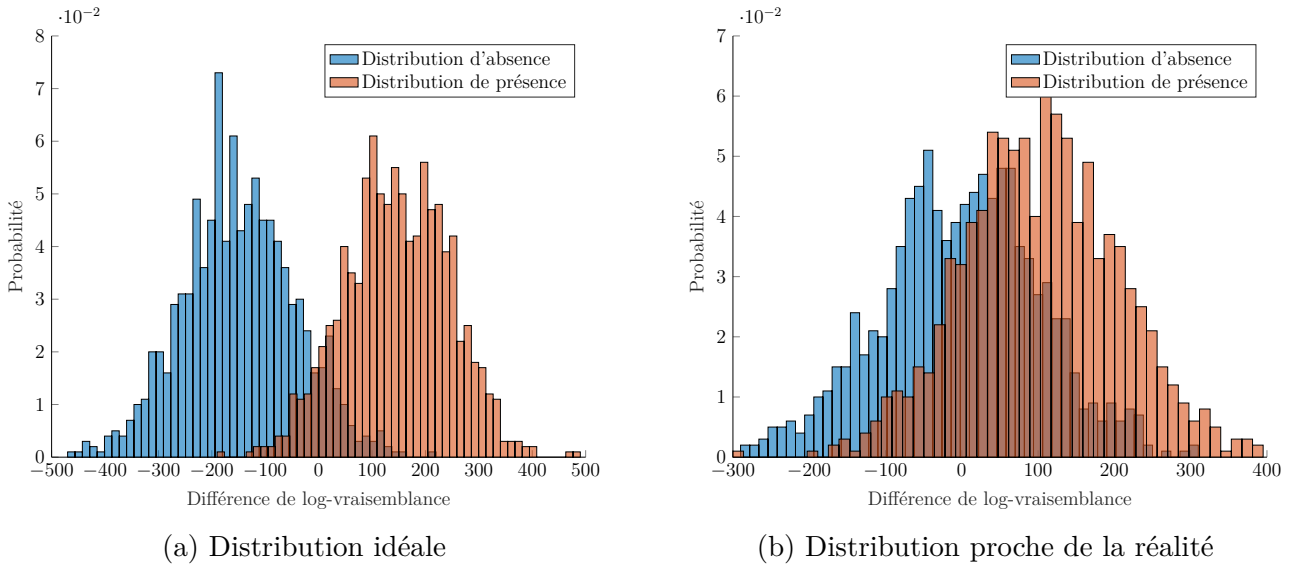


FIGURE 2.9 – Exemple théorique de distribution de la différence de log-vraisemblance pour l'ensemble des échantillons sur lesquels la classe est présente (dite distribution de présence) ou absente (dite distribution d'absence)

C'est ce que nous avons effectué afin de comparer les deux initialisations, aléatoire et par K-moyenne. En effet, le seuil de détection qui maximise les performances peut varier en fonction des différents paramètres qui influent sur l'apprentissage, il n'est donc pas toujours équitable de comparer deux méthodes à seuil fixe. On présente ainsi les taux d'erreur en fonction de la F-mesure issus du micro-moyennage pour des seuils de décision compris entre -100 et 100 par pas de 5 sur la figure 2.10a pour l'initialisation aléatoire et sur la figure 2.10b pour l'initialisation par K-moyenne. Pour plus de lisibilité, on indique seulement la légende d'un point sur deux. Dans les deux cas, on observe que, quand le seuil augmente, la F-mesure croît jusqu'à un certain seuil, puis décroît fortement ; le taux d'erreur, lui, décroît toujours avec l'augmentation du seuil, d'abord fortement, puis faiblement quand la F-mesure baisse significativement. Ainsi, il semble exister pour chacun de ces contextes une valeur du seuil de détection qui permet d'avoir une

F-mesure maximale avec un taux d'erreur relativement faible par rapport à son minimum. Pour l'initialisation aléatoire, il s'agit de la valeur 35 ; pour la K-moyenne, il s'agit de 5. A noter que les mêmes seuils peuvent être trouvés à partir des performances du macro-moyennage.

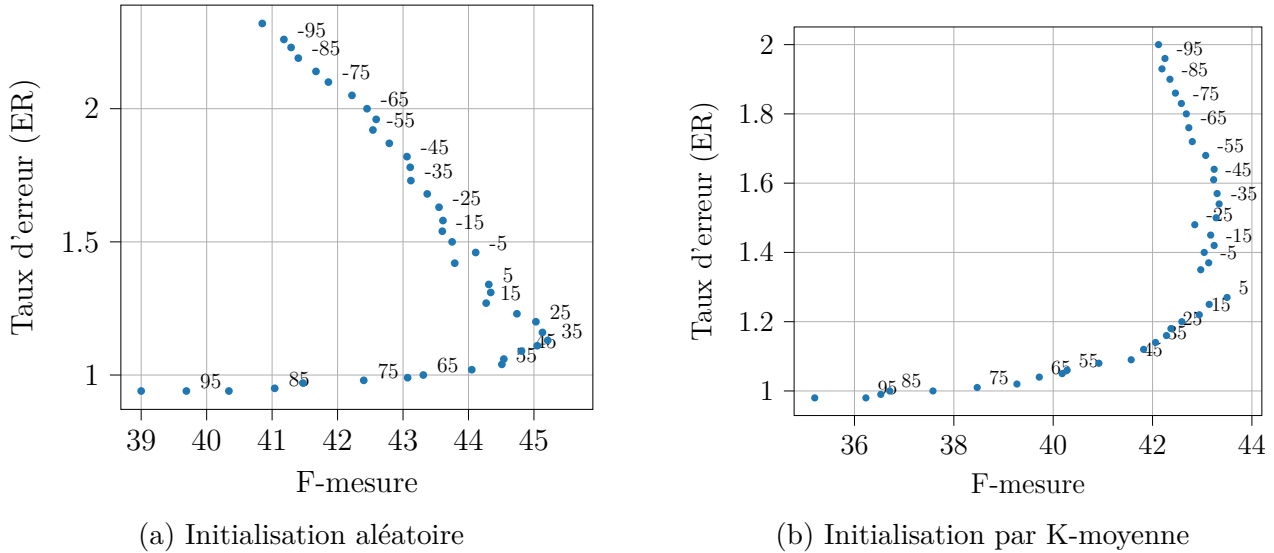


FIGURE 2.10 – Courbe du taux d'erreur en fonction de la F-mesure pour différents seuils de détection allant de -100 à 100 par pas de 5 (pour plus de lisibilité, seul un point sur deux est annoté)

Finalement, on compare les performances obtenues pour chacune de ces méthodes avec leurs seuils optimaux respectifs sur la table 2.8. Pour le micro-moyennage, la différence de performance entre les deux initialisations n'est pas très importante, avec +1.71% de F-mesure et -0.14 de taux d'erreur en faveur de l'initialisation aléatoire. En revanche, elle est davantage significative sur le macro-moyennage, avec +4.09% de F-mesure et -0.22 de taux d'erreur pour cette même initialisation. Le changement dans la F-mesure est principalement dû à la classe *children*, dont la F-mesure passe de 13.3% pour l'initialisation aléatoire à 0% pour l'initialisation par K-moyenne ; le changement du taux d'erreur, lui, est plutôt lié à la classe *large vehicle*, qui augmente pour ces mêmes contextes de 2.63 à 3.41.

	Init. Aléatoire		Init. K-moyenne	
	ER	F-mesure	ER	F-mesure
Total par micro-moyennage	1.13	45.21%	1.27	43.50%
Total par macro-moyennage	1.92	35.24%	2.14	31.15%
<i>brakes squeaking</i>	1.36	18.4%	1.35	15.6%
<i>car</i>	0.74	60.4%	0.74	62.4%
<i>children</i>	3.49	13.3%	3.64	0.0%
<i>large vehicle</i>	2.63	36.9%	3.41	33.0%
<i>people speaking</i>	2.03	30.1%	2.29	25.3%
<i>people walking</i>	1.26	51.9%	1.41	50.5%

TABLE 2.8 – Comparaison entre les taux d'erreur (ER) et F-mesure pour l'initialisation aléatoire (à gauche) avec un seuil de décision de 35, et par K-moyenne (à droite) avec un seuil de 5

Alors que l'initialisation par K-moyenne est instinctivement meilleure, puisqu'elle permet d'équilibrer les centroïdes avant même de commencer l'apprentissage des gaussiennes, il semble pourtant que l'initialisation aléatoire l'emporte pour toutes les classes. Ceci peut néanmoins s'expliquer par le fait que la K-moyenne ne prenne pas en compte la covariance des groupements d'individus. Or dans un GMM, celle-ci a un rôle critique que nous avons déjà abordé. Dans la suite des expériences, on choisira alors une initialisation aléatoire avec un seuil de détection à 35.

2.4.3 Etude du comportement

Nous avons établi le bon fonctionnement du GMM soustractif, ainsi que le seuil de décision optimal pour maximiser ses performances à paramètres fixés. D'autres paramètres peuvent néanmoins modifier le comportement du classifieur, comme le type de la matrice de covariance, le nombre de gaussiennes... En parallèle, la plupart de ces paramètres jouent un rôle primordial dans la complexité du test effectué par le classifieur, complexité qui déterminera ultérieurement la consommation d'énergie et le matériel d'un système de classification embarquée. Par exemple, l'entraînement d'un grand nombre de gaussiennes pour chaque GMM augmente la complexité du système à chacune de ses utilisations alors que la nature de l'initialisation, qui n'a lieu que pendant la phase d'apprentissage, ne l'augmentera pas (on estime que les paramètres, après la phase d'apprentissage, sont stockés dans la mémoire du système).

Si nous n'évaluerons pas dans ces travaux la complexité et la consommation du système de classification, nous étudions néanmoins dans cette section les paramètres qui les influencent lors de la phase de test, à savoir le nombre de gaussiennes et le type de la matrice de covariance des GMM, et la normalisation des descripteurs.

Choix du nombre de gaussiennes

Dans un premier temps, on va s'intéresser au nombre de composantes gaussiennes n_c dans les GMM. Ce paramètre, qui définit la précision de la modélisation des groupements d'individus, doit être choisi avec attention. En effet, si le nombre de gaussiennes est trop faible, celles-ci risquent d'avoir des difficultés à prendre en compte une grande quantité d'individus voire se retrouver entre deux groupements distincts ; mais si ce nombre est trop élevé, le moindre changement dans la base de test pourra ne pas être repéré en raison du sur-apprentissage. Le nombre optimal de composantes peut être déterminé à l'aide de plusieurs méthodes : par critère d'information (Akaike, bayésien...), par mélange de processus de Dirichlet ou par optimisation utilitaire [196]. Face à la complexité d'optimiser plusieurs GMM simultanément avec le même critère, on choisit cette dernière méthode, qui consiste à déterminer simplement la valeur de n_c qui maximise les performances que l'on souhaite mesurer.

Jusqu'à présent, nous utilisons des GMM à quatre composantes, qui sont connus dans la littérature pour être relativement efficaces. Souhaitant rester à un niveau de complexité relativement faible, on choisit d'observer les performances de classification pour un nombre de composantes compris entre 1 et 10. On conserve également un seuil de décision constant de 35 car, en observant les courbes de taux d'erreur en fonction de la F-mesure et du seuil pour chaque valeur de n_c , il semble que ce seuil reste un bon compromis entre les deux mesures. Les résultats sont ainsi présentés pour les valeurs de n_c considérées sur la figure 2.11. Ceux-ci sont sans équivoque : la F-mesure est bien meilleure pour $n_c = 4$ que pour les autres valeurs. Si la modification du seuil pourrait éventuellement accroître la F-mesure, elle aurait également

pour effet d'augmenter le taux d'erreur, qui est déjà équivalent à celui des GMM à quatre composantes pour $n_c > 6$, et supérieur pour les autres. Le choix du nombre de gaussiennes va donc se porter à 4.

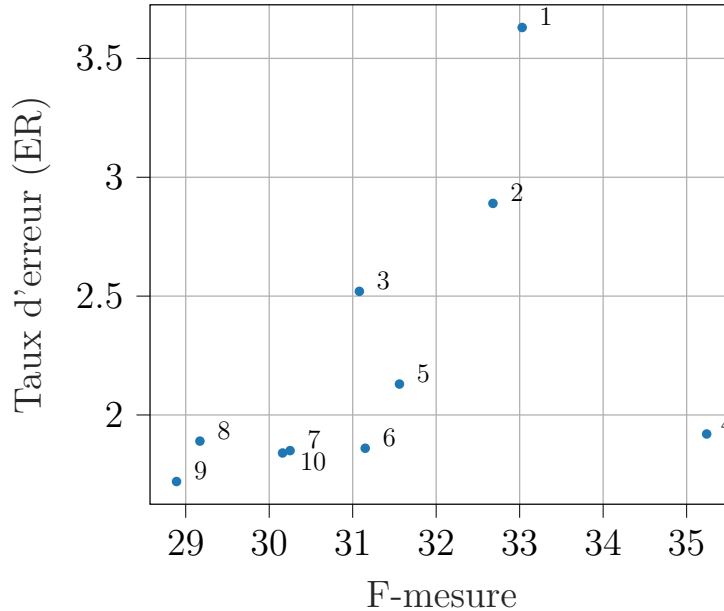


FIGURE 2.11 – Courbe du taux d'erreur en fonction de la F-mesure pour différentes valeurs du nombre de composantes gaussiennes allant de 1 à 10

Influence du type de la matrice de covariance

Dans un second temps, on étudie la forme de la matrice de covariance : celle-ci peut être pleine, ou considérée comme une matrice diagonale. Lorsqu'elle est pleine, pour n descripteurs, la complexité spatiale de la matrice est de n^2 , et la complexité calculatoire suit cette grandeur. Or nous savons que la matrice est censée être strictement définie positive [196] et que, dans le cas idéal où les descripteurs seraient effectivement indépendants les uns des autres, les valeurs qui ne sont pas sur la diagonale sont nulles ou très faibles. En faisant cette hypothèse d'indépendance, on peut choisir de simplifier la matrice à sa diagonale : on stocke alors seulement les n de la diagonale. L'inversion de la matrice de covariance et le calcul de son déterminant sont alors simplifiés, ce qui réduit la complexité d'expressions comme la vraisemblance. En pratique, les descripteurs ne sont pas parfaitement indépendants et, l'approximation étant partiellement correcte, les performances en seront dégradées. Cependant, dans le cadre de la réduction de la complexité de notre programme, il est intéressant de savoir à quel point les résultats sont diminués ; si l'impact n'est pas trop important, un compromis pourrait être accepté.

Dans cette expérience, on reprend les paramètres de la table 2.6, excepté le paramètre covtype que l'on fixe soit à *"full"* pour une matrice pleine, soit à *"diag"* pour une matrice diagonale. On présente les courbes de taux d'erreur en fonction de la F-mesure par micro-moyennage pour différents seuils allant de -100 à 100 par pas de 5 sur la figure 2.12a pour la matrice de covariance pleine et sur la figure 2.12b pour la matrice diagonale. La plus évidente observation à faire est la différence entre les performances des deux méthodes. En effet, si les taux d'erreur sont légèrement supérieurs pour la matrice diagonale, la F-mesure est elle bien

inférieure. Pour la matrice pleine, on a déjà choisi un seuil de décision de 35 ; pour la matrice diagonale, on se propose à nouveau de choisir le seuil permettant la plus grande F-mesure, qui est ici de -40 : on remarque d'ores et déjà une différence importante entre ces seuils. On note alors les performances par micro et macro-moyennage correspondantes dans la table 2.9. Par micro-moyennage, on confirme que la F-mesure chute de 9.0% et le taux d'erreur grimpe de 1.33 – ce qui représente plus que le double – pour la matrice diagonale. Par macro-moyennage, l'effet est moins important pour la F-mesure qui passe seulement de 35.24% à 30.24%, soit une perte de 5.0%, mais bien plus pour le taux d'erreur, qui augmente de 2.97. En simplifiant la matrice de covariance à une simple diagonale, on perd en effet des subtilités qui étaient importantes pour certaines classes. En particulier, la chute brutale de la F-mesure est liée à la baisse de performance de la seconde classe la plus importante, *people walking*, dont la F-mesure passe de 51.9% à 43.3%, et à la classe minoritaire dont la F-mesure passe de 13.3% à 4.7%. Le taux d'erreur, lui, augmente fortement pour toutes les classes sauf celle majoritaire, *car*. En effet, pour obtenir une F-mesure convenable, on a abaissé le seuil de décision, ce qui conduit à un bien plus grand nombre d'erreurs.

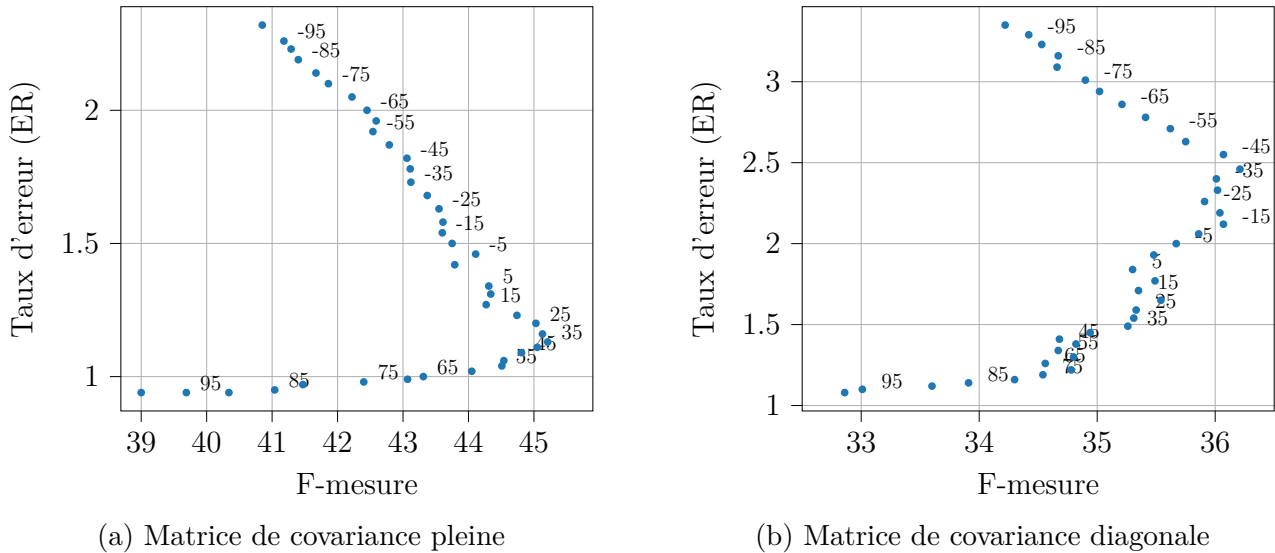


FIGURE 2.12 – Courbe du taux d'erreur en fonction de la F-mesure pour différents seuils de détection allant de -100 à 100 par pas de 5 et pour une matrice de covariance pleine ou diagonale

Finalement, la diminution des performances avec une matrice de covariance diagonale par rapport à une pleine est indéniable. Cependant, le gain en terme de complexité l'est tout autant. En fonction de notre intérêt, il n'est donc pas inintéressant de conserver l'idée d'une matrice de covariance diagonale, notamment en sachant qu'un travail sur les descripteurs permettra peut-être de fortifier l'hypothèse d'indépendance, et ainsi d'augmenter les performances.

Importance de la normalisation des descripteurs

Notre dernière expérience concerne la normalisation des descripteurs. Il s'agit d'une opération réalisée pour remettre toutes les valeurs des différentes variables sur une même échelle, et qui peut s'opérer de plusieurs manières. La plus connue - et celle que l'on va utiliser - est la normalisation centrée réduite, qui permet de réduire les variables à une loi normale par

	Matrice pleine		Matrice diagonale	
	ER	F-mesure	ER	F-mesure
Total par micro-moyennage	1.13	45.21%	2.46	36.21%
Total par macro-moyennage	1.92	35.24%	4.89	30.24%
<i>brakes squeaking</i>	1.36	18.4%	6.77	17.5%
<i>car</i>	0.74	60.4%	0.77	64.8%
<i>children</i>	3.49	13.3%	9.96	4.7%
<i>large vehicle</i>	2.63	36.9%	6.11	24.4%
<i>people speaking</i>	2.03	30.1%	3.38	26.7%
<i>people walking</i>	1.26	51.9%	2.37	43.3%

TABLE 2.9 – Comparaison entre les taux d’erreur (ER) et F-mesure pour une matrice de covariance pleine (à gauche) avec un seuil de décision de 35, et diagonale (à droite) avec un seuil de -40

l’opération

$$x_{\text{norm}} = \frac{x - \mu_x}{\Sigma_x} \quad (2.18)$$

avec μ_x la moyenne vectorielle des descripteurs d’apprentissage et Σ_x la matrice de covariance calculée sur ces mêmes descripteurs. En contrôlant ainsi la moyenne et la variance, on espère faciliter la discrimination au sein de chaque GMM, et donc des classes.

On considère dans un premier temps des GMM munis d’une matrice pleine et des paramètres détaillés précédemment. Sur la figure 2.13a, on retrouve les performances par micro-moyennage à l’issue de la classification sur la BDD de test pour différents seuils. Le seuil de décision choisi est toujours 35. Sur la figure 2.13b, on observe à présent les performances du GMM soustractif avec normalisation. Le seuil de décision à privilégier est toujours 35 en raison d’une combinaison entre une haute F-mesure et un bas taux d’erreur. En revanche, on observe que, quel que soit le seuil, la F-mesure est toujours plus basse et le taux d’erreur toujours plus haut pour le GMM soustractif entraîné sur des descripteurs normalisés. Nous pouvons en déduire que la seule distribution des descripteurs n’est pas forcément suffisante, et que les différences de niveaux jouent un rôle, même faible. Une rapide analyse des résultats avec ou sans normalisation et un seuil de décision de 35, présentés sur la table 2.10, nous confirme que la baisse de performance avec normalisation des descripteurs est assez faible mais présente, avec des classes majoritaires moins touchées que les classes minoritaires. Seule la classe *brakes squeaking* gagne légèrement en F-mesure.

On réitère dans un second temps une expérience analogue avec une matrice diagonale. Néanmoins, on remarque sur la figure 2.14 que les conclusions sont similaires au cas de la matrice pleine, à l’exception du seuil de décision qui peut changer (on conseille alors de prendre 20 si on privilégie le taux d’erreur, et -35 si on privilégie la F-mesure).

Alors que la normalisation des descripteurs semblait être un passage obligé, il semble que les performances soient similaires voire dégradées avec un tel système. Comme nous utilisons les MFCC, il est possible que la variabilité de leurs valeurs soit importante dans le processus d’apprentissage. Par la suite, on ne conservera donc pas cette opération dans notre algorithme.

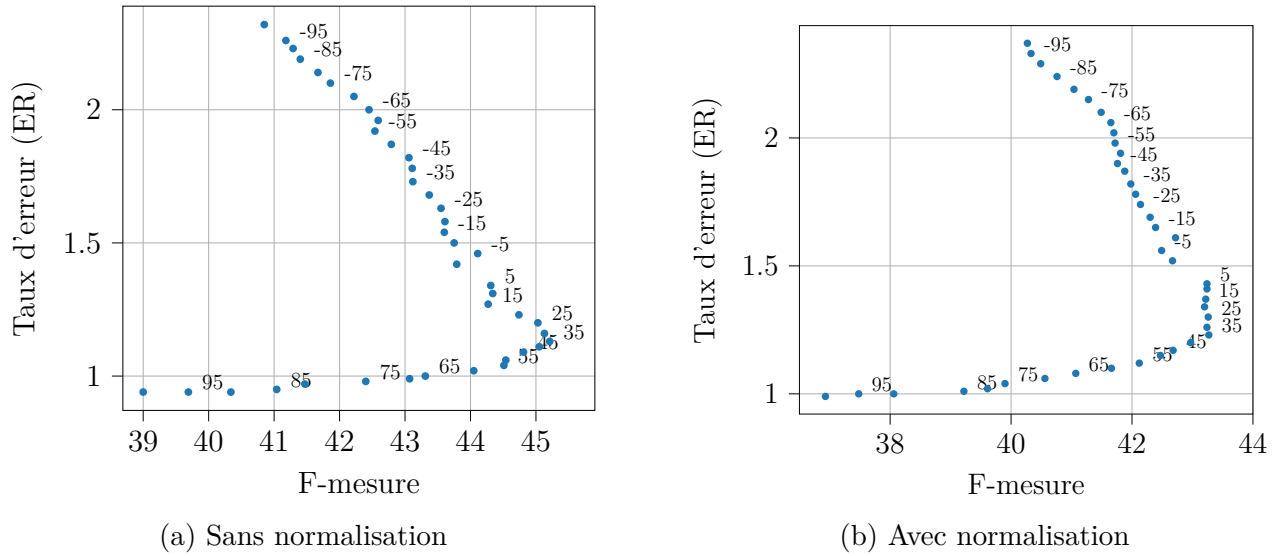


FIGURE 2.13 – Courbe du taux d'erreur en fonction de la F-mesure pour différents seuils de détection allant de -100 à 100 par pas de 5, sans ou avec normalisation pour une matrice pleine

	Sans normalisation		Avec normalisation	
	ER	F-mesure	ER	F-mesure
Total par micro-moyennage	1.13	45.21%	1.23	43.27%
Total par macro-moyennage	1.92	35.24%	2.06	32.50%
<i>brakes squeaking</i>	1.36	18.9%	1.35	20.1%
<i>car</i>	0.74	60.4%	0.74	60.4%
<i>children</i>	3.49	13.3%	3.47	3.7%
<i>large vehicle</i>	2.63	36.9%	3.02	34.2%
<i>people speaking</i>	2.03	30.1%	2.36	25.7%
<i>people walking</i>	1.26	51.9%	1.42	50.8%

TABLE 2.10 – Comparaison entre les taux d'erreur (ER) et F-mesure pour des descripteurs non normalisés (à gauche) avec un seuil de décision de 35, et normalisés (à droite) avec un seuil de 35

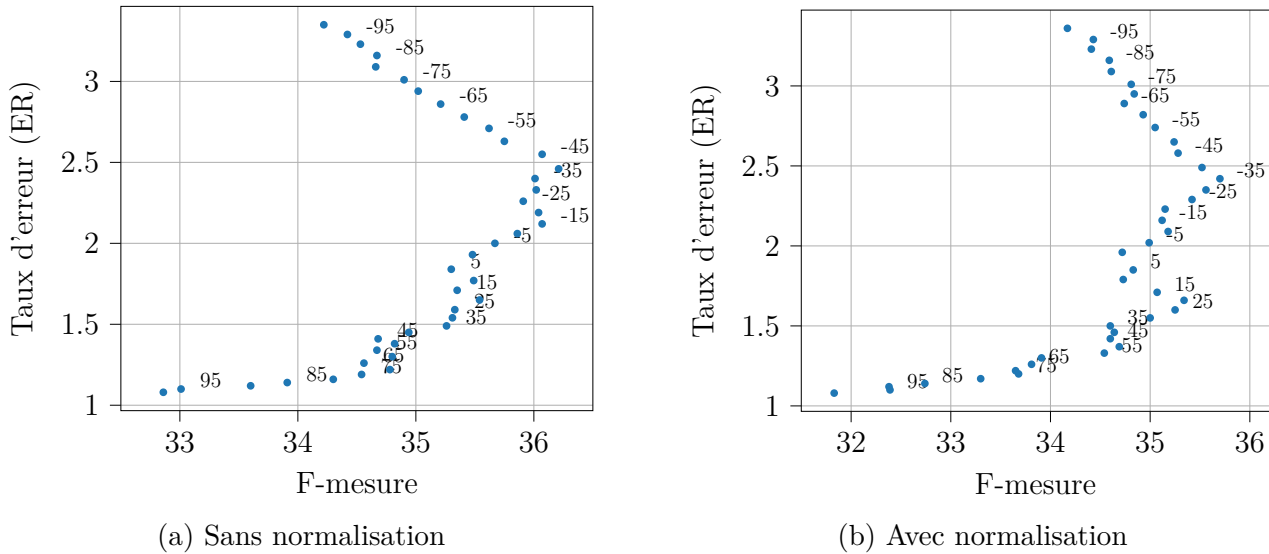


FIGURE 2.14 – Courbe du taux d’erreur en fonction de la F-mesure pour différents seuils de détection allant de -100 à 100 par pas de 5, sans ou avec normalisation pour une matrice diagonale

2.5 Le réseau neuronal convolutif et récurrent

Dans cette section, nous présentons le second classifieur utilisé dans nos travaux, à savoir un CRNN. Celui introduit ici est inspiré de la référence des tâches de détection des challenges DCASE 2019 et 2020 [191], et qui est encore régulièrement utilisé comme modèle. Comme on l’a indiqué dans la section 2.3.2, son entrée n’est non pas le vecteur des MFCC et de ses dérivées mais le mel-spectrogramme associé. Comme pour le GMM, nous cherchons principalement ici à reproduire l’état de l’art.

On commencera par décrire le fonctionnement de ce CRNN et l’intérêt de ses différentes couches. Par la suite, comme pour le GMM, on l’implémentera et on vérifiera qu’il fonctionne. Enfin, on étudiera le comportement du CRNN en relation avec différents paramètres qui influent sur la complexité de la phase de test du système.

2.5.1 Fonctionnement théorique d’un réseau neuronal convolutif et récurrent

A l’origine, les réseaux connexionnistes ont pour but de modéliser les systèmes cérébraux [38]. Ils se sont rapidement avérés très performants pour tout type de classification et se sont répandues en tant que méthode majoritaire en apprentissage automatique. Pour la classification acoustique, qui cherche à simuler l’écoute et la compréhension sonore humaine, il était donc logique d’utiliser également les réseaux neuronaux. Dès 2016, ils deviennent majoritaires dans les solutions proposées aux différentes tâches du challenge DCASE [128]. En 2017, le gagnant de la tâche 3 de détection en environnement réel utilise un CRNN qui s’impose peu à peu [3].

Le CRNN que nous utilisons est celui présenté sur les figures 2.15 et 2.16, issu du gagnant de la tâche 3 du challenge DCASE 2017 [3] et repris dans les éditions suivantes. Il est constitué de trois blocs convolutifs suivis de deux couches récurrentes. L’entrée de ce réseau est un ensemble de matrices constituées par les descripteurs d’une succession de n_{seq} trames, c’est-à-dire les

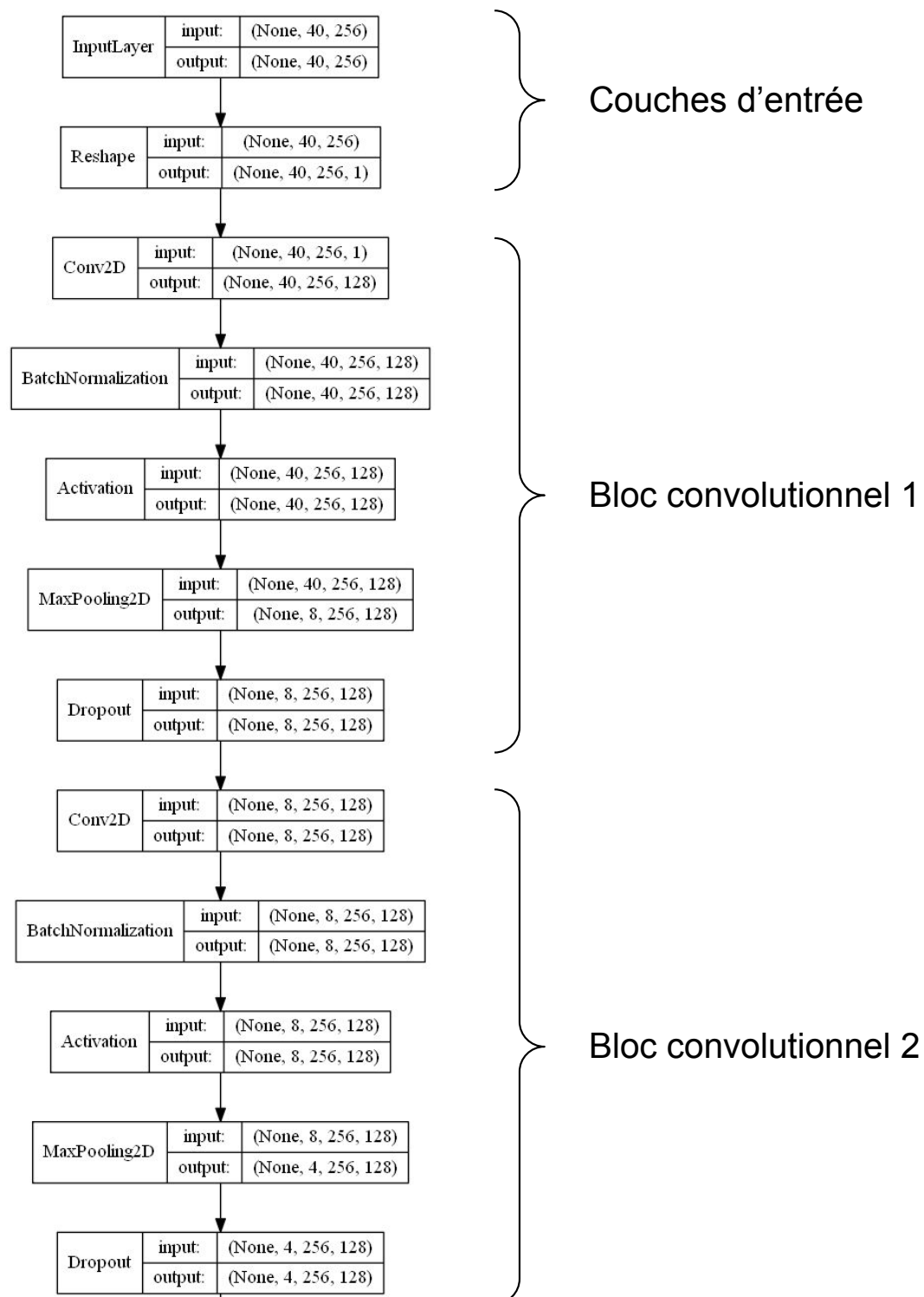


FIGURE 2.15 – Graphique des différentes couches du modèle du CRNN (partie 1)

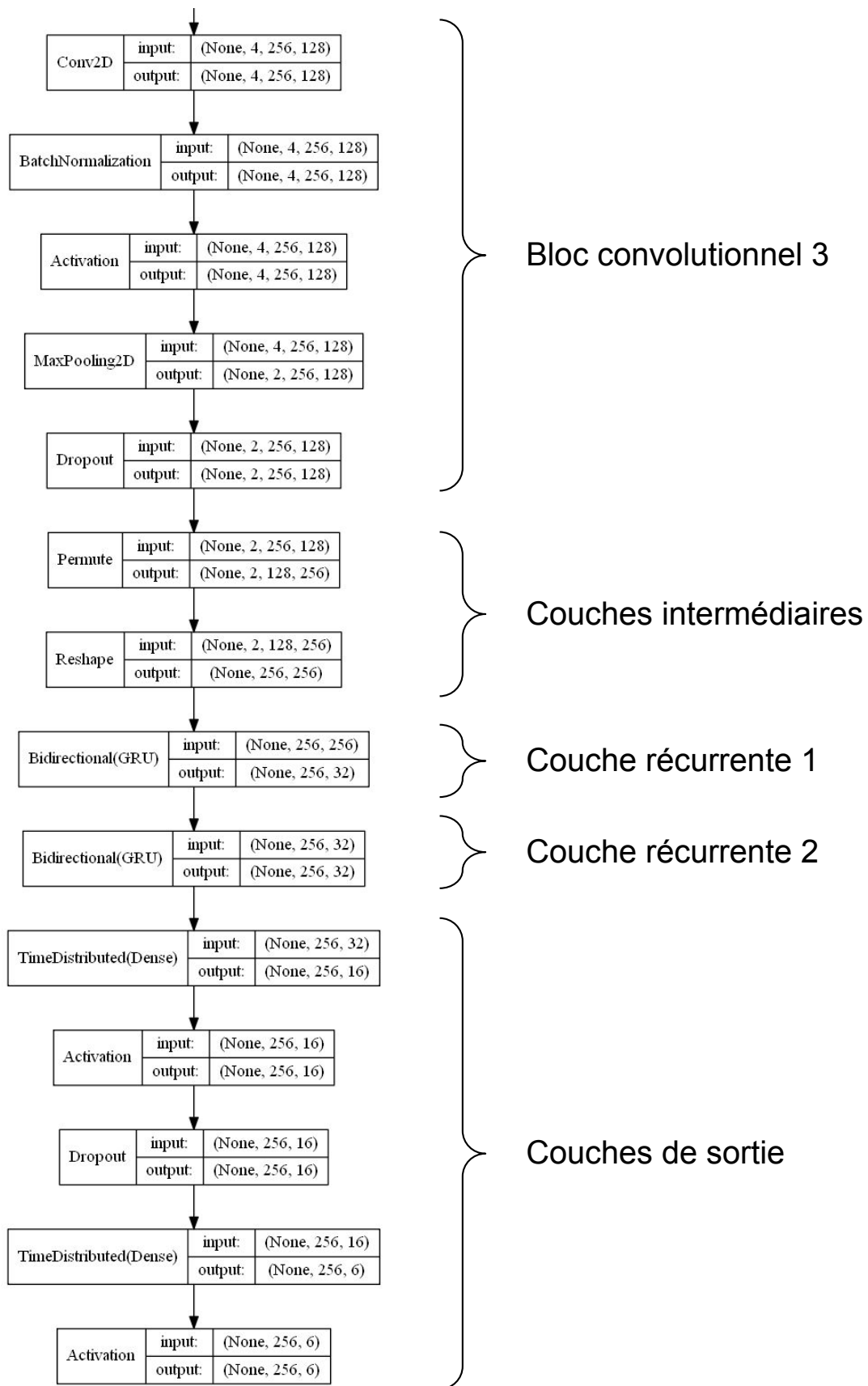


FIGURE 2.16 – Graphique des différentes couches du modèle du CRNN (partie 2)

vecteurs des 40 coefficients du mel-spectrogramme. À noter que ces descripteurs sont normalisés pour être centrés et réduits [26, 50]. Cet ensemble de matrices est découpé en plusieurs lots (*mini-batch* en anglais) qui sont successivement présentés en entrée du réseau de neurones. Les paramètres de celui-ci sont mis à jour après le passage de chaque lot.

Les blocs convolutifs L'utilisation de plusieurs trames permet aux couches convolutives d'apprendre à percevoir non seulement les informations importantes à chaque trame, mais également la relation avec les précédentes et les suivantes, comme le faisaient les Δ MFCC et Δ^2 MFCC pour le GMM. Ce phénomène est particulièrement utile pour les signaux audio qui présentent une corrélation temporelle court-terme, c'est-à-dire sur un petit nombre de trames successives [196]. Au sein de notre réseau, chaque bloc convolutif est formé par une succession de plusieurs couches indiquées sur les figures 2.15 et 2.16 :

Une convolution 2D Cette première couche réalise une convolution entre un filtre de taille t_f et une portion de même dimension de l'entrée (en glissant à chaque fois de n_p pixels sur l'image), et répète cette opération avec n_f filtres différents. La succession des couches convolutives entraîne un apprentissage de plus en plus fin sur les données [38]. À noter que, pour la première couche, le signal est remodelé par une couche nommée *Reshape* qui indique que nous n'utilisons qu'un seul canal (la couche convolutive peut théoriquement en traiter plusieurs).

Une normalisation *batch* Afin d'accélérer l'apprentissage grâce à des données correctement mises en forme, une normalisation est appliquée à la sortie de la couche convolutive. Il s'agit ici d'une normalisation par *batch*, c'est-à-dire que les données sont centrées et réduites par rapport à une moyenne et une variance calculées sur l'ensemble des données d'un même *mini-batch*.

Une couche d'activation Une couche d'activation non linéaire est ensuite appliquée pour activer seulement certains neurones du réseau. On utilise ici la fonction d'activation *rectified linear unit* (ReLU), qui a prouvé son efficacité dans les réseaux sans pré-entraînement [38].

Un *MaxPooling* en 2D Aussi appelée couche d'agrégation et de sous-échantillonnage, elle a pour but de réduire la taille des entrées tout en assurant leur robustesse au bruit et distorsions. Pour cela, on découpe chaque entrée à l'aide d'une fenêtre glissante (en glissant à chaque fois de n_m pixels sur l'image), et une valeur unique est calculée puis conservée à partir de ces portions d'entrée. Dans le cas du *MaxPooling*, on conserve la valeur maximale d'une fenêtre de taille t_m .

Une couche *Dropout* La dernière couche de ce bloc est celle de *dropout*, qui consiste à désactiver temporairement certains neurones pour permettre un meilleur apprentissage pour les autres (notamment les moins utilisés), ce qui permet d'éviter le surapprentissage. Cette couche est caractérisée par le taux t_d de neurones désactivés.

Les paramètres utilisés pour ce bloc sont indiqués dans la table 2.11.

Les couches récurrentes Les convolutions sont suivies de deux couches récurrentes qui sont des *gated recurrent unit* (GRU) bidirectionnelles. L'intérêt des couches récurrentes est de prendre en compte la temporalité des descripteurs. Contrairement aux couches convolutives, celles-ci déduisent les relations entre trames sur le long terme. Dans le réseau utilisé, les deux couches récurrentes ont deux propriétés :

Paramètre	Description	Valeur
t_f	Taille des filtres convolutifs	(3×3)
n_p	Fenêtre glissante de la couche convolutive	1
n_f	Nombre de filtres	128
n_m	Fenêtre glissante du <i>pooling</i>	1
t_m	Taille d'une fenêtre du <i>pooling</i>	(2×1)
t_d	Taux de <i>dropout</i>	0.4

TABLE 2.11 – Paramètres utilisés dans chaque bloc convolutif

Couches GRU Particulièrement adaptées aux longues séquences [38], leur sortie est calculée à partir de l'entrée actuelle, mais aussi des précédentes. Les neurones que traversent l'entrée actuelle et les neurones récurrents subissent également un *dropout* de taux t_a et t_r , que l'on fixe ici à 0.5 (*i.e.* on en désactive la moitié). Comme la couche convolutive, une fonction d'activation non linéaire est placée en sortie. Il s'agit ici d'une tangente hyperbolique, qui permet d'obtenir des valeurs de sortie entre -1 et 1.

Bidirectionnelles Les hyper-paramètres du réseaux apprennent des trames passées, mais aussi des trames futures, ce qui implique de travailler sur des séquences finies (les *mini-batch*).

Couches convolutives et récurrentes sont finalement complémentaires puisque les unes permettent d'extraire des descripteurs locaux, tandis que les autres déterminent des descripteurs globaux, c'est-à-dire sur une plus longue plage temporelle [132].

Les couches de sortie Plusieurs couches terminent le réseau, comme on peut le voir sur la figure 2.16.

Les couches denses Il s'agit de couches dans lesquelles tous les neurones sont connectés, sans aucun *dropout*. Ce sont les couches les plus simples et les plus utilisées. Comme on le remarque sur la figure 2.16, la première diminue la taille des matrices de $n_{seq} \times 32$ en entrée à $n_{seq} \times 16$ en sortie. La seconde diminue encore cette taille pour parvenir à une matrice indiquant les probabilités de chaque classe pour chaque trame, soit ici une taille de $n_{seq} \times 6$.

TimeDistributed Les deux couches denses utilisées sont enveloppées dans une couche *TimeDistributed* qui permet d'appliquer la couche à chacune des n_{seq} trames indépendamment.

Une couche Dropout Nous avons déjà décrit le fonctionnement d'une couche de *dropout*. Ici, le taux d'abandon est fixé à 0.1.

Une couche d'activation La dernière couche est une couche d'activation employant une fonction sigmoïde. Celle-ci, de forme similaire à la tangente hyperbolique, permet de répartir les données d'entrée entre 0 et 1, ce qui nous permet d'obtenir les probabilités de sortie.

Après un éventuel lissage temporel de la sortie, il est possible de seuiller ces probabilités pour obtenir nos décisions (0 : absence de la classe sur la trame ; 1 : présence de la classe). Ce seuil est généralement fixé à 0.5.

Lors de l'apprentissage, on cherche à optimiser les hyper-paramètres afin d'obtenir les probabilités les plus justes possibles. Pour cela, on minimise une fonction de coût par descente

du gradient. Ici, la descente du gradient est réalisée par la méthode Adam (pour *Adaptive Momentum estimation*). La fonction de coût, quant à elle, est l'entropie croisée catégorielle, que l'on utilise pour les tâches de classification multi-labels. En posant C le nombre de classes, et $y = (y_1, \dots, y_C)$ et $\hat{y} = (\hat{y}_1, \dots, \hat{y}_C)$ respectivement les étiquettes réelles et estimées d'une trame, l'entropie croisée catégorielle d'une trame est donnée par :

$$\mathcal{L} = - \sum_{i=1}^C y_i \ln(\hat{y}_i). \quad (2.19)$$

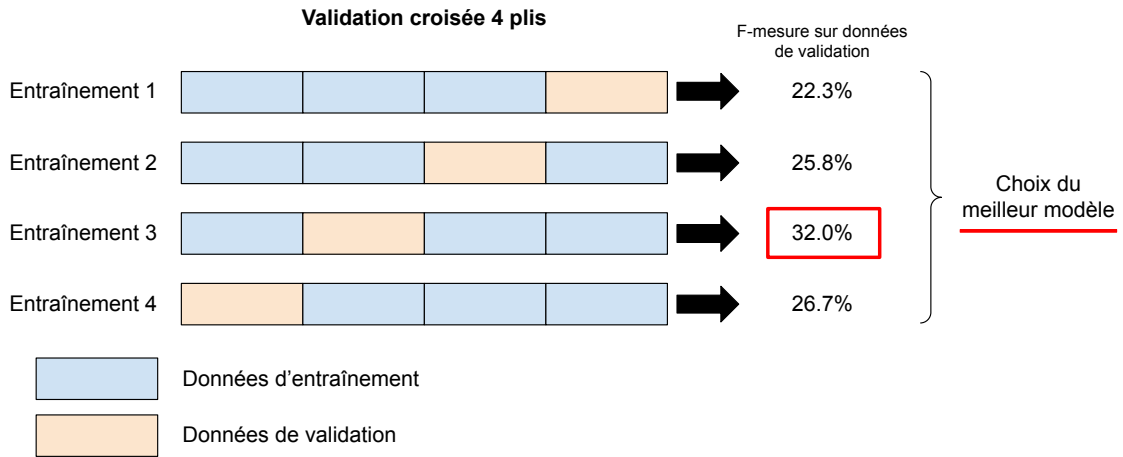


FIGURE 2.17 – Fonctionnement d'une validation croisée à 4 plis

Dans la suite de nos expériences, on fixe les paramètres aux valeurs de la table 2.12. La présence ou non de la normalisation des descripteurs est définie par le paramètre `norm`. Les données sont ensuite découpées en matrices de taille $40 \times n_{\text{seq}}$ avec 40 le nombre de descripteurs (*i.e.* le nombre de coefficients du mel-spectrogramme) et n_{seq} le nombre de trames prises en compte. Chaque lot est constitué de n_{batch} de ces matrices. Les lots constituent la base de développement de notre réseau ; on réalise alors un apprentissage à n_{fold} plis (comme décrit sur la figure 2.17) sur n_{epoch} époques. En raison de la variabilité des résultats en sortie des réseaux neuronaux, on effectue cet apprentissage n_{test} fois et on moyennera les performances obtenues lors des n_{test} tests pour déterminer les résultats de notre système.

Paramètre	Description	Valeur
n_{seq}	Nombre de trames d'une entrée	256
n_{batch}	Taille des lots	32
n_{epoch}	Nombres d'époques de l'entraînement	100
<code>norm</code>	Normalisation	vrai
n_{fold}	Nombre de plis	4
n_{test}	Nombre de tests réalisés	5

TABLE 2.12 – Paramètres utilisés dans le CRNN

2.5.2 Implémentation et vérification de l'algorithme

Dans cette section, on s'intéresse à vérifier le fonctionnement du CRNN. Comme pour le GMM, celui-ci est implémenté en Python 3 à l'aide des bibliothèques *dcase_util*, *librosa* et *sed_eval*. Cette fois-ci, on utilise en complément la bibliothèque *keras* pour l'apprentissage du réseau [37]. Les réseaux de neurones sont bien plus complexes que les GMM et comprennent des milliers d'hyper-paramètres (ici 366 582 précisément), il ne sera donc pas possible de les détailler.

Afin de s'assurer du fonctionnement du CRNN, on vérifiera que ses performances correspondent à celles de l'état de l'art. On s'intéressera ensuite à comprendre l'influence du nombre d'époques et du seuil de décision sur les performances.

Comparaison des performances du CRNN avec l'état de l'art

Notre système est une réimplémentation de celui proposé par [3]. Celui-ci propose le CRNN décrit dans la section 2.5.1 entraîné sur 100 époques. Les résultats présentés dans [3] et détaillés sur la page officielle de la tâche 3 du challenge DCASE 2017 sont obtenus en moyennant les performances de cinq tests successifs. Nous reprenons cette méthode pour comparer nos propres résultats à l'état de l'art.

La comparaison entre les performances obtenues par simulation et celles des références peut être réalisée à partir de la table 2.13. Les performances sont relevées pour un seuil de décision de 0.5 au bout de 100 époques d'entraînement. Les F-mesures globales connaissent des différences aux alentours de 2% entre l'implémentation et la référence. Celle obtenue pour notre implémentation par micro-moyennage est ainsi de 2.6% inférieure à la référence, et celle par macro-moyennage est supérieure de 1.5%. Ces modifications sont liées à la différence de reconnaissance entre les classes. En effet, la classe majoritaire *car* connaît une baisse de F-mesure de 3.1% et la classe *large vehicle* de 19.1%, ce qui conduit à une F-mesure micro-moyennée plus basse. En revanche, la classe *people walking* est un peu mieux détectée (avec une F-mesure qui croît de 7.1% par rapport à la référence), tandis que la classe *people speaking* passe tout simplement de 0.0% à 23.0% de F-mesure, ce qui mène à une augmentation du critère macro-moyenné. Dans le même temps, on observe une augmentation globale du taux d'erreur pour notre implémentation, plus ou moins forte selon les classes. Ainsi, si la classe *car* connaît une hausse de seulement 0.07 du taux d'erreur, la classe *large vehicle* passe de 1.07 pour la référence à 2.87 pour notre implémentation ; ceci est à mettre en relation avec la baisse de F-mesure, ce qui nous amène à considérer une forte baisse de performance pour cette classe. Ces différences sont probablement liées à la longueur du segment utilisé pour la décision. En effet, au vu de la longueur des plus courts événements (en dessous de 2s pour certaines classes d'après les tables 2.2 et 2.3), nous avons choisi de conserver des segments de 1s, comme pour le GMM soustractif, alors que [3] utilise une longueur de 2.5s.

Finalement, malgré certaines différences notables dans les performances des deux implémentations, on admet que notre implémentation est conforme aux attentes.

Validation des paramètres

Nous avons considéré jusqu'à présent un entraînement sur 100 époques car, d'après [3], le taux d'erreur obtenu ensuite ne décroît plus. Néanmoins, si les performances sont acceptables pour la base de développement, on ne sait pas réellement ce qu'il en est pour la base de test.

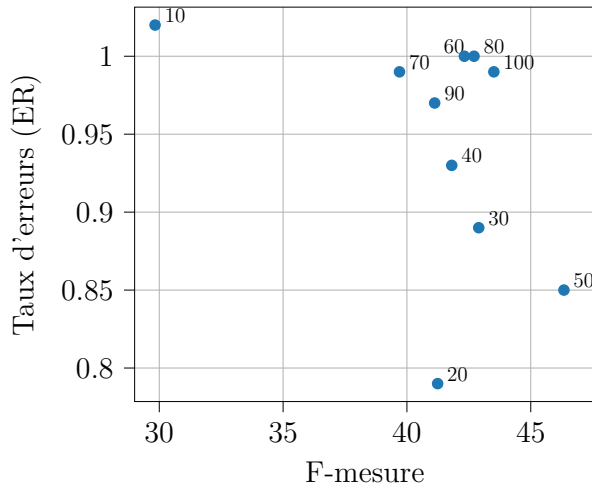
	Implémentation		Référence	
	ER	F-mesure	ER	F-mesure
Total par micro-moyennage	1.07	39.1%	0.79	41.7%
Total par macro-moyennage	1.67	25.2%	1.02	23.7%
<i>brakes squeaking</i>	1.0	0.0%	1.0	0.0%
<i>car</i>	0.84	51.5%	0.77	54.6%
<i>children</i>	1.97	0.7%	1.20	0.0%
<i>large vehicle</i>	2.87	30.1%	1.07	49.3%
<i>people speaking</i>	1.92	23.0%	1.04	0.0%
<i>people walking</i>	1.43	45.8%	1.03	38.7%

TABLE 2.13 – Comparaison entre les taux d'erreur (ER) et F-mesure de notre implémentation (à gauche) et de l'implémentation de référence (à droite) pour un seuil de décision de 0.5

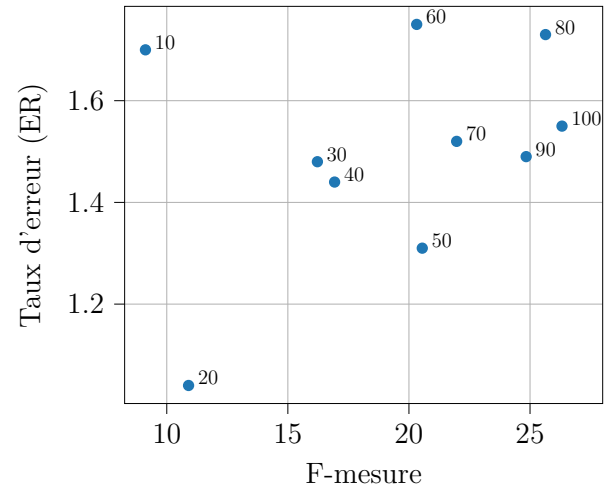
En effet, nous avons vu dans la section 2.3.1 que celle-ci a des caractéristiques quelques peu différentes de la base d'apprentissage, notamment en terme de longueur des événements ; on peut alors imaginer que le problème de sur-apprentissage intervient avant ces 100 époques.

Notre première expérience consiste donc à vérifier sur la base de test le nombre d'époques qu'il vaut mieux adopter pour maximiser les performances. Les résultats de cinq tests consécutifs sont moyennés comme précédemment. Le taux d'erreur moyen en fonction de la F-mesure moyenne, le tout micro-moyenné, est affiché sur la figure 2.18a pour des époques allant de 10 à 100 par pas de 10. Comme attendu, on remarque que les plus faibles performances sont obtenues pour 10 époques d'entraînement seulement. Au fur et à mesure de l'entraînement, on observe que les performances varient grandement avant de commencer à se stabiliser autour d'une même valeur à partir de 60 époques. Néanmoins, cette tendance intervient après que le réseau ait atteint sa plus forte valeur de F-mesure et son deuxième plus bas taux d'erreurs, au bout de 50 époques. Les performances présentes étant calculées à l'issue de l'entraînement de plusieurs CRNN, il ne s'agit probablement pas ici d'une coïncidence : on pourrait alors imaginer qu'il s'agisse de l'impact d'un sur-apprentissage. Néanmoins, la figure 2.18b indique que ce n'est pas le cas, car la F-mesure macro-moyennée continue d'augmenter. Finalement, on confirme l'entraînement de notre réseau sur $n_{\text{epoch}} = 100$ époques.

Le nombre d'époques pour l'entraînement fixé, nous nous intéressons, comme pour le GMM soustractif, à déterminer le seuil de décision optimal. Actuellement celui-ci est, d'après la théorie, fixé à 0.5. Néanmoins, la base de données étant déséquilibrée, le seuil peut être amené à être sous-optimal. La sortie étant exprimée sous forme de probabilité de présence, notre seconde expérience consiste donc à faire varier le seuil de décision entre 0 et 1 par pas de 0.05. Les performances moyennes à l'issue de cinq tests sont affichées sur la figure 2.19a pour le micro-moyennage et 2.19b pour le macro-moyennage. Pour plus de lisibilité, seul un point sur deux est annoté. Sur les deux figures, lorsque le seuil augmente, on observe d'abord une forte décroissance du taux d'erreur ; celui-ci se stabilise à partir d'un certain seuil et, au-delà, c'est la F-mesure qui décroît rapidement. En regardant attentivement les valeurs à cette jonction, on remarque que le seuil optimal se situe bien autour de 0.5 : non seulement le seuil pratique correspond à la théorie, mais on a une confirmation de la convergence du CRNN. Dans la suite, on conservera donc le seuil de décision de 0.5.

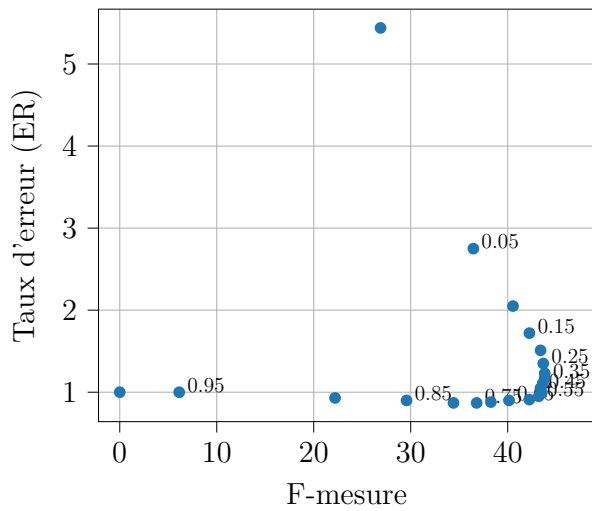


(a) Courbe pour des performances micro-moyennées

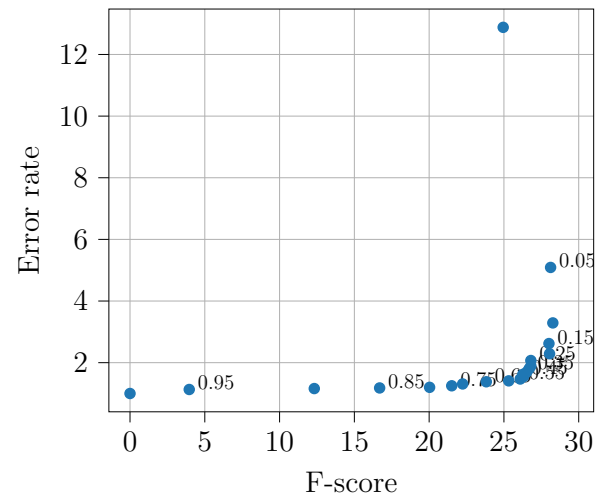


(b) Courbe pour des performances macro-moyennées

FIGURE 2.18 – Courbe du taux d'erreur en fonction de la F-mesure pour différentes époques allant de 10 à 100 par pas de 10



(a) Courbe pour des performances micro-moyennées



(b) Courbe pour des performances macro-moyennées

FIGURE 2.19 – Courbe du taux d'erreur en fonction de la F-mesure micro ou macro-moyenné pour différents seuils de décision allant de 0 à 1 par pan de 0.05 (pour plus de lisibilité, seuil un point sur deux est annoté)

	Avec normalisation		Sans normalisation	
	ER	F-mesure	ER	F-mesure
Total par micro-moyennage	1.07	39.1%	1.07	37.5%
Total par macro-moyennage	1.67	25.2%	1.65	24.5%
<i>brakes squeaking</i>	1.0	0.0%	1.0	0.0%
<i>car</i>	0.84	51.5%	0.83	48.0%
<i>children</i>	1.97	0.7%	1.76	0.0%
<i>large vehicle</i>	2.87	30.1%	2.76	32.9%
<i>people speaking</i>	1.92	23.0%	2.02	22.5%
<i>people walking</i>	1.43	45.8%	1.54	44.7%

TABLE 2.14 – Comparaison entre les taux d'erreur (ER) et F-mesure de notre implémentation avec (à gauche) ou sans normalisation (à droite) pour un seuil de décision de 0.5

2.5.3 Etude du comportement

Nous avons validé le fonctionnement du CRNN, ainsi que le nombre d'époques nécessaires à son entraînement et son seuil de détection. Bien d'autres paramètres permettent de personnaliser ce classifieur, que ce soit la taille des lots, celle des matrices d'entrée, le taux d'apprentissage, la taille des noyaux de chaque couche... En raison du nombre de ces paramètres et de la complexité du CRNN, il nous est matériellement difficile d'étudier en profondeur le comportement de ce réseau. Néanmoins, plusieurs travaux antérieurs ont porté sur l'étude du fonctionnement de tels CRNN, et notamment son comportement pour la classification d'événements sonores [2, 17, 27, 45].

L'étude spécifique des classifieurs n'étant pas notre thématique de recherche, nous nous contentons ici d'étudier l'intérêt de la normalisation des descripteurs. Pour cela, nous moyennons les performances à l'issue de cinq tests sans normalisation, pour un seuil de 0.5 et un entraînement sur 100 époques, que nous comparons avec ceux avec normalisation dans la table 2.14. Globalement, on observe peu de différence entre les deux implémentations. En effet, les performances sont similaires pour la majorité des classes. Seule les classes *car* et *people walking* sont réellement mieux reconnues avec normalisation, avec une augmentation de la F-mesure respectivement de 3.5% et 1.1%; au contraire, la classe *large vehicle* connaît une baisse de 2.8% de sa F-mesure. Finalement, la F-mesure par micro-moyennage augmente de 1.6% avec la normalisation – principalement car la classe majoritaire *car* est mieux reconnue – et celle par macro-moyennage croît de 0.7%. L'augmentation des performances induite par la normalisation n'est donc pas très grande, en particulier en regard de la complexité qu'elle occasionne. Néanmoins, dans la suite de nos expériences, nous choisissons de conserver cette normalisation.

2.6 Conclusion

Ce chapitre a présenté un état de l'art des méthodes et outils utilisés ces dernières années pour la détection d'événements sonores environnementaux. Pour cela, nous avons exposé les principales caractéristiques d'un système de classification acoustique. Nous avons ainsi abordé la nature des différentes classifications existantes, mais aussi celle des sources audio, des signaux et des annotations. Les étapes composant un tel système de classification ont également été décrites. L'une d'entre elle a particulièrement retenu notre attention : il s'agit de l'étape

d'apprentissage, réalisée par un classifieur.

Si celui-ci est souvent le cœur d'un système de classification, d'autres éléments restent indispensables. Nous en avons introduit certains qui nous ont semblé les plus importants : les bases de données, les descripteurs et les critères de performance. A cette occasion, nous avons également indiqué leur utilisation dans l'état de l'art et justifié nos choix pour la suite.

Ayant ainsi pris connaissance de tous les éléments de la chaîne de classification, nous avons pu présenter deux classifieurs bien connus du domaine : le GMM soustractif, accessible grâce à sa faible complexité, et un CRNN, apprécié pour ses performances. Nous avons détaillé le fonctionnement, mais aussi l'implémentation et le comportement de ces deux classifieurs.

L'ensemble de ces connaissances nous permet de disposer d'une base solide et reconnue dans le domaine afin de poursuivre nos travaux sur l'introduction d'un système de classification sonore embarqué. Dans la suite, et particulièrement dans le chapitre 5, il nous sera donc possible d'étudier le comportement d'un système de classification face aux contraintes de l'embarqué.

CONSOMMATION D'ÉNERGIE D'UN SYSTÈME COMMUNICANT

3.1 Introduction

Parmi les applications de la classification d'événements sonores, certaines nécessitent une importante réactivité, et donc un fonctionnement en temps réel. Dans le cas d'un système distribué, les nœuds comprenant les microphones vont communiquer avec les nœuds effectuant les calculs. Afin de respecter les contraintes temps-réel, cette communication doit être suffisamment rapide, ce qui nécessite de connaître ses caractéristiques. En particulier, on cherche à déterminer les éléments affectant le temps de fonctionnement du système. Une de ces caractéristiques est l'efficacité spectrale, qui indique le nombre de données binaires émises sur le canal de communication par ressource temps-fréquence. Les ressources fréquentielles étant souvent fixées par des normes strictes liées à leur usage [35, 36, 160], l'aspect temporel y joue un rôle prédominant. Cependant, transmettre de l'information a un coût qu'il est essentiel de prendre en compte, d'autant plus que les nœuds d'un réseau sans fil sont généralement limités en énergie. Il est donc également nécessaire de savoir comment calculer ce coût énergétique, et quelle influence ont sur lui les paramètres du système. Avec la connaissance de ces éléments, il devient alors possible de déterminer un compromis entre longévité de notre réseau de nœuds et réactivité de l'application.

Dans ce chapitre, nous proposerons d'abord de présenter les différents éléments qui composent une chaîne de communication point-à-point sans fil, de la modulation aux caractéristiques du canal physique, en passant par les blocs de la chaîne de transmission. On introduira ensuite les spécificités de la chaîne liées à une transmission coopérative. Nous détaillerons les différentes grandeurs qui permettent de déterminer les performances d'une communication, puis décrirons le modèle de consommation d'énergie issu de l'état de l'art. Forts de ces deux éléments, nous établirons le lien entre énergie par bit transmis et efficacité spectrale afin de déterminer s'il existe un compromis. L'ensemble de ces notions seront reprises dans le chapitre suivant pour analyse et approfondissement.

3.2 Transmission point-à-point au travers d'un canal de communication

Une communication point-à-point est définie comme la transmission d'information entre deux nœuds, l'émetteur étant appelé la source et le récepteur le destinataire. Nous nous intéressons au cas où chaque nœud ne peut émettre et recevoir simultanément : le canal de communication est alors *half-duplex*, ou à l'alternat [15]. Nous supposerons en outre que la source et le destinataire sont toujours représentés par les mêmes nœuds : le canal transporte

donc l'information dans un seul sens, et il est considéré comme simplex, ou monodirectionnel.

Le signal source traverse un système complexe composé de plusieurs blocs qui forment la chaîne de communication. Ces blocs, présentés dans cette section, permettent de transmettre le signal le plus efficacement et le plus exactement possible.

Dans cette section, on commencera par présenter le modèle mathématique général de la chaîne de transmission. Afin d'établir un contexte précis pour nos travaux, on détaillera ensuite les caractéristiques du canal de propagation et de la modulation.

3.2.1 Modèle d'une communication point-à-point

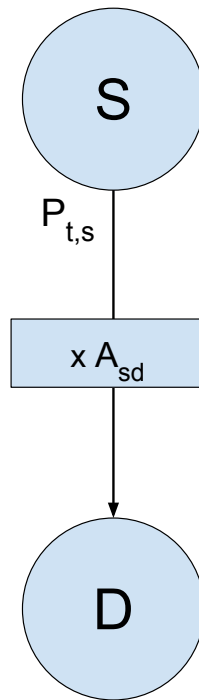


FIGURE 3.1 – Schéma simplifié du canal physique de communication : la source émet à une puissance initiale $P_{t,s}$ un signal qui est reçu par le destinataire après une atténuation A_{sd}

On considère un lien de transmission connectant deux nœuds sans fil, une source S et un destinataire D, ainsi que présenté sur la figure 3.1. Le signal est émis à une puissance $P_{t,s}$ et parvient au récepteur après avoir subi une atténuation A_{sd} que l'on va détailler.

L'information à transmettre est originalement un signal numérique binaire, qu'on considère de puissance unitaire. Il traverse la chaîne de communication présentée dans la figure 3.2 et subit des modifications jusqu'en sortie de la chaîne. Pour plus de simplicité, le codage de source – qui compresse le signal pour n'en retenir que les informations essentielles – et le codage canal – ou code correcteur d'erreurs – sont ignorés.

Pour assurer un débit maximum de données binaires sur le canal, le signal va bénéficier d'une modulation numérique, qui consiste à associer un symbole à chaque groupement de bits

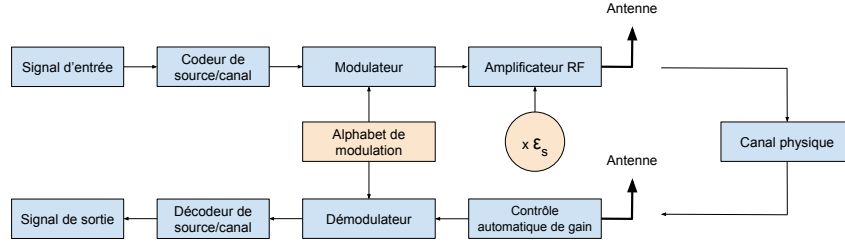


FIGURE 3.2 – Présentation des différents modules d'une communication sans fil

dont la taille dépend de la modulation choisie. La suite de symboles ainsi obtenue, une fois convertie en un signal analogique, constitue un signal x de puissance σ_x^2 . Afin d'atteindre la puissance d'émission $P_{t,s}$, le signal est amplifié par un facteur $\sqrt{\varepsilon_s}$, comme indiqué sur la figure 3.2; il est ensuite transmis sur le canal de propagation par l'antenne émettrice sous la forme $\sqrt{\varepsilon_s}x$.

Ce canal de propagation va déformer, atténuer le signal émis de manière à ce que le signal y reçu par l'antenne réceptrice soit

$$y = h_{sd}\sqrt{\varepsilon_s}x + n_{sd} \quad (3.1)$$

avec h_{sd} le gain du canal de propagation, qui suit une loi de probabilité dépendant de ce même canal, et n_{sd} le bruit additif gaussien centré. Le RSB au niveau de l'antenne de réception, défini comme le rapport entre la puissance du signal utile sur la puissance du bruit, est alors noté

$$\gamma_{sd} = \frac{\varepsilon_s |h_{sd}|^2 \sigma_x^2}{\sigma_{sd}^2} \quad (3.2)$$

avec $\sigma_{sd}^2 = \mathbb{E}[|n_{sd}|^2]$ la puissance du bruit sur la liaison source-destinataire.

Afin de retrouver la suite de données binaires originale, le signal va être numérisé puis démodulé. Cependant, les gains du canal et de l'amplificateur ont modifié la puissance du signal et peuvent donc mettre à mal le seuillage qui a lieu lors de la conversion analogique vers numérique. Le contrôle automatique du gain (CAG) vient résoudre ce problème en rétablissant le signal à sa puissance initiale [6]. Celui-ci est alors normalisé par un facteur z qui prend en compte la nature complexe du gain h_{sd} :

$$z = \frac{1}{\sqrt{\varepsilon_s}} \frac{h_{sd}^*}{|h_{sd}|^2} \quad (3.3)$$

avec h_{sd}^* le complexe conjugué et transposé de h_{sd} .

A noter que pour des raisons de simplicité d'écriture, dans la suite de nos travaux, nous considérerons $\sigma_x^2 = 1$, soit $\varepsilon_s = P_{t,s}$.

3.2.2 Le canal de propagation

Le canal de propagation est le support, physique ou non, dans lequel se propage le signal pour aller de l'émetteur au destinataire. Pendant son parcours au sein du canal, ce signal va subir des modifications de plusieurs sortes. En effet, la propagation n'est pas instantanée, ce

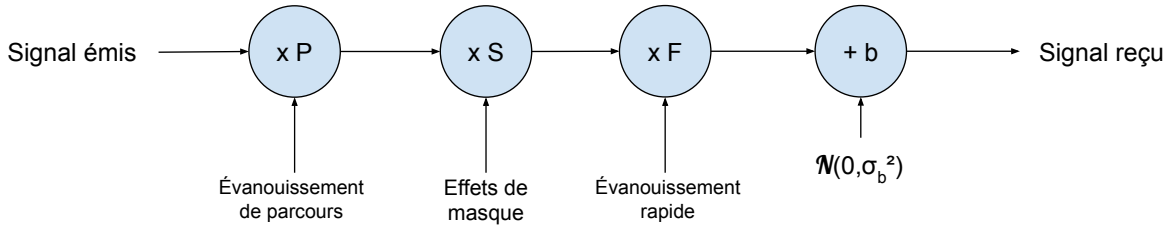


FIGURE 3.3 – Schéma des différentes déperditions subies par le signal sur le canal

qui peut conduire le signal à être plus ou moins atténué en fonction de la distance et de la nature du canal ; elle peut également être contrariée par des obstacles et des interférences : le signal peut expérimenter réflexions, réfractions, dispersions ou diffractions par exemple [172].

Dans le cas particulier des communications sans fil, ce canal est tout simplement l'atmosphère terrestre. Pour autant, cela ne signifie pas que le canal soit parfait, et il est nécessaire de modéliser les éléments cités précédemment. Sur la figure 3.3, les différentes déperditions du signal sont détaillées (affaiblissement de parcours, effets de masque et évanouissement), ainsi que le bruit, que l'on supposera ici additif et gaussien [39].

L'affaiblissement de parcours (ou *path loss*)

L'affaiblissement de parcours est un phénomène déterministe. Il est dû autant à l'éloignement du récepteur qu'aux phénomènes physiques tels que la réflexion ou la diffraction. Son effet est une réduction de la puissance du signal. La modélisation de cette atténuation peut être effectuée de manière déterministe ou empirique.

Les méthodes déterministes se basent sur des considérations géométriques pour modéliser les effets de l'affaiblissement. Le modèle le plus répandu, et le plus simple, est celui en air libre. Il s'agit du modèle idéal, dans lequel on estime qu'il existe une ligne de vue entre l'émetteur et le récepteur, et qu'aucun obstacle ne vient perturber le signal. Il est modélisé par la formule de Friis [143, 158, 181] :

$$PL = \frac{G_t G_r \lambda^2}{(4\pi)^2 d_{sd}^k} \quad (3.4)$$

avec G_t et G_r les gains des antennes d'émission et de réception, λ la longueur d'onde du signal, d_{sd} la distance entre le nœud émetteur et le nœud récepteur, et k l'exposant d'affaiblissement. Ce dernier est, pour les communications sans fil par onde radio, généralement compris entre 1 et 6. Il dépend du canal de propagation et peut être ajusté pour coller au plus près de la réalité.

A noter que ce modèle peut être légèrement complexifié pour prendre en compte une propagation extérieure dans laquelle un signal réfléchi parvient également au récepteur [160]. Il est alors nécessaire de prendre en compte des phénomènes additionnels tels que la diffraction comme par exemple au niveau des toits des habitations. La modélisation de ces phénomènes demeure assez complexe et des hypothèses simplificatrices (*i.e* toits de même taille, ou toits plats) sont utilisées [202].

Dans les situations où l'affaiblissement est particulier, ou quand on souhaite le déterminer avec précision, il est également possible de faire appel à des modèles dits empiriques, qui s'appuient sur plusieurs recherches dans le domaine. On peut citer entre autres le modèle de

Hata-Okumira, qui propose un ensemble de formules pour divers environnements extérieurs et pour des fréquences comprises entre 200MHz et 2GHz [71]. De nombreuses autres situations ont été étudiées, et l'on dispose aujourd'hui d'un grand panel de modèles [169, 224].

Dans la suite de nos travaux, sauf mention du contraire, nous considérons le modèle en espace libre avec un exposant $k = 2$, qui modélise l'affaiblissement de manière satisfaisante, en particulier en extérieur [133], et qui est le modèle le plus répandu dans l'état de l'art [40, 88].

Les effets de masque (ou *shadowing*)

L'effet de masque est un phénomène plus aléatoire que l'affaiblissement de parcours : local, il est lié à des atténuations successives généralement causées par de gros obstacles qui diffèrent au cours du temps (nuage, véhicule roulant...). Contrairement au phénomène précédent, il va donc varier d'une transmission à l'autre. Néanmoins, la moyenne de l'effet de masque correspond à l'affaiblissement.

En effet, l'effet de masque résultant de la somme de plusieurs évanouissements locaux successifs que l'on peut apparenter à une suite de variables aléatoires habituellement de nature gaussienne ou log-normale. Ce premier modèle est généralement préféré car, si on suppose ces variables comme étant indépendantes et identiquement distribuées, le théorème central limite nous indique que ce phénomène suit une loi gaussienne [172]. De par la nature des variations, on considère que la moyenne de chacune est nulle, et que la variance globale σ_L^2 de l'effet de masque dépend de l'environnement [39].

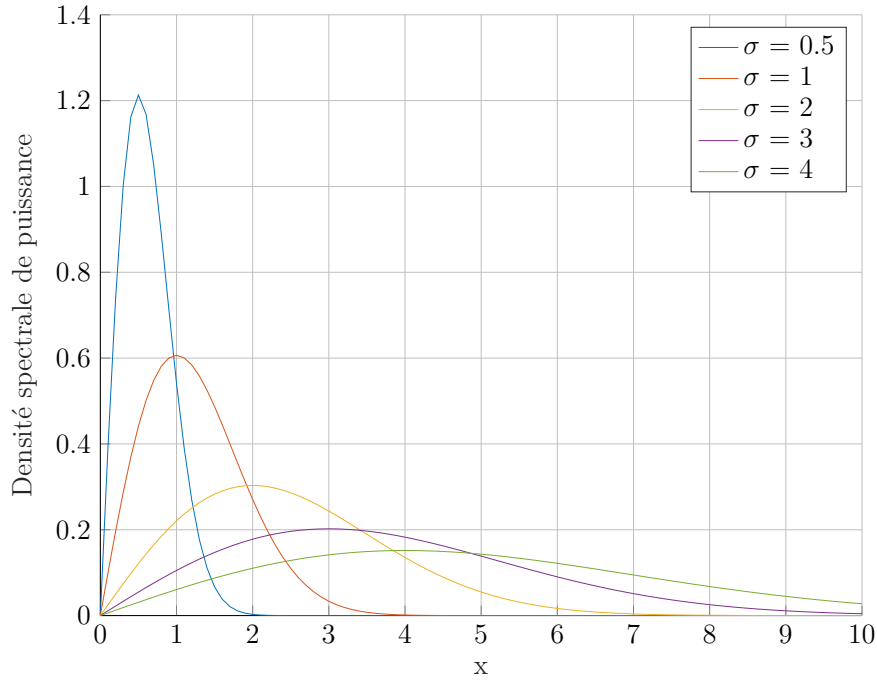
Les évanouissements rapides (ou *fast fading*)

Les dernières déperditions considérées sont les évanouissements rapides. Ceux-ci sont également un phénomène local qui suit une variable aléatoire et agit sur l'amplitude du signal. Ces interférences sont principalement liées à la propagation multi-trajets du signal. Le canal est finalement considéré comme étant à propagation linéaire et variant dans le temps [172].

Les évanouissements peuvent être plats ou sélectifs en fréquence. Dans le premier cas, toutes les fréquences de la bande passante subissent un même changement d'amplitude ; dans le second cas, ce changement peut être différent pour chaque bande de fréquence. Dans les communications sans fil pour les applications de l'Internet des objets, les bandes considérées sont généralement étroites si bien que les évanouissements rapides sont habituellement considérés comme plats en fréquence [155]. Parmi les principaux canaux plats en fréquences on peut alors citer :

Le canal à bruit blanc additif gaussien : Il s'agit du canal le plus simple possible. En réalité, le facteur d'évanouissement est ici considéré comme unitaire, c'est-à-dire que seuls l'affaiblissement de parcours et les effets de masque altèrent le signal. L'atténuation est alors invariante dans le temps.

Le canal de Rayleigh à évanouissement plat : Le canal de Rayleigh est très répandu dans le domaine des télécommunications [181, 168]. Il est utilisé pour modéliser un environnement comportant de nombreux trajets et sans ligne de vue dominante. Le module du coefficient du canal suit alors une distribution de Rayleigh, qui représente la norme d'un


 FIGURE 3.4 – Distribution de la loi de Rayleigh pour différentes valeurs du paramètre σ

vecteur gaussien bi-dimensionnel de coordonnées indépendantes. Sa densité de probabilité,

$$f(x|\sigma) = \frac{x}{\sigma^2} \exp\left(-\frac{x^2}{2\sigma^2}\right) \quad (3.5)$$

avec x le module du coefficient du canal, est représentée sur la figure 3.4 pour différentes valeurs de sa variance σ^2 .

Le canal de Rice : A la différence du canal de Rayleigh, le canal de Rice demande l'hypothèse d'une ligne de vue entre l'émetteur et le récepteur, c'est-à-dire que le signal direct domine sur ceux indirects. Le module du coefficient du canal suit alors une loi de Rice, qui est une généralisation de la loi de Rayleigh. Cette fois-ci, la distribution est décrite par deux paramètres, σ et ν , tels que :

$$f(x|\sigma, \nu) = \frac{x}{\sigma^2} \exp\left(-\frac{x^2 + \nu^2}{2\sigma^2}\right) I_0\left(\frac{x\nu}{\sigma^2}\right) \quad (3.6)$$

avec I_0 la loi de Bessel d'ordre 0. A noter que les paramètres σ et ν expriment le rapport K entre la puissance des signaux direct et indirects et la somme Ω des puissances de ces chemins, de telle manière que :

$$K = \frac{\nu^2}{2\sigma^2} \quad \text{et} \quad \Omega = \nu^2 + 2\sigma^2. \quad (3.7)$$

On peut alors considérer ν^2 comme étant la puissance (et la variance) du chemin direct, et 2σ celle des chemins indirects.

Le canal de Nakagami : Le canal de Nakagami est une généralisation des canaux susmentionnés. Il modélise de manière efficace la distribution des sources de diffusion du signal sous forme de groupes plutôt que sous forme d'individus indépendants. Il est ainsi utilisé pour modéliser des environnements généraux, tandis que les précédents canaux sont davantage applicables à des environnements spécifiques (urbains pour Rayleigh, suburbain pour Rice...) [172]. Le module du coefficient du canal suit cette fois-ci une distribution de Nakagami, dont la densité de probabilité dépend de deux paramètres m et Ω :

$$f(x|m, \Omega) = \frac{2m^m}{\Gamma(m)\Omega^m} x^{2m-1} \exp\left(-\frac{m}{\Omega} x^2\right). \quad (3.8)$$

où $\Gamma(\cdot)$ est la fonction Gamma, $\Omega = \bar{r}^2$ est la puissance moyenne et m le paramètre d'évanouissement [176]. À noter que quand $m = 1$, ce canal est équivalent à un canal de Rayleigh ; quand m tend vers l'infini, le comportement du canal s'apparente à celui du canal AWGN [172].

Le canal log-normal : Le canal log-normal est régulièrement utilisé dans le domaine des communications sans fil pour modéliser un environnement contenant de hauts obstacles comme des immeubles ou des montagnes, en particulier pour l'*ultra wideband*. Ce modèle est également très utilisé pour modéliser les communications en intérieur des bâtiments [102]. La densité de probabilité de ce canal est :

$$f(x|\sigma, \mu) = \frac{1}{\sigma\sqrt{2\pi}x} \exp\left(-\frac{(\ln x - \mu)^2}{2\sigma^2}\right) \quad (3.9)$$

avec σ^2 la variance et μ la moyenne logarithmique [187].

Le canal Rayleigh-lognormal : Le canal muni d'une distribution Rayleigh-lognormale est une autre modélisation utilisée dans les communications sans fil pour simuler un environnement multi-trajet avec un certain masquage [1]. La densité de probabilité de ce canal est un mélange de la distribution de Rayleigh et log-normale. Elle se définit ainsi :

$$f(x|\sigma, \mu) = \int_0^{+\infty} f_R(x|\lambda) f_N(\lambda|\sigma, \mu) d\lambda \quad (3.10)$$

avec σ^2 la variance et $\mu \geq 0$ la moyenne logarithmique, f_R la densité de probabilité de la loi de Rayleigh et f_N celle de la loi log-normale.

Dans nos travaux, nous utiliserons le canal AWGN, qui est en général utilisé pour modéliser les communications à très courtes distances (0-10m) [87], et un canal de Rayleigh à évanouissement plat, qui convient particulièrement aux communications terrestres longues distances (> 10m) [155].

3.2.3 Modulation du signal

Le signal transmis sur le canal de communication doit s'adapter aux caractéristiques de ce dernier (largeur de bande, fréquence porteuse...). Pour ce faire, on réalise une opération de modulation qui modifie les paramètres (amplitude et phase) de la fréquence porteuse de

manière à s'adapter au canal de transmission. On effectue donc une opération sur le signal, appelée modulation, afin de lui donner une forme adaptée à la transmission.

La modulation peut être analogique ou numérique en fonction de la nature du signal (lui-même analogique ou numérique). Pour une bande passante fixée, le choix d'un format de modulation consiste à trouver le compromis entre un débit binaire maximal et un taux d'erreurs minimal. Dans notre cas, on verra également qu'à cela s'ajoute la nécessité de limiter la consommation du système.

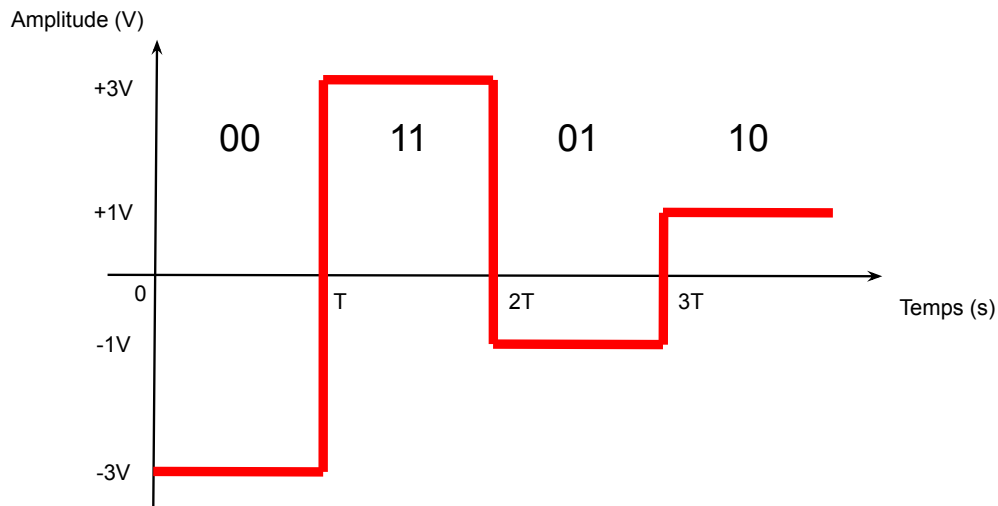


FIGURE 3.5 – Exemple de décomposition du signal binaire en symboles et créneaux associés. Ici, chaque symbole représente un couple de bits, ce qui nous donne quatre symboles possibles, que l'on décrit chacun par une amplitude particulière.

Transmission en bande de base

L'un des choix possibles est une transmission en bande de base. Dans ce cas, il n'y a pas de modulation à proprement parler : le signal binaire est décomposé en symboles, chacun associé à un créneau de durée et d'amplitude spécifique (voir figure 3.5). Le choix du filtre de mise en forme du signal (retour à zéro, bi-phase,...) permet d'adapter les propriétés de la suite de symboles ainsi créée (*i.e.* le code en ligne) à la spécificité du milieu de propagation ; on observe généralement deux ou trois niveaux d'amplitude, un plus grand nombre de niveaux nécessitant au récepteur un rapport RSB plus important pour atteindre la même qualité de service à puissance d'émission équivalente.

Cependant, la transmission en bande de base impose de transmettre l'information via les basses fréquences du spectre, ce qui est envisageable pour des transmissions en milieux clos tels que par câbles ou fibre optiques, mais ne l'est plus pour des milieux ouverts tels que des communications sans fil en raison des risques de la présence de bruits liés aux dispositifs industriels. On va alors préférer une transmission en bande transposée, c'est-à-dire permettant le choix de la fréquence d'émission, dite fréquence de porteuse.

Transmission en bande transposée

Moduler l'amplitude et/ou la phase d'une fréquence porteuse présente plusieurs avantages. D'une part, elle permet de transposer le signal autour d'une fréquence porteuse, c'est-à-dire de déplacer l'information à transmettre des basses fréquences vers une bande fréquentielle qui sera dédiée à l'application. Le signal est alors transmis dans le milieu de propagation à l'aide d'antennes et il est possible de découper le spectre en plusieurs bandes de fréquences, et donc transmettre plusieurs signaux simultanément sur plusieurs bandes de fréquences. D'autre part, en modifiant l'amplitude, la fréquence et la phase des symboles, il est possible d'augmenter la taille de l'alphabet des symboles tout en étant capable à la réception de discriminer les différents symboles car le RSB en sera plus avantageux comparé à la transmission en bande de base.

En bande transposée, le signal peut se décomposer suivant :

$$x(t) = x_I(t) \cos(2\pi f_0 t) - x_Q(t) \sin(2\pi f_0 t) \quad (3.11)$$

$$= \Re[\tilde{x}(t) \exp(j2\pi f_0 t)] \quad (3.12)$$

avec f_0 la fréquence porteuse et $\tilde{x}(t) = x_I(t) + jx_Q(t)$ l'enveloppe complexe du signal, dans laquelle $x_I(t)$ et $x_Q(t)$ sont respectivement les composantes en phase et en quadrature. En général, on pourra également noter $\tilde{x}(t) = \sum_k d_k h(t - kT)$ avec d_k la suite temporelle de symboles complexes, et $h(t)$ la réponse impulsionnelle du filtre de mise en forme du signal.

Il existe un grand nombre de modulations numériques possibles pour le signal, de plus en plus évoluées et de plus en plus spécifiques en fonction de l'application [122, 206]. Ces modulations complexes sont généralement basées sur le principe des modulations élémentaires que nous allons présenter.

La modulation par déplacement d'amplitude (MDA) (ou *amplitude-shift keying* (ASK))

La MDA est une modulation d'amplitude, c'est-à-dire qu'elle agit sur l'amplitude du signal en fonction des bits à coder. Elle est aussi appelée M-ASK en raison de ses M niveaux d'amplitudes. L'une des versions les plus simples, à deux états, est la modulation tout-ou-rien (ou *ook*) pour laquelle le signal émis a une valeur nulle ou de $+V$, pour une information binaire à transmettre respective de 0 ou 1.

Le signal issu de la M-ASK est réel, c'est-à-dire que

$$x_I(t) = \sum_k a_k h(t - kT) \quad (3.13)$$

$$x_Q(t) = 0 \quad (3.14)$$

avec a_k appartenant à l'alphabet $\{-(M-1), \dots, -3, -1, +1, +3, \dots, +(M-1)\}$ la valeur réelle d'amplitude associée au symbole émis.

La modulation par déplacement de fréquence (MDF) (ou *frequency-shift keying* (FSK))

La M-FSK (pour *multiple-FSK* ou encore *multilevel-FSK*) est une modulation numérique en fréquence, ses niveaux logiques sont associés à une variation de la fréquence porteuse f_0 . Le

signal modulé peut s'écrire suivant :

$$x(t) = A \cos \left(2\pi \sum_k ((f_0 + \delta_k)t + \phi_k) \right) \quad (3.15)$$

avec A la valeur de l'amplitude, δ_k la variation de la fréquence porteuse associée au symbole et ϕ_k la phase à l'origine.

La MAQ (ou *quadrature amplitude modulation* (QAM))

La MAQ, aussi appelée M-QAM, est une modulation numérique qui agit à la fois en phase et en amplitude. Ses composantes en phase et en quadrature s'écrivent :

$$x_I(t) = \sum_k a_k h(t - kT) \quad (3.16)$$

$$x_Q(t) = \sum_k b_k h(t - kT) \quad (3.17)$$

avec a_k et b_k appartenant à l'alphabet $\{-(M-1), \dots, -3, -1, +1, +3, \dots, (M-1)\}$ et $h(t)$ la réponse impulsionnelle du filtre de mise en forme d'onde. Chaque coefficient a_k ou b_k peut prendre $\sqrt{M}$ valeurs et chaque symbole complexe $a_k + jb_k$ véhicule alors une information binaire constituée de $\log_2(M)$ bits.

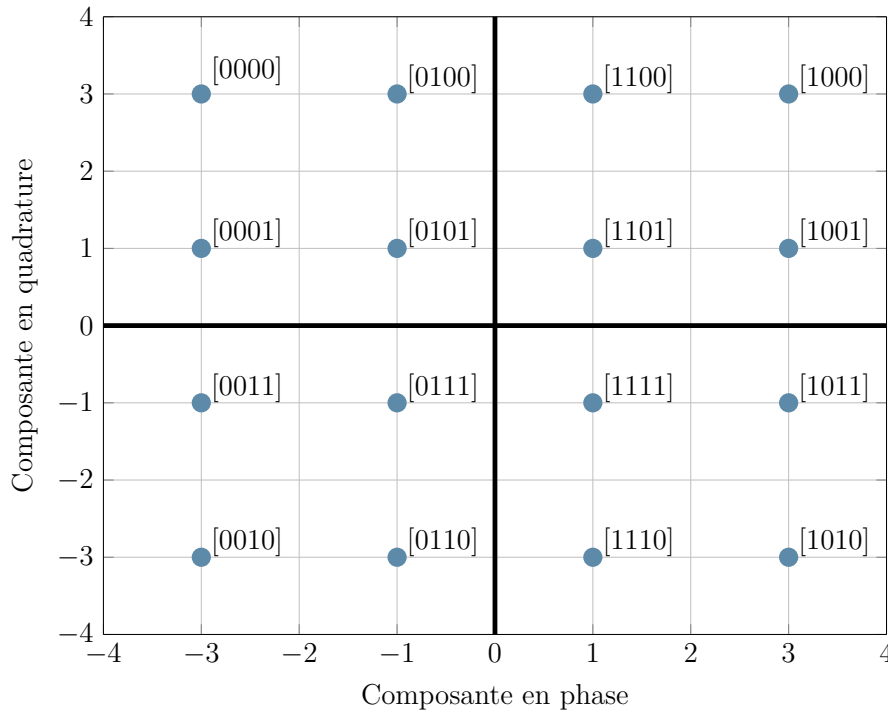


FIGURE 3.6 – Constellation de la modulation 16-QAM

Le choix d'un schéma particulier d'association entre éléments binaires et symboles de l'alphabet est habituellement effectué de manière à minimiser le taux d'erreur binaire, *i.e.* pour qu'une erreur de décodage sur un symbole se traduise par une seule erreur sur le mot binaire

correspondant. Ceci est obtenu en adoptant un codage de Gray pour l'association entre éléments binaires et symboles de la constellation. Par exemple, sur la figure 3.6, on peut observer la constellation obtenue pour une modulation 16-QAM. Ici, on note que les coefficients a_k et b_k prennent leurs valeurs dans $\{-3, -1, +1, +3\}$. Lorsqu'on prend par exemple le symbole $(-3, -1)$, qui code la séquence 0001, et que l'on observe ses voisins directs, on remarque que les autres séquences ont au maximum 1 bit de différence.

La modulation par déplacement de phase (MDP) (ou *phase-shift keying* (PSK))

La M-PSK est une modulation numérique de phase. La BPSK (*Binary Phase Shift Keying*, deux symboles) et la QPSK (*Quadrature Phase Shift Keying*, quatre symboles) en sont deux exemples très répandus. On utilise de préférence sa forme à coefficient complexe en définissant :

$$d_k = \exp(j\phi_k), \quad \phi_k = \frac{2\pi n}{M} + \theta, \quad n \in [0; M-1] \quad (3.18)$$

$$h(t) = V \text{ pour } t \in [0; T] \text{ et nul ailleurs} \quad (3.19)$$

avec θ une phase de rotation de la constellation comprise entre 0 et 2π .

La particularité de la PSK est qu'elle agit uniquement sur la phase : en effet, le symbole d_k est toujours de puissance unitaire, et tous les symboles de la constellation seront donc sur le cercle unité.

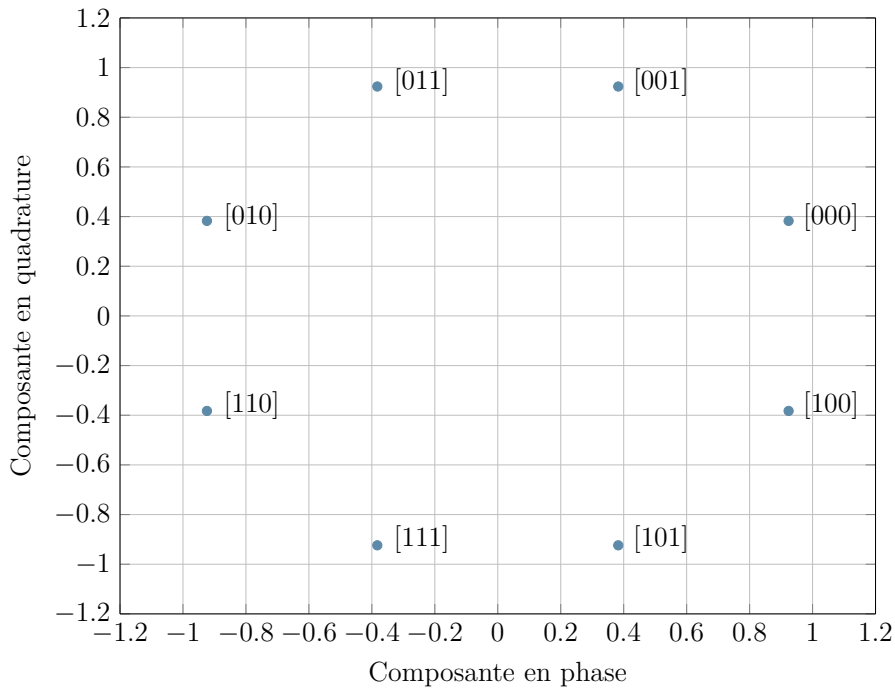


FIGURE 3.7 – Constellation de la modulation 8-PSK

On observe un exemple des symboles obtensibles sur la constellation de la modulation 8-PSK en figure 3.7. Sur cette figure, où l'on a considéré que $V = 1$, on observe que les symboles se trouvent bien tout le long du cercle unité. Les symboles représentent chacun un groupement de 3 bits car $2^3 = 8$. On note enfin, comme dans le cas de la M-QAM, que les symboles représentant les séquences les plus éloignées sont également très éloignées sur le cercle unité.

3.3 Transmission coopérative à l'aide d'un noeud-relai

Dans la partie précédente, nous avons étudié la manière dont communiquent deux nœuds entre eux, et quelles sont les particularités de cette communication. Nous avons notamment établi que le signal émis par la source est atténué par la distance et les obstacles avant d'arriver au niveau du destinataire, ce qui peut dégrader les informations reçues. Pour cette raison, il peut être intéressant de s'aider d'un troisième nœud, appelé nœud-relai, qui reçoit l'information et la retransmet au destinataire. Une telle communication, durant laquelle un voire plusieurs nœuds aident la source à transmettre son message, est appelée coopérative.

Dans cette section, on commencera par présenter les particularités du modèle mathématique de la communication coopérative. Par la suite, on s'intéressera aux différents protocoles selon lesquels le nœud-relai peut retransmettre l'information reçue. Enfin, nous détaillerons les méthodes grâce auxquelles le destinataire, qui reçoit des données issues de la source et du relai, combine ces informations.

3.3.1 Modèle d'une communication avec relai

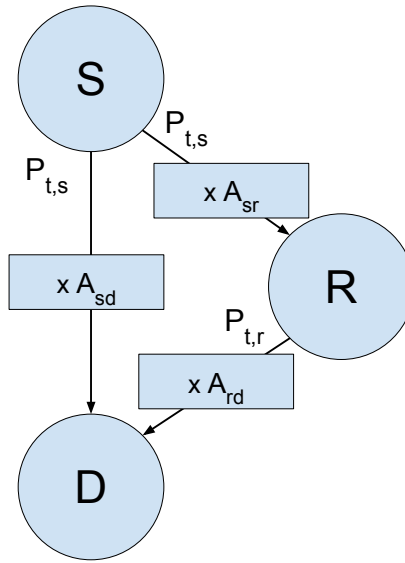


FIGURE 3.8 – Schéma simplifié du canal physique de communication : en plus du signal issu directement de la source, le nœud-destinataire reçoit un signal issu du nœud-relai. Le signal émis au nœud-relai est dépendant du protocole utilisé.

Considérons maintenant un schéma coopératif utilisant un nœud relai intermédiaire entre les nœuds source et destinataire comme représenté à la figure 3.8. Dans ce modèle, deux nouveaux chemins apparaissent : la liaison source-relai et la liaison relai-destinataire.

Comme auparavant, la source émet un signal à la puissance $P_{t,s}$. Ce signal parvient au nœud-destinataire atténué d'un coefficient A_{sd} ; de la même manière, le signal arrive au nœud-relai affecté par le facteur de pertes en amplitude A_{sr} sur la liaison source-relai différent de la valeur de A_{sd} car les chemins sont différents, les deux signaux n'ont donc pas subi les mêmes

transformations. Le signal en entrée du nœud-relai est ensuite modifié en fonction du protocole coopératif utilisé (protocoles qui seront abordés dans la section suivante) avant d'être retransmis à une puissance moyenne $P_{t,r}$ vers le nœud-destinataire et d'être atténué par le facteur de perte en amplitude A_{rd} sur la liaison relai-destinataire. Puisque nous considérons notre canal comme monodirectionnel, ce mécanisme a lieu en deux temps : le premier pendant lequel la source émet ses informations, et le second pendant lequel le relai les renvoie.

Pour décrire ce nouveau schéma coopératif, reprenons le modèle mathématique introduit précédemment en section 3.2.1. Le signal qui parvient au destinataire en provenance de la source s'écrit ainsi

$$y_{sd} = \sqrt{\varepsilon_s} h_{sd} x + n_{sd}. \quad (3.20)$$

avec x le signal émis, ε_s la puissance d'amplification à l'émetteur, h_{sd} le gain du canal de propagation sur la liaison source-destinataire et n_{sd} le bruit additif gaussien centré sur cette même liaison. Le RSB au niveau de l'antenne de réception se note

$$\gamma_{sd} = \frac{\varepsilon_s |h_{sd}|^2 \sigma_x^2}{\sigma_{sd}^2}. \quad (3.21)$$

Le signal reçu par le relai en provenance de la source se note de manière analogue :

$$y_{sr} = \sqrt{\varepsilon_s} h_{sr} x + n_{sr}. \quad (3.22)$$

avec h_{sr} le gain du canal de propagation entre la source et le relai, et n_{sr} le bruit additif gaussien sur ce même canal. Le signal y_{sr} subit alors une transformation qui dépend du protocole utilisé et vise à régénérer le signal x d'origine. Le signal en sortie de cette transformation se note x_r et est donc supposé de puissance σ_x . Il est alors émis par l'antenne du nœud-relai amplifié d'une puissance ε_r . Il est ensuite reçu au nœud destinataire sous cette forme :

$$y_{rd} = \sqrt{\varepsilon_r} h_{rd} x_r + n_{rd} \quad (3.23)$$

avec ε_r la puissance d'émission du relai, h_{sr} le gain du canal de propagation entre le relai et le destinataire, et n_{sr} le bruit additif gaussien sur ce même canal. Le RSB en réception de chacune de ces liaisons est alors

$$\gamma_{sr} = \frac{\varepsilon_s |h_{sr}|^2 \sigma_x^2}{\sigma_{sr}^2} \quad (3.24)$$

$$\gamma_{rd} = \frac{\varepsilon_r |h_{rd}|^2 \sigma_x^2}{\sigma_{rd}^2} \quad (3.25)$$

avec $\sigma_{sr}^2 = \mathbb{E}[|n_{sr}|^2]$ et $\sigma_{rd}^2 = \mathbb{E}[|n_{rd}|^2]$ respectivement la puissance du bruit sur les liaisons source-relai et relai-destinataire.

Les deux signaux y_{sd} et y_{rd} reçus par le nœud-destinataire sont ensuite combinés au niveau du récepteur selon des méthodes qui seront présentées dans une section ultérieure.

3.3.2 Protocoles de coopération

Lorsque le signal issu de la source parvient au niveau du relai, il n'est généralement pas retransmis tel quel ; il est traité puis retransmis vers le nœud-destinataire de manière à ce que le RSB en entrée du récepteur soit amélioré par rapport à celui du lien direct source-destinataire.

Plusieurs traitements sont possibles au nœud-relai. Les deux principaux protocoles sont l'*Amplify-and-Forward* et le *Decode-and-Forward* ; cependant, il existe également d'autres protocoles, bien souvent dérivés des deux précédents.

Amplify-and-Forward

AF signifie "amplifier et transmettre". Ainsi, ce protocole consiste simplement à multiplier le signal reçu par le relai par un gain G avant d'être réémis vers le destinataire. Le signal x_r en sortie de cette amplification va alors s'écrire [28]

$$x_r = G_{AF} y_{sr}. \quad (3.26)$$

Ce gain est généralement fixé à une valeur déterminée statistiquement au préalable en fonction des caractéristiques du système. Le but est de retrouver le signal x de puissance σ_x envoyé initialement par la source afin de le retransmettre au destinataire. Pour cela, l'opération à mener consiste à produire un signal x_r de puissance moyenne normalisée à σ_x , avant d'être re-amplifié par le gain ε_r de l'amplificateur du nœud-relai avant émission. On obtient donc

$$G_{AF} = \frac{\sigma_x}{\sqrt{\varepsilon_s |h_{sr}|^2 \sigma_x^2 + \sigma_{sr}^2}}. \quad (3.27)$$

Ce gain peut être intégré à l'équation 3.23 pour préciser le signal qui parvient au destinataire en provenance du nœud-relai :

$$\begin{aligned} y_{rd} &= \sqrt{\varepsilon_r} h_{rd} G_{AF} (\sqrt{\varepsilon_s} h_{sr} x + n_{sr}) + n_{rd} \\ &= \sqrt{\varepsilon_r \varepsilon_s} h_{rd} h_{sr} G_{AF} x + (\sqrt{\varepsilon_r} h_{rd} G_{AF} n_{sr} + n_{rd}) \end{aligned} \quad (3.28)$$

avec $\sqrt{\varepsilon_r} h_{rd} G_{AF} n_{sr} + n_{rd}$ le bruit de l'ensemble de la liaison source-relai-destinataire (aussi notée S-R-D). En supposant les bruits n_{sr} et n_{rd} centrés et indépendants, et h_{sr} et h_{rd} déterministes, le RSB instantané peut alors être noté

$$\begin{aligned} \gamma_{srd} &= \frac{\varepsilon_s \varepsilon_r |h_{sr}|^2 |h_{rd}|^2 G_{AF}^2 \sigma_x^2}{\mathbb{E}[\varepsilon_r |h_{rd}|^2 G_{AF}^2 |n_{sr}|^2 + |n_{rd}|^2]} \\ &= \frac{\varepsilon_s \varepsilon_r |h_{sr}|^2 |h_{rd}|^2 G_{AF}^2 \sigma_x^2}{\varepsilon_r |h_{rd}|^2 G_{AF}^2 \sigma_{sr}^2 + \sigma_{rd}^2}. \end{aligned} \quad (3.29)$$

En posant $\gamma_{rd} = \frac{\varepsilon_r |h_{rd}|^2 \sigma_x^2}{\sigma_{rd}^2}$, la relation précédente peut également s'écrire

$$\gamma_{srd} = \frac{\gamma_{sr} \gamma_{rd}}{\gamma_{sr} + \gamma_{rd} + 1}. \quad (3.30)$$

Le protocole AF est régulièrement utilisé dans la littérature en raison de sa simplicité et de sa faible complexité en comparaison d'autres techniques [53, 134].

Decode-and-Forward

decode-and-forward (DF) signifie "décoder et transmettre". En effet, dans ce deuxième protocole, la réception du signal au niveau du relai s'effectue de la même manière qu'au niveau du destinataire lors d'une communication directe, c'est-à-dire qu'il va le numériser et le décoder avant de le retransmettre. Le signal en entrée du nœud-relai est tout d'abord filtré par un filtre adapté optimal c'est à dire qu'il est multiplié par le coefficient :

$$z = \frac{1}{\sqrt{\varepsilon_s}} \frac{h_{sr}^*}{|h_{sr}|^2}. \quad (3.31)$$

puis le décodage se poursuit normalement. Finalement, le signal obtenu est réencodé pour donner un signal à émettre $x_r = G_{DF}y_{sr}$ de puissance supposée σ_x^2 avec G_{DF} un gain complexe liant l'entrée y_{sr} et la sortie x_r du protocole DF. En estimant que les signaux x et x_r sont de puissance moyenne égale, voire identiques dans le cas d'un décodage sans erreur, on peut écrire le signal qui parvient au destinataire en provenance du nœud-relai sous la forme :

$$y_{rd} = \sqrt{\varepsilon_r}h_{rd}x + n_{rd}. \quad (3.32)$$

Le RSB de ce signal est alors

$$\gamma_{rd} = \frac{\varepsilon_r|h_{rd}|^2\sigma_x^2}{\sigma_{rd}^2}. \quad (3.33)$$

Le point critique de ce protocole est le décodage du signal y_{sr} . En effet, le protocole DF fonctionne particulièrement bien quand le signal est correctement décodé au niveau du relai et qu'on obtient effectivement $x_r \approx x$, c'est-à-dire quand le canal source-relai a une faible probabilité d'erreur binaire [18, 114]. Il est cependant plus complexe que le protocole AF.

Autres protocoles

Il existe de nombreux autres protocoles. La plupart d'entre eux sont des variations des deux précédents. Pour le protocole AF, on peut notamment citer *orthogonal amplify-and-forward* (OAF), *non-orthogonal amplify-and-forward* (NAF) et *slotted amplify-and-forward* (SAF) [79, 101, 137, 212], qui se différencient principalement par la manière dont sont gérées les périodes d'émission et de réception du nœud-relai. Pour le protocole DF, *Laneman-Tse-Wornell decode-and-forward* (LTW DF), *Nabar-Bölcskei-Kneubühler decode-and-forward* (NBK DF) et *dynamic decode-and-forward* (DDF) [15, 79, 101, 137] sont des variations qui jouent également sur ces aspects.

Le protocole *compress-and-forward* (CF), ou "compresser et retransmettre", dont le protocole AF est un cas particulier, consiste à compresser – sans décoder – les données reçues par le relai avant de les réémettre. Il fonctionne particulièrement bien dans les cas où la liaison relai-destinataire est bonne [18].

D'autres méthodes consistent combiner les protocoles AF et DF. Une technique est par exemple de permettre au relai d'opérer de l'une ou l'autre des manières en fonction du RSB en réception, puisque le protocole DF fonctionne mal quand celui-ci est bas [113, 217]. [16] propose également une solution qui permet au relai de choisir entre une retransmission AF, une retransmission DF, ou pas de retransmission du tout.

Enfin des protocoles spécifiques sont développés pour certaines applications. Par exemple pour les nœuds réalisant de la récolte d'énergie, des protocoles alternant récolte et retransmission sont mis en œuvre [140, 213]

3.3.3 Combinaison des signaux à la réception

Le nœud-destinataire reçoit deux signaux en entrée : le signal y_{sd} issu de la source et y_{rd} celui issu du relai. Ces deux signaux ne sont généralement pas traités indépendamment, sous peine de doubler la quantité d'information dans le reste du système : ils sont combinés en entrée du nœud pour tenter de retirer la meilleure information possible, et en particulier augmenter le RSB global.

Dans nos travaux, on s'intéresse uniquement à la combinaison des deux signaux y_{sd} et y_{rd} , que l'on notera :

$$y_d = \alpha_{sd}^* y_{sd} + \alpha_{rd}^* y_{rd} \quad (3.34)$$

avec α_{sd}^* et α_{rd}^* deux coefficients complexes. Pour déterminer ces coefficients, il existe de nombreuses méthodes, dont on présente ici les principales : combinaison par sélection, à gains égaux et de rapport maximum.

A noter que dans le cas d'un réseau à multiples relais, il est possible de généraliser ces méthodes, voire de les combiner [141].

Combinaison par sélection

La combinaison par sélection consiste à sélectionner le signal dont le RSB par rapport à la source est le plus haut, c'est-à-dire que le coefficient pour l'un des chemins est de 1 et pour les autres de 0. Dans notre configuration, il s'agira généralement du relai ($\alpha_{sd} = 0$ et $\alpha_{rd} = 1$).

Le RSB en sortie de combinaison devient le maximum des RSB des deux chemins. Pour un protocole AF, ce RSB se note alors

$$\gamma_s = \max(\gamma_{sd}, \gamma_{srd}) \quad (3.35)$$

avec γ_{sd} exprimé à l'équation 3.2 et γ_{srd} à l'équation 3.30. Le CAG, bloc par lequel passe ensuite le signal, doit bien évidemment être adapté au signal sélectionné.

Si cette méthode est l'une des plus simples du point de vue de la complexité, elle n'est cependant pas la plus efficace [141] et ne travaillant qu'avec un seul relai, on se permettra d'utiliser d'autres méthodes.

Combinaison à gains égaux

Au contraire de la sélection, qui ne donne du poids qu'à un seul des signaux reçus, la combinaison à gains égaux (ou *equal gain combiner* (EGC)) donne le même poids à l'ensemble des chemins. Pour un système à $M - 1$ relais, le poids sera donc de $k = \frac{1}{M}$. Dans notre exemple, il sera de $\alpha_{sd} = \alpha_{rd} = \frac{1}{2}$.

Le RSB résultant d'un EGC ne peut être écrit simplement en fonction des RSB des signaux combinés, et a été détaillé dans [181]. On n'oublie pas non plus le CAG, qui doit être effectuée indépendamment pour chaque signal avant leur sommation.

L'EGC est une méthode dont la simplicité calculatoire est similaire à celle de la combinaison par sélection. Elle a également l'avantage de ne pas nécessiter de connaissance préalable sur les signaux que l'on combine. Elle est cependant mieux adaptée aux modulations à enveloppe constante comme la M-PSK. En outre, le calcul du RSB du signal disponible en sortie de combinaison n'est pas direct, ce qui en fait une méthode peu privilégiée dans les expériences où l'on utilise le RSB en réception et notamment, comme on le verra ultérieurement, dans le calcul de la consommation.

Combinaison *maximal ratio combiner* (MRC)

La combinaison de rapport maximum (ou MRC) est la solution qui, par définition, maximise le RSB en sortie du combineur.

Comme tout signal en entrée du noeud-destinataire peut s'écrire sous la forme $y_i = H_i x_s + N_i$ avec i le chemin suivi, on note :

$$y_d = \boldsymbol{\alpha}^H [\mathbf{H}x_s + \mathbf{N}] \quad (3.36)$$

avec $\boldsymbol{\alpha}^T = \begin{pmatrix} a_{sd} \\ a_{rd} \end{pmatrix}$, $\mathbf{H}^T = \begin{pmatrix} H_{sd} \\ H_{rd} \end{pmatrix}$ et $\mathbf{N}^T = \begin{pmatrix} N_{sd} \\ N_{rd} \end{pmatrix}$.

On cherche alors à maximiser le RSB :

$$\boldsymbol{\alpha}^{opt} = \underset{\boldsymbol{\alpha}}{\operatorname{argmax}} \left(\frac{\boldsymbol{\alpha}^H \mathbf{H} \mathbf{R}_{xx} \mathbf{H}^H \boldsymbol{\alpha}}{\boldsymbol{\alpha}^H \mathbf{R}_{nn} \boldsymbol{\alpha}} \right) \quad (3.37)$$

avec $\mathbf{R}_{xx} = \sigma_x^2 \mathbf{I}$ et $\mathbf{R}_{nn} = \mathbb{E}[\mathbf{N}\mathbf{N}^H]$.

En résolvant le problème, on détermine que la solution est :

$$\boldsymbol{\alpha}^{opt} = \frac{\mathbf{H}}{\mathbf{H}^H \mathbf{H}}. \quad (3.38)$$

Dans notre système à un seul relai avec protocole AF, on peut ainsi définir

$$\alpha_{sd} = \frac{h_{sd}}{\sqrt{\varepsilon_s} |h_{sd}|^2} \quad (3.39)$$

$$\alpha_{rd} = \frac{h_{sr} h_{rd}}{\sqrt{\varepsilon_s \varepsilon_r} |h_{sr}|^2 |h_{rd}|^2 G}. \quad (3.40)$$

Comme on peut le comprendre en comparant notre solution avec l'équation 3.3, ces coefficients incluent également le filtrage adaptatif du CAG, au contraire de ceux des méthodes précédentes. Le RSB en sortie de combinaison, pour un relai AF, est alors

$$\begin{aligned} \gamma_{MRC} &= \gamma_{sd} + \gamma_{srd} \\ &= \gamma_{sd} + \frac{\gamma_{sd} \gamma_{rd}}{\gamma_{sd} + \gamma_{rd} + 1} \end{aligned} \quad (3.41)$$

d'après l'équation 3.30.

La MRC est probablement la méthode de combinaison la plus utilisée dans l'état de l'art [33, 51, 134, 173, 223] en raison de sa maximisation du RSB. En effet, un plus grand RSB va notamment permettre de limiter la PEB en réception. Cependant, comme nous l'avons déjà fait remarquer, la MRC nécessite un plus grand nombre de calculs que les méthodes précédentes.

Autres combinaisons

Il existe bien d'autres combinaisons que nous n'aborderons ici que de manière superficielle.

Généralement utilisées dans des systèmes où le nombre de relais est strictement supérieur à 1, les méthodes hybrides combinent plusieurs méthodes, par exemple en effectuant une MRC sur une sélection de relais [141].

Dans certaines applications, il peut également être intéressant de maximiser le rapport signal sur interférence plutôt que le RSB, c'est ce qu'on appelle la combinaison optimum [181].

D'autres méthodes existent encore, impliquant une commutation dans les canaux sélectionnés, des variations de la MRC ou des combinaisons successives [181].

3.4 Mesurer les performances de la transmission

Une communication sans fil s'effectue habituellement sous un certain nombre de contraintes, aussi appelées qualité de service requise. En effet, il est possible de souhaiter une probabilité d'erreur maximale, un certain débit pour assurer un traitement en temps réel ou encore fixer un critère sur l'énergie. Il est nécessaire de savoir calculer tous ces critères afin de pouvoir établir une vue d'ensemble des performances du système communicant et de fixer les paramètres correspondants.

Dans cette section, on commencera par introduire la notion de probabilité d'erreur pour un ensemble de canaux présentés précédemment. Le concept d'efficacité de la communication est ensuite décrit pour mesurer les performances spectrales et énergétiques du système. On expliquera ensuite la notion de capacité d'un canal de transmission, que l'on mettra en relation avec le débit. Enfin, on mentionnera des critères spécifiques à une transmission coopérative.

3.4.1 Probabilité d'erreur en réception

L'objectif de la communication numérique est de transmettre une suite de données binaires transformées en une série de symboles, puis, à la réception, de construire le récepteur optimal réalisant ses décisions suivant le maximum de vraisemblance. Malheureusement, nous avons vu qu'au cours d'une transmission sans fil de nombreux phénomènes aléatoires peuvent impacter le signal analogique, et qu'il s'avère nécessaire de les estimer pour les contrer. Ainsi, des événements tels que l'ajout de bruit sur le canal ou l'atténuation aléatoire produite par un canal non-sélectif en fréquence peuvent causer des erreurs de décodage.

La PEB d'un canal est la probabilité qu'un bit $\hat{a}[n]$ estimé en réception soit différent du bit correspond $a[n]$ émis à l'origine, soit $P_{eb} = p(a[n] \neq \hat{a}[n])$ avec $a \in 0, 1$. La probabilité d'erreur symbole (PES), elle, correspond à la probabilité que l'estimation en réception $\hat{s}[k]$ d'un symbole émis $s[k]$ soit incorrecte, soit $P_{es} = p(s[k] \neq \hat{s}[k])$ avec s un symbole issu de l'alphabet de la modulation. Expérimentalement, on estime ces quantités sur des séquences de longueur finie. Lorsqu'on dispose à la fois des séquences binaires émises par l'émetteur et décidées au récepteur, il est possible de comptabiliser les erreurs binaires commises. On aboutit ainsi à la notion de taux d'erreur binaire (TEB), rapport entre le nombre nb_e de bits erronés et le nombre total nb_t de bits transmis : $T_{eb} = nb_e/nb_t$ avec T_{eb} le TEB. De même, cette mesure d'erreurs peut s'opérer en comparant la suite de symboles incorrectement transmis. On obtient alors la mesure T_{es} du taux d'erreur symbole (TES) : $T_{es} = ns_e/ns_t$ où ns_e et ns_t désignent respectivement le nombre de symboles incorrectement transmis et le nombre total de symboles émis. On estime qu'aux RSB considérés en réception, comme expliqué dans la section 3.2.3, une erreur de symbole équivaut en moyenne à une seule erreur binaire. On peut alors noter $T_{eb} = T_{es}/b$ avec b le nombre d'éléments binaires transmis par symbole modulé.

Le TEB varie en fonction de nombreux paramètres du système tels que la modulation, le type de canal ou le RSB en réception [40]. Il est possible de l'estimer par simulation de Monte-Carlo, c'est-à-dire en simulant un canal de communication sur lequel on émet plusieurs fois différentes suites de données binaires aléatoires. De nombreux travaux ont cherché à estimer les courbes obtenues par de telles simulations sous la forme de formules analytiques de la PEB ou de la PES à la fois précises et inversibles [40, 174, 177, 187, 174, 188, 216]. De telles formules permettent en effet de déterminer le RSB ciblé en fonction de la probabilité d'erreur souhaitée.

On sait que pour chaque paramétrage du système (canal, modulation, protocole coopéra-

tif...), il existe une courbe de probabilité d'erreur différente. Dans cette étude, on considère que l'émetteur utilise une modulation d'amplitude en quadrature (MAQ) (ou M-QAM) et transmet sur des canaux AWGN et de Rayleigh à évanouissement plat. On cherche à étudier la probabilité d'erreur à la fois pour une communication point-à-point et pour une communication coopérative avec protocole de relai AF et MRC.

La probabilité d'erreur est traditionnellement représentée par la PES en fonction du RSB par symbole modulé $\gamma_s = E_s/N_0$ [177, 188, 216]. On choisit de conserver ce format pour introduire les différentes formules analytiques de l'état de l'art ainsi que les courbes associées. Celles-ci sont comparées avec le TES obtenu par simulation de Monte-Carlo afin de valider leur usage.

Communication point-à-point

On commence par présenter les formules de la PES pour une communication point-à-point. Dans ce type de communication, le RSB symbole peut également être noté γ_{sd} .

Pour un canal AWGN, la formule exacte est donnée par [156] :

$$\text{Pes} = 1 - \left(1 - 2 \left(1 - \frac{1}{\sqrt{M}} \right) \mathcal{Q} \left(\sqrt{\frac{3\gamma_s}{M-1}} \right) \right)^2 \quad (3.42)$$

avec M l'ordre de modulation et $\mathcal{Q}(x) = \frac{1}{\sqrt{2\pi}} \int_x^\infty \exp\left(-\frac{x^2}{2}\right) dx$ la fonction Q-gaussienne [174]. Une approximation de cette formule est également donnée :

$$\text{Pes} \approx 4 \left(1 - \frac{1}{\sqrt{M}} \right) \mathcal{Q} \left(\sqrt{\frac{3\gamma_s}{M-1}} \right). \quad (3.43)$$

Puisque la fonction $\mathcal{Q}$ n'est pas forcément facile à manipuler, aussi on utilise bien souvent l'équivalence $\mathcal{Q}(x) \approx \frac{1}{2} \exp\left(-\frac{x^2}{2}\right)$ [40, 89]. Les deux précédentes formules deviennent alors :

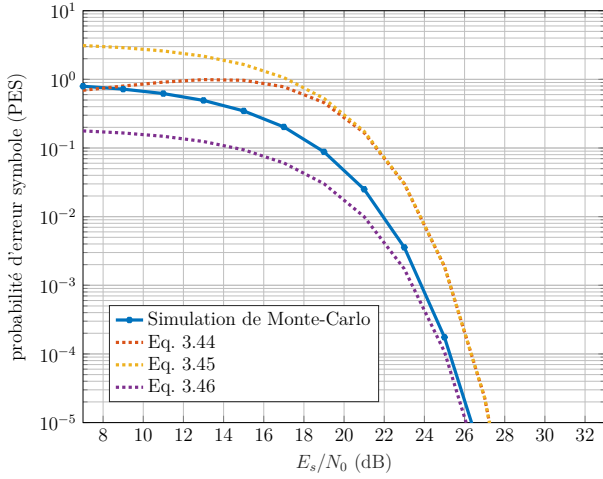
$$\text{Pes} \approx 1 - \left(1 - \left(1 - \frac{1}{\sqrt{M}} \right) \exp\left(-\frac{3\gamma_s}{2(M-1)}\right) \right)^2 \quad (3.44)$$

$$\text{Pes} \approx 2 \left(1 - \frac{1}{\sqrt{M}} \right) \exp\left(-\frac{3\gamma_s}{2(M-1)}\right). \quad (3.45)$$

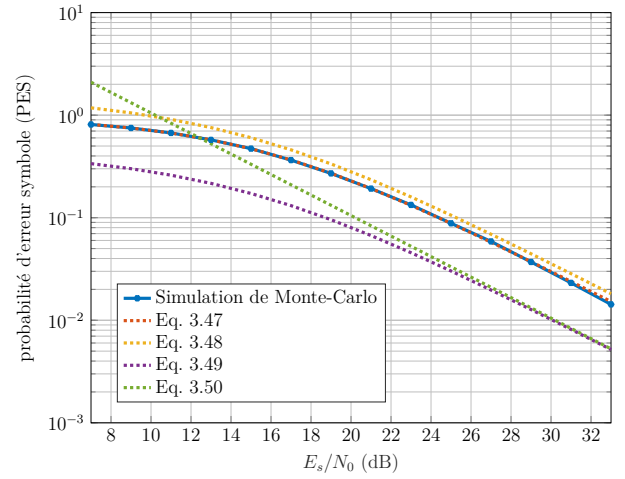
A partir de la même approximation de la fonction $\mathcal{Q}$, les auteurs de [87] proposent une nouvelle approximation, plus simple à manipuler, et donnée par :

$$\text{Pes} \approx \frac{1}{5} \exp\left(-\frac{3\gamma_s}{2(M-1)}\right). \quad (3.46)$$

On compare les courbes de PES obtenues à l'aide des équations 3.44, 3.45 et 3.46 sur la figure 3.9a pour un ordre de modulation arbitrairement fixé à $M = 2^6$, soit une 64-QAM. La courbe bleue représente la courbe obtenue par simulation de Monte-Carlo et qui est censée être très proche des courbes issues des formules 3.42 et 3.43. Cependant, on remarque que les courbes des équations 3.44 et 3.45 sont relativement éloignées de la courbe réelle, ce qui est lié à l'approximation de la fonction $\mathcal{Q}$. Néanmoins, ces courbes restent d'allure semblable à



(a) PES pour un canal AWGN



(b) PES pour un canal de Rayleigh

 FIGURE 3.9 – Probabilité d'erreur symbole en fonction du RSB symbole E_s/N_0 en dB, pour une communication point-à-point avec un signal modulé par 2^6 -QAM

la courbe réelle. L'équation 3.46, quant à elle, nous donne des résultats plutôt satisfaisants. En effet, pour de hauts RSB, l'allure de la courbe suit de près celle de la courbe réelle. Ceci est d'autant plus intéressant que, pour des PES inférieures à 10^{-2} qui sont traditionnellement utilisées, cette approximation est meilleure que celle des équations 3.44 et 3.45. Pour les RSB plus bas, il semble que les trois équations se valent.

On étudie à présent les formules de l'état de l'art pour le canal de Rayleigh à évanouissements plats rapides. L'équation exacte de la PES est donnée par [216] :

$$\text{Pes} = 1 - \frac{1}{M} - 2\rho \left(1 - \frac{1}{\sqrt{M}}\right) \left(\frac{1}{\sqrt{M}} + \frac{2}{\pi} \left(1 - \frac{1}{\sqrt{M}}\right) \tan^{-1}(\rho)\right) \quad (3.47)$$

avec $\rho = \sqrt{\frac{\gamma_s}{(2/3)(M-1)+\gamma_s}}$. Dans ce même article [216], les auteurs proposent une borne supérieure de la relation 3.47 donnée par :

$$\text{Pes} \leq 2 \left(1 - \frac{1}{\sqrt{M}}\right) (1 - \rho). \quad (3.48)$$

Ces deux formules restent relativement compliquées à utiliser, notamment pour calculer le RSB. L'équation 3.47 n'admet d'ailleurs aucune réciproque par rapport à γ_s , tandis que la fonction réciproque de 3.48 a une forme très complexe que l'on ne reproduira pas ici. En considérant que la borne supérieure 3.47 est atteinte, il est possible de présenter des approximations dont la réciproque est facilement calculable. Les auteurs de l'article [40] proposent l'approximation suivante, parmi les plus répandues :

$$\text{Pes} \approx \frac{1}{2}(1 - \rho). \quad (3.49)$$

Pour de forts RSB, on peut de plus considérer que $\text{Pes} \ll 1$ ce qui permet d'obtenir l'approximation [40] :

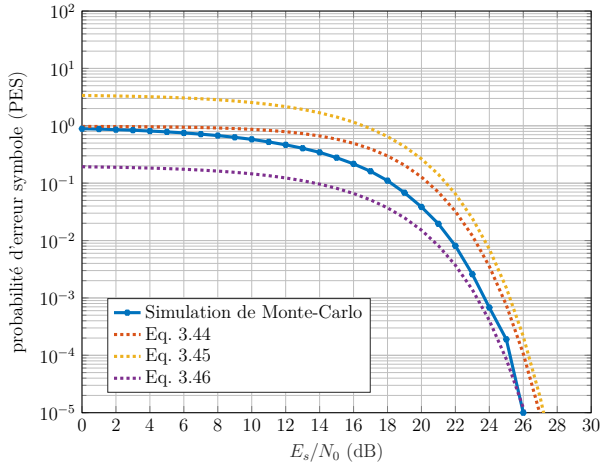
$$\text{Pes} \approx \frac{M-1}{6\gamma_s}. \quad (3.50)$$

A nouveau, on va comparer les courbes issues de l'ensemble de ces formules avec celle issue de la simulation de Monte-Carlo pour une 64-QAM. Les résultats disponibles sur la figure 3.9b nous montrent que la courbe orange issue de 3.47 modélise quasi-parfaitement la courbe expérimentale, et que sa borne supérieure 3.48, en jaune, reste une bonne approximation. En revanche, l'allure de la courbe violette issue de 3.49 a une allure similaire à celle de la courbe expérimentale, mais avec un décalage qui semble constant en échelle logarithmique. En effet, lorsqu'on calcule la différence de leur logarithme, on obtient pour toutes les valeurs de RSB un nombre qui tourne autour d'une même constante. Quant à l'approximation 3.50, on constate qu'elle est effectivement valable uniquement pour les hauts RSB.

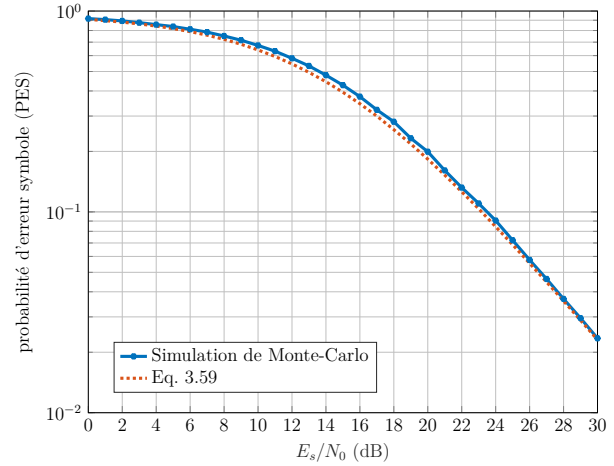
Communication coopérative

On présente dans un second temps les performances en termes de PES d'une communication coopérative avec protocole AF et combinaison des voies suivant MRC.

Dans un premier temps, on s'intéresse au canal AWGN. On sait que le produit de deux fonctions gaussiennes est une fonction gaussienne [22]. En raison de cette propriété, on peut estimer que le canal résultant de la succession des liaisons source-relai et relai-destinataire est gaussien. De la même manière, la somme de deux fonctions gaussiennes indépendantes est une fonction gaussienne, et on estime que chacun des trois canaux qui constituent notre système est indépendant des deux autres canaux. En raison de la combinaison de type MRC qui est une opération linéaire, le canal équivalent global du schéma coopératif peut donc être considéré lui aussi comme un canal AWGN. Ainsi, il est possible de reprendre les équations 3.44, 3.45 et 3.46 avec γ_s correspondant à γ_{MRC} donné à l'équation 3.41.



(a) PES pour un canal AWGN



(b) PES pour un canal de Rayleigh

FIGURE 3.10 – Probabilité d'erreur symbole en fonction du RSB symbole E_s/N_0 en dB, pour une communication coopérative avec protocole AF et un signal modulé par 2⁶-QAM

On fixe à nouveau l'ordre de modulation à $M = 64$. Comme on l'expliquera plus tard, il existe un lien linéaire entre le RSB sur la liaison source-destinataire et la liaison source-relai, que l'on notera pour le moment $\gamma_{sr} = D\gamma_{sd}$ avec D un coefficient réel positif, généralement supérieur à 1. On fixe ici arbitrairement $D = 2$. Il reste donc deux inconnues sur le canal : γ_{sd} et γ_{rd} . On choisit également de fixer arbitrairement $\gamma_{rd} = 10$ dB. En calculant γ_{MRC} à partir

de ces éléments, on trace ainsi les courbes des équations 3.44, 3.45 et 3.46 que l'on peut comparer avec la courbe expérimentale obtenue par simulation de Monte-Carlo, toutes les deux en fonction de $\gamma_{sd} = E_s/N_0$. On remarque que l'allure et les valeurs de l'ensemble des courbes sont relativement proches de la courbe expérimentale; l'approximation semble satisfaisante. Par ailleurs, les mêmes remarques que pour la communication point-à-point s'appliquent quant à la comparaison des différentes courbes.

Dans un second temps, on va étudier la PES pour un canal de Rayleigh à évanouissement plat. Au contraire d'un canal AWGN, la combinaison de différents canaux de Rayleigh ne résulte pas en un canal aux propriétés similaires; il est donc nécessaire de mener une analyse spécifique de la PES sur de tels canaux. Celle-ci a été largement étudiée dans [9] et [188] pour des communications à relais multiples. Nous simplifions ici la formulation pour ne considérer qu'un unique relais. On commence par définir les fonctions suivantes :

$$\mathcal{I}_1(c) = 0.5 \left(1 - \frac{c}{1+c} \right), \quad (3.51)$$

$$\mathcal{I}_3(c) = 0.25 \left(1 - \frac{4}{\pi} \frac{c}{1+c} \arctan \sqrt{\frac{c+1}{c}} \right), \quad (3.52)$$

$$N(c, \theta) = 2\sqrt{c(1+c)} \sin 2\theta, \quad (3.53)$$

$$D(c, \theta) = (1 + 2c) \cos 2\theta - 1, \quad (3.54)$$

$$\mathcal{T}(c, \theta) = 0.5 \arctan \left(\frac{N(c, \theta)}{D(c, \theta)} \right) + \frac{\pi}{2} \left(1 - \text{sign}(N(c, \theta)) \frac{1 - \text{sign}(D(c, \theta))}{2} \right). \quad (3.55)$$

On pose également les coefficients c_1 et c_2 respectivement liés au relais et à la source :

$$c_1 = g \frac{\gamma_{sr} \gamma_{rd}}{\gamma_{sr} + \gamma_{rd}}, \quad (3.56)$$

$$c_2 = \gamma_{sd} \quad (3.57)$$

avec $g = \frac{3}{2(M-1)}$. Finalement, [188] définit ainsi la PES sur un canal de Rayleigh avec relais AF et MAQ :

$$\text{Pes} \approx 4K \left(\mathcal{I}_1(c_1) - \frac{\mathcal{T}(c_2, \frac{\pi}{2})}{\pi} \sqrt{\frac{c_2}{1+c_2}} \frac{c_2}{c_2 - c_1} + \frac{\mathcal{T}(c_1, \frac{\pi}{2})}{\pi} \sqrt{\frac{c_1}{1+c_1}} \frac{c_2}{c_2 - c_1} \right) \quad (3.58)$$

$$+ 4K^2 \left(\mathcal{I}_3(c_1) - \frac{\mathcal{T}(c_2, \frac{\pi}{4})}{\pi} \sqrt{\frac{c_2}{1+c_2}} \frac{c_2}{c_2 - c_1} + \frac{\mathcal{T}(c_1, \frac{\pi}{4})}{\pi} \sqrt{\frac{c_1}{1+c_1}} \frac{c_2}{c_2 - c_1} \right). \quad (3.59)$$

La courbe correspondant à cette équation est tracée sur la figure 3.10b en reprenant les conditions expérimentales utilisées pour obtenir la figure 3.10a. Si son calcul est complexe, on obtient néanmoins une courbe qui suit de très près la courbe expérimentale obtenue par simulation de Monte-Carlo.

3.4.2 Efficacités du canal

L'efficacité d'un canal est une valeur qui mesure l'utilisation des ressources de ce canal. Ici, nous nous intéressons aux efficacités énergétique et spectrale, qui renseignent sur l'utilisation de l'énergie et de la bande passante.

Efficacité énergétique

Aucune définition de l'efficacité énergétique (EE) ne fait réellement consensus [87]. Pourtant, la communication sans fil est l'un des éléments les plus coûteux d'un système communicant [72, 119]. Afin de maîtriser cette consommation, il semble donc nécessaire de définir un critère pour la calculer. On estime par exemple qu'augmenter le nombre de bits émis par unité d'énergie permet de limiter la consommation, et notamment d'absorber la croissance du trafic avec la même énergie [119]. Ainsi, son usage le plus courant établit un lien entre le nombre de bits émis et l'énergie consommée par leur transmission [108]. Plus précisément, le critère énergétique η_{EE} le plus répandu [87, 59] est le rapport entre le débit R (en bits/s) et puissance moyenne P (en J/s) consommée pour transmettre ce débit, tel que

$$\eta_{EE} = \frac{R}{P} \text{ en bits/J.} \quad (3.60)$$

Obtenir une valeur importante de l'EE consiste alors à maximiser la valeur de η_{EE} , soit en augmentant le débit, soit en diminuant la puissance moyenne consommée. Par ailleurs, on note que l'énergie consommée par bit émis s'écrit

$$E_b = \frac{P}{R} = \frac{1}{\eta_{EE}}, \quad (3.61)$$

c'est-à-dire que maximiser η_{EE} revient à diminuer E_b .

Efficacité spectrale

L'ES exprime l'utilisation de la bande passante. Notée θ , elle est définie comme le rapport entre le débit R et la largeur de la bande passante B (en Hz), soit :

$$\theta = \frac{R}{B} \text{ en bits/s/Hz [87].} \quad (3.62)$$

Pour certaines constellations comme la M-PSK et la M-QAM d'ordre de modulation M , on peut établir le lien $\theta \approx b$ avec $b = \log_2(M)$, avec b exprimant le nombre de bits transmis par symbole [206].

Avec l'épuisement des bandes passantes disponibles et l'accroissement des usages, l'ES est un critère crucial dans beaucoup d'applications [215]. Par ailleurs, il existe un compromis entre ES et EE qui, en raison des impératifs d'économie d'énergie, gagne en importance [78, 108, 215]. Ce compromis est détaillé dans la section suivante.

3.4.3 Capacité de Shannon

La capacité $C = \sup I(x, y)$ d'un canal de propagation correspond à la borne supérieure de l'information mutuelle entre les signaux d'entrée $x[n]$ et de sortie $y[n]$ du canal. Elle est peut être interprétée comme la borne supérieure du débit R accessible sur un canal. Le théorème de Shannon l'établit pour une source gaussienne non codée et un canal AWGN, il s'agit de la capacité de Shannon [156] :

$$C = B \log_2 \left(1 + \frac{P}{N_0 B} \right) = B \log_2 \left(1 + \frac{E_b R}{N_0 B} \right) \quad (3.63)$$

avec P la puissance moyenne du signal émis, N_0 la densité spectrale monolatérale de bruit sur le canal, E_b l'énergie consommée par bit transmis et B la largeur de la bande passante. On sait en outre que l'efficacité spectrale peut se noter $\theta = \frac{R}{B} = b$, il est donc possible de réécrire l'équation 3.63 :

$$\begin{aligned} C &= B \log_2 \left(1 + \frac{E_b}{N_0} b \right) \\ &= B \log_2 \left(1 + \frac{E_s}{N_0} \right) \\ &= B \log_2 (1 + \gamma_s) \end{aligned} \quad (3.64)$$

avec E_s l'énergie consommée par symbole transmis et γ_s le RSB par symbole. On trace ainsi la courbe de la capacité de Shannon normalisée par la bande passante B sur la figure 3.11. L'aire en dessous de la courbe correspond aux valeurs d'efficacité spectrale (ou débit normalisé par B) réalisables car la capacité est la borne supérieure du débit. L'aire au-dessus de la courbe représente l'efficacité spectrale inaccessible sur un tel canal.

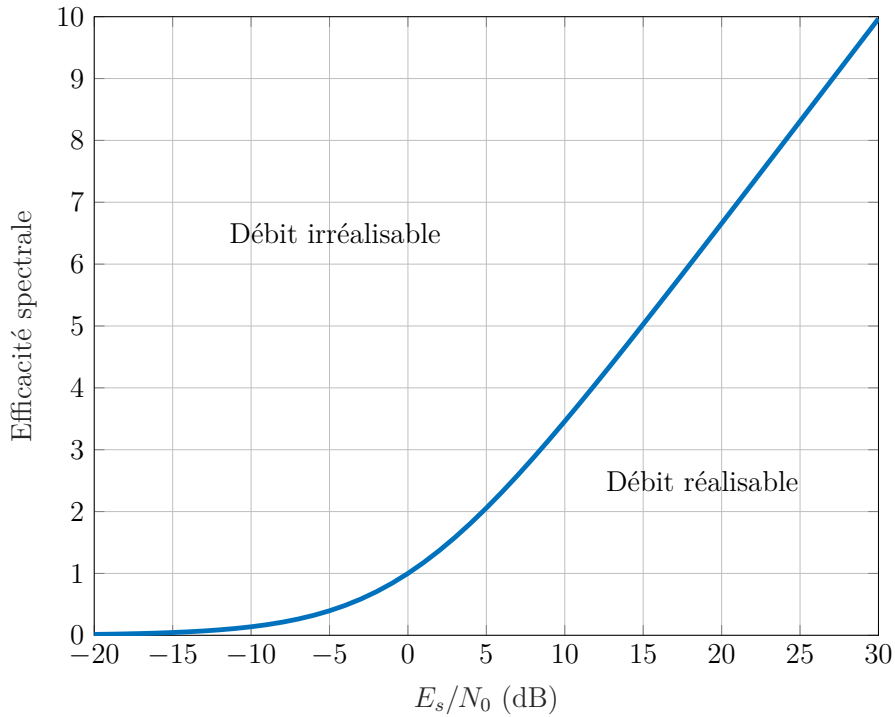


FIGURE 3.11 – Capacité de Shannon normalisée en fonction du RSB symbole

Par ailleurs, on suppose en général que le débit atteint sa limite, soit $R \approx C$. En normalisant l'expression 3.64 de la capacité, on obtient alors l'efficacité spectrale en fonction du RSB par symbole :

$$\theta = \log_2 (1 + \gamma_s). \quad (3.65)$$

On établit ainsi une relation entre la capacité et l'efficacité spectrale d'une part, et l'efficacité spectrale et le RSB par symbole d'autre part. On note de surcroît que cette expression

admet une réciproque, ce qui permet d'exprimer simplement le rapport RSB par symbole qui sera d'une grande importance dans les sections suivantes :

$$\gamma_s = 2^\theta - 1. \quad (3.66)$$

Enfin, la capacité nous permet de comprendre le lien qui unit les efficacités spectrale et énergétique. En effet, en reprenant l'équation 3.63 et en normalisant le débit par la bande passante, on obtient :

$$\begin{aligned} \theta &= \log_2 \left(1 + \frac{E_b R}{N_0 B} \right) \\ &= \log_2 \left(1 + \frac{\theta}{\eta_{EE} N_0} \right) \end{aligned} \quad (3.67)$$

avec toujours η_{EE} l'efficacité énergétique. On peut finalement déduire de cette équation que :

$$\eta_{EE} = \frac{\theta}{N_0(2^\theta - 1)}. \quad (3.68)$$

On trace cette efficacité énergétique, normalisée par $1/N_0$, sur la figure 3.12. Cette courbe confirme l'analyse que nous pouvons faire de l'équation 3.68 : quand l'efficacité spectrale θ augmente, l'efficacité énergétique η_{EE} diminue (à densité spectrale de puissance constante). Il existe donc bien un compromis à faire entre ces deux efficacités, comme mentionné dans [78, 108, 215].

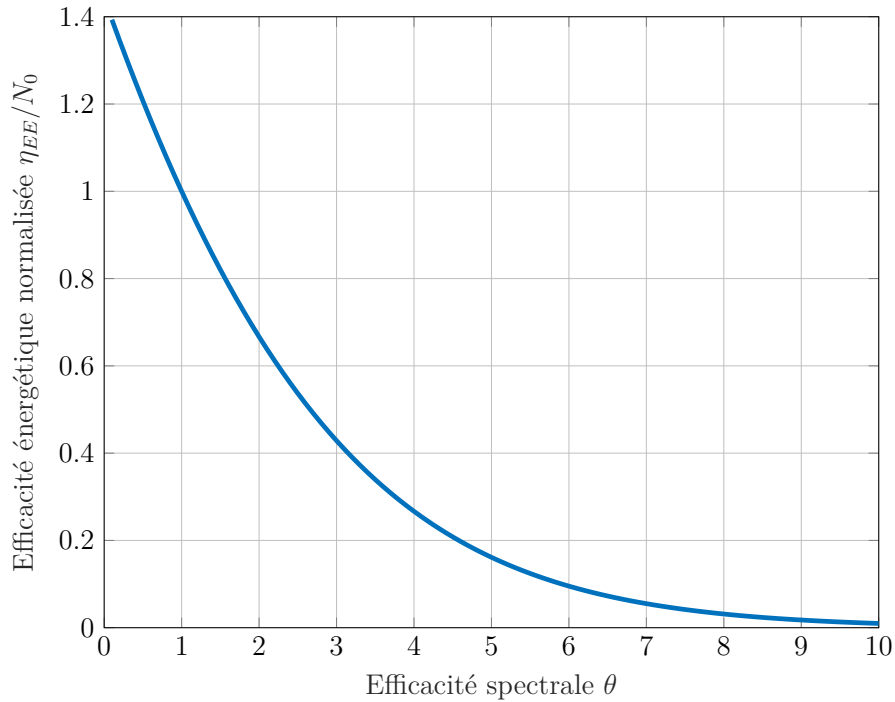


FIGURE 3.12 – Efficacité énergétique η_{EE} normalisée en fonction de l'efficacité spectrale θ

3.4.4 Métriques du coopératif

Les précédents critères peuvent s'appliquer tant à une communication point-à-point que coopérative afin de pouvoir comparer leurs performances. A ce titre, il existe à notre connaissance très peu de critères dédiés uniquement au coopératif.

[170] propose un critère pour comparer la puissance de réception pour deux schémas de transmission, par exemple le protocole AF avec le protocole DF.

De son côté, [33] définit le gain de coopération permettant de comparer la consommation d'énergie d'une transmission directe E_{direct} avec celle d'une transmission coopérative E_{coop} . Ce gain est alors noté

$$G_c = \frac{E_{direct}}{E_{coop}}. \quad (3.69)$$

En raison des valeurs de l'énergie mises en jeu, ce rapport est toujours strictement positif. S'il est supérieur à 1, alors le trajet direct est plus avantageux du point de vue de la consommation qu'un trajet coopératif. Dans le cas contraire, c'est l'inverse qui s'opère ; plus ce gain est petit, plus la transmission coopérative devient énergétiquement avantageuse.

Puisque nous étudions seulement la communication point-à-point et la communication coopérative munie du protocole AF, ce dernier critère sera utilisé pour étudier les avantages de la transmission multi-nœuds dans la suite de nos travaux.

3.5 Consommation d'énergie d'une communication

Nous avons déjà expliqué que la communication est souvent la partie la plus coûteuse énergétiquement dans un système communicant [72]. Afin de maîtriser cette consommation d'énergie, il est nécessaire de savoir la calculer. Un modèle précis de consommation permet d'analyser l'impact de chaque paramètre du système sur celle-ci, ce qui nous permettra à terme de réaliser un compromis entre efficacités spectrale et énergétique.

Dans cette section, nous détaillerons dans un premier temps le modèle traditionnel de l'état de l'art en matière de consommation point-à-point. Nous exposerons dans un second temps les spécificités et le calcul dans le cas d'un relai AF.

3.5.1 Modèle énergétique d'une transmission point-à-point

Dans cette première section, on souhaite calculer la consommation d'une communication point-à-point. Pour cela, nous nous intéressons uniquement à la partie analogique du système, qui consiste en un ensemble de modules présentés sur la figure 3.13. Ces blocs successifs ont pour but de traiter le signal de son passage vers l'analogique jusqu'à sa retranscription en numérique après la transmission.

Le convertisseur numérique-analogique (CNA) : Comme son nom l'indique, ce bloc convertit un signal numérique en un signal analogique continu dans le temps et dont les valeurs d'amplitude sont proportionnelles à celles des valeurs numériques.

Le convertisseur analogique-numérique (CAN) : Il s'agit du bloc réalisant l'opération inverse à celle du CNA, c'est-à-dire convertissant un signal analogique vers un signal numérique. Ce signal analogique, continu en temps et en amplitude, doit être discrétisé dans ces deux dimensions. Pour cela, il est important de connaître à la fois le nombre de

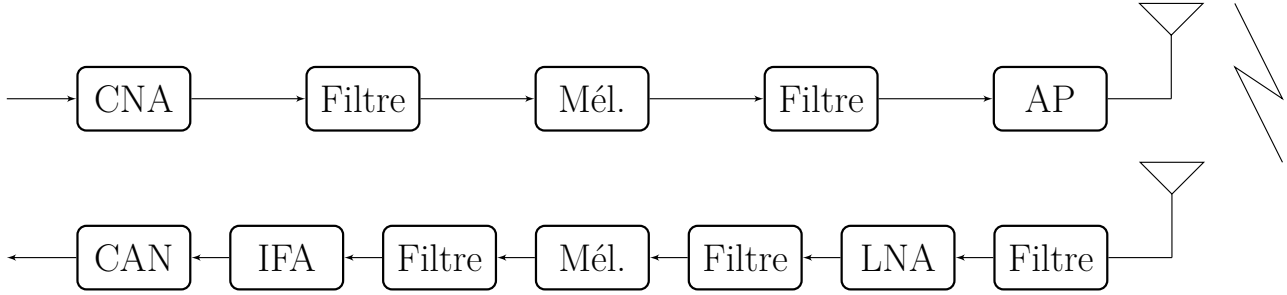


FIGURE 3.13 – Blocs du circuit de l'ensemble transmetteur-récepteur

valeurs discrètes possibles pour l'amplitude du signal et la fréquence à laquelle le signal doit être échantillonné [42]. Il faut aussi spécifier le seuil de saturation x_s du convertisseur ou de façon équivalente le facteur de charge du convertisseur défini par $\gamma_{CAN} = x_s/\sigma_x$ avec σ_x l'écart-type du signal à quantifier.

Les filtres : Au niveau du transmetteur et du récepteur, on trouve plusieurs filtres, notamment des filtres passe-bande anti-repliement et des filtres passe-bas anti-bruit [20, 112].

Les mélangeurs : Le rôle du mélangeur consiste à traduire le signal d'une fréquence porteuse à une autre. Au niveau du transmetteur, il conduit la fréquence porteuse intermédiaire du signal à sa fréquence porteuse radio fréquence (RF). Au niveau du récepteur, il le récupère en bande de base [56].

L'amplificateur de puissance (AP) : L'amplificateur de puissance est le dernier module avant l'émission du signal par l'antenne. Il agit simplement sur l'amplitude du signal afin que sa puissance atteigne la valeur souhaitée.

L'amplificateur faible bruit : Généralement appelé *low noise amplifier* (LNA), il s'agit d'un module présent en sortie de l'antenne de réception et qui a pour but d'amplifier le signal utile tout en minimisant le bruit du canal [171]. On note que le filtre passe-bas qui le suit permet d'éliminer le bruit faiblement amplifié.

L'amplificateur de fréquences intermédiaires : Appelé aussi *intermediate frequency amplifier* (IFA), ce module a pour but d'amplifier les parties du spectre attachées au signal utile. Il permet au signal de retrouver un niveau de puissance compréhensible pour le CAN. Il est généralement associé à un filtre passe-bande, qui se trouve juste avant dans notre modèle [197].

A chaque période T , l'émetteur transmet L bits au récepteur. Au cours de cette période, le système peut se trouver dans quatre états : en sommeil, en activité, en transition du sommeil vers l'activité et en transition de l'activité vers le sommeil. Ces états durent respectivement T_{sp} , T_{on} , T_{tr} et T_{off} tels que $T = T_{sp} + T_{on} + T_{tr} + T_{off}$. L'extinction des modules étant supposée très rapide, T_{off} peut être négligé [40].

La puissance est l'énergie consommée pour un certain temps. L'énergie totale consommée par le système est ainsi la somme des énergies consommées dans chaque état :

$$\begin{aligned} E_{totale} &= E_{on} + E_{sp} + E_{tr} \\ &= P_{on}T_{on} + P_{sp}T_{sp} + P_{tr}T_{tr} \end{aligned} \quad (3.70)$$

avec P_{on} la puissance en fonctionnement, P_{sp} la puissance en sommeil et P_{tr} la puissance en transition entre sommeil et activité. On peut également étudier cette énergie sous la forme

d'une énergie consommée par bit émis :

$$\begin{aligned} E_b &= E_{totale}/L \\ &= (P_{on}T_{on} + P_{sp}T_{sp} + P_{tr}T_{tr})/L. \end{aligned} \quad (3.71)$$

- En veille, on considère que plus aucun module du circuit n'est en marche, et que leur consommation d'énergie est donc nulle : $P_{sp} = 0$.
- Le temps de transition T_{tr} est fixé à la valeur de celui du synthétiseur de fréquence. En effet, ce module est le plus lent à mettre en marche. De plus, pour limiter le gaspillage d'énergie, les autres modules sont mis en fonction quelques instants seulement avant la fin d'une durée de transition : on considère alors que seuls les synthétiseurs du transmetteur et du récepteur consomment de l'énergie pendant ce temps, et donc $P_{tr} = 2P_{syn}$.
- Le temps de fonctionnement T_{on} est égal à $\frac{L}{bB}$ avec $b = \log_2(M)$, M la taille de l'alphabet de modulation. On sait par ailleurs que l'efficacité spectrale $\theta = \frac{R}{B} = \frac{L}{BT_{on}}$, ceci nous permet de confirmer que pour une M-QAM, $b \approx \theta$ [40]. La puissance de fonctionnement P_{on} , quant à elle, prend plusieurs éléments en compte dans son calcul. En effet, on peut décomposer P_{on} en la somme de deux puissances : P_t , la puissance de transmission du signal, et P_{c0} , la puissance de consommation du circuit. Cette dernière puissance peut elle-même être décomposée avec d'une part P_c la consommation du circuit à l'exception de l'amplificateur de puissance, qui est compris d'autre part avec une puissance $P_{amp} = \alpha P_t$ telle que $\eta = \frac{\xi}{1+\alpha}$ soit l'efficacité du drain de l'amplificateur de puissance, ξ étant le PAPR du signal à amplifier.

En résumé, on obtiendra :

$$\begin{aligned} P_{on} &= P_t + P_{c0} \\ &= P_t + P_{amp} + P_c \\ &= (1 + \alpha)P_t + P_c \end{aligned} \quad (3.72)$$

$$\leq P_{max}. \quad (3.73)$$

En effet, en pratique, la puissance mise à disposition par la batterie de l'appareil est finie : la puissance consommée en tout temps ne peut excéder une certaine valeur (le problème ne se pose pas en veille et en transition car les valeurs sont bien plus basses qu'en état de marche). Par ailleurs, on peut détailler P_c comme ceci :

$$P_c = 2P_{mix} + 2P_{syn} + P_{LNA} + P_{IFA} + P_{fil} + P_{DAC} + P_{ADC} \quad (3.74)$$

avec P_{mix} la puissance consommée par chaque mélangeur, P_{LNA} celle de l'amplificateur faible bruit, P_{IFA} celle de l'amplificateur à fréquence intermédiaire, P_{fil} celle de l'ensemble des filtres du circuit, P_{DAC} celle du CNA et P_{ADC} celle du CAN.

Intéressons-nous maintenant au calcul de P_t . De manière générale, on peut réécrire P_t sous la forme :

$$P_t = P_r G_d \quad (3.75)$$

avec P_r la puissance du signal reçu au destinataire et G_d le facteur de gain de puissance sur le canal de transmission. La puissance du signal reçu se définit elle-même par

$$P_r = BN_0 N_f \gamma_{sd} \quad (3.76)$$

avec B la bande passante, N_0 la densité spectrale de puissance monolatérale du bruit sur le canal, N_f la puissance du bruit liée au récepteur et γ_{sd} le RSB sur la liaison source-destinataire, que l'on peut déterminer en fonction de la probabilité souhaitée (voir section 3.4.1). Le facteur de gain de puissance, lui, est détaillé sous la forme :

$$G_d = \frac{(4\pi d_{sd})^2}{G_r G_t \lambda^2} \quad (3.77)$$

avec d_{sd} la distance entre la source et le destinataire, G_r et G_t les gains respectifs du récepteur et du transmetteur, et λ la longueur d'onde du signal. On reconnaîtra ici l'inverse de la formule de Friis (équation 3.4, pour laquelle on fixe un exposant d'affaiblissement de 2).

Finalement, le modèle de l'énergie consommée par bit transmis, que l'on nommera (E.trad), se résume par :

$$E_b = (1 + \alpha) \frac{N_0 N_f (4\pi d_{sd})^2 \gamma_{sd}}{G_r G_t \lambda^2 b} + \frac{P_c}{bB} + \frac{2P_{syn} T_{tr}}{L}. \quad (3.78)$$

Paramètre	Description	Valeur
f_c	Fréquence porteuse	2.4 GHz
B	Largeur de la bande passante	1 MHz
$G_r G_t$	Facteur des gains d'antenne	1
$N_0/2$	Demi-puissance du bruit du canal	10^{-16} W/Hz
N_f	Facteur de bruit en réception	10 dB
L	Nombre de bits émis	200 kb
T_{tr}	Temps de transition	$5 \mu s$
Peb	Probabilité d'erreur binaire	10^{-3}
P_{mix}	Puissance dissipée par un mélangeur	30.3 mW
P_{syn}	Puissance dissipée par un synthétiseur de fréquence	50 mW
P_{fil}	Puissance dissipée par l'ensemble des filtres	5 mW
P_{LNA}	Puissance dissipée par le <i>low noise amplifier</i> (LNA)	20 mW
P_{IFA}	Puissance dissipée par l'IFA	3 mW
P_{ADC}	Puissance dissipée par le CAN	6.7 mW
P_{DAC}	Puissance dissipée par le CNA	15.4 mW

TABLE 3.1 – Paramètres du modèle

Cette équation est la base de nos travaux suivants. On s'inspire de [40] pour présenter la consommation qui en résulte. On présente ici les courbes d'énergie par bit pour une M-QAM et pour un canal AWGN et un canal Rayleigh à évanouissement plat, ceci sur des distances de 0.1m, 1m et 10m, comme justifié dans la section 3.2.2. Les paramètres utilisés sont ceux présents dans le tableau 3.1, dont on justifie les valeurs ci-dessous.

La fréquence porteuse f_c est une fréquence communément utilisée pour les transmissions radio dans les bandes ISM (*Industrial Scientific Medical*), auxquelles s'intéressent nos travaux et dont les fréquences porteuses les plus utilisées en Europe correspondent aux bandes situées à 453 MHz, 868 MHz et 2.4GHz. [12, 80].

La largeur de la bande passante B doit, comme nous l'avons déjà expliqué, être limitée, en particulier dans les fréquences qui nous intéressent, qui sont déjà surchargées. Dans les applications Wi-fi à 2.4GHz, elle est généralement comprise entre 20 et 40MHz [211]. Dans les normes IoT orientées faible consommation, la bande passante est bien plus étroite, avec des largeurs allant jusqu'à 500kHz seulement pour LoRaWan [104], ou 1MHz pour Ingenu [210].

Les gains d'antenne en réception G_r et en transmission G_t sont usuellement d'un ordre de grandeur compris entre 1 et 10. Lorsqu'on travaille avec une seule antenne en émission et en réception, on peut considérer leur produit égal à 1 [11].

La valeur de N_0 , qui représente la puissance du bruit liée au canal, va définir la sensibilité au niveau du récepteur. Ce bruit radioélectrique est issu de divers phénomènes électromagnétiques dont les sources diffèrent notamment en fonction de la bande passante [80].

Le facteur de bruit N_f correspond à la dégradation du RSB en raison du bruit thermique du circuit [61], c'est-à-dire le rapport entre le RSB à l'entrée du récepteur et celui en entrée du CAN. En outre, on fixe généralement une valeur M_l comme une marge vis-à-vis des niveaux de bruit les plus importants, marge qui tourne usuellement autour de 10dB [75]. Cette marge est ici prise en compte dans le facteur N_f .

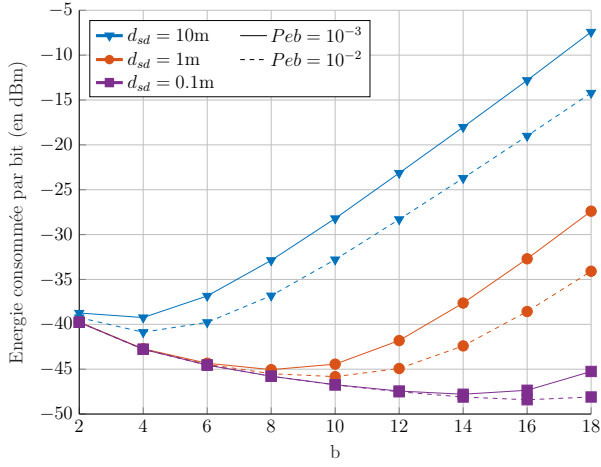
Le nombre de bits L émis dépend directement du besoin de l'application. Cependant, on l'estime généralement de l'ordre de 100kb [40, 189, 198].

Le temps de transition T_{tr} est basé sur celui de l'allumage du synthétiseur de fréquence. Celui-ci est de l'ordre de la micro-seconde, avec des valeurs actuellement en-dessous de $3\mu s$ [225]. Pour nos expériences, nous utilisons cependant la valeur de l'état de l'art qui est de $5\mu s$.

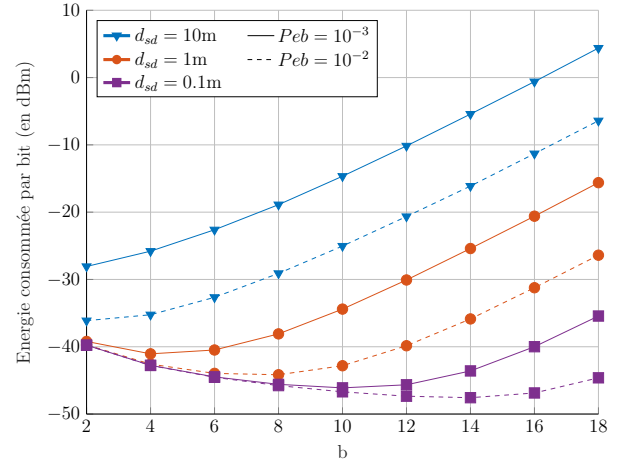
La probabilité d'erreur binaire P_{eb} correspond à la probabilité d'erreur binaire (sans codage correcteur d'erreurs) que nous fixerons à des valeurs de 10^{-2} et 10^{-3} . En effet, même si une opération de codage/décodage correcteur d'erreurs est toujours présent dans un systèmes de transmission réel, nous faisons le choix de ne pas en intégrer dans nos modèles car cela est hors des travaux de cette thèse. Nous considérerons en conséquence ce codage comme apportant un gain constant qui produit uniquement un décalage des courbes de la PEB en fonction du RSB.

Les puissances des différents modules sont toujours notées en mW en raison de leurs valeurs, que l'on peut calculer à partir de différents paramètres du circuit analogique [111].

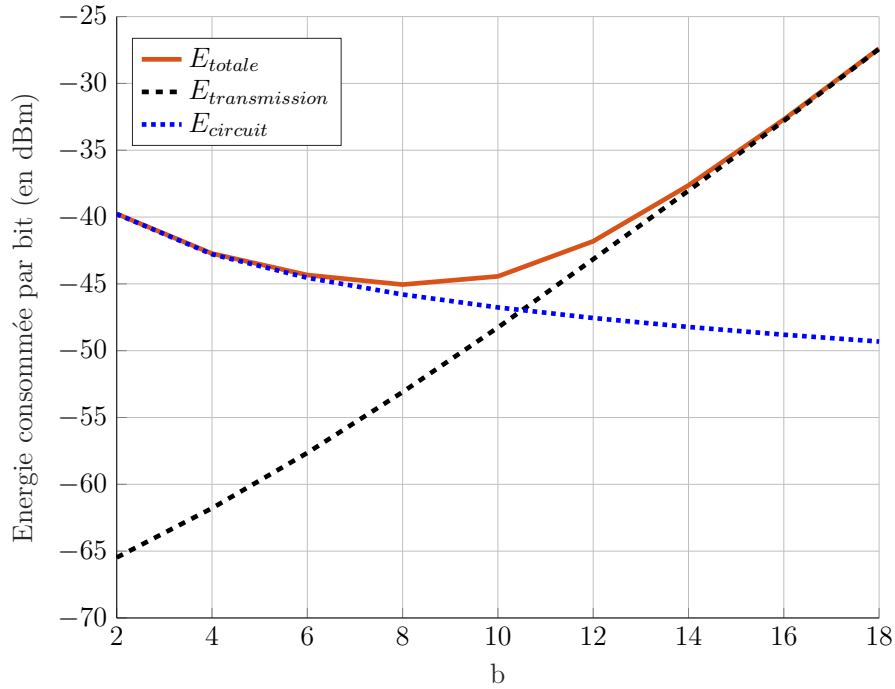
Les courbes de consommation d'énergie pour les différents paramétrages décrits sont présentées sur les figures 3.14a et 3.14b respectivement pour les canaux AWGN et de Rayleigh. Dans un premier temps, quelle que soit la nature du canal, on remarque que la consommation d'énergie augmente avec la distance, ce qui est tout naturel en raison du coefficient G_d qui est proportionnel à d_{sd}^2 avec d_{sd} la distance source-destinataire. Dans un deuxième temps, on peut également observer le même comportement quand on diminue la PEB. En effet, comme on l'observe dans la section 3.4.1, plus la PEB (ou PES) diminue, plus le RSB requis pour assurer une qualité de service donnée doit être important ; comme la puissance de réception est proportionnelle au RSB conformément à l'équation 3.76, ces résultats sont cohérents. Enfin, on peut noter que l'énergie consommée quand on modélise le canal à l'aide d'une distribution de Rayleigh est généralement plus haute que pour un canal AWGN, en particulier pour les longues distances et les hauts ordres de modulation.



(a) Consommation pour un canal AWGN



(b) Consommation pour un canal de Rayleigh

 FIGURE 3.14 – Energie consommée par bit pour une M-QAM ($M = 2^b$) et un canal AWGN ou de Rayleigh en fonction de l'ordre de modulation et pour différentes PEB

 FIGURE 3.15 – Energies totale, de transmission et du circuit consommées par bit pour une M-QAM($M = 2^b$), un canal AWGN et une distance $d_{sd} = 1\text{m}$ en fonction de la taille de constellation b de la modulation et pour PEB de 10^{-3}

Canal	AWGN		Rayleigh	
Peb	10^{-2}	10^{-3}	10^{-2}	10^{-3}
$d_{sd} = 10\text{m}$	4	4	2	2
$d_{sd} = 1\text{m}$	10	8	8	4
$d_{sd} = 0.1\text{m}$	16	14	14	10

 TABLE 3.2 – Valeurs de b^* en fonction du canal, de la distance source-destinataire et la PEB

Pour une distance de $d_{sd} = 10\text{m}$ et un canal de Rayleigh, la consommation d'énergie croît monotonement. Ce n'est pas le cas pour les distances très petites telles que 0.1m et 1m, ni avec un canal AWGN. La raison est illustrée en figure 3.15, qui présente des courbes pour un canal AWGN, une distance $d_{sd} = 1\text{m}$ et une PEB de 10^{-3} . Pour une PEB donnée, la diminution de l'ordre de modulation implique nécessairement une réduction du RSB requis au récepteur pour atteindre la qualité de service souhaitée, ce qui conduit à l'utilisation d'une puissance de transmission moyenne plus faible. Au contraire, l'énergie du circuit, qui s'exprime $\frac{P_c}{bB}$ comme vu sur l'équation 3.78, décroît linéairement avec b . Ainsi, on remarque que l'énergie de transmission est négligeable devant celle du circuit pour les faibles ordres de modulation, ce qui nous donne une énergie totale décroissante. Quand b augmente, l'énergie de transmission croît rapidement, et c'est l'énergie du circuit qui devient négligeable; l'énergie totale croît à son tour. Il existe donc une valeur optimale de b qui minimise l'énergie totale consommée par élément binaire. On remarque que cette valeur optimale ne correspond pas forcément à la valeur minimale de b (et donc de l'efficacité spectrale) car la loi de variation de E_{totale} selon le paramètre b n'est pas forcément monotone croissante.

Ce phénomène sera abordé plus amplement dans la section suivante. Néanmoins, la figure 3.14 nous apporte déjà quelques éléments. En effet, on peut observer l'ordre qui minimise chacune des courbes de la figure sur le tableau 3.2. On appelle cet ordre b^* . Il semble que b^* augmente quand la distance source-destinataire d_{sd} diminue. C'est le cas également quand la PEB croît. Enfin, b^* est toujours supérieur à paramètres équivalents pour le canal AWGN par rapport au canal de Rayleigh. Tous ces éléments sont à mettre en parallèle avec l'augmentation de l'énergie totale par bit, qui a lieu exactement dans les mêmes circonstances. Il semble donc que l'augmentation de b^* soit liée à une diminution de l'énergie de transmission.

3.5.2 Modèle énergétique d'une transmission coopérative

L'ajout d'un nœud-relai est à l'origine de la création de deux canaux de transmission supplémentaires, comme on a pu le voir sur la figure 3.8. Pour transmettre la même information, on ne doit cependant effectuer que deux émissions successives au sein d'une même période : l'une par le nœud-source, et l'autre par le nœud-relai, comme détaillé sur la figure 3.16. Chacune de ces transmissions va s'effectuer en une même durée T_{on} comme pour la communication point-à-point. On suppose alors dans nos travaux que $2T_{on} < T$, afin de pouvoir réaliser une communication coopérative toujours en temps réel.

Modèle de consommation

On détaille à présent la consommation pour chacune des émissions. A la première période, le système réalise une émission, celle de la source, et deux réceptions, par le relai et le destinataire.

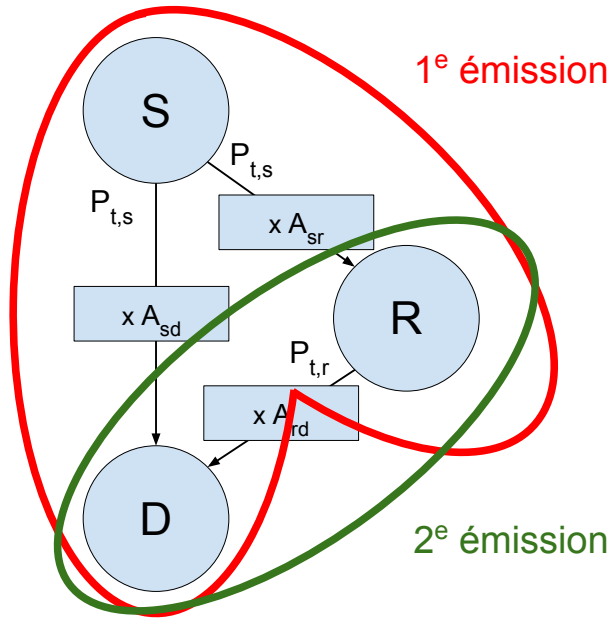


FIGURE 3.16 – Schéma simplifié du canal physique de communication coopérative : dans un premier temps T_{on} , la source transmet le signal au relai et au destinataire ; dans un second temps, c'est le relai qui transmet le signal au destinataire.

Si l'émission par la source et la réception par le destinataire s'effectuent comme décrit dans le modèle point-à-point, la réception par le relai est quelque peu différente : en effet, si on utilise le protocole AF, il n'est pas nécessaire de numériser le signal avant réémission, et on supprimera donc les modules CAN et CNA du circuit. Finalement, on peut noter la puissance de fonctionnement pour cette première transmission :

$$P_{on,1} = (1 + \alpha)P_{t,s} + P_{c,tx,s} + P_{c,rx,d} + P_{c,rx,r} \quad (3.79)$$

avec

- la puissance d'émission de la source $P_{t,s} = BN_0N_f \frac{(4\pi d_{sd})^2}{G\lambda^2} \gamma_{sd} = BN_0N_f \frac{(4\pi d_{sr})^2}{G\lambda^2} \gamma_{sr}$;
- la puissance du circuit émetteur de la source $P_{c,tx,s} = P_{mix} + P_{syn} + P_{fil}/2 + P_{DAC}$ en supposant une consommation des filtres similaire en émission et en réception ;
- la puissance du circuit récepteur du destinataire $P_{c,rx,d} = P_{mix} + P_{syn} + P_{LNA} + P_{IFA} + P_{fil}/2 + P_{ADC}$;
- la puissance du circuit récepteur du relai $P_{c,rx,r} = P_{mix} + P_{syn} + P_{LNA} + P_{IFA} + P_{fil}/2$.

La puissance de transition entre l'état de sommeil et celui d'activité correspond à la puissance du synthétiseur de fréquence pour chaque nœud, ce qui nous donne $P_{tr,1} = 3P_{syn}$.

Lors de la deuxième période, le système ne réalise plus qu'une seule émission, celle du relai, et une réception, celle du destinataire. La puissance de fonctionnement peut donc s'écrire :

$$P_{on,2} = (1 + \alpha)P_{t,r} + P_{c,tx,r} + P_{c,rx,d} \quad (3.80)$$

avec

- la puissance d'émission du relai $P_{t,r} = BN_0N_f \frac{(4\pi d_{rd})^2}{G\lambda^2} \gamma_{rd}$;

- la puissance du circuit émetteur du relai $P_{c,tx,r} = P_{mix} + P_{syn} + P_{fil}/2$;
- la puissance du circuit récepteur du destinataire $P_{c,rx,d}$ identique à précédemment.

De la même manière, la puissance de transition pour cette période ne prend en compte que le relai et le destinataire, soit $P_{tr,2} = 2P_{syn}$.

Finalement, la consommation totale d'énergie par bit transmis lors d'une communication coopérative AF, qui sera notée sous le nom (E.coop.trad), s'exprime par

$$E_b = \frac{(P_{on,1} + P_{on,2})T_{on} + (P_{tr,1} + P_{tr,2})T_{tr}}{L} = \frac{((1 + \alpha)(P_{t,s} + P_{t,r}) + P_{c,tx,s} + 2P_{c,rx,d} + P_{c,tx,s} + P_{c,rx,r})T_{on} + 5P_{syn}T_{tr}}{L}. \quad (3.81)$$

Le rapport signal sur bruit (RSB)

On travaille avec un protocole coopératif AF et une combinaison à la réception qui est la MRC. On a vu dans 3.3.2 que le RSB sur l'ensemble de la liaison source-relai-destinataire est donné par [32] :

$$\gamma_{srd} = \frac{\gamma_{sr}\gamma_{rd}}{\gamma_{sr} + \gamma_{rd} + 1}. \quad (3.82)$$

Considérant une combinaison MRC au destinataire, le RSB obtenu au récepteur en sortie du combineur MRC est celui vu à la section 3.3.3 [32, 188] :

$$\begin{aligned} \gamma_{MRC} &= \gamma_{sd} + \gamma_{srd} \\ &= \gamma_{sd} + \frac{\gamma_{sr}\gamma_{rd}}{\gamma_{sr} + \gamma_{rd} + 1}. \end{aligned} \quad (3.83)$$

De plus, comme le signal arrivant au relai et au destinataire est issu du même signal, ainsi montré dans l'équation 3.79, on peut obtenir la relation suivante :

$$d_{sd}^2 \gamma_{sd} = d_{sr}^2 \gamma_{sr}. \quad (3.84)$$

On obtient alors :

$$\gamma_{sr} = D\gamma_{sd} \quad (3.85)$$

avec $D = \frac{d_{sd}^2}{d_{sr}^2}$ le coefficient déjà présenté dans la section 3.4.1. En remplaçant γ_{sr} dans l'équation 3.83, on obtient alors :

$$\gamma_{MRC} = \gamma_{sd} + \frac{D\gamma_{sd}\gamma_{rd}}{D\gamma_{sd} + \gamma_{rd} + 1}. \quad (3.86)$$

Finalement, l'équation du RSB en réception ne dépend plus que de deux inconnues au lieu de trois. Or nous travaillons généralement avec γ_{MRC} fixé, il n'y a donc plus qu'une seule inconnue dans ce système. Le plus simple est de conserver γ_{sd} en tant qu'inconnue. On peut alors écrire :

$$\gamma_{rd} = (1 + D\gamma_{sd}) \left(\frac{\gamma_{MRC} - \gamma_{sd}}{D\gamma_{sd} - (\gamma_{MRC} - \gamma_{sd})} \right). \quad (3.87)$$

On se propose également de formuler une méthode directe exhaustive (ou brute-force) pour calculer le couple $(\gamma_{sd}, \gamma_{rd})$ correspondant à notre application, à savoir ici le couple conduisant à la consommation d'énergie E_b minimale. Pour cela, on commence par déterminer l'ensemble des couples $(\gamma_{sd}, \gamma_{rd})$ satisfaisant la PEB souhaitée à partir de la formule qui convient à notre canal. On calcule ensuite l'énergie totale par bit émis pour l'ensemble des couples, et on conserve l'énergie minimale.

Résultats expérimentaux

Dans un premier temps, on place le nœud-relai à la distance médiane entre le nœud-source et le nœud-destinataire. On considère toujours des transmissions à courte portée de 0.1m, 1m et 10m. On va s'intéresser d'une part à la consommation d'énergie par bit transmis, et d'autre part au gain de coopération présenté en section 3.4.4.

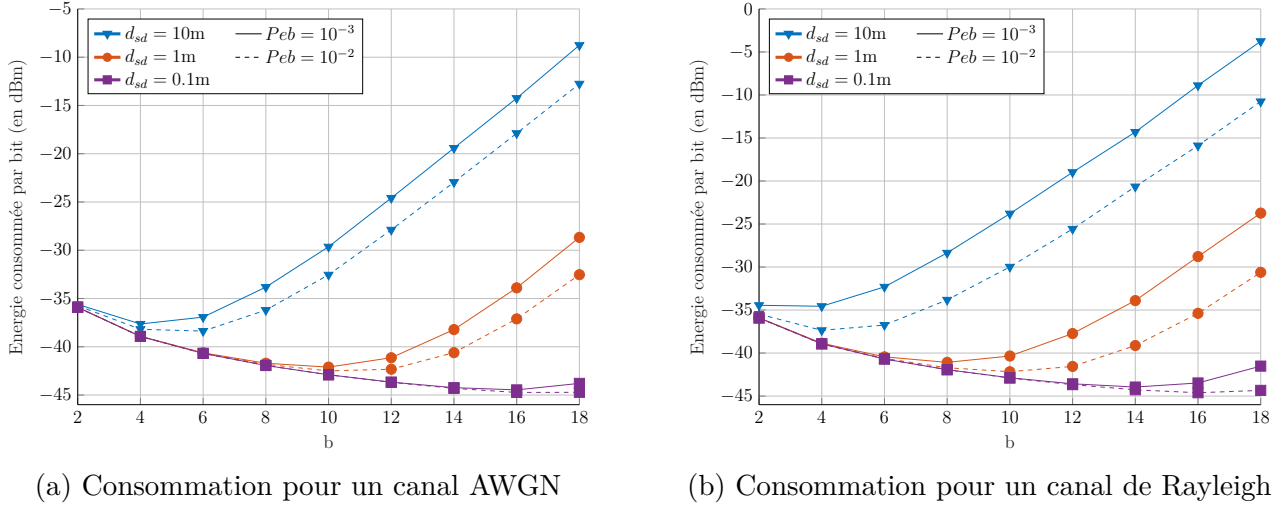
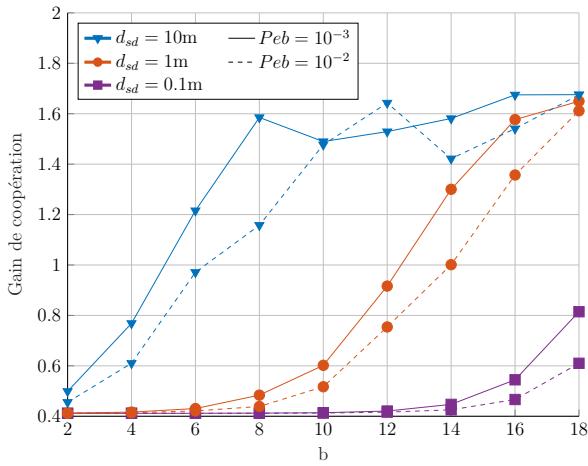


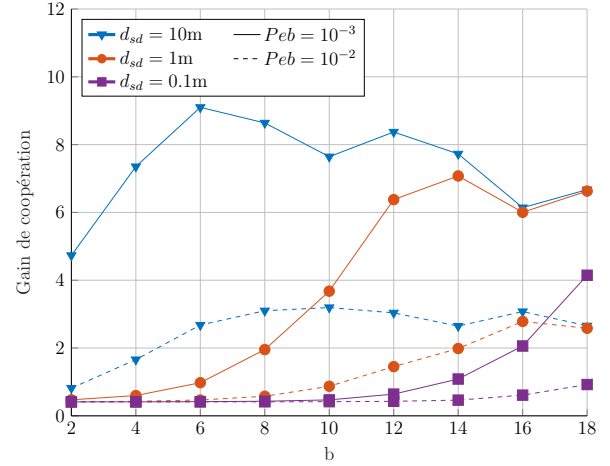
FIGURE 3.17 – Énergie consommée par bit lors d'une transmission coopérative par une M-QAM ($M = 2^b$) et un canal AWGN ou de Rayleigh en fonction de l'ordre de modulation et pour différentes PEB (nœud-relai à mi-distance source-destinataire)

La consommation d'énergie par bit pour un canal AWGN est tracée sur la figure 3.17a, et celle pour un canal de Rayleigh sur la figure 3.17b. On remarque à nouveau que, quand la PEB est plus grande, la consommation d'énergie totale diminue en raison de la baisse des RSB requis. Cette diminution de la puissance de transmission apparaît plus forte dans le cas du canal de Rayleigh, ce qui conduit à une augmentation de l'ordre b^* pour les courtes distances. En outre, on observe que ce sont pour ces mêmes distances que les courbes pour $Peb = 10^{-3}$ et $Peb = 10^{-2}$ se confondent le plus longtemps : la puissance de transmission est négligeable devant celle du circuit pour les valeurs de b jusqu'à 14 ou 16 selon la nature du canal, et c'est donc quand la situation s'inverse que l'on peut espérer trouver une valeur de b minimisant l'énergie. Pour les longues distances, en particulier à partir de 10m, la puissance de transmission est déjà importante en comparaison avec la puissance de circuit, et ce quelle que soit la PEB, par conséquent la valeur de b^* connaît moins de variation. On note également que la consommation d'énergie augmente avec la distance et que b^* diminue. Par exemple, pour le canal AWGN et $Peb = 10^{-3}$, $b^* = 4$ à 10m, $b^* = 10$ à 1m et $b^* = 16$ à 0.1m. Ceci s'explique pour les mêmes raisons, à savoir la croissance de la puissance de transmission avec la croissance. Enfin, on tient à préciser qu'en analysant les couples de RSB (γ_{sd}, γ_{rd}) pour des distances $d_{sd} \in [0.1, 1, 10, 50, 75, 100]$ m, on observe une constante : en effet, quelle que soit la distance, le couple (γ_{sd}, γ_{rd}) sera toujours le même pour une valeur de b donnée, ce qui nous donne à penser que ce couple ne dépend pas de la distance entre la source et le destinataire.

La connaissance de la consommation d'énergie par bit d'une communication coopérative, et notamment de sa consommation minimale, nous est utile pour le paramétrage de notre système. Cependant, les valeurs de b donnant cette consommation minimale sont très proches de celles



(a) Gain pour un canal AWGN

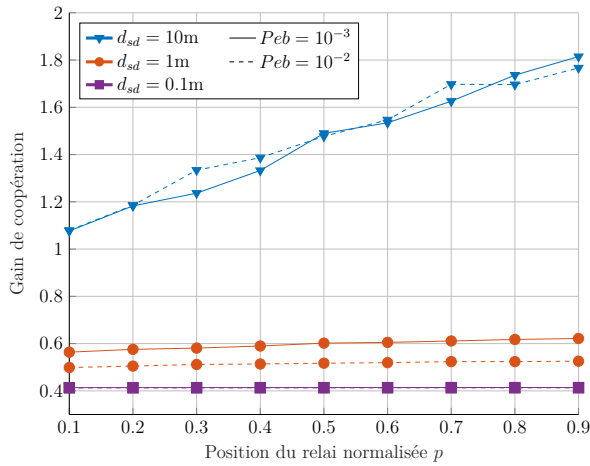


(b) Gain pour un canal de Rayleigh

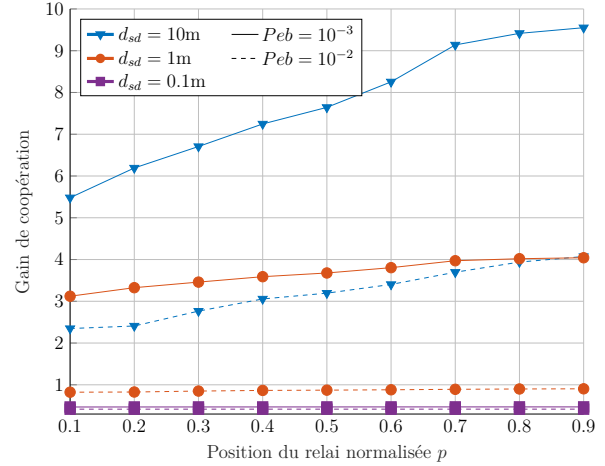
FIGURE 3.18 – Gain de coopération entre une transmission point-à-point et coopérative pour un canal AWGN ou de Rayleigh en fonction de l'ordre de modulation et pour différentes PEB (nœud-relai à mi-distance source-destinataire)

de la communication point-à-point, et nous souhaitons donc savoir s'il y a véritablement un intérêt dans l'utilisation d'un relai en fonction des paramètres sélectionnés. Dans cette optique, les figures 3.18a et 3.18b présentent le gain de coopération pour les canaux AWGN et de Rayleigh en fonction de l'ordre de modulation b . Pour le canal AWGN, on remarque que la différence de gain entre les deux PEB est assez faible, au contraire de celle pour le canal de Rayleigh. Par ailleurs, on remarque que, quelle que soit la distance, le canal ou la PEB, la tendance des courbes est plutôt à la hausse en fonction de b . Cependant pour les distances plus grandes, dans notre cas pour $d_{sd} = 10\text{m}$, il semble y avoir un palier du gain à partir d'une certaine valeur de b . Enfin, l'analyse des valeurs du gain est particulièrement intéressante. Pour le canal AWGN et pour notre gamme de valeurs de b , à 0.1m , le relai (lorsque placé à mi-distance entre la source et le destinataire) n'est jamais profitable d'un point de vue énergétique car les valeurs du gain sont toujours inférieures à 1. En revanche pour 1m et 10m , à $P_{eb} = 10^{-3}$, le gain dépasse l'unité respectivement pour $b = 14$ et $b = 6$. Pour des constellations QAM d'ordre élevé, la présence du relai s'avère donc énergétiquement avantageuse. Cependant, on remarque que b^* n'appartient pas aux valeurs de b qui minimise l'énergie dans chaque cas : le relai ne sera donc pas le meilleur choix dans ces cas si l'on souhaite une énergie minimale. En revanche, l'analyse du canal de Rayleigh nous donne des résultats plus engageants. En effet, si pour une distance de 0.1m , le gain de coopération ne dépasse l'unité qu'à partir de $b = 14$ pour $P_{eb} = 10^{-3}$ et jamais pour 10^{-2} , il le dépasse dès $b = 8$ pour $d_{sd} = 1\text{m}$ et $P_{eb} = 10^{-3}$, dès $b = 6$ pour $d_{sd} = 10\text{m}$ et $P_{eb} = 10^{-2}$, et toujours pour $d_{sd} = 10\text{m}$ et $P_{eb} = 10^{-3}$. En outre, les valeurs de gain atteintes sont bien plus grandes : pour $d_{sd} = 10\text{m}$ et $P_{eb} = 10^{-3}$, il peut atteindre jusqu'à 9.1 pour $b = 6$, et son minimum est à 4.7 quand $b = 2$. Excepté pour les très petites distances, il est donc quasiment toujours préférable d'un point de vue énergétique de choisir la solution avec relai AF.

Jusqu'à présent, le nœud-relai était placé au point médian du segment reliant source et destinataire. Afin de connaître l'influence de sa position le long de celui-ci, on va déplacer le relai à des points régulièrement espacés sur le segment. Le relai se place ainsi à une distance $d_{sr} = p d_{sd}$ de la source avec $p \in [0, 1]$ que nous nommerons position normalisée. En outre, on



(a) Gain pour un canal AWGN



(b) Gain pour un canal de Rayleigh

FIGURE 3.19 – Gain de coopération entre une transmission point-à-point et coopérative pour un canal AWGN ou de Rayleigh en fonction de la localisation du relai et pour différentes PEB

fixe arbitrairement la valeur de b à 10. Les résultats du gain de coopération en fonction de la position du relai normalisée par la distance totale d_{sd} sont affichés sur la figure 3.19a pour le canal AWGN et sur la figure 3.19b pour le canal de Rayleigh (0.1 étant la position la plus proche de la source, et 0.9 la plus proche du destinataire). A nouveau, on remarque que les courbes pour le canal AWGN sont très proches pour les deux PEB utilisées, alors que ce n'est pas le cas pour le canal de Rayleigh, en particulier pour 1m et 10m. Pour $d_{sd} \in [0.1, 1]\text{m}$, on remarque que le gain change très peu en fonction de la localisation du relai, ce qui peut s'expliquer par le fait que, comme on a pu le voir sur la figure 3.17, l'énergie du circuit est toujours majoritaire au sein de l'énergie totale pour $b = 10$. Ce n'est pas le cas pour $d_{sd} = 10\text{m}$, et on observe alors une augmentation du gain avec le déplacement du relai vers le destinataire, avec des valeurs toujours supérieures à 1 pour les deux canaux.

Un point sur la retransmission

Nous avons vu dans les expériences précédentes que le relai peut permettre d'économiser de l'énergie quand on est capable de moduler à la fois la puissance de transmission de la source et celle du relai. Néanmoins, l'utilisation du relai s'avère parfois inutile du point de vue énergétique. Une autre manière d'utiliser le relai consiste à ne lui faire retransmettre l'information qu'en cas de nécessité.

L'énergie totale peut ainsi être calculée d'après le modèle présenté par [33]. Dans ce système, la source émet un signal reçu à la fois par le relai et le destinataire. Si le RSB sur la liaison source-destinataire n'atteint pas un certain seuil, alors on considère l'information reçue comme incorrecte et la retransmission a lieu. Si un tel système n'a pas d'intérêt sur un canal AWGN au sein duquel les coefficients sont déterministes, il l'est pour le canal de Rayleigh à évanouissement plat, dont le module au carré $|h_{sd}|^2$ des coefficients de canal suit une loi exponentielle.

Les travaux présentés dans [33] propose de positionner le seuil γ_{th} par rapport au RSB symbole moyen en sortie du MRC $\overline{\gamma_{MRC}}$ tel que :

$$\gamma_{th} = bB \overline{\gamma_{MRC}} \quad (3.88)$$

avec B la largeur de la bande passante et $b = \log_2(M)$ la taille de la constellation de la M-QAM. Ici, on choisit de calculer $\overline{\gamma_{MRC}}$ à partir des formules présentées par [188] et rappelées sous leur forme généralisée dans la section 3.4.1. On considère alors le contexte de [188], c'est-à-dire un ensemble de canaux indépendants et identiquement distribués tels que les RSB moyens sont $\overline{\gamma_{sd}} = \overline{\gamma_{sr}} = \overline{\gamma_{rd}} = \gamma_0$, ce qui conduit à $\overline{\gamma_{MRC}} = 1.5\gamma_0$. A partir de la PEB, il est possible de retrouver γ_0 , et donc $\overline{\gamma_{MRC}}$.

On considère également la puissance du bruit comme équivalente sur toutes les liaisons, c'est-à-dire $\sigma_{sd}^2 = \sigma_{sr}^2 = \sigma_{rd}^2 = N_0$. On rappelle alors que le RSB instantané sur le lien direct source-destinataire, comme indiqué à l'équation 3.2, est :

$$\gamma_{sd} = \frac{P_t}{d_{sd}^2 N_0} |g_{sd}|^2 \quad (3.89)$$

avec $|g_{sd}|^2 d_{sd}^{-2} = |h_{sd}|^2$. Comme $|g_{sd}|^2$ suit une loi exponentielle, alors sa fonction de répartition nous indique que :

$$p(\gamma_{sd} \leq \gamma_{th}) = 1 - e^{-\frac{N_0 d_{sd}^2}{P_t} \gamma_{th}}. \quad (3.90)$$

En fixant une puissance d'émission identique au relai et au destinataire, c'est-à-dire $\varepsilon_r = \varepsilon_s = P_t$, on peut alors écrire le RSB instantané sur la liaison source-relai-destinataire :

$$\gamma_{srd} = \frac{\frac{P_t}{d_{sr}^2 N_0} |g_{sr}|^2 \frac{P_t}{d_{rd}^2 N_0} |g_{rd}|^2}{\frac{P_t}{d_{sr}^2 N_0} |g_{sr}|^2 + \frac{P_t}{d_{rd}^2 N_0} |g_{rd}|^2 + 1}, \quad (3.91)$$

ce qui nous donne un RSB instantané en sortie du MRC $\gamma_{MRC} = \gamma_{sd} + \gamma_{srd}$. Cette dernière variable a une fonction de répartition égale à :

$$p(\gamma_{sd} + \gamma_{srd} \leq \gamma_{th}) = \int_0^{\gamma_{th}} p(\gamma_{srd} \leq \gamma_{th} - \gamma_{sd} | \gamma_{sd}) f(\gamma_{sd}) d\gamma_{sd} \quad (3.92)$$

avec $f(\gamma_{sd}) = \frac{N_0 d_{sd}^2}{P_t} e^{-\frac{N_0 d_{sd}^2}{P_t} \gamma_{sd}}$, ce qui conduit à

$$p(\gamma_{srd} \leq \gamma_{th} - \gamma_{sd} | \gamma_{sd}) = 1 - \sqrt{\xi} \exp\left(-(\gamma_{th} - \gamma_{sd}) \left(\frac{N_0 d_{sr}^2}{P_t} + \frac{N_0 d_{rd}^2}{P_t}\right)\right) K_1(\sqrt{\xi}) \quad (3.93)$$

avec $\xi = \frac{4((\gamma_{th} - \gamma_{sd})^2 + (\gamma_{th} - \gamma_{sd})) N_0^2 d_{sr}^2 d_{rd}^2}{P_t^2}$ et K_1 la fonction de Bessel modifiée de deuxième espèce et du premier ordre.

Finalement, on peut établir la probabilité que le signal soit correctement transmis :

$$p_{succes} = 1 - p(\gamma_{sd} \leq \gamma_{th}) p(\gamma_{sd} + \gamma_{srd} \leq \gamma_{th}). \quad (3.94)$$

On peut également reformuler l'énergie moyenne consommée de l'équation 3.81. En effet, la deuxième transmission n'a lieu que dans le cas d'un échec à la première, c'est-à-dire que $\gamma_{sd} \leq \gamma_{th}$:

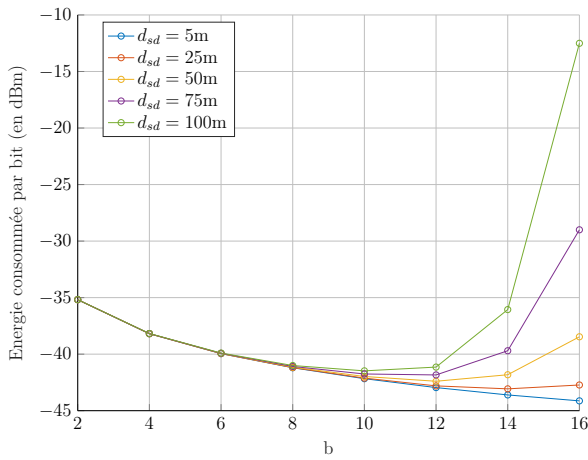
$$E_b = \frac{P_{on,1} T_{on} + P_{tr,1} T_{tr}}{L} p(\gamma_{sd} \geq \gamma_{th}) + \frac{P_{on,2} T_{on} + P_{tr,2} T_{tr}}{L} p(\gamma_{sd} \leq \gamma_{th}). \quad (3.95)$$

Avec ces données, il est possible de calculer ce qu'on appelle l'énergie de succès, c'est-à-dire l'énergie consommée par bit transmis avec succès. Celle-ci se calcule comme :

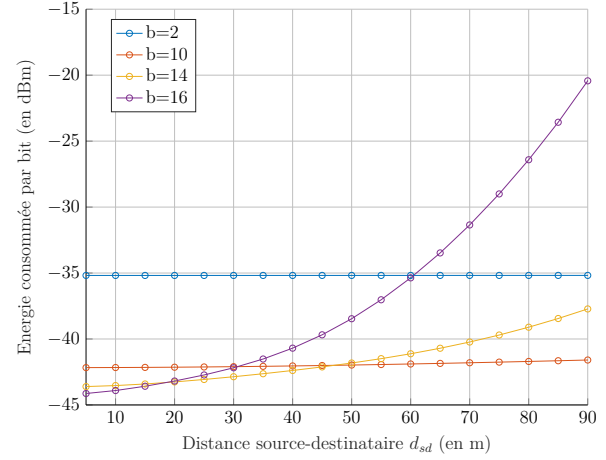
$$E_{succes} = \frac{E_b}{p_{succes}}. \quad (3.96)$$

A noter que pour calculer le gain de coopération, on utilise l'énergie de succès d'une transmission directe, qui se note alors :

$$E_{direct} = \frac{P_{on,1}T_{on} + P_{tr,1}T_{tr}}{L p(\gamma_{sd} \geq \gamma_{th})}. \quad (3.97)$$



(a) Énergie consommée en fonction de la taille de la constellation



(b) Énergie consommée en fonction de la distance source-destinataire

FIGURE 3.20 – Énergie consommée par bit transmis avec succès lors d'une communication coopérative sur des canaux de Rayleigh et avec un signal modulé par M-QAM

On souhaite étudier la consommation d'énergie qui résulte d'un tel système. On reprend les mêmes paramètres que dans les expériences précédentes, et on place le relai à la distance médiane entre la source et le destinataire. On fixe la puissance de transmission de la source et du relai à $P_t = 100mW$.

Dans un premier temps, on étudie uniquement la consommation de la communication coopérative. L'énergie consommée par bit transmis avec succès en fonction de la taille de la constellation est tracée sur la figure 3.20a, et en fonction de la distance source-destinataire sur la figure 3.20b. Pour des très petites distance inférieures à 5m, il n'est pas possible d'observer une différence dans la consommation ; on trace donc nos résultats pour des longues distances comprises entre 5 et 100m. Sur la figure 3.20a, on remarque pour les faibles distances entre la source et le relai, l'énergie consommée diminue inversement relativement à l'ordre de la M-QAM. Cela est lié au fait que, lorsqu'on augmente l'ordre b , le débit augmente et donc T_{on} diminue, ce qui a pour conséquence de diminuer la consommation d'une transmission. Cela augmente aussi le RSB de seuil γ_{th} , mais cela a peu d'impact pour les faibles distances puisque le RSB de réception reste haut. En revanche, plus la distance augmente, plus celui-ci baisse (car nous transmettons à puissance constante), et la probabilité qu'il passe sous le seuil d'acceptabilité augmente. Ainsi, la probabilité d'une retransmission, assez coûteuse, croît elle aussi,

ce qui conduit à une augmentation de l'énergie consommée. La figure 3.20b met cette fois-ci en perspective l'évolution de l'énergie en fonction de la distance de transmission, et ceci pour plusieurs ordres de modulation. Pour les tailles de constellation inférieures à $b = 12$, comme on peut l'observer sur la figure précédente, la distance a peu d'importance car la probabilité de succès est toujours importante grâce à un γ_{th} relativement faible. Pour les larges constellations, ce n'est plus le cas et la distance augmente l'énergie consommée.

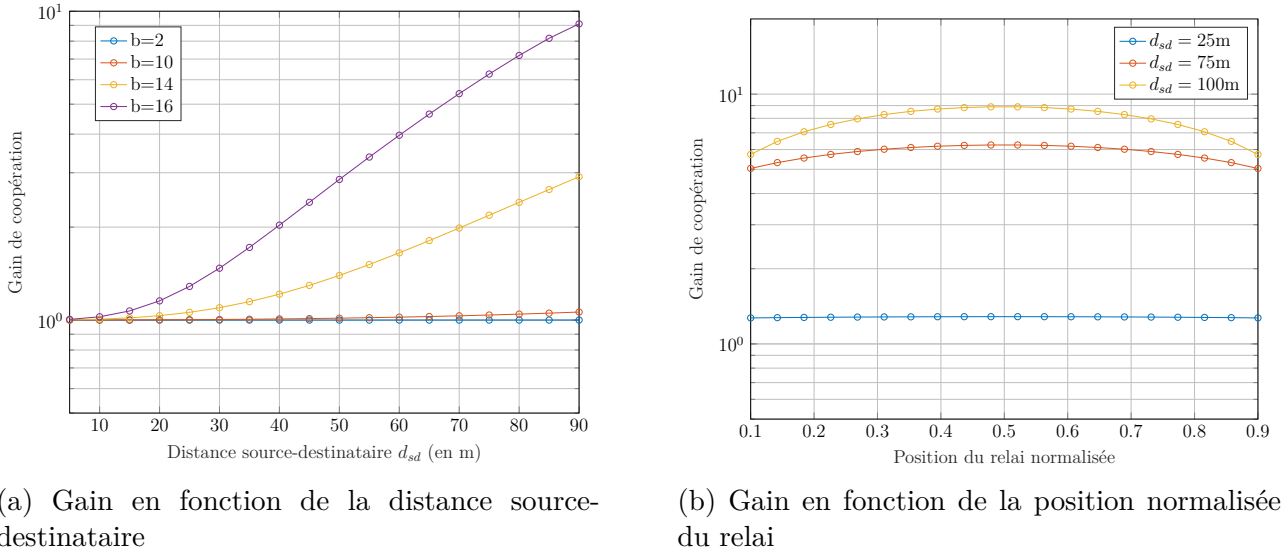


FIGURE 3.21 – Gain de coopération pour l'énergie de transmission avec succès lors d'une communication sur des canaux de Rayleigh et avec un signal modulé par M-QAM

Dans un second temps, on étudie le gain de coopération pour différentes distances et constellations. Sur la figure 3.21a, le relai est considéré positionné à la distance médiane entre la source et le destinataire, et on observe le gain en fonction de la distance source-destinataire. Lorsque l'efficacité spectrale b est inférieure à 12, on a déjà vu que l'énergie reste similaire en fonction de la distance car la liaison source-destinataire est fiable ; il n'y a donc qu'une seule transmission d'effectuée et l'énergie consommée est similaire à celle dans le cas d'une transmission directe. Lorsque l'efficacité spectrale b est supérieure à 12, il y a retransmission du signal via le relai, et l'énergie consommée croît de manière plus importante. Cependant, cette retransmission permet un taux de bonne réception supérieur à celui de la transmission directe : comme nous considérons l'énergie consommée pour une transmission avec succès, la consommation du système avec relai, bien que supérieure dans l'absolu, reste avantageuse par rapport à celle du système sans relai. Ces résultats se retrouvent à nouveau sur la figure 3.21b. Cette fois-ci, nous fixons $b = 16$ pour pouvoir observer des différences d'énergie en fonction de la distance, et on fait varier la position du relai sur l'axe source-destinataire. On affiche les courbes du gain coopératif pour plusieurs distances en fonction de la position normalisée du relai. Pour une faible distance, la transmission directe suffit à une réception fiable, donc le gain est unitaire. En revanche, quand la distance entre la source et le destinataire augmente, le gain de coopération va également croître. Sur les courbes pour $d = 75m$ et $d = 100m$, on observe que le gain est symétrique par rapport à la position médiane, avec une croissance puis une décroissance ; le gain maximal est ainsi obtenu pour une position normalisée du relai à 0.5.

3.6 Lien entre énergie et efficacité spectrale

Dans la section 3.4.3, nous avons établi qu'il existe un lien entre efficacités spectrale et énergétique. L'efficacité énergétique (EE) pouvant elle-même être calculée à partir de l'énergie E_b , il est donc possible de lier la consommation d'énergie par bit transmis et l'efficacité spectrale (ES). En effet, en approximant $\theta \approx b$ avec θ l'ES et $b = \log_2(M)$ la taille de constellation de la M-QAM, nous avons pu observer empiriquement dans la section 3.5.1 qu'il existe une valeur spécifique de l'ES qui maximise l'EE. Cette valeur semble changer en fonction de l'importance de l'énergie de transmission au sein de l'énergie totale. Ceci semble cohérent au vu du lien entre ordre de modulation et RSB établi dans la section 3.4.1, RSB dont dépend l'énergie de transmission.

Dans cette section, on s'intéresse uniquement à une communication point-à-point. Dans un premier temps, on va explicitement présenter le lien entre énergie et ES en s'aidant de la capacité de Shannon. Celle-ci n'étant valable que pour une source gaussienne non codée et un canal AWGN, on étendra ensuite cette réflexion au cas d'un signal modulé par M-QAM et pour un canal AWGN ou de Rayleigh à évanouissement plat. Enfin, on expliquera comment trouver la valeur d'ES minimisant l'énergie totale par bit, c'est-à-dire maximisant l'EE.

3.6.1 Lien entre consommation d'énergie et capacité

Pour Shannon, la capacité C d'un canal AWGN dont l'entrée est une source gaussienne s'exprime $C = B \log_2 \left(1 + \frac{P}{N_0 B}\right)$ avec P la puissance au niveau de l'antenne réceptrice, qui correspond dans notre modèle à P_r/N_f [156]. On obtient alors

$$C = B \log_2 \left(1 + \frac{P_r}{N_f N_0 B}\right). \quad (3.98)$$

L'équation 3.76 nous donne en outre la relation $P_r = B N_0 N_f \gamma_{sd}$. On peut donc toujours établir une relation entre la capacité et le RSB, comme expliqué dans la section 3.4.3 :

$$C = B \log_2 (1 + \gamma_{sd}). \quad (3.99)$$

Si on considère que le débit R peut atteindre la capacité, soit $R = C$ et l'équation 3.62 de l'ES $\theta = R/B$, alors le RSB peut s'exprimer

$$\gamma_{sd} = 2^\theta - 1. \quad (3.100)$$

En considérant la formule 3.78 de la consommation d'énergie par bit totale, on peut finalement définir l'énergie consommée en fonction de l'ES :

$$E_b = (1 + \alpha) G_d N_0 N_f \frac{2^\theta - 1}{\theta} + \frac{P_c}{\theta B} + \frac{P_{tr} T_{tr}}{L} \quad (3.101)$$

avec α , G_d , N_0 , N_f , L , P_{tr} et T_{tr} définis dans les sections précédentes. Ce modèle sera également appelé (A.gauss.trad). On remarque alors que, outre l'énergie de transmission, l'énergie du circuit est elle aussi dépendante de l'ES et impacte donc le compromis entre ES et EE, comme cela a déjà été démontré [40].

Pour les valeurs d'ES considérées, on a déjà expliqué dans la section 3.5.1 que l'énergie de transmission croît avec θ et que l'énergie du circuit décroît avec ce même paramètre, résultant en une énergie totale qui semble admettre un minimum global (voir figure 3.15). En effet, on peut prouver que l'énergie E_b est une fonction convexe de θ . Ce résultat, qui a déjà été démontré par [87], est redémontré en annexe C.1.

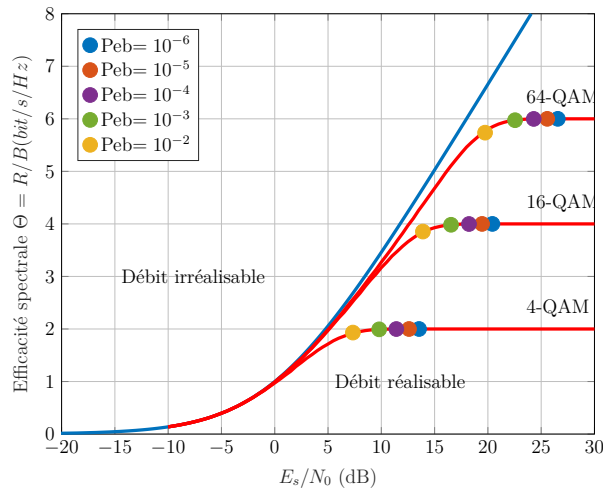
3.6.2 Capacité pour un signal modulé

La capacité de Shannon s'exprime pour un canal AWGN présentant à son entrée une source gaussienne, il ne prend donc pas en compte la modulation ou les évanouissements rapides. Cependant, la modulation ou le choix du canal peuvent tous les deux affecter la capacité du canal [109]. Dans nos travaux, nous considérons une M-QAM avec $M = 2^b$ niveaux d'amplitude. Comme indiqué précédemment, nous faisons l'hypothèse que le rapport E_s/N_0 est suffisamment fort pour considérer que l'égalité $\theta \approx b$ [89] est valide. Aussi, dans la suite de ce chapitre, nous utiliserons indistinctement b ou θ pour désigner l'efficacité spectrale.

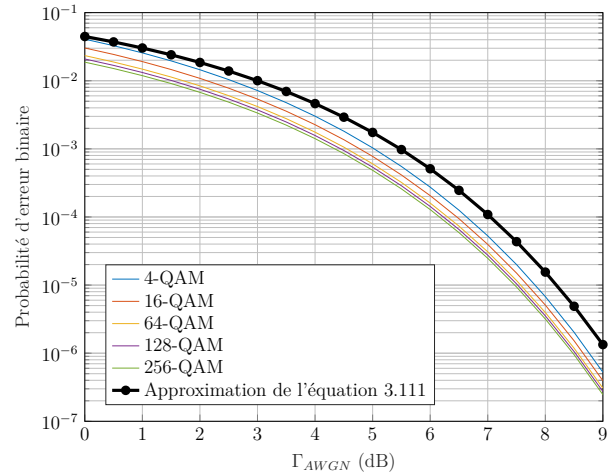
Capacité pour un canal AWGN

Dans un premier temps, nous cherchons à connaître la différence entre la capacité d'un canal AWGN avec ou sans modulation. Dans cette optique, nous allons estimer le RSB symbole moyen en réception d'un canal AWGN avec une M-QAM, et ce quel que soit l'ordre de modulation. Nous rappelons que, conformément à l'équation 3.99, l'ES θ peut être approximée par :

$$\theta \approx \log_2(1 + \gamma_{sd}). \quad (3.102)$$



(a) Efficacité spectrale en fonction du RSB symbole pour une source gaussienne (en bleu) ou une source modulée par M-QAM (en rouge)



(b) PEB en fonction de la pénalité en RSB symbole Γ_{AWGN} pour une source modulée dans un canal AWGN

FIGURE 3.22 – Efficacité spectrale en fonction du RSB symbole et PEB en fonction de la pénalité en RSB Γ_{AWGN} pour une source modulée par M-QAM dans un canal AWGN

La capacité instantanée d'un canal AWGN est définie comme l'information mutuelle maximale entre l'entrée et la sortie du canal :

$$\begin{aligned} C_{AWGN} &= \max_{p(X)} I(X; Y) \\ &= \max_{p(X)} \int \int p(x, y) \log_2 \left(\frac{p(x, y)}{p(x)p(y)} \right) \partial x \partial y \end{aligned} \quad (3.103)$$

où X et Y désignent les processus aléatoires en entrée et sortie du canal de transmission, $I(.;.)$ l'opération d'information mutuelle, et $p(x)$, $p(y)$ et $p(x, y)$ respectivement les densités de

probabilité du signal source X , de la sortie Y et de leur loi conjointe (X, Y) . En explicitant ces densités de probabilité, on obtient finalement une capacité

$$C_{AWGN} = \log_2(M) - \frac{1}{M} \sum_{k=1}^M \mathbb{E} \left(\log_2 \left(\sum_{p=1}^M \exp \left(-\frac{\|x_k - x_p + n\|^2 - \|n\|^2}{N_0} \right) \right) \right) \quad (3.104)$$

avec n un bruit blanc gaussien complexe suivant une loi $\mathcal{N}(0, N_0/2) + j\mathcal{N}(0, N_0/2)$ et les symboles x_k et x_p appartenant tous deux à l'alphabet des symboles modulés (ici par une modulation M-QAM).

La simulation de l'ES, ou capacité C_{AWGN} normalisée par B , est tracée en fonction du RSB symbole sur la figure 3.22a pour des ordres de modulation M de 4, 16 et 64 à l'aide de l'équation 3.104. Les courbes résultantes sont comparées avec l'ES obtenue pour une source gaussienne à l'équation 3.100. Pour chaque RSB symbole et chaque ordre de modulation, on peut calculer la PEB sur le canal; on spécifie sur chaque courbe le point correspondant à $\text{Peb} \in [10^{-2}, 10^{-3}, 10^{-4}, 10^{-5}, 10^{-6}]$. On remarque que chaque courbe de θ suit la courbe de la source gaussienne jusqu'à atteindre une limite égale à $\log_2(M)$. En outre, ce seuil semble être atteint pour tout $\text{Peb} < 10^{-3}$, c'est-à-dire pour les faibles valeurs de PEB, ce qui est cohérent avec notre approximation $\theta \approx b$.

Dans un second temps, on s'intéresse à la formule du PEB pour notre canal et à son lien avec la capacité de Shannon. Conformément à l'équation 3.43, la PEB pour un canal AWGN modulé par M-QAM peut s'écrire :

$$\text{Peb} \approx \frac{4}{b} \left(1 - \frac{1}{\sqrt{2^b}} \right) \mathcal{Q} \left(\sqrt{\frac{3\gamma_{sd}}{2^b - 1}} \right). \quad (3.105)$$

Il en résulte un RSB :

$$\gamma_{sd} = (2^b - 1) \frac{1}{3} Q^{-1} \left(b \frac{\text{Peb}}{4} \left(\frac{\sqrt{2^b}}{\sqrt{2^b} - 1} \right) \right)^2. \quad (3.106)$$

Dans cette dernière équation, nous pouvons noter la présence du terme $2^b - 1$ du RSB obtenu au récepteur sur un canal AWGN pour une source gaussienne à son entrée selon l'équation 3.100. En définissant la pénalité en RSB par rapport au cas d'une source gaussienne par

$$\Gamma_{AWGN}(\text{Peb}, b) = \frac{1}{3} Q^{-1} \left(b \frac{\text{Peb}}{4} \left(\frac{\sqrt{2^b}}{\sqrt{2^b} - 1} \right) \right)^2, \quad (3.107)$$

on obtient une nouvelle forme du RSB :

$$\gamma_{sd} = (2^b - 1) \Gamma_{AWGN}(\text{Peb}, b). \quad (3.108)$$

Cependant, Γ_{AWGN} dépend de b et est difficile à manipuler, notamment en raison de la fonction $\mathcal{Q}$. Pour contrer ce problème, [89] propose d'utiliser une forme indépendante de b et inspirée de la formule 3.46 de la PEB :

$$\Gamma_{AWGN}(\text{Peb}) = \frac{2}{3} \ln \left(\frac{1}{5\text{Peb}} \right). \quad (3.109)$$

En outre, la pénalité Γ_{AWGN} en RSB étant définie comme le RSB du canal AWGN avec source M-QAM normalisé par $2^b - 1$, on a l'égalité :

$$\Gamma_{AWGN}(\text{Peb}) = \frac{\gamma_{sd}}{2^b - 1} \quad (3.110)$$

On obtient finalement une formule de la PEB :

$$\text{Peb} = \frac{1}{5} \exp\left(-\frac{3}{2}\Gamma_{AWGN}\right). \quad (3.111)$$

La PEB est tracée en fonction de $\Gamma_{AWGN} \in [0; 9]$ pour différents ordres de modulation sur la figure 3.22b. La courbe de PEB obtenue à partir de 3.111 est également tracée en comparaison. On remarque en premier lieu que la normalisation du RSB par $2^b - 1$ réduit beaucoup la dépendance de la PEB à l'ordre de modulation car les courbes expérimentales sont toutes très proches l'une de l'autre. Ceci confirme la possibilité d'utiliser une formule de la pénalité indépendante de cet ordre. En second lieu, la formule proposée par [89] semble modéliser correctement cette pénalité, et ce malgré un certain décalage pour les plus hautes valeurs de RSB.

Capacité pour un canal de Rayleigh à évanouissement plat

Nous souhaitons mener une analyse similaire à celle décrite au paragraphe précédent mais en considérant maintenant un canal de Rayleigh à évanouissement plat, dont la capacité a déjà été étudiée. Ainsi, les travaux menés dans [8] proposent une borne supérieure de cette capacité dans le cas d'une source gaussienne :

$$\theta \leq \log_2(e) E_1(1/\gamma_{sd}) e^{1/\gamma_{sd}} \quad (3.112)$$

avec $E_1(x) = \int_1^{+\infty} \frac{e^{-xt}}{t} dt$. [110] nous rappelle alors les bornes inférieures et supérieures suivantes :

$$\frac{1}{2} \ln(1 + 2\gamma_{sd}) \leq E_1(1/\gamma_{sd}) e^{1/\gamma_{sd}} \leq \ln(1 + \gamma_{sd}) \quad (3.113)$$

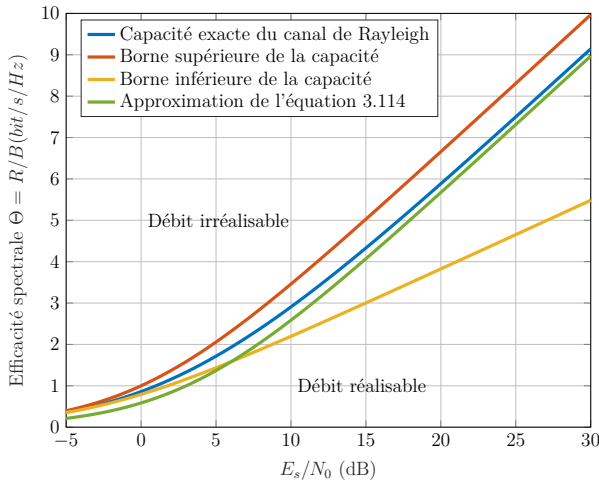
puis, sur le même modèle, suggère que :

$$E_1(1/\gamma_{sd}) e^{1/\gamma_{sd}} \approx \ln\left(1 + \frac{\gamma_{sd}}{2}\right). \quad (3.114)$$

La formulation exacte de l'équation 3.112, ses bornes supérieure et inférieure de l'équation 3.113 et l'approximation 3.114 sont tracées en fonction du RSB symbole sur la figure 3.23a. D'une part, les courbes nous indiquent que la capacité exacte est bien comprise entre ses bornes inférieure et supérieure. D'autre part, l'équation 3.114 semble une approximation satisfaisante de la capacité, excepté pour les faibles valeurs de RSB, pour lesquelles la borne inférieure est plus proche de la courbe exacte.

En acceptant l'équation 3.114 comme approximation de l'efficacité spectrale pour un signal non codé, on obtient alors une nouvelle forme du RSB pour le canal de Rayleigh :

$$\gamma_{sd} = 2(2^\theta - 1). \quad (3.115)$$



(a) Efficacité spectrale en fonction du RSB symbole pour une source gaussienne

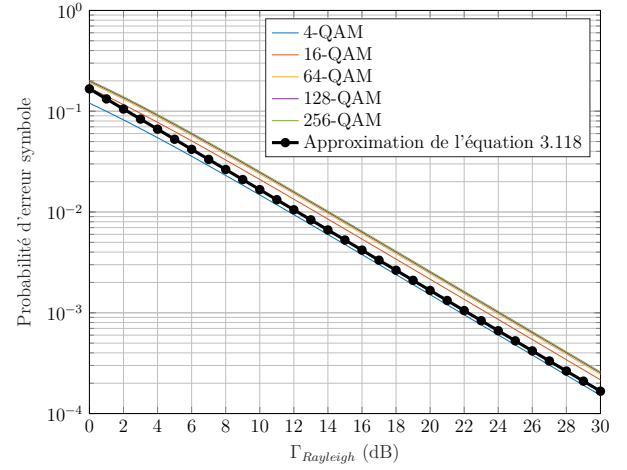

 (b) PEB en fonction de la pénalité en RSB symbole $\Gamma_{Rayleigh}$ pour une source modulée dans un canal de Rayleigh

FIGURE 3.23 – Efficacité spectrale en fonction du RSB symbole pour une source gaussienne et PEB en fonction de la pénalité du RSB liée à l'utilisation d'une source M-QAM au lieu d'une source gaussienne pour un canal de Rayleigh à évanouissement plat

Cette forme est le double de celle obtenue pour le canal AWGN. On cherche finalement une formule de la pénalité en RSB $\Gamma_{Rayleigh}(\text{Peb})$ définie de manière analogue à $\Gamma_{AWGN}(\text{Peb})$:

$$\Gamma_{Rayleigh}(\text{Peb}) = \frac{\gamma_{sd}}{2(2^b - 1)}. \quad (3.116)$$

Néanmoins, à notre connaissance, il n'existe aucune forme de $\Gamma_{Rayleigh}$ dépendante de la PEB. [87] propose d'utiliser une pénalité dépendante de la PES :

$$\Gamma_{Rayleigh}(\text{Pes}) = \frac{1}{6\text{Pes}}. \quad (3.117)$$

On obtient alors une PES issue de l'équation 3.50 :

$$\text{Pes} = \frac{1}{6\Gamma_{Rayleigh}}. \quad (3.118)$$

Sur la figure 3.23b nous avons représenté les courbes donnant la PES en fonction de la pénalité en RSB symbole pour différentes tailles d'alphabet ainsi que son approximation obtenue à partir de la relation 3.118. Comme pour le canal AWGN, on remarque que toutes les courbes expérimentales sont très proches les unes des autres, et ce quel que soit l'ordre de modulation.

En outre, l'approximation de la PES donnée par la relation 3.118 semble correctement modéliser celle-ci, notamment pour des modulations M-QAM de tailles supérieures à 128-QAM. Cependant, des écarts de l'ordre de 2 dB peuvent être observés sur la valeur de la pénalité pour des tailles de constellation proches de 2 et ceci quelle que soit la valeur considérée de la PES cible.

3.6.3 Minimisation de la consommation

Nous avons établi le lien qui unit énergie et ES, ainsi que celui qui unit ES et RSB symbole. Il est donc possible de définir une expression de l'énergie consommée par bit transmis qui dépende simplement de la probabilité d'erreur pour un signal source modulé par M-QAM et transmis sur un canal AWGN ou de Rayleigh. Pour le canal AWGN, l'énergie peut être exprimée en fonction de la PEB :

$$E_b = (1 + \alpha)G_d N_0 N_f \Gamma_{AWGN}(\text{Peb}) \frac{(2^b - 1)}{b} + \frac{P_c}{bB} + \frac{P_{tr} T_{tr}}{L} \quad (3.119)$$

et pour le canal de Rayleigh, l'énergie peut être exprimée en fonction de la PES :

$$E_b = (1 + \alpha)G_d N_0 N_f \Gamma_{Rayleigh}(\text{Pes}) \frac{2(2^b - 1)}{b} + \frac{P_c}{bB} + \frac{P_{tr} T_{tr}}{L}. \quad (3.120)$$

Ces modèles de l'énergie, que l'on nommera respectivement (A.qam.trad) et (R.qam.trad), peuvent être généralisés à tout type de canal à condition de pouvoir déduire la pénalité associée en fonction de la probabilité d'erreur [87]. Ils permettent d'obtenir une expression analytique de la consommation d'énergie par bit émis sans avoir à exprimer la fonction réciproque de la probabilité d'erreur pour obtenir le RSB symbole. Ceci est d'autant plus important que certaines relations n'admettent pas de fonctions réciproques, comme l'équation 3.47. Pour vérifier que ces modèles sont fidèles à l'équation 3.78, on compare les courbes issues du modèle exact (E.trad) et du modèle avec pénalité (A.qam.trad) ou (R.qam.trad) sur la figure 3.24. Pour le canal AWGN, on fixe $\text{Peb} = 10^{-3}$ et pour le canal de Rayleigh, on fixe $\text{Pes} = 10^{-3}$. On note que les deux courbes correspondant au même contexte sont toujours très proches, avec une différence maximale de 2.57dBm pour le canal AWGN et de 2.58dBm pour le canal de Rayleigh.

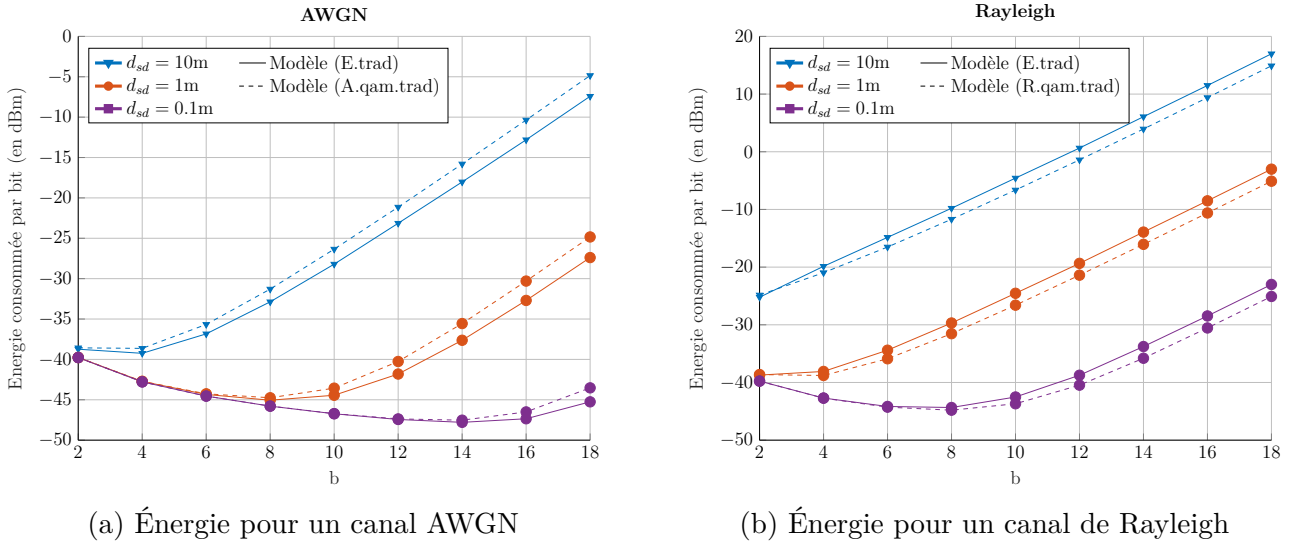


FIGURE 3.24 – Comparaison de l'énergie consommée par bit pour le modèle exact et le modèle calculé à partir de la pénalité

Les équations 3.119 et 3.120, grâce à leur manipulation aisée, nous permettent de mieux détailler l'existence d'une ES θ^* (ou b^* pour la M-QAM) minimisant l'énergie totale, ainsi que l'influence des différents paramètres sur cette quantité. En effet, on remarque que l'énergie

de transmission augmente exponentiellement avec l'ES, tandis que l'énergie du circuit décroît linéairement avec elle. Dans la suite de cette section, nous souhaitons calculer analytiquement cette valeur b^* , qui sera donc définie par :

$$b^* = \underset{b}{\operatorname{argmin}} E_b(b). \quad (3.121)$$

Dans un premier temps, nous devons déterminer si b^* est unique ou non. C'est dans cette optique que [87] prouve la quasi-convexité de l'équation 3.119 par rapport à b . Cette démonstration, reproduite dans l'annexe C.2, nous permet non seulement de s'assurer de l'unicité de la solution, mais surtout de l'existence de celle-ci dans tous les cas. En outre, comme l'équation 3.120 diffère de 3.119 seulement d'un facteur 2 dans le premier terme de la somme, la quasi-convexité est conservée. Dans un second temps, comme les deux formules de l'énergie sont différentiables, alors l'argument du minimum global de E_b est l'unique solution de

$$\frac{\partial E_b}{\partial b}(b) = 0. \quad (3.122)$$

Pour un canal AWGN, la solution égalité, qui sera nommée (Bopt.A.trad), est donnée par [87]

$$b^* = \frac{1}{\ln(2)} \left(1 + W \left(\frac{c - a}{ae} \right) \right) \quad (3.123)$$

avec $a = (1 + \alpha)N_0N_f\Gamma_{AWGN}(\text{Peb})\frac{(4\pi d_{sd})^2}{G\lambda^2}$, $c = \frac{P_c}{B}$, $e = \exp(1)$ et W la fonction W de Lambert. En fixant $a' = (1 + \alpha)N_0N_f2\Gamma_{Rayleigh}(\text{Pes})\frac{(4\pi d_{sd})^2}{G\lambda^2}$, nous pouvons trouver la solution pour un canal de Rayleigh, qui sera nommée (Bopt.R.trad), de manière analogue :

$$b^* = \frac{1}{\ln(2)} \left(1 + W \left(\frac{c - a'}{a'e} \right) \right). \quad (3.124)$$

La démonstration de ces solutions, déjà présentée par [87], est également donnée dans l'annexe **bonne réf.**

En théorie, l'ES peut prendre toutes les valeurs réelles positives. En pratique, l'ordre de modulation de la M-QAM est choisi comme $M = 2^b$ avec b un entier pair strictement positif. La transformation pour obtenir une valeur de l'ES parmi celles autorisées est alors :

$$b_d^* = \max(2, 2 \left\lfloor \frac{b^*}{2} + \frac{1}{2} \right\rfloor). \quad (3.125)$$

La figure 3.25 illustre l'évolution de l'ES discrète optimale b_d^* minimisant la consommation d'énergie par bit d'une transmission avec M-QAM en fonction de la distance d_{sd} entre la source et le destinataire. Pour le canal AWGN, on compare les résultats issus d'un algorithme de recherche exhaustive avec ceux issus de l'équation 3.123 pour une valeur de la PEB de 10^{-3} . On observe que l'ES diminue avec la distance, comme on en avait fait l'hypothèse en visualisant les courbes de la consommation. On remarque que les deux courbes sont très proches l'une de l'autre, avec seulement un léger décalage au passage d'une valeur de l'ES à une autre, qui se manifeste par une diminution plus rapide de l'expression analytique 3.123 que pour l'algorithme de recherche. En outre, ce décalage augmente avec la distance. Pour le canal de Rayleigh et une valeur de la PES de 10^{-3} , on observe cette même décroissance de l'ES. Cependant, les résultats de l'expression analytique sont beaucoup moins satisfaisants : en effet, il existe un décalage important entre la courbe théorique et celle de l'algorithme de recherche, ce qui conduit à des valeurs de l'ES qui ne coïncident avec les valeurs réelles qu'à partir d'une distance de 1.6m, quand l'ES atteint de toute manière sa valeur minimale.

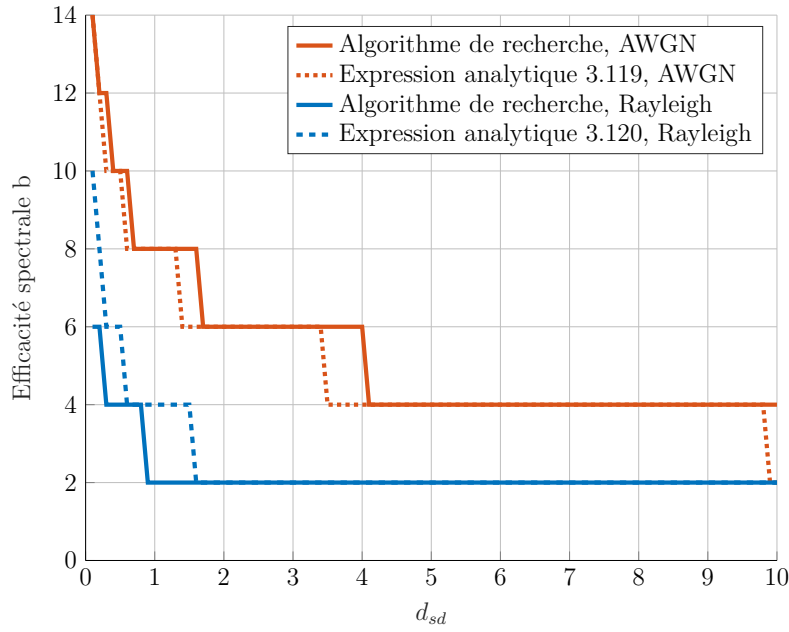


FIGURE 3.25 – Évolution de l'efficacité spectrale b^*_d minimisant la consommation d'énergie par bit pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction de la distance d_{sd}

3.7 Conclusion

Dans ce chapitre, nous avons présenté un état de l'art exhaustif de la consommation d'énergie pour une communication sans fil. Nous y avons détaillé le modèle énergétique d'une communication, qu'elle soit point-à-point ou coopérative. Auparavant, pour comprendre la nature et l'importance des paramètres qui impactent cette consommation, nous avons décrit les éléments qui composent une communication point-à-point puis coopérative. Nous avons notamment abordé ces notions importantes que sont la modulation, la nature du canal ou encore le protocole de coopération.

Afin de pouvoir minimiser l'énergie tout en gardant une qualité de service minimale, nous avons également énuméré les différents critères de performance qui définissent la qualité de service, et notamment l'efficacité spectrale qui est d'une grande importance à une période où les ressources en bande passante sont limitées. Par la suite, la capacité du canal nous a permis de déterminer un lien entre efficacité spectrale et efficacité énergétique, deux notions qui nécessitent un compromis l'une avec l'autre. Ce lien est dépendant du système utilisé, en particulier de la modulation et du canal. Pour une M-QAM et des canaux AWGN et de Rayleigh à évanouissement plat, il nous a ainsi été possible de présenter une formule analytique de l'efficacité spectrale maximisant l'efficacité énergétique, c'est-à-dire minimisant l'énergie par bit transmis lors de la communication.

L'ensemble de ces éléments sont d'une grande importance dans la maîtrise de la consommation d'un système de classification sonore embarqué. Néanmoins, nous avons vu que le calcul de l'énergie consommée se base sur de nombreuses approximations, comme celui de la probabilité d'erreur, du coefficient de l'amplificateur de puissance, ou encore de la puissance des CNA et CAN. Dans le chapitre suivant, on va s'intéresser à détailler ce modèle pour en obtenir une interprétation plus réaliste.

ANALYSE D'UN MODÈLE DE CONSOMMATION RÉALISTE

4.1 Introduction

Dans le chapitre précédent, nous avons présenté le modèle de consommation traditionnel d'une communication point-à-point ou coopérative. Néanmoins, ce modèle présente un certain nombre d'imprécisions qui peuvent mettre à mal le calcul de la consommation. Par exemple, la saturation de la puissance de transmission n'est jamais prise en compte, et de nombreux travaux agissent comme si on pouvait alimenter l'amplificateur de puissance de manière illimitée. Ces modèles intègrent des puissances du CAN et du CNA constantes alors qu'elles dépendent de plusieurs paramètres peu utilisés de ces deux modules [41, 87, 111]. Le coefficient de l'AP est lui aussi considéré comme fixe alors que l'efficacité du drain et le PAPR qui le composent ont été prouvés dépendants des paramètres du système [44, 120, 207, 227]. Quelques recherches ont utilisé ces formes, parfois simplifiées [33, 41, 97], mais jamais encore ces éléments n'ont été rassemblés ensemble dans un unique modèle de consommation. En outre, le coefficient de l'AP dépendant du système n'a jamais été utilisé dans le cadre du calcul de l'ES maximisant l'EE pour une consommation d'énergie point-à-point. En s'inspirant des travaux menés par [87] et détaillés dans la section 3.6.3, nous allons chercher à exprimer la notion de l'ES en utilisant un modèle de consommation plus réaliste. Ce sera aussi l'occasion d'interroger le modèle de la capacité pour les canaux AWGN et de Rayleigh. En effet, l'absence de formule de la pénalité dépendant de la PEB pour le canal de Rayleigh est une lacune qu'il convient de combler.

Ce chapitre sera donc organisé comme ci-après. Dans la première section, nous nous intéresserons aux précisions qui ont été apportées au modèle traditionnel de la consommation, que ce soit pour l'AP ou le CAN et le CNA. Dans une deuxième section, nous adapterons les expressions analytiques de l'ES minimisant la consommation d'énergie eu sein de ce nouveau modèle énergétique. Nous analyserons plus particulièrement la capacité d'un canal de transmission (AWGN ou de Rayleigh à évanouissement plat) pour une source modulée QAM et nous proposerons des expressions analytiques précises de la pénalité en RSB (par rapport à une source gaussienne). Ces expressions nous permettront de caractériser l'efficacité spectrale optimale minimisant l'énergie consommée par un capteur pour une distance source-destinataire fixée. Enfin, nous aborderons aussi les notions d'ES et d'EE pour une consommation coopérative en questionnant l'influence de la position du relai.

4.2 Vers un modèle de consommation plus réaliste

Dans un grand nombre d'études, le modèle de consommation s'appuie sur la simplification de certains éléments. C'est le cas notamment du PAPR et de l'efficacité du drain dont dépend le coefficient de l'AP et qui sont considérés comme fixes même dans certains travaux récents

[33, 57, 87]. Quand ce n'est pas le cas, leur utilisation reste tout de même incomplète [77, 120, 135]. La puissance du CAN et du CNA, si elle a déjà été étudiée, reste également un élément du système peu pris en compte. Enfin, on travaille généralement à puissance de transmission infinie. Or dans un système réel, l'AP ne peut être alimenté à l'infini et connaît donc un seuil de saturation qui peut avoir un impact sur les performances de la communication.

Dans cette section, on s'attachera ainsi à détailler ces différents éléments. Dans un premier temps, on donnera une explication précise sur l'origine et le calcul de PAPR et de l'efficacité du drain, et l'impact d'un modèle dépendant du système sur la consommation totale. Dans un deuxième temps, on présentera l'état de l'art des méthodes de calcul des puissances du CNA et du CAN, ainsi que l'influence des différents paramètres sur ces puissances. Enfin, on s'attardera sur les conséquences potentielles de la saturation de la puissance de transmission.

4.2.1 Coefficient de l'amplificateur de puissance

Pour rappel, la puissance consommée par l'AP s'écrit $P_{amp} = \alpha P_t$ avec P_t la puissance de transmission et $\alpha = \frac{\xi}{\eta} - 1$ le coefficient de l'AP. ξ est le PAPR et η l'efficacité du drain. Jusqu'à présent, on a considéré ces éléments comme constants aux valeurs $\xi = 1$ et $\eta = 0.35$, ce qui correspond au paramétrage le plus fréquent dans les travaux de recherche portant sur l'efficacité énergétique des réseaux de capteurs sans fil. On détaille à présent chacun d'entre eux.

Le PAPR ξ

Le PAPR est défini comme le rapport entre la puissance instantanée maximale du signal en entrée de l'AP et sa puissance moyenne [96]. Autrement dit,

$$\xi = \frac{\max(|x(t)|^2)}{\mathbb{E}(|x(t)|^2)} \quad (4.1)$$

avec $x(t)$ le signal en entrée de l'AP.

Avant de prendre cette forme $x(t)$, le signal en entrée de l'émetteur subit des transformations telles que la modulation, des filtrages tels que le filtrage de mise en forme d'onde, une transposition en fréquence porteuse... Le PAPR dépend donc de ces différents éléments. On peut décomposer l'expression du PAPR en un produit de deux facteurs : le facteur ξ_{mod} correspondant au PAPR mesuré sur la suite de symboles générés à l'émetteur et lié au type et à l'ordre de modulation, et le facteur ξ_ρ lié au filtre de mise en forme [120] :

$$\xi = \xi_{mod} \xi_\rho. \quad (4.2)$$

Dans la majorité des modèles, le filtrage n'est pas pris en compte, seule la modulation l'est [96]. Par ailleurs, le PAPR est fixé à 1 la plupart du temps [40, 135], ce qui correspond en réalité à choisir une modulation BPSK ou QPSK [120]. Le sujet du PAPR a néanmoins été étudié pour de nombreuses modulations, en monoporteuse mais surtout en multiporteuse, modulations pour lesquelles le PAPR a une très grande importance [67]. Pour une M-QAM, on s'intéresse à la puissance des symboles pour tout ordre $M = 2^b$ avec b un entier pair non nul. Lorsque le filtrage de mise en forme d'onde n'est pas pris en compte, on a donc $\max(|x(t)|^2) = 2(\sqrt{M} - 1)^2$ et $\mathbb{E}(|x(t)|^2) = \frac{2}{3}(M - 1)$ en considérant une distribution uniforme des symboles dans le signal. Finalement, pour une M-QAM, on peut écrire [41, 120, 207] :

$$\xi_{mod} = 3 \frac{\sqrt{M} - 1}{\sqrt{M} + 1}. \quad (4.3)$$

Pour connaître la partie qui concerne le filtrage, il faut s'intéresser à des études plus théoriques. En télécommunication, l'équirépartition entre émetteur/récepteur du filtrage vérifiant le critère de Nyquist conduit à utiliser principalement un filtre en racine de cosinus surélevé, ou SRRC. Celui-ci est paramétré par le facteur de roll-off ρ , dont l'origine est abordé en annexe D.1.

Certains travaux [44] ont porté sur la caractérisation du facteur PAPR du signal présent en sortie d'un filtre SRRC. Il est établi que ce paramètre dépend de l'ordre de la M-QAM, du facteur de suréchantillonnage du CNA, du facteur de roll-off et de la longueur du filtre. Nous considérons que les coefficients du filtre SRRC sont normalisés de manière à conserver la puissance moyenne du signal, aussi nous considérons que la puissance moyenne du signal en sortie de filtre SRRC est identique à celle du signal en entrée du filtre, *i.e.* à celle de la relation 4.3. Par conséquent, nous nous intéressons uniquement au terme de la puissance instantanée maximale du signal en sortie du filtre SRRC.

Pour un taux de suréchantillonnage de 1, on a simplement $\max(|x(t)|^2) = 2a_{max}^2$ avec $a_{max} = \sqrt{M} - 1$ l'amplitude maximale par voie des symboles de la M-QAM, ce qui correspond finalement au PAPR sans filtre de mise en forme, c'est-à-dire avec $\xi_\rho = 1$. En revanche, quand ce taux est supérieur à 1, la dépendance aux autres paramètres apparaît. Ainsi, pour un facteur de roll-off de valeur $\rho \leq 0.4$, la puissance maximale du signal est :

$$\max(|x(t)|^2) = 8a_{max}^2 \left(\sum_{k=1}^n \left| \frac{A(k) + B(k)}{\frac{\pi}{2}(1 - 2k)(1 - (2(1 - 2k)\rho)^2)} \right| \right)^2 = 2a_{max}^2 \xi_\rho \quad (4.4)$$

avec $A(k) = (-1)^k \cos\left(\rho \frac{\pi}{2}(1 - 2k)\right)$ et $B(k) = 2\rho(1 - 2k)(-1)^{k+1} \sin\left(\rho \frac{\pi}{2}(1 - 2k)\right)$, et avec la longueur du filtre de mise en forme donnée par $L = 2nT_s$ (où T_s est la durée d'un symbole modulé). Pour un facteur de roll-off de valeur $\rho > 0.4$, la puissance maximale du signal est :

$$\max(|x(t)|^2) = 2a_{max}^2 \left(1 - \rho + \frac{4\rho}{\pi} + 2 \sum_{k=1}^n \left| \frac{C(k) + D(k)}{k\pi(1 - (4k\rho)^2)} \right| \right)^2 = 2a_{max}^2 \xi_\rho \quad (4.5)$$

avec $C(k) = (-1)^{k+1} \sin(\rho k\pi)$ et $D(k) = 4k\rho(-1)^k \cos(ak\pi)$.

La figure 4.1 présente une comparaison entre les valeurs théoriques du PAPR, obtenues à l'aide des relations 4.4 et 4.5, et celles obtenues expérimentalement par simulation de Monte-Carlo en fonction de ρ pour une 4-QAM et pour une 64-QAM. On peut observer sur les simulations de Monte-Carlo que, pour une M-QAM, la valeur du PAPR n'est jamais égale à 1 comme dans les modèles traditionnels, et augmente même avec l'ordre de modulation, d'où l'importance d'une modélisation dépendante du contexte. En outre, on remarque que la courbe de la borne supérieure théorique suit correctement celle de la simulation pour une 4-QAM, mais qu'une différence significative apparaît pour la 64-QAM. Malgré cette différence, qui augmente avec l'ordre de modulation, on peut considérer cette approximation comme suffisante en cela qu'elle approche beaucoup mieux la réalité du PAPR que l'approximation par $\xi = 1$.

Les expériences de [44] montrent également que le facteur de suréchantillonnage au-delà de 2 et la longueur du filtre au-delà de 8 symboles n'ont pas d'influences majeures sur la valeur du PAPR. Dans la suite de nos travaux, on considèrera donc uniquement l'impact du facteur de roll-off ρ pour comprendre comment les modules de l'émetteur agissent sur la consommation. La taille du filtre de mise en forme sera également fixée à une valeur de $L = 6T_s$ (soit $n = 3$).

A notre connaissance, un tel PAPR théorique n'a jamais été intégré dans un modèle de consommation énergétique alors qu'il influence notablement les valeurs d'énergie obtenues. Pour

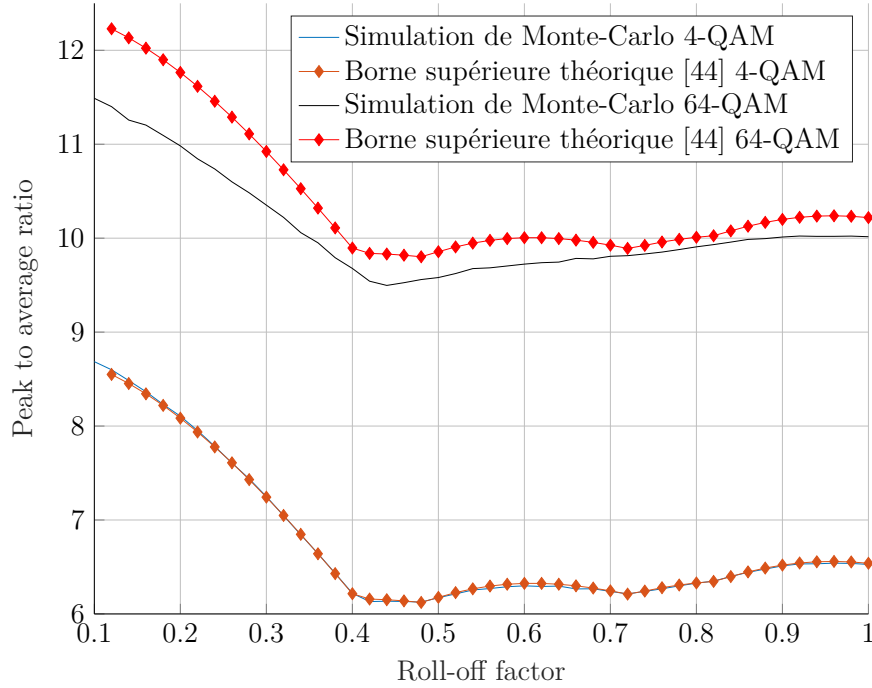
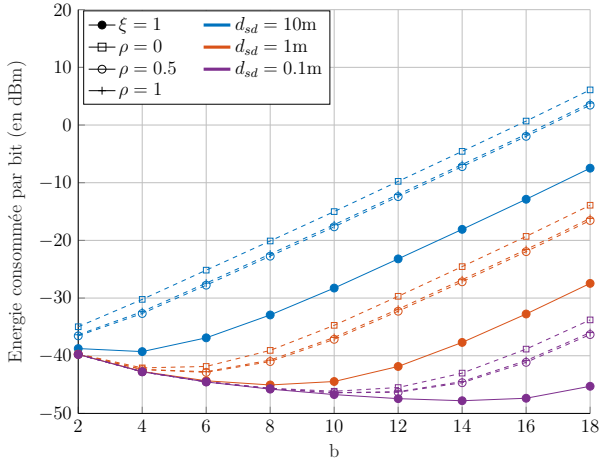


FIGURE 4.1 – PAPR en fonction du facteur de roll-off ρ pour une 4-QAM et une 64-QAM ($M = 4$, suréchantillonnage = 8, longueur du filtre $L = 6T_s$ ($n = 3$))

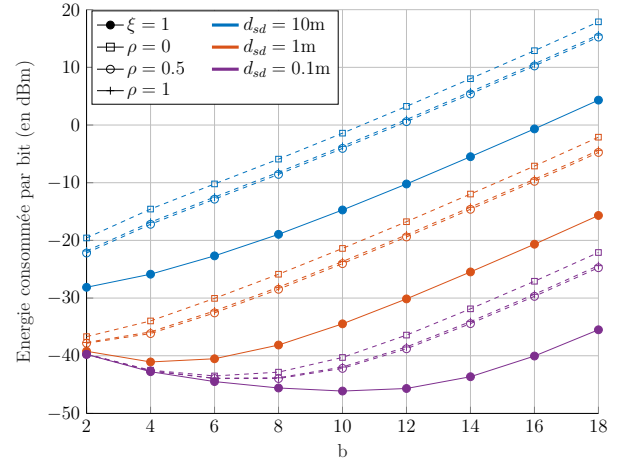
une communication point-à-point, l'impact de ce PAPR pour différentes valeurs du roll-off sur la consommation est présentée sur la figure 4.2a pour le canal AWGN et sur la figure 4.2b pour le canal de Rayleigh. Pour une communication coopérative, des résultats analogues sont présentés sur les figures 4.3a et 4.3b respectivement pour les canaux AWGN et de Rayleigh. Le facteur de suréchantillonnage est fixé à 4, l'efficacité du drain η à 0.35 et la longueur du filtre à 6 symboles. Le roll-off prend ses valeurs dans $\{0; 0.5; 1\}$.

Sur chaque figure, on observe exactement le même comportement. D'une part, on remarque que les courbes pour un PAPR dépendant du facteur de roll-off sont très proches pour un même contexte, qui produit une différence maximale d'énergie consommée d'environ 2.28 dBm. Cette différence est liée aux valeurs observées sur la figure 4.1. En effet, pour tout $\rho > 0.35$, les valeurs du PAPR sont très proches les unes des autres, et l'impact sur la consommation est assez faible. En outre, nos résultats montrent que la modification de l'ES b_d^* par le facteur de roll-off fait figure d'exception. D'autre part, les valeurs des courbes utilisant un PAPR dépendant du système sont bien supérieures à celles pour le PAPR fixé à 1, avec une différence allant jusqu'à 13.6 dBm pour $d_{sd} = 10\text{m}$ entre les courbes pour $\xi = 1$ et ξ fonction du roll-off avec $\rho = 0$. En effet, les valeurs du PAPR étant bien supérieures à 1, la puissance de transmission augmente, ainsi que l'énergie totale, ce qui influence de façon importante l'ES. Si on prend l'exemple de la communication point-à-point avec canal AWGN, on observe qu'à $d_{sd} = 10\text{m}$, l'efficacité spectrale optimale b_d^* est égale à 4 pour un PAPR constant de $\xi = 1$ alors qu'elle ne vaut que $b_d^* = 2$ pour ξ variant; de même pour $d_{sd} = 1\text{m}$, $b_d^* = 8$ à $\xi = 1$, mais diminue à $b_d^* = 6$ (respectivement 4) quand la valeur du roll-off vaut 1 ou 0.5 (respectivement 0); enfin, pour de faibles distances $d_{sd} = 0.1\text{m}$, on obtient $b_d^* = 14$ (respectivement 10) pour un PAPR de $\xi = 1$ (respectivement ξ variant en fonction du roll-off ρ).

Ces simulations nous indiquent que la modélisation du PAPR par une expression dépendante

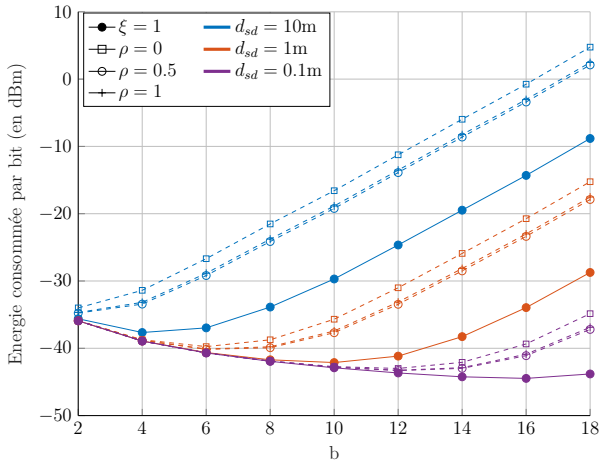


(a) Consommation pour un canal AWGN

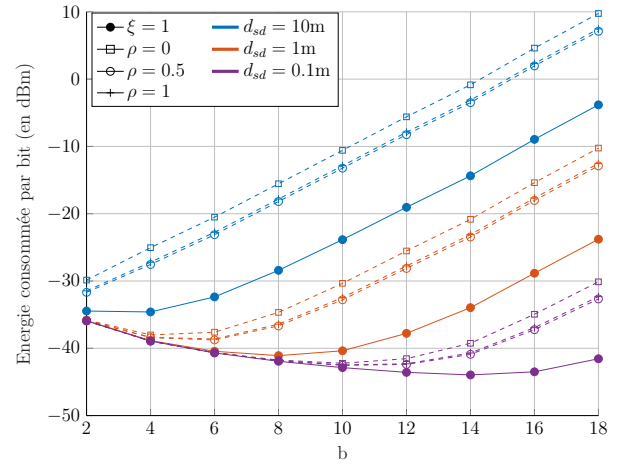


(b) Consommation pour un canal de Rayleigh

FIGURE 4.2 – Comparaison de l'énergie consommée par bit lors d'une communication point-à-point pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction de l'ES b pour un PAPR $\xi = 1$ ou dépendant du facteur de roll-off pour différentes valeurs de celui-ci



(a) Consommation pour un canal AWGN



(b) Consommation pour un canal de Rayleigh

FIGURE 4.3 – Comparaison de l'énergie consommée par bit lors d'une communication coopérative (protocole AF) pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction de l'ES b pour un PAPR $\xi = 1$ ou dépendant du facteur de roll-off pour différentes valeurs de celui-ci

des paramètres du système, et notamment le facteur de roll-off, a une grande importance dans le calcul de la consommation, ainsi que dans le compromis entre l'ES et l'EE. Etant donné que les valeurs habituelles du facteur de roll-off sont autour de 0.5, on peut cependant considérer que ce facteur a en lui-même une influence mesurée sur la consommation.

L'efficacité du drain η

L'efficacité du drain est définie comme le rapport entre la puissance RF en sortie – aussi appelée puissance de transmission – et la puissance délivrée par l'AP. Notée η , elle décrit le rendement de cet amplificateur. Sa valeur est généralement fixée à 0.35 ou 0.30 en considérant un amplificateur de classe A [120, 135]. Le détail de son calcul et de ses propriétés est donné dans l'annexe D.2.

Des modèles alternatifs de la consommation utilisent une modélisation de l'efficacité du drain en racine carrée [67, 120] ou en suivi de l'enveloppe de la puissance de l'amplificateur (*envelop tracking PA*, ou ET-PA) [77]. Nous nous intéressons ici au modèle généraliste formulé par [227] :

$$\eta = \eta_{max} \left(\frac{P_t}{P_{t,max}} \right)^\beta \quad (4.6)$$

avec η_{max} l'efficacité du drain maximale atteinte quand on opère à puissance de sortie maximale [120], P_t la puissance de transmission, $P_{t,max}$ la puissance de transmission maximale et β l'exposant de la puissance, compris entre 0 et 1. Finalement, l'efficacité du drain ne peut être constante lorsque P_t varie en raison de la relation

$$\eta \propto P_t^\beta. \quad (4.7)$$

L'exposant de la puissance β peut être réglé afin de s'adapter aux valeurs de l'efficacité du drain de dispositifs radio réels tels que les émetteurs-récepteurs CC2420 et CC1000 de Chipcon. Les courbes correspondant à l'efficacité du drain normalisée $\eta_n = \frac{\eta}{\eta_{max}}$ sont tracées sur la figure 4.4 pour une valeur optimale de $\beta \approx 0.5$ pour le CC2420 (respectivement $\beta \approx 0.44$ pour le CC1000) et comparées aux valeurs données dans les fiches techniques de ces deux appareils [35, 36]. L'équation 4.7 fournit ici un modèle de l'efficacité du drain réaliste, avec des courbes qui suivent de très près la tendance des valeurs réelles. On considèrera donc à l'avenir l'équation 4.7 dans les différents modèles de consommation utilisés.

Canal	AWGN				Rayleigh			
Modèle	η	β			η	β		
Valeur	0.35	0.25	0.45	0.65	0.35	0.25	0.45	0.65
$d_{sd} = 10\text{m}$	4	4	4	6	2	2	4	6
$d_{sd} = 1\text{m}$	8	8	8	8	4	6	6	6
$d_{sd} = 0.1\text{m}$	14	14	14	12	10	10	10	10

TABLE 4.1 – Efficacité spectrale discrète b_d^* d'une communication point-à-point pour une M-QAM ($M = 2^b$) et un canal AWGN ou de Rayleigh en fonction du modèle de l'efficacité du drain

On compare la consommation d'énergie par bit obtenue en fonction de la taille de constellation de la M-QAM pour une valeur constante de l'efficacité du drain $\eta = 0.35$ issue du modèle

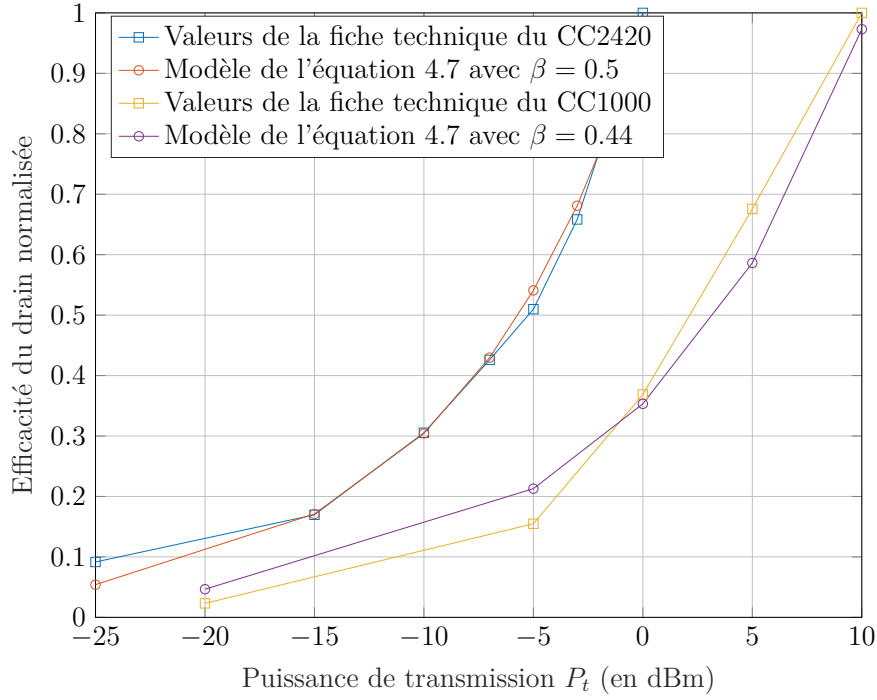


FIGURE 4.4 – Evolution de l'efficacité du drain normalisée en fonction de la puissance de transmission pour des émetteurs-récepteurs TI CC2420 et TI CC1000

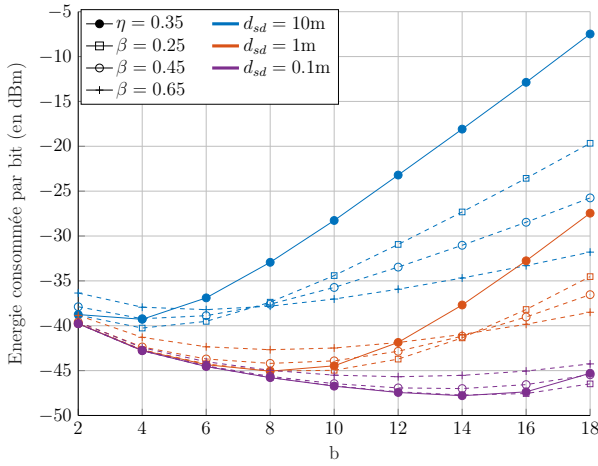
traditionnel et pour une efficacité issue de l'équation 4.7 avec $\beta \in \{0.25, 0.45, 0.65\}$. Le PAPR est à nouveau fixé à 1 de manière à ne faire varier qu'un seul paramètre à la fois. La consommation est tracée sur la figure 4.5a pour un canal AWGN et sur la figure 4.5b pour un canal de Rayleigh à évanouissement rapide. L'efficacité spectrale discrète minimisant l'énergie totale consommée est également indiquée dans la table 4.1.

On commence par observer que plus l'ordre de modulation est grand, plus il y a d'impact du modèle sur la consommation. On remarque que c'est également le cas avec la distance d_{sd} . En effet, l'efficacité du drain influence uniquement la puissance de l'AP qui dépend de la puissance de transmission, laquelle est plus importante pour les grandes distances ou ordres de modulation. Par ailleurs, on remarque que, cette efficacité croissant avec la taille de la constellation, elle va diminuer l'énergie de fonctionnement pour de fortes valeurs de l'ES b par rapport au modèle avec $\eta = 0.35$. Tout ceci contribue à modifier l'ES minimisant la consommation : pour une distance d_{sd} de 1 ou 10m, il semble pour les deux canaux que b_d^* augmente entre le modèle traditionnel et le modèle avancé, et quand l'exposant β croît ; au contraire, pour 0.1m, elle tend à diminuer dans les mêmes conditions. Cette tendance mérite en effet une étude plus poussée car, comme on le remarque sur la figure 4.5, la consommation totale d'énergie ne suit pas une évolution monotone en fonction de β . Tout ceci sera donc étudié dans les sections suivantes.

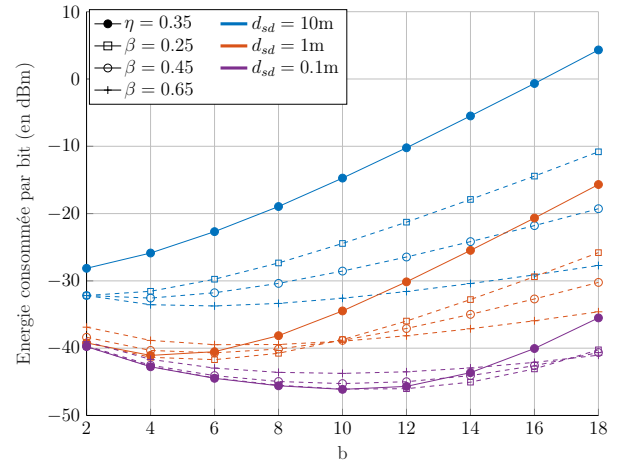
A noter également que les mêmes observations peuvent être faites pour une communication coopérative, que nous avons donc choisi ici de ne pas illustrer.

Le coefficient de l'amplificateur de puissance α

La modélisation du PAPR et de l'efficacité du drain, si elles sont connues dans le domaine de la communication sans fil, n'avaient jamais été étudiées à notre connaissance pour le calcul



(a) Consommation pour un canal AWGN



(b) Consommation pour un canal de Rayleigh

FIGURE 4.5 – Comparaison de l'énergie consommée par bit lors d'une communication point-à-point pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction de l'ES b pour une efficacité du drain $\eta = 0.35$ ou dépendante de la puissance de transmission pour différentes valeurs de l'exposant β

de la consommation d'énergie par bit transmis en fonction de la taille de constellation de la M-QAM, ni pour l'analyse du compromis entre ES et EE. En outre, leur utilisation conjointe pour modéliser le fonctionnement de l'AP reste à ce jour inédite. Nous souhaitons donc analyser la consommation d'une énergie par bit transmis impliquant ce coefficient α variant en fonction des paramètres du système, et la comparer avec celle impliquant un coefficient fixe. Ceci nous permettra de déterminer s'il existe un intérêt à utiliser ce modèle avancé de la puissance de l'AP dans le compromis entre ES et EE.

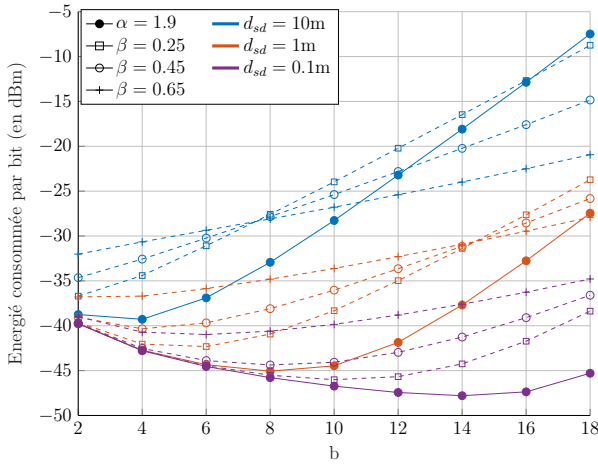
Pour une communication point-à-point, l'utilisation conjointe des relations 4.3, 4.4 et 4.5 décrivant le PAPR, et 4.7 décrivant l'efficacité de drain de l'AP nous permet d'obtenir un modèle avancé de l'énergie consommée par bit transmis, appelé dans la suite (E.proposé) :

$$E_b = 3 \frac{2^{b/2} + 1}{2^{b/2} - 1} \frac{\xi_\rho}{b} \left(\frac{N_0 N_f (4\pi d_{sd})^2 \gamma_{sd}}{G_r G_t \lambda^2} \right)^{1-\beta} + \frac{P_c}{bB} + \frac{2P_{syn} T_{tr}}{L}. \quad (4.8)$$

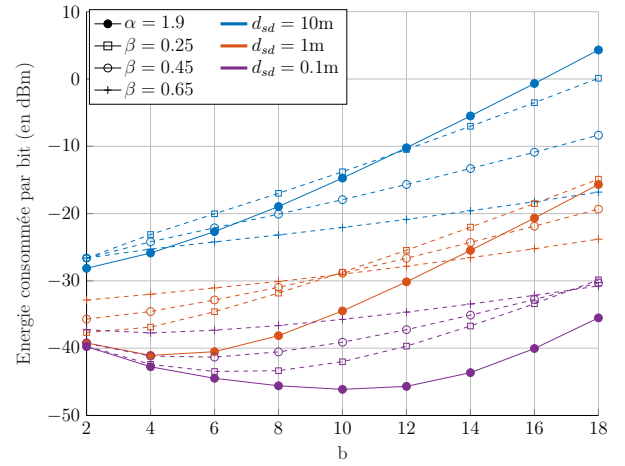
Pour une communication coopérative, on simplifiera sur le modèle de l'équation 3.81 pour un obtenir un modèle dit (E.coop.proposé) :

$$E_b = \frac{(\xi((P_{t,s})^{1-\beta} + (P_{t,r})^{1-\beta}) + P_{c,tx,s} + 2P_{c,rx,d} + P_{c,tx,s} + P_{c,rx,r})T_{on} + 5P_{syn}T_{tr}}{L}. \quad (4.9)$$

On s'intéresse dans un premier temps à la consommation d'énergie pour une communication point-à-point. On compare le modèle traditionnel (E.trad) avec un coefficient de l'AP égal à 1.9 avec le modèle avancé (E.proposé) pour un facteur de roll-off fixé à 0.5 et différentes valeurs de l'exposant β . Les courbes résultantes sont présentées sur la figure 4.6a pour le canal AWGN et sur la figure 4.6b pour le canal de Rayleigh. On remarque dès lors que pour les très courtes distances ($d_{sd} = 0.1m$ ou $d_{sd} = 1m$), l'énergie totale est généralement plus grande pour le modèle (E.proposé) que pour le modèle traditionnel. Ceci conduit, comme nous l'avons déjà remarqué, à une réduction de l'ES b_d^* discrète minimisant l'énergie consommée. En outre,



(a) Consommation pour un canal AWGN

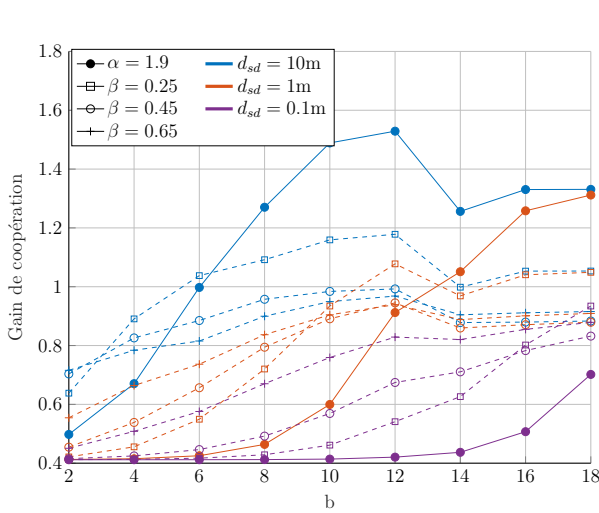


(b) Consommation pour un canal de Rayleigh

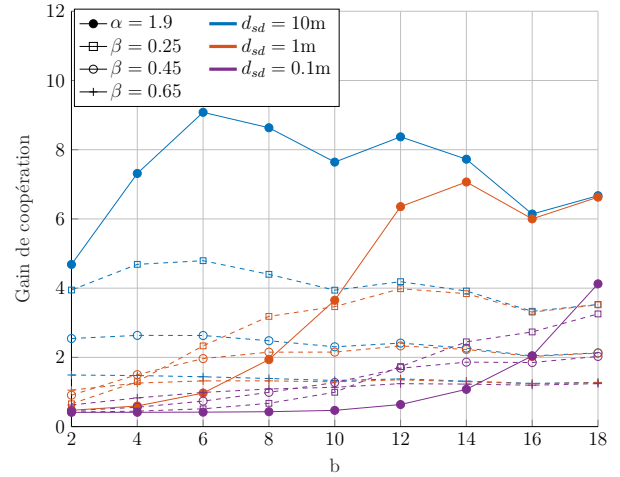
FIGURE 4.6 – Comparaison de l'énergie consommée par bit lors d'une communication point-à-point pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction de l'ES b pour un modèle du coefficient de l'AP constant ou dépendant du système pour différentes valeurs de l'exposant β

l'évolution de l'énergie en fonction de l'exposant β reste similaire à celle que nous avons détaillée précédemment.

Dans un second temps, on étudie l'influence du modèle avancé (E.proposé) pour une communication coopérative. Les courbes de consommation nous donnant les mêmes conclusions que celles de la communication point-à-point, on analyse le gain de coopération. La figure 4.7a nous présente le gain pour un canal AWGN. On y observe que, pour un même contexte (c'est-à-dire le même modèle de consommation ou le même exposant β), le gain augmente avec la distance. En revanche, il n'est pas possible d'exprimer une tendance fixe entre les différents contextes ; par exemple, les valeurs de gain pour un coefficient α constant sont plus faibles à $d_{sd} = 0.1\text{m}$ que pour un coefficient variant avec le système, et celles-ci croissent avec l'exposant β , mais c'est le contraire qui s'opère à $d_{sd} = 10\text{m}$ pour la plupart des tailles de constellation. Il est cependant intéressant de remarquer que ces variations peuvent modifier l'utilité du relai dans le système. En effet, pour $d_{sd} = 1\text{m}$, le gain pour un coefficient α constant dépasse l'unité à partir de $b = 14$, c'est-à-dire que le relai devient énergétiquement avantageux à partir de cette taille de constellation. En revanche, pour le modèle (E.proposé), le gain ne dépasse l'unité que pour $\beta = 0.25$ quand $b = 12$ ou $b > 16$, mais pas avec les autres valeurs de l'exposant. Le choix d'une communication avec ou sans relai peut donc en être affecté. La communication sur canal de Rayleigh, dont le gain de coopération est présenté sur la figure 4.7b, est également impactée par ce choix. Le problème se pose cependant différemment puisque le gain de coopération est quasiment toujours supérieur à 1. En particulier, pour le modèle (E.proposé), le gain dépasse la valeur unitaire pour des tailles de constellation bien plus faibles que pour le modèle (E.trad). On note par ailleurs que, pour chaque valeur de l'exposant β , les courbes tendent vers une même valeur du gain, et ce quelle que soit la distance d_{sd} . Ceci rejoint la tendance que semble suivre les courbes pour $\alpha = 1.9$ (on peut observer que les courbes pour $d_{sd} = 10\text{m}$ et $d_{sd} = 1\text{m}$ se rejoignent à partir de $b = 16$), mais avec une convergence plus rapide. Cela signifie que, à partir d'un certain ordre de modulation, celui-ci joue un rôle plus important dans la différence de consommation entre communication point-à-point et coopérative que la distance



(a) Gain de coopération pour un canal AWGN



(b) Gain de coopération pour un canal de Rayleigh

 FIGURE 4.7 – Gain de coopération sur un canal AWGN ou de Rayleigh en fonction de l'ES b pour un modèle du coefficient de l'AP constant ou dépendant du système pour différentes valeurs de l'exposant β

source-destinataire.

Dans la suite de nos travaux, sauf mention du contraire, nous utiliserons le modèle avancé (E.proposé) en fixant $\beta = 0.5$ et $\rho = 0.5$.

4.2.2 Puissance des convertisseurs

Jusqu'à présent, la puissance du circuit P_c a été considérée comme une valeur constante, dépendant des modules du circuit mais d'aucun des paramètres introduits. On a vu cependant que ces paramètres peuvent avoir un impact sur la puissance de l'amplificateur. De la même manière, en réalité, P_c varie en fonction de la largeur de la bande passante [41, 89] et du taux de suréchantillonnage [111] au travers du CAN et du CNA.

Puissance du CNA

On considère un CNA à pondération binaire et à orientation en courant, dont le schéma interne simplifié est présenté sur la figure 4.8. On y observe R_L et C_L respectivement la résistance et la conductance du module, I_0 la source de courant unitaire correspondant au bit le moins significatif (LSB), et V_{dd} la tension d'alimentation. Les n_1 branches du convertisseurs intègrent chacune une source de courant alimentée ou non grâce à un commutateur, lui-même actionné par l'élément b_i correspondant issu du mot binaire – ou symbole – $(b_0, b_1, \dots, b_{n_1-1})$.

La puissance moyenne consommée par le CNA est composée de deux parties : la puissance statique P_s et la puissance dynamique P_d . La puissance statique est la puissance consommée par ces sources de courant et donnée par [41]

$$P_s = \mathbb{E} \left[V_{dd} I_0 \sum_{i=0}^{n_1-1} 2^i b_i \right] = \frac{1}{2} V_{dd} I_0 (2^{n_1} - 1) \quad (4.10)$$

en supposant la probabilité de chaque élément binaire $p(b_i = 1) = p(b_i = 0) = \frac{1}{2}$.

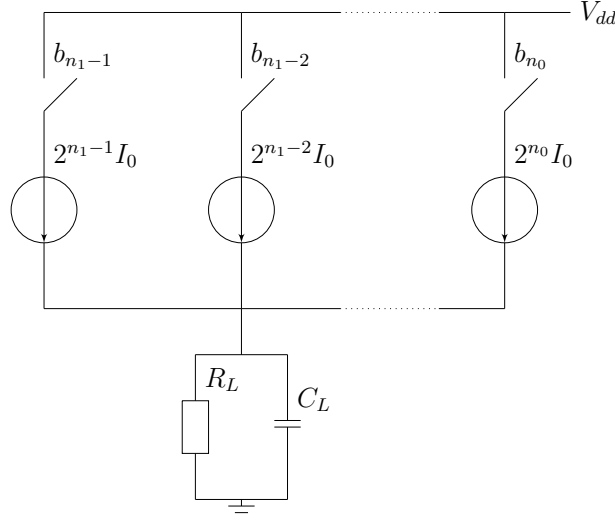


FIGURE 4.8 – Schéma interne simplifié du convertisseur numérique-analogique (CNA)

La puissance dynamique P_d correspond à la puissance consommée par le commutateur quand il est connecté ou déconnecté par un changement de valeur de l'élément binaire b_i associé. Elle est estimée à

$$P_d = \frac{1}{2} n_1 C_p f_{switch} V_{dd}^2 \quad (4.11)$$

avec C_p la capacité parasite de chaque commutateur, $\frac{1}{2}$ le facteur de commutation, c'est-à-dire la probabilité pour le commutateur de changer d'état à chaque nouveau symbole, et f_{switch} la fréquence de commutation, généralement fixée à la fréquence d'échantillonnage f_s . Cette dernière équivaut par ailleurs à $f_s = 2(B + f_{corr})$ avec B la largeur de la bande passante et f_{corr} la fréquence d'angle du bruit en $1/f$ [41]. On voit alors apparaître au niveau de la puissance du CNA une dépendance à la largeur de la bande passante.

En outre, si on considère le facteur de suréchantillonnage OSR, le lien entre fréquence d'échantillonnage et bande passante devient $f_s = \text{OSR} \times B$ [111]. Finalement, la puissance consommée par le CNA s'écrit

$$P_{DAC} = \frac{1}{2} V_{dd} I_0 (2^{n_1} - 1) + \frac{1}{2} n_1 C_p \text{OSR} B V_{dd}^2 \quad (4.12)$$

Puissance du CAN

[103] décrit le fonctionnement du CAN rapide à fréquence de Nyquist. Ce module est constitué d'un circuit de pré-traitement qui fonctionne à la fréquence f_i du signal en entrée, et de comparateurs qui décodent ce signal à la fréquence d'échantillonnage f_s du système numérique. La puissance du CAN se résume à

$$P_{ADC} = V_{dd}^2 (B + f_i) C \quad (4.13)$$

avec C la capacitance des comparateurs. Cette capacitance est proportionnelle à la longueur minimale du canal L_{min} pour la technologie CMOS donnée. Une régression linéaire de la puissance du CAN pour différentes valeurs de bits de résolution n_2 , c'est-à-dire la longueur du

symbole en sortie du CAN, permet finalement de déterminer sa puissance sous la forme [103]

$$P_{ADC} = \frac{3V_{dd}^2 L_{min} (B + f_i)}{10^{-0.1525n_2 + 4.838}}. \quad (4.14)$$

Cette fois-ci, on considère $f_s \approx 2(B + f_i)$. En prenant en compte le facteur de suréchantillonnage, on a finalement la puissance

$$P_{ADC} = \frac{3V_{dd}^2 L_{min} \cdot OSR \cdot B}{2 \cdot 10^{-0.1525n_2 + 4.838}}. \quad (4.15)$$

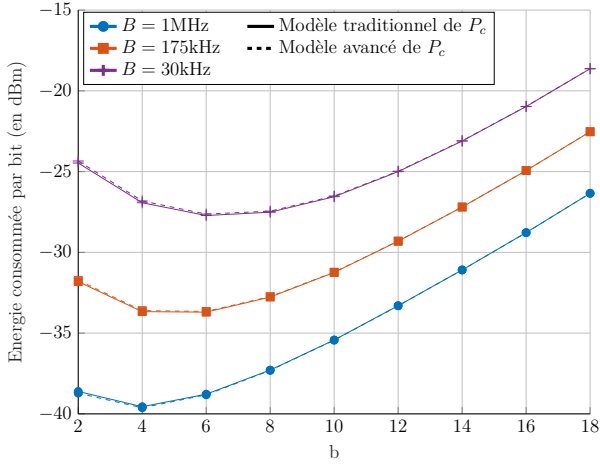
Analyse du modèle avancé de P_c

Finalement, on définit comme le modèle avancé de la puissance P_c la relation 3.74 pour laquelle P_{DAC} et P_{ADC} suivent les modèles exprimés respectivement par les équations 4.12 et 4.15, et non pas par une valeur constante. Dans les prochaines simulations conduites dans nos travaux, sauf mention du contraire, on utilisera ce modèle avancé de P_c avec les paramètres indiqués dans la table 4.2.

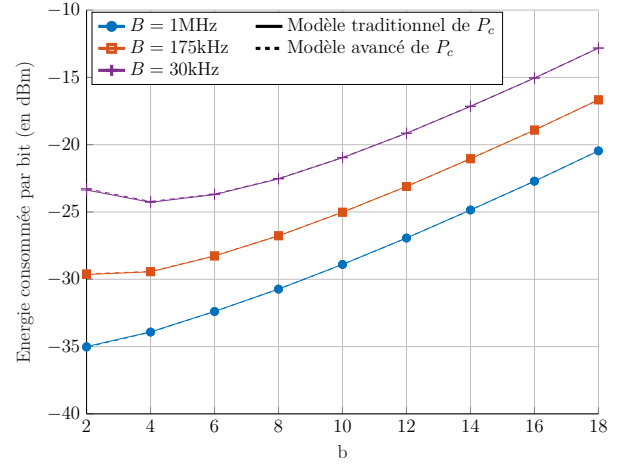
Paramètre	Description	Valeur
V_{dd}	Tension d'alimentation	3 V
I_0	Courant du LSB	10 μ A
C_p	Capacité parasite des commutateurs	1 pF
L_{min}	Longueur minimale du canal	0.5 μ A
B	Largeur de la bande passante	1 MHz
OSR	Facteur de suréchantillonnage	4
n_1	Résolution binaire du CNA	10 bits
n_2	Résolution binaire du CAN	10 bits

TABLE 4.2 – Paramètres du CAN et du CNA

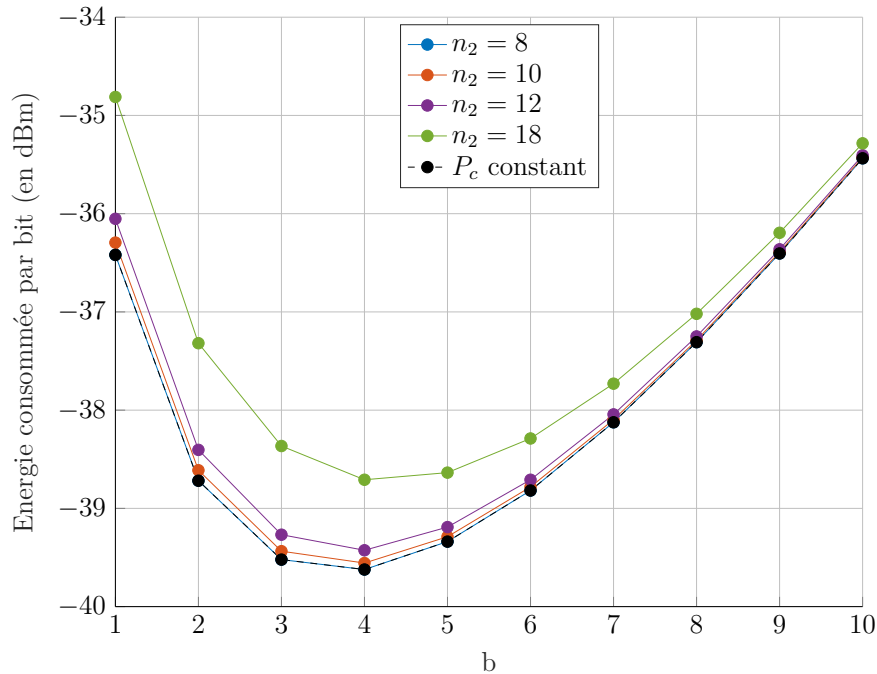
Jusqu'à présent, nous avons travaillé avec une largeur de bande passante fixée à $B = 1$ MHz. Dans un premier temps, on fait varier cette largeur pour connaître son influence sur l'énergie totale et valider notre modèle. On fixe la distance source-destinataire à $d_{sd} = 1$ m et on observe l'énergie consommée pour une communication point-à-point en fonction de la taille de constellation sur la figure 4.9. On utilise ici le modèle de consommation avec le coefficient de l'amplificateur de puissance α dépendant du système. L'énergie totale est tracée pour des valeurs de largeur de bande usuelles : 1 MHz, qu'utilise notamment le transmetteur-receveur CC2420 [35], et 175 kHz et 30 kHz, que l'on peut par exemple retrouver pour le CC1000 [36]. Que ce soit pour le canal AWGN sur la figure 4.9a ou pour le canal de Rayleigh sur la figure 4.9b, on observe très peu de différence entre les courbes du modèle traditionnel et celles du modèle avancé. En effet, on remarque un léger décalage pour les faibles tailles de constellations, c'est-à-dire quand la puissance de transmission est négligeable devant la puissance du circuit. Cependant, même au sein de la puissance du circuit, les variations de la bande passante semblent avoir peu d'importance. On peut néanmoins noter que, pour $B = 1$ MHz et $B = 175$ kHz, les valeurs de l'énergie pour le modèle avancé sont légèrement supérieures à celles du modèle traditionnel ; pour $B = 30$ kHz, c'est l'inverse qui se produit. Pour une communication coopérative,



(a) Consommation pour un canal AWGN



(b) Consommation pour un canal de Rayleigh

 FIGURE 4.9 – Comparaison de l'énergie consommée par bit lors d'une communication point-à-point pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction de l'ES b pour différents modèles de P_c et différentes valeurs de la largeur de bande passante

 FIGURE 4.10 – Comparaison de l'énergie consommée par bit lors d'une communication point-à-point pour une M-QAM ($M = 2^b$) sur un canal AWGN en fonction de l'ES b pour différents modèles de P_c et différentes valeurs de la résolution n_2

les puissances du CAN et du CNA ont une influence encore moins grande sur l'énergie totale, l'impact est donc moindre.

Une deuxième simulation, non présentée ici, nous indique que la valeur de n_1 a également peu d'influence sur l'énergie consommée. En revanche, comme on peut l'apercevoir sur la figure 4.10, pour laquelle nous avons fixé la distance d_{sd} à 1m, des valeurs importantes de la résolution du CAN n_2 ont un impact significatif sur l'énergie, en particulier pour les faibles tailles de constellation comme expliqué précédemment. En effet, la puissance du CAN augmente exponentiellement avec n_2 , et non pas linéairement comme c'était le cas pour B ou n_1 , ce qui influe plus fortement sur l'énergie totale.

Pour les valeurs observées de n_2 , l'ES minimisant l'énergie consommée n'est pas modifiée. En revanche, puisque seule P_c augmente avec ce paramètre, il est possible qu'elle le soit avec des valeurs plus grandes de la résolution du CAN. Une étude plus détaillée sera conduite à ce sujet dans la suite de nos travaux.

4.2.3 Limitation de la puissance de transmission

Un système réaliste ne peut pas travailler à puissance instantanée illimitée : il existe P_{max} tel que, à tout instant, $\max(P_{on}, P_{sp}, P_{tr}) \leq P_{max}$. Comme $P_{on} \gg P_{sp}$ et $P_{on} \gg P_{tr}$, on s'intéresse uniquement à la valeur de P_{on} [40]. En outre, la puissance de l'amplificateur de puissance peut elle-même être saturée, ce qui restreint également la puissance d'émission. Comme la puissance du circuit P_c est fixée par les différents paramètres du système, c'est souvent cette puissance d'émission qui sera considérée au lieu de la puissance de fonctionnement [157].

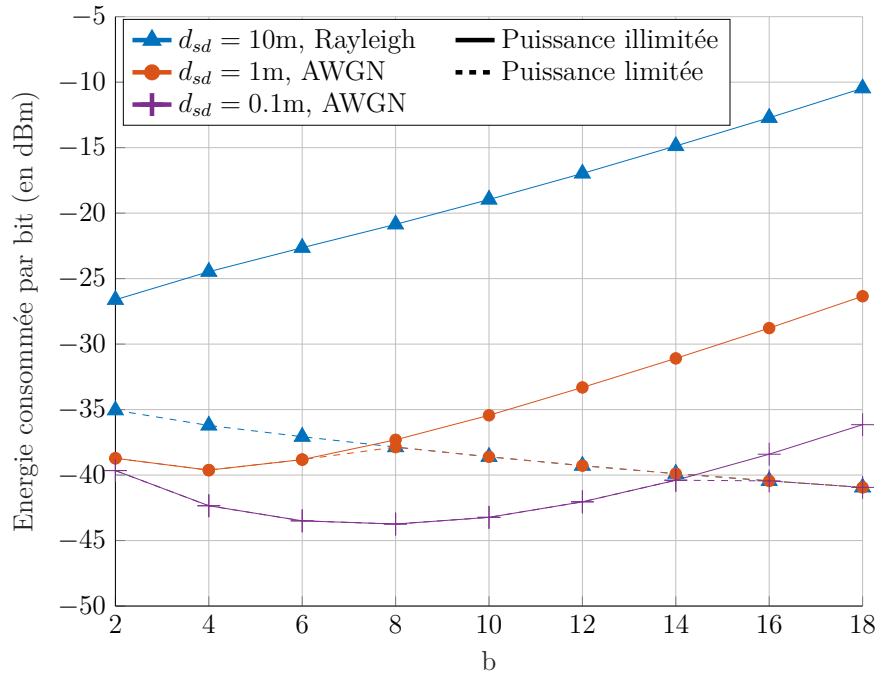


FIGURE 4.11 – Comparaison de l'énergie consommée par bit lors d'une communication point-à-point pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction de l'ES b et pour une puissance de transmission P_t limitée ou non

On prend exemple du niveau maximal en entrée de l'amplificateur de puissance du CC2420, qui est fixé à $P_{t,max} = 10$ dBm [35]. On compare l'énergie consommée par bit pour un modèle

théorique avec énergie limitée et illimitée sur la figure 4.11. On observe que les courbes avec puissance limitée (en pointillés) sont confondues avec celles permettant une puissance illimitée (en ligne continue) jusqu'à une certaine valeur b . Pour $d_{sd} = 1\text{m}$ et $d_{sd} = 0.1\text{m}$, ces valeurs sont respectivement de $b = 8$ et $b = 14$. Au-delà de cette valeur de b , la puissance P_t est saturée à $P_{t,max}$; l'énergie de transmission prend alors une valeur constante, qui ne dépend plus ni de b ni de la distance d_{sd} , ce qui conduit les courbes en pointillés à être confondues (à noter que ces courbes décroissent alors avec b du fait de l'énergie du circuit). Pour des distances plus longues, la puissance de transmission est toujours plus grande que $P_{t,max}$, ce qui résulte en une consommation d'énergie qui suit une courbe maximale.

Avec une puissance de transmission limitée, le RSB en réception de la communication est plus faible qu'attendu, ce qui conduit la probabilité d'erreur cible à ne pouvoir être atteinte. En parallèle, les grandes tailles de constellation conduisent à un temps de fonctionnement court tandis que la puissance de transmission reste à la puissance maximale, ce qui provoque une décroissance de l'énergie totale. Afin de préserver la qualité de service visée, il n'est donc pas toujours pertinent de choisir l'ordre de modulation qui offre une consommation d'énergie moindre.

4.3 Une meilleure estimation de l'efficacité spectrale b^*

Dans le chapitre précédent, nous avons présenté des formules analytiques du compromis entre ES et EE issues de l'état de l'art. Ces formules, déduites de la capacité, permettent de déterminer l'ES minimisant la consommation d'énergie par bit transmis calculée à partir d'un modèle simplifié de cette énergie. Nous souhaiterions obtenir des expressions analytiques équivalentes en considérant maintenant le modèle avancé de consommation d'énergie qui a été présenté dans la section précédente. Bien entendu, ce modèle faisant intervenir des paramètres dépendants de b , en particulier le coefficient de l'AP qui dépend lui-même de l'ES, la résolution du problème de minimisation de l'énergie consommé se révèle plus délicate. En outre, pour des raisons d'homogénéité et pour bien comprendre l'impact de la transmission sur la classification d'événements sonores, nous souhaitons travailler à partir de la PEB ; ceci n'est actuellement pas possible pour le canal de Rayleigh, pour lequel nous ne disposons pas de formule de la pénalité en RSB.

Dans un premier temps, on s'intéressera à déterminer une forme de la pénalité en RSB la plus proche des simulations, tant pour le canal de Rayleigh que pour le canal AWGN dans le cadre d'une communication point-à-point. Dans un deuxième temps, on tentera de formuler de nouvelles expressions analytiques de l'ES minimisant la consommation calculée à l'aide du modèle avancé de l'énergie. Enfin, on présentera les variations de l'ES optimisant l'énergie en fonction des différents paramètres du système.

4.3.1 Capacité du canal

Cette première section est l'occasion d'étudier l'estimation de la capacité ergodique des canaux AWGN et de Rayleigh pour une M-QAM, et d'y apporter notre contribution.

Capacité pour un canal AWGN

Dans le chapitre 3, nous avons utilisé une PEB dépendant uniquement de la pénalité du RSB Γ_{AWGN} , c'est-à-dire l'équation 3.111 issue de [89], que nous rappelons ici :

$$P_{eb} = \frac{1}{5} \exp\left(-\frac{3}{2}\Gamma_{AWGN}\right). \quad (4.16)$$

Cette équation représente une bonne approximation de la borne supérieure de la PEB en fonction de la pénalité du RSB. Inspirés par ce modèle, nous proposons de déterminer une meilleure approximation des courbes sous la forme

$$P_{eb} = \alpha_1 \exp(-\alpha_2 \Gamma_{AWGN}) \quad (4.17)$$

telle que α_1 et α_2 soient les coefficients minimisant l'erreur quadratique moyenne entre ce modèle et la moyenne des courbes pour les alphabets M-QAM considérés. Une approximation des moindres carrés par l'algorithme de Levenberg-Marquardt nous donne finalement les paramètres $\alpha_1 = 0.1231$ et $\alpha_2 = 2.62208$. La figure 4.12 nous présente à nouveau le tracé de la PEB pour différents ordres de modulation de la M-QAM, ainsi que la borne supérieure de l'équation 4.16. Additionnellement, on présente en rouge la courbe de l'équation 4.17. Celle-ci atteste que l'approximation donnée par l'équation 4.17 est très proche des courbes réelles de la M-QAM.

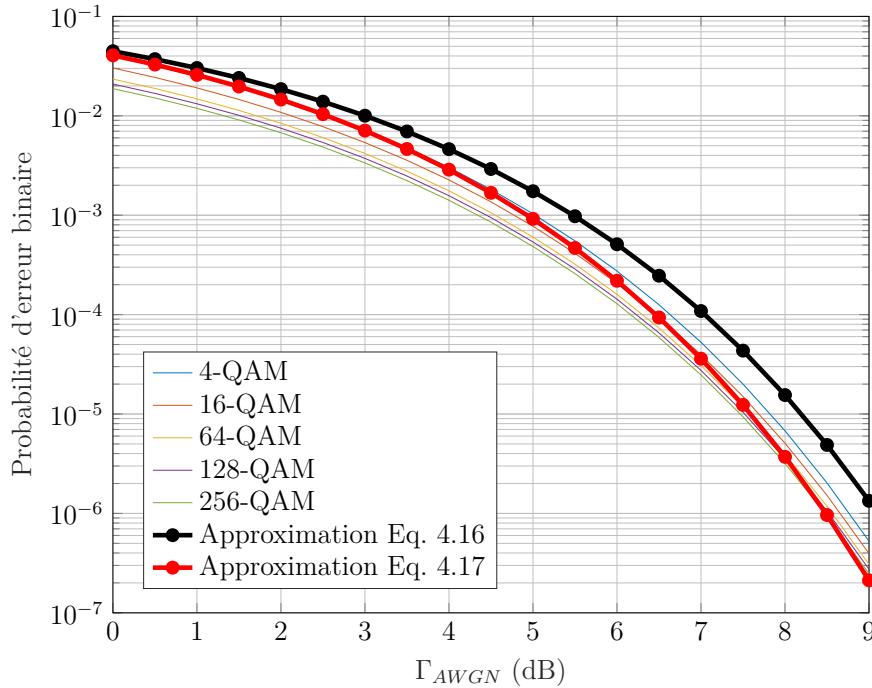


FIGURE 4.12 – Probabilité d'erreur binaire d'une M-QAM sur un canal AWGN en fonction de la pénalité Γ_{AWGN} du RSB par symbole pour différents ordres de modulation et différentes approximations

Capacité pour un canal de Rayleigh à évanouissement plat

Dans le chapitre 3, un encadrement et une approximation de la borne supérieure de l'ES 3.112 ont été données aux équations 3.113 et 3.114. Cette dernière, qui s'exprime sous la forme

[110]

$$\theta \approx \log_2 \left(1 + \frac{\gamma_{sd}}{2} \right) \quad (4.18)$$

est, pour rappel, une approximation satisfaisante pour les hautes valeurs de RSB, mais s'avère moins bonne que les bornes supérieure et inférieure de l'équation 3.113 pour les faibles RSB. On se propose d'améliorer cette approximation en s'inspirant des équations précédentes pour poser un modèle paramétrique dépendant de deux réels positifs α_1 et α_2 :

$$\theta = \alpha_1 \log_2 (1 + \alpha_2 \gamma_{sd}). \quad (4.19)$$

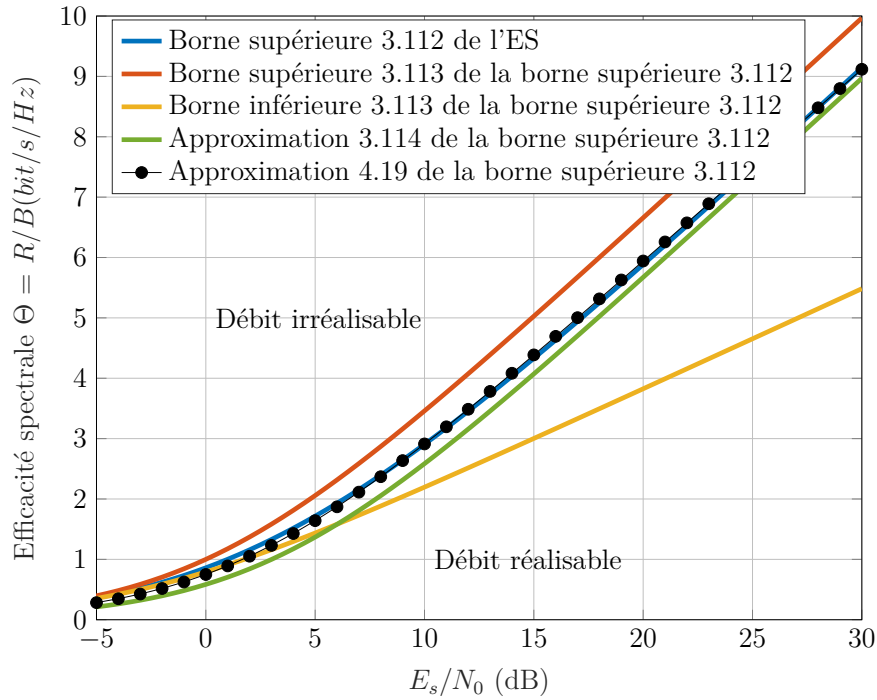


FIGURE 4.13 – Efficacité spectrale en fonction du RSB symbole pour différentes approximations

De manière analogue à la PEB pour le canal AWGN, on utilise une approximation des moindres carrés obtenue par l'algorithme de Levenberg-Marquardt pour déterminer les valeurs optimales de $\alpha_1 = 0.961$ et $\alpha_2 = 0.7168$. La courbe résultante est ajoutée à celles des précédentes approximations sur la figure 4.13. On remarque que cette courbe suit avec précision celle de la borne supérieure 3.112 de l'ES, ce qui nous indique qu'elle en représente une très bonne approximation. Finalement, le RSB symbole pour un signal gaussien transmis sur un canal de Rayleigh est donné par

$$\gamma_{sd} = \frac{2^{\theta/\alpha_1} - 1}{\alpha_2}. \quad (4.20)$$

On considère à présent le cas d'un signal modulé par M-QAM. Par rapport au cas d'un signal gaussien de l'équation 4.20, le RSB doit être corrigé par une pénalité $\Gamma_{Rayleigh}$ dépendant de la probabilité d'erreur. Le RSB s'écrit donc

$$\gamma_{sd} = \frac{2^{\theta/\alpha_1} - 1}{\alpha_2} \Gamma_{Rayleigh}. \quad (4.21)$$

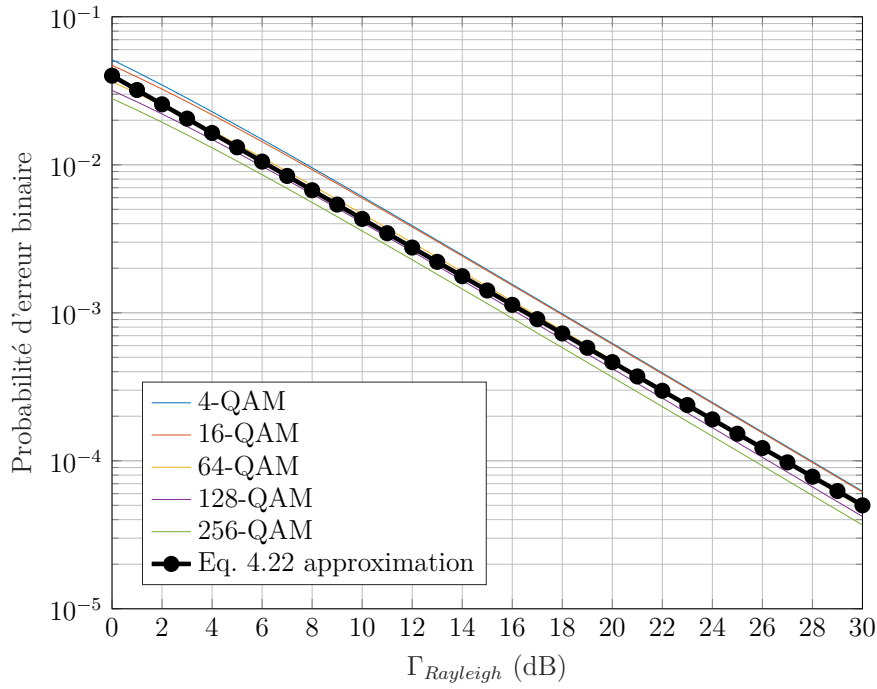


FIGURE 4.14 – Probabilité d'erreur binaire d'une M-QAM sur un canal de Rayleigh en fonction de la pénalité $\Gamma_{Rayleigh}$ du RSB par symbole pour différents ordres de modulation et pour l'approximation 4.22

Dans le chapitre 3, nous avons expliqué qu'il existe une forme de cette pénalité $\Gamma_{Rayleigh}$ dépendant de la PES, mais aucune dépendant de la PEB dans l'état de l'art. Pour des raisons de cohérence, nous allons donc chercher une nouvelle forme de $\Gamma_{Rayleigh}$ dépendant de la PEB. Pour rappel, la recherche d'une telle pénalité du RSB pour une source modulée par M-QAM par rapport à une source gaussienne est d'une grande importance, car elle nous permet d'obtenir une expression de l'énergie consommée par bit transmis plus simple à manipuler que lorsque le RSB est déterminé à partir de la réciproque de la formule de la PEB. Il devient alors possible d'exprimer l'ES qui minimise cette énergie.

Sur la figure 4.14, on trace alors les courbes de la PEB en fonction de la pénalité du RSB en décibels pour différents ordres de modulation, conformément à l'équation 3.47 qui exprime la PEB obtenue sur un canal de Rayleigh en utilisant une modulation M-QAM. L'ensemble de ces courbes, dont l'ordonnée est tracée en échelle logarithmique, restent proches les unes des autres et semblent avoir un comportement linéaire. On se propose alors de réaliser une régression linéaire sur le logarithme de la moyenne de ces courbes afin d'obtenir une approximation de la PEB sous la forme

$$\ln(\text{Peb}) = \beta_1 10 \log(\Gamma_{Rayleigh}) + \beta_2 \quad (4.22)$$

avec β_1 et β_2 deux réels. Le calcul de ces coefficients nous donne alors $\beta_1 = -0.2228$ et $\beta_2 = -3.2189$. La courbe correspondante, tracée sur la figure 4.14, suit de très près les courbes théoriques. Le modèle étant satisfaisant, on utilisera la pénalité $\Gamma_{Rayleigh}(\text{Peb})$ suivante dans la suite de nos travaux :

$$\Gamma_{Rayleigh}(\text{Peb}) = 10^{\frac{\ln(\text{Peb}) - \beta_2}{10\beta_1}}. \quad (4.23)$$

En insérant l'équation 4.23 dans la relation 4.21, puis en remplaçant l'expression ainsi obtenue du RSB γ_{sd} dans le modèle de consommation donné par 3.120 au chapitre 3, nous obtenons le nouveau modèle suivant, nommé (R.qam.trad.peb) :

$$E_b = (1 + \alpha)G_d N_0 N_f \Gamma_{Rayleigh}(\text{Peb}) \frac{2^{b/\alpha_1} - 1}{\alpha_2 b} + \frac{P_c}{bB} + \frac{P_{tr} T_{tr}}{L}. \quad (4.24)$$

Le fait d'avoir exprimé la PEB d'une source M-QAM sur un canal de Rayleigh non pas directement à partir de l'expression 3.47, mais à l'aide de la pénalité $\Gamma_{Rayleigh}$ en RSB par symbole (relativement au cas d'une source gaussienne sur un canal de Rayleigh), nous permet de faire apparaître explicitement dans l'équation 4.24 toutes les dépendances en fonction du paramètre b . Ainsi, il est maintenant plus aisé de différencier cette expression pour déterminer la valeur de l'efficacité spectrale optimale b^* minimisant l'énergie consommée.

En suivant le même raisonnement que celui de la section 3.6.3, nous obtenons l'expression analytique suivante, qui sera nommé (Bopt.R.trad.proposé) :

$$b^* = \frac{\alpha_1}{\ln(2)} \left(1 + W \left(\frac{c - a'}{a' e} \right) \right) \quad (4.25)$$

avec $a' = (1 + \alpha)N_0 N_f \Gamma_{Rayleigh}(\text{Peb}) \frac{(4\pi d_{sd})^2}{\alpha_2 G \lambda^2}$ et $c = \frac{P_c}{B}$.

Comparaison des différents modèles de la capacité

On souhaite à présent comparer les courbes de l'ES b^* minimisant l'énergie consommée par bit transmis conformément aux équations 3.119 et 4.24.

On reprend la figure 3.25 qui présente les courbes de b_d^* en fonction de la distance source-destinataire d_{sd} , et on y ajoute celles des expressions analytiques étudiées dans cette section sur la figure 4.15. Pour le canal AWGN, on compare les courbes issues de l'expression analytique 3.123 de la solution (Bopt.A.trad) pour une pénalité $\Gamma_{AWGN}(\text{Peb})$ 3.109 issue de l'état de l'art ou 4.17 proposée avec celles obtenues au moyen d'un algorithme de recherche. Les deux courbes issues de l'expression analytique sont très proches l'une de l'autre ; on observe cependant une petite amélioration à l'aide de notre modèle de la pénalité. Pour le canal de Rayleigh, on calcule à présent b^* à partir de la PEB, on ne peut donc pas utiliser les résultats du chapitre 3. On compare alors cette fois-ci les valeurs exactes de b_d^* avec celles obtenues à l'aide de l'expression analytique 4.25 de la solution (Bopt.R.trad.proposé) pour une pénalité $\Gamma_{AWGN}(\text{Peb})$ 4.23 proposées dans ces travaux. Les résultats, contrairement à ceux de la figure 3.25, sont extrêmement satisfaisants. En effet, les deux courbes sont presque parfaitement confondues, ce qui confirme la bonne modélisation de notre système.

Nous disposons finalement de modèles approximant efficacement le RSB sur la liaison source-destinataire pour un canal AWGN ou de Rayleigh traversé par un signal modulé par M-QAM.

4.3.2 Efficacité spectrale pour un modèle de consommation plus réaliste

Dans le modèle traditionnel (E.trad) de la consommation, la dépendance à l'ES provient de la présence à la fois du RSB et du débit. Néanmoins, nous avons vu dans la section 4.2.1 que le coefficient α de l'amplificateur de puissance dépend lui aussi de l'ES, tant par le PAPR que

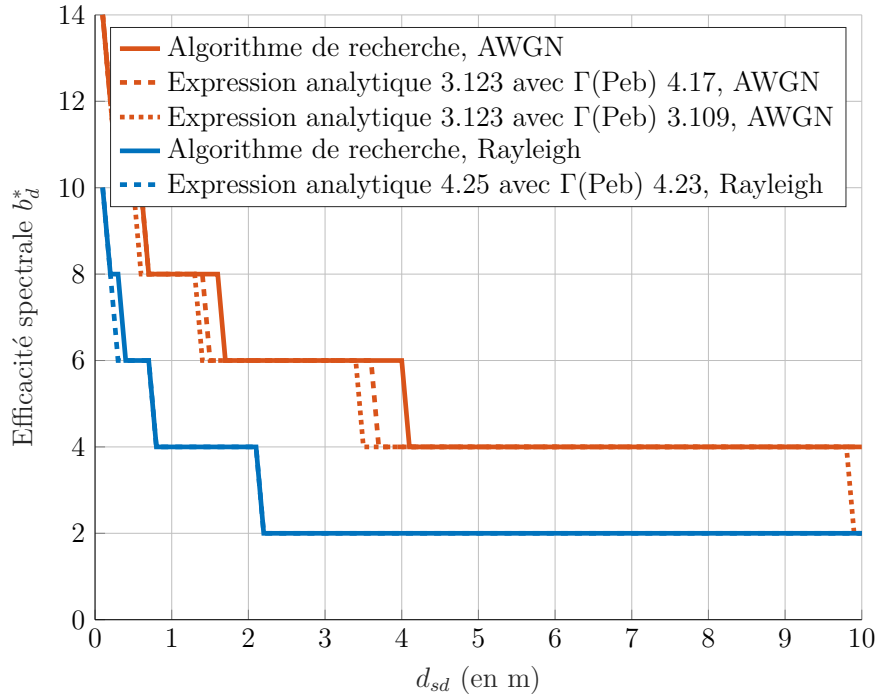


FIGURE 4.15 – Évolution de l'efficacité spectrale (ES) b^*_d minimisant la consommation d'énergie par bit pour une M-QAM ($M = 2^b$) et un canal AWGN ou de Rayleigh en fonction de la d_{sd}

par l'efficacité du drain. Le calcul de l'ES b^* minimisant l'énergie va donc se complexifier au même titre que celle-ci.

En combinant les équations 4.8 de l'énergie par bit transmis pour le modèle de consommation avancé (E.proposé), 3.108 du RSB en fonction de l'ES et 4.17 de la PEB en fonction de la pénalité du RSB, il est possible de réécrire la formule de l'énergie consommée par élément binaire en fonction de l'ES pour le modèle de consommation avancé et le canal AWGN parcouru par une source modulée par M-QAM :

$$E_{b,AWGN}(b) = A \frac{2^{b/2} - 1}{2^{b/2} + 1} \frac{(2^b - 1)^{1-\beta}}{b} + \frac{C}{b} + D \quad (4.26)$$

avec $A = 3\xi_\rho \frac{(BG_d N_0 N_f \Gamma_{AWGN}(Peb))^{1-\beta}}{B}$, $C = \frac{P_c}{B}$ et $D = \frac{P_{tr} T_{tr}}{L}$. De la même manière, on peut combiner les équations 4.8, 4.21 et 4.23 pour être en mesure d'exprimer l'énergie consommée par élément binaire pour une source modulée par M-QAM sur un canal de Rayleigh :

$$E_{b,Rayleigh}(b) = A' \frac{2^{b/2} - 1}{2^{b/2} + 1} \frac{(2^{b/\alpha_1} - 1)^{1-\beta}}{b} + \frac{C}{b} + D \quad (4.27)$$

en posant $A' = 3\xi_\rho \frac{(BG_d N_0 N_f \Gamma_{Rayleigh}(Peb))^{1-\beta}}{\alpha_2 B}$.

Ces deux modèles seront par la suite notés respectivement (A.qam.proposé) et (R.qam.proposé). L'énergie calculée à partir du modèle avancé (E.proposé), comme nous avons pu l'observer dans la section 4.2.1, semble admettre un unique minimum global par rapport à b , au même titre que celle calculée à l'aide du modèle traditionnel (E.trad). Pour s'assurer de ce résultat, on démontre la convexité de l'équation 4.26 en annexe C.3. Comme α_1 et α_2 sont strictement positifs, une démonstration analogue de la convexité peut-être réalisée pour l'équation 4.27.

Les équations 4.26 et 4.27 sont dérivables et admettent un minimum global pour $b \in]0; +\infty[$. Ce minimum est l'unique solution de

$$\frac{\partial E_b}{\partial b}(b) = 0. \quad (4.28)$$

Pour un canal AWGN, la dérivée de l'énergie consommée par élément binaire transmis, donnée par l'équation 4.26, peut s'écrire sous la forme :

$$\frac{\partial E_b}{\partial b}(b) = \frac{A (2^b - 1)^{1-\beta}}{b^2 (2^{b/2} + 1)^2} (2^b ((1 - \beta)b \ln(2) - 1) + 2^{b/2} b \ln(2) + 1) - \frac{C}{b^2}. \quad (4.29)$$

En solutionnant l'égalité 4.28, nous obtenons donc :

$$\begin{aligned} 2^b ((1 - \beta)b \ln(2) - 1) + 2^{b/2} b \ln(2) + 1 &= \frac{C (2^{b/2} + 1)^2}{A (2^b - 1)^{1-\beta}} \\ \Leftrightarrow 2^b ((1 - \beta)b \ln(2) - 1) &= \frac{C (2^{b/2} + 1)^2}{A (2^b - 1)^{1-\beta}} - 2^{b/2} b \ln(2) - 1 \\ \Leftrightarrow \exp(b \ln(2) - \frac{1}{1-\beta})(b \ln(2) - \frac{1}{1-\beta}) &= \frac{\exp \frac{1}{1-\beta}}{1-\beta} \left(\frac{C (2^{b/2} + 1)^2}{A (2^b - 1)^{1-\beta}} - 2^{b/2} b \ln(2) - 1 \right) \\ \Leftrightarrow b \ln(2) - \frac{1}{1-\beta} &= W \left(\frac{\exp \frac{1}{1-\beta}}{1-\beta} \left(\frac{C (2^{b/2} + 1)^2}{A (2^b - 1)^{1-\beta}} - 2^{b/2} b \ln(2) - 1 \right) \right) \end{aligned}$$

avec W la fonction de Lambert. Finalement, l'ES b^* qui minimise l'énergie consommée, que nous nommerons (Bopt.A.proposé), peut s'écrire

$$b^* = \frac{1}{\ln(2)} \left(\frac{1}{1-\beta} + W \left(\frac{e^{-\frac{1}{1-\beta}}}{1-\beta} \left(\frac{C (2^{b^*/2} + 1)^2}{A (2^{b^*} - 1)^{1-\beta}} - 2^{b^*/2} b^* \ln(2) - 1 \right) \right) \right). \quad (4.30)$$

Pour le canal de Rayleigh, une résolution similaire nous donne la solution (Bopt.R.proposé) suivante :

$$b^* = \frac{\alpha_1}{\ln(2)} \left(\frac{1}{1-\beta} + W \left(\frac{e^{-\frac{1}{1-\beta}}}{1-\beta} \left(\frac{C 2^{b^*/2} + 1}{A' 2^{b^*/2} - 1} (2^{b^*} - 1)^\beta - 2^{b^*/2} b^* \ln(2) \frac{2^{b^*/\alpha_1} - 1}{2^{b^*} - 1} - 1 \right) \right) \right). \quad (4.31)$$

Néanmoins, on remarque que les deux membres des équations 4.30 et 4.31 dépendent de b^* , l'ES optimale ne peut donc être calculée directement. Dans l'ensemble de nos simulations, un algorithme de Newton est utilisé pour obtenir la solution b^* .

4.3.3 Résultats expérimentaux

Dans cette section, on étudie l'évolution de l'ES optimale b^* minimisant l'énergie consommée par bit transmis pour une M-QAM en fonction de différents paramètres. Bien que les valeurs pratiques de l'ES b_d^* soient des entiers naturels pairs, les résultats seront ici présentés pour un b^* continu, ceci afin d'analyser avec précision les tendances des courbes même pour de faibles variations de l'ES, pour déterminer les principaux paramètres qui agissent sur l'énergie, et pour

vérifier l'exactitude de nos formules. Les paramètres considérés sont toujours ceux des tables 3.1 et 4.2.

Nous allons analyser l'impact des paramètres suivants : la distance entre le transmetteur et le récepteur d_{sd} , la largeur de bande B , le coefficient de l'efficacité du drain β , le facteur de roll-off ρ , la PEB sur le canal P_{eb} , et les résolutions binaires n_1 du CNA et n_2 du CAN. Dans chaque expérience, on compare des valeurs de référence issue d'un algorithme de recherche brut avec les résultats de l'expression 4.30 de la solution (Bopt.A.proposé) pour $\Gamma(P_{eb})$ de l'état de l'art 3.109 ou proposé en 4.17 pour le canal AWGN, et avec les résultats de l'expression 4.31 de la solution (Bopt.R.proposé) pour $\Gamma(P_{eb})$ proposé en 4.23 pour le canal de Rayleigh.

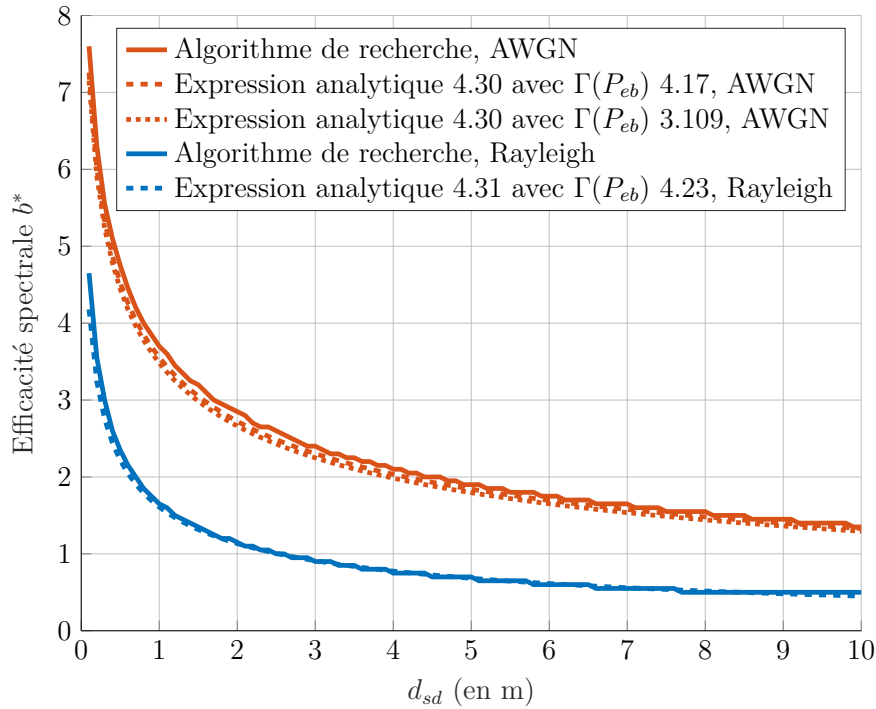


FIGURE 4.16 – Efficacité spectrale b^* minimisant l'énergie consommée par bit transmis pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction de la distance d_{sd}

Dans un premier temps, on étudie l'influence de la distance source-destinataire d_{sd} sur la figure 4.16. Comme dans la figure 4.15, l'ES optimale b^* décroît exponentiellement avec la distance car la puissance de transmission, qui en dépend conformément à l'équation 3.75, gagne en importance par rapport à la puissance du circuit. Quand l'ES continue b^* est inférieure à 3, l'ES discrète b_d^* est égale à 2, c'est-à-dire à sa valeur minimale ; en prenant en compte cette relation, on constate que la modulation optimale est une 4-QAM pour toute transmission avec $d_{sd} > 1.7\text{m}$ pour un canal AWGN, et $d_{sd} > 0.3\text{m}$ pour un canal de Rayleigh. Cela signifie que la 4-QAM devrait être la modulation la plus courante pour minimiser l'énergie totale consommée. En outre, on constate que la courbe pour un canal de Rayleigh est toujours en dessous de celle pour un canal AWGN, en partie à cause de $\Gamma_{Rayleigh}(P_{eb})$. Finalement, la conclusion la plus importante à reconnaître est l'approximation rigoureuse des courbes réelles par les expressions analytiques proposées.

Dans un deuxième temps, on s'intéresse à l'influence de la largeur de bande passante B . Dans cette simulation, on pose $d_{sd} = 1\text{m}$ pour les deux canaux AWGN et de Rayleigh. La figure 4.17

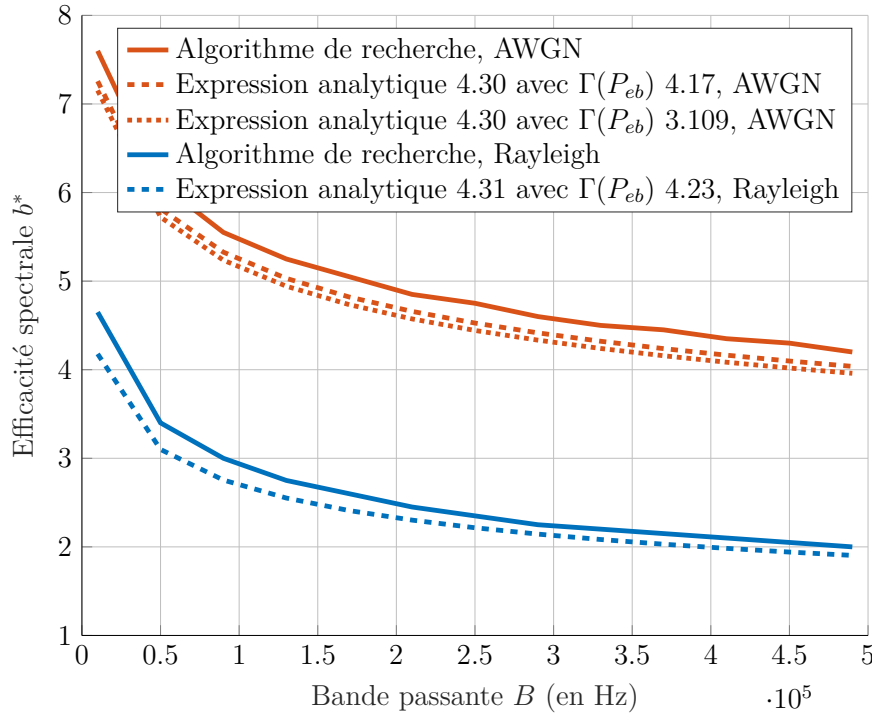


FIGURE 4.17 – Efficacité spectrale b^* minimisant l'énergie consommée par bit transmis pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction de la largeur de bande B

montre une décroissance lente de l'ES, ce qui est cohérent avec la présence du terme B à la fois dans le terme de l'énergie lié à la puissance de transmission et celui lié à la puissance du circuit dans les équations 4.26 et 4.27. Malgré une petite différence entre les valeurs analytiques et de référence, les deux courbes résultantes suivent la même tendance. Pour la gamme de largeurs de bande présentée, le paramètre optimal b_d^* est égal à 2 ou 4 pour un canal de Rayleigh, et est compris entre 4 et 8 pour un canal AWGN.

Dans un troisième temps, on étudie l'impact du coefficient $\beta \in [0, 1]$ de puissance de l'efficacité du drain de l'AP, issu de l'équation 4.7. Comme nous nous intéressons à des valeurs de β réalistes, c'est-à-dire en accord avec des transmetteurs-émetteurs RF existants, seule la gamme de valeurs $[0.2, 0.8]$ sera analysée. La figure 4.18 nous montre l'ES optimale b^* obtenue à partir des équations 4.30 et 4.31 pour une distance source-destinataire spécifique $d_{sd} = 1\text{m}$. Des résultats similaires à ceux déjà identifiés pour le paramètre B peuvent être observés, c'est-à-dire une courbe à décroissance lente. Néanmoins, le déclin est ici presque linéaire avec l'augmentation du paramètre β .

Dans un quatrième temps, l'influence du facteur de roll-off $\rho \in [0, 1]$ est analysé sur la figure 4.19, sur laquelle seules des valeurs réalistes entre 0.2 et 0.8 sont considérées. Puisque le facteur de roll-off n'a d'impact que sur le PAPR, on peut s'attendre à un effet négligeable, comme on l'avait observé sur la consommation dans la section 4.2.1. En effet, le paramètre continu b^* varie entre 1.33 et 1.65 pour le canal de Rayleigh, et entre 3.15 et 3.75 pour le canal AWGN, ce qui en fait des variations bien plus faibles que précédemment. En conséquence, la modulation choisie pour ces valeurs du facteur de roll-off est toujours une 4-QAM pour le canal de Rayleigh, et une 16-QAM pour le canal AWGN. Pour le canal AWGN, la comparaison entre les valeurs de la

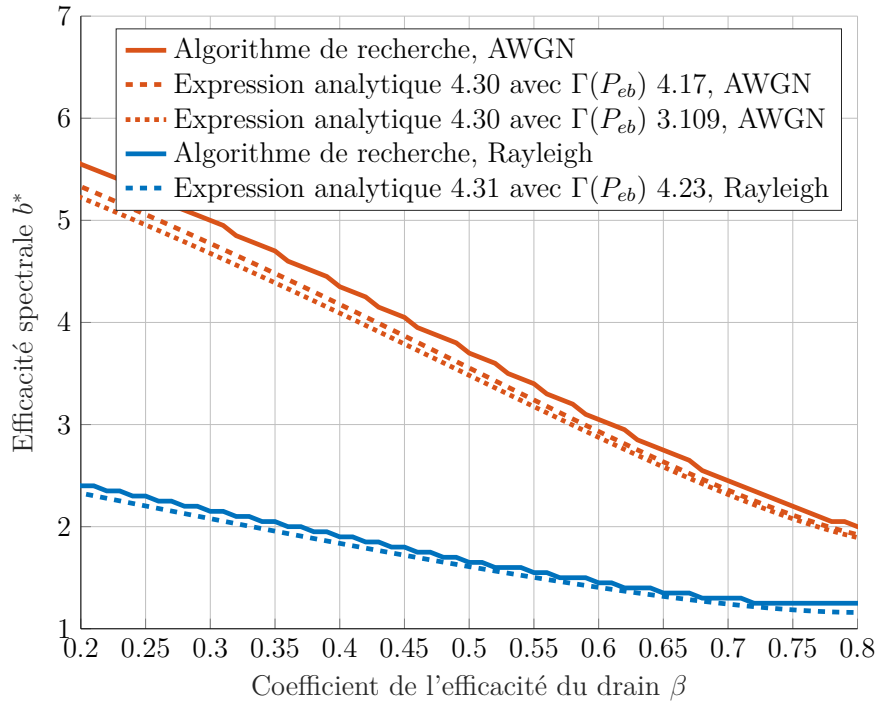


FIGURE 4.18 – Efficacité spectrale b^* minimisant l'énergie consommée par bit transmis pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction du coefficient de l'efficacité du drain β

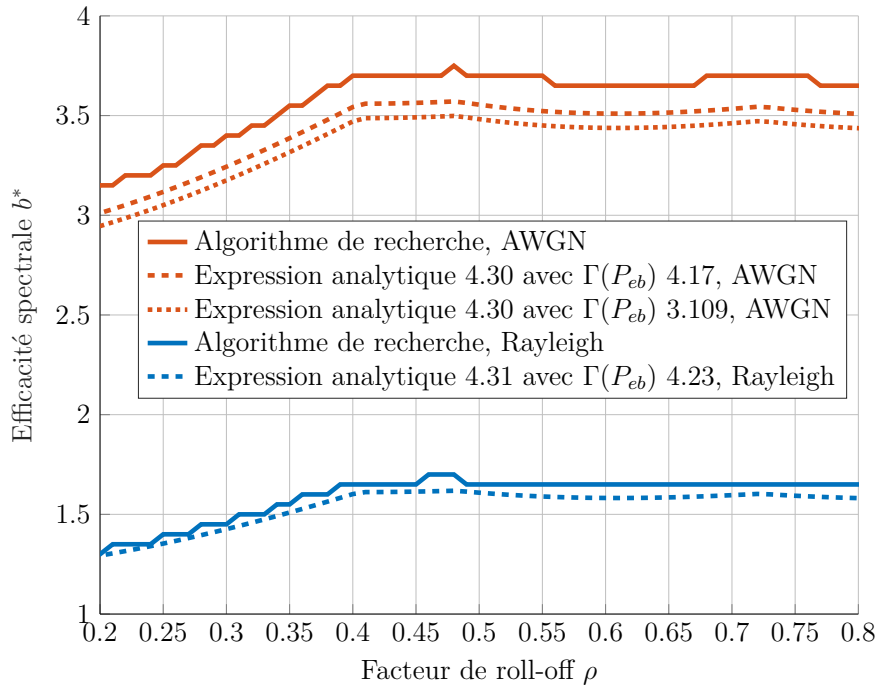


FIGURE 4.19 – Efficacité spectrale b^* minimisant l'énergie consommée par bit transmis pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction du facteur roll-off ρ

solution optimale 4.30 en utilisant une pénalité en RSB $\Gamma(P_{eb})$ exprimée par l'équation 3.111 (issue de l'état de l'art) ou 4.17 (proposée) nous montre une légère amélioration en faveur du modèle proposé en 4.17. Pour un canal de Rayleigh à évanouissement plat, les solutions trouvées avec 4.31 en utilisant la pénalité en RSB de l'équation 4.23 nous montrent des performances satisfaisantes quant à la capacité du modèle proposé à modéliser une consommation réaliste.

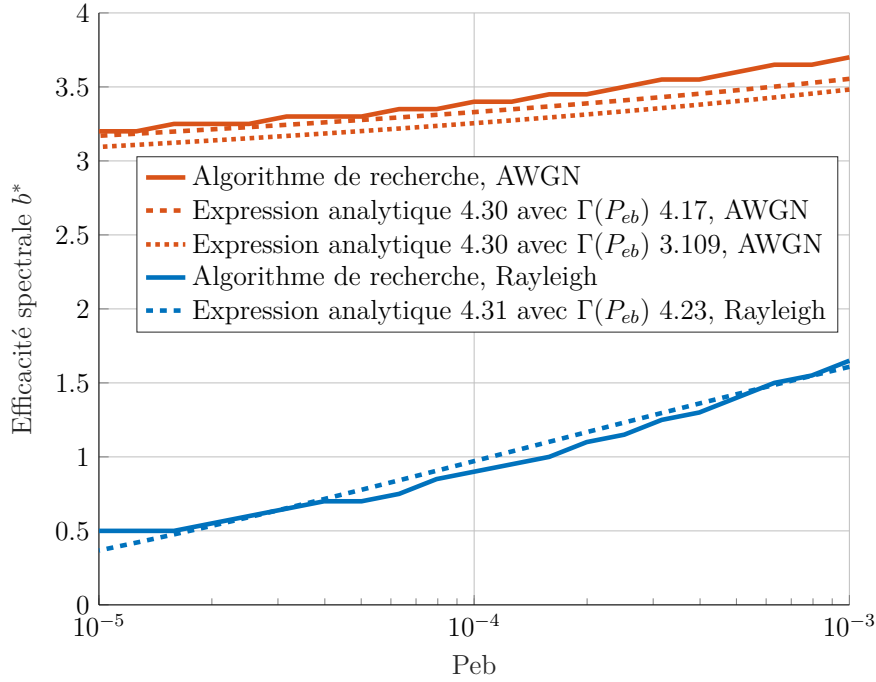


FIGURE 4.20 – Efficacité spectrale b^* minimisant l'énergie consommée par bit transmis pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction de la PEB

Cinquièmement, on cherche à comprendre l'influence de la PEB. Ayant noté sur la figure 3.22a que l'approximation $\theta \approx b$ est valable pour les PEB inférieures à 10^{-3} pour un canal AWGN, on fait varier sur la figure 4.20 la PEB dans la gamme de valeurs $[10^{-5}, 10^{-3}]$ pour analyser les changements dans l'ES b^* . Comme on peut le voir, l'impact d'une telle variation est modéré, avec b^* croissant de 3.2 à 3.4 pour un canal AWGN, ce qui signifie que le b_d^* correspondant est toujours égal à 4. Pour un canal de Rayleigh, l'accroissement de b^* est plus important mais ses valeurs restent toujours inférieures à 2, ce qui conduit à une 4-QAM pour toutes les valeurs du paramètre P_{eb} . En regard des expressions analytiques (Bopt.A.proposé) et (Bopt.R.proposé) proposées par les équations 4.30 et 4.31, on observe que les courbes correspondantes suivent fidèlement les valeurs réelles obtenues par algorithme de recherche brut.

Finalement, nous étudions les effets des résolutions binaires n_1 (pour le CNA) et n_2 (pour le CAN) introduites au sein du modèle de la puissance du circuit dans la section 4.2.2. Les figures 4.21 et ?? révèlent que, comme dans les simulations précédentes, les expressions analytiques 4.30 et 4.31 donnent une bonne approximation des valeurs réelles de b^* . Pour le canal AWGN, la figure 4.21 indique une augmentation marquée de l'ES avec l'accroissement de n_1 , ce qui conduit à une ES discrète progressant de 4 (quand $n_1 = 10$) à 8 (quand $n_1 = 16$). Pour n_2 , l'évolution est bien plus faible, et la taille de constellation b_d^* ne sera pas modifiée par le nombre de bits de décision. Pour le canal de Rayleigh, des observations analogues peuvent être faites.

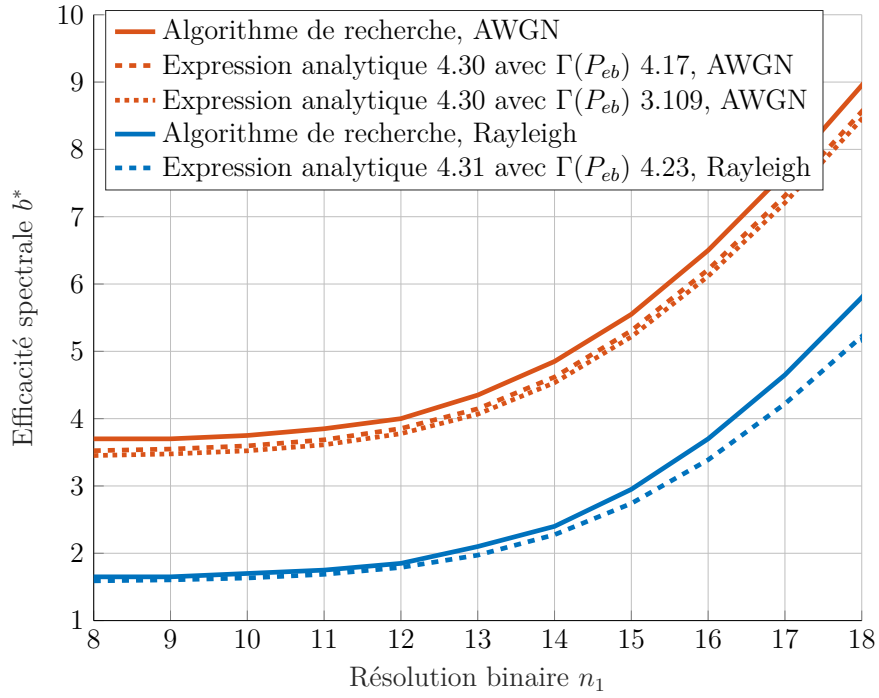
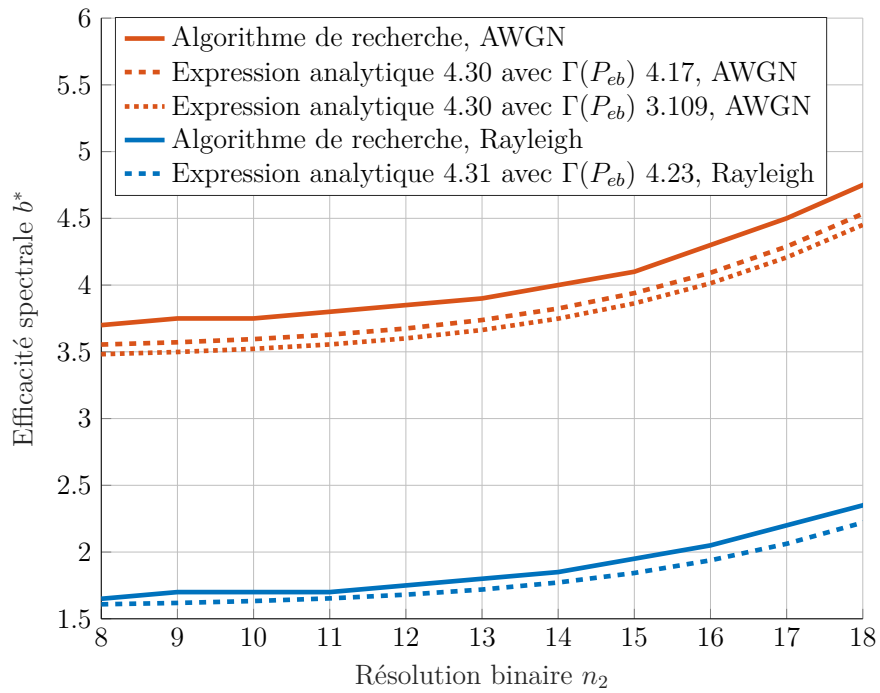


FIGURE 4.21 – Efficacité spectrale b^* minimisant l'énergie consommée par bit transmis pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction de la résolution binaire n_1 du CNA



4.4 Influence de la position du relai

Jusqu'à présent, nous avons majoritairement étudié l'énergie et l'ES d'une communication sans fil point-à-point. Pour une communication coopérative, le lien entre la capacité du canal équivalent et le RSB de chacune des liaisons est beaucoup moins direct puisque la capacité dépend du RSB en sortie du MRC tandis que l'énergie dépend des RSB sur les liens source-destinataire et relai-destinataire [134].

La section 3.5.2 a été l'occasion de présenter des courbes expérimentales de la consommation d'une transmission coopérative et l'influence de la position du relai le long du segment source-destinataire. Si ces études ont permis quelques premières conclusions, elles ne reflètent néanmoins pas la réalité avec précision. En effet, outre l'utilisation d'un modèle de consommation simplifié, les relais sont rarement positionnés dans la lignée des nœuds source et destinataire pour les applications à faible portée car, à cette distance, la présence de ces relais peut impacter physiquement le canal. En outre, il n'est pas toujours possible de fixer le nœud à cet emplacement spécifique ; par exemple, si une source et un destinataire sont placés de part et d'autre d'une rue, il n'est pas envisageable de positionner un relai au centre. Par conséquent, il est d'un intérêt crucial de connaître l'influence de la position du nœud-relai dans un espace au moins à deux dimensions.

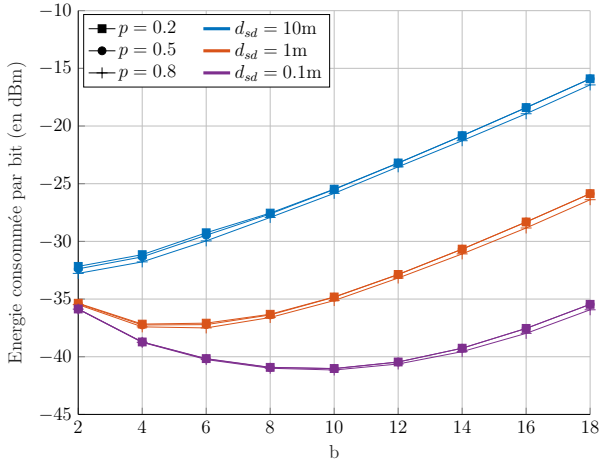
Dans un premier temps, on cherchera à retrouver les résultats de la consommation pour un relai le long du segment source-destinataire avec le modèle avancé de la consommation. Dans un second temps, on s'intéressera à l'influence de la position du relai dans un espace à deux dimensions sur la consommation d'énergie totale et sur l'ES.

4.4.1 Variation du relai sur l'axe source-destinataire

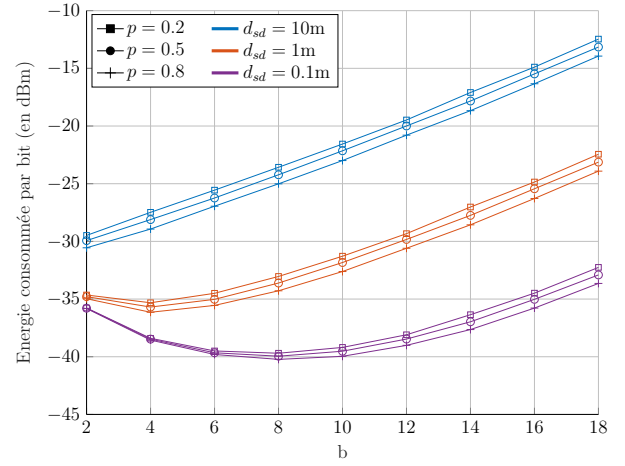
Dans la section 3.5.2, le modèle de consommation d'énergie d'une communication coopérative a été présenté, ainsi que la manière de calculer l'énergie minimale en choisissant le meilleur couple $(\gamma_{sd}, \gamma_{rd})$ de RSB sur les liaisons source-destinataire et relai-destinataire. Suivant cette même méthode, la consommation et le gain de coopération pour le modèle avancé de l'énergie ont été étudiés dans la section 4.2.1. Ici, nous cherchons à analyser l'influence de la position du relai le long de l'axe source-destinataire. Sauf mention du contraire, les paramètres sont toujours identiques à ceux des expériences précédentes.

L'impact de la position normalisée du relai p , définie dans la section 3.5.2 comme étant le rapport entre les distances source-relai et source-destinataire d_{sr}/d_{sd} , est étudié sur la figure 4.22 pour des valeurs de $p \in \{0.2, 0.5, 0.8\}$. Pour le canal AWGN, dont les résultats sont présentés sur la figure 4.22a, on remarque très peu de différence entre les courbes, et ce quelle que soit la distance d_{sd} ; la différence maximale est de 0.53 dBm pour $d_{sd} = 10\text{m}$ et $b = 18$ entre les courbes pour $p = 0.2$ et $p = 0.8$. On remarque toutefois que, plus le relai est proche du destinataire - c'est-à-dire plus la position normalisée p est proche de 1 -, plus la consommation est faible. Une remarque similaire peut être faite en observant les courbes pour le canal de Rayleigh sur la figure 4.22b. Néanmoins, on remarque ici que la position du relai a une plus grande importance, avec une différence allant jusqu'à 1.47dBm dans le même contexte que précédemment.

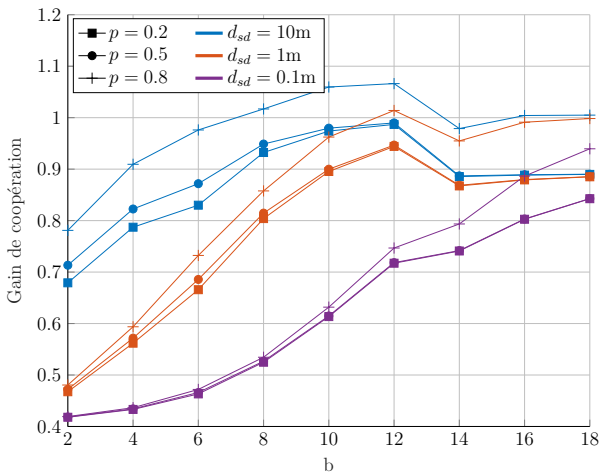
L'énergie consommée nous permet de calculer le gain de coopération, dont on affiche les courbes pour différentes valeurs de position normalisée p sur la figure 4.23. Sans surprise, on observe pour le canal AWGN sur la figure 4.23a que les courbes de gain coopératif sont très proches les unes des autres pour une même distance d_{sd} , en particulier pour les constellations



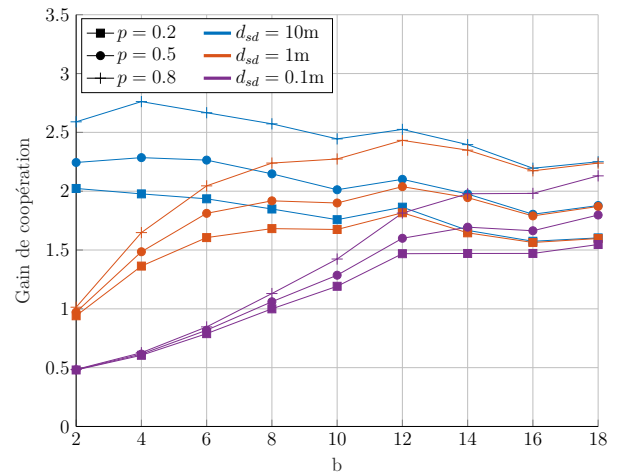
(a) Consommation pour un canal AWGN



(b) Consommation pour un canal de Rayleigh

 FIGURE 4.22 – Énergie consommée par bit lors d'une communication coopérative pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction de l'ES pour différentes positions du relai sur l'axe source-destinataire


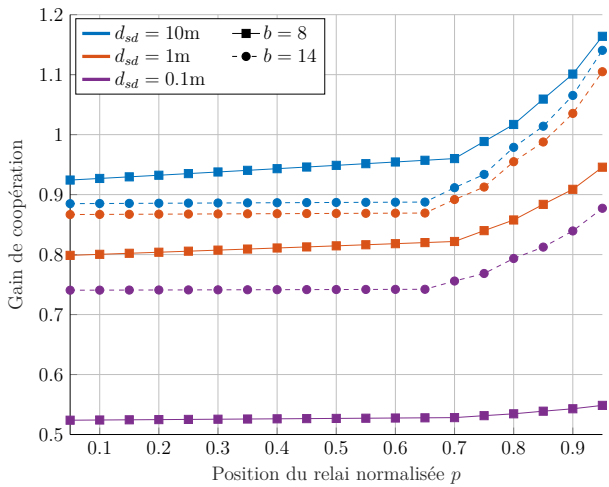
(a) Gain coopératif pour un canal AWGN



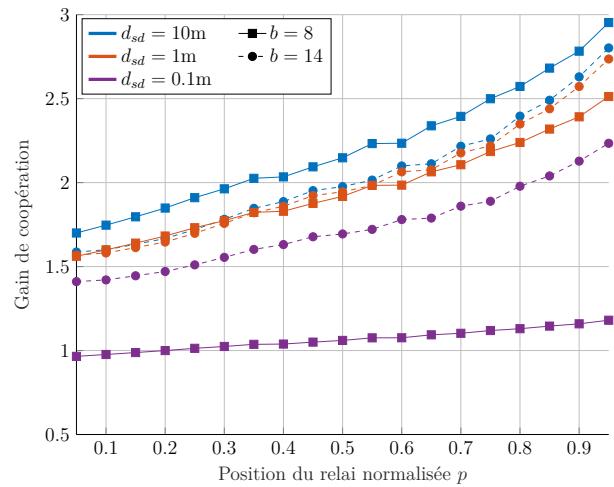
(b) Gain coopératif pour un canal de Rayleigh

 FIGURE 4.23 – Gain de coopération pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction de l'ES pour différentes positions du relai sur l'axe source-destinataire

de petites tailles, quand la puissance de transmission est faible par rapport à la puissance du circuit. Pour une même valeur de la position normalisée p , il semble par ailleurs que le gain de coopération converge vers une valeur particulière. Néanmoins, comme on l'avait remarqué dans la section 4.2.1, le gain de coopération nous indique que le relai n'est pas avantageux énergétiquement pour des communications faible portée sur un canal AWGN. Pour le canal de Rayleigh, des remarques analogues peuvent être soulignées, excepté concernant l'utilité énergétique du relai. En effet, le gain de coopération est pratiquement toujours supérieur à 1 pour $d_{sd} = 10\text{m}$ ou 1m , et dès $b = 8$ pour $d_{sd} = 0.1\text{m}$.



(a) Gain coopératif pour un canal AWGN

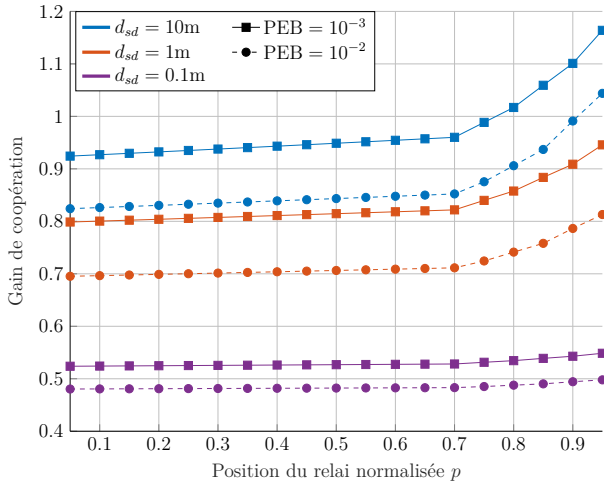


(b) Gain coopératif pour un canal de Rayleigh

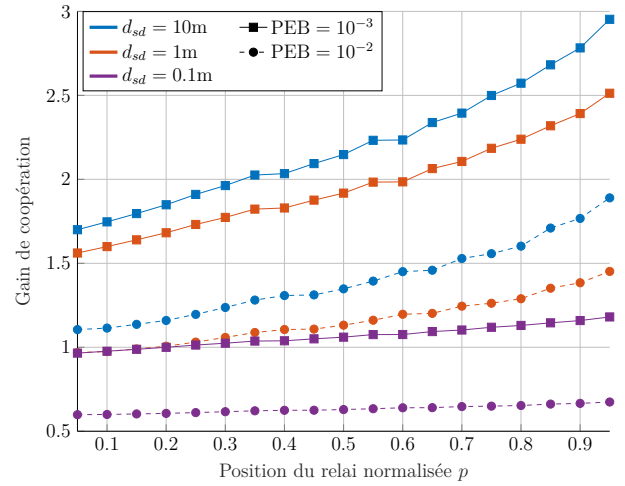
FIGURE 4.24 – Gain de coopération pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction de la position normalisée du relai pour différentes tailles de constellation

On cherche à confirmer nos hypothèses sur l'influence de la position du relai au moyen de la figure 4.24. Celle-ci présente le gain coopératif pour différentes tailles de constellation $b \in \{8, 14\}$ en fonction de la position normalisée p du relai. On commence par remarquer que, pour les deux canaux AWGN sur la figure 4.24a et de Rayleigh sur la figure 4.24b, le gain de coopération est le plus fort quand le relai s'approche du destinataire. En outre, on note que pour $d_{sd} = 1\text{m}$ et $d_{sd} = 0.1\text{m}$, le gain de coopération est plus grand pour $b = 14$ que pour $b = 8$, alors que l'inverse se produit quand $d_{sd} = 10\text{m}$. Ce phénomène apparaît également sur les courbes de la figure 4.23 car le gain a tendance à croître avec la taille de constellation pour $d_{sd} = 1\text{m}$ et $d_{sd} = 0.1\text{m}$, tandis qu'il diminue à partir d'une certaine valeur de b pour $d_{sd} = 10\text{m}$.

Enfin, on s'intéresse également à l'influence de la PEB sur le gain de coopération et la localisation optimale du relai. Pour cela, on fixe $b = 8$. Comme on peut l'observer sur la figure 4.25, le gain croît quand le PEB diminue. Le relai est donc plus intéressant énergétiquement pour les faibles probabilités d'erreur. En revanche, la tendance du gain en fonction de la position reste identique aux expériences précédentes : pour les distances d_{sd} de 10m ou 1m, le gain croît fortement quand le relai approche du destinataire ; pour 0.1m, le gain augmente faiblement et reste quasiment constant.



(a) Gain coopératif pour un canal AWGN



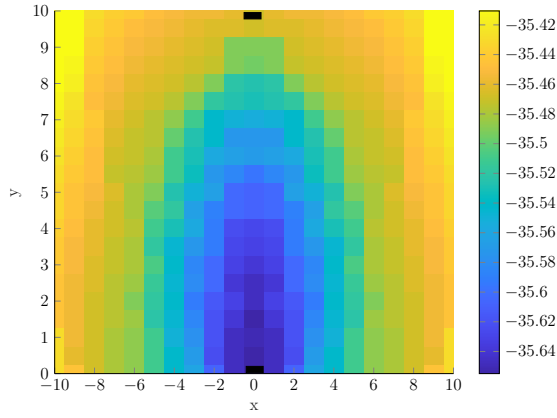
(b) Gain coopératif pour un canal de Rayleigh

 FIGURE 4.25 – Gain de coopération pour une M-QAM ($M = 2^b$) sur un canal AWGN ou de Rayleigh en fonction de la position normalisée du relai pour différentes PEB

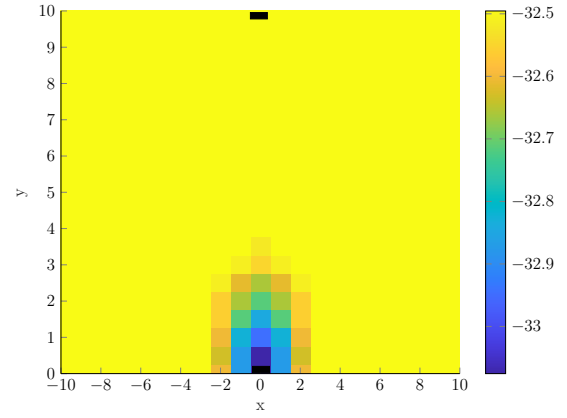
4.4.2 Variation du relai dans un espace à deux dimensions

On étudie à présent l'influence de la position du relai dans un environnement à deux dimensions. En effet, dans un système multi-nœuds embarqué, les nœuds seront placés en fonction des emplacements disponibles. Pour les expériences suivantes, on considère un espace à deux dimensions de longueur maximale d_{sd} sur l'axe des abscisses et celui des ordonnées. Le destinataire est alors positionné au point $(0,0)$, la source en $(0, d_{sd})$, et le relai en (x,y) , $-d_{sd} < x < d_{sd}$ et $0 < y < d_{sd}$. Les distances d_{sr} et d_{rd} seront déduites de cette position. On fait alors l'hypothèse d'une symétrie axiale des résultats le long de l'axe source-destinataire puisque seules les mesures de distance sont utilisées dans le calcul de la consommation. On fixe la longueur entre la source et le destinataire à $d_{sd} = 10\text{m}$, le nombre de bits par symbole QAM à $b = 2$ et la PEB à $\text{Peb} = 10^{-3}$. On comparera pour chaque simulation les modèles simplifié et avancé de la consommation d'énergie.

On commence par étudier l'énergie minimale consommée pour chacune des positions décrites du relai fonctionnant selon le protocole AF. Les résultats sont présentés sur la figure 4.26 pour le canal AWGN, et sur la figure 4.27 pour le canal de Rayleigh. Sur chaque figure, la source est placée tout en haut au milieu (rectangle noir en $(0, d_{sd})$), et le destinataire, en bas au milieu (rectangle noir en $(0, 0)$). Sur la figure 4.26a, on peut observer la consommation sur un canal AWGN pour un coefficient multiplicateur de la puissance de l'amplificateur α fixé à 1.9 (modèle (E.coop.trad)), et sur la figure 4.26b pour α variant en fonction du système (modèle (E.coop.proposé)). Dans un premier temps, on remarque, comme noté dans la section 4.2.1, que pour toute position du relai, l'énergie consommée par le modèle énergétique avancé (E.coop.proposé) est supérieure à celle du modèle simplifié (E.coop.trad). Dans un second temps, on confirme qu'il existe une symétrie par rapport à l'axe source-destinataire. Enfin, on remarque que, sur les deux figures, la consommation minimale est obtenue quand le relai est au plus proche du destinataire. Néanmoins, cette conclusion est à nuancer : en effet, les variations de consommation entre les différentes positions sont très faibles, moins de 0.5 dBmJoule. Ceci nous indique qu'en vérité, quelle que soit la position du relai (du moins dans l'espace où nous travaillons), si les puissances d'émission de la source et du relai sont correctement adaptées,

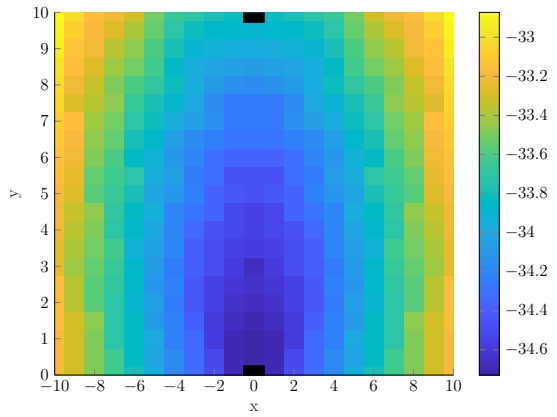


(a) Consommation pour le modèle (E.coop.trad)

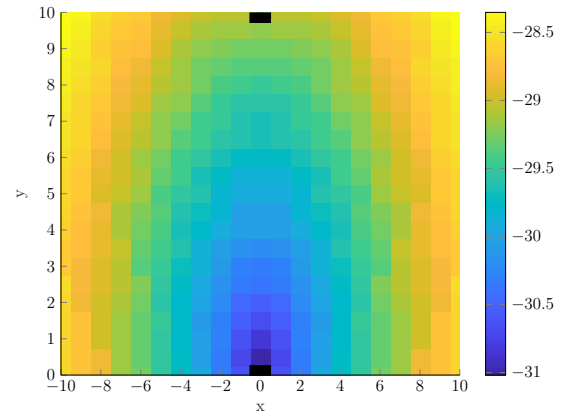


(b) Consommation pour le modèle (E.coop.proposé)

FIGURE 4.26 – Energie minimale consommée par bit lors d’une communication coopérative AF en fonction de la position du relai pour un canal AWGN, une ES $b = 2$ et une distance source-destinataire $d_{sd} = 10\text{m}$



(a) Consommation pour le modèle (E.coop.trad)



(b) Consommation pour le modèle (E.coop.proposé)

FIGURE 4.27 – Energie minimale consommée par bit lors d’une communication coopérative AF en fonction de la position du relai pour un canal de Rayleigh, une ES $b = 2$ et une distance source-destinataire $d_{sd} = 10\text{m}$

alors il n'y a pas de grande différence de consommation d'énergie. Enfin, on souligne que les conclusions restent les mêmes en faisant varier les valeurs de d_{sd} et de b , seule la valeur moyenne de l'énergie change.

De l'observation des figures 4.27a et 4.27b présentant des résultats analogues pour un canal de Rayleigh, il est possible d'aboutir aux mêmes conclusions.

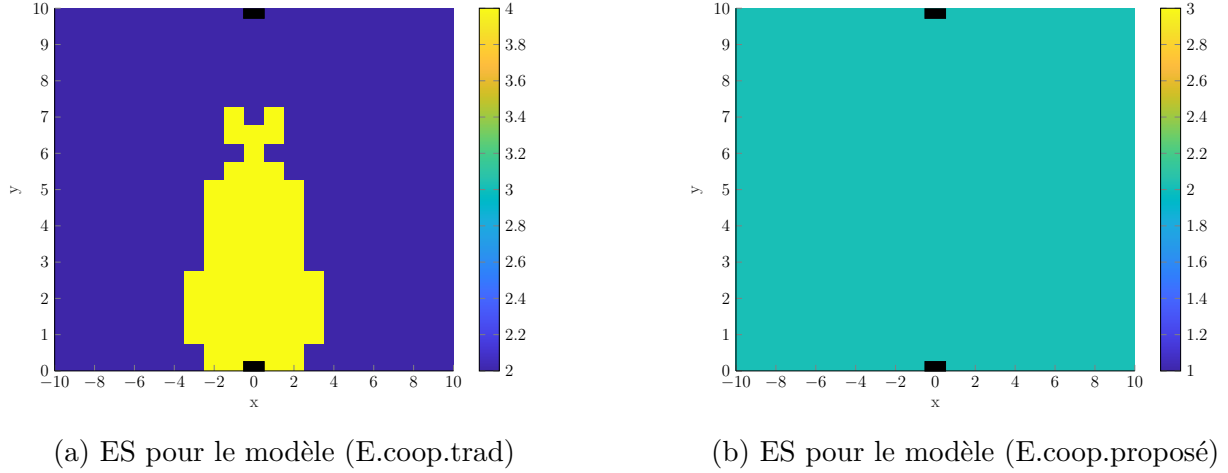


FIGURE 4.28 – Efficacité spectrale minimisant l'énergie consommée par bit émis lors d'une transmission coopérative AF en fonction de la position du relai pour un canal de Rayleigh et une distance source-destinataire $d_{sd} = 10\text{m}$

On étudie à présent l'ES discrète b_d^* qui minimise la consommation en chaque position du relai. La distance entre d_{sd} est toujours fixée à 10m. Sur la figure 4.28a, on observe la variation de b_d^* pour le modèle énergétique simplifié (E.coop.trad) et un canal de Rayleigh. On remarque que les positions proches du destinataire permettent d'obtenir une ES supérieure, de $b_d^* = 4$ contre $b_d^* = 2$ partout ailleurs. Pour le modèle énergétique avancé (E.coop.proposé) sur un canal de Rayleigh, la valeur de l'ES affichée sur la figure 4.28b est uniformément égale à 2.

Nous ne présentons pas les figures pour un canal AWGN, qui ne comportent pas non plus de variation : pour le modèle (E.coop.proposé), b_d^* est toujours égal à 2, et pour le modèle (E.coop.trad), b_d^* est égal à 4. Par ailleurs quand la distance d_{sd} augmente, l'ES diminue jusqu'à sa valeur minimale de 2.

Ces deux premières expériences laissent à penser que ni la position ni même la présence du relai n'a d'intérêt. Cependant, l'étude du gain de coopération souligne la pertinence de ces deux éléments.

Dans un premier temps, on observe sur la figure 4.29 le gain de coopération pour un canal AWGN avec $b = 2$ et toujours une distance $d_{sd} = 10\text{m}$. Sur les figures 4.29a et 4.29b montrant respectivement les gains pour un modèle énergétique simplifié (E.coop.trad) et avancé (E.coop.proposé), les gains sont toujours inférieurs à 1. On note toutefois qu'ils gagnent en importance quand on s'approche du destinataire. De plus, on note que les gains sont plus importants – entre 0.73 et 0.84 – pour le modèle avancé que pour le modèle simplifié – entre 0.47 et 0.52. Néanmoins, en augmentant la distance d_{sd} , le gain va lui aussi augmenter progressivement. A 50m par exemple, il présente déjà un avantage pour les positions proches du destinataire (positions en vert et jaune sur les figures 4.29a et 4.29b). En continuant d'augmenter la distance d_{sd} , il sera donc possible d'élargir le champ des positions possibles du relai

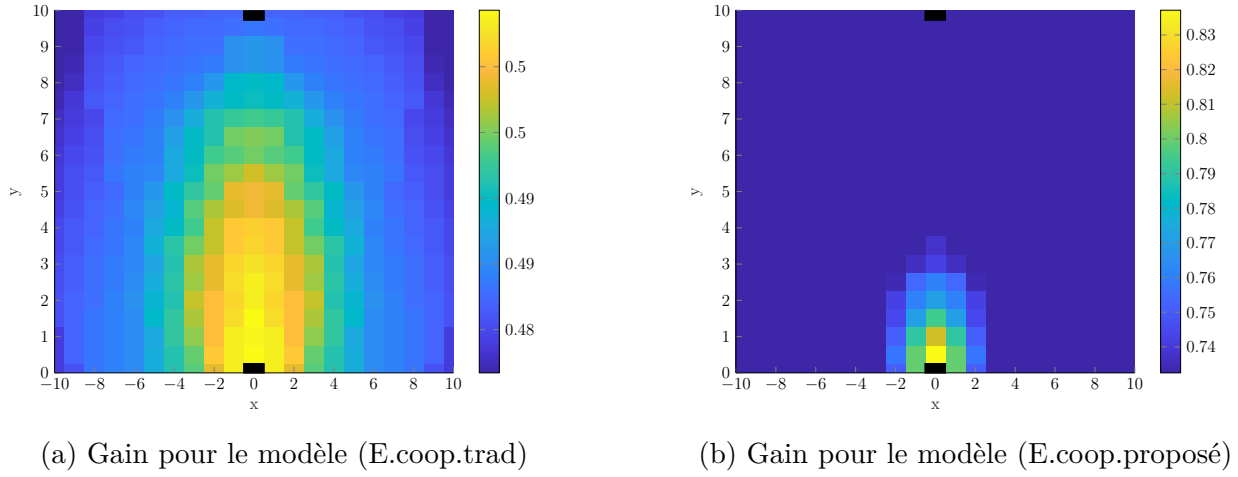


FIGURE 4.29 – Gain de coopération en fonction de la position du relai pour un canal AWGN, une ES $b = 2$ et une distance source-destinataire $d_{sd} = 10\text{m}$

pour lesquelles l'utilisation du relai est intéressante. De même, il est également possible de faire varier la valeur de b . Les résultats, de tendances similaires à ce que l'on peut observer sur la figure 4.29 mais de valeurs différentes, montrent alors qu'il existe des ordres de modulation pour lesquels l'existence du relai est avantageuse. Par exemple, quand $b = 8$, le gain est toujours supérieur à 1 sur notre domaine de définition.

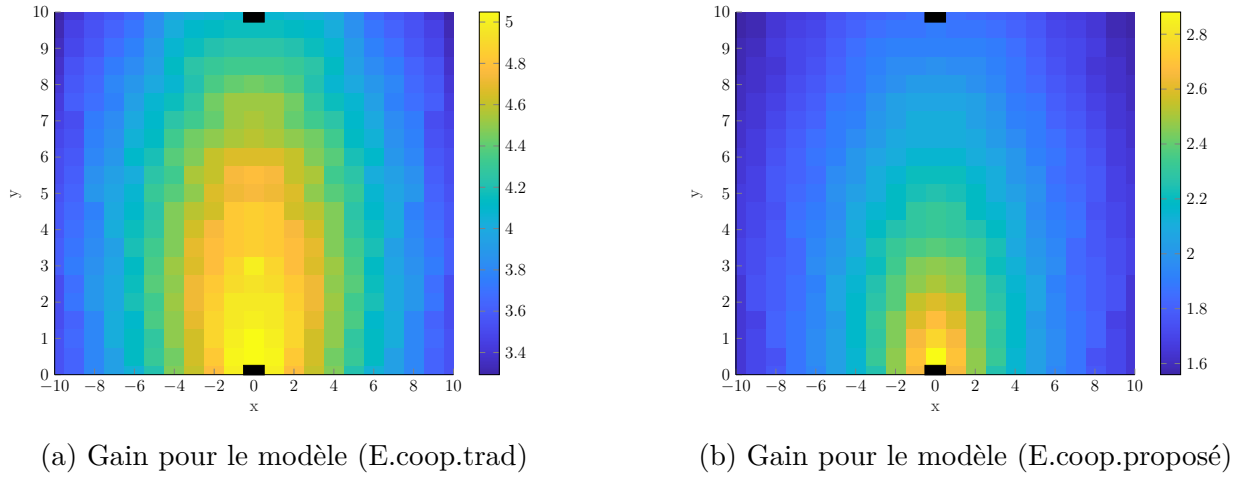


FIGURE 4.30 – Gain de coopération en fonction de la position du relai pour un canal de Rayleigh, une ES $b = 2$ et une distance source-destinataire $d_{sd} = 10\text{m}$

Quand on étudie le canal de Rayleigh, l'avantage du relai se présente pour des distances bien plus faibles. Par exemple, on note sur les figures 4.30a et 4.30b, obtenues avec $d_{sd} = 10\text{m}$ et $b = 2$, que les gains sont déjà très supérieurs à 1 pour toutes les positions. En augmentant la distance, on remarque une croissance du gain comme pour le canal AWGN. En revanche, quand on accroît la valeur de la taille de constellation b , le changement n'est pas aussi catégorique que pour le canal AWGN : les gains peuvent être plus grands ou plus petits en fonction du b choisi ; ce n'est pas linéaire.

4.5 Conclusion

Dans ce chapitre, nous nous sommes attachés à étudier un modèle de consommation plus proche d'un système réel pour des communications point-à-point et coopératives. Pour cela, nous avons dans un premier temps présenté un état de l'art des techniques qui modélisent avec plus de réalisme certains éléments de la consommation, comme les puissances du CNA et du CAN, le PAPR ou encore l'efficacité du drain. Chacune de ces améliorations a été ensuite intégrée conjointement au sein d'un même modèle avancé de consommation d'énergie par bit transmis, ce qui nous a conduit à mener une étude originale sur les différents paramètres influençant cette consommation d'énergie. L'impact de la contrainte correspondant à la limitation de la puissance de transmission, encore jamais prise en compte à notre connaissance, a également été étudiée.

Disposant d'un modèle de consommation plus avancé, les équations de l'ES pour une communication point-à-point présentées dans le chapitre 3 se sont révélées caduques en raison de la dépendance du coefficient de l'AP à l'ES pour une M-QAM. Une nouvelle expression analytique de l'ES prenant en compte ce nouveau modèle a donc été exprimée. Des résultats expérimentaux indiquant l'influence des différents paramètres du système sur l'ES minimisant la consommation d'énergie ont ensuite été présentés.

Finalement, nous nous sommes intéressés à la consommation d'une communication coopérative en conduisant des analyses sur l'impact de la position du relai dans un environnement à une ou deux dimensions. Ceci nous a notamment permis de comprendre que, quels que soient la distance ou le modèle de consommation, avec notre système, l'énergie consommée est toujours plus faible quand le relai approche du destinataire.

Au cours des deux dernières sections de ce chapitre, nous avons ainsi pu étudier le modèle de consommation d'une communication point-à-point ou coopérative dans un environnement réaliste, permettant notamment de lier consommation et PEB. La chaîne de transmission sans fil étant ainsi simulée, nous allons pouvoir nous intéresser à introduire la notion de système embarqué dans une application de classification d'événements sonores.

DÉTECTION D'ÉVÉNEMENTS SONORES AU SEIN D'UN SYSTÈME EMBARQUÉ MULTICAPTEURS

5.1 Introduction

Le développement de méthodes de traitement du signal fonctionnelles provoquent régulièrement un intérêt pour une implémentation matérielle, et la détection d'événements sonores ne déroge pas à cette règle. Ces dernières années, un grand nombre de recherches se sont penchées sur le sujet de la classification sonore à visées applicatives [17, 74, 121, 228]. En raison de la diffusion de l'usage des smartphones, équipés de microphones, le problème est généralement vu sous le prisme d'une application mobile dont les calculs sont réalisés via un *cloud* [17, 34, 121]. Cependant, le développement de l'Internet des objets permet aujourd'hui à des applications bien plus diverses d'exister sur des réseaux de capteurs variés [84, 180, 182]. Il est ainsi possible de travailler avec des systèmes embarqués multi-capteurs qui laissent apparaître de nouvelles perspectives de recherche. Parmi elles, la captation du signal par plusieurs nœuds simples ou multi-canaux permet d'envisager une amélioration des performances, par exemple en sélectionnant ou fusionnant les informations déduites de chaque nœud ou canal [46, 93], ou en utilisant des méthodes de spatialisation [3, 209].

Néanmoins, l'implémentation sur systèmes embarqués impose plusieurs contraintes. En effet, de nombreux problèmes matériels peuvent survenir : microphones mal étalonnés, différents ou déplacés, variation du contexte d'enregistrement dans le temps [190]... La transmission des informations qu'oblige la distribution des calculs impose le souci de la sécurité des données [30, 121], de même que celui du coût énergétique [72, 121]. En outre, viser une application dans un environnement précis nécessite souvent une BDD particulière et les annotations fortes correspondantes pour entraîner le système de détection. Malheureusement, on a déjà vu que la création d'une telle BDD nécessite des coûts humains, temporels et financiers importants. Pour cette raison, de nombreuses solutions cherchent à contrer ces problèmes. Par exemple, certains systèmes de classification tentent de détecter les classes à partir de données faiblement annotées [74, 123, 190]. Une autre solution, déjà régulièrement utilisée en détection d'images, est d'adapter un modèle pré-entraîné en lui faisant subir une nouvelle phase d'apprentissage sur nos propres données, disponibles en plus faible quantité [123]. On parle aussi d'adaptation de domaine [205]. Enfin, une autre alternative pour améliorer un modèle obtenu à partir d'une faible quantité de données est d'utiliser l'apprentissage par renforcement [138], qui se développe particulièrement pour le sous-titrage audio [126, 208].

Il ne s'agit pas de la seule contrainte, car les limites physiques du matériel en apportent également. L'une d'entre elles est la latence, spécifique aux applications temps-réel. Il faut alors se soucier de la rapidité d'exécution du code et de la complexité, et en particulier assurer un

traitement par trame plutôt que par piste audio [17, 121, 184]. Une seconde est tout simplement liée à l'implémentation elle-même. En effet, les systèmes embarqués sont généralement limités en énergie. Or une implémentation faible consommation nécessite que les données et les calculs soient codés en virgule fixe. Ces dernières années, de nombreuses recherches se sont ainsi penchées sur la question de la quantification des paramètres des classifieurs [34, 74, 123]. En raison de la limitation de puissance de calcul, les réflexions se sont également portées sur la complexité des tests [30, 54, 63, 93]. Enfin, le stockage des paramètres des classifieurs – et en particulier des réseaux de neurones – est une question prometteuse de l'actualité [24, 74, 123]. Néanmoins, à ce jour, personne ne semble aborder ces problèmes sous le prisme de la transmission d'information. En effet, bien que ce sujet soit déjà partie intégrante du domaine de la reconnaissance de parole distribuée [178, 185], ce n'est pas encore le cas, à notre connaissance, en détection d'événements sonores.

Dans ce chapitre, nous souhaitons nous intéresser aux différents problèmes induits par l'implémentation d'un système de détection d'événements sonores dans un réseau de capteurs sans fil. Dans un premier temps, nous tenons à présenter en détail le principe de ce système, en amenant notamment la problématique de la distribution des calculs. Dans un deuxième temps, l'intérêt de la présence du multi-capteurs est explicitée, et plusieurs solutions pour la valoriser seront présentées. Dans la suite, la limitation de la quantité de données est rappelée. On abordera alors la question de l'augmentation de données, qui pose toujours problème aujourd'hui. La quatrième section est consacrée à la quantification des descripteurs utilisés dans le chapitre 2 et l'étude de l'influence de cette quantification sur les performances de la détection d'événements acoustiques. Enfin, la dernière section de ce chapitre sera l'occasion de se questionner sur l'impact des erreurs de transmission des descripteurs sur la sortie des classifieurs.

5.2 Contexte d'une classification sonore embarquée

Passer du logiciel à l'implémentation matérielle d'un système de classification audio au sein d'un réseau de capteurs sans fil demande un certain nombre d'adaptations. Celles-ci nécessitent une connaissance du matériel, de son fonctionnement et de ses possibilités. En raison de l'énergie et de la puissance de calcul limitées de chaque nœud, il n'est néanmoins pas toujours possible de réaliser toutes les opérations sur le même nœud. En outre, il reste essentiel que le moniteur central puisse connaître les sorties des classifieurs pour chaque signal traité afin de poursuivre son application.

Dans cette section, nous commençons par présenter les différentes manières d'inclure un dispositif de classification au sein d'un système embarqué et d'en distribuer les calculs. Ensuite, nous proposons d'introduire un exemple d'émetteur-récepteur couramment utilisé dans les réseaux de capteurs sans fil et de détailler les possibilités de transmission du signal.

5.2.1 Introduction d'un système distribué multi-nœuds

De nouveaux dispositifs avantageux...

Dans la majorité des dispositifs de détection d'événements sonores, un unique signal est utilisé pour la classification ; tout au mieux dispose-t-il de deux canaux rarement utilisés ensemble, ou alors moyennés [85, 175]. Quelques initiatives proposent de s'intéresser au multi-capteur : c'est le cas de la base SINS [46], une BDD issue de l'enregistrement d'une scène intérieure (ha-

bitation) à l'aide de plusieurs microphones positionnés dans chaque pièce. Ces dernières années, la présence du multi-capteurs a néanmoins été facilitée par plusieurs avancées technologiques. D'une part, la présence de microphones dans de nombreux appareils permet d'enregistrer un signal audio en toute occasion. En raison de la nature différente des microphones dans chaque appareil, il est vrai que cela peut dégrader les performances de la classification ; cependant, on pourrait également imaginer des manières de rendre l'apprentissage plus robuste grâce à cela. D'autre part, la diminution des coûts des microphones a permis la création de plateformes multi-nœuds et multi-capteurs facilement déployables dans un environnement donné. De tels dispositifs sont de grand intérêt pour la classification acoustique. Tout d'abord, certaines nœuds "entendent" mieux que d'autres (par exemple si un capteur est situé proche de la source, tandis qu'au contraire un obstacle est placé entre la source et un deuxième microphone), ce qui permet de choisir le signal le plus pertinent pour la classification [46]. Si un nœud est situé proche d'une source de bruit, il peut également servir à la reconnaître et à procéder au débruitage des signaux issus d'autres nœuds [167]. Enfin, le multi-capteur permet d'apporter un caractère spatialisé à la scène. Il est ainsi possible de séparer les différentes sources présentes dans un même signal [186, 201], voire de localiser chacune d'entre elles [153, 161]. Ce dernier traitement, en complément de la classification, est crucial pour certaines applications comme la détection de chants d'oiseaux dans leur environnement [182] ou l'utilisation d'aides auditives [19].

... Mais qui apportent leur lot de contraintes

Depuis l'arrivée des réseaux de neurones, les classifieurs utilisés sont très performants, mais également très complexes. Cette complexité est problématique pour la mise en place au sein d'un système embarqué, limité en ressources et en énergie. Ainsi, la majorité des recherches considèrent une application dans laquelle le signal est enregistré par un microphone externe puis envoyé à un moniteur central via une communication sans fil, sans cependant considérer une quelconque altération du signal par la transmission. À ce sujet, une première étude sur l'impact de la compression audio sur la classification a été réalisée. Les auteurs de cette étude [159] montrent alors que, sans prise en compte spécifique de la compression et décompression du signal, les performances de la classification vont en réalité fortement décroître. Un tel système est également préoccupant vis-à-vis de la protection des données. En effet, dans certains réseaux de capteurs, il est possible de récupérer le signal émis, signal qui peut contenir des informations privées (activités, identité des locuteurs, conversations...). En outre, transmettre le signal brut nécessite une grande quantité d'énergie car la communication est la partie la plus coûteuse d'un système électronique [72]. Considérant un système multi-capteur, cela nous donne une quantité considérable d'information à transmettre puis à traiter par le moniteur central.

Distribution des calculs

On s'intéresse ainsi à la distribution des calculs au sein du système embarqué, mécanisme qui consiste à déléguer une partie du traitement à chaque nœud de capteurs. Ce problème, déjà connu depuis de nombreuses années dans le domaine des télécommunications – et notamment en reconnaissance de la parole distribuée – [58, 185, 221], est très peu interrogé en classification d'événements sonores. Il a déjà été abordé sous le prisme de la consommation d'énergie d'un tel dispositif [218], mais l'avènement des réseaux de neurones et la priorité donnée à la recherche de meilleures performances de classification l'avait mis de côté. Récemment, des études ont été réalisées pour tenter d'introduire les contraintes de l'embarqué dans des systèmes employant

des réseaux de neurones. Lors des challenges DCASE 2020 et 2021, les participants se sont ainsi penchés sur la question du stockage des paramètres obtenus à l'issue de l'entraînement des réseaux [74, 123]. Par ailleurs, afin de préserver la vie privée des utilisateurs, on sait aujourd'hui extraire les descripteurs sans avoir à enregistrer les signaux bruts [218].

Décision finale au moniteur central

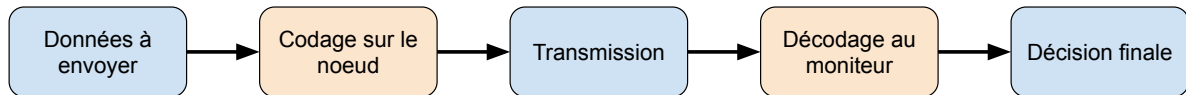


FIGURE 5.1 – Étapes de la transmission des données d'un nœud de capteur vers le moniteur central

Dans le cadre de l'utilisation du multi-capteurs, il est nécessaire que le moniteur central réalise la décision finale qui lui permettra d'exécuter la suite de l'application. La figure 5.1 présente les différentes étapes à réaliser dans le cadre de la communication entre un nœud de capteurs et le moniteur central. Les données à transmettre doivent être encodées en fonction de leur nature pour pouvoir être transmises sur le canal de communication. À la réception par le moniteur central, l'opération inverse de décodage est réalisée. Après un éventuel traitement préalable (classification, fusion d'informations, etc), le moniteur va pouvoir procéder à la décision finale. Il est alors nécessaire de choisir quelles données envoyer. Dans notre cas, on estime que trois éléments peuvent être transmis par les différents nœuds de capteurs : les descripteurs, les probabilités en sortie du classifieur, ou encore la décision réalisée par celui-ci.

Émission des descripteurs L'émission des descripteurs permet de réduire le jeu de données transmis par rapport au signal brut. Elle a l'avantage de requérir peu de capacité de calcul et de stockage sur le nœud, même si un microcontrôleur adapté à la TFR s'avère conseillé pour calculer le spectrogramme [218]. En revanche, les descripteurs peuvent être dégradés par la communication, ce qui risque de diminuer les performances de la classification. En outre, pour des raisons de sécurité, il est nécessaire de faire attention à ce que le signal ne puisse être reconstruit à partir des descripteurs, ou alors correctement protégé. Par exemple, en raison du caractère destructif (*i.e* il y a perte d'information si on effectue cette opération puis sa réciproque) de la TCD, les MFCC ne permettent pas de retrouver le signal original ; en revanche, le mel-spectrogramme ne bénéficie pas de cet avantage.

Émission des probabilités Si on effectue la classification sur le nœud, il est possible de ne transmettre que les probabilités en sortie du classifieur (ou les log-vraisemblances dans le cas du GMM). En fonction du nombre de classes étudiées, le jeu des données émises peut alors être réduit par rapport à celui des descripteurs (par exemple ici, si on transmet 40 coefficients du filtre mel contre 6 probabilités pour toutes les classes, le nombre de données émises est largement réduit). À noter que l'on pourrait également imaginer réduire ces données en n'envoyant par exemple que les trois probabilités les plus hautes et en fixant les autres à 0, ou un mécanisme similaire. Lorsque la classification est effectuée directement sur le nœud, il n'y a pas

de dégradation introduite sur les descripteurs (car ils ne sont pas transmis sur le lien radio), ce qui est un autre avantage. Néanmoins, cela demande de la mémoire et des capacités de calcul supplémentaires pour la classification.

Émission des décisions Enfin, la transmission des décisions binaires de chaque classe permet de réduire plus fortement encore le coût de la communication, et donc la consommation du nœud. Cependant, outre les inconvénients donnés pour la transmission des probabilités, la fusion d'information au moniteur central peut être moins pertinente. Par exemple, en supposant deux nœuds de capteurs de même importance, si le premier transmet une décision à 0 (correspondant à absence de la classe sur la trame) et le second une décision à 1 (présence de la classe), on ne sait quelle décision prendre. En revanche, si le premier nœud émet une probabilité de 0.1 et le deuxième de 0.55 (le seuil de décision étant fixé à 0.5), une solution facile est de moyenner les deux probabilités et conclure de l'absence de la classe sur cette trame.

Dans ces travaux, nous faisons l'hypothèse d'une étude sur la classification multi-nœuds, c'est-à-dire qui nécessite de connaître les informations de tous les nœuds *avant* la classification. On choisit alors de transmettre les descripteurs de nos deux classifieurs décrits dans le chapitre 2. Les conséquences de cette transmission sont étudiées dans les sections 5.5.3 et 5.6.3.

5.2.2 Exemples de systèmes utilisables

Dans cette section, on se propose de présenter un exemple de plateformes matérielles utilisables pour mettre en œuvre notre système, et en particulier transmettre les descripteurs. On peut utiliser une plateforme comprenant un émetteur-récepteur RF TI CC2420 recourant au protocole de communication ZigBee [35], déjà étudié dans le chapitre 4. Il s'agit d'un émetteur-récepteur couramment utilisé dans les réseaux de capteurs sans fil, et dont la vitesse de transmission est de 250 kbits/s.

Dans ce chapitre, nous considérons que le nœud émetteur doit transmettre toutes les 10 ms vers le nœud récepteur (ou le relai éventuel) l'information relative à une trame audio d'analyse. En effet, au sein du nœud émetteur, une nouvelle trame de 40 ms est disponible toutes les 20 ms au niveau du buffer en sortie du CAN. Si on souhaite pouvoir couvrir le cas d'un éventuel relai, il est nécessaire de réaliser deux transmissions de temps égal pendant ces 20 ms : une émission du nœud-source, et une autre du nœud-relai. Finalement, pour maintenir un traitement en temps réel, une transmission de données toutes les 10 ms minimum est nécessaire.

A 250 kbits/s, avec une période de transmission de 10 ms, il est possible de transmettre jusqu'à 2500 bits par trame, soit 312 octets. La fiche technique du transmetteur-récepteur TI CC2420 présente une vue schématique de la trame émise, reproduite sur la figure 5.2. On y remarque que 11 octets sont nécessaires à l'écriture des en-têtes (*header*) et du pied (*footer*), et qu'il est possible d'avoir entre 0 et 20 octets d'adressage (*address information*), ce qui nous amène jusqu'à 31 octets de données utilitaires. Cela laisse alors jusqu'à 281 octets de données utiles (*payload*). Il ne faut cependant pas oublier que les données utiles ne servent pas forcément toutes à transmettre directement les descripteurs : il faut également prendre en compte les bits de sécurité, de redondance, etc. Par ailleurs, en pratique, le nombre n de données utiles ne sera pas forcément fixé au maximum.

On se propose ainsi de présenter deux cas théoriques.

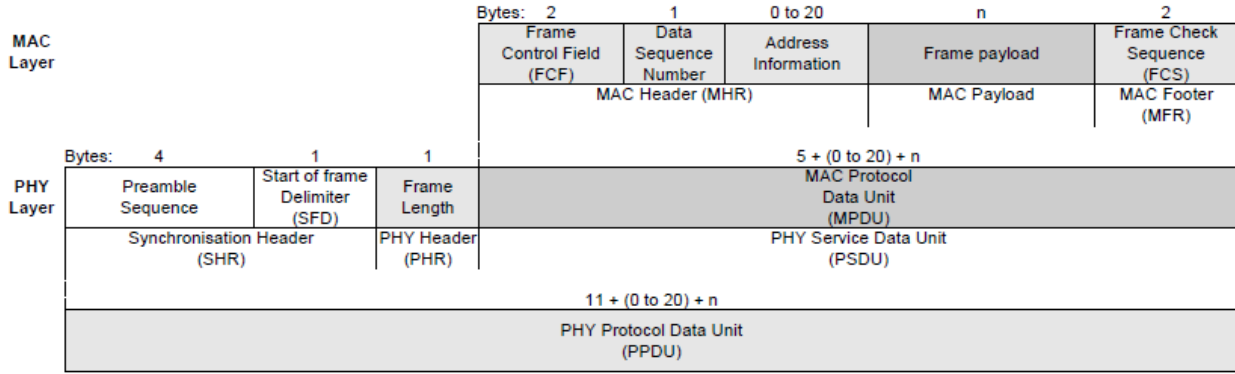


FIGURE 5.2 – Vue schématique du format d'émission d'une trame selon la norme IEEE 802.15.4 [81], issue de la figure 17 de la fiche technique du TI CC2420 [35]

- Dans le premier cas, on fixe à $n = 40$ octets le paquet de données utiles, sans sécurité, et à 4 octets les données d'adressage. En y ajoutant les 11 octets utilitaires, on obtient un total de $L = 55$ octets par paquet émis, ce qui nous donne environ 44 kbits/s. Le protocole ZigBee propose une bande passante jusqu'à 1 MHz. En fixant une largeur de bande $B = 25$ kHz et un temps de fonctionnement T_{on} maximal à 10 ms, la relation $b = \frac{L}{BT_{on}}$ nous donne une efficacité spectrale $b \approx 2$. A l'aide du chapitre précédent, on peut ainsi déterminer dans quel contexte cette configuration est opportune.
- Dans un second cas, on utilise 40 octets de données utiles, mais on y ajoute cette fois-ci 6 octets de sécurité, ce qui nous donne $n = 46$ octets. On emploie également 16 octets d'adressage, ce qui nous donne un total de $L = 73$ octets par paquet émis, soit environ 58 kbits/s. En fixant une bande passante plus large de $B = 32.5$ kHz et toujours $T_{on} = 10$ ms, on obtient encore une fois une efficacité spectrale $b \approx 2$.

Si on diminue T_{on} ou qu'on augmente le nombre n de données utiles, la relation $b = \frac{L}{BT_{on}}$ engendre une augmentation de l'efficacité spectrale b dans les deux cas. Néanmoins, le protocole ZigBee est fait pour des transmissions supérieures à 10 mètres, pour lesquelles l'efficacité spectrale minimisant l'énergie consommée par la communication est, dans la plupart des paramètres, $b^* = 2$ (voir le chapitre 4). Ces deux cas seraient donc idéaux dans un tel dispositif.

Les exemples présentés ici sont très succincts, et mériteraient une étude ultérieure afin de déterminer précisément la manière de coder et de transmettre les descripteurs via un canal de communication sans fil. Cependant, une telle étude se situe en dehors du périmètre de nos travaux.

5.3 Introduction d'éléments de spatialisation dans la classification

La majorité des recherches dans le domaine de la classification d'événement sonores n'exploitent pas la spatialisation acoustique. Pourtant, nous avons abordé dans la section 5.2.2 l'intérêt que pourraient présenter le multi-canal et le multi-nœud, et les possibilités permises par les avancées technologiques de ces dernières années. Il existe en effet des descripteurs multi-canaux, déjà utilisés en séparation et localisation de sources [146]. A cela s'ajoute un récent

regain d'intérêt pour des dispositifs de combinaison de signaux, issus du même nœud ou non [47].

Si, dans ces travaux, nous n'étudierons pas plus amplement le phénomène de spatialisation, nous souhaitons dans cette section présenter une vue d'ensemble de l'état de la recherche dans le domaine. Nous commençons par exposer les différents descripteurs spatiaux présents dans l'état de l'art. Ensuite, nous proposons un aperçu des méthodes de classification utilisant plusieurs nœuds ou canaux.

5.3.1 Descripteurs spatiaux

Des études ont prouvé depuis longtemps que notre cerveau extrait des informations spatiales des sons perçus par nos deux oreilles pour distinguer, localiser et reconnaître des sources audio [117]. La nature de ces informations, appelées indices spatiaux (*spatial cues*), n'a cependant été étudiée que bien plus tard [144].

Différents descripteurs spatiaux

Les descripteurs spatiaux sont pratiquement tous issus de la localisation de source, en particulier binaurale (*i.e.* basée sur l'oreille humaine), et sont donc calculés à partir de deux canaux. Les quatre descripteurs liés à des différences entre canaux sont les suivants :

- la différence de niveau intercanale (ICLD pour *interchannel level difference*) – aussi appelée interaurale (ILD pour *interaural level difference*) – représente la différence d'amplitude en décibel entre deux signaux. En posant X_1 et X_2 les spectres respectifs de deux canaux audio, on calcule alors

$$ICLD = 20 \log \left(\left| \frac{X_1}{X_2} \right| \right); \quad (5.1)$$

- la différence de magnitude intercanale (IMD pour *interaural magnitude difference*), quant à elle, est similaire à l'ICLD mais dans le domaine linéaire [209], c'est-à-dire :

$$IMD = |X_1| - |X_2|. \quad (5.2)$$

Ce dernier descripteur semble apporter des informations bénéfiques pour la classification, comme le démontre l'augmentation des performances dans [209] ;

- la différence de phase intercanale (ICPD pour *interchannel phase difference*) – aussi appelée interaurale (IPD pour *interaural phase difference*) – agit en complément de l'ICLD puisqu'elle représente la différence de phase d'arrivée entre deux signaux :

$$ICPD = \angle \left(\frac{X_1}{X_2} \right) \quad (5.3)$$

avec $\angle(\cdot)$ l'opération donnant la phase.

- enfin, l'IPD est proportionnelle à la différence de temps d'arrivée (ITD pour *interaural time difference*), aussi appelée *time difference of arrival* (TDOA) [136]. Les TDOA sont communément calculées à partir de l'intercorrélacion entre les signaux, qui consiste en un produit pondéré de leurs spectres [142], et que l'on va décrire ci-après.

Récemment, cette intercorrélation a été utilisée dans plusieurs études pour la reconnaissance et la localisation conjointes d'événements en entrée d'un CNN [29]. Tous ces descripteurs nécessitent une synchronisation spatiale pour fonctionner, c'est-à-dire qu'on doit connaître leur position les uns par rapport aux autres. Malheureusement, ceci n'est pas toujours évident à réaliser avec précision. Pour remédier à cela, [82] propose un cepstre spatial dont les valeurs sont calculées pour chaque couple temps-canal plutôt que temps-fréquence.

Expérimentation sur les TDOA

A titre d'exemple, nous proposons d'analyser l'intérêt d'ajouter des TDOA au vecteur des descripteurs des GMM soustractifs présentés dans le chapitre 2. En effet, les couches convolutives qui composent le CRNN également décrit considèrent leur entrée comme une image et donc l'ajout de descripteurs spatiaux, de nature différente, risque d'être néfaste aux performances de classification. Les TDOA sont extraits des deux canaux d'après la méthode issue de [4] et décrite dans l'encadré 5.1.

Cette méthode permet de déterminer un TDOA sur chacune des cinq bandes de fréquence, et donc de potentiellement repérer de multiples TDOA. En effet, si ceux-ci sont différents selon le segment du spectre observé, cela signifie que plusieurs sources occupant différentes parties du spectre sont présentes sur le signal.

Dans ces travaux, on réalise plusieurs expériences correspondant à divers usages des TDOA comme décrits dans [4]. Pour chaque expérience, les TDOA calculées sont concaténées à la suite des 59 descripteurs déjà utilisés.

TDOA simples : Dans cette première expérience, on estime cinq TDOA, soit un coefficient pour chaque bande de fréquence mel, et ceci pour toutes les trames de 40 ms du signal.

TDOA médianes : Dans cette deuxième expérience, on cherche à reproduire partiellement l'étude menée dans [4]. Les TDOA sont alors calculées sur trois fenêtres différentes : 120, 240 et 480 ms au lieu de 40 ms précédemment, mais toujours avec un saut de 20 ms (donc un nombre de trames identique). La variation de longueur de fenêtres est censée permettre d'apprécier plus précisément les TDOA correspondant à des événements eux-mêmes de longueurs très variables. On prend ensuite la médiane de ces trois valeurs pour éviter les valeurs aberrantes et les micro-variations des résultats.

Concaténation des TDOA sur les trois fenêtres : La deuxième proposition de [4] est celle de concaténer les TDOA obtenus sur les trois fenêtres plutôt que d'en retirer la médiane. Ainsi, au lieu d'avoir cinq descripteurs supplémentaires par trame, nous en ajoutons quinze. On prend alors le risque que nos valeurs soient bien plus variables au vu des fenêtres de tailles différentes, mais la finesse des résultats sera alors entièrement retranscrite.

Dérivées des TDOA : Les TDOA permettent d'apporter un élément de spatialisation en donnant la direction d'arrivée d'une source par rapport au nœud de capteurs. Cependant, en fonction de l'environnement, il peut y avoir peu d'événement fixes, ou peu venant d'une même direction pour une même classe. Par exemple dans notre base de données, les enregistrements sont faits en extérieur, et les événements sont principalement liés à des sources mobiles : véhicules ou humains se déplaçant. Ainsi, sur le modèle des MFCC, la vitesse et l'accélération des TDOA pourraient se révéler intéressantes : certaines sources se déplacent à des vitesses ou des accélérations plus ou moins constantes, comme les

Extraction des TDOA

1. On calcule l'interspectre pondéré par la méthode *PHAT* (pour *Phase Transform*) :

$$\Phi(k, t) = \frac{X_1(k, t)X_2^*(k, t)}{|X_1(k, t)||X_2(k, t)|} \quad (5.4)$$

avec $\Phi(k, t)$ l'interspectre pour la trame d'indice t et la fréquence d'indice k , et $.$ l'opérateur conjugué complexe.

2. On pondère successivement cette cohérence par cinq filtres mel différents, soit :

$$C_b(k, t) = H_b(k)\Phi(k, t) \quad (5.5)$$

avec $b \in \llbracket 1, 5 \rrbracket$ l'indice du filtre mel, $C_b(k, t)$ la cohérence pondérée par le b^e filtre mel pour le binôme temps-fréquence (t, k) et $H_b(k)$ la valeur du b^e filtre mel à la fréquence d'indice k .

3. On réalise une TFR inverse sur l'ensemble des n_{fft} fréquences pour chaque filtre, ce qui nous donne l'intercorrélation :

$$R_b(\Delta, t) = \sum_{k=0}^{n_{\text{fft}}-1} C_b(k, t) \exp(i \frac{2\pi k}{n_{\text{fft}}} \Delta) \quad (5.6)$$

avec $\Delta \in [-\tau_{\text{max}}, \tau_{\text{max}}]$ les délais possibles (limités ici par $\tau_{\text{max}} = 30$) et $R_b(\Delta, t)$ l'intercorrélation pondérée par le b^e filtre mel pour la trame t .

4. Finalement, les TDOA sont calculés comme étant l'argument de l'amplitude maximale de cette corrélation, soit

$$\tau(b, t) = \arg \max_{\Delta} (|R_b(\Delta, t)|). \quad (5.7)$$

TABLE 5.1 – Méthode d'extraction des TDOA

piétons qui marchent entre 3 et 5 km/h et les véhicules qui roulent entre 30 et 50 km/h en fonction de l'environnement.

Les performances de la classification à l'issue de chacune des ces expériences sont relevées dans les tables 5.2 et 5.3.

La table 5.2 présente une comparaison des performances entre (i) les 59 descripteurs originaux (les MFCC et leurs dérivées premières et secondes) et celles pour l'expérience (ii) des TDOA simples et (iii) des TDOA avec coefficients dérivés. A chaque fois, le seuil de décision est choisi pour donner la plus haute F-mesure micro-moyennée possible. Les seuils choisis ici sont donc 35 pour (i), 30 pour (ii) et 15 pour (iii). L'ajout des TDOA simples semble apporter peu d'information : en effet, les différences entre le taux d'erreur et la F-mesure pour (i) et (ii) sont très faibles, même si on observe quelques augmentations de F-mesure pour les classes *brakes squeaking* (+1.8%) et *car* (+1.4%). En revanche, une différence notable est observée pour la classe *children*, pour laquelle on assiste à une très forte baisse de la F-mesure, qui passe de 13.3% à 4.4%. Pour les TDOA avec leurs dérivés, les augmentations de F-mesure obtenues avec (i) ne sont plus observables (avec même une baisse de -3.2% pour *brakes squeaking* entre les expériences (i) et (iii)). Toutefois, la baisse significative des performances de *children* est toujours présente. On remarque également une dégradation généralisée du taux d'erreur.

La table 5.3 présente une comparaison des performances entre (i) les 59 descripteurs originaux et celles pour l'expérience (ii) des valeurs médianes des TDOA calculées sur trois fenêtres temporelles et (iii) de la concaténation des TDOA sur ces trois fenêtres. Là encore, on fixe le seuil de décision conduisant à la plus haute F-mesure, soit 35 pour (i), 75 pour (ii) et 55 pour (iii). Dans ces expériences comme dans les précédentes, on observe une dégradation des performances générales. En ajoutant les valeurs médianes des TDOA calculées sur trois fenêtres temporelles, on remarque une forte diminution de la F-mesure par rapport à l'expérience (i) pour les classes *brakes squeaking* (-4.7%), *children* (-7.5%) et *people speaking* (-10.9%). En revanche, la F-mesure gagne 2.4% pour la classe *people walking*. Le taux d'erreur suit approximativement les mêmes dégradations que la F-mesure, c'est-à-dire que, quand celle-ci diminue, le taux d'erreur augmente. De manière surprenante, la concaténation de toutes ces TDOA a un effet différent de la valeur médiane sur les classes. Par exemple, la classe *brakes squeaking* observe une augmentation de sa F-mesure (+10.2% par rapport à (1)), de même que la classe *car* (+5.7%). En revanche, la classe *children* tombe à 0.0% de F-mesure et *large vehicle* passe de 36.9% à 27.5%, en parallèle d'une forte augmentation du taux d'erreur pour les deux classes.

Comme l'avait également étudié [4], il semble que l'effet des TDOA sur la classification soit mitigé. Dans quelques cas de cibles mouvantes suivant un schéma régulier, comme les classes *car* ou *people walking*, il semble se dégager une possibilité d'amélioration grâce à ces descripteurs spatiaux. En revanche, la direction d'arrivée du signal est clairement néfaste pour des événements qui peuvent avoir lieu en tout point de la scène acoustique, comme ici la classe *children* (les effets sont d'autant plus importants que celle-ci est la classe minoritaire). L'intégration d'éléments de spatialisation au sein de la classification apparaît ainsi comme un procédé beaucoup plus complexe que le simple ajout de descripteurs spatiaux.

5.3.2 Utilisation des différents canaux et nœuds au sein du système

Plutôt que d'extraire des descripteurs à partir de deux canaux, il est possible d'adapter l'architecture du système de classification pour prendre en compte plusieurs signaux. La fusion des informations issues de ces signaux peut être réalisée à deux niveaux : avant ou après la

	MFCC seuls		TDOA simples		TDOA dérivées	
	ER	F-mesure	ER	F-mesure	ER	F-mesure
Total par micro-moyennage	1.13	45.2%	1.15	45.3%	1.23	43.6%
Total par macro-moyennage	1.92	35.2%	1.98	34.4%	2.13	32.6%
<i>brakes squeaking</i>	1.36	18.9%	1.31	20.7%	1.10	15.7%
<i>car</i>	0.74	60.4%	0.70	61.8%	0.72	60.3%
<i>children</i>	3.49	13.3%	3.82	4.4%	4.31	4.9%
<i>large vehicle</i>	2.63	36.9%	2.73	37.0%	2.85	36.2%
<i>people speaking</i>	2.03	30.1%	2.07	30.8%	2.24	29.8%
<i>people walking</i>	1.26	51.9%	1.27	51.7%	1.56	48.7%

TABLE 5.2 – Comparaison entre les taux d'erreur (ER) et F-mesure pour une classification par GMM soustractifs prenant en entrée (i) uniquement les MFCC et leurs coefficients dérivés, (ii) les TDOA simples concaténés aux descripteurs de (i), et (iii) les TDOA simples et leurs coefficients dérivés concaténés aux descripteurs de (i). Pour chacune des méthodes, les seuils de décision ont été adaptés pour sélectionner celui conduisant à la meilleure F-mesure micro-moyennée (seuils sélectionnés : (i) 35, (ii) 30 et (iii) 15).

	MFCC seuls		TDOA médianes		TDOA concaténées	
	ER	F-mesure	ER	F-mesure	ER	F-mesure
Total par micro-moyennage	1.13	45.2%	1.15	43.9%	1.46	42.9%
Total par macro-moyennage	1.92	35.2%	1.97	31.7%	2.51	32.1%
<i>brakes squeaking</i>	1.36	18.9%	1.24	14.2%	1.27	29.1%
<i>car</i>	0.74	60.4%	0.77	60.7%	0.78	66.1%
<i>children</i>	3.49	13.3%	3.62	5.8%	4.67	0.0%
<i>large vehicle</i>	2.63	36.9%	2.93	36.0%	4.61	27.5%
<i>people speaking</i>	2.03	30.1%	2.13	19.2%	2.36	21.7%
<i>people walking</i>	1.26	51.9%	1.11	54.3%	1.39	48.4%

TABLE 5.3 – Comparaison entre les taux d'erreur (ER) et F-mesure pour une classification par GMM soustractifs prenant en entrée (i) uniquement les MFCC et leurs coefficients dérivés, (ii) les TDOA médians calculés sur trois fenêtres concaténés aux descripteurs de (i), et (iii) les TDOA calculés sur trois fenêtres concaténés aux descripteurs de (i). Pour chacune des méthodes, les seuils de décision ont été adaptés pour sélectionner celui conduisant à la meilleure F-mesure micro-moyennée (seuils sélectionnés : (i) 35, (ii) 75 et (iii) 55).

classification.

L'intégration des canaux avant la classification donne lieu à des opérations de différentes natures sur les signaux issus de chacun. Par exemple, dans le cadre d'un réseau de N capteurs, les auteurs de [92] proposent de sélectionner seulement deux d'entre eux pour la classification. Considérant qu'un canal peu corrélé avec les autres signifie que ce signal est peu pertinent pour la détection (en raison d'un nœud défaillant, d'une source de bruit dominante, etc), ils sélectionnent le couple de canaux les plus fortement corrélés. Une combinaison linéaire sur les signaux est également possible. Dans le cas de deux canaux issus du même nœud, le choix de l'un ou de l'autre apporte des informations similaires ; en revanche, il est envisageable d'en calculer la moyenne ou la différence [226]. Les descripteurs des canaux sélectionnés peuvent ensuite être concaténés avant passage dans un unique classifieur, ou traverser chacun leur propre dispositif. Dans le deuxième cas, une fusion d'information en sortie de la classification est nécessaire.

Cette fusion est par ailleurs au centre des méthodes de la deuxième solution d'intégration des canaux. En effet, dans ce cas, le signal de chaque canal traverse son propre classifieur – que tous les dispositifs soient identiques ou non [116] – qui renvoie une probabilité d'apparition pour chaque classe. Une moyenne de ces probabilités – ou toute autre combinaison linéaire pondérée – est ensuite réalisée [83, 226].

5.4 Amélioration des données dans un système contraint

Les réseaux de neurones, qui sont aujourd'hui majoritaires en classification d'événements sonores, nécessitent un grand nombre de données d'entraînement pour atteindre des performances satisfaisantes. En outre, toutes ces données doivent être annotées fortement pour pouvoir réaliser une détection d'événements, ce qui induit un coût humain, temporel et financier important.

Plusieurs techniques présentées en introduction de ce chapitre peuvent être mises en œuvre pour se soustraire à ces problèmes. Une autre méthode consiste à augmenter artificiellement le nombre de données d'entraînement en faisant subir diverses modifications et combinaisons à celles-ci : il s'agit de l'AD. Pour des signaux monophoniques, il est également possible d'augmenter uniquement le nombre de données des classes minoritaires pour obtenir des données moins déséquilibrées [83]. Cependant, cette technique n'est pas applicable en classification polyphonique car les classes minoritaires sont généralement présentes en simultané avec les classes majoritaires ; à notre connaissance, aucune méthode d'AD ne permet d'équilibrer une BDD polyphonique, sauf à éventuellement séparer les sources au préalable. D'autres techniques sont alors mise en œuvre.

Dans cette section, on s'intéresse ainsi à présenter uniquement les techniques d'AD applicables aux signaux polyphoniques. On montrera ensuite les résultats d'expériences simples sur les données.

5.4.1 Augmentation de données compatibles avec la polyphonie

L'augmentation de données (AD) est un sujet récurrent dès qu'intervient l'utilisation de réseaux de neurones. L'usage de réseaux convolutionnels dont l'entrée est généralement un spectrogramme (que l'échelle soit linéaire, mel, etc) a donc tout naturellement conduit à importer des méthodes issues du traitement d'image. C'est le cas du *cutout* (ou découpage), qui consiste à masquer une partie de l'image d'entrée, ici le spectrogramme [52, 203]. D'autres méthodes s'en inspirent, comme les masquages temporel (*time masking*) et fréquentiel (*frequency*

masking), qui demandent de masquer respectivement toute une bande de temps ou de fréquence du spectrogramme. La récente technique du *mixup* (ou confusion) réalise une combinaison linéaire entre deux entrées du classifieur [219]. Elle a initialement été appliquée à la classification d'images et de parole.

En effet, de manière analogue à ce que nous avons noté pour des éléments de la classification comme l'extraction des descripteurs, plusieurs méthodes d'AD proviennent directement du domaine de la reconnaissance de la parole. C'est le cas de *SpecAugment* [147], aujourd'hui l'une des méthodes d'AD les plus efficaces et les plus populaires en classification de sons environnementaux [74, 123, 203, 220]. Également inspirée du *cutout*, elle combine masquage temporel et fréquentiel, ainsi que du *time warping* (ou distorsion temporelle) – une technique qui consiste à déformer le signal sur l'axe temporel.

Les autres méthodes principalement utilisées sont spécifiques au signal audio. Par exemple, le *pitch shifting* (ou changement de tonalité) a pour but de modifier la fréquence du signal sans altérer sa vitesse [90]. Au contraire, le *time stretching* (ou étirement du temps) vise à créer de nouvelles données en accélérant ou ralentissant le signal ; elle est particulièrement efficace en complément d'autres méthodes [165].

Une manière de transformer le signal peut être d'en enlever certains morceaux, continus ou non : il s'agit du *cropping* (ou recadrage) [52]. Le *block mixing* (ou mélange de blocs), de manière analogue au *mixup*, consiste à découper le signal en plusieurs blocs et à les additionner simplement entre eux [132]. Jouant sur l'intérêt d'entraîner le classifieur sur des données imparfaites, le *noise mixing* (ou mélange de bruit) ajoute du bruit sur le signal pour modifier son RSB. Cependant, en détection de sons environnementaux, certaines classes peuvent avoir des caractéristiques similaires à ce bruit, telle que la classe "climatisation" dans [165], ce qui décroît ses performances.

Enfin, pour les signaux multi-canaux, il est possible de réaliser une inversion des canaux de la piste audio (*channel swapping*) [203].

5.4.2 Expériences sur le recouvrement et la technique du *mixup*

On souhaite réaliser quelques expériences pour tenter de comprendre l'intérêt et les difficultés que pose l'augmentation de données.

Modification du recouvrement

Dans un premier temps, on augmente le nombre de données par une méthode facilement applicable à tout type de classifieur : l'augmentation du recouvrement. Jusqu'à présent, un recouvrement de 50% était appliqué, c'est-à-dire que pour des trames de 40 ms, une nouvelle trame était générée toutes les 20 ms. Ici, on se propose d'évaluer l'intérêt d'un recouvrement de 75% sur les performances du GMM soustractif et du CRNN. Pour une trame de longueur 40 ms, une nouvelle trame est générée toutes les 10 ms. Ceci conduit à doubler la taille de la base d'apprentissage présentée en entrée des classifieurs sans avoir à opérer de transformation sur le signal. On conserve pour les expériences suivantes les paramètres donnés au chapitre 2, en vérifiant bien évidemment au préalable la convergence de nos modèles.

GMM soustractif Une première comparaison des performances pour un recouvrement de 50% (à gauche) et de 75% (à droite) en sortie du GMM soustractif est présentée dans la table

	Recouvrement de 50%		Recouvrement de 75%	
	ER	F-mesure	ER	F-mesure
Total par micro-moyennage	1.13	45.2%	1.32	40.6%
Total par macro-moyennage	1.92	35.2%	2.28	29.1%
<i>brakes squeaking</i>	1.36	18.9%	1.51	14.2%
<i>car</i>	0.74	60.4%	0.73	60.4%
<i>children</i>	3.49	13.3%	3.69	0.0%
<i>large vehicle</i>	2.63	36.9%	3.65	32.4%
<i>people speaking</i>	2.03	30.1%	2.26	21.8%
<i>people walking</i>	1.26	51.9%	1.82	45.5%

TABLE 5.4 – Comparaison entre les taux d'erreur (ER) et F-mesure pour une classification par GMM soustractif pour un recouvrement entre les trames de 50% (à gauche) ou de 75% (à droite) avec un seuil de décision de 35

	Recouvrement de 50%		Recouvrement de 75%	
	ER	F-mesure	ER	F-mesure
Total par micro-moyennage	1.07	39.1%	1.10	38.0%
Total par macro-moyennage	1.67	25.2%	1.78	24.5%
<i>brakes squeaking</i>	1.0	0.0%	1.0	0.0%
<i>car</i>	0.84	51.5%	0.86	51.9%
<i>children</i>	1.97	0.7%	2.64	0.0%
<i>large vehicle</i>	2.87	30.1%	2.87	32.0%
<i>people speaking</i>	1.92	23.0%	1.98	18.0%
<i>people walking</i>	1.43	45.8%	1.3	45.3%

TABLE 5.5 – Comparaison entre les taux d'erreur (ER) et F-mesure pour une classification par CRNN pour un recouvrement entre les trames de 50% (à gauche) ou de 75% (à droite) avec un seuil de décision de 0.5

5.4. On assiste à une dégradation globale des performances, que ce soit pour la F-mesure ou pour le taux d'erreur. On observe cependant une différence dans l'évolution des performances entre les classes. Par exemple, la classe majoritaire *car* n'a pas été impactée par le changement de recouvrement, avec une F-mesure identique et un taux d'erreur diminué de seulement 0.01. En revanche, la classe minoritaire *children* a été fortement influencée par ce changement, avec une F-mesure qui décroît de 13.3% à 0.0% (mais un taux d'erreur qui n'augmente que de 0.20). Il apparaît ainsi que le recouvrement de 75% n'est pas à privilégier pour le GMM soustractif.

CRNN Une seconde comparaison est réalisée pour le CRNN sur la table 5.5 entre un recouvrement de 50% (à gauche) et de 75% (à droite). Là encore, les performances semblent être plutôt à la baisse. En effet, la plupart des classes ont des performances similaires (*brakes squeaking*, *car* et *people walking*), mais d'autres connaissent un certain changement. C'est le cas pour la classe *children*, qui voit son taux d'erreur augmenter de 1.97 à 2.64 (sa F-mesure restant aux alentours de 0.0%), et *people speaking*, qui connaît une diminution de 5.0% de sa F-mesure. Seule la classe *large vehicle* connaît une légère augmentation de sa F-mesure (+1.9%). En raison de la variabilité des résultats d'un entraînement à l'autre, il est finalement possible que

Mixup

1. Comme pour le dispositif sans AD, on découpe l'ensemble des signaux d'entraînement en blocs de même dimension, ce qui constitue un premier vecteur Y_1 de blocs.
2. On mélange aléatoirement l'ordre des blocs du vecteur Y_1 pour former un second vecteur de blocs Y_2 .
3. Pour chaque bloc du vecteur Y_1 , on génère un réel λ qui suit une loi $\text{Beta}(\alpha, \alpha)$ avec $\alpha \in \mathbb{R}^+$. Ici, on pose $\alpha = 1.5$ suivant la valeur donnée par [204] pour une classification audio.
4. On réalise la combinaison linéaire $y = \lambda y_1 + (1 - \lambda)y_2$ avec y_1 le bloc du vecteur Y_1 et y_2 le bloc du vecteur Y_2 d'indice correspondant.
5. On ajoute ce bloc y aux données d'origine.
6. On répète cette opération pour tous les blocs du vecteur Y_1 .

TABLE 5.6 – Description de la technique du *mixup*

les performances pour ces deux recouvrements soient en réalité similaires. Toutefois, compte tenu de la forte augmentation de la complexité d'un recouvrement à 75% par rapport à 50% (on double le nombre de trames à classifier), le recouvrement de 75% n'apparaît pas comme pertinent à adopter.

Méthode du *mixup*

Dans un second temps, on se propose de vérifier les performances d'une méthode d'AD relativement populaire au début de nos travaux : le *mixup* [55, 164, 199, 204]. Décrite pour la première fois par [219], elle consiste à réaliser une combinaison linéaire entre deux entrées de même dimension. Cette technique est donc davantage adaptée au CRNN, qui prend en entrée des blocs de taille constante, tandis que le GMM traite le signal trame par trame.

On choisit de multiplier la taille de notre BDD par dix. Pour augmenter les données d'entrée du CRNN, on procède selon la méthode décrite dans l'encadré 5.6. La boucle décrite est réalisée neuf fois pour multiplier la taille de notre BDD par dix.

Les performances sur la base de test du CRNN pour un entraînement réalisé avec la BDD d'origine (à gauche) ou avec une AD par *mixup* (à droite) sont consignées dans la table 5.7. À noter que la convergence pour l'entraînement avec des données augmentées a lieu à partir de 230 époques d'entraînement et pour un taux d'apprentissage (qui détermine la vitesse de convergence, une valeur proche de 1 indiquant une plus grande vitesse) de 0.1, alors que l'apprentissage sans *mixup* a été réalisé en 100 époques et pour un taux d'apprentissage de 10^{-4} . En premier lieu, on remarque que la F-mesure par micro-moyennage a diminué de 3.4%, alors que celle par macro-moyennage a augmenté de 0.6%. Le taux d'erreur, lui, a réduit dans les deux cas. Ces variations sont expliquées par les changements présents pour chaque classe. En effet, on observe pour la majorité des classes une forte augmentation des performances. Par

	Sans <i>mixup</i>		Avec <i>mixup</i>	
	ER	F-mesure	ER	F-mesure
Total par micro-moyennage	1.07	39.1%	0.96	35.7%
Total par macro-moyennage	1.67	25.2%	1.41	25.8%
<i>brakes squeaking</i>	1.0	0.0%	1.0	0.0%
<i>car</i>	0.84	51.5%	0.92	23.3%
<i>children</i>	1.97	0.7%	1.36	6.3%
<i>large vehicle</i>	2.87	30.1%	2.56	35.0%
<i>people speaking</i>	1.92	23.0%	1.56	36.2%
<i>people walking</i>	1.43	45.8%	1.08	54.0%

TABLE 5.7 – Comparaison entre les taux d'erreur (ER) et F-mesure pour une classification par CRNN dont l'entraînement est réalisé avec la BDD d'origine (à gauche) ou avec une AD par technique du *mixup* (à droite), pour un seuil de décision de 0.5

exemple, la classe *children*, qui n'était pratiquement pas détectée, atteint maintenant 6.3% de F-mesure. Les progressions les plus importantes sont observées pour *people speaking* (+13.2% de F-mesure) et *people walking* (+8.2%). La croissance de ces F-mesure est corrélée avec une baisse du taux d'erreur de chaque classe. La classe *brakes squeaking* – deuxième classe minoritaire – n'observe quant à elle aucun changement, en n'étant jamais détectée. Seule la classe *car* connaît une importante baisse de performance, avec une F-mesure qui diminue de 51.5% à 23.3%, soit -28.2%. En effet, il s'agit de la classe majoritaire, présente sur une part très importante de la base d'entraînement. La technique du *mixup*, en combinant plusieurs échantillons, a probablement conduit à additionner ensemble plusieurs échantillons sur lesquels l'événement *car* a lieu, ce qui pourrait perturber la reconnaissance (l'événement *car* ayant lieu plusieurs fois). En outre, il est possible que la reconnaissance de *car* ait été faite aux dépens de celle des autres classes, et que l'AD ait diminué ce phénomène. Finalement, si les performances totales ne sont pas convaincantes quant aux avantages du *mixup*, l'étude de celles des classes nous montre bien l'intérêt que l'AD peut avoir pour la détection d'événements sonores. Des conclusions similaires avaient déjà été portées par les auteurs de [165], qui décrivent un comportement différent des classes analysées face à plusieurs méthodes d'AD.

5.5 Encodage des descripteurs en virgule fixe

Dans cette section, on s'intéresse pour la première fois aux problèmes que posent une implémentation sur systèmes embarqués. En effet, alors que la majorité des études considèrent des variables codées en virgule flottante sur 64 bits, ce n'est pas forcément le cas sur *field-programmable gate array* (FPGA). Pour exécuter un programme efficacement et avoir une complexité moindre, on choisit généralement de travailler en virgule fixe et sur un nombre de bits moins important. En outre, la transmission de descripteurs du nœud de capteurs à un moniteur central réalisant la classification nécessite également un codage d'information. En conséquence, les descripteurs de test précédemment extraits devront être encodés avant leur utilisation dans le reste de la chaîne de traitement.

Pour cela, on réalise une quantification, qui est un procédé permettant d'approcher les valeurs continues d'une variable par une valeur discrète prise parmi un sous-ensemble des valeurs

d'origine. La qualité de la quantification dépend directement du nombre de bits utilisés pour celle-ci. En effet, pour un nombre de bits n , il est possible de coder 2^n valeurs différentes. Plus la longueur de quantification n sera grande, plus le nombre de valeurs possibles sera important, approchant le codage de la valeur réelle.

Dans un premier temps, on va présenter l'algorithme LBG, qui permet de réaliser la quantification d'une variable multi-dimensionnelle. On présentera ensuite les différentes méthodes assurant le passage d'un vecteur de descripteurs en précision infinie à un vecteur quantifié. Enfin, on étudiera l'influence de cette quantification sur la détection d'événement sonores.

5.5.1 L'algorithme LBG

L'algorithme LBG est un processus de quantification vectorielle initialement conçu pour la compression des descripteurs lors d'une transmission sans fil de la parole [115]. Devenu très populaire pour la compression de données, son usage s'est progressivement étendu à celui de la classification acoustique [151, 167].

Son principe est dérivé de celui de l'algorithme des K-moyennes incluant un processus itératif sur le nombre de centroïdes, c'est-à-dire suivant l'algorithme suivant :

- (0) À l'origine, on initialise deux centroïdes à partir des données (dans notre cas, on choisit aléatoirement deux individus parmi les entrées).
- (1) On réalise un regroupement des individus d'entrée autour des deux centroïdes selon l'algorithme des K-moyennes.
- (2) Quand les centroïdes ont convergé vers leur optimum, on double leur nombre en créant un nouveau centroïde légèrement décalé pour chaque centroïde d'origine.

On réalise les étapes (1) et (2) n fois jusqu'à l'obtention des 2^n centroïdes souhaités. A noter que l'étape (2) n'est pas exécutée lors de la dernière itération.

Pour classer un nouvel individu, il suffit finalement de l'associer au centroïde le plus proche selon une norme euclidienne.

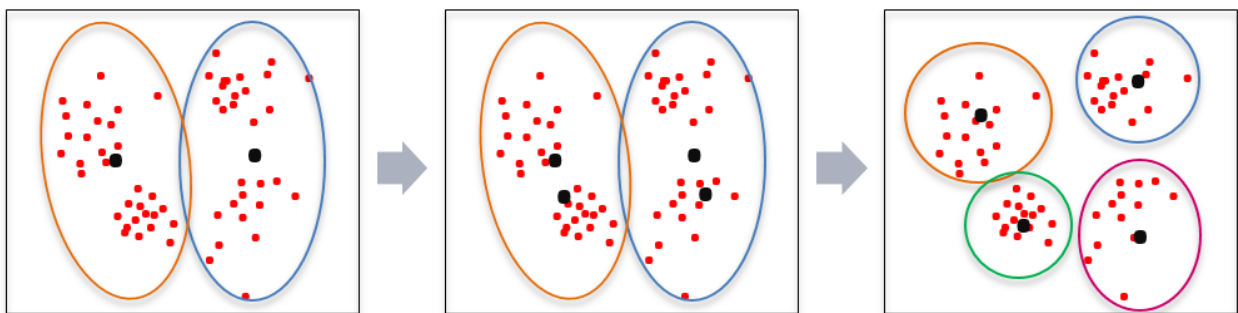


FIGURE 5.3 – Principe de l'algorithme LBG. Dans un premier temps, on réalise un algorithme de regroupement autour de deux centroïdes (en noir), ce qui sépare les individus en deux classes (entourées en bleu et orange). Dans un second temps, on double le nombre de centroïdes, ce qui nous amène à séparer nos individus en quatre nouveaux groupes.

5.5.2 Quantification des descripteurs

Un vecteur de descripteurs peut être quantifié selon deux méthodes : la quantification vectorielle, ou la quantification vectorielle fractionnée (*split vector quantization*). La première technique consiste à quantifier l'intégralité du vecteur des descripteurs sur un certain nombre de bits. Au lieu de quantifier chaque coefficient indépendamment, l'association du vecteur à ses valeurs quantifiées est multi-dimensionnelle. Bien que très simple à mettre en place, la détermination des valeurs discrètes peut néanmoins être très coûteuse. Par exemple, l'algorithme LBG permet d'associer un vecteur de descripteur au centroïde le plus proche, qui correspond à sa valeur quantifiée. Pour un codage sur n bits, il est nécessaire d'entraîner 2^n centroïdes. Or, chaque étape de l'algorithme LBG consiste en une K-moyenne, dont la complexité calculatoire croît linéairement avec le nombre de centroïdes dans la majorité des implémentations (dites de Lloyd ou d'Elkan) [70]. Le nombre d'opérations pour le calcul des centroïdes augmente donc exponentiellement avec la longueur de quantification n .

Pour contourner ce problème, les domaines des télécommunications et de la reconnaissance de la parole distribuée privilégient la quantification vectorielle fractionnée [58, 185]. Celle-ci consiste à diviser le vecteur des descripteurs en plusieurs sous-vecteurs, généralement des paires, qui sont chacun quantifiés indépendamment selon la méthode de quantification vectorielle. En comparaison avec la méthode précédente, le coût calculatoire de l'entraînement est souvent bien plus faible car le nombre total de bits de codage n est partagé entre tous les sous-vecteurs. Prenons l'exemple d'un vecteur de 8 descripteurs à coder sur 32 bits. Dans le cas de la quantification sur l'ensemble du vecteur, la dernière étape de l'algorithme LBG a une complexité proportionnelle à 2^{32} . En revanche, si on réalise une quantification fractionnée par paire de descripteurs et en supposant une répartition uniforme des bits de quantification, chacun des quatre couples de coefficients va être codé sur 8 bits. La complexité calculatoire de la dernière étape de l'algorithme LBG est alors proportionnelle à 2^8 . Même en réalisant cette opération quatre fois, cette grandeur n'atteindra jamais celle de la complexité de la quantification intégrale.

Bien qu'étant la norme en transmission de la parole, l'algorithme LBG n'est pas l'unique méthode de quantification. La plupart des méthodes sont néanmoins également à base de livres-code [58], tels que les cartes de Kohonen [48] ou les sacs de mots [145]. Plus récemment, quelques études s'intéressent à l'apprentissage profond pour tenter d'améliorer la quantification [13].

Dans nos travaux, nous nous intéressons à la quantification vectorielle fractionnée à partir de l'algorithme LBG. En effet, en raison de la complexité calculatoire de la quantification vectorielle, il ne nous sera pas possible de comparer les résultats des deux méthodes.

5.5.3 Résultats expérimentaux

Nous réalisons la quantification des descripteurs présentés en entrée du GMM soustractif et du CRNN selon la méthode détaillée dans les sections 5.5.1 et 5.5.2.

La transmission des MFCC pour la classification d'événements acoustiques a parfois été abordée par le biais des sacs de descripteurs [100, 145]. Ce phénomène est cependant davantage abordé dans le domaine de la reconnaissance de la parole, où la quantification vectorielle fractionnée est généralement mise en œuvre [58, 178]. En revanche, à notre connaissance, aucune étude ne détaille la transmission de la vitesse et de l'accélération des MFCC. La transmission étant l'élément le plus coûteux d'un système embarqué, il convient de transmettre le minimum d'information sur le canal de communication. Or les dérivées pouvant être calculées à partir des MFCC eux-mêmes, il apparaît logique que l'on ne les transmette pas. On propose ainsi de

quantifier uniquement les 20 premiers MFCC. On dérive ensuite les données reçues au moniteur central pour obtenir la vitesse et l'accélération, puis le vecteur de 59 descripteurs est reconstitué en supprimant le premier MFCC. On économise ainsi la transmission de 39 descripteurs, soit près des deux tiers du vecteur.

Pour le CRNN, une telle économie n'est pas possible et il nous faut alors transmettre les 40 coefficients du spectre mel, soit 20 paires de descripteurs. À notre connaissance, une telle quantification n'a encore jamais été réalisée.

Dans un premier temps, on suppose un entraînement des deux classifieurs réalisé à partir des descripteurs codés en virgule flottante sur 64 bits chacun, comme dans le chapitre 2. Seule la phase de test est alors exécutée avec des descripteurs quantifiés. On propose d'observer les performances pour des paires de descripteurs de test codées en virgule fixe successivement sur 1, 2, 4, 6 et 8 bits, et de les comparer aux performances obtenues pour des descripteurs codés en virgule flottante sur 64 bits. Une étude analogue des performances quand l'entraînement est réalisé sur données quantifiées sera effectuée par la suite. Parmi les différents systèmes dont on présente les résultats dans les tables 5.8, 5.9, 5.10, 5.11, 5.12 et 5.13, le taux d'erreur minimal et la F-mesure maximale pour chaque ligne sont mis en valeur en gras.

Apprentissage virgule flottante sur 64 bits

GMM soustractif avec une matrice de covariance pleine On commence par relever les performances pour le GMM soustractif paramétré à l'aide d'une matrice de covariance pleine sur la table 5.8. Le seuil de détection est toujours fixé à 35. On observe une augmentation globale de la F-mesure quand croît la longueur de quantification, ce qui semble logique puisque les valeurs des descripteurs quantifiés approchent davantage des valeurs d'origine (en virgule flottante sur 64 bits chacun). Toutefois, deux classes échappent à cette règle : *large vehicle* et *people speaking*. En effet, pour ces classes, la F-mesure croît jusqu'à une certaine longueur de quantification (2 bits pour *large vehicle*, 6 bits pour *people speaking*) puis décroît pour se stabiliser. Finalement, les valeurs du taux d'erreur et de la F-mesure pour une quantification fractionnée par paires de descripteurs sur 8 bits sont très proches des performances d'origine, voire sont légèrement supérieures pour la plupart d'entre elles : il n'est donc pas nécessaire de quantifier les données sur davantage de bits. Quant au taux d'erreur, il n'évolue pas de manière monotone quand on augmente le nombre de bits de quantification. En effet, notre seuil a initialement été choisi pour maximiser la F-mesure en virgule flottante, des variations peuvent donc apparaître. En outre, une faible F-mesure est associée à un faible taux de détection, qui conduit bien souvent à un faible nombre de vrais positifs, mais également de faux positifs.

GMM soustractif avec une matrice de covariance diagonale Dans l'optique de réduire la consommation énergétique du système, il est cohérent de s'intéresser tant au codage des descripteurs qu'à la complexité de la classification. À ce titre, nous choisissons d'étudier les conséquences de la quantification sur les performances du GMM soustractif paramétré par une matrice de covariance diagonale, qui diminue la complexité calculatoire. Le seuil de détection est celui déterminé dans le chapitre 2, c'est-à-dire -40. Les taux d'erreur et les F-mesures relevées dans la table 5.10 nous présentent un comportement assez différent des évolutions remarquées pour la matrice pleine. En effet, seules deux classes (*brakes squeaking* et *car*) affichent une F-mesure qui croît avec l'augmentation du nombre de bits de quantification. Pour les autres classes, la F-mesure tend plutôt à évoluer de manière non monotone, c'est-à-dire à croître puis

décroître. En outre, la F-mesure atteint à de nombreuses reprises sa valeur maximale pour une quantification des paires de descripteurs sur 1 bit, tandis que le taux d'erreur est minimal pour cette même longueur de quantification sur toutes les classes exceptée la classe majoritaire *car*. Ceci est probablement la conséquence d'un seuil de décision trop faible (il est en effet de -40, c'est-à-dire bien en dessous du seuil théorique de 0). En comparant avec les performances pour une matrice pleine, on remarque également que, outre une complexité plus faible, il semble que l'entraînement à partir d'une matrice diagonale apporte une robustesse vis-à-vis de la quantification des descripteurs. En effet, on note que les variations dans la F-mesure sont bien moindre pour la matrice diagonale (avec une variation maximale de 11.6% entre le codage sur 1 bit et sur 8 bits pour la classe *brakes squeaking*) que pour la matrice pleine (la différence minimale est de 7.8% entre les codages sur 1 bit et 2 bits pour *large vehicle*, et peut monter jusqu'à 23.2% pour *car*, alors que les performances des deux tableaux ne sont pas si éloignées). Finalement, on observe comme pour la matrice pleine que, quand la longueur de quantification des paires est de 8 bits, alors les performances sont similaires à celles d'origine.

CRNN Enfin, on s'intéresse à l'impact de la quantification des coefficients du spectre mel sur les performances de classification du CRNN. Les taux d'erreur et les F-mesures à l'issue des différentes expériences sont recensés dans la table 5.12. Une première chose à noter est que l'évolution des résultats en fonction de la longueur de quantification est moins claire pour le CRNN que pour le GMM. En effet, la F-mesure maximale et le taux d'erreur minimal de chaque ligne n'est pas systématiquement obtenu pour une même expérience : pour la classe *large vehicle*, la F-mesure maximale est atteinte pour les descripteurs d'origine ; pour *car*, avec un codage sur 1 bit ; pour *people walking*, avec un codage sur 8 bits, etc. Néanmoins, ce phénomène est à relativiser car en réalité, toutes les classes ne sont pas fortement impactées par la quantification : les classes *car* (majoritaire) et *brakes squeaking* (qui n'est jamais reconnue) connaissent peu de changement quelle que soit la longueur de quantification, alors que la F-mesure de *children* et *people speaking* peut doubler. À noter également que, même en codant les paires de descripteurs sur 8 bits, la F-mesure diffère toujours des performances d'origine de 1 ou 2% environ.

Apprentissage sur données quantifiées

Pour augmenter les performances de classification dont les signaux de test sont bruités, on préconise généralement d'utiliser des données bruitées pour l'entraînement. Or la quantification agit comme un bruit additif sur les données. Dans un second temps, on propose alors de comparer les performances avec un entraînement réalisé sur les données quantifiées.

GMM soustractif avec une matrice de covariance pleine Dans la table 5.9 sont présentées les performances pour un GMM soustractif muni d'une matrice de covariance pleine. Au contraire des résultats attendus, on observe une dégradation générale des performances en comparaison avec celles obtenues pour un classifieur entraîné sur les données non quantifiées ; seule la F-mesure pour une quantification sur 1 bit, qui était très dégradée, a vu ses valeurs augmenter. En outre, si l'évolution de la F-mesure et du taux d'erreur – augmentation ou diminution des valeurs – suit des tendances similaires au cas précédent (*i.e.* avec un apprentissage non spécifique aux données quantifiées), leurs valeurs ne tendent pas toujours vers celles d'origines. Pour un codage des paires de descripteurs sur 8 bits, toutes les classes montrent une différence entre leur taux d'erreur et celui d'origine inférieure à 0.10, mais seules quatre ont

atteint une F-mesure proche de celle d'origine (moins de 2.5% de différence ici). En effet, la classe *people speaking* affiche toujours une différence de 4.1%, mais c'est la classe minoritaire *children* qui est la plus touchée : elle n'est plus détectée correctement et sa F-mesure passe de 13.3% à 0.0%. Finalement, pour ce premier contexte de classification, l'apprentissage spécifique à partir de données quantifiées semble être une mauvaise solution, il vaut donc mieux conserver le classifieur d'origine.

GMM soustractif avec une matrice de covariance diagonale Puisque le comportement des performances pour une matrice de covariance pleine ou diagonale s'est avéré distinct dans l'expérience précédente, nous analysons également l'influence d'un entraînement spécifique pour un GMM soustractif muni d'une matrice de covariance diagonale. Les taux d'erreur et les F-mesures pour les différents contextes de quantification sont relevés dans la table 5.11. Malheureusement, l'apprentissage spécifique s'avère une nouvelle fois être à l'origine d'une dégradation des performances. En effet, si la F-mesure de certaines classes (*people walking*, *brakes squeaking*, *car*), en particulier pour de courtes longueurs de quantification, est positivement impactée par cet apprentissage, la règle générale semble être une dégradation légère. En revanche, le taux d'erreur est lui fortement influencé, notamment pour les classes *brakes squeaking* ou *children*. On remarque toutefois que ces différentes dégradations restent modestes en comparaison avec celles observées pour une matrice de covariance pleine.

CRNN Enfin, on étudie l'influence d'un apprentissage spécifique sur le CRNN en analysant les performances relevées dans la table 5.13. Comme pour le GMM soustractif, on observe que l'apprentissage sur les données quantifiées est favorable dans le cas d'un codage des paires de descripteurs sur 1 bit. En revanche, au-delà, les résultats sont mitigés. En effet, ce contexte est avantageux pour certaines classes qui voient leur F-mesure augmenter – mais leur taux d'erreur croître également – (*large vehicle* et *people walking*), et préjudiciables pour d'autres (*children* et *people speaking*). Pour les classes restantes, la différence de performances reste faible. À noter cependant que, pour un codage sur 8 bits, la F-mesure et le taux d'erreur approchent davantage les performances d'origine que pour un apprentissage sur les données non quantifiées.

Finalement, en prenant une longueur de quantification assez importante, il est possible d'obtenir des performances similaires à celles obtenues pour des descripteurs codés en virgule flottante sur 64 bits. En outre, il n'apparaît pas nécessaire d'adapter l'apprentissage. Nous pouvons également remarquer que le GMM soustractif semble plus robuste à la quantification de ses descripteurs que le CRNN.

	Sans quantif.		Codage 1 bit		Codage 2 bits		Codage 4 bits		Codage 6 bits		Codage 8 bits	
	ER	F-mesure	ER	F-mesure	ER	F-mesure	ER	F-mesure	ER	F-mesure	ER	F-mesure
Total micro-moy.	1.13	45.2%	1.29	28.1%	1.52	39.7%	1.39	41.7%	1.20	44.9%	1.12	45.7%
Total macro-moy.	1.92	35.2%	2.63	21.9%	2.64	30.6%	2.57	33.0%	2.08	34.7%	1.86	35.6%
<i>brakes squeaking</i>	1.36	18.9%	1.02	2.2%	1.08	4.0%	1.64	13.3%	1.44	14.8%	1.35	19.0%
<i>car</i>	0.74	60.4%	0.89	37.2%	0.83	54.7%	0.83	55.4%	0.77	58.6%	0.74	60.3%
<i>children</i>	3.49	13.3%	8.00	5.3%	6.60	9.7%	6.33	8.4%	4.22	11.2%	3.13	13.5%
<i>large vehicle</i>	2.63	36.9%	1.94	35.1%	1.76	42.9%	2.01	41.6%	2.53	37.6%	2.68	36.2%
<i>people speaking</i>	2.03	30.1%	2.14	20.3%	3.18	30.9%	2.70	34.2%	2.04	36.1%	1.96	32.6%
<i>people walking</i>	1.26	51.9%	1.78	31.5%	2.40	41.1%	1.91	45.2%	1.49	50.0%	1.30	52.1%

TABLE 5.8 – Comparaison des taux d'erreur (ER) et F-mesure à l'issue d'un test sur différentes longueurs de quantification (virgule flottante sur 64 bits, et quantification fractionnée par paires en virgule fixe sur 1, 2, 4, 6 ou 8 bits) pour une classification par GMM soustractif paramétré par une **matrice pleine** et dont l'apprentissage a été réalisé pour des descripteurs codés en virgule flottante sur 64 bits. Le seuil de détection est fixé à 35.

	Sans quantif.		Codage 1 bit		Codage 2 bits		Codage 4 bits		Codage 6 bits		Codage 8 bits	
	ER	F-mesure	ER	F-mesure	ER	F-mesure	ER	F-mesure	ER	F-mesure	ER	F-mesure
Total micro-moy.	1.13	45.2%	1.91	30.6%	1.53	33.4%	1.20	42.4%	1.15	42.6%	1.12	43.5%
Total macro-moy.	1.92	35.2%	5.56	25.7%	2.99	27.2%	2.19	31.4%	2.13	32.5%	1.91	31.6%
<i>brakes squeaking</i>	1.36	18.9%	8.53	15.5%	2.23	19.0%	1.47	9.8%	1.39	18.7%	1.26	16.5%
<i>car</i>	0.74	60.4%	0.85	58.0%	0.89	42.9%	0.79	54.7%	0.75	56.6%	0.74	59.0%
<i>children</i>	3.49	13.3%	17.78	4.5%	6.87	1.3%	4.91	0.9%	4.80	0.9%	3.49	0.0%
<i>large vehicle</i>	2.63	36.9%	1.96	16.0%	3.73	32.6%	2.27	40.9%	2.52	37.4%	2.60	36.9%
<i>people speaking</i>	2.03	30.1%	2.21	23.0%	2.33	25.9%	2.29	31.5%	2.09	29.5%	2.09	26.0%
<i>people walking</i>	1.26	51.9%	2.02	36.9%	1.88	41.7%	1.41	50.6%	1.23	51.7%	1.26	51.4%

TABLE 5.9 – Comparaison des taux d'erreur (ER) et F-mesure à l'issue d'un test sur différentes longueurs de quantification (virgule flottante sur 64 bits, et quantification fractionnée par paires en virgule fixe sur 1, 2, 4, 6 ou 8 bits) pour une classification par GMM soustractif paramétré par une **matrice pleine** et dont l'apprentissage a été réalisé sur les données quantifiées correspondantes. Le seuil de détection est fixé à 35.

	Sans quantif.		Codage 1 bit		Codage 2 bits		Codage 4 bits		Codage 6 bits		Codage 8 bits	
	ER	F-mesure	ER	F-mesure	ER	F-mesure	ER	F-mesure	ER	F-mesure	ER	F-mesure
Total micro-moy.	2.46	36.2%	1.81	36.2%	2.77	34.9%	2.93	34.6%	2.71	35.2%	2.55	36.0%
Total macro-moy.	4.89	30.2%	2.90	25.5%	5.02	28.0%	5.28	28.7%	5.03	29.5%	4.86	30.1%
<i>brakes squeaking</i>	6.77	17.5%	4.26	6.9%	8.38	10.2%	8.48	16.4%	7.38	17.7%	6.82	18.5%
<i>car</i>	0.77	64.8%	0.98	59.3%	0.92	62.4%	0.89	63.0%	0.81	63.7%	0.78	64.6%
<i>children</i>	9.96	4.7%	3.40	6.1%	8.02	2.7%	8.44	3.6%	8.87	3.9%	9.11	4.7%
<i>large vehicle</i>	6.11	24.4%	4.72	27.3%	6.03	24.4%	6.56	23.1%	6.47	23.5%	6.33	23.9%
<i>people speaking</i>	3.38	26.7%	2.01	19.3%	4.05	28.1%	4.41	26.9%	3.88	28.2%	3.57	27.4%
<i>people walking</i>	2.37	43.3%	2.02	33.8%	2.70	39.9%	2.92	39.5%	2.79	40.0%	2.57	41.7%

TABLE 5.10 – Comparaison des taux d'erreur (ER) et F-mesure à l'issue d'un test sur différentes longueurs de quantification (virgule flottante sur 64 bits, et quantification fractionnée par paires en virgule fixe sur 1, 2, 4, 6 ou 8 bits) pour une classification par GMM soustractif paramétré par une **matrice diagonale** et dont l'apprentissage a été réalisé pour des descripteurs codés en virgule flottante sur 64 bits. Le seuil de détection est fixé à -40.

	Sans quantif.		Codage 1 bit		Codage 2 bits		Codage 4 bits		Codage 6 bits		Codage 8 bits	
	ER	F-mesure	ER	F-mesure	ER	F-mesure	ER	F-mesure	ER	F-mesure	ER	F-mesure
Total micro-moy.	2.46	36.2%	3.49	31.4%	3.10	29.1%	3.05	32.1%	2.68	34.5%	2.64	35.8%
Total macro-moy.	4.89	30.2%	8.13	28.7%	6.99	27.3%	6.05	27.9%	5.28	29.3%	5.24	29.8%
<i>brakes squeaking</i>	6.77	17.5%	5.84	16.3%	12.07	13.0%	9.85	13.9%	7.16	16.2%	7.16	17.8%
<i>car</i>	0.77	64.8%	0.82	66.6%	0.88	55.6%	0.86	60.7%	0.80	62.6%	0.81	65.5%
<i>children</i>	9.96	4.7%	25.76	3.5%	14.71	3.8%	11.49	5.5%	10.87	4.3%	11.13	4.2%
<i>large vehicle</i>	6.11	24.4%	8.62	18.8%	7.68	20.6%	6.89	22.3%	5.90	25.1%	6.20	24.3%
<i>people speaking</i>	3.38	26.7%	5.38	26.8%	4.80	27.8%	4.28	26.6%	4.38	25.7%	3.54	25.9%
<i>people walking</i>	2.37	43.3%	2.39	40.3%	1.80	43.2%	2.96	38.3%	2.55	42.0%	2.62	41.0%

TABLE 5.11 – Comparaison des taux d'erreur (ER) et F-mesure à l'issue d'un test sur différentes longueurs de quantification (virgule flottante sur 64 bits, et quantification fractionnée par paires en virgule fixe sur 1, 2, 4, 6 ou 8 bits) pour une classification par GMM soustractif paramétré par une **matrice diagonale** et dont l'apprentissage a été réalisé sur les données quantifiées correspondantes. Le seuil de détection est fixé à -40.

	Sans quantif.		Codage 1 bit		Codage 2 bits		Codage 4 bits		Codage 6 bits		Codage 8 bits	
	ER	F-mesure	ER	F-mesure	ER	F-mesure	ER	F-mesure	ER	F-mesure	ER	F-mesure
Total micro-moy.	1.07	39.1%	1.16	34.2%	1.12	38.4%	1.06	38.8%	1.07	38.6%	1.06	38.2%
Total macro-moy.	1.67	25.2%	1.57	19.3%	1.05	23.9%	1.54	24.7%	1.60	24.8%	1.64	24.6%
<i>brakes squeaking</i>	1.0	0.0%	1.0	0.0%	1.0	0.0%	1.0	0.0%	1.0	0.0%	1.0	0.0%
<i>car</i>	0.84	51.5%	1.03	52.8%	0.91	51.5%	0.85	51.6%	0.86	50.5%	0.87	49.9%
<i>children</i>	1.97	0.7%	1.01	1.0%	1.01	2.1%	1.48	2.1%	1.86	1.6%	2.28	2.1%
<i>large vehicle</i>	2.87	30.1%	3.04	20.4%	2.51	21.8%	2.46	26.9%	2.56	27.5%	2.50	27.9%
<i>people speaking</i>	1.92	23.0%	1.56	13.0%	1.86	25.9%	1.86	23.7%	1.89	24.0%	1.83	21.6%
<i>people walking</i>	1.43	45.8%	1.76	26.7%	1.76	41.9%	1.58	43.8%	1.44	45.5%	1.33	46.2%

TABLE 5.12 – Comparaison des taux d'erreur (ER) et F-mesure à l'issue d'un test sur différentes longueurs de quantification (virgule flottante sur 64 bits, et quantification fractionnée par paires en virgule fixe sur 1, 2, 4, 6 ou 8 bits) pour une classification par **CRNN** dont l'apprentissage a été réalisé pour des descripteurs codés en virgule flottante sur 64 bits. Le seuil de détection est fixé à 0.5.

	Sans quantif.		Codage 1 bit		Codage 2 bits		Codage 4 bits		Codage 6 bits		Codage 8 bits	
	ER	F-mesure	ER	F-mesure	ER	F-mesure	ER	F-mesure	ER	F-mesure	ER	F-mesure
Total micro-moy.	1.07	39.1%	0.87	41.9%	1.19	34.0%	1.17	35.4%	1.05	39.9%	1.11	38.1%
Total macro-moy.	1.67	25.2%	1.14	20.8%	1.99	22.3%	1.82	23.2%	1.59	24.9%	1.70	25.0%
<i>brakes squeaking</i>	1.0	0.0%	1.0	0.0%	1.0	0.0%	1.0	0.0%	1.0	0.0%	1.0	0.0%
<i>car</i>	0.84	51.5%	0.86	60.5%	0.85	50.8%	0.83	48.4%	0.85	54.3%	0.83	50.8%
<i>children</i>	1.97	0.7%	1.11	0.0%	3.82	0.0%	2.36	0.6%	1.63	0.0%	2.07	1.9%
<i>large vehicle</i>	2.87	30.1%	1.31	21.7%	2.33	31.7%	2.81	29.8%	2.39	27.2%	2.73	31.6%
<i>people speaking</i>	1.92	23.0%	1.35	9.8%	2.32	16.7%	2.23	19.5%	2.24	19.5%	2.16	19.8%
<i>people walking</i>	1.43	45.8%	1.21	33.1%	1.60	34.7%	1.68	41.0%	1.43	48.4%	1.45	45.7%

TABLE 5.13 – Comparaison des taux d'erreur (ER) et F-mesure à l'issue d'un test sur différentes longueurs de quantification (virgule flottante sur 64 bits, et quantification fractionnée par paires en virgule fixe sur 1, 2, 4, 6 ou 8 bits) pour une classification par **CRNN** dont l'apprentissage a été réalisé sur les données quantifiées correspondantes. Le seuil de détection est fixé à 0.5.

5.6 Influence de la transmission sur la classification d'événements sonores

Dans la section précédente, nous avons étudié l'influence de la quantification des descripteurs sur les performances de la classification pour le GMM soustractif et le CRNN. Ces analyses nous ont permis de conclure que la réalisation de l'apprentissage sur des données dégradées (c'est-à-dire encodées à la longueur de quantification souhaitée pour le test, qui correspond finalement à un ajout de bruit) est la plupart du temps néfaste pour les performances. En outre, nous avons observé que la transmission des données quantifiées permet de retrouver des performances similaires à celles obtenues avec les descripteurs codés en précision infinie quand la longueur de quantification est suffisante. Cependant, ces études supposent une transmission sans erreur sur le canal de communication, ce qui est rarement le cas en pratique.

Dans cette section, la possibilité d'erreurs de décodage en réception sera donc prise en compte. Nous présenterons tout d'abord le principe de transmission des données, et comment les erreurs de décodage peuvent affecter celui-ci. Nous analyserons ensuite l'impact des erreurs de transmission sur la détection d'événements sonores dans différents contextes d'apprentissage. Enfin, pour contrer les dégradations des performances liées aux erreurs de canal, nous proposons une nouvelle méthode d'AD basée sur l'ajout de bruit, et nous analysons son intérêt.

5.6.1 Présentation du système

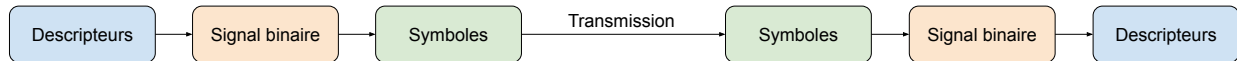


FIGURE 5.4 – Schéma simplifié des différentes formes par lesquelles passent les descripteurs lors de leur transmission sur le canal de communication

Le vecteur des descripteurs constitue un signal numérique que l'on peut transmettre sur un canal de communication sans fil selon les méthodes décrites au chapitre 3 au même titre que tout autre signal.

La chaîne de transmission qui nous intéresse ici est schématisée sur la figure 5.4. Les descripteurs dont les n paires sont quantifiées sur 2^b niveaux de valeurs sont transformés en une suite de $n \times b$ données binaires (ne sont ici pas pris en compte les bits de redondance, de sécurité, etc). Une bijection est ensuite réalisée entre les données binaires et les symboles de l'alphabet de modulation choisi. Les symboles sont transmis sur le canal physique, puis décodés en réception. Un signal binaire est alors reconstruit grâce à la réciproque de la bijection précédente. Finalement, on retrouve les descripteurs quantifiés originaux, puis on continue notre traitement audio.

Dans ces travaux, nous considérons uniquement le paramètre qu'est la PEB en réception. En effet, la plupart des paramètres introduits dans le chapitre 3 influent sur la PEB, qui en constitue finalement une synthèse. Il existe des méthodes pour limiter l'impact de la PEB sur le vecteur des descripteurs. La première consiste à associer les données binaires aux symboles correspondants de manière à minimiser le nombre d'erreurs sur la suite binaire en cas de mauvais

décodage d'un symbole. C'est le principe du codage du Gray, introduit dans le chapitre 3, qui minimise le changement d'entropie entre deux états adjacents, *i.e.* la confusion d'un symbole avec le plus proche symbole de l'alphabet entraîne une unique erreur sur la séquence binaire associée. La même opération peut être réalisée pour l'association entre les paires quantifiées et leur représentation binaire. L'algorithme LBG contribue à réaliser une telle bijection en raison de la création successive de centroïdes proches les uns des autres. Cependant, nous ne détaillerons pas davantage cette association.

Dans les expériences qui suivent, l'ajout de bruit est effectué directement sur le vecteur des données binaires en considérant que le décodage de chaque bit a une probabilité d'erreur qui sera notée P_{eb} .

5.6.2 Résultats expérimentaux

Dans cette partie, nous allons simuler des erreurs de transmission aléatoires sur les données binaires obtenues à partir du vecteur des descripteurs. Le contexte est similaire à celui de la section 5.5.3, c'est-à-dire qu'on relève les performances pour les classifieurs que sont le GMM soustractif (avec une matrice pleine ou diagonale) et le CRNN, et dont l'apprentissage est réalisé sur les données quantifiées ou non. On s'intéresse ici à deux longueurs de quantification des paires de descripteurs : 1 et 8 bits. La PEB est fixée successivement à 10^{-3} et 10^{-5} . Pour simplifier notre étude, on compare uniquement la F-mesure obtenue (et non le taux d'erreur) à celles des classifications sur données quantifiées de la section 5.5.3 et sur données non quantifiées. Enfin, en raison du caractère aléatoire des erreurs sur le canal, les résultats présentés sont les F-mesures moyennes à l'issue de cinq transmissions et classifications pour chacun des paramétrages abordés.

GMM soustractif avec une matrice de covariance pleine Dans un premier temps, on relève les performances de l'expérience décrite au paragraphe précédent pour le GMM soustractif muni d'une matrice de covariance pleine sur la figure 5.5 et pour un seuil de détection de 35. Sur les figures 5.5a (pour des paires des descripteurs de test quantifiés sur 1 bit) et 5.5b (pour des paires quantifiées sur 8 bits), on présente l'histogramme de la F-mesure pour les totaux par micro-moyennage (Micro) et macro-moyennage (Macro), ainsi que pour l'ensemble des classes : *brakes squeaking* (Brakes), *car* (Car), *children* (Child.), *large vehicle* (Large), *people speaking* (Speak.) et *people walking* (Walk.). La comparaison des performances originales sans quantification (première barre, en noir) est réalisée avec celles sur données quantifiées sans entraînement spécifique pour une PEB de 0, 10^{-3} et 10^{-5} (respectivement les deuxième, troisième et quatrième barres en rouge, vert foncé et bleu foncé), ainsi que sur données quantifiées avec entraînement spécifique dans les mêmes conditions (respectivement les cinquième à septième barres en orange, vert clair et bleu clair). Sur la figure 5.5a qui présente les résultats pour des paires de descripteurs codées sur 1 bit, on retrouve des observations déjà réalisées dans la section 5.5.3, c'est-à-dire que les performances sont fortement dégradées par la quantification par rapport aux résultats originaux, dégradation qui peut être compensée pour la plupart des classes par l'apprentissage sur données quantifiées. En analysant les nouveaux résultats que sont les F-mesures moyennes après ajout d'erreur sur le canal, on remarque que ces erreurs n'ont pas d'impact sur la F-mesure (voire peuvent l'augmenter légèrement pour certaines classes), et ce quelle que soit la nature des données d'apprentissage. Ceci peut notamment être expliqué par le fait que, les données binaires étant peu nombreuses pour chaque vecteur de descripteurs, il y

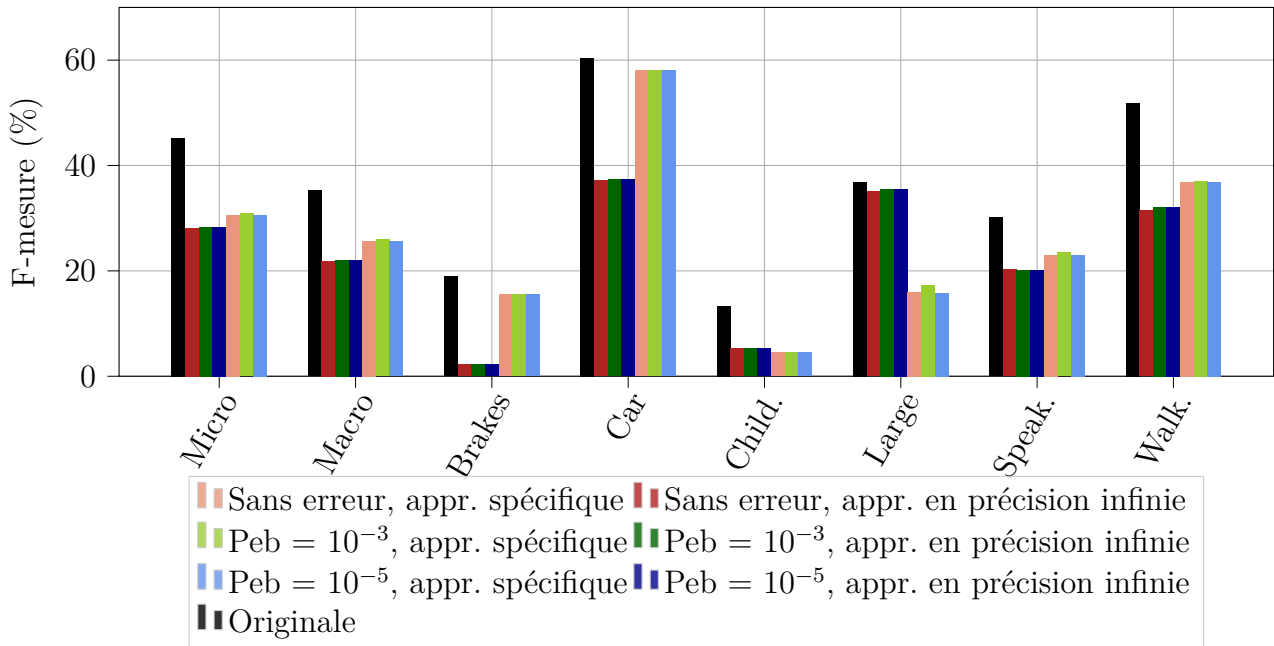
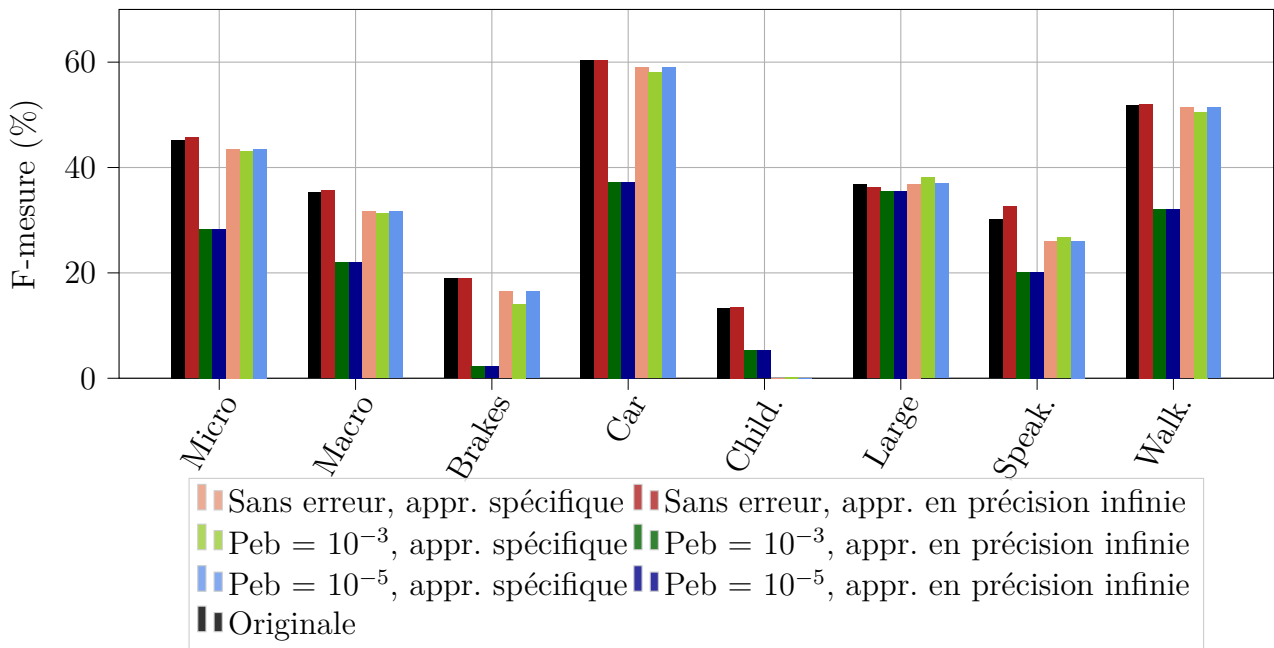
(a) F-mesure pour des paires de descripteurs codées sur $n = 1$ bit(b) F-mesure pour des paires de descripteurs codées sur $n = 8$ bits

FIGURE 5.5 – Évolution de la F-mesure pour un GMM soustractif muni d'une **matrice de covariance pleine** et dont les paires de descripteurs sont codées sur n bits, en fonction de l'apprentissage (sur données quantifiées ou non) et pour différentes PEB en réception (sans erreur, 10^{-3} et 10^{-5}). Les résultats sont présentés pour les totaux par micro-moyennage (Micro) et macro-moyennage (Macro), ainsi que pour l'ensemble des classes : *brakes squeaking* (Brakes), *car* (Car), *children* (Child.), *large vehicle* (Large), *people speaking* (Speak.) et *people walking* (Walk.). La comparaison est réalisée avec les performances originales sans quantification.

a peu d'erreurs quantitativement sur celles-ci. En outre, les performances sont déjà largement dégradées par la quantification. Enfin, cette classification est également réalisée sur les dérivées des MFCC, ce qui réduit probablement l'impact des erreurs, en particulier sur 1 bit, car les dérivées sont réalisées sur une succession de plusieurs trames (dont peu seront impactées par les erreurs). À l'exception de ce dernier argument, ce n'est pas le cas pour des paires de descripteurs codées sur 8 bits, et dont les F-mesures pour toutes les expériences sont relevées sur la figure 5.5b. En effet, à ce niveau de précision dans la quantification, les performances sont pratiquement identiques à celles de la classification sur les données originales. En revanche, quand le canal de transmission est source d'erreurs, la F-mesure pour un apprentissage sur données en précision infinie diminue largement. À noter par ailleurs que ces F-mesures sont identiques pour $P_{eb} = 10^{-3}$ et $P_{eb} = 10^{-5}$. L'apprentissage spécifique révèle alors son intérêt : en effet, on observe très peu de différence entre la F-mesure avec ou sans erreur sur le canal de transmission. On peut déduire que cet apprentissage permet au GMM soustractif avec matrice de covariance pleine d'être plus robuste face à ces erreurs de décodage.

GMM soustractif avec une matrice de covariance diagonale On renouvelle ces mêmes expériences en prenant cette fois-ci en compte une matrice de covariance diagonale pour le GMM soustractif et un seuil de détection de -40. Pour une longueur de quantification de 1 bit pour chaque paire de descripteurs, on retrouve sur la figure 5.6a la même observation que pour une matrice pleine : les erreurs de décodage dégradent peu voire pas du tout la F-mesure, et ce pour les deux types d'entraînement (spécifique ou non). En revanche, des conclusions différentes peuvent être tirées pour un codage sur 8 bits. En effet, sur la figure 5.6b, la dégradation de la F-mesure quand l'apprentissage du classifieur est réalisé sur des données non quantifiées est bien moindre que pour la matrice pleine pour l'ensemble des classes, et on note même une augmentation de la F-mesure pour les classes *children* et *large*. Ceci nous donne une F-mesure totale macro-moyennée plus faible (-4.7% par rapport au canal sans erreur) mais micro-moyennée identique (+0.0%). Toutefois, on remarque là encore que l'apprentissage sur données quantifiées permet de réduire voire annuler l'impact des erreurs de transmission sur la détection de chaque classe.

CRNN Enfin, on s'intéresse aux résultats qui concernent le CRNN. Comme nous l'avons remarqué dans la section 5.5.3, le CRNN réagit de manière différente à la quantification des descripteurs selon les classes. Pour une quantification sur 1 bit, dont les résultats sont présentés sur la figure 5.7a, on rappelle que l'apprentissage spécifique permet d'augmenter la F-mesure y compris par rapport aux données originales en raison de la forte reconnaissance de la classe majoritaire *car*. Ces différences avec le GMM soustractif se retrouvent dans le traitement des erreurs par le CRNN. En effet, si les performances de la classe *car* sont largement dégradées par les erreurs sur le canal, ce n'est pas le cas pour les classes *children*, *large vehicle* et *people walking* qui voient leur F-mesure augmenter avec l'ajout d'erreur pour les deux types d'apprentissage (pour rappel, pour $n = 1$ bit, le GMM n'observe pas ces variations de F-mesure).

Au vu de l'importance de la classe *car* dans les données (majoritaire), cela résulte en un maintien des performances pour la F-mesure totale macro-moyennée pour l'apprentissage sur données en précision infinie, mais une baisse notable de -4.4% à $P_{eb} = 10^{-3}$ pour la F-mesure totale micro-moyennée. Pour un apprentissage spécifique, on observe une très faible baisse des performances pour la F-mesure totale macro-moyennée, mais une diminution plus importante de -8.0% à $P_{eb} = 10^{-3}$ pour la F-mesure totale micro-moyennée. En observant les F-mesures

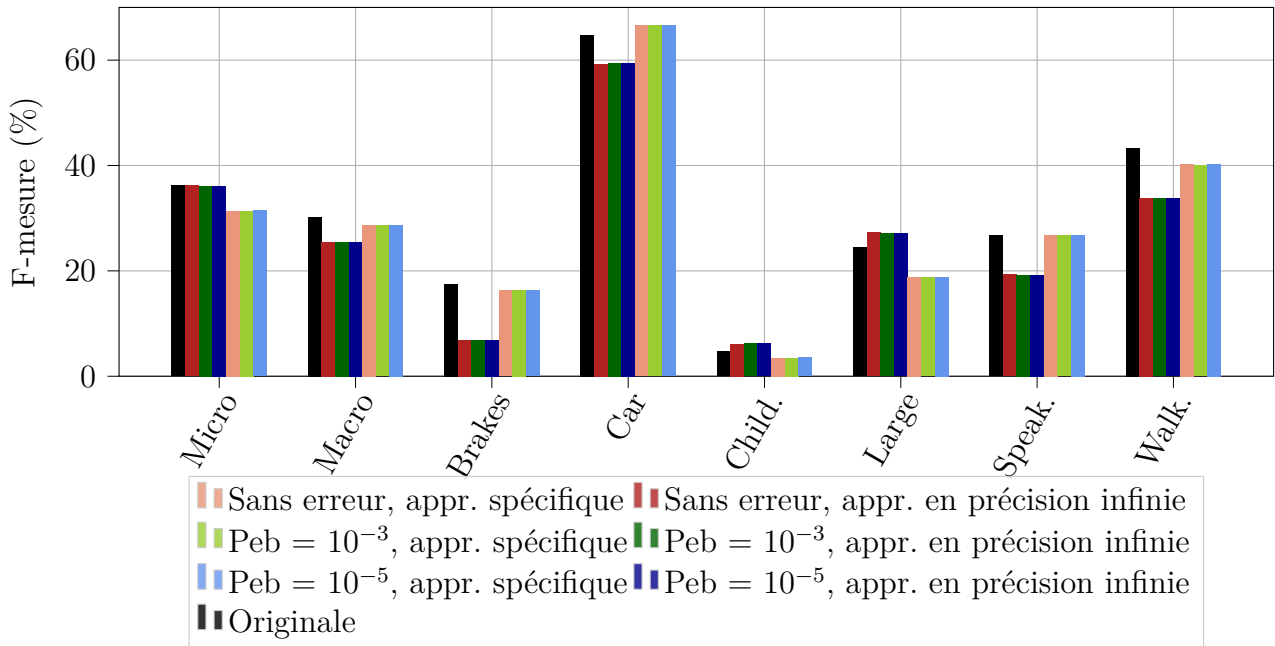
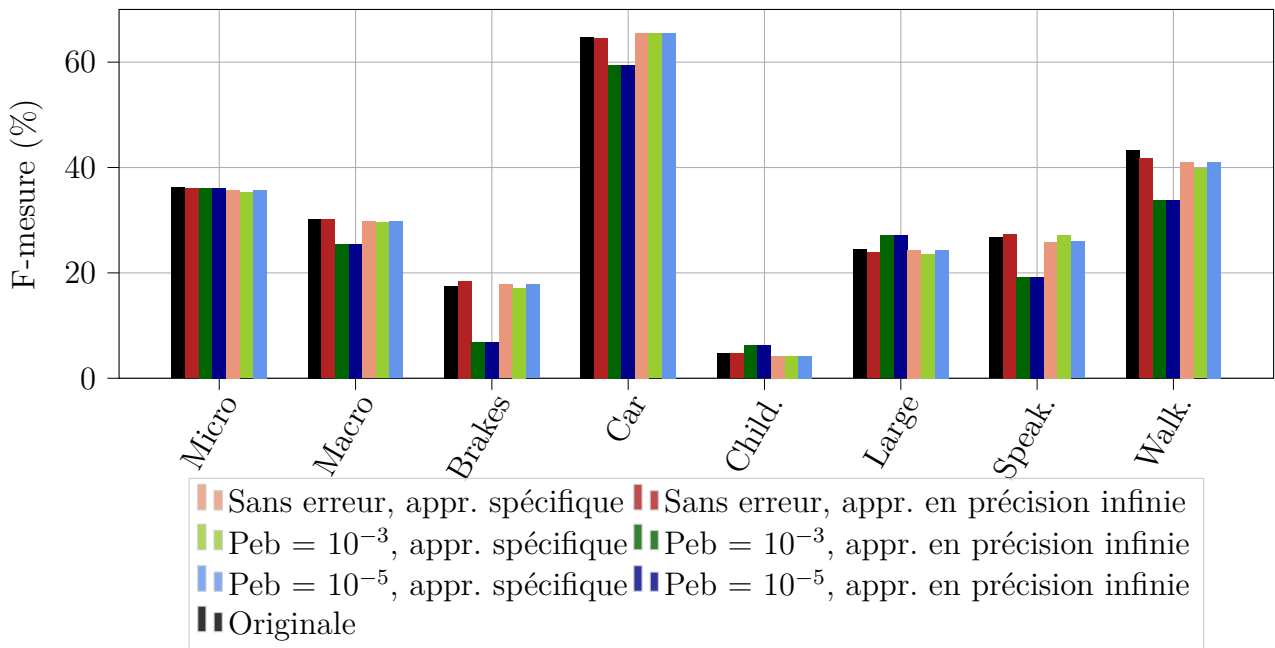
(a) F-mesure pour des paires de descripteurs codées sur $n = 1$ bit(b) F-mesure pour des paires de descripteurs codées sur $n = 8$ bits

FIGURE 5.6 – Évolution de la F-mesure pour un GMM soustractif muni d'une **matrice de covariance diagonale** et dont les paires de descripteurs sont codées sur n bits, en fonction de l'apprentissage (sur données quantifiées ou non) et pour différentes PEB en réception (sans erreur, 10^{-3} et 10^{-5}). Les résultats sont présentés pour les totaux par micro-moyennage (Micro) et macro-moyennage (Macro), ainsi que pour l'ensemble des classes : *brakes squeaking* (Brakes), *car* (Car), *children* (Child.), *large vehicle* (Large), *people speaking* (Speak.) et *people walking* (Walk.). La comparaison est réalisée avec les performances originales sans quantification.

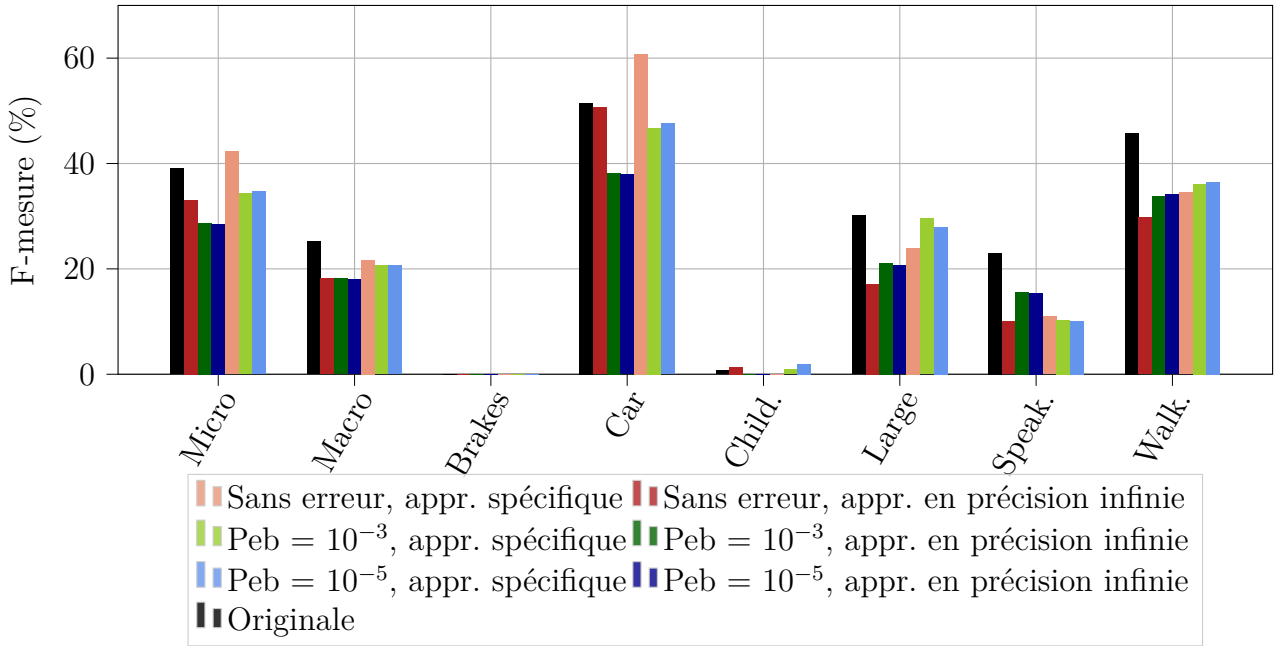
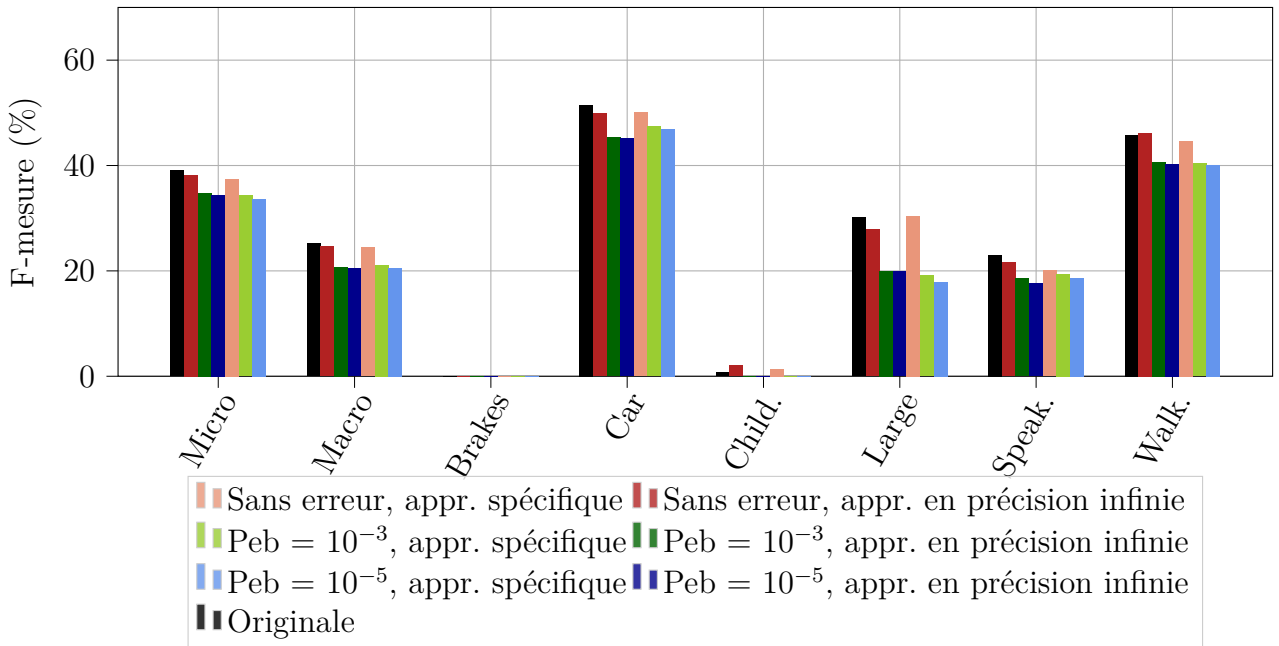

 (a) F-mesure pour des paires de descripteurs codées sur $n = 1$ bit

 (b) F-mesure pour des paires de descripteurs codées sur $n = 8$ bits

FIGURE 5.7 – Évolution de la F-mesure pour un **CRNN** et dont les paires de descripteurs sont codées sur n bits, en fonction de l'apprentissage (sur données quantifiées ou non) et pour différentes PEB en réception (sans erreur, 10^{-3} et 10^{-5}). Les résultats sont présentés pour les totaux par micro-moyennage (Micro) et macro-moyennage (Macro), ainsi que pour l'ensemble des classes : *brakes squeaking* (Brakes), *car* (Car), *children* (Child.), *large vehicle* (Large), *people speaking* (Speak.) et *people walking* (Walk.). La comparaison est réalisée avec les performances originales sans quantification.

pour $n = 8$ bits sur la figure 5.7b, on retrouve, comme pour le GMM soustractif, le fait que les performances à l'issue d'une classification sur données quantifiées sans erreur de décodage s'approchent de celles sur les données en précision infinie. Néanmoins, au contraire du GMM, l'apprentissage sur données quantifiées ne semble pas réellement atténuer l'influence des erreurs de transmission par rapport à celui réalisé en précision infinie. Cela peut s'expliquer par la sensibilité généralement plus grande du CRNN aux différences entre les données d'apprentissage et de test. En outre, le vecteur des descripteurs transmis est celui utilisé directement en entrée du classifieur, au contraire du GMM soustractif dont nous déduisons les dérivées sur plusieurs trames successives, donc chaque erreur de transmission a peut-être davantage d'importance que pour le GMM. Néanmoins, on peut remarquer que l'impact sur la classification reste relativement faible : -3.1% et -3.4% à $P_{eb} = 10^{-3}$ sur la F-mesure micro-moyennée pour un apprentissage effectué sur des données quantifiées ou non, et -3.3% et -3.8% sur la F-mesure macro-moyennée pour conditions d'apprentissage identiques.

Finalement, pour le GMM soustractif, il apparaît que la transmission des vecteurs de descripteur sur le canal et l'ajout d'erreurs qui en découle ne soit pas un obstacle à la détection d'événements sonores si on adapte la manière de réaliser l'apprentissage. En revanche, la F-mesure du CRNN est dégradée par ces erreurs.

5.6.3 Augmentation de données pour la transmission de descripteurs

La différence entre les données d'apprentissage et celles de test constitue une source de dégradation des performances de la classification d'événements sonores. Pour contrer ce problème, il est conseillé d'entraîner le classifieur utilisé directement sur les données bruitées [132].

C'est ce que nous proposons de réaliser pour tenter de prévenir la chute de la F-mesure observée pour le CRNN dans la section précédente. Nous introduisons alors une nouvelle technique d'augmentation de données audio – que l'on nommera *BinError* pour plus de facilité – basée sur l'ajout d'erreurs sur les données binaires obtenues à partir des descripteurs quantifiés. La méthode réalisée suit les étapes de l'encadré 5.14.

Dans notre expérience, nous proposons d'effectuer cette boucle deux fois pour chaque $P_{eb} \in \{10^{-2}, 10^{-3}, 10^{-4}, 10^{-5}\}$, ce qui nous amène à multiplier le nombre de données total par neuf. On entraîne le CRNN sur ces données augmentées, puis on observe ses performances sur les données de test transmises successivement sur un canal paramétré par $P_{eb} = 10^{-3}$ puis $P_{eb} = 10^{-5}$. À noter que le protocole reste le même que dans la section précédente, à savoir $n_{\text{test}} = 5$ entraînements différents pour une validation à $n_{\text{fold}} = 4$ plis. Pour chacun des 5 modèles obtenus, on effectue chacun cinq tests (conduisant à un total de 25 tests pour chaque paramétrage), dont on moyenne finalement les performances.

On fixe le seuil de détection à 0.5, puis on relève la F-mesure obtenue pour une quantification des paires de descripteurs sur 1 bit sur la figure 5.8a et sur 8 bits sur la figure 5.8b. On y présente toujours l'histogramme de la F-mesure pour les totaux par micro-moyennage (Micro) et macro-moyennage (Macro), ainsi que pour l'ensemble des classes : *brakes squeaking* (Brakes), *car* (Car), *children* (Child.), *large vehicle* (Large), *people speaking* (Speak.) et *people walking* (Walk.). La comparaison des performances originales sans quantification (première barre, en noir) et de celles sur données quantifiées avec un entraînement réalisé sur ces mêmes données pour une PEB de 0, 10^{-3} et 10^{-5} (respectivement les deuxième, troisième et quatrième barres

Etapas de *BinError*

1. On quantifie les descripteurs d'une base de données audio.
2. Ces descripteurs sont transformés en une suite de données binaires selon la méthode expliquée à la section 5.6.1.
3. On génère un vecteur X de données binaires aléatoires suivant une loi uniforme de probabilité $P(X = 1) = \text{Peb}$ ($\text{Peb} \in [0, 1]$ une valeur constante) et de même longueur que nos données utiles. Chaque présence de la valeur unité représente une erreur à venir sur les données binaires à l'indice donné.
4. On additionne modulo 2 les deux vecteurs binaires. On obtient alors un unique vecteur des données utiles dont certains bits sont erronés en raison du vecteur X .
5. On décode le vecteur ainsi obtenu pour retrouver celui des descripteurs sous forme numérique (et quantifiée).
6. On traite les descripteurs selon la technique en usage (ici séquençage en lots) puis on concatène ces nouveaux descripteurs aux originaux.

Cette boucle peut être réalisée plusieurs fois et pour plusieurs PEB de valeur Peb afin d'obtenir des données différentes.

TABLE 5.14 – Description de la méthode d'AD proposée

en orange, vert clair et bleu clair), qui proviennent des expériences de la section 5.6.2, est réalisée avec celles obtenues à la suite d'une augmentation de données *BinError* pour PEB de 10^{-3} puis 10^{-5} (respectivement les cinquième et sixième barres en jaune et rouge). Pour 1 bit comme 8 bits, on remarque sur les figures 5.8a et 5.8b que les F-mesures totales micro et macro-moyennée pour $\text{Peb} \in \{10^{-3}, 10^{-5}\}$ ont rejoint celle de la transmission sans erreur quand *BinError* est employée lors de l'apprentissage. Cette amélioration est plus particulièrement due à l'augmentation drastique de la F-mesure pour la classe *car*, qui avait fortement décru avec les erreurs de transmission. Sur la figure 5.8a, il s'agit en effet de la seule classe à avoir réellement subi des modifications. En revanche, pour 8 bits de quantification sur la figure 5.8b, on remarque une évolution positive pour toutes les classes à l'exception de *people speaking*, dont les performances décroissent encore. On note d'ailleurs que, pour quelques classes, cette augmentation de données apparaît comme avantageuse même par rapport aux tests originaux. En effet, la classe *brakes squeaking* est enfin reconnue (même faiblement, à hauteur d'une F-mesure de 0.56% pour les deux PEB), la classe *car* gagne par exemple +1.3% et *large vehicle* +2.7% quand $\text{Peb} = 10^{-3}$ par rapport à leur F-mesure originales respectives.

Finalement, la méthode d'augmentation de données présentée semble être efficace dans la prévention de la dégradation des performances suite à la transmission des descripteurs sur un canal de communication sans fil.

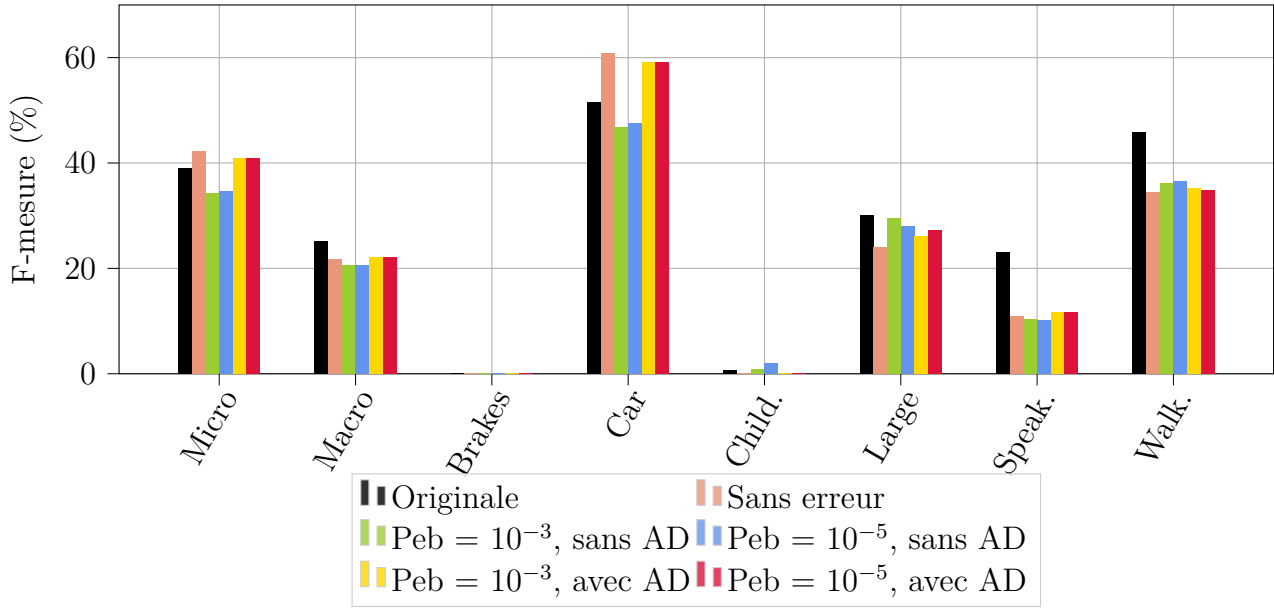
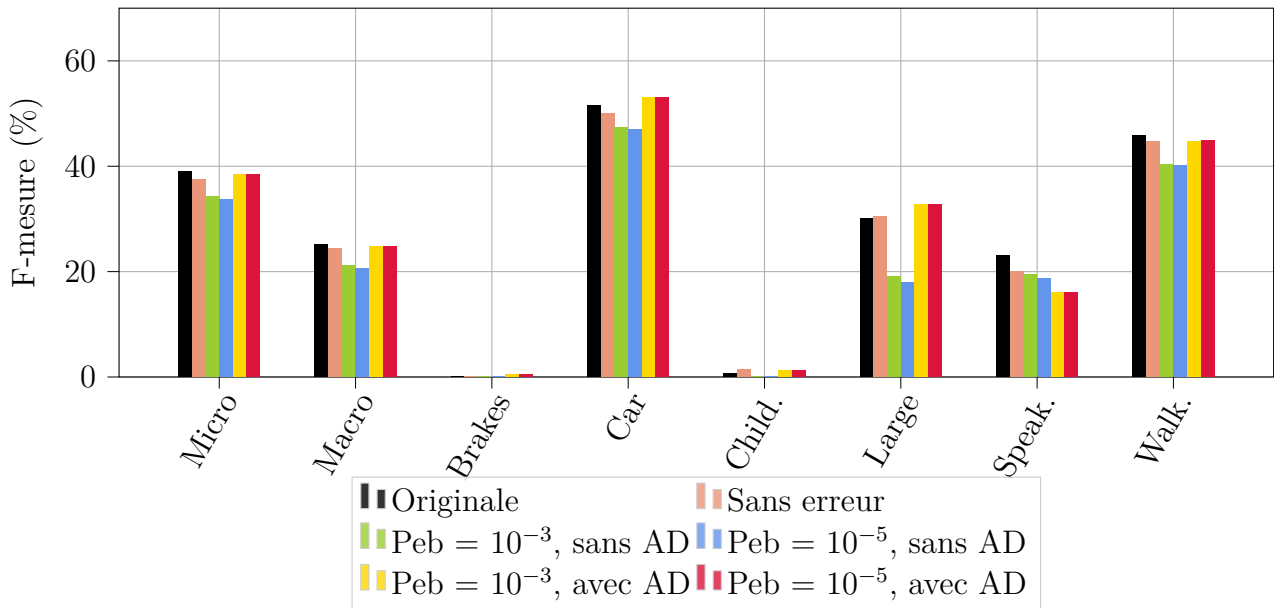
(a) F-mesure pour des paires de descripteurs codées sur $n = 1$ bit(b) F-mesure pour des paires de descripteurs codées sur $n = 8$ bits

FIGURE 5.8 – Évolution de la F-mesure pour un **CRNN** et dont les paires de descripteurs sont codées sur n bits, pour différentes PEB en réception (sans erreur, 10^{-3} et 10^{-5}) et un apprentissage réalisé avec ou sans augmentation de données. Les résultats sont présentés pour les totaux par micro-moyennage (Micro) et macro-moyennage (Macro), ainsi que pour l'ensemble des classes : *brakes squeaking* (Brakes), *car* (Car), *children* (Child.), *large vehicle* (Large), *people speaking* (Speak.) et *people walking* (Walk.). La comparaison est réalisée avec les performances originales sans quantification et avec celles pour un apprentissage sur données quantifiées sans erreur de décodage lors du test.

5.7 Conclusion

Dans ce chapitre, nous avons souhaité présenter une vue d'ensemble des techniques permettant d'adapter voire de valoriser l'implémentation d'un système de classification distribué au sein d'un réseau de capteurs. Pour cela, nous avons commencé par introduire les spécificités et les obstacles qu'un tel dispositif peut apporter. Nous avons en particulier abordé le problème de la distribution des calculs entre les différents nœuds du système, ce problème étant exposé sous le prisme d'un cas applicatif dans lequel communiquent un nœud-source et un moniteur central. Des explications sur la faisabilité de l'implémentation ont également été apportées.

L'intérêt d'un réseau de nœuds multi-capteurs réside en la possibilité d'exploiter les signaux issus des différents nœuds ou canaux pour potentiellement améliorer la classification. Une présentation des principaux descripteurs spatiaux et des possibilités d'emploi de plusieurs nœuds a alors été réalisée. A cette occasion, nous avons réalisé plusieurs expériences sur l'usage des TDOA, des descripteurs spatiaux qui s'intéressent à la position de la source dominante. Malheureusement, conformément à d'autres études menées sur le sujet, aucune amélioration de la détection d'événements sonores n'a été observée.

Nous nous sommes ensuite intéressés à la contrainte de la quantité de données d'apprentissage, qui représente un problème majeur dans l'usage des RN, qui en nécessitent un grand nombre. L'une des solutions consiste à réaliser une augmentation artificielle des données. Il existe pour ce faire un grand nombre de techniques, dont nous avons décrit les principales compatibles avec l'usage de signaux polyphoniques. Nous avons également conduit des expériences pour visualiser l'intérêt de deux d'entre elles, l'augmentation du recouvrement et le *mixup*. La première méthode s'est révélée être néfaste pour les performances des deux classifieurs utilisés, tandis que la seconde a montré des limitations dues à l'usage de signaux polyphoniques.

Finalement, nous avons évoqué la quantification des descripteurs, nécessaire à l'implémentation en virgule fixe sur FPGA du dispositif de calcul, ainsi qu'à leur transmission entre le nœud de capteurs et le moniteur central. Nous inspirant de l'usage en télécommunications et en reconnaissance de la parole distribuée, nous avons étudié l'impact de la quantification fractionnée et de sa prise en compte éventuelle sur la capacité de nos classifieurs à détecter correctement chaque classe. Nous avons ensuite simulé l'ajout inédit d'erreurs sur les données binaires émises sur le canal physique lors d'une transmission entre deux nœuds. Nos analyses montrent que l'entraînement des classifieurs sur des données quantifiées permet de limiter fortement l'impact de ces erreurs sur la détection des événements acoustiques, et ce quelle que soit la longueur de quantification de chaque paire de descripteurs. Les dégradations restantes semblent pouvoir être empêchées par l'augmentation de données d'apprentissage en se basant sur l'ajout d'erreurs artificielles. Ces résultats sont très encourageants pour une éventuelle implémentation matérielle d'un tel dispositif.

CONCLUSIONS ET PERSPECTIVES

6.1 Conclusion

Depuis quelques années, la recherche théorique en classification d'événements sonores a atteint un tel niveau d'efficacité qu'il devient envisageable de mettre en pratique de tels dispositifs tout en conservant des performances satisfaisantes. Cela a conduit à la multiplication des applications possibles, allant des domaines de la sécurité à la santé en passant par la recherche ornithologique. En parallèle, l'usage des objets connectés, et en particulier des réseaux de capteurs sans fil, a été facilité par la baisse des coûts de l'électronique, notamment par les microphones *micro-electro-mechanical system* (MEMS) capables de capter un signal audio de bonne qualité pour un très faible coût [7].

Il en résulte une forte volonté d'implémenter des dispositifs de détection d'événements sonores sur des réseaux de capteurs sans fil. Malheureusement, une telle transposition du code de la théorie à la pratique nécessite des adaptations bien précises qui sont encore peu étudiées. La distribution des calculs induite par l'existence de plusieurs nœuds en est l'une des problématiques principales. En effet, la communication entre les nœuds de capteurs ainsi qu'avec un éventuel moniteur central constitue l'étape la plus coûteuse énergétiquement malgré la création de protocoles de communication faible consommation tels que Sigfox, LoRa, Ingenu ou NB-IOT (à titre d'exemple, on estime que les batteries des appareils communiquant à l'aide de l'un des protocoles *low power wireless access network* (LPWAN) a une durée de vie moyenne au moins supérieure de 10 ans à celle d'un appareil fonctionnant sur un réseau cellulaire classique [91]). Une faible consommation permet alors de prolonger la vie des nœuds de capteurs et d'alléger la maintenance du système, le rendant à la fois plus sûr et moins contraignant.

Dans nos travaux, on se place dans le contexte d'un réseau de capteurs sans fil utilisant ces technologies dans le cadre de communications à courte portée. Cette thèse a été successivement conduite sur deux axes de recherche. Le premier concerne la consommation d'une transmission sans fil afin de comprendre comment la minimiser tout en maintenant un débit satisfaisant. Ce faisant, nous nous sommes également intéressés à la mise en place d'un dispositif de détection d'événements sonores distribuée sur système embarqué et à l'influence de la distribution des calculs sur celui-ci.

Consommation d'une communication sans fil

L'estimation de la consommation réelle d'une communication sans fil suppose un modèle au plus proche du fonctionnement du matériel. Malheureusement, nous avons remarqué que la plupart des modèles utilisés en recherche approximent plusieurs paramètres comme étant des constantes, à l'instar du *peak to average power ratio* (PAPR) et de l'efficacité du drain qui permettent de calculer la puissance consommée par l'amplificateur de puissance, ou encore les puissances du convertisseur numérique-analogique (CNA) et du convertisseur analogique-

numérique (CAN). Une première partie de notre travail a donc consisté à réunir en un unique modèle toutes les expressions de ces éléments dépendant du système, ce qui n'avait encore jamais été réalisé. Le modèle de consommation d'énergie par bit émis proposé prend ainsi en compte un PAPR variant en fonction des facteurs de roll-off et de suréchantillonnage, une efficacité du drain croissant avec la puissance de transmission, et des puissance des CNA et CAN fonction notamment de la bande passante et du facteur de suréchantillonnage. Nos analyses montrent que les consommations estimées par les modèles traditionnel et proposé sont différentes. En particulier, le choix de l'un ou de l'autre des modèles joue sur l'efficacité spectrale qui minimise l'énergie consommée.

Ce dernier point a justement été notre deuxième sujet d'intérêt. En effet, les contraintes de temps-réel et de performance minimale imposent souvent de transmettre une certaine quantité d'information binaire à un débit minimal (et donc à une efficacité spectrale minimale) et à une probabilité d'erreur binaire (PEB) fixée. Or la PEB est une fonction du RSB en réception, lui-même paramètre de la consommation d'énergie par bit émis sur le canal. L'estimation de cette énergie nécessite donc le calcul de la réciproque de la PEB, réciproque qui se révèle bien souvent complexe et difficile à manipuler pour l'approximation de l'efficacité spectrale minimisant la consommation. Une deuxième solution pour estimer la consommation consiste à passer par la capacité du canal. En effet, celle-ci dépend également du rapport signal sur bruit (RSB), et ce de manière tout à fait évidente dans le cadre d'une source gaussienne émise sur un canal AWGN. Lorsqu'on souhaite travailler à partir d'une autre source, il est possible de déterminer une pénalité du RSB entre la source cible et la source gaussienne qui ne dépende que de la PEB. Dans nos travaux, nous avons ainsi proposé une nouvelle expression de la capacité pour un canal de Rayleigh, ainsi que des approximations inédites de la pénalité d'une source modulée par M-QAM pour des canaux AWGN et de Rayleigh lors d'une communication point-à-point. Notre étude montre que l'évolution de la consommation d'énergie par bit transmis en exprimant celle-ci via la réciproque de la PEB ou via la réciproque de la capacité du canal est pratiquement identique en fonction de l'efficacité spectrale, ce qui confirme les approximations proposées. Finalement, il devient possible d'exprimer simplement l'efficacité spectrale optimale minimisant la consommation. Nous proposons ensuite une analyse de cette efficacité spectrale optimale en fonction de différents paramètres du système tels que la PEB, la bande passante ou encore le facteur de roll-off.

Enfin, les réseaux de capteurs sans fil permettent d'envisager l'utilisation d'un nœud-relai pour transmettre l'information au destinataire en plus du nœud-source. En effet, la puissance de transmission requise à PEB fixé croît exponentiellement avec la distance entre les nœuds source et destinataire. L'insertion d'un nœud-relai à une distance plus courte du destinataire peut donc permettre de diminuer la puissance de transmission de la source, et peut-être ainsi la consommation d'énergie totale. Notre thèse propose une étude de la consommation totale d'énergie par bit émis pour une transmission coopérative avec un protocole AF et pour un nœud-relai pouvant être placé à différentes positions dans un espace à une puis deux dimensions. Nos analyses tendent à montrer que, en fonction de l'efficacité spectrale et du canal choisis, il est possible de diminuer la consommation à l'aide d'un nœud-relai par rapport à une transmission point-à-point. En outre, la position du relai joue une grande importance dans cette consommation : il semble en effet que les énergies les plus basses soient atteintes lorsque le relai est proche du destinataire.

Détection d'événements sonores

Notre second axe de recherche répond à notre envie de comprendre les enjeux d'une implémentation d'un système de détection d'événements sonores sur un réseau de capteurs sans fil. Le premier élément que nous avons abordé dans nos travaux est l'idée d'une spatialisation de la détection d'événements sonores, ce qui pourrait éventuellement apporter de l'information supplémentaire. Nous avons regroupé les techniques utilisées dans l'état de l'art en deux blocs distincts de méthodes : les descripteurs spatiaux d'une part, les architectures multi-canaux d'autre part. Les descripteurs spatiaux sont généralement calculés à partir de canaux issus d'un même nœud multi-capteurs, et ont pour objectif de mettre en évidence le lien entre classe et position ou déplacement de l'événement. Malheureusement, les quelques études sur le sujet indiquent une grande difficulté à établir véritablement une corrélation qui permette une augmentation des performances. Les techniques modifiant l'architecture du système de détection d'événements sonores pour y inclure plusieurs canaux (issus du même nœud ou non) semblent plus prometteuses en cela qu'elles apportent une redondance plutôt que de nouvelles informations. La multiplication des points de vue – en mettant en valeur des différences entre les canaux ou en employant plusieurs classifieurs – permet de choisir les informations les plus pertinentes et de pallier le défaut (éventuellement temporaire) de détection d'une partie du système – à l'instar d'un système d'écoute humaine, qui fonctionnerait mieux en employant plusieurs personnes.

L'écoute humaine est justement très liée à la classification acoustique, car elle est à l'origine des étiquettes utilisées en classification supervisée. Or nous avons vu que cet étiquetage des données est à la fois long, fastidieux et coûteux, ce qui empêche bien souvent d'obtenir une base de données suffisamment conséquente pour entraîner efficacement certains classifieurs (en particuliers les réseaux de neurones) lors d'applications pratiques. Le recours à des méthodes d'augmentation de données, qui créent artificiellement de nouvelles données légèrement différentes à partir de celles d'origine, est ainsi devenu systématique. Au fil des années, ces techniques ont largement fait leurs preuves pour accroître significativement les performances de la détection. Parmi les méthodes qui ont gagné en influence, on retrouve notamment celles issues du traitement de l'image grâce à l'utilisation fréquente des couches convolutionnelles dans les réseaux de neurones.

Les réseaux de capteurs à faible consommation impliquent généralement de coder l'ensemble des variables en virgule fixe plutôt qu'en virgule flottante. Ceci est également rendu obligatoire par la transmission des informations entre les nœuds de capteur et un moniteur central. Quelques études se sont intéressées à la quantification des hyper-paramètres des réseaux de neurones ou à l'impact de la transmission d'un signal brut sur la classification. Néanmoins, à notre connaissance, aucune étude antérieure n'a été menée sur les questions de l'encodage et de la transmission des descripteurs ce qui permet de limiter les calculs sur le nœud-source tout en permettant de préserver la vie privée en cas de fuite des données (par rapport à un signal brut). En nous inspirant des techniques employées en télécommunication pour la transmission de la parole, nous avons proposé une quantification fractionnée des vecteurs de descripteurs par paires à l'aide de l'algorithme Linde-Buzo-Gray (LBG). Afin de vérifier l'intérêt de la prise en compte de la quantification pour l'apprentissage, nous avons réalisé des comparaisons des performances pour des données de test quantifiées et des données d'apprentissage quantifiées ou non. Nos expériences ont été réalisées pour deux classifieurs : un GMM soustractif entraîné sur des MFCC et leurs dérivées premières et secondes, issu de l'état de l'art de l'époque pré-réseaux de neurones, et un CRNN considéré comme parmi les plus performants pour la détection d'évé-

nements sonores au début de ces travaux, et entraîné sur un mel-spectrogramme. En étudiant l'influence de différentes longueurs de quantification pour chaque paire, nous avons observé qu'un apprentissage spécifique sur données quantifiées est bénéfique pour les courtes longueurs de quantification (1 ou 2 bits), mais qu'il n'est pas nécessaire lorsque celle-ci atteint un niveau suffisant (à partir de 8 bits). Cependant, lorsqu'on simule une transmission bruitée des descripteurs sur un canal de communication sans fil, ce qui conduit à l'adjonction d'erreurs lors du décodage des descripteurs, les résultats sont tout autre. En effet, pour le GMM soustractif, nous avons démontré que l'apprentissage spécifique permet de prévenir la décroissance des performances qu'un apprentissage sur les données d'origine n'offre pas. En revanche, pour le CRNN, la F-mesure relevée décroît quelle que soit la nature de l'apprentissage lorsque certains descripteurs sont erronés. Pour contrer cette diminution des performances, nous avons proposé une méthode d'augmentation de données qui simule la transmission des descripteurs sur un canal de communication pour plusieurs PEB en réception. L'apprentissage sur ces données artificiellement erronées a permis de retrouver une F-mesure au niveau d'origine sans erreur.

6.2 Perspectives

Ces travaux ont tenté d'interroger la problématique de l'implémentation d'un système de détection d'événements sonores sur un réseau de capteurs à faible consommation. Nous pensons avoir répondu à plusieurs questions pour avancer dans ce sens, mais bien d'autres persistent encore. Dans cette section, nous proposons de présenter plusieurs points qui nous semblent d'intérêt dans des travaux futurs.

Consommation d'une communication sans fil

Nous commençons par introduire les sujets qui concernent la communication sans fil. L'élargissement le plus évident de notre contexte consiste à prendre en compte davantage d'éléments de la chaîne de communication tels que le codage de canal. On pourrait alors tenter de savoir s'ils jouent un rôle dans la consommation d'énergie d'une transmission, et notamment par des modifications au niveau du RSB en réception. L'aspect retransmission des données en cas d'erreur de décodage a également été abordé très succinctement, et pourrait donc être davantage intégré dans nos recherches pour la diminution de l'énergie consommée.

Néanmoins, l'axe le plus pertinent semble être celui de la transmission coopérative, que nous n'avons pu traiter en détail. Par exemple, nous avons étudié uniquement le protocole AF et la combinaison MRC pour le calcul de la consommation. L'étude et la comparaison de différents protocoles et combinaisons en réception nous permettrait de comprendre comment la consommation peut être diminuée selon le contexte. En outre, l'efficacité spectrale minimisant l'énergie par bit émis n'a pas été exprimée analytiquement. Une analyse menée brièvement nous a indiqué que le lien entre consommation d'énergie et capacité est beaucoup plus complexe pour une communication coopérative que point-à-point, il pourrait donc s'agir d'une piste intéressante.

Afin de mener cette étude, il nous semble pertinent de s'interroger d'abord sur les couples de RSB source-destinataire et relai-destinataire $(\gamma_{sd}, \gamma_{rd})$. En effet, dans nos travaux, nous avons majoritairement considéré que le couple $(\gamma_{sd}, \gamma_{rd})$ choisi pour le calcul de la consommation était celui la minimisant a posteriori. Cependant, cela induit de calculer l'énergie totale pour tous les couples acceptables à PEB fixée, ce qui représente une forte complexité calculatoire et

peu de praticité. Or nous avons justement remarqué que, à efficacité spectrale et PEB fixés, le couple $(\gamma_{sd}, \gamma_{rd})$ choisi est toujours le même, et ce quelle que soit la distance source-destinataire. On peut donc supposer qu'il existe une manière de déterminer analytiquement ce couple en fonction des paramètres du système. Néanmoins, il semble également judicieux de s'interroger sur l'intérêt réel de ce couple pour la longévité des nœuds de capteurs. En effet, la minimisation de la consommation d'énergie totale par bit émis ne garantit pas que cette consommation soit équitablement partagée. Ainsi, si la consommation d'un des nœuds est bien supérieure à l'autre, sa longévité sera également moindre, conduisant à une mise en panne plus rapide du système global par rapport au cas d'une consommation uniforme suivant les nœuds.

Enfin, il pourrait être intéressant de prendre en compte le cas où un nœud-relai est lui-même source d'informations et apporte une information complémentaire aux autres nœuds. Ainsi, pour une même période de transmission des données, une nouvelle transmission entre le nœud-relai et le nœud-destinataire serait nécessaire (en plus des deux transmissions déjà présentes dans le protocole AF) de manière à transmettre cette information complémentaire. Dans un tel schéma, l'utilisation de cette information additionnelle par le nœud destinataire est également un axe potentiel de recherche.

Détection d'événements sonores

Nous abordons à présent la question des perspectives dans le domaine de la détection d'événements sonores. Là encore, un des éléments que nous n'avons pas traité, alors qu'il s'agit d'une thématique en pleine expansion, est la problématique de la reproductibilité du code et des résultats. En effet, nous pourrions nous pencher sur la normalisation de nos codes par rapport aux usages contemporains du domaine, et les publier.

Une deuxième zone d'intérêt, peu abordée dans nos travaux, est la complexité calculatoire des tests de détection et du calcul des descripteurs. Sur les réseaux de capteurs à très faible consommation, le compromis entre performance de la détection d'événements sonores et complexité calculatoire du système est très important et mériterait d'être étudié attentivement, d'autant plus que plusieurs travaux se sont déjà emparés de la question de la réduction de ces complexités.

L'introduction du multi-capteurs, comme nous l'avons vu, est encore limitée mais prend de l'ampleur d'année en année. Forts de notre compréhension de la consommation d'énergie des communications sans fil, nous pourrions tenter d'analyser les compromis à faire lors de l'utilisation et de la combinaison de signaux issus de différents nœuds pour l'amélioration de la détection d'événements sonores, par exemple en choisissant de n'utiliser certains nœuds que dans des situations précises et de faire évoluer ce contexte au cours du temps.

Enfin, la recherche sur la quantification des différents éléments du dispositif de détection d'événements sonores n'étant qu'à leurs balbutiements, cela semble être une occasion idéale d'imaginer une multitude de nouvelles expériences. En premier lieu, nos analyses n'ont été réalisées que sur deux classifieurs qui ne font usage que de deux types de descripteurs associés, ce qui ne correspond plus tout à fait au contexte des dispositifs actuels les plus avancés. Cela pourrait donc être intéressant de vérifier nos conclusions sur d'autres classifieurs, d'autres bases de données, et d'y inclure des techniques actuellement utilisées d'augmentation de données. D'un point de vue plus spécifique, nous avons choisi de quantifier toutes les paires des descripteurs sur la même longueur de quantification. En pratique, nous savons par exemple que tous les coefficients MFCC n'ont pas la même importance. Les normes en télécommunication

et reconnaissance de la parole distribuée proposent d'ailleurs d'encoder toutes les paires sur 6 bits, exceptée la première (constituée de la log-énergie et du premier coefficient mel) sur 8 bits. Il pourrait ainsi être pertinent de savoir s'il est possible de répartir différemment les données binaires entre les paires de descripteurs, et éventuellement d'économiser sur les données transmises. En outre, nos études considèrent uniquement une quantification en virgule fixe des descripteurs alors que pour une implémentation matérielle, toutes les variables sont stockées et les calculs réalisés en virgule fixe. Ceci pourrait finalement nous conduire à mener des simulations suffisamment avancées pour comprendre le comportement d'une implémentation matérielle, puis de la réaliser pour en faire l'expérience pratique.

6.3 Publications

Les travaux menés dans cette thèse ont mené aux publications suivantes :

- M. Lacroix, R. Rocher et P. Scalart, "Realistic power amplifier model for energy optimization in wireless networks", *Inter. Symp. on Wireless Comm. Systems (ISWCS)*, 2021.
- M. Lacroix, R. Rocher et P. Scalart, "Advanced energy model and spectral efficiency optimization in short-range wireless sensor networks", *Inter. Conf. on Green Computing and Comm. (GreenCom)*, 2021.
- M. Lacroix, R. Rocher et P. Scalart, "Realistic spectral efficiency analysis for energy-efficient wireless communications", à soumettre.
- M. Lacroix, R. Rocher et P. Scalart, "Spectral and energy efficiencies tradeoff analysis from an advanced energy model in wireless communication", à soumettre.

LISTE DES DESCRIPTEURS DITS "ARTISANAUX"

Descripteurs temporels

Les descripteurs temporels sont généralement calculés directement sur le signal temporel. Parmi eux, on trouve les enveloppes pour détecter les silences, le taux de passage par zéro, l'auto-corrélation, les moments temporels [150]...

Descripteurs spectraux

Certains descripteurs sont calculés à partir de la représentation fréquentielle du signal. De manière similaire aux descripteurs temporels, on peut obtenir l'énergie, l'enveloppe, les moments, la pente, etc, mais aussi la fréquence fondamentale ou période [150].

De transformations plus importantes peuvent être apportées pour obtenir ce que l'on appelle des coefficients spectraux, comme les MFCC, les *gammatone-frequency cepstral coefficients* (GFCC) ou les *constant-Q cepstral coefficients* (CQCC), respectivement calculés à partir des spectrogrammes mel, gammatone ou Q-constant [196].

Descripteurs perceptuels

Les descripteurs perceptuels, eux, tentent d'imiter le traitement de l'oreille humaine. Parmi eux, on cite l'intensité sonore et ses moments, ou la netteté [150].

Descripteurs issus du traitement d'image

Finalement, quelques descripteurs proviennent du traitement d'image, comme l'histogramme de gradient orienté, la distribution de puissance dans les sous-bandes ou les motifs binaires locaux.

DESCRIPTION DE L'INITIALISATION ALÉATOIRE D'UN GMM

L'initialisation aléatoire consiste à attribuer des poids aléatoires à l'ensemble des n_c gaussiennes et des N individus de la base d'apprentissage $X = (x_1, \dots, x_N)$, puis d'en déduire les paramètres du GMM. Pour cela, on procède en plusieurs étapes.

1. On génère une matrice $P = (p_{k,j})$ de taille $n_c \times N$ avec $k \in \llbracket 1, n_c \rrbracket$, $j \in \llbracket 1, N \rrbracket$ et $p_{k,j} \sim \mathcal{N}(0, 1)$.
2. On normalise cette matrice par rapport à l'axe des gaussiennes, ce qui nous donne la matrice $W = (w_{k,j})$ avec

$$w_{k,j} = \frac{p_{k,j}}{\sum_{i=1}^{n_c} p_{i,j}}. \quad (\text{B.1})$$

W correspond à la matrice des poids des gaussiennes pour chaque individu j , *i.e.* $\sum_{i=1}^{n_c} w_{i,j} = 1$.

3. On moyennant le poids des individus, on obtient le vecteur des poids de chaque gaussienne k :

$$\pi_k = \frac{1}{N} \sum_{j=1}^N w_{k,j}. \quad (\text{B.2})$$

4. Le vecteur moyenne de chaque gaussienne est ensuite calculé en attribuant simplement le poids $w_{k,j}$ à chacun des individus x_j correspondant. On obtient le vecteur

$$\mu_k = \frac{\sum_{j=1}^N w_{k,j} x_j}{\sum_{j=1}^N w_{k,j}}. \quad (\text{B.3})$$

5. Enfin, la matrice de covariance de chaque gaussienne est estimée de manière similaire à celle de l'algorithme EM, soit

$$\Sigma_k = \frac{\sum_{j=1}^N w_{k,j} (x_j - \mu_k)(x_j - \mu_k)^\top}{\sum_{j=1}^N w_{k,j}}. \quad (\text{B.4})$$

PREUVES DE CONVEXITÉ ET CALCUL DE b^*

C.1 Démonstration de la convexité de l'énergie totale pour un modèle simplifié issue de l'équation 3.101

Pour rappel, l'expression de l'équation 4.26 pour un modèle simplifié de l'énergie totale sur un canal AWGN et pour une source modulée par M-QAM est donnée par :

$$E_b(b) = A \underbrace{\frac{2^b - 1}{b}}_{h_1(b)} + \underbrace{\frac{C}{b}}_{h_2(b)} + D \quad (\text{C.1})$$

avec $b \in [2, +\infty[$. Puisque D est un terme constant et que A est strictement positif, si h_1 et h_2 sont convexes, alors E_b est convexe en tant que somme de fonctions convexes. C étant lui aussi strictement positif, h_2 est convexe pour $b \in [2, +\infty[$. Pour connaître la convexité de h_1 , il suffit de montrer que sa dérivée seconde est toujours strictement positive. le calcul de cette dérivée seconde nous donne :

$$(h_1)''(b) = \frac{\ln(2)^2 2^b b^2 - 2 \ln(2) 2^b b + 2(2^b - 1)}{b^3}. \quad (\text{C.2})$$

Pour $b \in [2, +\infty[$, b^3 est strictement positif. De plus, on vérifie par simulation que $\ln(2)^2 2^b b^2 - 2 \ln(2) 2^b b + 2(2^b - 1)$ est strictement positif pour $b \in [2, +\infty[$. Finalement, la dérivée seconde de h_1 est strictement positive, donc h_1 est strictement convexe pour $b \in [2, +\infty[$.

On a ainsi prouvé la stricte convexité de l'équation C.1.

C.2 Calcul de b^* pour un modèle simplifié de l'énergie

L'efficacité spectrale optimale b^* qui minimise l'énergie totale consommée par la transmission d'une source modulée par M-QAM sur un canal AWGN est l'unique solution de

$$\frac{\partial E_b}{\partial b}(b) = 0 \quad (\text{C.3})$$

avec $E_b(b)$ issue de l'équation 3.119 pouvant être simplifiée en

$$E_b(b) = a \frac{(2^b - 1)}{b} + \frac{c}{b} + d. \quad (\text{C.4})$$

On obtient ainsi la dérivée de l'énergie totale

$$\frac{\partial E_b}{\partial b}(b) = \frac{a(b \ln(2) 2^b - 2^b + 1) - c}{b^2}. \quad (\text{C.5})$$

On considère des valeurs de l'efficacité spectrale $b \in [2, +\infty[$, donc :

$$\begin{aligned}
& \frac{\partial E_b}{\partial b}(b) = 0 \\
& \Leftrightarrow a(b \ln(2) 2^b - 2^b + 1) - c = 0 \\
& \Leftrightarrow 2^b(b \ln(2) - 1) + 1 = \frac{c}{a} \\
& \Leftrightarrow \exp(b \ln(2))((b \ln(2) - 1) = \frac{c - a}{a} \\
& \Leftrightarrow \exp(b \ln(2) - 1)((b \ln(2) - 1) = \frac{c - a}{ae} \\
& \Leftrightarrow b \ln(2) - 1 = W\left(\frac{c - a}{ae}\right) \\
& \Leftrightarrow b = \frac{1}{\ln(2)} \left(1 + W\left(\frac{c - a}{ae}\right)\right)
\end{aligned}$$

avec $e = \exp(1)$ et W la fonction W de Lambert.

On retrouve finalement l'expression de l'efficacité spectrale optimale

$$b^* = \frac{1}{\ln(2)} \left(1 + W\left(\frac{c - a}{ae}\right)\right). \quad (\text{C.6})$$

La démonstration de l'expression de b^* pour le canal de Rayleigh est réalisée de manière analogue.

C.3 Validation de la convexité de l'énergie totale pour un modèle avancé sur un canal AWGN issue de l'équation 4.26

On commence par rappeler l'expression de l'équation 4.26 pour un modèle avancé de l'énergie totale sur un canal AWGN et pour une source modulée par M-QAM :

$$E_{b,AWGN}(b) = A \frac{2^{b/2} - 1}{2^{b/2} + 1} \frac{(2^b - 1)^{1-\beta}}{b} + \frac{C}{b} + D \quad (\text{C.7})$$

$$= A \underbrace{(2^{b/2} - 1)^2}_{f_1(b)} \underbrace{\frac{(2^b - 1)^{-\beta}}{b}}_{f_2(b)} + \underbrace{\frac{C}{b}}_{f_3(b)} + D \quad (\text{C.8})$$

avec $b \in [2, +\infty[$. Puisque D est un terme constant, $E_{b,AWGN}$ est convexe si $A f_1 f_2 + f_3$ est convexe. En outre, comme C est strictement positif, alors f_3 est convexe pour $b \in [2, +\infty[$. A étant lui aussi strictement positif, $E_{b,AWGN}$ est finalement convexe en tant que somme de fonctions convexes si et seulement si $f_1 f_2$ est convexe pour $b \in [2, +\infty[$.

Que deux fonctions soient convexes n'implique pas que leur produit soit convexe. En conséquence, il est nécessaire de prouver que la dérivée seconde de $f_1 f_2$ est strictement positive sur

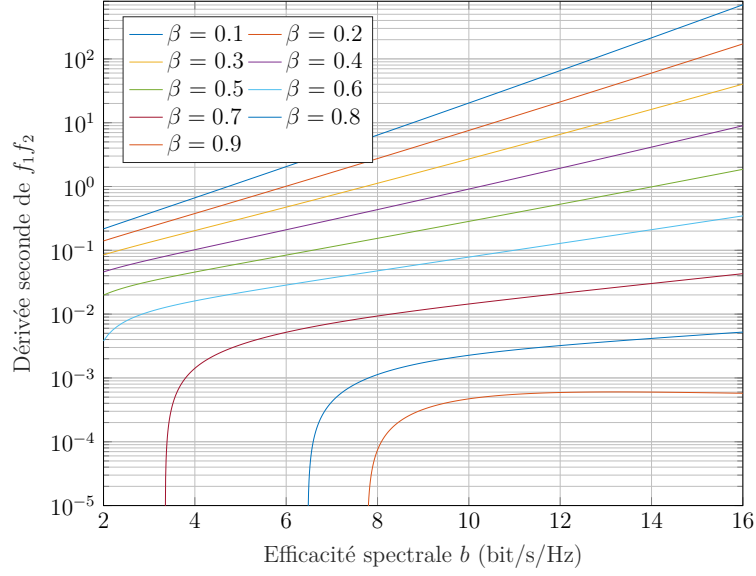


FIGURE C.1 – Dérivée seconde de f_1f_2 en fonction de l'efficacité spectrale b

$[2, +\infty[$. Le calcul de cette dérivée nous donne :

$$\begin{aligned}
 (f_1f_2)''(b) = & \frac{2(2^{b/2} - 1)^2}{b^3(2^b - 1)^\beta} + \frac{2^b \log(2)^2}{2b(2^b - 1)^\beta} - \frac{2 \cdot 2^{b/2} \log(2)(2^{b/2} - 1)}{b^2(2^b - 1)^\beta} + \frac{2^{b/2} \log(2)^2(2^{b/2} - 1)}{2b(2^b - 1)^\beta} \\
 & + \frac{2 \cdot 2^{b/2} \beta \log(2)(2^{b/2} - 1)^2}{b^2(2^b - 1)^{(\beta+1)}} - \frac{2^b \beta \log(2)^2(2^{b/2} - 1)^2}{b(2^b - 1)^{\beta+1}} - \frac{2 \cdot 2^{b/2} 2^b \beta \log(2)^2(2^{b/2} - 1)}{b(2^b - 1)^{\beta+1}} \\
 & + \frac{2^{2b} \beta \log(2)^2(2^{b/2} - 1)^2(\beta+1)}{b(2^b - 1)^{\beta+2}}
 \end{aligned} \tag{C.9}$$

Il est difficile de prouver que l'équation C.9 est strictement positive. Nous proposons donc de valider cette hypothèse expérimentalement, c'est-à-dire en traçant la dérivée seconde de f_1f_2 en fonction de b pour différentes valeurs du paramètre β . Les courbes résultantes sont tracées sur la figure C.1. On observe alors que les dérivées croissent quand b augmente, et ce pour toutes les valeurs du paramètre β . Les courbes, tracées en échelle logarithmique sur l'axe des ordonnées, indiquent que la dérivée seconde est toujours positive quand $\beta \leq 0.6$. En revanche, f_1f_2 apparaît comme non définie positive pour des valeurs $\beta > 0.6$. Ces observations nous permettent valider la convexité de l'équation 4.26 uniquement pour $\beta \leq 0.6$.

DÉTAILS SUR LE FACTEUR DE ROLL-OFF ET L'EFFICACITÉ DU DRAIN

D.1 Origine du facteur de roll-off

Le facteur de roll-off ρ mesure l'excès de bande passante d'un filtre, c'est-à-dire la largeur de bande occupée au-delà de la demi-bande passante de Nyquist $B_{1/2} = \frac{1}{2T_S}$ avec T_S la période d'un symbole. La demi-bande passante réellement occupée par le filtre se note alors $B_{1/2,r} = \frac{\rho+1}{2T_S}$. Un exemple de l'occupation de la bande pour un filtre passe-bas est donné sur la figure D.1.

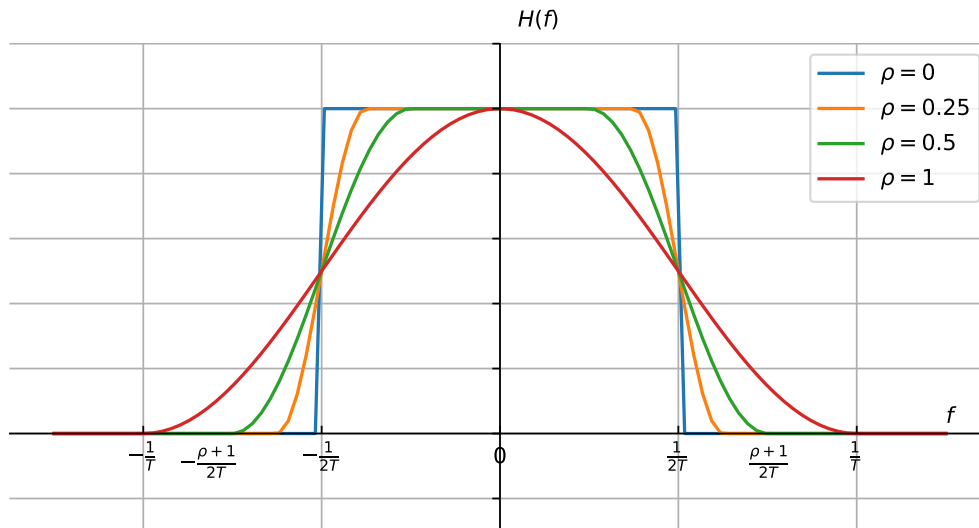


FIGURE D.1 – Exemple du rôle du facteur de roll-off ρ pour un filtre SRRC. Le dépassement de la bande passante est indiqué pour $\rho = 0.5$.

Quand ρ tend vers 0, le dépassement est quasiment nul et le filtre peut être considéré comme rectangulaire. En outre, de faibles valeurs de ρ conduisent généralement à une efficacité spectrale plus haute. Quand ρ tend vers 1, l'occupation de la bande est maximale [68].

Le facteur de roll-off a ainsi un lien particulier avec les filtres de mise en forme du signal. Il revêt par exemple d'une grande importance pour le filtre SRRC qui limite l'interférence entre symboles.

D.2 Calcul et propriétés de l'efficacité du drain

Un amplificateur de puissance est caractérisé par différentes puissances comme indiquées sur la figure D.2, qui représente un transistor amplificateur (G en est la grille, D le drain et S

la source). Les puissances entrantes sont : P_{in} la puissance du signal RF en entrée (aussi appelé puissance de transmission P_t dans nos travaux) et P_{DC} la puissance délivrée par l'alimentation continue. Les puissances sortantes sont : P_{out} la puissance du signal RF en sortie et P_{diss} la puissance dissipée par l'amplificateur. Ces puissances satisfont au principe de conservation de l'énergie [14], à savoir

$$P_{in} + P_{DC} = P_{diss} + P_{out}. \quad (D.1)$$

Généralement exprimée en pourcentage par rapport à l'efficacité maximale théorique η_{max} , l'efficacité énergétique du drain désigne alors le rapport entre les puissances de sortie et d'alimentation [14, 106], c'est-à-dire

$$\eta = \frac{P_{out}}{P_{DC}}. \quad (D.2)$$

Une importante valeur de η correspond à une grande efficacité du drain. En effet, en transformant l'équation D.1 de conservation de l'énergie pour inclure η , on obtient :

$$\frac{P_{in}}{P_{DC}} + 1 = \frac{P_{diss}}{P_{DC}} + \eta. \quad (D.3)$$

A puissance d'entrée P_{in} constante, une valeur de η importante signifie que la puissance dissipée est faible, d'où la notion d'efficacité.

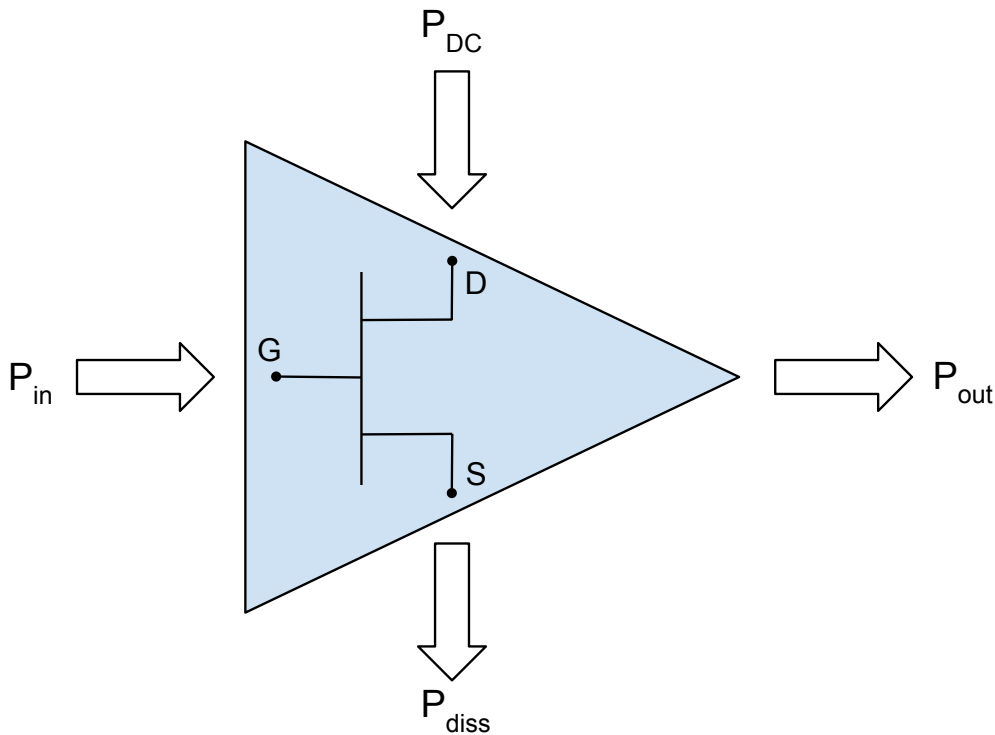


FIGURE D.2 – Schéma d'un transistor et des différents puissances impliquées

Il existe différentes classes d'amplificateurs, réparties en deux groupes : les amplificateurs dits linéaires (classes A, AB, B et C), et ceux dits non-linéaires (classes D, E et F), divisés en fonction de la tension d'alimentation choisie au point de fonctionnement.

L'ouverture du drain, qui influe sur le courant du drain, est un deuxième paramètre caractérisant les classes d'amplificateurs. L'angle d'ouverture du transistor, compris entre 0 et π radians, agit ainsi sur l'efficacité théorique idéale. Pour les amplificateurs linéaires, plus cet angle est faible, plus l'efficacité théorique idéale est grande. A 0 radians (angle correspondant à une classe d'amplificateur C), on estime alors qu'elle pourrait atteindre les 100% ; en revanche à π radians (angle correspondant à une classe A), elle tombe à 50% seulement. Ceci conduit à vouloir privilégier l'utilisation d'amplificateurs de classe C. Néanmoins, pour des faibles angles d'ouverture, la linéarité de l'amplificateur est dégradée, ce qui peut être à l'origine d'interférences sur le signal [14]. C'est notamment le cas pour des signaux modulés par M-QAM, qui sont impactés par la non-linéarité de l'amplificateur. En pratique, on ne peut donc pas toujours choisir l'amplificateur ayant la plus haute efficacité du drain idéale.

BIBLIOGRAPHIE

- [1] A. ABDI et M. KAVEH, « Comparison of DPSK and MSK Bit Error Rates for K and Rayleigh-Lognormal Fading Distributions », *IEEE Commun. Lett.*, t. 4, 4, p. 122-124, 2000.
- [2] S. ADAVANNE, « Sound event localization, detection, and tracking by deep neural networks », thèse de doct., Tampere University, 2020.
- [3] S. ADAVANNE, P. PERTILÄ et T. VIRTANEN, « Sound event detection using spatial features and convolutional recurrent neural network », in *Int. Conf. Acoust. Speech Signal Process.*, 2017, p. 771-775. arXiv : 1706.02291v1.
- [4] S. ADAVANNE et al., « Sound event detection in multichannel audio using spatial and harmonic features », in *Detect. Classif. Acoust. Scenes Events*, Budapest, 2016. arXiv : 1706.02293.
- [5] S. ADAVANNE et al., « Sound event localization and detection of overlapping sources using convolutional recurrent neural networks », p. 1-13, 2018. arXiv : 1807.00129.
- [6] J. P. ALEGRE PÉREZ, S. S. PUEYO et B. C. LOPEZ, *Automatic gain control*, SPRINGER, éd. 2011.
- [7] F. ALÍAS et R. M. ALSINA-PAGÈS, « Review of Wireless Acoustic Sensor Networks for Environmental Noise Monitoring in Smart Cities », *J. Sensors*, t. 2019, 2019.
- [8] M. S. ALOUINI et A. J. GOLDSMITH, « Capacity of Rayleigh fading channels under different adaptive transmission and diversity-combining techniques », *IEEE Trans. Veh. Technol.*, t. 48, 4, p. 1165-1181, 1999.
- [9] M. S. ALOUINI et M. K. SIMON, « An MGF-based performance analysis of generalized selection combining over rayleigh fading channels », *IEEE Trans. Commun.*, t. 48, 3, p. 401-415, 2000.
- [10] S. AMIRIPARIAN et al., « Deep unsupervised representation learning for abnormal heart sound classification », in *Int. Conf. IEEE Eng. Med. Biol. Soc.*, t. 2018-July, 2018, p. 4776-4779.
- [11] J. B. ANDERSEN, « Antenna arrays in mobile communications : Gain, diversity, and channel capacity », *IEEE Antennas Propag. Mag.*, t. 42, 2, p. 12-16, 2000.
- [12] ANFR, « Tableau national de répartition des bandes de fréquences », rapp. tech., 2017, p. 286.
- [13] M. ARVINTE, S. VISHWANATH et A. H. TEWFIK, « Deep learning-based quantization of L-Values for gray-coded modulation », in *IEEE Globecom Work.*, 2019. arXiv : 1906.07849.
- [14] M. EL-ASMAR, « Amélioration de la linéarité et du rendement énergétique des amplificateurs de puissance de topologie à deux branches pour les communications sans fil ; cas d'un amplificateur Linc », thèse de doct., Université du Québec, 2009.

-
- [15] K. AZARIAN, H. EL GAMAL et P. SCHNITER, « On the achievable diversity-multiplexing tradeoff in half-duplex cooperative channels », *IEEE Trans. Inf. Theory*, t. 51, 12, p. 4152-4172, 2005. arXiv : 0506018 [cs].
- [16] Z. BAI et al., « Performance analysis of SNR-based incremental hybrid decode-amplify-forward cooperative relaying protocol », *IEEE Trans. Commun.*, t. 63, 6, p. 2094-2106, 2015.
- [17] D. BATTAGLINO, « Acoustic Scene Classification : Contributions To Fundamental and Applied Research », thèse de doct., Télécom ParisTech, 2017.
- [18] S. BERGER et al., « Recent advances in amplify-and-forward two-hop relaying », *IEEE Commun. Mag.*, t. 47, 7, p. 50-56, 2009.
- [19] A. BERTRAND, « Applications and trends in wireless acoustic sensor networks : A signal processing perspective », in *IEEE Symp. Commun. Veh. Technol.*, 2011.
- [20] M. A. S. BHUIYAN et al., « Design of a nanoswitch in 130 nm CMOS technology for 2.4 GHz wireless terminals », *Bull. Polish Acad. Sci. Tech. Sci.*, t. 62, 2, p. 399-406, 2014.
- [21] V. BISOT, « Apprentissage de représentations pour l'analyse de scènes sonores », thèse de doct., Télécom ParisTech, 2018.
- [22] P. A. BROMILEY, « Products and convolutions of gaussian probability density functions », rapp. tech., 2014.
- [23] M. BROUSMICHE, J. ROUAT et S. DUPONT, « SECL-UMons database for sound event classification and Localization », in *IEEE Int. Conf. Acoust. Speech Signal Process.*, IEEE, 2020, p. 756-760.
- [24] L. BYTTEBIER et al., « Small-footprint acoustic scene classification through 8-bit quantization aware training and pruning of resnet models », in *Detect. Classif. Acoust. Scenes Events*, 2021.
- [25] R. CAI et al., « A flexible framework for key audio effects detection and auditory context inference », *IEEE Trans. Audio, Speech Lang. Process.*, t. 14, 3, p. 1026-1038, 2006.
- [26] E. CAKIR, « Deep neural networks for sound event detection », thèse de doct., Tampere University, 2019.
- [27] E. CAKIR et al., « Convolutional recurrent neural networks for polyphonic sound event detection », *IEEE/ACM Trans. Audio, Speech, Lang. Process.*, t. 25, 6, p. 1291-1303, 2017. arXiv : 1807.03043.
- [28] R. CAO et L. YANG, « The affecting factors in resource optimization for cooperative communications : A case study », *IEEE Trans. Wirel. Commun.*, t. 11, 12, p. 4351-4361, 2012.
- [29] Y. CAO et al., « Polyphonic sound event detection and localization using a two-stage strategy », in *Detect. Classif. Acoust. Scenes Events*, 2019, p. 30-34. arXiv : 1905.00268.
- [30] G. CERUTTI et al., « Neural network distillation on IoT platforms for sound event detection », in *Int. Speech Commun. Assoc.*, 2019, p. 3609-3613.
- [31] T. K. CHAN et C. S. CHIN, « A comprehensive review of polyphonic sound event detection », *IEEE Access*, t. 8, p. 103 339-103 373, 2020.

-
- [32] B. CHEN, « Audio recognition with distributed wireless sensor networks », thèse de doct., University of Victoria, 2010, p. 1-48.
 - [33] Q. CHEN et M. C. GURSOY, « Energy efficiency analysis in amplify-and-forward and decode-and-forward cooperative networks », in *IEEE Wirel. Commun. Netw. Conf.*, IEEE, 2010, p. 1-6. arXiv : 1004.0897.
 - [34] W. CHEN et al., « Fixed-point Gaussian Mixture Model for analysis-friendly surveillance video coding », *Comput. Vis. Image Underst.*, t. 142, p. 65-79, 2015.
 - [35] CHIPCON, *CC2420, 2.4 GHz IEEE 802.15.4 / ZigBee-ready RF Transceiver*, 2020.
 - [36] CHIPCON, *CC1000, Single Chip Very Low Power RF Transceiver*.
 - [37] F. CHOLLET et al., *Keras*, <https://keras.io>, 2015.
 - [38] A. CORNUÉJOLS, L. MICLET et V. BARRA, *Apprentissage artificiel*. Eyrolles, 2018.
 - [39] M. COULON, *Canal de propagation*, 2008.
 - [40] S. CUI, A. J. GOLDSMITH et A. BAHAI, « Modulation optimization under energy constraints », *IEEE Int. Conf. Commun.*, t. 4, p. 2805-2811, 2003.
 - [41] S. CUI, A. J. GOLDSMITH et A. BAHAI, « Energy-constrained modulation optimization », *IEEE Trans. Wirel. Commun.*, t. 4, 5, p. 2349-2360, 2005.
 - [42] W. DARGIE et C. POELLABAUER, *Fundamentals of wireless sensor networks*, WILEY, éd. 2010.
 - [43] S. DAS, S. PAL et M. MITRA, « Acoustic feature based unsupervised approach of heart sound event detection », *Comput. Biol. Med.*, t. 126, p. 103 990, 2020.
 - [44] S. DAUMONT, B. RIHAWI et Y. LOUT, « Root-raised cosine filter influences on PAPR distribution of single carrier signals », *Int. Symp. Commun. Control. Signal Process.*, p. 841-845, 2008.
 - [45] D. DE BENITO-GORRON, D. RAMOS et D. T. TOLEDANO, « A multi-resolution CRNN-based approach for semi-supervised sound event detection in DCASE 2020 challenge », *IEEE Access*, t. 9, p. 89 029-89 042, 2021.
 - [46] G. DEKKERS et al., « The SINS database for detection of daily activities in a home environment using an acoustic sensor network », in *Detect. Classif. Acoust. Scenes Events*, 2017.
 - [47] G. DEKKERS et al., « Dcase 2018 challenge - task 5 : monitoring of domestic activities based on multi-channel acoustics », in *Detect. Classif. Acoust. Scenes Events*, 2018, p. 2-5. arXiv : arXiv:1807.11246v2.
 - [48] C. B. P. DEL NOTARIO et A. MALANDA-TRIGUEROS, « Kohonen-GLA algorithms for efficient vector quantization », *Proc. IEEE Int. Conf. Electron. Circuits, Syst.*, t. 2, p. 671-674, 2000.
 - [49] J. R. DELGADO-CONTRERAS et al., « Feature selection for place classification through environmental sounds », in *Int. Conf. Emerg. Ubiquitous Syst. Pervasive Networks*, t. 37, 2014, p. 40-47.
 - [50] L. DELPHIN-POULAT et C. PLAPOUS, « Mean teacher with data augmentation for DCASE 2019 Task 4 », in *Detect. Classif. Acoust. Scenes Events*, 2019.

-
- [51] X. DENG et A. M. HAIMOVICH, « Power allocation for cooperative relaying in wireless networks », *IEEE Commun. Lett.*, t. 9, 11, p. 994-996, 2005.
 - [52] T. DEVRIES et G. W. TAYLOR, « Improved regularization of convolutional neural networks with cutout », 2017. arXiv : 1708.04552.
 - [53] M. DING et X. CHENG, « Fault tolerant target tracking in sensor networks », *Proc. Int. Symp. Mob. Ad Hoc Netw. Comput.*, p. 125-134, 2009.
 - [54] P. L. DOGNIN et al., « A fast, accurate approximation to log likelihood of Gaussian mixture models », *IEEE Int. Conf. Acoust. Speech Signal Process.*, 3, p. 3817-3820, 2009.
 - [55] M. DORFER et al., « Acoustic scene classification with fully convolutional neural networks and I-vectors », in *Detect. Classif. Acoust. Scenes Events*, 2018.
 - [56] J.-M. DUTERTRE, « Design radio fréquence », 2007.
 - [57] N. EL HODA DJIDI et al., « Enhancing wake-up radio range through minimum energy coding », *IEEE Int. Conf. Commun.*, p. 1-6, 2021.
 - [58] ETSI, « Speech processing, Transmission and Quality aspects (STQ) ; Distributed speech recognition ; Front-end feature extraction algorithm ; Compression algorithms », rapp. tech., 2003, p. 1-22.
 - [59] D. FENG et al., « A survey of energy-efficient wireless communications », *IEEE Commun. Surv. Tutorials*, t. 15, 1, p. 167-178, 2013.
 - [60] P. FOSTER et al., « Chime-home : a dataset for sound source recognition in a domestic environment », in *IEEE Work. Appl. Signal Process. to Audio Acoust.*, 2015, p. 4-8.
 - [61] H. T. FRIIS, « Noise figures of radio receivers », in *Proc. IRE*, t. 32, 1944, p. 419-422.
 - [62] W. S. GAN et J.-W. CHOI, *Spatial audio*, MDPI, éd. 2017, p. 206.
 - [63] M. GENOVESE et al., « FPGA implementation of gaussian mixture model algorithm for 47 fps segmentation of 1080p video », *J. Electr. Comput. Eng.*, t. 2013, 2013.
 - [64] D. GIANNOULIS et al., « A database and challenge for acoustic scene classification and event detection », in *Eur. Signal Process. Conf.*, IEEE, 2013, p. 1-5.
 - [65] R. GONZALEZ, « Better than MFCC audio classification features », 2013.
 - [66] M. W. W. van GROOTEL, A. T.C. et J. D. KRIJNDERS, « DARES-G1 : database of annotated real-world everyday sounds », in *NAG/DAGA Meet.*, 2009, p. 4.
 - [67] S. H. HAN et J. H. LEE, « An overview of peak-to-average power ratio reduction techniques for multicarrier transmission », *Modul. Coding Signal Process. Wirel. Commun.*, p. 56-65, 2005.
 - [68] K. HARAKO et al., « Roll-off factor dependence of Nyquist pulse transmission », *Opt. Express*, t. 24, 19, p. 21 986, 2016.
 - [69] D. C. HARRISON, W. K. G. SEAH et R. RAYUDU, « Rare event detection and propagation in wireless sensor networks », *ACM Comput. Surv.*, t. 48, 4, p. 1-22, 2016.
 - [70] J. A. HARTIGAN et M. A. WONG, « Algorithm AS 136 : A K-Means clustering algorithm », *J. R. Stat. Soc. Ser. C (Applied Stat.)*, t. 28, 1, p. 100-108, 1979.

-
- [71] M. HATA, « Empirical formula for propagation loss in land mobile radio services », *IEEE Trans. Veh. Technol.*, t. 29, 3, p. 317-325, 1980.
 - [72] W. R. HEINZELMAN, A. CHANDRAKASAN et H. BALAKRISHNAN, « Energy-efficient communication protocol for wireless microsensor networks », in *Int. Conf. Syst. Sci.*, 2000, p. 1-10.
 - [73] T. HEITTOLA, *DCASE*, <http://dcase.community/>.
 - [74] T. HEITTOLA, A. MESAROS et T. VIRTANEN, « Acoustic scene classification in DCASE 2020 Challenge : generalization across devices and low complexity solutions », in *Detect. Classif. Acoust. Scenes Events*, 2020. arXiv : 2005.14623.
 - [75] T. HELLER, E. COHEN et E. SOCHER, « A 102-129-GHz 39-dB Gain 8.4-dB noise figure I/Q receiver frontend in 28-nm CMOS », *IEEE Trans. Microw. Theory Tech.*, t. 64, 5, p. 1535-1543, 2016.
 - [76] S. HORNAUER et al., « Unsupervised discriminative learning of sounds for audio event classification », in *Int. Conf. Acoust. Speech Signal Process.*, 2021, p. 3035-3039. arXiv : 2105.09279.
 - [77] M. M. HOSSAIN et R. JÄNTTI, « Impact of efficient power amplifiers in wireless access », in *Conf. Green Commun.*, IEEE, 2011, p. 36-40.
 - [78] S. HUANG et al., « Energy efficiency and spectral-efficiency tradeoff in amplify-and-forward relay networks », *IEEE Trans. Veh. Technol.*, t. 62, 9, p. 4366-4378, 2013.
 - [79] C. HUCHER, « Definition and performance analysis of new cooperative protocols », thèse de doct., Télécom ParisTech, 2009.
 - [80] C. HUEBNER et al., « Long-range wireless sensor networks with transmit-only nodes and software-defined receivers », *Wirel. Commun. Mob. Comput.*, t. 13, p. 1499-1510, 2011. arXiv : arXiv:1405.1155v1.
 - [81] IEEE COMPUTER SOCIETY, *IEEE standards for Information technology, Telecommunications and information exchange between systems, Local and metropolitan area networks Specific requirements*. 2003, t. 12, p. 98-98.
 - [82] K. IMOTO et N. ONO, « Spatial cepstrum as a spatial feature using a distributed microphone array for acoustic scene analysis », *IEEE/ACM Trans. Audio, Speech, Lang. Process.*, t. 25, 6, p. 1335-1343, 2017.
 - [83] T. INOUE et al., « Domestic activities classification based on CNN using shuffling and mixing data augmentation », in *Detect. Classif. Acoust. Scenes Events*, 2018.
 - [84] D. ISTRATE, M. BINET et S. CHENG, « Real time sound analysis for medical remote monitoring », in *Int. IEEE EMBS Conf.*, 2008, p. 4640-4643.
 - [85] D. ISTRATE et al., « Multichannel smart sound sensor for perceptive spaces », *Complex Syst. Intell. Mod. Technol. Appl.*, p. 691-696, 2014.
 - [86] S. JAMALI et J. AHMADPANA, « Bird audio detection using supervised seighted NMF », rapp. tech., 2018.
 - [87] R. JAOUADI, « Compromis efficacité énergétique et efficacité spectrale pour les objets communicants autonomes », thèse de doct., Université de Rennes 1, 2017.

-
- [88] R. JAOUADI et al., « Modulation scaling for energy efficiency », in *Journées Sci. URSI Fr.*, 2016, p. 105-111.
 - [89] R. JAOUADI et al., « Rate optimization for energy efficient system with M-QAM », *Int. Conf. Comput. Netw. Commun.*, p. 975-979, 2017.
 - [90] N. JUILLERAT et B. HIRSBRUNNER, « Low latency audio pitch shifting in the frequency domain », *Int. Conf. Audio, Lang. Image Process. Proc.*, November, p. 16-24, 2010.
 - [91] Y. KABALCI et M. ALI, « Emerging LPWAN Technologies for Smart Environments : An Outlook », *Proc. - 2019 IEEE 1st Glob. Power, Energy Commun. Conf. GPECOM 2019*, p. 24-29, 2019.
 - [92] H. G. KIM et J. Y. KIM, « Environmental sound event detection in wireless acoustic sensor networks for home telemonitoring », *China Commun.*, t. 14, 9, p. 1-10, 2017.
 - [93] Y. KIM, S. JEONG et D. KIM, « A GMM-based target classification scheme for a node in wireless sensor networks », *IEICE Trans. Commun.*, t. E91-B, 11, p. 3544-3551, 2008.
 - [94] P. KLEIN, « Non-Intrusive Information Sources for Activity Analysis in Ambient Assisted Living Scenarios », thèse de doct., Université de Haute-Alsace, 2016.
 - [95] Y. KOIZUMI et al., « Description and Discussion on DCASE2020 Challenge Task2 : Un-supervised Anomalous Sound Detection for Machine Condition Monitoring », in *Detect. Classif. Acoust. Scenes Events*, 2020. arXiv : 2006.05822.
 - [96] K. KOUASSI, G. ANDRIEUX et J. F. DIOURIS, « PAPR distribution for single carrier M-QAM modulations », *Wirel. Pers. Commun.*, t. 104, 2, p. 727-738, 2019.
 - [97] D. KUDAVITHANA et al., « Energy modelling and optimization of amplify-and-forward relay transmission », in *IEEE Wirel. Commun. Netw. Conf. Work.*, IEEE, 2017, p. 1-6.
 - [98] A. KUMAR et B. RAJ, « Audio event detection using weakly labeled data », in *ACM Int. Conf. Multimed.*, 2016, p. 1038-1047. arXiv : 1605.02401.
 - [99] A. KUMAR et al., « Event detection in short duration audio using Gaussian Mixture Model and Random Forest Classifier », *Eur. Signal Process. Conf.*, p. 1-5, 2013.
 - [100] J. KÜRBY et al., « Bag-of-features acoustic event detection for sensor networks », in *Detect. Classif. Acoust. Scenes Events*, Budapest, 2016, p. 1-5.
 - [101] J. N. LANEMAN, D. N. TSE et G. W. WORNELL, « Cooperative diversity in wireless networks : Efficient protocols and outage behavior », *IEEE Trans. Inf. Theory*, t. 50, 12, p. 3062-3080, 2004.
 - [102] A. LAOURINE, A. STEPHENNE et S. AFFES, « Capacity of log-normal fading channels », *Int. Wirel. Commun. Mob. Comput. Conf.*, t. 2, 1, p. 13-17, 2007.
 - [103] E. LAUWERS et G. GIELEN, « Power estimation methods for analog circuits for architectural exploration of integrated systems », *IEEE Trans. Very Large Scale Integr. Syst.*, t. 10, 2, p. 155-162, 2002.
 - [104] A. LAVRIC et V. POPA, « A LoRaWAN : long range wide area networks study », *Int. Conf. Electromechanical Power Syst.*, t. 2017-Janua, p. 417-420, 2017.
 - [105] H. LEE, P. PHAM et A. Y. NG, « Unsupervised feature learning for audio classification using convolutional deep belief networks », in *Adv. Neural Inf. Process. Syst.*, 2009, p. 1096-1104.

-
- [106] T. H. LEE, *The Design of CMOS Radio-Frequency Integrated Circuits*. 2004.
 - [107] J. W. LEWIS et al., « Human brain regions involved in recognizing environmental sounds », *Cereb. Cortex*, t. 14, 9, p. 1008-1021, 2004.
 - [108] G. Y. LI et al., « Energy-efficient wireless communications : Tutorial, survey, and open issues », *IEEE Wirel. Commun.*, t. 18, 6, p. 28-35, 2011.
 - [109] J. LI, A. BOSE et Y. Q. ZHAO, « Rayleigh flat fading channels' capacity », *Annu. Commun. Networks Serv. Res. Conf.*, t. 2005, p. 214-217, 2005.
 - [110] Y. LI, B. BAKKALOGLU et C. CHAKRABARTI, « A comprehensive energy model and energy-quality evaluation of wireless transceiver front-ends », *IEEE Work. Signal Process. Syst. SiPS Des. Implement.*, t. 2005, p. 262-267, 2005.
 - [111] Y. LI, B. BAKKALOGLU et C. CHAKRABARTI, « A system level energy model and energy-quality evaluation for integrated transceiver front-ends », *IEEE Trans. Very Large Scale Integr. Syst.*, t. 15, 1, p. 90-102, 2007.
 - [112] Y. LI, D. QIAO et Y. ZHANG, « Energy efficient video transmission over fast fading channels », *Eurasip J. Wirel. Commun. Netw.*, t. 2010, 2010.
 - [113] Y. LI et al., « An improved relay selection scheme with hybrid relaying protocols », in *IEEE Glob. Telecommun. Conf.*, 2007, p. 3704-3708.
 - [114] F. LIN et al., « Impact of relay location according to ser for amplify-and-forward cooperative communications », *IEEE Int. Work. Anti-counterfeiting, Secur. Identif.*, 1, p. 324-327, 2007.
 - [115] Y. LINDE, A. BUZO et R. GRAY, « An algorithm for vector quantizer design », *IEEE Trans. Commun.*, t. 28, 1, p. 84-95, 1980.
 - [116] H. LIU et al., « An ensemble system for domestic activity recognition », in *Detect. Classif. Acoust. Scenes Events*, 2018, p. 2-4.
 - [117] O. P. R. LORD RAYLEIGH, « On our perception of sound direction », *London, Edinburgh, Dublin Philos. Mag. J. Sci.*, t. 5982, p. 214-232, 1907.
 - [118] B. LUITEL, Y. V. MURTHY et S. G. KOOLAGUDI, « Sound event detection in urban soundscape using two-level classification », *Int. Conf. Distrib. Comput. VLSI, Electr. Circuits Robot.*, p. 259-263, 2016.
 - [119] R. MAHAPATRA et al., « Energy efficiency tradeoff mechanism towards wireless green communication : A survey », *IEEE Commun. Surv. Tutorials*, t. 18, 1, p. 686-705, 2016.
 - [120] F. MAHMOOD, E. PERRINS et L. LIU, « Energy-efficient wireless communications : from energy modeling to performance evaluation », *IEEE Trans. Veh. Technol.*, t. 68, 8, p. 7643-7654, 2019.
 - [121] P. MAJI et R. MULLINS, « On the reduction of computational complexity of deep convolutional neural networks », *Entropy*, t. 20, 4, 2018.
 - [122] T. MAO et al., « Novel index modulation techniques : A survey », *IEEE Commun. Surv. Tutorials*, t. 21, 1, p. 315-348, 2019.
 - [123] I. MARTÍN-MORATÓ et al., « Low-complexity acoustic scene classification for multi-device audio : analysis of DCASE 2021 Challenge systems », in *Detect. Classif. Acoust. Scenes Events*, t. 28, 2021. arXiv : 2105.13734.

-
- [124] B. MCFEE et al., « librosa : Audio and music signal analysis in Python », in *Python Sci. Conf.*, t. 3, 2015, p. 18-25.
 - [125] N. MEHDIYEV et al., « Evaluating forecasting methods by considering different accuracy measures », in *Complex Adapt. Syst.*, t. 95, The Author(s), 2016, p. 264-271.
 - [126] X. MEI et al., « An encoder-decoder based audio captioning system with transfer and reinforcement learning », *Detect. Classif. Acoust. Scenes Events*, 2021. arXiv : 2108.02752.
 - [127] A. MESAROS, T. HEITTOLA et T. VIRTANEN, « Metrics for polyphonic sound event detection », *Appl. Sci.*, t. 6, 6, p. 162, 2016.
 - [128] A. MESAROS, T. HEITTOLA et T. VIRTANEN, « TUT database for acoustic scene classification and sound event detection », in *Eur. Signal Process. Conf.*, 2016, p. 1128-1132.
 - [129] A. MESAROS et al., « DCASE 2017 Challenge setup : Tasks, datasets and baseline », in *Detect. Classif. Acoust. Scenes Events*, 2017.
 - [130] A. MESAROS et al., « Detection and classification of acoustic scenes and events : Outcome of the DCASE 2016 challenge », *IEEE/ACM Trans. Audio, Speech, Lang. Process.*, t. 26, 2, p. 379-393, 2018. arXiv : 1604.07160.
 - [131] A. MESAROS et al., « Sound event detection in the DCASE 2017 Challenge », *IEEE/ACM Trans. Audio, Speech, Lang. Process.*, 2019.
 - [132] A. MESAROS et al., « Sound event detection : A tutorial », *IEEE Signal Process. Mag.*, t. 38, 5, p. 67-83, 2021. arXiv : 2107.05463.
 - [133] J. MIRANDA et al., « Path loss exponent analysis in wireless sensor networks : experimental evaluation », in *IEEE Int. Conf. Ind. Informatics*, IEEE, 2013, p. 54-58.
 - [134] M. M. MOLU et N. GOERTZ, « Amplify and forward relaying ; channel model and outage behaviour », in *Int. Symp. Wirel. Commun. Syst.*, IEEE, 2012, p. 206-210.
 - [135] A. MORAGREGA, C. IBARS et Y. GENG, « Energy efficiency of a cooperative wireless sensor network », in *Int. Work. Cross Layer Des.*, IEEE, 2009, p. 1-5.
 - [136] J. C. MURRAY, H. ERWIN et S. WERMTER, « Robotics sound-source localization and tracking using interaural time difference and cross-correlation », in *AI Work. NeuroBotics*, 2004, p. 89-97.
 - [137] R. U. NABAR, H. BÖLCSKEI et F. W. KNEUBÜHLER, « Fading relay channels : Performance limits and space-time signal design », *IEEE J. Sel. Areas Commun.*, t. 22, 6, p. 1099-1109, 2004.
 - [138] E. R. NASCIMENTO et al., « Acoustic-driven interior vehicle adaptation based on deep reinforcement learning to improve driver's comfort », in *IROS Work. Human/Robot In-the-loop Mach. Learn.*, 2018, p. 3-6.
 - [139] B. NASERSHARIF et A. AKBARI, « SNR-dependent compression of enhanced mel sub-band energies for compensation of noise effects on MFCC features », *Pattern Recognit. Lett.*, t. 28, 11, p. 1320-1326, 2007.
 - [140] A. A. NASIR et al., « Relaying protocols for wireless energy harvesting and information processing », *IEEE Trans. Wirel. Commun.*, t. 12, 7, p. 3622-3636, 2013.

-
- [141] T.-d. NGUYEN, « Stratégies de MIMO coopératif pour les réseaux de capteurs sans fil contraints en énergie », thèse de doct., 2010.
 - [142] U. T. NGUYEN et T. V. PHAM, « Performance assessment of generalized cross-correlation based algorithms for multisource point-based localization and detection », *Int. Conf. Adv. Technol. Commun.*, p. 303-306, 2011.
 - [143] H. NURMINEN et al., « Statistical path loss parameter estimation and positioning using RSS measurements in indoor wireless networks », *Int. Conf. Indoor Position. Indoor Navig.*, November, 2012.
 - [144] J. van OPSTAL, *The auditory system and human sound-localization behavior*. Academic Press, 2016.
 - [145] S. PANCOAST et M. AKBACAK, « Softening quantization in bag-of-audio-words », in *IEEE Int. Conf. Acoust. Speech Signal Process.*, IEEE, 2013, p. 778-782.
 - [146] P. P. PARADA et al., « Robust statistical processing of TDOA estimates for distant speaker diarization », *Eur. Signal Process. Conf.*, p. 86-90, 2017.
 - [147] D. S. PARK et al., « SpecAugment : A simple data augmentation method for automatic speech recognition », *Annu. Conf. Int. Speech Commun. Assoc.*, t. 2019-Sept, p. 2613-2617, 2019. arXiv : 1904.08779.
 - [148] N. P. PATEL et M. S. PATWARDHAN, « Identification of most contributing features for audio classification », *Int. Conf. Cloud Ubiquitous Comput. Emerg. Technol.*, p. 219-223, 2013.
 - [149] F. PEDREGOSA et al., « Scikit-learn : machine learning in Python », *J. Mach. Learn. Res.*, 12, p. 2825-2830, 2011.
 - [150] G. PEETERS, « A large set of audio features for sound description (similarity and classification) in the CUIDADO project », rapp. tech., 2004, p. 1-25.
 - [151] J. POHJALAINEN, « Methods of automatic audio content classification », thèse de doct., Helsinki University of Technology, 2007.
 - [152] G. E. POLINER et D. P. W. ELLIS, « A discriminative model for polyphonic piano transcription », *EURASIP J. Adv. Signal Process.*, 2007.
 - [153] A. POLITIS, S. ADAVANNE et T. VIRTANEN, « A dataset of reverberant spatial sound scenes with moving sources for sound event localization and detection », 2020. arXiv : 2006.01919.
 - [154] J. PORTÊLO et al., « Non-speech audio event detection », *IEEE Int. Conf. Acoust. Speech Signal Process.*, p. 1973-1976, 2009.
 - [155] M. G. PRASAD et al., « Analysis of fast fading in wireless communication channels », *Int. J. Syst. Technol.*, t. 3, 1, p. 139-145, 2010.
 - [156] J. G. PROAKIS et M. SALEHI, *Communication systems engineering*. Prentice Hall, 2002, t. 293, p. 149.
 - [157] W. QIN, M. HEMPSTEAD et Y. WOODWARD, « A realistic power consumption model for wireless sensor network devices », *IEEE Commun. Soc. Sens. Adhoc Commun. Networks*, t. 1, p. 286-295, 2006.

-
- [158] R. RAMESH et al., « LoRaWAN for smart cities : experimental study in a campus deployment », *LPWAN Technol. IoT M2M Appl.*, p. 327-345, 2020.
- [159] N. K. S. RAO et N. PETERS, « On the effect of coding artifacts on acoustic scene classification », in *Detect. Classif. Acoust. Scenes Events*, 2021.
- [160] H. K. RATH et al., « Realistic indoor path loss modeling for regular WiFi operations in India », in *arXiv*, 2017. arXiv : 1707.05554.
- [161] J. REX, « Audio event detection and localisation for directing video surveillance - a survey of potential techniques », in *Int. Conf. Imaging Crime Detect. Prev.*, 2012, P21-P21.
- [162] P. RUVOLO, I. FASEL et J. R. MOVELLAN, « A learning approach to hierarchical feature selection and aggregation for audio classification », *Pattern Recognit. Lett.*, t. 31, 12, p. 1535-1542, 2010.
- [163] M. RYYNÄNEN et A. K LAPURI, « Polyphonic Music Transcription Using Note Event Modeling », in *IEEE Work. Appl. Signal Process. to Audio Acoust.*, IEEE, 2005, p. 319-322.
- [164] Y. SAKASHITA et M. AONO, « Acoustic scene classification by ensemble of spectrograms based on adaptative temporal divisions », in *Detect. Classif. Acoust. Scenes Events*, t. 25, 2018, p. 29-29.
- [165] J. SALAMON et J. P. BELLO, « Deep convolutional neural networks and data augmentation for environmental sound classification », *IEEE Signal Process. Lett.*, t. 24, 3, p. 279-283, 2017. arXiv : 1608.04363.
- [166] J. SALAMON et al., « Scaper : A library for soundscape synthesis and augmentation », *IEEE Work. Appl. Signal Process. to Audio Acoust.*, t. 2017-Octob, p. 344-348, 2017.
- [167] E. L. SALOMONS et P. J. HAVINGA, « A survey on the feasibility of sound classification on wireless sensor nodes », *Sensors*, t. 15, 4, p. 7462-7498, 2015.
- [168] S. SAMPEI et T. SUNAGA, « Rayleigh fading compensation for QAM in land mobile radio communications », *IEEE Trans. Veh. Technol.*, t. 42, 2, 1993.
- [169] M. G. SÁNCHEZ, E. L. CID et A. V. ALEJOS, « Empirical modeling of radiowave angular power distributions in different propagation environments at 60 GHz for 5G », *Electron.*, t. 7, 12, 2018.
- [170] A. SENDONARIS, E. ERKIP et B. AAZHANG, « User cooperation diversity - Part II : Implementation aspects and performance analysis », *IEEE Trans. Commun.*, t. 51, 11, p. 1939-1948, 2003.
- [171] D. K. SHAEFFER et T. H. LEE, « A 1.5-V, 1.5-GHz CMOS low noise amplifier », *IEEE J. Solid-State Circuits*, t. 32, 5, p. 745-759, 1997.
- [172] P. M. SHANKAR, *Fading and shadowing in wireless systems*. Springer, 2012.
- [173] S. SHRESTHA et K. H. CHANG, « Analysis of outage capacity performance for cooperative DF and AF relaying in dissimilar rayleigh fading channels », in *IEEE Int. Symp. Inf. Theory - Proc.*, 2008, p. 494-498.
- [174] M. K. SIMON et M. S. ALOUINI, *Digital communication over fading channels : A unified approach to performance analysis*. Wiley, 2000.

-
- [175] S. SIVASANKARAN et K. M. M. PRABHU, « Robust features for environmental sound classification », in *IEEE Int. Conf. Electron. Comput. Commun. Technol.*, 2017.
 - [176] N. SOOD, A. K. SHARMA et M. UDDIN, « BER performance of OFDM-BPSK and -QPSK over Nakagami-m fading channels », *IEEE Int. Adv. Comput. Conf.*, 1, p. 88-90, 2010.
 - [177] H. SOURY, F. YILMAZ et M. S. ALOUINI, « Exact symbol error probability of square M-QAM signaling over generalized fading channels subject to additive generalized Gaussian noise », in *IEEE Int. Symp. Inf. Theory - Proc.*, 2013, p. 51-55.
 - [178] N. SRINIVASAMURTHY et al., « Towards efficient and scalable speech compression schemes for robust speech recognition applications », in *IEEE Int. Conf. Multi-Media Expo*, 2000, p. 249-252.
 - [179] D. STOWELL et al., « Detection and classification of acoustic scenes and events », *IEEE Trans. Multimed.*, t. 17, 10, p. 1733-1746, 2015. arXiv : [arXiv:1708.03211v2](https://arxiv.org/abs/1708.03211v2).
 - [180] D. STOWELL et al., « Bird detection in audio : a survey and a challenge », in *IEEE Int. Work. Mach. Learn. Signal Process.*, 2016.
 - [181] G. L. STÜBER, *Principles of mobile communication*. Springer, 2017, p. 1-830.
 - [182] S. SUMITANI et al., « Extracting the relationship between the spatial distribution and types of bird vocalizations using robot audition system HARK », in *Int. Conf. Intell. Robot. Syst.*, 2018, p. 2485-2490.
 - [183] N. SURAMPUDI, M. SRIRANGAN et J. CHRISTOPHER, « Enhanced feature extraction approaches for detection of sound events », in *IEEE Int. Conf. Adv. Comput.*, 2019, p. 223-229.
 - [184] H. TABANI et al., « An ultra low-power hardware accelerator for acoustic scoring in speech recognition », in *Parallel Archit. Compil. Tech.*, 2017, p. 41-52.
 - [185] Z. H. TAN, P. DALSGAARD et B. LINDBERG, « Automatic speech recognition over error-prone wireless networks », *Speech Commun.*, t. 47, 1-2, p. 220-242, 2005.
 - [186] R. TANABE et al., « Multichannel acoustic scene classification by blind dereverberation, blind source separation, data augmentation, and model ensembling », in *Detect. Classif. Acoust. Scenes Events*, 2018.
 - [187] C. TANG, H. LIU et G. PAN, « Performance analysis of log-normal and Rayleigh-lognormal fading channels », in *Int. Conf. Signal Process.*, IEEE, 2014, p. 1557-1560.
 - [188] M. TORABI, W. AJIB et D. HACCOUN, « Performance analysis of amplify-and-forward cooperative networks with relay selection over Rayleigh fading channels », in *IEEE Veh. Technol. Conf.*, 2009.
 - [189] L.-Q.-V. TRAN, « Energy-efficient cooperative relay protocols for wireless sensor networks », thèse de doct., Université de Rennes 1, 2012.
 - [190] N. TURPAULT, « Analyse des problématiques liées à la reconnaissance de sons ambiants en environnement réel », thèse de doct., Université de Lorraine, 2021.
 - [191] N. TURPAULT et al., « Sound event detection in domestic environments with weakly labeled data and soundscape synthesis », in *Detect. Classif. Acoust. Scenes Events*, 2019, p. 253-257.

-
- [192] N. TURPAULT et al., « Improving sound event detection in domestic environments using sound separation », in *Detect. Classif. Acoust. Scenes Events*, 2020.
- [193] M. VACHER et al., « The Sweet-Home speech and multimodal corpus for home automation interaction », in *Lang. Resour. Eval. Conf.*, 2014, p. 4499-4506.
- [194] G. VALENZISE et al., « Scream and gunshot detection and localization for audio-surveillance systems », *IEEE Conf. Adv. Video Signal Based Surveill.*, p. 21-26, 2007.
- [195] T. VIRTANEN, « Tutorial on detection and classification of acoustic scenes and events », in *IEEE Int. Conf. Acoust. Speech Signal Process.*, 2019.
- [196] T. VIRTANEN, M. D. PLUMBLEY et D. P. W. ELLIS, *Computational analysis of sound scenes and events*. Springer, 2018.
- [197] H. R. WALKER, « Intermediate-Frequency Amplifiers », *Encycl. RF Microw. Eng.*, 2005.
- [198] H. WAN, « Energy efficiency, synchronization and power control for the implementation of cooperative communications in wireless sensor networks », thèse de doct., Université de Nantes, 2010.
- [199] D. WANG et al., « A CRNN-based system with mixup technique for large-scale weakly labeled sound event detection », in *Detect. Classif. Acoust. Scenes Events*, 2018.
- [200] H. WANG, D. LYMBEROPOULOS et J. LIU, « Local business ambience characterization through mobile audio sensing », in *Int. Conf. World Wide Web*, 2014, p. 293-303.
- [201] J. C. WANG et al., « Mixed sound event verification on wireless sensor network for home automation », *IEEE Trans. Ind. Informatics*, t. 10, 1, p. 803-812, 2009.
- [202] W. WANG, T. JOST et R. RAULEFS, « A semi-deterministic path loss model for in-harbor LoS and NLoS environment », *IEEE Trans. Antennas Propag.*, t. 65, 12, p. 7399-7404, 2017.
- [203] K. N. WATCHARASUPAT et al., « Improving polyphonic sound event detection on multichannel recordings with the Sørensen-Dice coefficient loss and transfer learning », in *Detect. Classif. Acoust. Scenes Events*, 2021, p. 1-5. arXiv : 2107.10471.
- [204] S. WEI et al., « Sample mixed-based data augmentation for domestic audio tagging », in *Detect. Classif. Acoust. Scenes Events*, 2018. arXiv : 1808.03883.
- [205] W. WEI et al., « A-CRNN : A domain adaptation model for sound event detection », in *ICASSP 2020 - 2020 IEEE Int. Conf. Acoust. Speech Signal Process.*, 2020, p. 276-280.
- [206] M. WEN et al., « A survey on spatial modulation in emerging wireless systems : research progresses and applications », *IEEE J. Sel. Areas Commun.*, t. 37, 9, p. 1949-1972, 2019. arXiv : 1907.02941.
- [207] F. XIONG, « M-ary amplitude shift keying OFDM system », *IEEE Trans. Commun.*, t. 51, 10, p. 1638-1642, 2003.
- [208] X. XU et al., « A CRNN-GRU based reinforcement learning approach to audio captioning », in *Detect. Classif. Acoust. Scenes Events*, 2020, p. 225-229.
- [209] Y. XU et al., « Convolutional gated recurrent neural network incorporating spatial features for audio tagging », in *2017 Int. Jt. Conf. Neural Networks*, IEEE, 2017, p. 3461-3466.

-
- [210] K. YAMAMOTO et T. SUGIYAMA, « Separation distance performances between Ingenu and Wi-Fi systems with concatenated FEC », in *Int. Conf. Inf. Commun. Technol. Converg.*, 2021, p. 413-415.
 - [211] C. YANG et H. R. SHAO, « WiFi-based indoor positioning », *IEEE Commun. Mag.*, t. 53, 3, p. 150-157, 2015.
 - [212] S. YANG et J.-c. BELFIORE, « Towards the optimal amplify-and-forward », *IEEE Trans. Inf. Theory*, t. 53, 9, p. 3114-3126, 2007.
 - [213] Y. YE et al., « Improved hybrid relaying protocol for DF relaying in the presence of a direct link », *IEEE Wirel. Commun. Lett.*, t. 8, 1, p. 173-176, 2019.
 - [214] H. YEGANEH et al., « Weighting of mel sub-bands based on SNR/entropy for robust ASR », in *IEEE Int. Symp. Signal Process. Inf. Technol.*, 2008, p. 292-296.
 - [215] L. YOU et al., « Spectral efficiency and energy efficiency tradeoff in massive MIMO down-link transmission with statistical CSIT », *IEEE Trans. Signal Process.*, t. 68, p. 2645-2659, 2020. arXiv : 2004.03092.
 - [216] K. YU, J. EVANS et I. COLLINGS, « Performance analysis of pilot symbol aided QAM for rayleigh fading channels », *IEEE Int. Conf. Commun.*, t. 3, p. 1731-1735, 2002.
 - [217] M. YU et J. LI, « Is amplify-and-forward practically better than decode-and-forward or vice versa ? », in *IEEE Int. Conf. Acoust. Speech Signal Process.*, IEEE, 2005, p. 365-368.
 - [218] Y. ZHAN, « Low-complex environmental sound recognition algorithms for power-aware wireless sensor networks », thèse de doct., Keio University, 2012.
 - [219] H. ZHANG et al., « MixUp : Beyond empirical risk minimization », in *Int. Conf. Learn. Represent.*, 2018, p. 1-13. arXiv : arXiv:1710.09412v2.
 - [220] J. ZHANG, W. DING et L. HE, « Data augmentation and prior knowledge-based regularization for sound event localization and detection », in *Detect. Classif. Acoust. Scenes Events*, 2019, p. 1-5.
 - [221] W. ZHANG et al., « The study on distributed speech recognition system », in *IEEE Int. Conf. Acoust. Speech Signal Process.*, 2000, p. 1431-1434.
 - [222] X. J. ZHANG et Y. GONG, « Joint power allocation and relay positioning in multi-relay cooperative systems », *IET Commun.*, t. 3, 10, p. 1683-1692, 2009.
 - [223] Y. ZHAO, R. ADVE et T. J. LIM, « Symbol error rate of selection amplify-and-forward relay systems », *IEEE Commun. Lett.*, t. 10, 11, p. 757-759, 2006.
 - [224] H. ZHENG et al., « Path loss models for urban macro cell scenario at 3.35, 4.9 and 5.4 GHz », in *Int. Symp. Pers. Indoor Mob. Radio Commun.*, IEEE, 2015, p. 2229-2233.
 - [225] G. ZHIQING et al., « A low power fast-settling frequency-presetting PLL frequency synthesizer », *J. Semicond.*, t. 31, 8, 2010.
 - [226] J. ZHOU, « Sound event detection in multichannel audio LSTM network », in *Detect. Classif. Acoust. Scenes Events*, 2017. arXiv : 1412.6980.
 - [227] J. ZHU, « On the power efficiency and optimal transmission range of wireless sensor nodes », in *IEEE Int. Conf. Electro/Information Technol.*, 2009, p. 277-281.
 - [228] X. ZHUANG et al., « Real-world acoustic event detection », *Pattern Recognit. Lett.*, t. 31, 12, p. 1543-1551, 2010.

Titre : Détection d'événements sonores environnementaux dans les réseaux de capteurs sans fil à faible consommation

Mot clés : détection d'événements sonores, réseaux de capteurs sans fil, faible consommation, efficacité spectrale, quantification

Résumé : Ces dernières années, le gain d'intérêt pour la détection d'événements sonores et la réduction des coûts de l'électronique ont conduit à s'intéresser au déploiement de dispositifs de détection distribués sur des réseaux de capteurs sans fil. Cependant, la transmission d'informations binaires qu'elle requiert induit à la fois une forte consommation d'énergie et une dégradation des données émises. L'objectif de cette thèse est de travailler sur la minimisation de la consommation d'énergie d'une communication sans fil, et d'analyser l'impact de la transmission des descripteurs sur les performances de la détection. Tout d'abord, nous étudions les différents modèles énergétiques d'une transmis-

sion afin d'approcher au mieux la consommation réelle. Nous proposons alors une analyse de l'efficacité spectrale minimisant cette énergie. Dans un second temps, nous étudions la faisabilité d'un dispositif de détection d'événements sonores distribué, avec ses avantages et ses contraintes. Nous introduisons les méthodes utilisées pour coder en virgule fixe les descripteurs audio sur une longueur de quantification limitée. Nous commençons par analyser l'influence de cette quantification sur les performances d'un GMM ou d'un CRNN. Nous étudions finalement l'impact des erreurs de transmission sur cette même détection, et proposons une méthode d'augmentation de données pour le limiter.

Title: Environmental sound event detection from low power wireless sensor network

Keywords: sound event detection, wireless sensor networks, low power, spectral efficiency, quantization

Abstract: In recent years, increased interest in sound event detection and reduced electronics cost have led to a focus in the deployment of distributed detection devices on wireless sensor networks. However, the transmission of binary information that it requires induces both high energy consumption and degradation of the transmitted data. This thesis aims to work on minimizing the energy consumption of wireless communication, and to analyse the impact of features transmission on detection performance. First of all, we study the different transmission energy models in order to better approximate the real con-

sumption. We then propose an analysis of the spectral efficiency that minimizes this energy. In a second step, we study the feasibility of a distributed sound event detection system, with its advantages and constraints. We introduce the methods used to fixed-point encode the audio features over a limited quantization length. We start by analyzing the influence of this quantization on the performance of a GMM or CRNN. We finally study the impact of transmission errors on this same detection, and propose a data augmentation method to limit it.