



Codage de canal et codage réseau pour les CPL-BE dans le contexte des réseaux Smart Grid

Wendyida Abraham Wendyida Abraham Kabore

► To cite this version:

Wendyida Abraham Wendyida Abraham Kabore. Codage de canal et codage réseau pour les CPL-BE dans le contexte des réseaux Smart Grid. Traitement du signal et de l'image [eess.SP]. Université de Limoges, 2016. Français. NNT : 2016LIMO0038 . tel-01377897

HAL Id: tel-01377897

<https://theses.hal.science/tel-01377897>

Submitted on 11 Oct 2016

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

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

UNIVERSITÉ DE LIMOGES

ÉCOLE DOCTORALE S2I

Laboratoire XLIM C²S₂

Thèse

pour obtenir le grade de

**DOCTEUR DE L'UNIVERSITÉ DE
LIMOGES**

**Discipline : Électronique des Hautes Fréquences,
Photonique et Systèmes/Télécommunications**

présentée et soutenue par

Abraham KABORE

le 9 mars 2016

**Codage de canal et codage réseau pour
les CPL-BE dans le contexte des réseaux
Smart Grid**

**Thèse dirigée par Vahid MEGHDADI et Jean-Pierre
Cances**

Jury :

Prof. Philippe GABORIT	Université de Limoges	Président
Prof. Catherine DOUILLARD	Télécom Bretagne	Rapporteur
Prof. Jean-Marie GORCE	INSA Lyon	Rapporteur
Prof. Andrea TONELLO	University of Udine, Italie	Examineur
M. Jean-Luc GAUTIER	ERDF - Directeur Haute Vienne	Examineur
Prof. Jean-Pierre CANCES	ENSIL Limoges	Co-Directeur
Prof. Vahid MEGHDADI	ENSIL Limoges	Directeur



**Université
de Limoges**

*Pour mes parents
et mes frères.*

Remerciements

Tout d'abord je tiens à remercier Vahid MEGHDADI pour l'opportunité qui m'ait été offerte de pouvoir réaliser cette thèse sous sa direction. Sa gentillesse et la qualité de son encadrement ont facilité et rendu agréable la réalisation de cette thèse. J'ai beaucoup appris grâce à son expertise sur les techniques de traitement numérique du signal et sa forte perspicacité concernant différentes pistes que l'on a exploré ensemble. Je remercie Jean-Pierre CANCES, mon co-directeur de thèse, pour ses conseils en temps opportun et ses encouragements continuels.

Je voudrais aussi profiter de cette occasion pour remercier le directeur du laboratoire XLIM, Dominique BAILLARGEAT, et le responsable du département C2S2, Bernard JARRY, qui m'ont permis de pouvoir réaliser ma thèse dans le laboratoire XLIM.

Je suis également très reconnaissant envers Philippe GABORIT et Olivier RUATTA pour leur suivi technique concernant les codes à métrique rang.

Je tiens aussi à exprimer ma gratitude envers l'ensemble des membres de mon jury de soutenance : Andrea TONELLO, Professeur à l'Université d'Udine (Italie), Jean-Luc GAUTIER, Directeur Territorial Haute Vienne ERDF, Philippe GABORIT, Professeur à l'Université de Limoges, et en particulier Jean-Marie GORCE, Professeur à INSA de Lyon et Catherine DOUILLARD, Professeure à Télécom Bretagne qui m'ont fait l'honneur d'être rapporteurs de ma thèse. Merci d'avoir accepté d'évaluer ce travail et pour les commentaires pertinents et constructifs qu'ils ont fait sur ce manuscrit, le rendant de meilleure qualité et plus lisible.

Je remercie aussi tous mes collègues qui sont passés dans les locaux de l'équipe ESTE pendant la durée de ma thèse. Ce groupe a été pour moi une source de collaboration, de discussions passionnantes et d'amitié.

Enfin je remercie ma famille et mes amis pour leur amour et leurs encouragements.

Table des matières

Liste des figures	iii
Liste des tableaux	vii
Liste des algorithmes	ix
Glossaire	xii
1 Introduction générale	1
1.1 Contexte de l'étude	1
1.2 Objectif de la thèse	2
1.3 Organisation de la thèse	3
2 Réseaux CPL-BE dans le SG	7
2.1 Introduction	8
2.2 Introduction aux Courant Porteur en Ligne à Bande Étroite (CPL-BE) pour le Smart Grid (SG)	8
2.2.1 Nouveaux paradigmes des réseaux électriques	8
2.2.2 SG dans le réseau de distribution	10
2.2.3 CPL-BE pour les applications SG	12
2.3 Description des perturbations sur le canal CPL-BE	14
2.3.1 Propriétés et généralités du canal CPL-BE	14
2.3.2 Fonction de transfert du canal CPL-BE	17
2.3.3 Bruit sur le canal CPL-BE	21
2.4 Technologies et architecture CPL-BE	28
2.4.1 Architecture et topologie des CPL-BE dans le domaine d'accès	28
2.4.2 Spécification des standards et normes CPL-BE existants.	29
2.5 Conclusion	31

3	Codes correcteurs d'erreurs pour les réseaux CPL-BE	33
3.1	Introduction	34
3.2	Théorie des codes correcteurs d'erreurs	34
3.2.1	Notions essentielles en théorie des codes	34
3.3	Codes correcteurs d'erreurs pour les CPL-BE	40
3.3.1	Codes de Reed-Solomon	40
3.3.2	Codes convolutifs	42
3.3.3	Concaténation série de codes	44
3.3.4	Code à métrique rang	46
3.3.5	Protocole ARQ	48
3.3.6	Combinaison entre le codage correcteur d'erreurs et les protocoles de retransmission	49
3.4	Codes fontaines	51
3.4.1	Codes fontaines aléatoires	53
3.4.2	Codes LT	54
3.4.3	Codage réseau numérique	56
3.5	Conclusion	60
4	Codes à métrique rang pour les CPL-BE	61
4.1	Introduction	62
4.2	Description du système proposé	62
4.2.1	Description du mappage	62
4.2.2	Motifs d'erreurs	63
4.2.3	Construction d'un code Gabidulin sur $\mathbb{F}_{2^4}$	67
4.2.4	Code Gabidulin concaténé avec un code convolutif	68
4.3	Performances du schéma proposé	70
4.3.1	Résultats de simulation	70
4.3.2	Analyse et comparaison de la complexité des codes RS et Gabidulin	76
4.4	Conclusion	78
5	Relayage de codes fontaines pour les CPL-BE	79
5.1	Introduction	80
5.2	Relayage coopératif de codes fontaines : sans codage réseau	80
5.2.1	Codage coopératif	82
5.2.2	Stratégies de relayage	82
5.3	Codes fontaines combinés avec le codage réseau pour les CPL-BE	90
5.3.1	Revue de la littérature du codage réseau pour les CPL- BE	91
5.3.2	Optimisation de la matrice de probabilité jointe	101
5.3.3	Génération du degré de sortie	102

TABLE DES MATIÈRES

5.3.4	Algorithme de ré-allocation	103
5.3.5	Résultats de simulation	108
5.4	Conclusion	113
6	Conclusion générale et perspectives	115
6.1	Conclusion	115
6.2	Perspectives	116
	Bibliographie	120

Table des figures

2.1	Structure du réseau électrique européen	15
2.2	Illustration de la segmentation du réseau.	19
2.3	Exemple de trace de l'atténuation en fonction de la fréquence.	21
2.4	Bruit sur le canal CPL-BE	22
2.5	Variance de l'amplitude instantannée du bruit selon le modèle de Katayama calibré pour deux environnements A et B	25
2.6	Exemples de réalisations du bruit selon le modèle de Katayama calibré pour deux environnements A et B	25
2.7	Spectrogramme du bruit selon le modèle de Nassar calibré pour une zone donnée.	26
3.1	Modélisation de la chaîne de transmission avec un encodeur et décodeur	39
3.2	Exemple de représentation d'un code LDPC avec le graphe de Tanner	46
3.3	Les courbes du TEB en fonction de la longueur des trames transmises pour les environnements A et B	50
3.4	Distributions d'entrée et de sortie représentées avec le graphe bipartite des codes LT	54
3.5	Exemple de mise en œuvre de l'algorithme de propagation de croyance pour le décodage LT sur un canal à effacement	56
3.6	Représentation d'un code raptor	57
3.7	Codage réseau au niveau d'un nœud	58
4.1	Chaîne de transmission du code à métrique rang Gabidulin	64
4.2	Occurrence d'erreurs dûs au bruit impulsif ou au bruit à bande étroite	65
4.3	Motifs d'erreurs criss-cross rencontrés sur le canal CPL-BE	66
4.4	Schéma d'un code concaténé : code Gabidulin extérieur et convolutif intérieur.	68
4.5	Code produit intérieur appliqué aux matrices transmises	69

4.6	Mappage des symboles du code Reed-Solomon (RS) pour la transmission multi-porteuses	69
4.7	TEB du code Gabidulin comparé au code RS avec différents nombres de sous-porteuses affectées par l'interférence à bande étroite.	71
4.8	TEB du code Gabidulin comparé au code RS avec différents nombres de symboles OFDM affectés par le bruit impulsif . . .	72
4.9	TEB du code Gabidulin comparé au code RS avec différents nombres de sous-porteuses affectées par l'interférence à bande étroite et en présence de bruit impulsif	73
4.10	TEB des codes concaténés Gabidulin-Convolutif et RS-Convolutif en présence uniquement du bruit de fond	75
4.11	TEB des codes concaténés Gabidulin-Convolutif et RS-Convolutif en présence de différents nombres de symboles OFDM affectés par le bruit impulsif	75
4.12	TEB des codes concaténés Gabidulin-Convolutif et RS-Convolutif en présence de différents nombres de sous-porteuses affectées par l'interférence à bande étroite	76
5.1	Relayage de codes fontaines sur un réseau mono-chémin	81
5.2	Loi de probabilité du nombre requis de transmissions (pour décoder) pour une transmission directe, ou avec l'utilisation d'un ou de deux relais ($K = 20$).	88
5.3	Taux d'erreurs sur les trames de PRIME avec différentes distances	89
5.4	Probabilité d'échec du décodage en fonction de l'overhead sur un canal CPL pour différents schémas de transmission ($K = 20$)	89
5.5	La distribution de nœuds inspiré du modèle test IEEE 34 nœuds	92
5.6	(a) Modèle de collecte de données CPL-BE pour le SG , (b) Représentation équivalente comme dans le modèle test IEEE-34-nœuds	93
5.7	Topologie d'un réseau en Y	96
5.8	Topologie d'un réseau linéaire	96
5.9	Taux de succès du décodage pour des codes LT en terme d'overhead lorsque les paquets codés ne sont pas choisis de façon aléatoire et uniforme (avec K paquets source, K_1 paquets au relais, $c = 0.05$ et $\delta = 0.5$)	98
5.10	Taux de succès du processus de décodage pour les matrices de fusion $\mathbf{P}_o$ et $\mathbf{P}^*$ en terme d'overhead (avec $K_1 = K_2 = 250$, $c = 0.05$, $\delta = 0.5$)	109

TABLE DES FIGURES

5.11	Taux de succès du processus de fusion, en fonction de l'overhead, comparé à différentes stratégies dans l'état de l'art (avec $K_1 = 250$, $K_2 = 250$, $c = 0.05$, $\delta = 0.5$)	110
5.12	Taux de succès du processus de fusion proposé en fonction de l'overhead (avec $K_1 = 450$, $K_2 = 150$, $c = 0.05$, $\delta = 0.5$)	111
5.13	Taux de succès du processus de fusion proposé en fonction de l'overhead pour un scénario mis en cascade (avec 3 sauts $K_1 = 100$, $c = 0.05$, $\delta = 0.5$)	112

Liste des tableaux

2.1	Paramètres de la fonction de transfert hybride selon [1].	21
2.2	Paramètres du modèle de Katayama calibré pour deux environnements différents	24
2.3	Comparaison des principaux paramètres des spécifications de produits et standards CPL-BE multi-porteuses existants . . .	30
4.1	Analyse des complexités de décodage des codes RS et Gabidulin [2]	77

Liste des algorithmes

1	Algorithme de fusion des packets au relais	104
2	Calcul de P_o , la matrice de fusion	105

Glossaire

Liste des acronymes

AMM Automated Meter Management	11
AMR Automated Meter Reading	11
ARIB Association of Radio Industries and Businesses	13
ARQ Automatic reQuest Response.....	48
BBAG Bruit Blanc Additif Gaussien.....	43
BSC Binary Symmetric Channel	43
CENELEC Comité Européen de Normalisation en Électronique et en Électronique.....	13
CPL-BE Courant Porteur en Ligne à Bande Étroite	i
CPL-BUE Courant Porteur en Ligne à Bande Ultra Étroite	12
CRC Cyclic Redundancy Check.....	51
DLT Distributed LT	93
DSR Distribution de Soliton Robuste	93
DR Demande Response	11

SR Soliton Robuste	83
FCC Federal Communications Commission	13
FEC Forward Error Correction	29
FFT Fast Fourier Transform	30
HA Home Automation	11
HAN Home Area Network	13
LDPC Low Parity Check Code	34
LT Luby Transform	54
LUT Look Up Table	112
MAC Medium Access Control	29
MAP Maximum A Posteriori	43
NAN Neighbour Area Network	13
PEV Plug-in Electric Vehicles	11
PRIME Powerline Related Intelligent Metering Evolution	29
RF Radio Fréquence	12
RLNC Random Linear Network Coding	57
RS Reed-Solomon	vi
RSB Rapport Signal sur Bruit	87

Liste des acronymes

SCADA Supervisory Control and Data Acquisition	9
SDLT Selective Distributed LT.....	94
SG Smart Grid	i
SLRC Soliton Like Rateless Codes	94
SR Soliton Robuste	83
TEM Transverse Electro Magnétique.....	18
UIT Union Internationale des Télécommunications	29
WAN Wide Area Network	13
XOR eXclusive OR	93
RC Rank Code.....	70

Liste des symboles

q	puissance d'un nombre premier.
$\mathcal{C}(n, k)$	code en bloc de dimension k , de longueur n .
$Q(\cdot)$	fonction d'erreur de Gauss.
$\mathbb{F}_{q^m}$	le corps fini (corps de Galois) à q^m éléments.
$\ \cdot\ $	norme L^2 d'un vecteur.
$\mathbb{F}_q^k$	ensemble des k -upples d'éléments de $\mathbb{F}_q$.
$\mathbb{F}_q[z]$	le polynôme de l'indéterminée z défini dans $\mathbb{F}_q$.
$P_e^{S_i S_j}$	indique le taux d'effacement entre les noeuds S_i et S_j .
$ \mathcal{E} $	désigne le cardinal de $\mathcal{E}$.
$\mathbb{1}_n$	représente un vecteur $n \times 1$ de 1.
$triu(\mathbf{P})$	renvoie la partie triangulaire supérieure de $\mathbf{P}$.
$\mathbb{1}_A$	représente la fonction indicatrice de domaine de définition A .
$\mathcal{N}(z; 0; \gamma_k)$	dist. de prob. d'une v.a. Gaussienne de moyenne 0 et de variance γ_k .
$\mathbb{E}_X[x]$	espérance mathématique de la v.a. x .
$rand()$	renvoie un nombre aléatoire dans $[0,1]$, choisi de façon uniforme.
$\lfloor x \rfloor$	la partie entière par défaut de x .
$v(0 : n)$	notation indicielle pour les éléments $v(0), \dots, v(n)$ du vecteur v .

Chapitre 1

Introduction générale

Ce chapitre constitue l'introduction générale du manuscrit. Nous commençons par décrire la problématique scientifique de l'étude : la fiabilisation des communications par CPL-BE dans les SG. Ensuite nous présentons l'organisation de la thèse et ses différentes contributions originales.

1.1 Contexte de l'étude

Un SG est un réseau électrique qui intègre les technologies de l'information et de la communication (TIC) pour permettre une amélioration de son fonctionnement. Cette mise à jour du réseau électrique devra répondre par la même occasion à de nouveaux enjeux. Il s'agit de l'adaptation vis-à-vis de la transition énergétique, qui nécessite l'intégration adéquate d'énergies renouvelables produites de manière intermittente et décentralisée. Et aussi la maîtrise de la consommation itinérante avec par exemple l'introduction à grande échelle de voitures électriques [3].

L'implémentation d'un SG nécessite un ensemble hétérogène de supports physiques de transmission étant donné qu'aucune solution n'est adaptée à tous les scénarios de communication [4]. La liste est longue des technologies de communication sans fil et filaires complémentaires et parfois concurrentes qui pourront être utilisées dans le déploiement des liens de communication dans les SG depuis le réseau cœur jusqu'au réseau d'accès. Les communications

par Courant Porteur en Ligne à bande étroite (CPL-BE)¹ sont viables et attrayantes comme médium de communication pour les SG au niveau du réseau d'accès [5] comme en témoigne la pléthore de standards CPL-BE [6]. L'avantage indiscutable des CPL étant la possibilité de fournir une infrastructure qui est beaucoup plus vaste et plus omniprésente que toute autre solution filaire ou même sans fil.

Cependant, le canal CPL présente des caractéristiques très défavorables à la transmission de données en raison du fait qu'il n'a pas été conçu pour propager des signaux au delà de 1 kHz. En effet le canal CPL est affecté par plusieurs déficiences telles qu'une forte atténuation, une sélectivité en temps et en fréquence et un bruit non Gaussien dominé par une composante impulsive périodique [7]. Ceci constitue l'un des principaux facteurs qui limitent la performance des communications des systèmes CPL-BE.

C'est dans ce contexte que s'inscrivent les travaux de cette thèse, où nous proposons des techniques de codage correcteur d'erreurs pour les CPL-BE.

1.2 Objectif de la thèse

Afin de permettre une communication fiable (malgré une limitation de puissance²) sur un canal bruyant, pratiquement tous les systèmes de communication sont pourvus de codes détecteurs et correcteurs d'erreurs. Dans cette thèse, nous étudions l'application de ces codes aux communications par CPL-BE pour les SG. L'implémentation efficace des codes correcteurs d'erreurs sur la couche physique des réseaux CPL mérite une attention particulière, du fait que les perturbations (bruit) sur le canal CPL-BE sont inhabituelles [9].

Le canal CPL peut être considéré comme un canal à effacement de trames du fait, entre autres, de la présence de bruit. Donc, en plus de la correction d'erreurs directe sur la couche physique (mécanismes de type Forward Error

1. La bande étroite se réfère généralement à la gamme de fréquences de 3 kHz à moins de 500 kHz

2. Les modems CPL ont des contraintes strictes de puissance : ils fonctionnent généralement avec une puissance inférieure à celle d'une ampoule électrique [8].

Correction), un protocole de retransmission sur les couches supérieures est un moyen d'assurer des communications fiables. Ce schéma peut conduire à un nombre élevé de retransmissions. Ceci justifie l'usage des codes fontaines [10] à la place des protocoles de retransmission traditionnels. En outre, en raison de la nature multi-sauts des réseaux CPL-BE pour le SG, ceux-ci pourraient bénéficier de protocoles de communication coopératifs et du codage réseau [11, 12]. Aussi nous verrons comment appliquer ces techniques aux CPL-BE pour les communications SG. Nous effectuons une optimisation conjointe du codage correcteur d'effacements et du codage réseau dans le but d'augmenter le débit tout en maintenant une complexité acceptable.

1.3 Organisation de la thèse

Le corps du manuscrit de thèse est organisé en quatre chapitres :

- Les deux premiers chapitres présentent les concepts qui sont nécessaires à la compréhension des résultats dans la suite de ce manuscrit. Tout d'abord, dans le deuxième chapitre, nous introduisons le contexte technologique de la thèse c'est-à-dire les communications par CPL-BE pour le SG. Nous exposons les contraintes physiques du canal CPL-BE ; celui-ci sera caractérisé précisément et les modèles qui le décrivent seront présentés. Nous terminons ce chapitre par une présentation de l'architecture et des standards des systèmes de transmission par CPL-BE pour le SG.
- Ensuite après une introduction sur la théorie du codage, le chapitre III présente les différents types de codage, au niveau de la couche physique et des couches plus hautes, qui ont été étudiés dans le cadre du travail de thèse et qui sont utilisés dans les chapitres suivants : codes correcteurs en blocs, codes convolutifs, codes fontaines, codage réseau.
- Le quatrième chapitre concerne les transmissions point-à-point dans les CPL-BE. Nous nous intéressons aux types de bruit présents dans ces systèmes : interférence à bande étroite, bruit impulsif et bruit de fond qui entraînent des motifs d'erreurs criss-cross sur la nappe temps-fréquence lors des transmissions multi-porteuses. Nous abordons dans

ce chapitre le design et les performances d'un schéma robuste face aux motifs d'erreurs de type criss-cross rencontrés dans le CPL-BE. Ce schéma est constitué d'un code en bloc Gabidulin, basé sur la métrique rang, concaténé avec un code convolutif. Les performances du code Gabidulin sont d'abord comparées à celles du code RS. Si le code Gabidulin ne présente pas d'intérêt particulier en présence du bruit de fond seul, il se comporte mieux que le code RS lorsque le nombre de sous-porteuses affectées par l'interférence à bande étroite et/ou le nombre de symboles OFDM affectés par le bruit impulsif deviennent suffisamment élevés. Dans un second temps, les performances des deux codes sont comparées dans un schéma de concaténation avec un code convolutif, la concaténation (RS + convolutif) ayant été adoptée dans la plupart des standards de transmission CPL-BE. Cette partie constitue la première contribution du travail de thèse et a été publiée dans [13].

- Le cinquième chapitre s'intéresse aux communications CPL-BE avec relais, où la transmission d'une information entre une source et un destinataire peut être améliorée grâce à l'usage d'un nœud relais. Ce chapitre contient en réalité deux parties relativement indépendantes. La première partie porte sur l'application du codage coopératif pour une communication CPL à plusieurs sauts, ici les codes fontaines sont préconisés pour le codage coopératif. Le protocole de relayage considéré est le Decode-and-Forward coopératif en ce sens qu'un nœud relais situé entre la source et la destination, qui bénéficie d'un meilleur canal, accumule les paquets reçus. Lorsqu'il est en mesure de décoder, il remplace la source dans la transmission vers la destination. Les performances théoriques pour une liaison à trois sauts sont étudiées et validées par simulations. Les conclusions de cette partie ont été publiées dans [14].

Dans la deuxième partie, nous présentons des algorithmes pour combiner le codage fontaine et le codage réseau pour les réseaux CPL-BE. La combinaison de ces deux techniques entraîne la construction de codes fontaines distribués.

Nous considérons le cas où le relais est également une source. Il

doit donc décider de combiner ses propres paquets avec ceux de la source initiale. Nous concevons soigneusement les algorithmes dans les nœuds relais afin d’assurer un bon décodage³ des codes fontaines en se conformant à certaines de leurs propriétés critiques qui maximisent la probabilité de décodage. Deux approches sont envisagées : la résolution en utilisant un outil d’optimisation convexe et la proposition d’une solution sous-optimale plus simple à mettre en œuvre. Les combinaisons de codes obtenues sont ensuite simulées et les performances obtenues sont comparées à des scénarios de référence et à des algorithmes de la littérature. Il s’avère que les deux approches proposées présentent des performances équivalentes à celles d’un code LT unique et se comportent mieux, en termes de taux de succès vs overhead, qu’un multiplexage conventionnel et que les algorithmes de fusion trouvés dans la littérature. Les conclusions de cette dernière partie ont été présentées à IEEE TON.

- Le cinquième chapitre conclut le manuscrit. Il résume le contenu de la thèse et ses principales contributions puis présente quelques pistes d’approfondissement de l’étude.

3. par propagation de croyance

Chapitre 2

Réseaux CPL-BE dans le SG

Sommaire

2.1	Introduction	8
2.2	Introduction aux CPL-BE pour le SG	8
2.2.1	Nouveaux paradigmes des réseaux électriques	8
2.2.2	SG dans le réseau de distribution	10
2.2.3	CPL-BE pour les applications SG	12
2.3	Description des perturbations sur le canal CPL-BE	14
2.3.1	Propriétés et généralités du canal CPL-BE	14
2.3.2	Fonction de transfert du canal CPL-BE	17
2.3.3	Bruit sur le canal CPL-BE	21
2.4	Technologies et architecture CPL-BE	28
2.4.1	Architecture et topologie des CPL-BE dans le domaine d'accès	28
2.4.2	Spécification des standards et normes CPL-BE existants.	29
2.5	Conclusion	31

2.1 Introduction

À cause de leur ancienneté et de l'apparition de nouveaux enjeux, les réseaux électriques doivent être modernisés en réseaux SG. Ceux-ci s'appuieront sur une infrastructure de communication pour leur gestion plus "intelligente". Aujourd'hui, les communications CPL-BE font déjà partie intégrante des débuts du déploiement du SG avec les compteurs communicants¹. Ils constituent à ce jour la technologie la plus adoptée pour la mise en œuvre de ces compteurs [15].

Ce premier chapitre qui traite de l'environnement des communications par CPL-BE pour les applications SG se décompose en trois sections.

Dans la première section, après avoir décrit le réseau électrique et présenté les limites de son fonctionnement actuel, le concept de SG est introduit.

Ensuite, dans une deuxième section, la caractérisation des CPL-BE en tant que moyen de transmission d'information est étoffée : les différentes modélisations usuelles décrivant le bruit et la fonction de transfert du canal de propagation des CPL-BE sont présentées.

Deux aspects des CPL-BE sont abordés dans la troisième section. Premièrement, la description des architectures et des topologies des réseaux CPL-BE pour le SG est effectuée. Ensuite, les normes et standards actuels CPL-BE sont présentés.

2.2 Introduction aux CPL-BE pour le SG

2.2.1 Nouveaux paradigmes des réseaux électriques

Le fonctionnement actuel des réseaux électriques est basé sur le principe suivant : des centrales produisent de l'électricité à partir de l'énergie fossile, nucléaire, éolienne.... Ensuite un ensemble de câbles électriques appelé réseau de transport transporte l'énergie en très haute tension (HT) depuis les endroits de production vers les premières zones de distribution. À la suite

1. Un compteur communicant est un compteur ayant la capacité de fournir de manière détaillée et précise le comportement de consommation d'un utilisateur du réseau à l'exploitant du réseau.

du réseau de transport, le réseau de distribution (MT pour moyenne tension, et BT pour basse tension) achemine l'électricité aux consommateurs finaux. L'interface entre les réseaux HT et MT (resp. MT et BT) est assurée par des équipements qui abaissent la tension : postes de transformation HT/MT (resp. MT/BT).

Un système de contrôle et d'acquisition des données Supervisory Control and Data Acquisition (SCADA) assure la surveillance et le contrôle du réseau. Ce centre de contrôle communique avec les différentes sous-stations du réseau MT et BT par le biais d'un backhaul [16].

Dans son architecture traditionnelle, le réseau électrique a une répartition inégale des dispositifs de communication et des moyens de contrôle. Le réseau de transport est muni de moyens de communication hauts débits pour permettre des fonctions de télé-gestion et de surveillance. À l'opposé, les réseaux de distribution, dont la fonction principale est exclusivement l'acheminement de l'énergie des transformateurs vers les consommateurs, n'ont pas été outillés de moyens de communication conséquents. Par ailleurs, les utilisateurs finaux munis d'appareils électro-ménagers constituent des charges que l'exploitant du réseau peut estimer mais ne peut pas piloter à distance. Ces utilisateurs ne sont pas actifs dans le fonctionnement du réseau.

Ainsi, après plusieurs décennies de lentes évolutions, nous assistons aujourd'hui à une nécessité de modernisation du réseau électrique. Cette nécessité est une conséquence (entre autres) de plusieurs faits :

- L'ancienneté du réseau électrique et l'augmentation constante de la consommation énergétique, qui appellent soit à un renforcement du réseau électrique, soit à une gestion plus efficace de celui-ci. La première solution n'est ni économiquement ni écologiquement viable.
- Le réseau électrique est rarement exploité au plus près de sa capacité réelle. Ceci est dû à la conjugaison de deux facteurs : une consommation qui n'est pas lisse et le mode de fourniture d'énergie actuel qui a tendance à être maintenu à une capacité en excès. Une connaissance plus fine du réseau électrique, surtout du réseau de distribution où se trouvent les consommateurs, pourra permettre de mieux maîtriser et de mieux "lisser" la demande [17].

- La transition énergétique qui nécessite l'intégration adéquate d'énergies renouvelables pour répondre à des enjeux économiques et écologiques liés à la crise énergétique. Le mode de production et de transport de l'énergie change, hier il était centralisé avec un sens unidirectionnel et aujourd'hui distribué et bidirectionnel [18].
- L'introduction de nouveaux modes de consommation : le consommateur devient actif avec la possibilité de moduler sa charge et de stocker de l'énergie. La consommation devient itinérante aussi, avec l'introduction à grande échelle de voitures électriques et/ou hybrides.
- Une nécessaire amélioration de la réactivité du réseau, en cas de pannes par exemple, et un apport en sécurité et en protection des données qui transitent sur le réseau électrique.
- Enfin, les TIC représentent aujourd'hui un tremplin permettant de fournir de nouveaux services aux différents acteurs du réseau électrique.

Tous ces facteurs bouleversent les schémas traditionnels de gestion du réseau électrique et appellent à une mutation profonde de ce réseau.

2.2.2 SG dans le réseau de distribution

Il y a une absence de consensus sur une définition unique et universellement admise du concept de « smart grid ». La définition retenue dans cette thèse est celle de [19] : « un SG est un réseau électrique qui peut intégrer les actions de tous les utilisateurs qui y sont connectés afin d'assurer de façon efficace, l'approvisionnement durable, à faibles pertes et avec un niveau de qualité et de sécurité élevés les consommateurs d'électricité ». Le concept étant défini, les moyens pour atteindre cet objectif concernent entre autres plusieurs aspects des TIC : la gestion de bases de données et la science des données de façon générale, le traitement de signal pour une transmission efficace de données

Aujourd'hui, les efforts pour la mise en place des réseaux SG sont principalement axés sur le développement du domaine des compteurs dits communicants². Les CPL-BE constituent à ce jour la technologie la plus

2. Un compteur communicant est un compteur ayant la capacité de fournir de manière

adoptée pour la mise en œuvre de ces compteurs [15]. Deux niveaux de dispositifs de comptage peuvent être distingués :

- La lecture automatisée des compteurs Automated Meter Reading (AMR) est destinée uniquement à la collecte des données des compteurs vers les centres de contrôle. L'AMR a évolué dans le sens de la réalisation des SG vers la gestion automatique des compteurs Automated Meter Management (AMM).
- L'AMM est un système de communication bidirectionnel par lequel les exploitants du réseau électrique peuvent obtenir des informations en temps réel sur la demande des clients et aussi encourager certains comportements de consommation. L'AMM permet aussi la télégestion de certaines tâches liés au comptage qui auparavant étaient manuelles [3, 5].

D'autres applications des SG concernent :

- La demande/réponse (Demande Response (DR)), qui permet de rendre les consommateurs actifs vis-à-vis de la gestion du réseau électrique. Par exemple, le réseau pourrait inciter les consommateurs à réduire leur consommation pour mieux gérer l'équilibre offre/demande.
- L'automatisation des maisons (Home Automation (HA)) par l'intermédiaire du réseau domestique et d'une intelligence embarquée dans certains appareils peut permettre une meilleure gestion d'énergie au niveau local.
- La gestion des véhicules électriques (Plug-in Electric Vehicles (PEV)) se fait aussi dans le cadre d'applications smart grid.

Dans cette thèse nous nous sommes intéressé aux CPL-BE comme support de communication pour les applications smart grid au niveau des réseaux d'accès BT.

détaillée et précise le comportement de consommation d'un utilisateur du réseau à l'exploitant du réseau.

2.2.3 CPL-BE pour les applications SG

2.2.3.1 Bref historique des CPL

Cet aperçu historique met en lumière certains jalons expliquant la manière dont les CPL ont vu le jour et se sont développés.

L'idée de superposition d'un signal de télécommunication sur la forme d'onde principale du courant électrique remonte depuis la fin du 19^{ème} siècle. En 1838 Edward Davy proposait l'utilisation de l'infrastructure du réseau électrique, pour la surveillance du niveau de tension de batteries situées dans des sites distants non habités entre Londres et Liverpool, pour le système Télégraphique [20]. Le premier brevet connu, portant sur les communications par CPL a été enregistré en 1897. Les opérateurs des réseaux électriques ont utilisé les technologies CPL sur les réseaux de distributions HT pour leur propres systèmes de téléphonie depuis les années 20 [21]. Dans les années 30 le contrôle de charge dans l'objectif de moduler la demande d'énergie en période de pic employait la technologie CPL. Un moyen informel qui a permis d'accélérer la maturation des CPL aura été leur utilisation massive pendant la seconde guerre mondiale par certains radio-amateurs lorsque leurs activités sur le spectre Radio Fréquence (RF) avaient été limitées [22].

L'intérêt pour ce médium de communication s'est accru dans les années 80 et s'est intensifié surtout pendant les années 90. Et plus récemment depuis 2007 on assiste à la montée des technologies CPL multi-porteuses à bande étroite pour des applications smart grid [23, 6].

2.2.3.2 Classification des communications par CPL

Les réseaux CPL sont généralement classés en fonction des plages de fréquences de fonctionnement. Ainsi nous distinguons :

- Les Courant Porteur en Ligne à Bande Ultra Étroite (CPL-BUE) qui opèrent dans la bande de fréquence $[0.3 - 3]$ kHz et fournissent de l'ordre de 100 bit/s sur de longues distances (plus de 150 km). Elles offrent un débit insuffisant pour les futures communications SG [4].
- Les CPL-BE qui utilisent la gamme de fréquences $[3 - 500]$ kHz. En fonction des pays, les organismes de normalisation attribuent des

bandes différentes. La bande en dessous de 150 kHz en Europe avec la bande A du Comité Européen de Normalisation en Électronique et en Électronique (CENELEC) : 3 à 148,5 kHz. Pour les États-Unis, la commission fédérale des communications (Federal Communications Commission (FCC)) alloue la bande de 10 à 490 kHz. Pour le Japon, l'association japonaise des industries et entreprises de radio (Association of Radio Industries and Businesses (ARIB)) utilise la bande ([10 – 450] kHz). Et enfin la Chine utilise la bande [3 – 500] kHz.

Les CPL-BE sont en mesure de livrer quelques kbit/s en utilisant les communications mono-porteuses, ou plusieurs centaines de kbit/s en utilisant des communications multi-porteuses. Les CPL-BE et à haut débit ont suscité énormément d'intérêt en tant que moyen de communication pour les applications de télémétrie et pour les SG en général [8].

- Pour les CPL à Large Bande (CPL-LB) et à haut débit, la borne inférieure de la bande de fréquence est typiquement de 1 MHz, et la borne supérieure peut aller jusqu'à 30 ou 60 MHz. L'utilisation de fréquences supérieures, de l'ordre de centaines de MHz a fait l'objet d'étude et d'implémentation [24].

Nous distinguons également dans cette section plusieurs types de réseaux télécoms supportés par le réseau électrique en fonction de leur localisation et de leur taille :

- Les réseaux privés domestiques Home Area Network (HAN) qui se situent derrière le compteur électrique et relient les appareils et les capteurs sur les lignes électriques à l'intérieur des bâtiments.
- Les réseaux de quartier Neighbour Area Network (NAN) qui relient les compteurs intelligents aux concentrateurs de données du réseau électrique.
- Les réseaux étendus Wide Area Network (WAN) qui sont en fait des liens de communication acheminant le trafic entre les concentrateurs de données et les centres de contrôle du réseau électrique.

☞ Remarques :

- En Europe, les concentrateurs de données sont généralement déployés sur les lignes BT. La communication entre les concentrateurs de données et les compteurs n'a donc pas besoin de traverser les transformateurs MT-BT.
- Parce qu'elles sont basses fréquences ($3 - 500$ kHz), certaines technologies CPL-BE sont capables de communiquer à travers des transformateurs au prix d'un rapport signal sur bruit assez élevé [5].

La figure 2.1 présente les différents segments d'une topologie typique de réseau électrique européen. Dans cette thèse, nous nous sommes intéressés à la partie du réseau reliant les compteurs intelligents aux concentrateurs de données, c'est la dire le réseau d'accès.

2.3 Description des perturbations sur le canal CPL-BE

Dans cette partie nous décrivons le comportement du canal CPL-BE. L'idée étant de présenter des modèles de canal suffisamment précis pour évaluer l'impact réel du canal sur les performances des communications.

La première section présente des généralités et des propriétés des canaux CPL. Dans la deuxième section, le canal de propagation est modélisé suivant deux approches, l'une basée sur le modèle paramétrique de l'évanouissement multi-trajets et l'autre sur les méthodes standards de la théorie des lignes de transmission. Ensuite dans une troisième section, nous présentons le bruit du canal CPL-BE.

2.3.1 Propriétés et généralités du canal CPL-BE

Nous rappelons ici quelques propriétés générales des canaux CPL :

- Une propriété importante du canal CPL qui a été prouvée mathématiquement et expérimentalement [25] est sa symétrie.
- Une autre caractéristique du canal CPL est son comportement variant dans le temps. Les consommateurs d'électricité, dans le cadre d'une utilisation normale du réseau, et leurs appareils électriques commutent

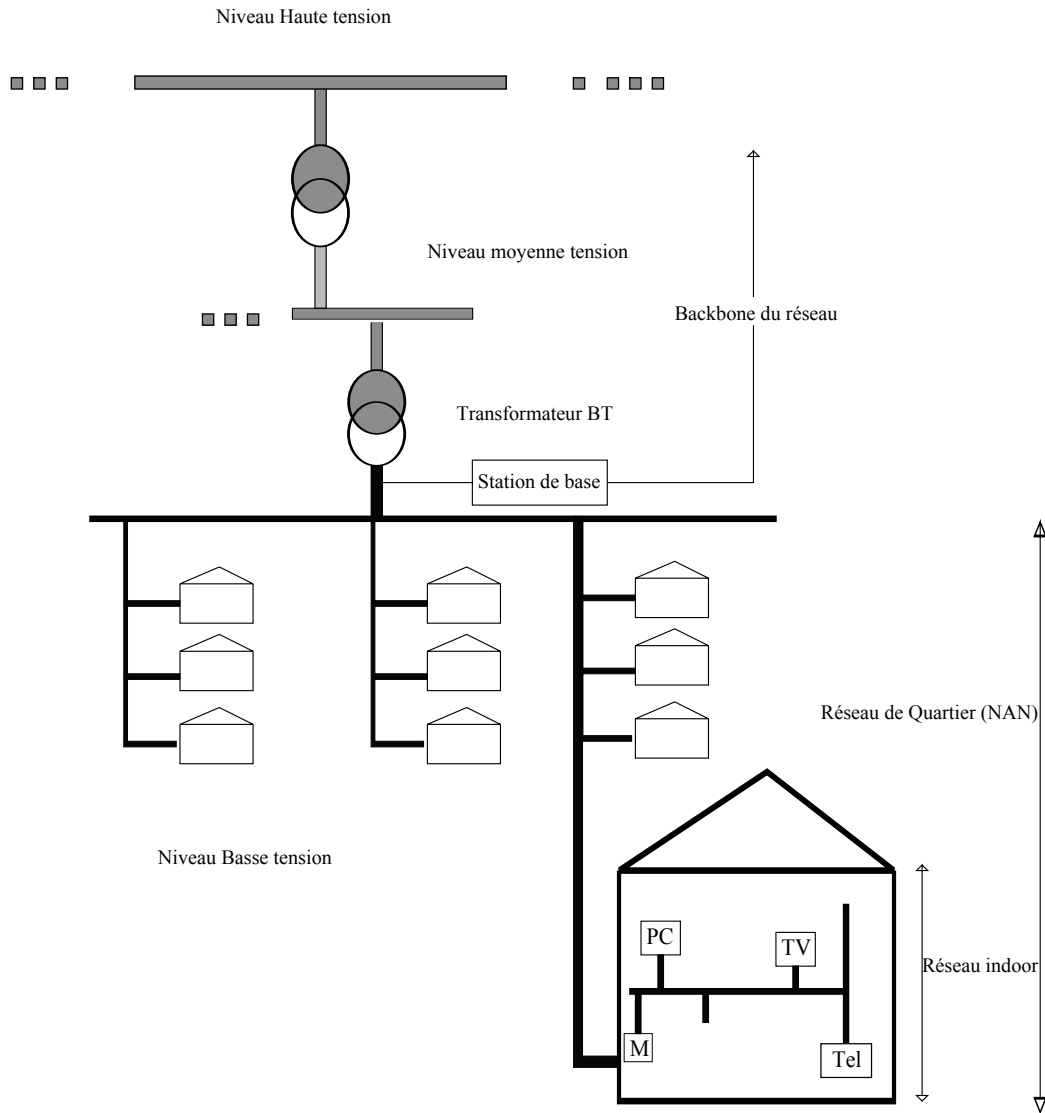


FIGURE 2.1: Structure du réseau électrique européen

les charges d'une manière incontrôlée et imprévisible, provoquant des fluctuations d'impédance soudaines et à grande dynamique. Ces changements entraînent des variations lentes du canal. En outre, le canal CPL présente des variations synchrones à la fréquence du secteur car les paramètres haute fréquence des appareils électriques qui sont connectés au réseau dépendent de l'amplitude instantanée de la tension du réseau. Le bruit injecté dans la ligne CPL par les appareils connectés

au réseau dépend lui aussi de l'amplitude instantanée de la tension du réseau. Par conséquent, le bruit et la réponse fréquentielle du canal présentent un comportement cyclo-stationnaire³ en temps [23].

- Les caractéristiques des canaux CPL-BE dépendent des pratiques de câblage, de la topologie et de la densité des charges connectées au réseau électrique. Ces paramètres varient d'un pays à un autre et pour un même pays ils changent en fonction du segment du réseau considéré (HT, MT ou BT), de l'environnement dans lequel l'électricité est distribué (rural ou urbain, résidentiel ou industriel) et indépendamment de ces facteurs également avec le temps. Malheureusement, il n'existe pas de canaux de référence pour les communications CPL-BE dans la partie basse tension⁴.

Le canal CPL-BE peut être complètement décrit par sa fonction de transfert et le bruit additif qui l'affecte.

2.3.1.1 Atténuation

L'atténuation croît avec la fréquence, elle augmente aussi avec la distance, mais est surtout impactée par la densité de charges sur le réseau électrique [27]. L'atténuation serait normalement d'impact « mineur » pour assurer une communication sans erreur. Aussi longtemps que celle-ci est connue, et qu'un module d'amplification approprié permet d'assurer un bon niveau de puissance du signal à la réception.

Cependant, il y a en Europe un niveau de signal maximal et une densité spectrale de puissance limitée par la norme EN 50065-1. Dans la bande FCC, alors qu'il n'y a pas de limite explicite, le comportement non-linéaire des amplificateurs utilisés impose une limite pratique.

2.3.1.2 Impédance d'accès du réseau CPL-BE

L'impédance d'accès Z_A , ensemble avec l'impédance équivalente de sortie du modem CPL Z_T , bien qu'ils n'aient pas une influence directe sur la qualité

3. périodiquement stationnaire

4. Un canal de référence est un ensemble de paramètres qui permet de décrire une réalisation de canal représentative des caractéristiques réelles du canal [26].

de propagation du signal, ils déterminent le niveau de perte de couplage des transmissions sur le réseau CPL. En fait, pour un transfert de puissance optimal entre les ports de sortie du modem CPL et la ligne électrique, l'impédance de sortie du modem et l'impédance d'entrée de la ligne électrique doivent être adaptées. L'impédance d'accès dépend des nœuds (impédances) connectés, elle varie avec le temps, la fréquence et l'endroit de mesure. Toutes ces caractéristiques rendent l'adaptation d'impédance non triviale.

2.3.2 Fonction de transfert du canal CPL-BE

Les effets de distorsion du canal comprennent les aspects usuels d'atténuation mais aussi de sélectivité fréquentielle à cause des désadaptations rencontrées sur les lignes de transmission. Dans cette section, nous introduisons trois approches populaires pour la modélisation de la fonction de transfert des canaux CPL.

2.3.2.1 Modèle paramétrique : « top-down » :

Elle a d'abord été proposée par Phillips dans [28] et plus tard améliorée et validée par Zimmermann et Dostert dans [29]. L'idée ici étant de fournir une fonction paramétrique qui décrit tous les aspects du comportement du canal CPL : ses caractéristiques passe-bas, son atténuation, sa sélectivité fréquentielle. La fonction de transfert paramétrique s'écrit :

$$H(f) = \sum_{i=1}^N g_i(f) e^{-\alpha(f)l_i} e^{-j2\pi f\tau_i}, \quad (2.1)$$

où N est le nombre de trajets dit distinguables, g_i est le gain complexe sur chaque trajet, l_i la longueur du i -ème trajet, τ_i le délais sur le i -ème trajet et $\alpha(f)$ représente l'atténuation en fonction de la fréquence.

$$\alpha(f) = a_0 + a_1 \cdot f^k, \quad (2.2)$$

où a_0 , a_1 et k sont les paramètres de l'atténuation, $\alpha(f)$ prend en compte les effets de peau et de perte diélectrique [29]. Le modèle nécessite des campagnes

de mesures de la réponse du canal pour déterminer les paramètres N , $\alpha(f)$ et (g_i, l_i, τ_i) . Ce modèle nécessite beaucoup de paramètres, spécialement quand N est grand et est sensible à la taille de N en termes de précision. Par exemple, des valeurs de N comprises entre 15 et 44 sont nécessaires pour un ajustement d'un lien de 110m [9, 29]. Enfin, le modèle est peu lié à la réalité physique du réseau qui souvent peut être complexe. Et une fois défini il est peu flexible, en d'autres termes pour un réseau donné il est sensible à tout changement de topologie. Cependant, cette approche est utile lorsque la topologie du réseau n'est pas disponible.

2.3.2.2 Modélisation structurelle : « bottom-up » :

« Bottom-up » en anglais se traduit littéralement de bas en haut pour indiquer que la fonction de transfert s'obtient depuis les caractéristiques basses (physiques) du réseau électrique. Cette approche est basée sur la théorie des lignes de transmission à deux ou à plusieurs conducteurs. Celle-ci nécessite une connaissance de la structure complète du réseau (les paramètres linéiques des câbles électriques, la topologie du réseau, les impédances connectées au réseau...). La théorie des lignes de transmission permet l'analyse de la propagation d'ondes en milieu filaire. Elle se base sur la décomposition de la ligne en plusieurs éléments infinitésimaux décrits chacun par les équations des télégraphistes [30]. La théorie des lignes de transmission à deux conducteurs est incapable de décrire parfaitement tous les phénomènes physiques se produisant dans les câbles électriques dans la pratique (comme le couplage de mode). Cependant, en première approche, elle donne un aperçu sur une grande majorité des caractéristiques du canal de propagation CPL [23].

Dans cette partie, nous décrivons la fonction de transfert des liens CPL en suivant la théorie des lignes de transmission à deux conducteurs. Pour une analyse complète des réseaux CPL suivant la théorie des lignes de transmission multi-conducteurs, voir [31]. Pour les fréquences considérées dans les CPL-BE, le mode principal se propageant sur les lignes électriques est le mode Transverse Electro Magnétique (TEM) ou quasi-TEM [30]. En

fonction de la technique utilisée, la propagation des signaux sur les lignes de transmission peut être décrite par les paramètres S, le rapport de tension, les paramètres ABCD.

L'approche du rapport de tension, permet de calculer la fonction de transfert $H(f)$ comme étant le rapport entre la tension du récepteur V_{Rx} et la tension de l'émetteur V_{Tx} .

$$H(f) = \frac{V_{Rx}(f)}{V_{Tx}(f)} \quad (2.3)$$

La fonction de transfert entre un émetteur et un récepteur situés sur la ligne électrique se calcule de la façon suivante :

- D'abord il faut trouver le chemin (backbone) le plus court qui relie l'émetteur et le récepteur. Diviser la ligne en plusieurs segments élémentaires le long de ce chemin, de préférence au niveau des discontinuités du réseau.

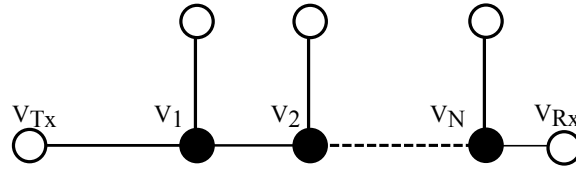


FIGURE 2.2: Illustration de la segmentation du réseau.

- Une seconde étape consiste à décrire la propagation des signaux dans un segment élémentaire du réseau électrique sur la base de la théorie des lignes de transmission. Chaque segment est supposé avoir des caractéristiques homogènes. On note la fonction de transfert de chaque segment par :

$$H_n(f) = \frac{V_{n+1}(f)}{V_n(f)}, \quad (2.4)$$

qui est le rapport entre la tension d'entrée et de sortie du segment. Les segments finaux sont ceux associés à l'émetteur et au récepteur ($V_{N+1} = V_{Rx}$ et $V_0 = V_{Tx}$).

- Enfin la fonction de transfert entre l'émetteur et le récepteur est calculée en mettant en cascade les différentes fonctions de transfert des

segments élémentaires intermédiaires. La fonction de transfert totale peut s'écrire :

$$H(f) = \frac{V_{N+1}}{V_0} = \prod_{n=0}^N H_n(f). \quad (2.5)$$

L'avantage de cette approche est qu'elle permet de rendre compte de la corrélation des fonctions de transfert pour des nœuds (émetteurs ou récepteurs) situés sur la même ligne de transmission.

2.3.2.3 Modélisation hybride :

La forme paramétrique définie plus haut peut donner lieu à une description statistique du canal en affectant à ses paramètres des lois de distribution statistiques. Ceci à été surtout accompli pour les canaux CPL-LB [32, 33]. Dans un modèle hybride, on pourrait définir un ensemble de topologies réalistes qui peuvent être considérées comme représentatives de la majorité de ce qui peut être trouvé dans la réalité. Sur la base de ces topologies et pratiques courantes, et à l'aide de la modélisation structurale, on peut définir un ensemble de fonctions de transfert pour les canaux CPL-BE dans la partie BT. Les lois de variations des paramètres de la fonction de transfert $H(f)$ (défini dans 2.1) sont données par l'organisme de standardisation IEEE dans [1] et définis comme :

- l_i est une variable aléatoire gaussienne de moyenne l_a et de déviation standard d_s qui représente la longueur des trajets de propagation.
- ν représente la vitesse de propagation de l'onde.
- a_0 et a_1 sont les paramètres d'atténuation qui dépendent des caractéristiques physiques de la ligne de transmission.
- g_i , le gain sur le trajet i , est une variable aléatoire gaussienne de moyenne nulle et de variance 1, qui est mis à l'échelle par 10000.
- N est le nombre de trajets significatifs.

Les paramètres l_i , l_a , l_s , l_m sont présentés dans le tableau 2.1, et sont uniquement pertinents pour les réseaux BT. La figure 2.3 illustre bien les phénomènes de fading qui sont une conséquence des multi-trajets causés par les désadaptation d'impédances sur la ligne électrique.

paramètres	valeurs	descriptions
l_a	1000	Trajet moyen entre Trx et Rx
l_s	400	déviation standard des trajets entre Trx et Rx
l_m	100	distance minimale des l_i
ν	$3 \cdot 10^4/4$	vitesse de propagation de l'onde
k	$1 \quad (0.5 \leq k \leq 1)$	pente de l'atténuation en fonction de la fréq.
a_0	$1e^{-3}$	dépend des caractéristique des lignes
a_1	$2.5e^{-9}$	dépend des caractéristiques des lignes
N	$5 \leq N \leq 50$	Nombre de trajets significatifs

Tableau 2.1: Paramètres de la fonction de transfert hybride selon [1].

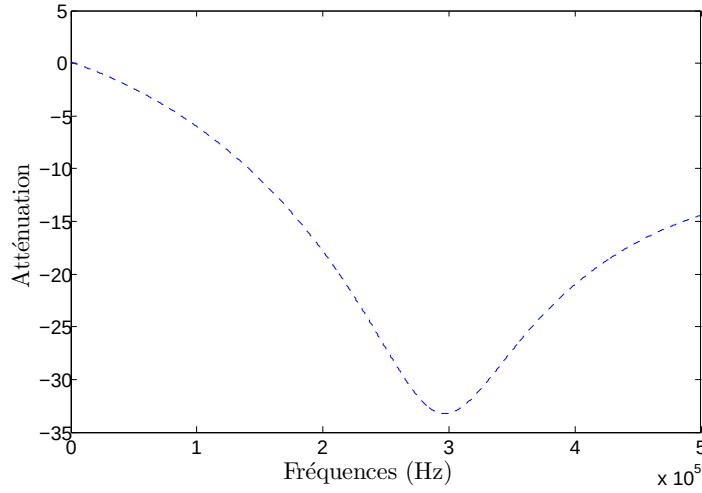


FIGURE 2.3: Exemple de trace de l'atténuation en fonction de la fréquence.

2.3.3 Bruit sur le canal CPL-BE

Le bruit sur les CPL-BE provient d'une part de la somme de plusieurs formes d'onde produites et émises sur les lignes par des appareils connectés au réseau électrique. Et d'autre part, du rayonnement d'autres sources opérant aux mêmes fréquences que les systèmes CPL-BE, qui ne sont pas connectés au réseau mais dont les émissions sont couplées aux systèmes CPL-BE. Le bruit résultant est non-blanc avec un contenu spectral qui présente une décroissance en fonction de la fréquence en raison de la diminution de la concentration des sources de bruit avec la fréquence [23]. De façon générale, le

bruit affectant les communications CPL est classé en trois catégories comme présenté dans la figure 2.4 : le bruit de fond coloré, le bruit impulsif et le bruit à bande étroite. Le bruit impulsif est généralement classé en trois catégories : le bruit impulsif périodique asynchrone à la fréquence du secteur principalement causé par les alimentations commutées. Le bruit impulsif périodique synchrone à la fréquence du réseau causé par la commutation des diodes de redressement d'alimentations. Et le bruit impulsif asynchrone principalement causé par des événements de commutation.

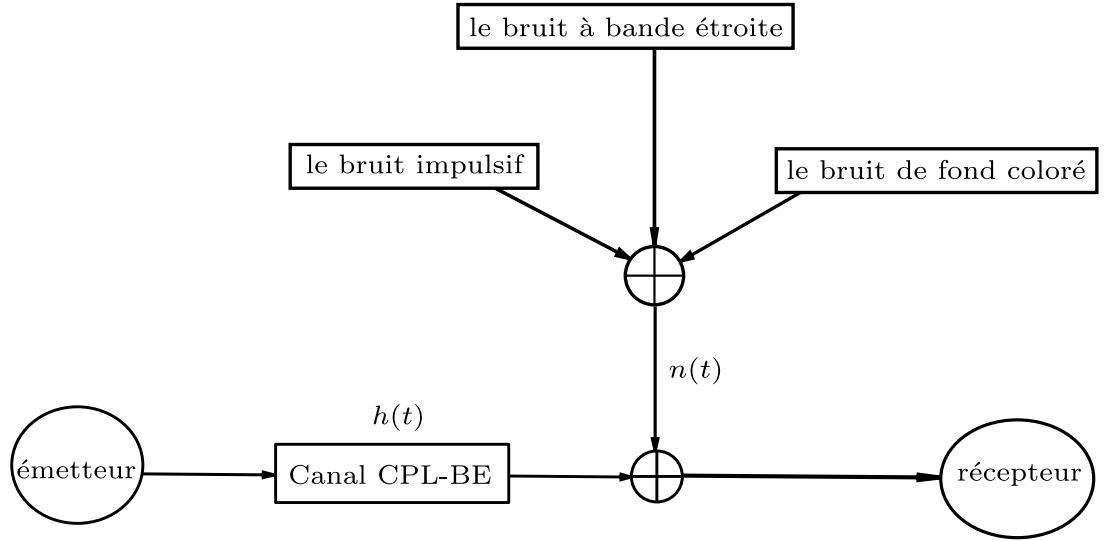


FIGURE 2.4: Bruit sur le canal CPL-BE

Différentes campagnes de mesure ont montré que la composante dominante du bruit sur le canal CPL-BE est le bruit impulsif synchrone à la tension du secteur. La fréquence des caractéristiques périodiques de l'enveloppe du bruit synchrone est la même que le double de la fréquence du secteur. La variance du bruit $\sigma^2(t)$ est alors une fonction périodique du temps :

$$\sigma^2(t) = \sigma^2(t + k \cdot T_{AC}/2) \quad (2.6)$$

où T_{AC} est la durée du cycle du courant alternatif, et $k \in \mathbb{Z}$. Le comportement du bruit CPL-BE a été caractérisé par les modèles cyclo-stationnaires de Nassar et Katayama [34, 35].

2.3.3.1 Modèle de Katayama et Al.

Le bruit est exprimé sous la forme d'un processus Gaussien dont la variance instantanée est une fonction temporelle périodique de période égale à la moitié du cycle du courant alternatif ; c'est-à-dire :

$$n_t \sim \mathcal{N}(0, \sigma^2(t, f)), \quad (2.7)$$

$$\sigma^2(t, f) = \sigma^2(t) \cdot \alpha(f), \quad (2.8)$$

$$\alpha(f) = \frac{a}{2} \exp(-a \cdot f), \quad (2.9)$$

$$\sigma^2(t) = \sum_{l=0}^{L-1} A_l \left| \sin\left(\frac{2\pi t}{T_{AC}} + \theta_l\right) \right|^{n_l}, \quad (2.10)$$

où f , la fréquence est en Hz. Les diverses composantes du bruit CPL-BE sont généralement décrites par la somme de 3 puissances de sinusoides :

- i) Le bruit invariant continu est pris en compte par la composante ($l = 0$) de l'équation.
- ii) Le bruit cyclo-stationnaire continu à variations lentes est pris en compte par la sinusoides correspondante à ($l = 1$).
- iii) Le bruit impulsif cyclo-stationnaire est décrit par la dernière partie de l'équation ($l = 2$).

☞ Remarques :

- L'avantage de cette modélisation est qu'elle permet d'avoir une expression exacte de la loi de probabilité du bruit des CPL-BE.
- Un des inconvénients majeurs de cette modélisation est lié au nombre potentiellement important des paramètres $3L$ pour approximer le bruit. En outre, le choix de la forme sinusoidale de la variance du bruit est arbitraire et ne provient pas des campagnes de mesures ; ce qui rend l'ajustement des paramètres du modèle particulièrement complexe quand la complexité du bruit augmente.
- Enfin le modèle de Katayama découple les variations spectrales et temporelles du bruit. On sait, grâce aux campagnes de mesure, que le comportement spectral du bruit varie avec le temps.

2.3 Description des perturbations sur le canal CPL-BE

l	A_l	$\Theta_l[deg]$	n_l		l	A_l	$\Theta_l[deg]$	n_l
0	0.13	-	-		0	0.23	-	-
1	2.8	128	9.3		1	1.38	-6	1.91
2	16	161	5.3×10^5		2	7.17	-35	1.57×10^5
(a) Paramètres pour l'environnement A					(b) Paramètres pour l'environnement B			

Tableau 2.2: Paramètres du modèle de Katayama calibré pour deux environnements différents

Les paramètres $\{A_l, \Theta_l, n_l\}_{l=0,1,2}$, calibrés pour deux environnements différents A et B⁵, sont décrits dans le tableau 2.2. Les figures 2.5 et 2.6 montrent un exemple des variations de la variance et de l'amplitude instantanée du bruit pour les paramètres définis dans le tableau 2.2.

2.3.3.2 Modèle de NASSAR et Al.

Les auteurs dans [9] ont proposé un modèle (qui a d'ailleurs été adopté par la norme IEEE 1901.2) dans lequel chaque période de bruit est divisée en M régions temporelles au cours desquelles le bruit est considéré comme coloré mais stationnaire. Ainsi, chaque région est caractérisée par une forme spectrale et un filtre de mise en forme correspondant. Selon ce modèle, le bruit CPL-BE est modélisé comme la convolution d'un bruit blanc Gaussien avec un filtre linéaire variant périodiquement dans le temps de réponse temporelle $h[k, \tau]$ et est exprimé comme suit :

$$n[k] = \sum_{\tau} h[k, \tau] s[\tau] = \sum_{i=1}^M \mathbb{1}_{k \in \mathcal{R}_i} \sum_{\tau} h_i[\tau] s[\tau], \quad (2.11)$$

avec $0 \leq k \leq N - 1$ et $\mathbb{1}_{\mathbb{A}}$ représente la fonction indicatrice de domaine de définition $\mathbb{A}$. Les filtres linéaires invariants dans le temps h_i correspondent aux filtres de mise en forme spectrale pour chaque région du spectrogramme. Ce modèle a besoin d'un grand nombre de coefficients pour les filtres (par exemple, pour atteindre une résolution de 1 kHz à un taux d'échantillonnage

⁵. Ces deux environnements correspondent à deux emplacements arbitraires du laboratoire des auteurs de [34]

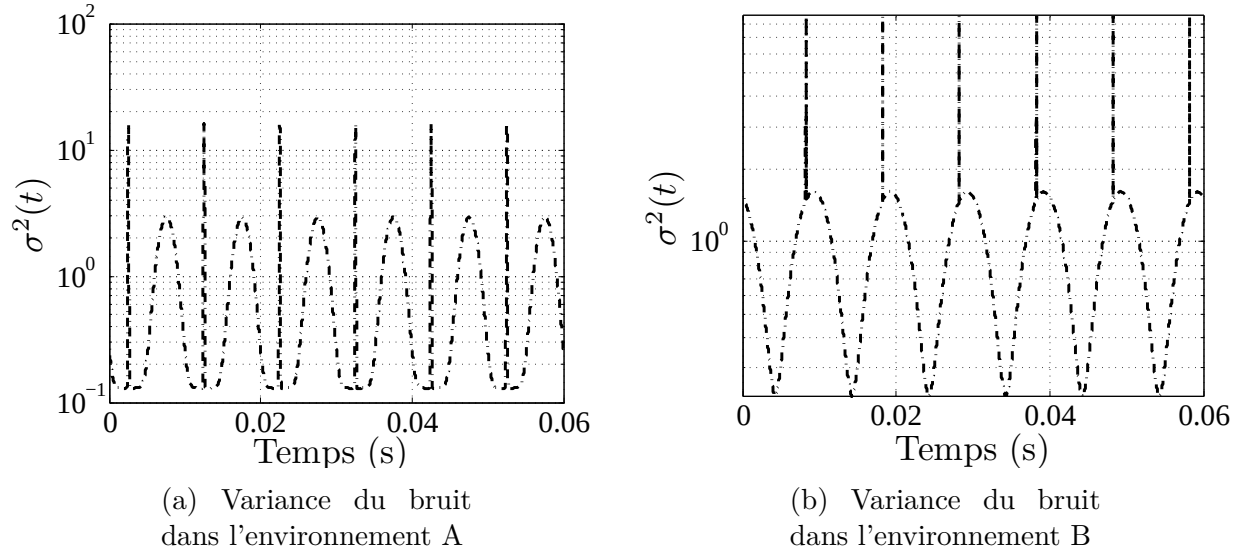


FIGURE 2.5: Variance de l'amplitude instantannée du bruit selon le modèle de Katayama calibré pour deux environnements A et B

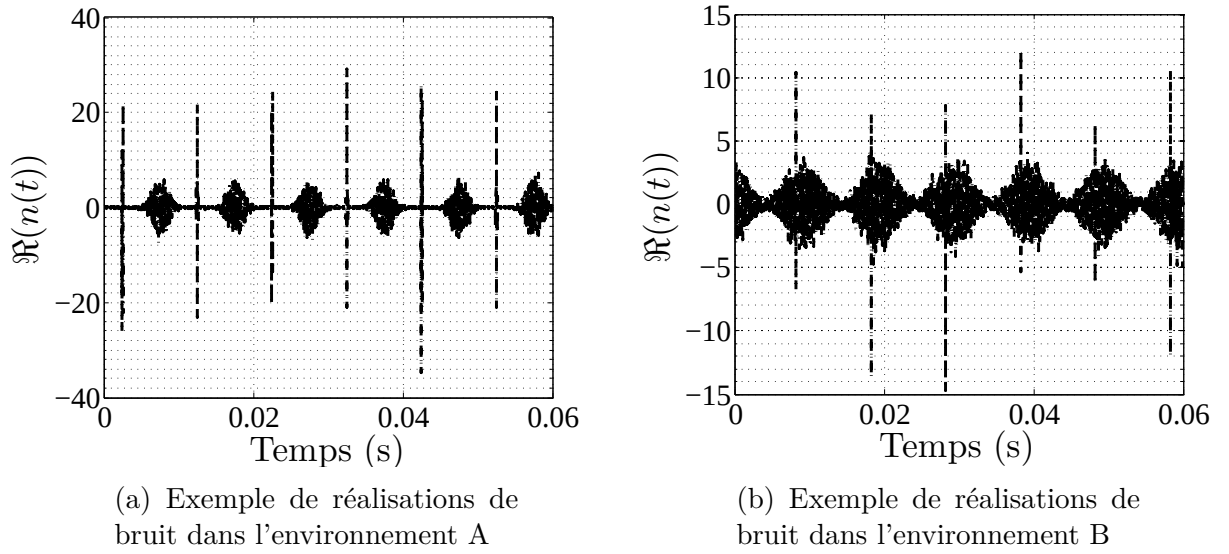


FIGURE 2.6: Exemples de réalisations du bruit selon le modèle de Katayama calibré pour deux environnements A et B

de 2 Mbps, un filtre numérique a besoin de 2000 coefficients [36]).

Un avantage de ce modèle est qu'il permet de capturer d'office une partie

des brouilleurs à bande étroite présents dans un site donné sans avoir besoin de les introduire explicitement.

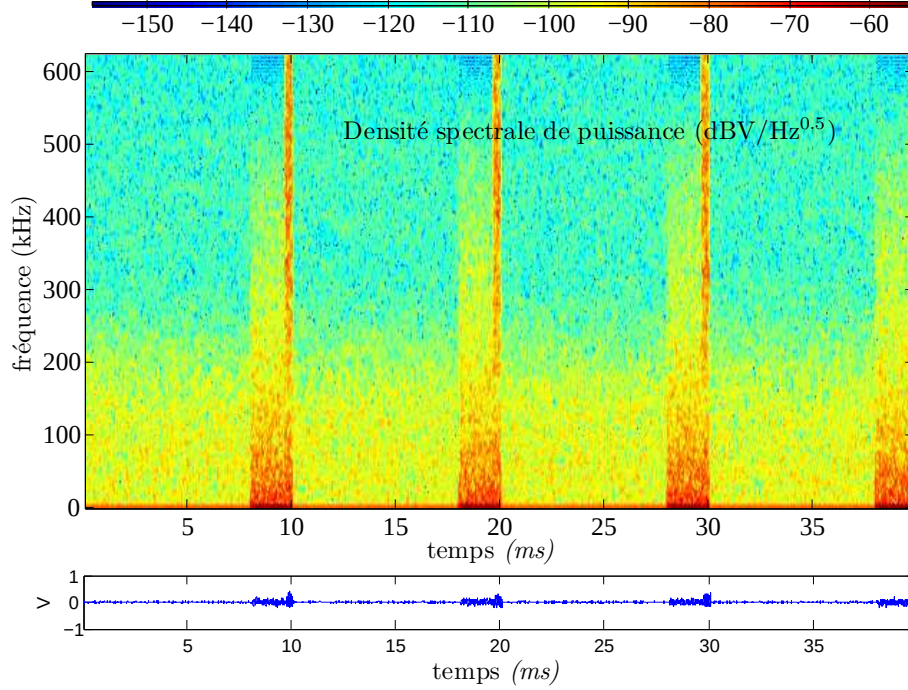


FIGURE 2.7: Spectrogramme du bruit selon le modèle de Nassar calibré pour une zone donnée.

La figure 2.7 présente un exemple de réalisation du bruit selon le modèle de Nassar calibré dans un environnement CPL-BE outdoor urbain.

2.3.3.3 Interférence apériodique

En absence d'une uniformisation des normes et standards, certains systèmes CPL peuvent causer des interférences sur les canaux CPL-BE. Comme l'interférence occasionnée peut être complètement aléatoire elle pourrait ne pas être totalement prise en compte par les modèles de bruit discutés ci-dessus.

Pour décrire les distributions statistiques de l'amplitude instantanée du bruit impulsif apériodique, différents modèles ont été utilisés dans la

littérature : le mélange de Gaussiennes, les distributions Middleton de catégorie A [37], le modèle de Markov [38], pour ne citer que quelques uns. En caractérisant les émissions aléatoires d'interférences dans un réseau CPL comme suivant temporellement un processus de Poisson, les auteurs dans [39] par dérivation analytique ont montré que les interférences vues depuis un récepteur dans un réseau CPL peuvent être modélisées par un mélange de Gaussiennes ou un modèle Middleton de classe A.

Definition 1. *Modèle de mélange de Gaussiennes : une variable aléatoire Z suit une distribution correspondante à un mélange de Gaussiennes si sa loi de probabilité est la somme pondérée de différentes distributions Gaussiennes, à savoir,*

$$f_Z(z) = \sum_{k=0}^K \pi_k \cdot \mathcal{N}(z; 0; \gamma_k),$$

où $\mathcal{N}(z; 0; \gamma_k)$ est la distribution de probabilité d'une variable aléatoire Gaussienne de moyenne nulle et de variance γ_k . π_k est la probabilité que la k -ième composante Gaussienne intervienne dans le mélange.

Definition 2. *Le modèle de Middleton de classe A est paramétré par un facteur de recouvrement A et un rapport bruit de fond et bruit impulsif Ω . Le modèle Middleton classe A peut être considéré comme un cas particulier de la distribution de mélange Gaussienne avec $\pi_k = \exp(-A) \frac{A^k}{k!}$ et $\gamma_k = \frac{k/A + \Omega}{1 + \Omega}$, $K \rightarrow \infty$. Le paramètre A est défini comme le nombre moyen d'impulsions par unité de temps multiplié par le temps de durée moyenne de ces impulsions.*

$$\sigma_m^2 = \sigma_i^2 \frac{m}{A} + \sigma_g^2$$

où σ_i^2 est la variance du bruit impulsif et σ_g^2 est la variance de bruit de fond.

Le modèle de Middleton de classe A permet de décrire le bruit rencontré dans le canal CPL sans la prise en compte de ses caractéristiques temporelles [34].

2.3.3.4 Bruit à bande étroite

Le bruit à bande étroite est la conséquence de la captation par les lignes électriques des émissions de radiodiffusion : les transmissions AM, les radioamateurs, des alimentations à découpage, etc [9, 7]. Leurs enveloppes peuvent restées invariables dans le temps sur de longues périodes ou peuvent changer périodiquement de façon synchrone avec la tension du réseau. Il a également été rapporté par [40] que le bruit à bande étroite est présent de préférence à des fréquences inférieures à 140 kHz ou supérieure à 410 kHz et que la moyenne de leur bande passante est d'environ 3 kHz.

2.4 Technologies et architecture CPL-BE

2.4.1 Architecture et topologie des CPL-BE dans le domaine d'accès

Nous décrivons ici la topologie typique du réseau électrique MT/BT européen (230V/3 phases). Les foyers fournis par un poste de distribution sont organisés en cellule servie par un transformateur. Ces cellules peuvent comporter jusqu'à 350 ménages. La topologie est généralement arborescente et chaque cellule peut inclure jusqu'à 10 branches et chaque branche sert environ 35 ménages. Les câbles de connexion entre le transformateur et les ménages est d'environ 100 m de longueur [23]. Le concentrateur de données est généralement situé à la fin du segment BT avant le transformateur dans les réseaux européens. Le modèle de base d'un réseau CPL-BE pour les applications SG est donné dans la figure 2.1. Par exemple, pour le télécomptage, le compteur communicant situé à la maison recueille les informations d'utilisation et envoie ces données à travers le segment BT au concentrateur de données. Par la suite, le routeur MT est chargé de transmettre les informations au gestionnaire du réseau électrique par l'intermédiaire d'un réseau étendu. Bien que ce soit une topologie de bus qui est représentée sur la figure 2.1, d'autres topologies d'accès comme celle en étoile et en anneau sont possibles [9, 22].

2.4.2 Spécification des standards et normes CPL-BE existants.

Jusqu'à présent, deux générations de CPL-BE ont été développées :

- La première génération est basée sur une modulation à une ou 2 porteuses et a été proposée dans le cadre de standards comme LonWork (local operation network), KNX, CEBus.
- La seconde génération a commencé avec Powerline Related Intelligent Metering Evolution (PRIME) en 2007. PRIME est une solution CPL-BE supportant les modulations multi-porteuses comme l'OFDM, et en fonction du mode de transmission il intègre un module de codage correcteur d'erreurs (Forward Error Correction (FEC)).

Dans le même esprit, G3 a vu le jour en 2011. Il définit une solution CPL-BE supportant les modulations multi-porteuses et a été conçu pour être très robuste, à l'opposé de PRIME.

La norme IEEE 1901.2 basée sur G3, décrit uniquement les couches PHY et Medium Access Control (MAC) des systèmes CPL-BE multi-porteuses. IEEE 1901.2 inclut des mécanismes de coexistence permettant à des systèmes CPL-BE qui peuvent être non interopérables de partager la même bande de fréquence.

Une nouvelle technologie CPL-BE a été développée par l'Union Internationale des Télécommunications (UIT) en coopération avec les membres des Alliances G3 et PRIME : il s'agit de G.hnem. Cette norme définit une technologie basée sur l'OFDM en ciblant de multiples applications SG [\[41\]](#).

Un résumé des paramètres des normes et standards CPL-BE est disponible dans le tableau 2.3.

Tableau 2.3: Comparaison des principaux paramètres des spécifications de produits et standards CPL-BE multi-porteuses existants

Paramètre	Prime	G3	IEEE 1901.2	G.hnem
Bande de fréq.	CEN A	42-89 kHz	35.9-90.6 kHz	35.9-90.6 kHz
	FCC	/	159.4-478.1 kHz	34.4-478.1 kHz
Fréq. échantillonnage	CEN A	250 kHz	400 kHz	200 kHz
	FCC	/	1.2 MHz	800 kHz
Taille de la Fast Fourier Transform (FFT)	CEN A	512	256	128
	FCC	/	256	256
Durée du PC	CEN A	48 (192 μ s)	30 (75 μ s)	20/32 (100 μ s/160 μ s)
	FCC	/	30 (25 μ s)	40/64 (50 μ s/80 μ s)
Taille de la fenêtre	CEN A	0	8	8
	FCC	/	8	16
Durée du PC effectif	CEN A	48 (192 μ s)	22 (55 μ s)	12/24 (60 μ s/120 μ s)
	FCC	/	22 (18.3 μ s)	24/48 (30 μ s/60 μ s)
Espacement des SP	CEN A	488 Hz	1.5625 kHz	1.5625 kHz
	FCC	/	4.6875 kHz	3.125 kHz
Durée symboles OFDM	CEN A	2240 μ s	695 μ s	700/760 μ s
	FCC	/	231.7 μ s	350/380 μ s
Modulation	DPSK	DPSK	DPSK (QAM)	QAM
	M=2,4,8	M=2,4,8	M=2,4,8,16	M=2,4,8,16
FEC	Conv (Optional)	Conv et RS	Conv et RS	Conv et RS
	CEN A	61.4123 kbits/s	45 kbits/s	101.3 kbits/s
Débit Maximum	FCC	/	207.6 kbits/s	821.1 kbits/s
	FCC	/	203.2207.6 kbits/s	

2.5 Conclusion

Ce premier chapitre a présenté le concept de SG, nous avons introduit les communications par CPL-BE pour répondre aux besoins de liens de communication pour le domaine d'accès SG. Ce chapitre a été l'occasion de nous familiariser avec le contexte des SG et des CPL.

Toutefois, les développeurs de systèmes CPL sont mis au défi par les caractéristiques très défavorables du canal CPL-BE. Il est nécessaire de lui adjoindre des techniques de traitement de signal comme le codage correcteur d'erreurs pour fiabiliser les transmissions sur ce médium. La conception et la mise en œuvre de ces codes devra tenir compte des caractéristiques particulières du bruit rencontré sur les canaux CPL : le bruit à motif criss-cross par exemple.

Chapitre 3

Codes correcteurs d'erreurs pour les réseaux CPL-BE

Sommaire

3.1	Introduction	34
3.2	Théorie des codes correcteurs d'erreurs	34
3.2.1	Notions essentielles en théorie des codes	34
3.3	Codes correcteurs d'erreurs pour les CPL-BE .	40
3.3.1	Codes de Reed-Solomon	40
3.3.2	Codes convolutifs	42
3.3.3	Concaténation série de codes	44
3.3.4	Code à métrique rang	46
3.3.5	Protocole ARQ	48
3.3.6	Combinaison entre le codage correcteur d'erreurs et les protocoles de retransmission	49
3.4	Codes fontaines	51
3.4.1	Codes fontaines aléatoires	53
3.4.2	Codes LT	54
3.4.3	Codage réseau numérique	56
3.5	Conclusion	60

3.1 Introduction

Ce chapitre présente les concepts fondamentaux des techniques de codage de canal déployées pour fiabiliser les communications sur les systèmes télécoms. Nous introduisons différents codes correcteurs d'erreurs préconisés et/ou implémentés dans les systèmes CPL-BE.

Dans une première partie, la théorie des codes détecteurs-correcteurs d'erreurs est présentée.

Nous présentons ensuite, dans une seconde partie, différentes techniques de codage utilisées dans la couche physique des systèmes télécoms. Il s'agit des codes RS, des codes convolutifs, des codes Low Parity Check Code (LDPC) et des codes à métrique rang.

Enfin dans une troisième partie nous introduisons des techniques de codage adaptées à des canaux à effacements. Dans ce cas précis, il s'agit du codage réseau et des codes fontaines qui constituent une excellente alternative aux techniques traditionnelles de codage FEC pour des canaux non stationnaires ou dont les caractéristiques sont inconnues à l'émetteur.

3.2 Théorie des codes correcteurs d'erreurs

Dans son papier publié en 1948 [42], Claude E. Shannon a formalisé mathématiquement beaucoup de problèmes fondamentaux concernant la transmission (et le stockage) de données. Il a élaboré la théorie mathématique de la transmission de l'information, solutionnant entre autres, le problème de communications fiables sur des canaux non fiables.

3.2.1 Notions essentielles en théorie des codes

3.2.1.1 Définition de l'information selon Shannon

Pour donner une définition quantitative de l'information on adopte les notations suivantes, que l'on gardera dans la suite de ce manuscrit. On définit le triplet $(x, \mathcal{A}_X, \mathcal{P}_X)$, où x est une réalisation d'une variable aléatoire discrète, qui prend ses valeurs dans l'ensemble $\mathcal{A}_X = \{a_1, a_2, \dots, a_I\}$ avec

les probabilités $\mathcal{P}_X = \{p_1, p_2, \dots, p_I\}$, avec $\Pr(X = x_i) = p_i$, $p_i \geq 0$ et $\sum_{x_i \in \mathcal{A}_X} P(X = x_i) = 1$. L'ensemble $\mathcal{A}_X$ est appelé alphabet, les éléments de $\mathcal{A}_X$ sont appelés symboles et une chaîne finie d'éléments de $\mathcal{A}_X$, $a_1 a_2 \dots a_n$ est appelée mot (de longueur n). Généralement, l'alphabet est un corps de Galois : ($\mathcal{A}_X = \mathbb{F}_{2^m} (m \geq 1)$). De la même manière, on définit le triplet $(y, \mathcal{B}_Y, \mathcal{P}_Y)$ de même type que $(x, \mathcal{A}_X, \mathcal{P}_X)$, où y prend ses valeurs dans l'ensemble $\mathcal{B}_Y = \{b_1, b_2, \dots, b_J\}$ avec les probabilités $\mathcal{P}_Y = \{q_1, q_2, \dots, q_J\}$. Avec $\Pr(Y = y_j) = q_j$, $q_j \geq 0$ et $\sum_{y_j \in \mathcal{B}_Y} P(y = y_j) = 1$. La quantité d'information concernant la réalisation de l'évènement $X = x$ est une fonction f de la probabilité $p(x)$:

$$I(x) = f(p(x)). \quad (3.1)$$

La mesure de l'information (propre) f est une fonction décroissante de $p(x)$. L'information est une grandeur additive : si deux évènements x et y sont statistiquement indépendants alors l'information totale qu'ils peuvent fournir est la somme des informations associées à chacun de ces évènements $f(p(x, y)) = f(p(x) \times p(y)) = f(p(x)) + f(p(y))$ [43]. Toutes ces propriétés sont bien intégrées par la définition de I .

De façon qualitative, Shannon assimile l'information à la notion de surprise. Si un évènement rare se produit, sa réalisation apporte beaucoup d'information ; mais si au contraire un évènement certain se produit, sa réalisation n'apporte aucune information. La mesure de l'information (propre) apportée par la réalisation de l'évènement $X = x$ est donnée par :

$$I(x) = -\log_2(P(x)). \quad (3.2)$$

Pour une variable aléatoire X , la moyenne de son information propre correspond à son entropie $H(X)$:

$$H(X) = \mathbb{E}_X(I(x)) = - \sum_{x \in \mathcal{A}_X} P(x) \log_2(P(x)). \quad (3.3)$$

L'entropie mesure l'incertitude inhérente à la distribution d'une variable aléatoire. L'entropie conditionnelle de X sachant Y est la moyenne, par rapport à Y , de l'entropie de X étant donné y .

$$H(X|Y) = - \sum_{x \in \mathcal{A}_X} \sum_{y \in \mathcal{B}_Y} P(x, y) \log(P(x|y)) \quad (3.4)$$

Celle-ci mesure l'incertitude moyenne qui reste à propos de x lorsque y est connue. L'information mutuelle entre X et Y est définie par :

$$I(X; Y) = H(X) - H(X|Y) = H(Y) - H(Y|X). \quad (3.5)$$

Elle mesure la réduction moyenne dans l'incertitude à propos de x qui résulte de l'apprentissage de la valeur de y [44].

Un canal de transmission discret est complètement décrit par un alphabet d'entrée $\mathcal{A}_X$, un alphabet de sortie $\mathcal{B}_Y$ et les lois de transition probabilistes des différents mots qui transitent sur ce canal. Le canal est dit sans mémoire si pour tout $(a_1, \dots, a_n)$ transmis et $(b_1, \dots, b_n)$ reçus à travers ce canal,

$$P(b_1, \dots, b_n | a_1, \dots, a_n) = P(b_1 | a_1) \cdots P(b_n | a_n).$$

Definition 3. Une propriété importante des canaux de communication est leur capacité C . La capacité C d'un canal est le rendement avec lequel l'information peut être transmise à travers ce canal avec une probabilité d'erreur arbitrairement faible.

$$C = \max_{\mathcal{P}_X} I(X; Y). \quad (3.6)$$

C est défini comme le maximum de l'information mutuelle moyenne $I(X; Y)$, X et Y sont des variables aléatoires associées à l'entrée et à la sortie du canal.

3.2.1.2 Théorème fondamental du codage de canal [42]

- Associé à chaque canal discret sans mémoire, il existe une quantité non négative C (appelée capacité du canal) ayant la propriété suivante.

Pour tout $\epsilon > 0$ et $R < C$, pour N suffisamment grand, il existe un code en bloc de taille N et de rendement $\geq R$ et un algorithme de décodage, tel que la probabilité maximale d'erreurs bloc soit $< \epsilon$.

- Si une probabilité d'erreur p_b est acceptable, des rendements de transmission allant jusqu'à $R(p_b)$ sont réalisables, où

$$R(p_b) = \frac{C}{1 - H_2(p_b)}. \quad (3.7)$$

H_2 étant la fonction d'entropie d'une source binaire (ayant deux sorties possibles avec les probabilités x et $1 - x$) et est définie par $H_2(x) = x \log_2(1/x) + (1 - x) \log_2(1/(1 - x))$.

- Pour tout p_b , des rendements plus grand que $R(p_b)$ ne sont pas réalisables.

Le concept de capacité de canal même s'il n'est pas traité de façon exhaustive dans ce manuscrit, permet de rendre compte des limitations du canal en matière de rendement pour le codage correcteur d'erreurs. La capacité des canaux CPL-BE à été largement abordé dans [45].

✎ Remarque : Les canaux CPL-BE où le bruit est cyclo-stationnaire peuvent obtenir de meilleures performances que sous un canal à bruit blanc Gaussien. Il faut pour cela que le système soit conçu pour s'adapter aux propriétés statistiques et temporelles du bruit : en comptabilisant de façon optimale les propriétés périodiques du canal et du bruit.

3.2.1.3 Codage en bloc linéaires

Avant d'entamer la description des codes en bloc linéaires, nous présentons un résumé des propriétés des corps finis qui sont utilisés pour le codage. Les auteurs dans [46, 47], entre autres, donnent une description complète et détaillée de ces concepts.

Il existe un seul corps fini à un isomorphisme près $\mathbb{F}_{p^m}$ ayant p^m éléments. p étant nécessairement un nombre premier et m un entier ≥ 1 .

Les entiers modulo p munis de l'addition et de la multiplication modulo p forment un corps fini noté $\mathbb{F}_p$. Les polynômes $\mathbb{F}_p[x]$ définis sur $\mathbb{F}_p$ modulo un

polynôme irréductible $g(x) \in \mathbb{F}_p[x]$ de degré m , forment un corps fini de p^m éléments quand ils sont munis de l'addition et de la multiplication modulo $g(x)$.

Muni de l'addition, $\mathbb{F}_{p^m}$ est isomorphe à l'espace vectoriel $(\mathbb{F}_p)^m$. Muni de la multiplication, les éléments non nuls de $\mathbb{F}_{p^m}$ forment un groupe cyclique $\{1, \alpha, \dots, \alpha^{p^m-2}\}$ généré par un élément primitif $\alpha \in \mathbb{F}_{p^m}$.

Definition 4. Un code en bloc $\mathcal{C}(n, k)$ sur un alphabet $\mathcal{B}$ de q symboles est un ensemble de mots de longueur n appelés mots de code. Associé au code, se trouve l'opération d'encodage qui associe à chaque k -uplet $m = (m_0, m_1, \dots, m_{k-1}) \in \mathcal{A}^k$, appelé message, un n -uplet $(c_0, c_1, \dots, c_{n-1}) \in \mathcal{B}^n$ correspondant au mot de code [48].

$$\begin{aligned} \phi &: \mathcal{A}^k \rightarrow \mathcal{C} \subseteq \mathcal{B}^n \\ m &\mapsto c = \phi(m) \end{aligned}$$

Definition 5. Un code $\mathcal{C}(n, k)$ est dit linéaire si et seulement si ses q^k mots de codes, forment un k sous-espace vectoriel de l'espace vectoriel $\mathcal{B}^n$ des n -uplets de $\mathcal{B}$.

La matrice génératrice de $\mathcal{C}$ est une matrice de taille $k \times n$ dont les lignes forment une famille génératrice de $\mathcal{C}$. Pour un code donné, le décodeur optimal est celui qui minimise la probabilité d'erreur en bloc. Pour ce décodeur, une sortie y est décodée par l'entrée s qui a la probabilité a posteriori maximale $P(s|y)$. C'est-à-dire la probabilité d'avoir s connaissant l'observation y maximale. La probabilité $P(y|s)$ est la vraisemblance de l'observation, c'est-à-dire la probabilité d'observer y lorsque s est émis.

$$P(s|y) = \frac{P(y|s)P(s)}{P(y)} \quad (3.8)$$

$$s_{\text{optimal}}^* = \arg \max P(s|y) \quad (3.9)$$

Une distribution uniforme des s est généralement supposée, dans un tel cas, le décodeur optimal est également le décodeur à maximum de vraisemblance, à savoir, le décodeur qui associe la sortie y à l'entrée s qui a le maximum de

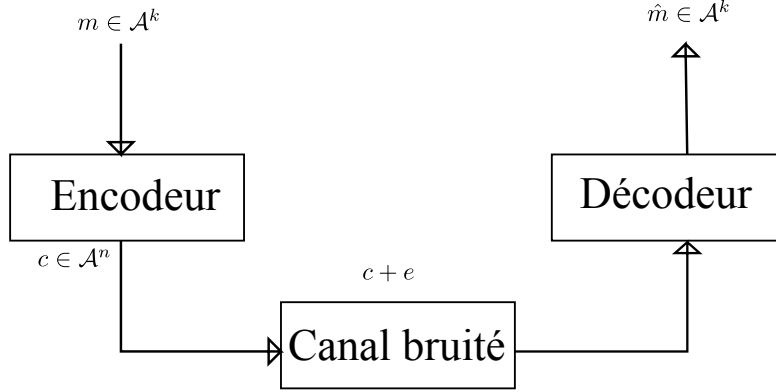


FIGURE 3.1: Modélisation de la chaîne de transmission avec un encodeur et décodeur

vraisemblance $P(y|s)$ [48]. Le rendement du code est $R = \frac{k}{n}$ et il correspond au ratio du nombre de transmissions de symboles d'information utile k et le nombre d'utilisations du canal n .

Definition 6. Soit $\mathcal{A}^n$ l'ensemble des mots de longueur n sur l'alphabet $\mathcal{A}$, pour $x^{(1)} = (a_1^1, \dots, a_n^1) \in \mathcal{A}^n$ et $x^{(2)} = (a_1^2, \dots, a_n^2) \in \mathcal{A}^n$. La distance de Hamming entre $x^{(1)}$ et $x^{(2)}$ est définie par :

$$d_H(x^{(1)}, x^{(2)}) = \left| \{i \mid 1 \leq i \leq n; a_i^1 \neq a_i^2\} \right|, \quad (3.10)$$

L'alphabet $\mathcal{A}$ est muni d'une structure de groupe et 0 représente son élément neutre. Le poids de Hamming $w_H(x)$ d'un mot $x \in \mathcal{A}^n$ est le nombre de ses coordonnées non nulles. La distance minimale d_{min} d'un code est la plus petite distance non nulle entre les éléments du code considérés deux à deux. Elle est définie par :

$$d_{min} = \min_{x^{(1)}, x^{(2)} \in \mathcal{A}^n} d_H(x^{(1)}, x^{(2)}). \quad (3.11)$$

Pour un code linéaire d_{min} correspond au poids du mot le plus creux. Un code linéaire $\mathcal{C}$ de longueur n , de distance minimale d_{min} défini sur un alphabet $\mathcal{A}$ de cardinal q vérifie la propriété $|\mathcal{C}| \leq q^{n-d_{min}+1}$. Cette borne est appelée borne de singleton. Les codes qui satisfont l'égalité sont appelés

codes à maximum de distance séparable. La distance minimale d_{min} est un paramètre primordial pour décrire les caractéristiques d'un code. Par exemple sur le canal BSC, la capacité de détection (le nombre de symboles d'erreurs que celui-ci est capable de détecter) est égal à $d_{min} - 1$ et la capacité de correction (le nombre de symboles d'erreurs que le code est capable de corriger) correspond à $\lfloor \frac{d_{min} - 1}{2} \rfloor$.

3.3 Codes correcteurs d'erreurs pour les CPL-BE

3.3.1 Codes de Reed-Solomon

Les codes Reed-Solomon ont été inventés par Irving Reed et Gus Solomon et présentés dans [49]. Il existe deux façons (équivalentes) de définir ces codes. Nous commençons par présenter la première définition, qui est aussi la définition originale de [49]. Pour $(n = q)$ les codes RS peuvent être définis comme une application de l'ensemble $\mathbb{F}_q^k$ des k -upples d'éléments de $\mathbb{F}_q$ vers l'ensemble $\mathbb{F}_q^n$ des n -upples d'éléments de $\mathbb{F}_q$. Pour chaque ensemble de k symboles d'information à envoyer $(f_0, f_1, \dots, f_{k-1})$, où $f_j \in \mathbb{F}_q$ et $0 \leq j \leq k - 1$, posons par

$$f(z) = f_0 + f_1 z + \dots + f_{k-1} z^{k-1} \in \mathbb{F}_q[z],$$

le polynôme de l'indéterminée z associé au message à envoyer.

Soient $\beta_1, \beta_2, \dots, \beta_q$ les q éléments de $\mathbb{F}_q$. Le polynôme $f(z)$ associe à chaque k -upple d'éléments de $\mathbb{F}_q$ un n -upple $(f(\beta_1), f(\beta_2), \dots, f(\beta_n)) \in \mathbb{F}_q^n$ [49]. Les $f(\beta_i)$ correspondent à l'évaluation du polynôme $f(z)$ sur tous les éléments β_i du corps $\mathbb{F}_q$:

Une autre méthode de représentation des codes RS consiste à les interpréter de la façon suivante. Pour tout $m = (m_0, m_1, \dots, m_{k-1})$ constitué de k symboles d'information à transmettre et son polynôme correspondant

$$m(x) = m_0 + m_1x + \cdots + m_{k-1}x^{k-1},$$

$$c(x) = m(x) \cdot x^{n-k} - R_{g(x)}[m(x) \cdot x^{n-k}],$$

où $R_{g(x)}$ désigne l'opération qui consiste à prendre le reste après division par $g(x)$. Et le polynôme $g(x)$ est définie par : $g(x) = \prod_{i=1}^{(n-k)} (x - \alpha^i)$ [48] et α un élément primitif de $\mathbb{F}_q$.

Ces codes sont à maximum de distance séparable, il est possible à partir de n'importe quel nombre k de symboles codés reçus -correctement- de pouvoir effectuer le décodage.

Une série d'algorithmes de décodage unique des codes RS portant les noms de Peterson [50], Berlekamp-Massey [47, 51], Berlekamp-Welch [52], Euclide ont été développés pour corriger les erreurs et/ou les effacements¹. Le décodage consiste généralement à trouver les positions d'erreurs et la valeur des erreurs dans ces positions. Nous détaillerons la méthode itérative utilisant l'algorithme de Berlekamp-Massey.

On exprime le n -upple r reçu à travers le canal comme étant la somme (sur $\mathbb{F}_q$) d'un mot de code et d'un n -upple inconnu e qui est l'erreur. Supposons qu'il y'a ν erreurs dans les localisations $i_1, i_2, \dots, i_\nu$ ($\nu \leq \lfloor \frac{n-k}{2} \rfloor$), et en notant par $e_{i_j} \neq 0$ les erreurs correspondants à ces localisations. La première étape de décodage des codes RS commence par le calcul des syndromes :

$$S_j = c(\alpha^j) + e(\alpha^j) = e(\alpha^j), \quad (3.12)$$

$$S_j = \sum_{l=1}^{\nu} e_{i_l} (\alpha^j)^{i_l} = \sum_{l=1}^{\nu} e_{i_l} (\alpha^{i_l})^j, \quad (3.13)$$

$$S_j = \sum_{l=1}^{\nu} e_{i_l} X_l^j. \quad (3.14)$$

Avec $X_l = \alpha^{i_l}$, $j = 1, 2, \dots, 2t$, $t = \lfloor (n-k)/2 \rfloor$. Les syndromes ne dépendent pas des données transmises mais uniquement des caractéristiques des erreurs.

1. Pour les codes définis avec la métrique de Hamming, un effacement correspond à une position d'erreur dont on ignore la valeur.

On définit le polynôme localisateur d'erreurs $\Lambda(x)$ par :

$$\Lambda(x) = \prod_{l=1}^{\nu} (1 - X_l x). \quad (3.15)$$

Le polynôme évaluateur d'erreurs est défini par :

$$\Gamma(x) \triangleq S(x)\Lambda(x) \mod x^{2t} \quad (3.16)$$

$$= \sum_{i=1}^{\nu} e_i X_i \frac{\Lambda(x)}{1 - X_i x} \quad (3.17)$$

Les syndromes et les coefficients du polynôme localisateur sont donc reliés par la relation de récurrence suivante :

$$S_j = - \sum_{i=1}^{\nu} \Lambda_i S_{j-i} \quad \nu \leq j \leq 2t - 1. \quad (3.18)$$

L'algorithme de Berlekamp-Massey par exemple, permet de construire de façon efficace un registre à décalage à rétroaction linéaire pour résoudre l'ensemble des équations permettant le décodage des codes RS. Pour la recherche des racines du polynôme localisateur l'implantation utilise généralement la recherche de Chien qui est manière pratique de faire une recherche exhaustive.

Enfin l'algorithme de Forney [48], calcule les valeurs des erreurs comme suit :

$$e_{i_k} = - \frac{\Lambda(X_k^{-1})}{\Gamma'(X_k^{-1})}. \quad (3.19)$$

3.3.2 Codes convolutifs

Les codes convolutifs sont des codes linéaires qui ont été proposés en premier par Peter Ellias [53]. Ces codes à l'opposé des codes en bloc opèrent sur un flux continu de données. Ils peuvent facilement être adaptés pour des transmissions de paquets grâce à des techniques dites de fermeture de treillis [54]. Les séquences d'entrée et les fonctions de transfert de ces codes sont

représentées par des séries de puissances de la variable D (comme Delay²). Le message à transmettre m peut être représenté sous forme de série de Laurent $m(D) = \sum_{l=-\infty}^{+\infty} m_l D^l$ et les fonctions de transfert qui à partir de $m(D)$ génèrent $c(D)$ peuvent être représentées sous forme $G(D) = \frac{N(D)}{D(D)}$ où, $N(D)$ et $D(D)$ sont des polynômes définis sur $\mathbb{F}_2$ et $G(D)$ n'est pas factorisable.

Trois représentations sont possibles pour les codes convolutifs : l'arbre, le treillis et la machine à états. La machine à états représente les différentes transitions entre les états du codeur. Elle a l'avantage de permettre de déterminer la fonction de transfert et le spectre de distances du code convolutif : toutes choses utiles pour prédire les performances d'un code. Un paramètre important des codes convolutifs est la longueur de contrainte : c'est le cardinal de l'ensemble des entrées m_i qui définissent l'état de sortie du codeur à un moment donné.

Le plus populaire des algorithmes de décodage à maximum de vraisemblance des codes convolutifs est l'algorithme de Viterbi. Il s'appuie sur la représentation en treillis des codes [55] et permet de trouver, à partir de la séquence bruitée des symboles reçus, de l'état initial du codeur, des transitions autorisées, la séquence d'états dans le treillis la plus probable.

L'idée basique d'un décodeur de Viterbi est la suivante. Une séquence codée est transmise à travers un canal, et à cause du bruit, la séquence reçue r peut ne pas correspondre exactement à un chemin correspondant à un mot de code dans le treillis. Le décodeur trouve un chemin dans le treillis qui est le plus proche (proche dans le sens d'une métrique adaptée : pour un canal Bruit Blanc Additif Gaussien (BBAG), cette métrique est la distance euclidienne, et pour un canal binaire symétrique Binary Symmetric Channel (BSC), la métrique est la distance de Hamming.) de la séquence envoyée. Pour une séquence donnée, les calculs de métriques et les comparaisons sont faites de façons récursives pour limiter la complexité.

Une autre façon de décoder les codes convolutifs se fait avec le décodage au sens Maximum A Posteriori (MAP). Basé sur l'algorithme BCJR du nom

2. Délai en anglais, c'est l'opérateur dont l'effet est de retarder chaque élément d'une séquence par une unité de temps ; à savoir, $m' = Dm$ signifie $m'_k = m_{k-1}$ pour tout $k \in \mathbb{Z}$.

de ses auteurs [56], il estime les données transmises en associant à chaque bit une probabilité (la fiabilité sur la décision de chaque bit) en se basant sur la séquence reçue.

3.3.3 Concaténation série de codes

Pour obtenir des codes à plus grande complexité, Forney [57] a proposé la concaténation de codes : un premier code (code extérieur) fournit un mot de code qui est ensuite re-encodé par le second codeur (code intérieur). Les systèmes CPL-BE utilisent les codes concaténés série RS pour le code extérieur et convolutif pour le code intérieur. Ces deux codes sont généralement précédés par un entrelaceur en temps et en fréquence pour éviter les erreurs en rafales dus au bruit impulsif et au bruit à bande étroite qui sont difficiles à corriger par les codes convolutifs. Le code convolutif utilisé est un code de rendement $1/2$ qui utilise la méthode de zéro tailing avec 6 bits qui permettent de terminer l'encodage dans un état donné en mettant tous les bits à zéro [26]. Le décodeur intérieur est le décodeur de Viterbi, qui tire profit des données de fiabilités fournies par le démodulateur. Le décodeur extérieur RS est bien adapté aux erreurs en rafales.

☞ Remarques : La concaténation de code peut être aussi parallèle, et le processus de décodage itératif, avec un échange de fiabilité entre codeur intérieur et codeur extérieur qui affinent à chaque fois leurs décisions sur les valeurs des bits décodés. L'exemple le plus parlant de cette stratégie est constitué par la classe de codes très performants que sont les Turbo codes [54]. Les Turbo codes sont considérés comme une des ruptures technologiques dans le domaine des communications numériques ; ils ont permis de mettre en évidence l'efficacité du décodage itératif.

3.3.3.1 Codes LDPC

Ils ont été découverts par Gallager, oubliés pendant un certain temps avant leur renaissance dû à Mackay et Neal [58] et -de façon indépendante- par Sipser et Spielman [59], juste après la découverte des turbo codes qui ont permis de montrer le potentiel du decodage iteratif comme moyen

d'approcher la capacité. Ils se sont également révélés comme étant les codes les plus à même d'approcher le plus près possible de la limite de Shannon sur des canaux BBAG. Ces codes ne sont pas protégés par un brevet à l'opposé des turbo codes.

Un code LDPC $\mathcal{C}(n, k)$ binaire est basé sur la représentation de la matrice de contrôle de parité; à savoir l'ensemble de tous les n -uple binaires qui satisfont les $n - k$ équations de contrôle de parité :

$$x \cdot H^T = 0, \quad (3.20)$$

où H de taille $(n - k) \times n$ est la matrice de contrôle de parité. L'idée de base d'un code LDPC est que d'une part n doit être grand et H doit être creuse : la densité de 1 dans H doit être de l'ordre $c \cdot n$ plutôt que n^2 . Ce caractère creux rend possible dans le sens pratique le décodage de C de façon itérative par l'algorithme somme-produit en un temps linéaire.

Dans le graphe représentant les codes LDPC appelé graphe de Tanner, les nœuds de variable sont les nœuds associés aux bits du message à coder et les nœuds de parité sont associés aux équations de parité. Si $H\{m, n\} = 1$, alors une branche relie le nœud de variable n au nœud de parité m . La figure 3.2 représente un graphe de Tanner avec 6 nœuds de variable et 4 nœuds de parités.

On distingue plusieurs types de codes LDPC, les codes LDPC réguliers et les codes LDPC irréguliers. Pour les codes LDPC réguliers, le nombre de « 1 » dans les lignes (d_c) et dans les colonnes (d_v) de la matrice de contrôle de parité est constant. Cette contrainte n'est pas nécessaire pour un code LDPC irrégulier. Ce degré de liberté supplémentaire des codes LDPC irréguliers les rend plus efficacement optimisables, d'où leurs meilleures performances par rapport aux codes LDPC réguliers.

La construction de H est effectuée en minimisant le nombre de cycles; elle peut se faire de façon déterministe ou aléatoire. Dans la représentation de Tanner du code, un cycle commence à un nœud et se termine par le même nœud en passant par d'autres nœuds différents. Étant donné que les codes LDPC sont définis à partir de leurs matrices de parité, leur encodage peut se

faire par calcul de la matrice génératrice G avec une complexité $\mathcal{O}(n^2)$. Mais la taille importante des mots de code rend cette méthode rédhitoire. Des méthodes d'encodage utilisées sont généralement ah hod [54].

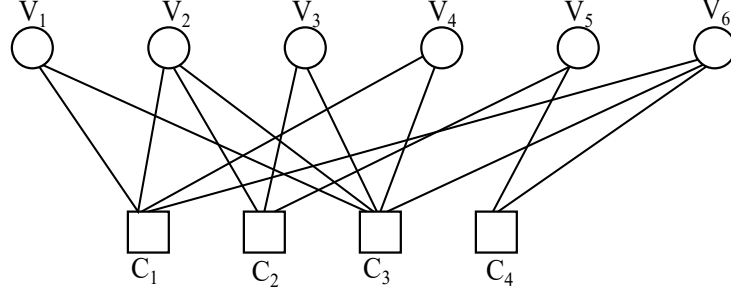


FIGURE 3.2: Exemple de représentation d'un code LDPC avec le graphe de Tanner

L'algorithme de décodage repose sur l'échange d'informations (message passing) des probabilités sur les fiabilités des décisions sur chaque bit, de façon itérative, entre les nœuds de variable et les nœuds de parité. Le message passing est une classe de décodeurs itératifs qui opèrent sur la représentation graphique d'un code. Le comportement de l'algorithme de décodage peut être analysé plutôt précisément à l'aide d'une technique appelée « évolution de densité » qui permet d'optimiser les performances des codes, du moins leurs caractéristiques asymptotiques pour un canal donné.

3.3.4 Code à métrique rang

La métrique rang a été introduite en théorie du codage par Delsarte [60] et a été largement développée par E. Gabidulin [61]. Cette métrique est adaptée, entre autres, à un modèle de canal où les mots de code peuvent être considérés comme des matrices d'éléments sur un corps fini et où les erreurs se produisent en bloc soit sur les lignes et/ou sur les colonnes.

Gabidulin [61] et Roth [62] ont introduit des codes à métrique rang avec la propriété de distance rang minimale capable de corriger un certain nombre de lignes et/ou colonnes de la matrice transmise.

Soit $\mathbb{F}_q$ un corps de Galois et $\mathbb{F}_{q^N}$ son extension de degré N . $\mathbb{F}_{q^N}$ peut être considéré comme un espace vectoriel sur $\mathbb{F}_q$. Soit $\mathcal{B}$ la base de $\mathbb{F}_{q^N}$ sur

$\mathbb{F}_q$. Considérons le mappage bijectif, pour $n \leq N$:

$$\mathcal{A} : (\mathbb{F}_{q^N})^n \longrightarrow A_N^n, \quad (3.21)$$

qui associe le vecteur $a = (a_1, \dots, a_n)$; $a_i \in \mathbb{F}_{q^N}$ à une matrice A de taille $(N \times n)$ représentant les coordonnées de a_i sur $(\mathbb{F}_q)^N$ décomposé selon la base $\mathcal{B}$. Le rang de a sur q est défini comme $r(a|q) = \text{rang}(A|q)$. Pour deux vecteurs x et y de longueur n avec des éléments dans $\mathbb{F}_{q^N}$, $d_r(x, y) = r(x-y|q)$. On vérifie que d_r est bien une distance. La distance rang pour un code linéaire sur $\mathbb{F}_{q^N}$ est définie comme étant la distance d_r minimale entre toutes les paires de mots de code

Definition 7. *Un polynôme linéarisé sur $\mathbb{F}_{q^N}$ est un polynôme de la forme :*

$$L(x) = \sum_{p=0}^{N(L)} L_p x^{q^p},$$

où $L_p \in \mathbb{F}_{q^N}$ et $N(L)$ caractérise le plus grand p tel que $L_p \neq 0$. $\otimes$ est définie comme étant la composition ou la multiplication symbolique de polynômes.

$$F(x) \otimes G(x) = F(G(x)), \quad (3.22)$$

Le produit symbolique est associatif et distributif, mais non commutatif. Ces polynômes sont linéaires sur $\mathbb{F}_q$.

$$F(\lambda_1 \beta_1 + \lambda_2 \beta_2) = \lambda_1 F(\beta_1) + \lambda_2 F(\beta_2), \quad (3.23)$$

pour tout $\lambda_1, \lambda_2 \in \mathbb{F}_q$ et $\beta_1, \beta_2 \in \mathbb{F}_{q^N}$.

Posons $[i] = q^i$ et soit G une matrice de taille $k \times n$ défini sur le corps $\mathbb{F}_{q^N}$. Soient $(g_1, \dots, g_n) \in (\mathbb{F}_{q^N})^n$ n éléments linéairement indépendants sur $\mathbb{F}_q$,

$$G = \begin{bmatrix} g_1 & g_2 & \cdots & g_n \\ g_1^{[1]} & g_2^{[1]} & \cdots & g_n^{[1]} \\ g_1^{[2]} & g_2^{[2]} & \cdots & g_n^{[2]} \\ \cdots & \cdots & \cdots & \cdots \\ g_1^{[k-1]} & g_2^{[k-1]} & \cdots & g_n^{[k-1]} \end{bmatrix}.$$

G ainsi définie est la matrice génératrice du principal type de code à métrique rang : les codes Gabidulin. On appelle $g = (g_1, \dots, g_n)$ support du code. Les codes Gabidulin de longueur n , de dimension k et de support g sont constitués par l'ensemble des mots obtenus par l'évaluation des polynômes linéarisés encore appelés q -polynômes de degré au plus égal à $k - 1$ sur n points.

Les codes Gabidulin sont semblables aux codes RS pour la métrique de Hamming, aussi leur décodage utilise des algorithmes inspirés du décodage des codes RS comme les algorithmes de Welch-Berlekamp, de Berlekamp-Massey, d'Euclide étendu et de Peterson-Gornstein-Zierler [63, 64, 65, 62].

3.3.5 Protocole ARQ

Il existe une probabilité d'échec du décodage du code correcteur d'erreurs direct (FEC), P_{out} qui est une fonction du taux d'erreurs binaire des trames reçues. Pendant les événements d'échec du décodage, le protocole Automatic reQuest Response (ARQ) permet la détection et la retransmission des données qui ont été incorrectement reçues. Dans les systèmes CPL-BE, le protocole ARQ est mis en œuvre sur la base de retransmissions acquittées ou non acquittées. Le principe est simple : au cours d'un échange de données, l'émetteur envoie une trame de données à la fois. Ensuite il attend de recevoir un message d'accusé de réception qui confirme la bonne transmission de la trame envoyée avant d'envoyer une trame supplémentaire. Si l'émetteur ne reçoit pas l'accusé de réception avant l'expiration d'un temps prédéfini (appelé temporisateur), il retransmet la trame précédemment envoyée. Ce protocole appelé stop-and-wait ARQ est habituellement utilisé pour les systèmes CPL-BE [66].

Selon les technologies CPL-BE utilisées, les stratégies sont différentes quant à la combinaison du codage correcteur d'erreurs direct (FEC) et des protocoles de retransmission (ARQ). Par exemple, PRIME allège les mécanismes de correction d'erreurs et compte sur un système de retransmissions pour garantir des transmissions fiables. Le protocole ARQ est préféré au codage correcteur d'erreurs direct pour le contrôle d'erreurs dans PRIME. En revanche, G3 utilise une correction d'erreurs forte sur la couche physique.

3.3.6 Combinaison entre le codage correcteur d'erreurs et les protocoles de retransmission

La connaissance du comportement du taux d'erreur binaire en fonction des durées de trames permet une meilleure compréhension de l'avantage d'une sélection optimale de durée de trames. Ceci pourrait permettre une optimisation intercouches des mécanismes de correction d'erreurs pour les canaux CPL-BE. En fait, à partir d'une certaine taille de trames, augmenter les longueurs de trames n'améliore pas forcément les performances en termes de débit utile.

Pour le bruit blanc additif Gaussien, le $TEB = f(F_l)$ est une fonction invariante de la longueur (F_l) des trames envoyées. Pour les canaux CPL-BE, le comportement du TEB selon les durées des trames transmises n'est plus évidente en raison des caractéristiques statistiques et temporelles du bruit. Nous proposons d'analyser le TEB en fonction de la longueur des trames IEEE 1901.2 [1] transmises.

Habituellement, de courtes durées de trame se traduisent par de meilleures temps de latence. De longues durées de trame (qui offrent souvent une meilleure protection contre les dégradations de canal) ont tendance à provoquer des périodes de traitement supplémentaires. Ce qui représente une perte d'efficacité pour le système.

Les longues durées de trame sont souvent associées à des profondeurs d'entrelacement plus grandes. Quand cela est fait intelligemment, tel que préconisé dans [67], le mécanisme d'entrelacement peut aider à apporter de

la diversité dans les transmissions et séparer les paquets d'erreurs rencontrés dans les canaux CPL-BE. La meilleure performance est généralement obtenue pour les grandes tailles de profondeurs d'entrelacement. Il est nécessaire de trouver des longueurs optimales pour les trames émises sur réseaux CPL-BE. Pour chaque longueur de trame, en utilisant des méthodes de Monte Carlo,

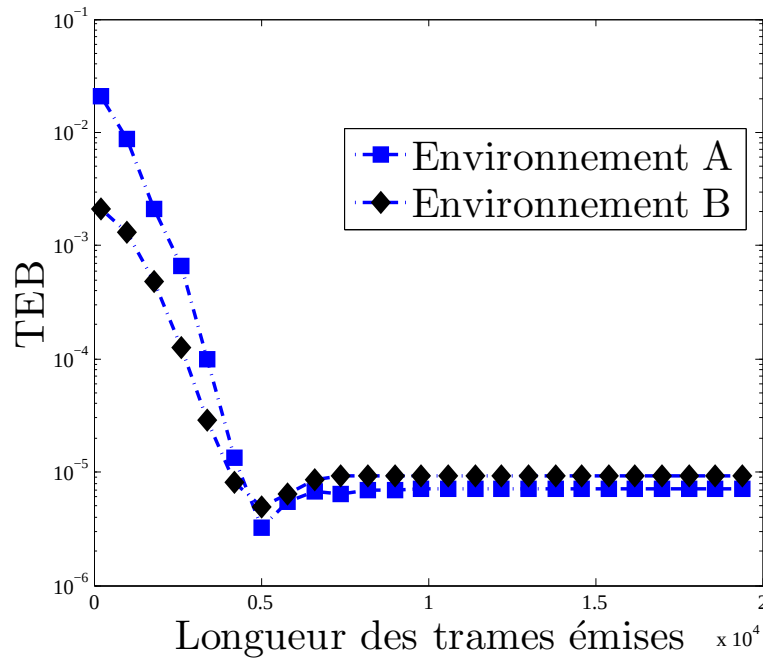


FIGURE 3.3: Les courbes du TEB en fonction de la longueur des trames transmises pour les environnements A et B

le TEB des trames émises est calculé pour un SNR fixe ($\text{SNR} = 3 \text{ dB}$) et un taux d'échantillonnage $f_s = 10^6 \text{ Hz}$.

L'analyse de l'allure des courbes de la figure 3.3 met en évidence deux parties. Le TEB diminue avec les longueurs des trames émises jusqu'à son minimum qui est atteint pour 5000 symboles transmis. Puis le TEB augmente légèrement avec la taille des trames émises avant de se stabiliser.

Cela est dû au fait que, pour de petites longueurs de trame, soit les

trames couvrent les périodes pendant lesquelles le bruit est faible (minima de la fonction de variance du bruit cyclo-stationnaire), auquel cas le TEB est bon ; où les trames couvrent les périodes pendant lesquelles le bruit est élevé (maximum de la fonction de variance du bruit cyclo-stationnaire). Pour de longues durées de trames, l'effet de la présence d'échantillons du bruit avec une puissance élevée est pondéré par les échantillons (de bruit) suivants ayant une faible puissance.

Ainsi à partir d'une certaine taille, étendre les longueurs de trames n'augmente pas forcément le TEB moyen.

3.4 Codes fontaines

Étant donné que les canaux CPL-BE sont non-stationnaires, pour garantir la fiabilité des communications, les rendements des codes correcteurs d'erreurs devraient soit être fixés - de façon non optimale - pour les conditions les moins favorables du canal. Ou bien ces rendements devraient être plus grands que ceux correspondants aux pires conditions du canal. Dans ce cas, il faut coupler le mécanisme de codage correcteur d'erreurs directe (FEC : Forward Error Correction) au niveau de la couche physique à d'autres mécanismes dans les couches supérieures pour terminer de fiabiliser le canal. C'est ce dernier schéma que les systèmes CPL-BE actuels utilisent. Ainsi les événements où le rendement du code est supérieur à la capacité du canal sont possibles³ et dans ces cas, occasionnellement, le décodeur peut délivrer des erreurs en sortie. Ces erreurs peuvent être détectées grâce à l'utilisation du mécanisme Cyclic Redundancy Check (CRC).

Quand le récepteur est muni d'un module de détection des paquets erronés, un paquet transmis à travers le canal est soit correctement reçu ou erroné au quel cas il est considéré perdu : on parle alors de canal à effacements. Les protocoles de transfert de données sur des canaux à effacement de paquets standards envoient un fichier constitué de plusieurs paquets en transmettant de façon répétée chaque paquet jusqu'à ce qu'ils

3. Ces événements correspondent à la survenue du bruit impulsif par exemple

soient tous reçus avec succès. Un paquet ici étant une unité élémentaire, constitué de l bits, qui est soit effacé ou correctement reçu à travers le canal à effacement. Suivant ce schéma, un canal de retour est généralement utilisé pour faire connaître à l'émetteur quel paquet a besoin d'être retransmis. Mais selon Shannon, il n'y a pas forcément besoin d'un canal de retour : la capacité d'un canal à effacement de paquets étant $(1 - f)l$ avec ou sans utilisation d'un canal de retour [68]. Les méthodes qui s'appuient sur l'utilisation systématique du canal de retour ne sont donc pas optimales, spécialement pour les canaux multicast qui possèdent des taux d'effacement différents pour chaque récepteur. Car il faudrait alors retransmettre à toutes les destinations un paquet que l'une d'entre elles n'aurait pas correctement reçu. En fait la mise à l'échelle pour le multicast n'est pas optimal pour les procédés de retransmissions. En effet pour les scénarios de transmission broadcast et/ou multicast, les messages de feedback pour des codes à rendement fixe sont excessifs.

Des codes à rendement fixe comme les codes RS ou LDPC pourraient être utilisés sur les canaux à effacement. L'inconvénient des codes RS est une complexité de décodage - rédhibitoire - pour des grandes tailles de paquets. La raison de cette complexité est que chaque symbole codé RS dépend de tous les symboles d'information. En outre, il faut que le rendement de ces codes soit calculé et calibré pour les caractéristiques d'un canal donné avant toute transmission.

Nous aimerions avoir des codes utilisant faiblement le canal de retour et qui à l'opposé des codes à rendement fixe pourraient s'adapter à n'importe quelles statistiques d'effacement du canal. Les codes fontaines, introduits dans [69], ont été proposés en tant que solution respectant ces contraintes en plus de celui d'être optimaux pour des canaux à effacements multicast.

Le principe de ces codes est simple : l'émetteur envoie des combinaisons linéaires aléatoires de paquets. Le nombre de paquets qui peut être généré à partir des données à l'émetteur est potentiellement illimité : ces codes sont dits sans rendement. Dans ce cas, l'outage n'est virtuellement jamais expérimenté par le récepteur étant donné que les transmissions peuvent être infinies (du moins jusqu'à ce que la destination accumule suffisamment

d'information mutuelle pour pouvoir décoder). Pour le récepteur, le nombre de paquets perdus et la composition des paquets importent peu, seul le nombre de paquets correctement reçus est important pour assurer le décodage d'où la métaphore de fontaine.

Un algorithme de décodage d'un code fontaine est un algorithme qui permet de récupérer les paquets d'origine (K packets) à partir de n'importe quel ensemble n de paquets codés avec une forte probabilité. Pour les bons codes fontaine la valeur de n est très proche de K et le décodage est presque linéaire en K .

3.4.1 Codes fontaines aléatoires

Pour cette famille de codes fontaines, chaque paquet codé est généré à partir de la combinaison linéaire de paquets choisis de façon uniforme et aléatoire. Pour chaque paquet codé, les paquets source interviennent dans la composition du paquet codé avec une probabilité de $1/2$. La distribution de sortie des degrés est alors binomiale. Le décodage s'effectue par inversion matricielle à partir de la réception de K paquets. Pour un excédent de E packets, la probabilité de succès du décodage est d'au moins $(1 - \delta)$, où $\delta = 2^{-E}$ [68].

Le graphe représenté dans la figure 3.4 permet de mieux représenter la mise en œuvre du codage fontaine. C'est un graphe bipartite composé de nœuds de variable (qui correspondent aux messages à coder et aussi aux lignes de la matrice de contrôle parité) et de nœuds de parité (qui correspondent aux contraintes et aussi aux colonnes définies dans la matrice de contrôle de parité). Une arête relie un nœud de variable à un nœud de parité si une entrée non nulle existe dans l'intersection correspondante à la ligne de ce nœud de variable et à la colonne de ce nœud de parité. Un polynôme permet de décrire le nombre d'arêtes qui entrent dans les nœuds de variable. Celui-ci est appelé distribution d'entrée des degrés. Le polynôme qui décrit le nombre d'arêtes qui entrent dans les nœuds de parité est appelé distribution de sortie des degrés.

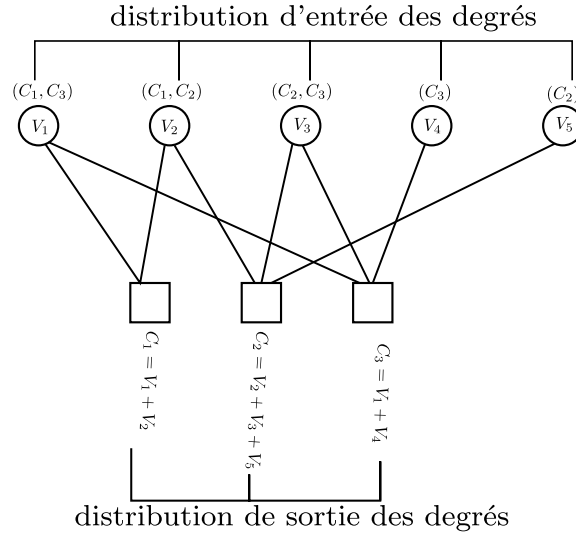


FIGURE 3.4: Distributions d'entrée et de sortie représentées avec le graphe bipartite des codes LT

3.4.2 Codes LT

Dans cette section, une description des codes Luby Transform (LT) est donnée, les distributions d'entrée optimales et de sortie des paquets sont décrites par [70]. Le codage LT est effectué paquet par paquet. Supposons que nous avons à transmettre K paquets d'informations. Le paquet codé t_n est produit à partir du bloc source M , qui comprend K paquets ($M = [M_1, M_2, \dots, M_K]$), par :

- i) Génération d'une variable aléatoire notée d_n à partir d'une distribution de degrés prédéfinie.
- ii) Transmission du paquet codé t_n , qui est obtenu par la somme bits par bits (mod 2) de d_n paquets *choisis uniformément de façon aléatoire* parmi les K paquets. Dans la pratique, les informations sur les paquets impliqués dans la combinaison peuvent être incluses dans une partie du paquet codé (en-tête par exemple), ou elles peuvent être obtenues par l'intermédiaire de moyens de synchronisation entre l'émetteur et le récepteur.

Du côté du décodeur, à la réception d'un paquet codé, le décodeur à propagation de croyance exécute l'algorithme suivant :

- i) Si le paquet est de degré 1, le paquet est considéré « découvert ». Ensuite, tous les paquets précédemment reçus et tous les futurs paquets impliquant ce paquet découvert sont « xoré » avec le paquet découvert. Ceci dans le but d'enlever son effet et d'obtenir des paquets de degré inférieurs.
- ii) Si pendant le processus de l'étape (i), des paquets de degré 1 sont générés, répéter l'étape (i). L'étape (i) est répétée jusqu'à ce que tous les K packets soient découverts.

Le design des codes LT utilise des graphes de densité logarithmique en nombre d'arêtes reliant nœuds de variable et nœuds de parité. La performance du décodeur à propagation de croyance est très dépendante de la distribution de degrés à partir de laquelle les degrés des paquets codés sont choisis. La distribution de sortie des degrés optimale en ce sens qu'il réduit au minimum l'overhead et la probabilité d'échec du décodeur est la distribution Robust Soliton définie par μ_K :

$$\mu_K(d) = \frac{\rho(d) + \tau(d)}{(\sum_{i=1}^K \rho(i) + \tau(i))}, \quad (3.24)$$

$$\rho(d) = \begin{cases} 1/K, & \text{si } d = 1 \\ 1/(d \cdot (d - 1)) & \text{autrement.} \end{cases} \quad \text{et} \quad (3.25)$$

$$\tau(d) = \begin{cases} S/K \cdot 1/d, & \text{pour } d = 1 \cdots \lfloor K/S \rfloor - 1. \\ S \cdot \ln(S/\delta)/K, & \text{pour } d = \lfloor K/S \rfloor. \\ 0, & \text{autrement.} \end{cases} \quad (3.26)$$

Où K est le nombre de paquets à envoyer, d est le degré envoyé, $S = c \cdot \ln(K/\delta) \cdot \sqrt{K}$, les paramètres c et δ sont utilisés pour ajuster la performance du code LT [70]. Dans le reste de ce manuscrit, μ_K désigne une distribution de soliton robuste de taille K , avec $\mu_K(0) = 0$. Outre la distribution de sortie des degrés, pour un décodage à propagation de croyance optimal, la distribution d'entrée des degrés doit être uniforme et aléatoire. Cette exigence produit une distribution binomiale pour la distribution d'entrée des degrés. Celle-ci est bien approchée par une distribution de Poisson [71].

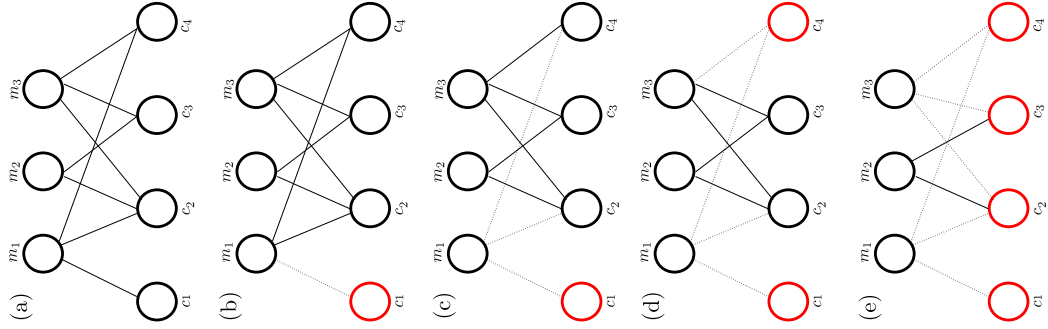


FIGURE 3.5: Exemple de mise en œuvre de l'algorithme de propagation de croyance pour le décodage LT sur un canal à effacement

La clé de conception des codes LT est basée sur l'introduction et l'analyse du processus LT. Ce processus permet d'assurer la présence - en moyenne - à chaque itération de l'algorithme de propagation de croyance d'un ensemble de paquets permettant le décodage du code LT, i.e, des paquets de degré 1.

Il existe d'autres classes de codes fontaines comme les codes Raptor [71] qui sont une extension des codes LT, ces codes peuvent être systématiques et possèdent un plancher d'erreurs plus intéressant que les codes LT. Les codes Raptor améliorent les complexités d'encodage et de décodage des codes LT de logarithmique à constant en fonction de K . Ils ont été brevetés et standardisés sous les noms de R10 et RQ [72]. Le graphe des codes LT a besoin d'avoir de l'ordre de $K \cdot \log(K)$ arêtes afin d'assurer que tous les nœuds d'entrée sont couverts avec une forte probabilité. La solution apportée par les codes Raptor consiste à d'abord encoder l'information avec un code correcteur d'effacement linéaire traditionnel à rendement élevé et ensuite à utiliser un code LT dont la distribution des degrés est tronquée à un degré maximal D . Tronquer la distribution de degrés à une valeur indépendante de K réduit le degré moyen des paquets encodés. L'idée clé de codage Raptor est de relâcher la contrainte que tous les symboles d'entrée doivent être récupérés.

3.4.3 Codage réseau numérique

Introduit en 2000 by R. Ahlswede [73], dans le codage réseau les flux d'un seul ou de différents utilisateurs sont mélangés dans une station relais pour

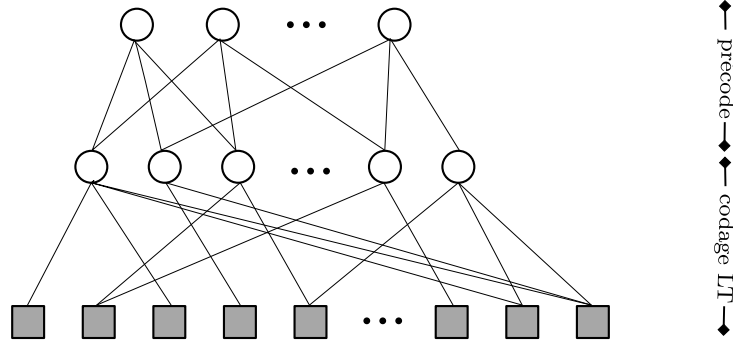


FIGURE 3.6: Représentation d'un code raptor

obtenir de nouveaux messages. Le codage réseau contrairement au routage traditionnel -store and forward- atteint la capacité multicast [74]. Le choix des combinaisons au niveau des relais peut être linéaire et aléatoire, ceci évite de fixer les combinaisons au niveau de chaque nœud du réseau et nous affranchi de la nécessité de connaître la topologie du réseau. L'émetteur injecte un certain nombre de paquets de longueur fixe dans le réseau⁴. Ces paquets se propagent à travers le réseau en passant par un certain nombre de nœuds intermédiaires entre émetteur et récepteur. Chaque fois qu'un nœud intermédiaire a l'opportunité d'envoyer un paquet, il crée une combinaison $\mathbb{F}_q$ -linéaire aléatoire des paquets dont il dispose et transmet cette combinaison aléatoire. Enfin, le récepteur recueille les paquets ainsi générés et essaie de déduire l'ensemble des paquets injectés dans le réseau.

Le codage réseau linéaire aléatoire (Random Linear Network Coding (RLNC)) ne tient pas compte de la topologie du réseau sous-jacent : aucune hypothèse n'est effectuée sur la topologie du réseau. Ici, l'information est codée à l'émetteur sous forme d'un sous-espace vectoriel (et pas de vecteur), et chaque sous-espace vectoriel est acheminé par la transmission d'un ensemble de vecteurs qui génère ce sous-espace vectoriel. Pour chaque paquet codé transmis, un entête permet de communiquer au récepteur les paquets qui ont été utilisés pour former cette combinaison.

Il existe aussi le codage réseau physique ou codage réseau analogique qui

4. chaque paquet peut être considéré comme un vecteur ligne de longueur N sur un corps fini $\mathbb{F}_q$

s'effectue naturellement par la super imposition de signaux physiques [75].

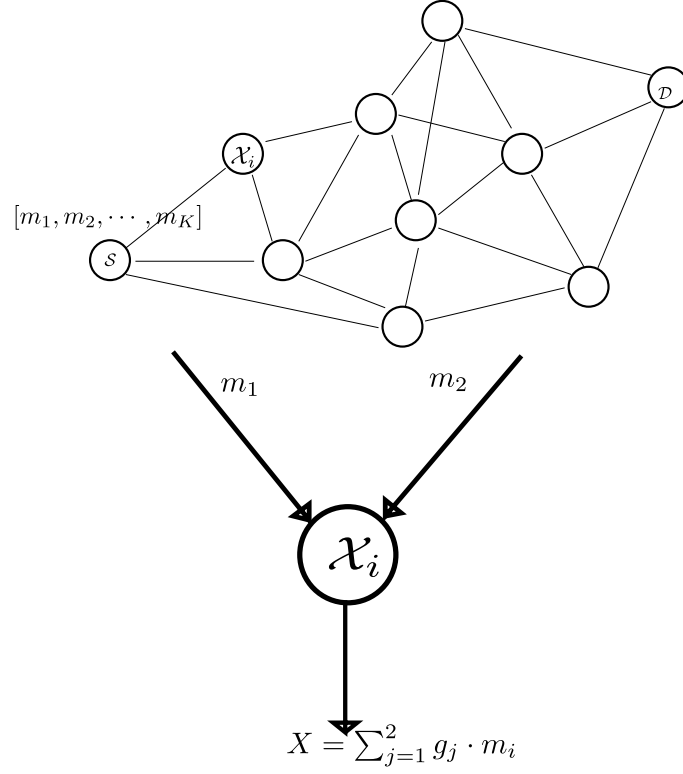


FIGURE 3.7: Codage réseau au niveau d'un nœud

Par exemple, un block M constitué de K paquets $m_1, \dots, m_K$ est transmis par la source, avec $m_i \in \mathbb{F}_q^N$. Avec un codage réseau linéaire, chaque nœud intermédiaire ($\mathcal{X}_i$) dans le réseau est associé à un vecteur de codage g_i de coefficients $g_{\{i,j\}} \in \mathbb{F}_q$ aléatoire. Le récepteur, pour être capable de décoder doit obtenir au moins K combinaisons linéaires indépendantes des paquets source.

Supposons que le récepteur, après les combinaisons linéaires effectuées dans le réseau reçoive n paquets, $Y = [Y_1, \dots, Y_n]$. Appelons A , la matrice de transition du réseau, qui résume toutes les combinaisons linéaires effectuées sur les paquets $m_1, \dots, m_K$.

$$y_i = \sum_{j=1}^K g_{\{i,j\}} m_j,$$

Des erreurs intentionnelles ou non (liés au bruit) sont injectées dans le réseau, notons par $Z = Z_1, \dots, Z_L$ les paquets d'erreurs. Dans ce cas :

$$y_i = \sum_{j=1}^K g_{\{i,j\}} \cdot m_j + \sum_{j=1}^L b_{\{i,j\}} \cdot Z_j$$

Ceci s'écrit sous forme matricielle :

$$Y = G \times M + B \times Z,$$

où G et B correspondent à la transformation globale à travers le codage réseau subit respectivement par X et Z . La topologie du réseau va certainement imposer la structure des matrices G et B .

Comme G est une matrice aléatoire, seul l'espace vectoriel des lignes de M peut être conservé lors de la transmission à travers le réseau. Les auteurs dans [76] ont envisagé la transmission de l'information non via le choix de M , mais plutôt par le choix de l'espace vectoriel engendré par les lignes ou les paquets de M . Ils ont ainsi proposé un modèle de canal qui offre l'abstraction d'un tel scénario d'injection de sous-espace vectoriel. Ils ont proposé ensuite une construction de codes à partir d'un tel modèle en se basant sur l'extension des codes Gabidulin.

3.5 Conclusion

Dans ce chapitre nous avons introduit les concepts et principes de base des codes correcteurs d'erreurs utilisés sur la couche PHY et/ou dans les couches supérieures des réseaux de télécommunication.

Des techniques de codage directe (FEC) usuels comme les codes RS, convolutifs, LDPC et plus atypiques comme les codes à métrique rang ont été définis. D'autres techniques de codage adaptées aux canaux à effacements comme les codes fontaines ont également été analysées.

En se basant sur ces pré-requis, les chapitres suivants sont consacrés au design de systèmes robustes face aux perturbations rencontrées sur les canaux CPL-BE dans le contexte d'une transmission point-à-point et multi-point à point.

Chapitre 4

Codes à métrique rang pour les CPL-BE

Sommaire

4.1	Introduction	62
4.2	Description du système proposé	62
4.2.1	Description du mappage	62
4.2.2	Motifs d'erreurs	63
4.2.3	Construction d'un code Gabidulin sur $\mathbb{F}_{2^4}$	67
4.2.4	Code Gabidulin concaténé avec un code convolutif	68
4.3	Performances du schéma proposé	70
4.3.1	Résultats de simulation	70
4.3.2	Analyse et comparaison de la complexité des codes RS et Gabidulin	76
4.4	Conclusion	78

4.1 Introduction

Le bruit à bande étroite et le bruit impulsif sont considérés comme les plus problématiques pour les transmissions sur les systèmes CPL-BE. Pour lutter contre le bruit impulsif et/ou les interférences à bande étroite dans le CPL-BE, les schémas classiques rencontrés dans la littérature reposent sur l'utilisation des codes à métrique de Hamming comme les codes RS et/ou les codes convolutifs, les codes à répétition ou les codes à permutation [77]. Parmi ces codes, seuls les codes à permutation sont immunisés contre les erreurs de type “criss-cross”¹ rencontrées sur les canaux CPL-BE. Les codes à permutation sont efficaces pour gérer les erreurs de type “criss-cross”, mais au détriment d’une transmission avec un faible rendement.

Nous abordons dans ce chapitre le design et les performances d’un schéma robuste face aux motifs d’erreurs de type criss-cross rencontrés dans le CPL-BE. Ce schéma est constitué d’un code en bloc à métrique rang (codes de Gabidulin) concaténé avec un code convolutif.

Ce chapitre comprend deux parties. Dans une première partie, nous décrivons le système proposé. La seconde partie présente les performances du schéma proposé et sa complexité comparées aux architectures déjà implémentées.

4.2 Description du système proposé

4.2.1 Description du mappage

Les signaux des systèmes de transmission multi-porteuses peuvent être « naturellement » représentés sous une forme de matrice, où chaque colonne de la matrice est utilisée pour générer un symbole OFDM. Prenons un vecteur $\mathbf{V}$ d’éléments dans le corps d’extension $\mathbb{F}_{q^N}$, $\mathbf{V} = (\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_n)$, le signal à envoyer. $\mathbb{F}_{q^N}$ peut être considéré comme un espace vectoriel sur $\mathbb{F}_q$; appelons $\mathcal{B}$ une base de $\mathbb{F}_{q^N}$ sur $\mathbb{F}_q$. Maintenant, nous pouvons représenter le vecteur $\mathbf{V}$

1. Lorsque les données transmises sont présentées sous forme de tableaux, et que les erreurs sont confinées sur des lignes et/ou des colonnes : ces motifs d’erreurs sont appelés erreurs criss-cross.

comme une matrice $\underline{\mathbf{V}}$ avec des entrées dans le corps fini $\mathbb{F}_q$ en décomposant les éléments $\mathbf{v}_i$ de $\mathbf{V}$ suivant la base $\mathcal{B}$,

$$\underline{\mathbf{V}} = \left[\begin{array}{c} \begin{bmatrix} v_{11} \\ v_{21} \\ \vdots \\ v_{N1} \end{bmatrix} \begin{bmatrix} v_{12} \\ v_{22} \\ \vdots \\ v_{N2} \end{bmatrix} \cdots \begin{bmatrix} v_{1n} \\ v_{2n} \\ \vdots \\ v_{Nn} \end{bmatrix} \end{array} \right]. \quad (4.1)$$

Les éléments de la matrice $\underline{\mathbf{V}}$ sont ensuite modulés avec des techniques de modulation appropriées. Le schéma de modulation dépend de q , le nombre d'éléments dans le corps de base. Par exemple pour $q = 2$, nous utiliserons une modulation binaire.

En utilisant un code de Gabidulin défini sur $\mathbb{F}_{2^N}$, le vecteur de code est défini par :

$$\mathbf{c} = \mathbf{m} \times G = (\mathbf{c}_1, \mathbf{c}_2, \dots, \mathbf{c}_n),$$

où $\mathbf{m} = (\mathbf{m}_1, \dots, \mathbf{m}_k) \in (\mathbb{F}_{2^N})^k$ est le message à envoyer et G la matrice génératrice. Nous pouvons présenter le vecteur $\mathbf{c}$ comme étant une matrice $\mathbf{C}$ avec des entrées dans $\mathbb{F}_2$;

$$\mathbf{C} = \begin{pmatrix} c_{1,1} & c_{1,2} & \cdots & c_{1,n} \\ c_{2,1} & c_{2,2} & \cdots & c_{2,n} \\ \vdots & \vdots & \ddots & \vdots \\ c_{N,1} & c_{N,2} & \cdots & c_{N,n} \end{pmatrix}.$$

Nous appliquons ici une modulation binaire, en l'occurrence le BPSK. Le modulateur est conçu pour assurer la transmission de la séquence codée avec une largeur de bande efficace.

4.2.2 Motifs d'erreurs

En supposant que le récepteur est muni d'un détecteur à seuil, le signal reçu à travers le canal est démodulé pour générer une matrice de la même taille que la matrice transmise.

Dans notre modèle d'erreurs, nous allons supposer que les entrées de

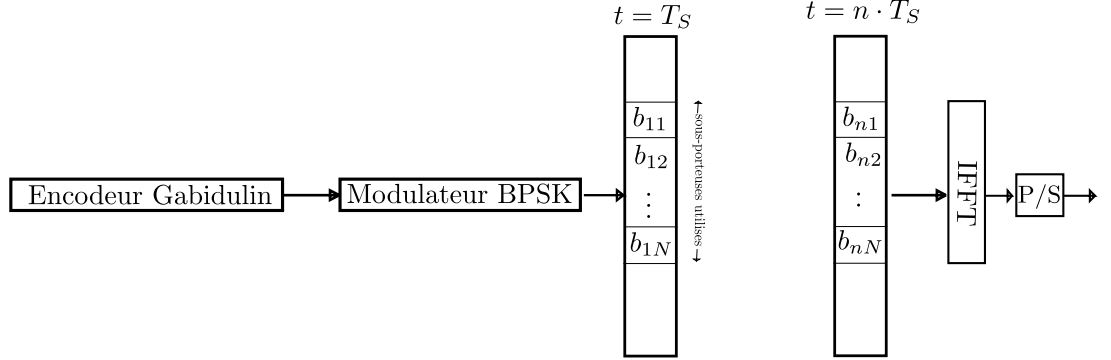


FIGURE 4.1: Chaîne de transmission du code à métrique rang Gabidulin

la matrice transmise $\mathbf{C}$ ont subit les effets du canal, résultant à une autre matrice $\bar{\mathbf{C}}$. La matrice d'erreurs est alors définie par $\mathbf{E} = \mathbf{C} - \bar{\mathbf{C}}$. Les erreurs criss-cross surviennent lors des événements du bruit impulsif ou du bruit à bande étroite et affectent la matrice reçue de la façon suivante :

- On suppose que le bruit impulsif, lorsqu'il se produit, affecte toutes les sous-porteuses utilisées lors de la transmission d'un symbole OFDM. Cela se traduit par la matrice des erreurs qui contient des erreurs en rafales sur une colonne donnée. Dans notre analyse, nous supposons que l'opération FFT au côté récepteur randomise complètement le bruit impulsif sur le symbole OFDM.
- La description donnée par [78] de l'interférence à bande étroite, où le signal d'interférence pourrait être modélisée comme affectant une ou plusieurs sous porteuses OFDM et où la fréquence et la phase du signal interférent a été supposée stochastique avec une distribution uniforme, a été retenue pour les fins de simulations. Lors de la survenue du bruit à bande étroite, pendant un certain nombre d'instances de temps, une puissance est introduite sur une ou plusieurs sous-porteuses données. La matrice des symboles OFDM transmise sera donc affectée par ce bruit dans la sous-porteuse correspondante. Cela se traduit par la matrice des erreurs qui contient des erreurs en rafales sur une ligne donnée.

En plus des erreurs criss-cross, le bruit de fond inverse la valeur d'un bit ²

2. qui n'est pas touché par le bruit impulsif ou le bruit à bande étroite

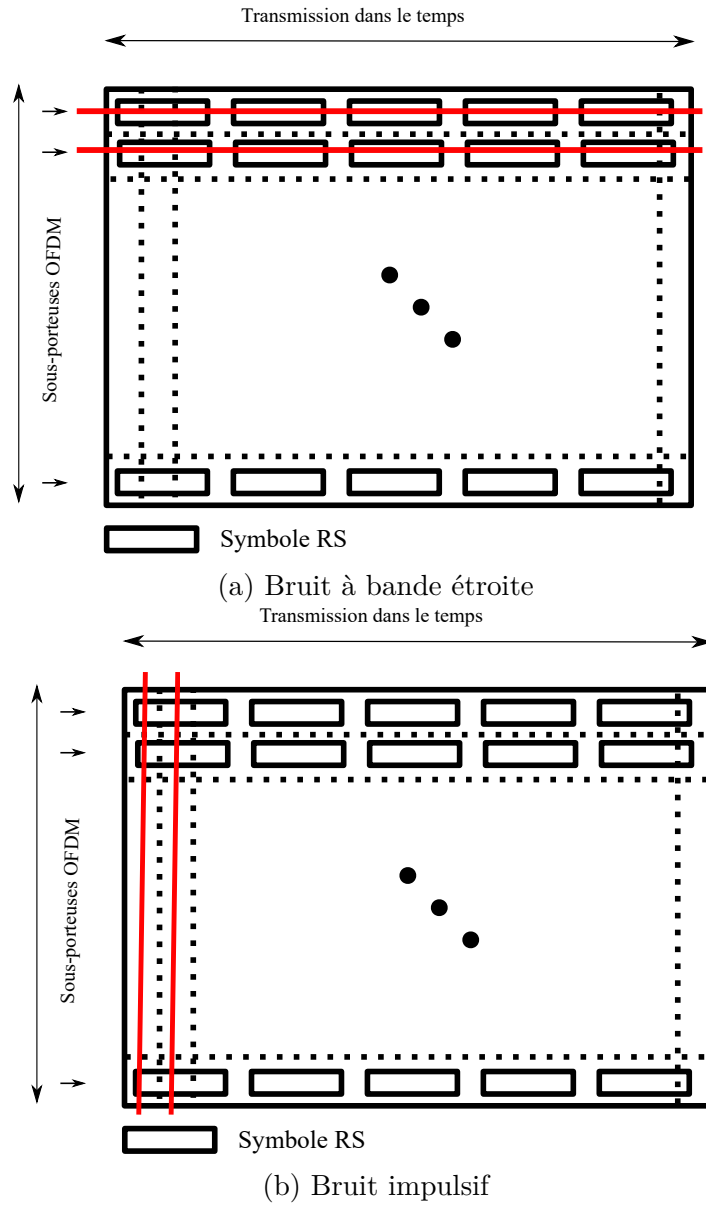


FIGURE 4.2: Occurrence d'erreurs dûs au bruit impulsif ou au bruit à bande étroite

transmis en temps/fréquence avec une probabilité $P_b = Q\left(\sqrt{\frac{2E_b}{N_0}}\right)$, où N_0 est la densité spectrale du bruit de fond et E_b l'énergie par symbole binaire transmis.

Des campagnes de mesure ont permis de montrer que le bruit impulsif périodique a un temps d'inter-arrivée de 8 ms (pour un réseau à 60 Hz) ou 10 ms (pour un réseau à 50 Hz). Pour un temps d'inter-arrivée de 8 ms, la durée de bruit impulsif varie entre 0,5 ms et 2.8 ms ; et pour un temps d'inter-arrivée de 10 ms, la durée de bruit impulsif varie de 0,625 ms à 3.5 ms. Ces durées permettent de calculer le nombre d'occurrence du bruit impulsif pendant la durée totale de transmission d'un mot de code lors des simulations. Les auteurs dans [79] ont par ailleurs proposé un modèle d'interférence à bande étroite pour les CPL-BE adapté aux systèmes OFDM.

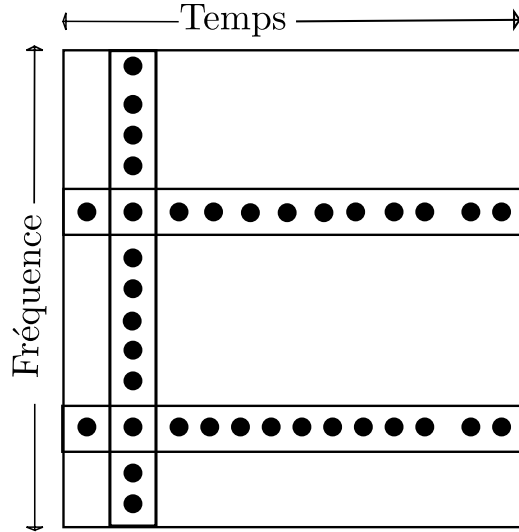


FIGURE 4.3: Motifs d'erreurs criss-cross rencontrés sur le canal CPL-BE

La figure 4.3 donne une représentation des motifs d'erreurs lors de l'envoi des données sous forme de matrices transmises en temps/fréquence. Elle permet de mettre en évidence les motifs d'erreurs criss-cross rencontrés dans le canal CPL-BE.

4.2.3 Construction d'un code Gabidulin sur $\mathbb{F}_{2^4}$

On construit un code Gabidulin $C(4,3)$ défini sur le corps $\mathbb{F}_{2^4}$, de dimension 3, de longueur 4 et de rendement $3/4$. $\mathbb{F}_{2^4} = \mathbb{F}_2(\alpha)$ avec α tel que $\alpha^4 = \alpha + 1$. Nous choisissons comme base de $\mathbb{F}_{2^4}$ sur $\mathbb{F}_2$ la base $g = (1, \alpha, \alpha^2, \alpha^3)$.

La matrice génératrice du code de Gabidulin est définie par :

$$G = \begin{pmatrix} 1, \alpha, \alpha^2, \alpha^3 \\ 1, \alpha^2, \alpha^4, \alpha^6 \\ 1, \alpha^4, \alpha^8, \alpha^{12} \end{pmatrix}.$$

Soit $\mathbf{m} = (\alpha, 1, 0)$ le message à envoyer, $\mathbf{c} = \mathbf{m} \times G = (1 + \alpha, 0, 1 + \alpha + \alpha^3, 1 + \alpha + \alpha^2 + \alpha^3)$, nous formons la matrice 4×4 de code C défini sur $\mathbb{F}_2$ en décomposant les c_i sur la base de $\mathbb{F}_{2^4}/\mathbb{F}_2$.

On obtient

$$C = \begin{bmatrix} 1 & 0 & 1 & 1 \\ 1 & 0 & 1 & 1 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 1 \end{bmatrix}.$$

Par conséquent, de cette manière, nous obtenons $C^0, C^1, C^2, C^3, \dots, C^{(2^4)^3-1}$ mots de codes, alors qu'il y'a $(2^4)^4$ possibilités.

Dans ce contexte, on peut se rendre compte de l'effet de divers types de bruit sur une petite trame, transmise en temps (le long des colonnes) et en fréquence (le long des lignes). Ces matrices correspondent aux configurations d'erreurs, la valeur "e" correspond à l'erreur et des emplacements 0 indiquent l'absence d'erreur.

0	e	0	0	0	0	e	0	0
0	0	0	0	0	0	e	0	0
e	0	0	0	0	0	e	0	0
0	0	0	0	0	0	e	0	0
Bruit de fond				Bruit impulsif				

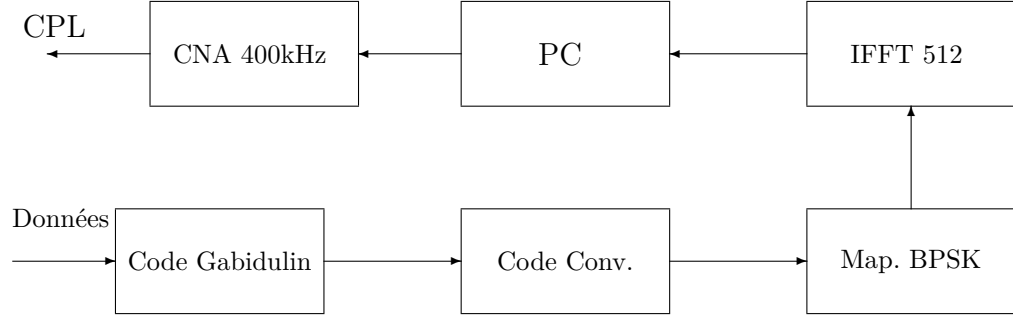


FIGURE 4.4: Schéma d'un code concaténé : code Gabidulin extérieur et convolutif intérieur.

$$\begin{array}{cccc}
 0 & 0 & 0 & 0 \\
 \text{e} & \text{e} & \text{e} & \text{e} \\
 0 & 0 & 0 & 0 \\
 0 & 0 & 0 & 0
 \end{array}
 \begin{array}{cccc}
 0 & 0 & 0 & 0 \\
 0 & 0 & 0 & 0 \\
 \text{e} & \text{e} & \text{e} & \text{e} \\
 0 & 0 & 0 & 0
 \end{array}
 .$$

Bruit BE Fading

4.2.4 Code Gabidulin concaténé avec un code convolutif

Les codes à métrique rang sont optimaux pour la correction des erreurs criss-cross, mais ils sont inadaptés pour gérer les erreurs isolées.

La figure 4.5 illustre une matrice de code Gabidulin de taille $k_v \times k_h$ d'élément binaires, sur laquelle on applique le code produit composé de deux codes convolutifs. Nous proposons d'utiliser un code produit convolutif de taille $n_v \times n_h$, formé de la façon suivante :

Un premier code $\mathcal{C}_h [n_h, n_h - r_h]$ permet de corriger la matrice transmise contre les erreurs qui surviennent sur les lignes, et un code $\mathcal{C}_v [n_v, n_v - r_v]$ est appliqué aux colonnes pour lutter contre les erreurs en rafales sur celles-ci. n_h, n_v sont les longueurs des codes et r_h, r_v les niveaux de redondances utilisées.

La figure 4.4 présente la chaîne considérée. Un mappage simple décrit à la figure 4.6 est considéré dans nos simulations.

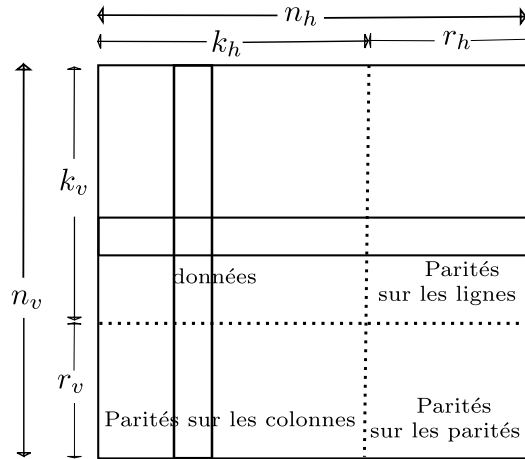


FIGURE 4.5: Code produit intérieur appliqué aux matrices transmises

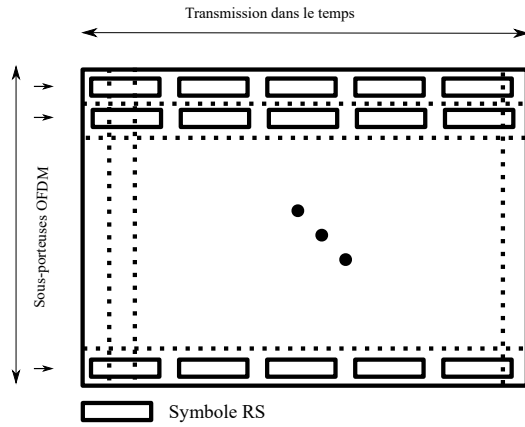


FIGURE 4.6: Mappage des symboles du code RS pour la transmission multi-porteuses

4.3 Performances du schéma proposé

4.3.1 Résultats de simulation

4.3.1.1 Comparaison codes à métrique rang et RS

Nous présentons les résultats de simulation qui montrent la performance du schéma de codage proposé. Nous avons simulé le canal CPL-BE en tenant compte de toutes ses caractéristiques mentionnées dans le premier chapitre. Les standards en vigueur utilisent souvent des rendements pour le codage de l'ordre de $\frac{1}{2}$ (typiquement la concaténation série d'un code convolutif et d'un code RS, quand ceux-ci ne sont pas associés à un code à répétition). Dans nos simulations nous utilisons le même rendement.

Nous comparons ainsi un code Gabidulin $(46, 23)$ sur $\mathbb{F}_{2^{46}}$ avec un code RS $(255, 127)$ sur $\mathbb{F}_{2^8}$. Le mot de code Gabidulin est une matrice (46×46) d'éléments binaires mappés en temps-fréquence avec la modulation OFDM. Nous notons que les deux codes sont de tailles (en bits) sensiblement égale. Dans les différents résultats de simulation, Rank Code (RC) (i, j) désigne un code à métrique rang avec i symboles OFDM affectés par le bruit impulsif et j sous-porteuses touchées par l'interférence à bande étroite. Le même principe de notation est adopté pour le code RS. Nous utilisons dans les simulations les modulations BPSK.

La figure 4.7 présente la performance du code à métrique rang par rapport au code RS en présence à la fois du bruit de fond et des interférences à bande étroite qui peuvent affecter jusqu'à 4 sous-porteuses. Tout d'abord, nous comparons les deux codes en l'absence de bruit impulsif et sans interférence à bande étroite ; c'est-à-dire $RC(0, 0)$ et $RS(0, 0)$: la seule perturbation étant le bruit de fond. Il résulte de la figure 4.7 que le code RS est d'environ 2,8 dB meilleur que le code à métrique rang à un TEB de 10^{-4} . En effet, la structure d'erreurs dans les matrices codées par un code à métrique rang n'est pas optimale quand il y'a peu de bruit à bande étroite et/ou de bruit impulsif. Ceci est dû à la présence des erreurs isolées qui augmentent de façon non optimale le rang de la matrice d'erreurs et réduisent par conséquent les possibilités de correction du code à métrique rang. Les codes à métrique rang

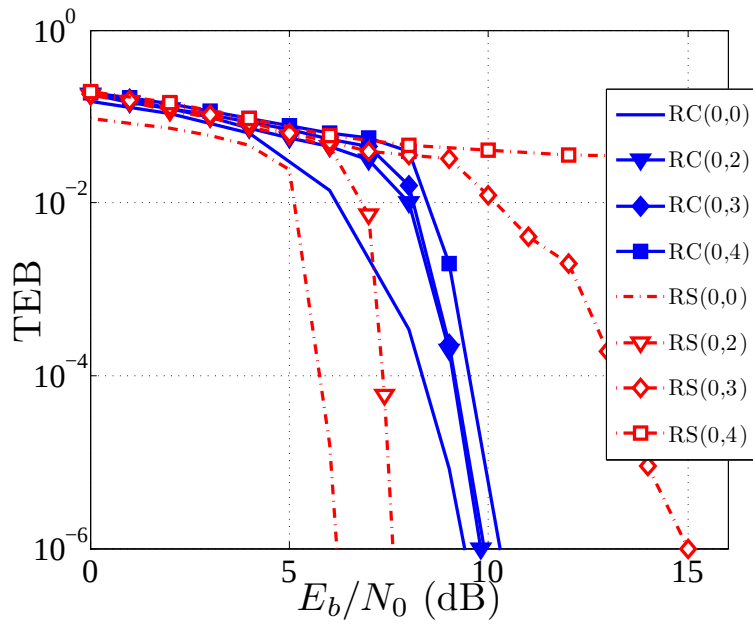


FIGURE 4.7: TEB du code Gabidulin comparé au code RS avec différents nombres de sous-porteuses affectées par l'interférence à bande étroite.

sont plus efficaces lorsque les erreurs sont confinées dans les lignes et/ou dans les colonnes.

Le gain de codage du code RS n'est pas très important par rapport à celui du code à métrique rang, mais la performance du code RS commence à se détériorer de façon significative lorsque trois sous-porteuses sont affectées par le bruit à bande étroite. Lorsqu'il y a quatre sous-porteuses affectées par l'interférence à bande étroite, la performance du code RS devient très dégradée et le code RS n'est plus en mesure de corriger les erreurs de manière efficace. En outre, on constate une faible sensibilité du code à métrique rang avec le nombre croissant d'évènements de bruit à bande étroite.

La figure 4.8 présente la performance du code à métrique rang par rapport au code RS en présence de bruit impulsif qui peut affecter jusqu'à 10 symboles OFDM. Le choix du nombre de symboles OFDM touchés par le bruit impulsif

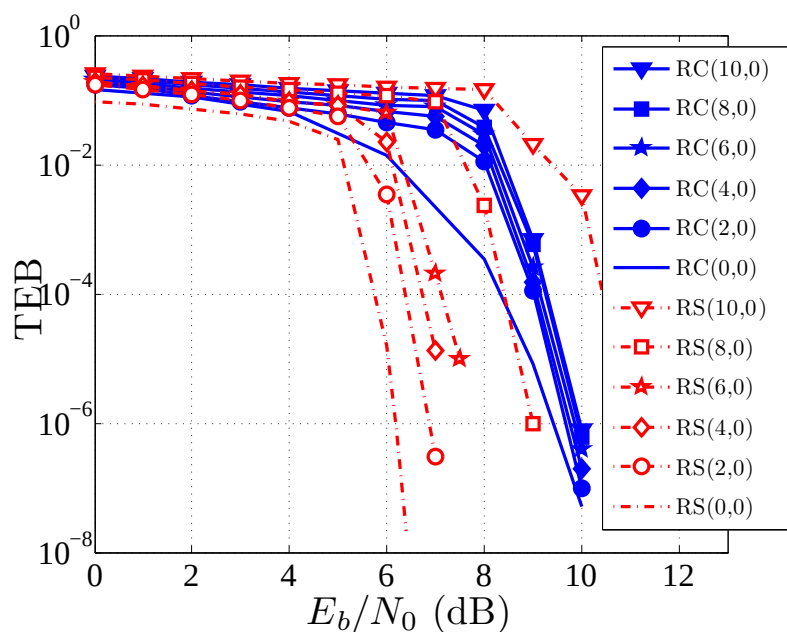


FIGURE 4.8: TEB du code Gabidulin comparé au code RS avec différents nombres de symboles OFDM affectés par le bruit impulsif

dans la simulation est en relation avec le taux d'arrivée du bruit impulsif, la durée de celui-ci et la longueur totale du mot de code transmis. Nous

observons que les codes RS sont environ 2 dB meilleurs que les codes à métrique rang lorsque le nombre de symboles OFDM touchés par le bruit impulsif est de moins de 6. Ceci est prévisible compte tenu du mappage adopté pour les codes RS lors de la transformation série parallèle (S/P) avant la transmission sur le canal. Cependant l'écart entre les deux codes diminue avec le nombre d'évènements du bruit impulsif. L'écart est significativement réduit en présence de 8 symboles OFDM affectés par le bruit impulsif et lorsque le nombre de symboles OFDM affectés par le bruit impulsif est supérieur à 8, le code à métrique rang est nettement meilleur que le code RS.

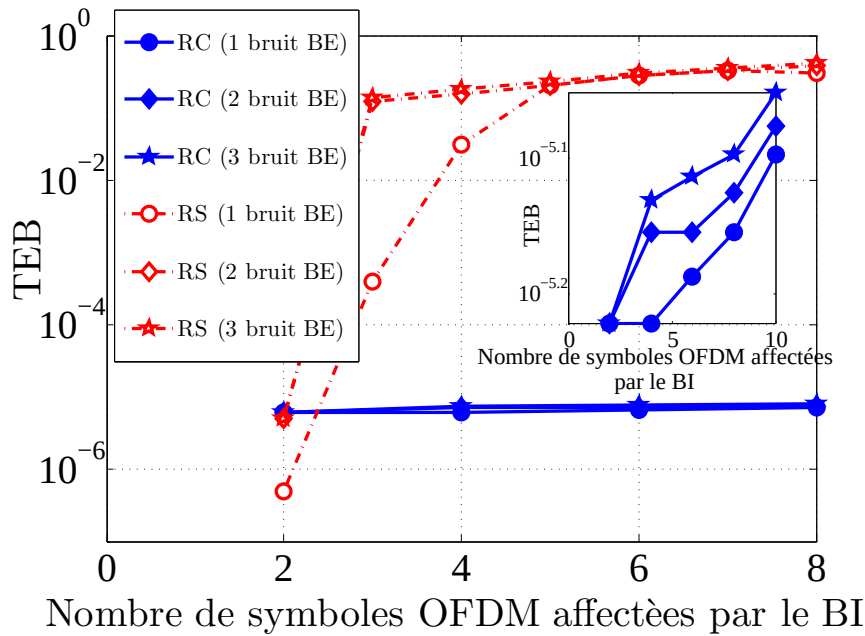


FIGURE 4.9: TEB du code Gabidulin comparé au code RS avec différents nombres de sous-porteuses affectées par l'interférence à bande étroite et en présence de bruit impulsif

La figure 4.9 présente les courbes du TEB des deux codes en fonction du nombre de lignes affectées par le bruit impulsif et de colonnes par l'interférence à bande étroite - en même temps - pour un SNR de 8 dB. Nous

nous rendons compte que la performance du code à métrique rang varie légèrement en fonction du nombre d'évènements du bruit impulsif et ceci quel que soit le nombre d'interférences à bande étroite touchant la matrice transmise. Ici, nous percevons le bon comportement du code à métrique rang en présence simultanée de bruit impulsif, de bruit à bande étroite et de bruit de fond touchant le canal CPL-BE.

4.3.1.2 Comparaison codes concaténés RC-convolutif et RS-convolutif

Ici, la comparaison est effectuée entre un code RS (255, 215) concaténé en série avec un code convolutif de rendement $\frac{1}{2}$, de longueur de contrainte 7. Par ailleurs, un procédé d'entrelacement temps-fréquence (deux dimensions) est utilisé avant le code convolutif pour réduire l'impact des erreurs en rafales qui sont difficiles à corriger par les codes convolutifs [67]. Ces paramètres sont choisis conformément aux normes CPL-BE, et avec les objectifs d'équité dans la comparaison entre code à métrique rang et code RS. Nous nous assurons d'avoir le même rendement pour les codes comparés. Nous comparons ainsi le schéma explicité plus haut avec un code Gabidulin (46,38) concaténé avec un code produit convolutif.

En l'absence d'interférence à bande étroite ou de bruit impulsif (voir figure 4.10), nous notons que les deux codes opèrent presque à l'identique. Ceci est expliqué par le fait que le code convolutif intérieur arrive à gérer les erreurs introduites par le bruit de fond. Les paquets d'erreurs résiduels sont gérés par les codes en bloc extérieurs (Gabidulin ou RS).

En présence de bruit impulsif (voir figure 4.11), les deux codes ont à peu près la même performance. Ceci s'explique par le fait que les deux codes comparés (Gabidulin et RS) sont deux codes en bloc qui ont des capacités semblables quand à la gestion des erreurs qui affectent une seule dimension de la matrice transmise.

Toutefois, en présence de l'interférence à bande étroite et du bruit de fond (voir figure 4.12), on peut voir que le code concaténé Gabidulin-convolutif surpasse le système constitué du code convolutif, du code RS et

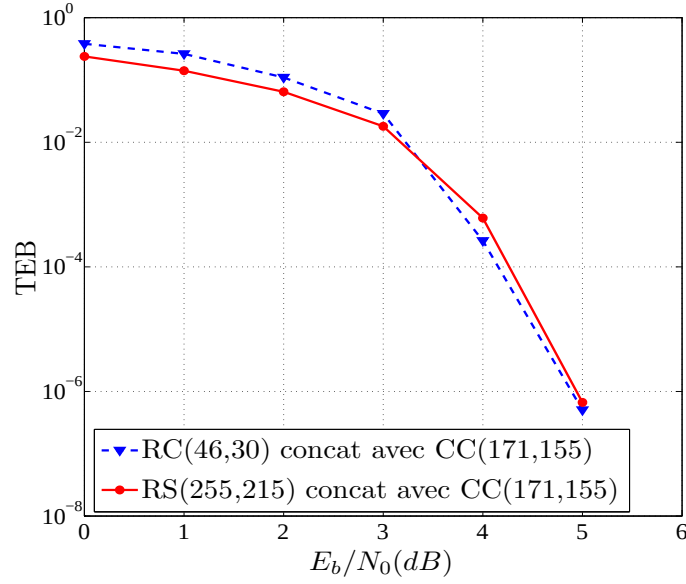


FIGURE 4.10: TEB des codes concaténés Gabidulin-Convolutif et RS-Convolutif en présence uniquement du bruit de fond

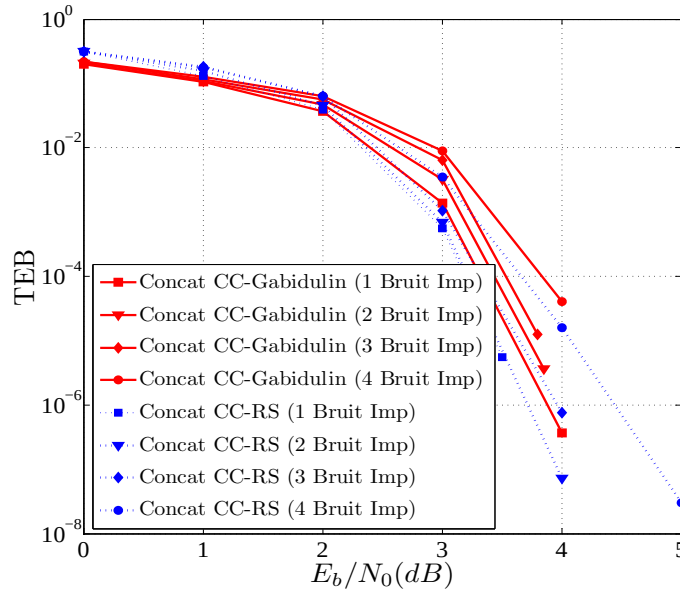


FIGURE 4.11: TEB des codes concaténés Gabidulin-Convolutif et RS-Convolutif en présence de différents nombres de symboles OFDM affectés par le bruit impulsif

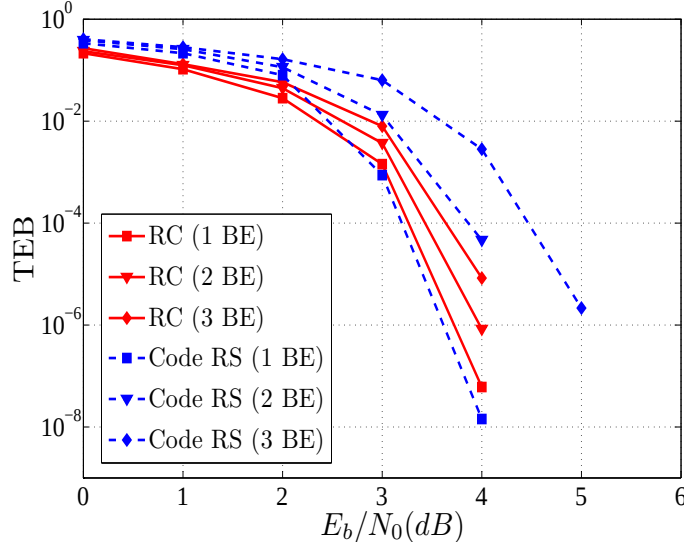


FIGURE 4.12: TEB des codes concaténés Gabidulin-Convolutif et RS-Convolutif en présence de différents nombres de sous-porteuses affectées par l'interférence à bande étroite

de l'entrelaceur. La survenue de l'interférence à bande étroite (qui affecte les lignes de la matrice transmise) est correctement corrigée par le code Gabidulin extérieur si le code convolutif intérieur n'a pas réussi à le corriger.

4.3.2 Analyse et comparaison de la complexité des codes RS et Gabidulin

La comparaison effectuée des deux codes nécessite aussi une analyse de leur complexité de décodage. En effet pour les codes algébriques que sont les codes RS et les codes Gabidulin, toute la complexité de la chaîne de transmission réside dans l'opération de décodage. Pour les codes RS (resp. Gabidulin) si n (resp. N) est le nombre de symboles par mot de code, m le nombre de bits par symbole, k la dimension du code et t son pouvoir de correction. Nous obtenons pour les paramètres considérés que les deux codes fonctionnent avec sensiblement la même complexité, voir tableau 4.1.

Tableau 4.1: Analyse des complexités de décodage des codes RS et Gabidulin [2]

Le code RS	Le code Gabidulin
Paramètres typiques du code RS $q = 2, n = 255,$ $m = 8, t = 64$	Paramètres typiques du code RC $q = 2, N = 46,$ $m = 46, t = 12$
Complexité de décodage standard $\mathcal{O}(tnm^2)$ dans $\mathbb{F}_q$	Complexité de décodage standard $\mathcal{O}(tNm^2)$ dans $\mathbb{F}_q$

☞ Remarque : Une autre problématique résolue par les codes à métrique rang concerne les effacements³ et la propagation des erreurs pour le codage réseau linéaire aléatoire. Le problème du contrôle d'erreur pour le codage réseau linéaire aléatoire a été considéré dans [76] par Koetter et Kschischang. Une transmission non cohérente est supposée où ni l'émetteur ni le récepteur ne sont supposés avoir connaissance de la fonction de transfert du réseau. Si l'on voit les paquets envoyés comme des vecteurs d'un certain espace vectoriel, le codage réseau à l'avantage de toujours préserver le sous espace vectoriel, pour un ensemble de paquets transmis. Ainsi ce codage par injection de sous-espace vectoriel est bien adapté à la métrique rang pour un codage réseau aléatoire.

3. Pour le codage réseau linéaire aléatoire les effacements arrivent quand la destination n'a pas reçu suffisamment de paquets pour décoder

4.4 Conclusion

Nous avons proposé un système qui est robuste face au bruit impulsif et à l'interférence à bande étroite. Nous avons étudié la performance des codes à métrique rang avec la mise en œuvre complète du codage et du décodage en tenant compte de l'environnement bruyant du canal CPL-BE.

Tout d'abord, le code Gabidulin $(46, 23)$ défini sur $\mathbb{F}_{2^{46}}$ est mis en œuvre et nous avons comparé ses performances avec un code RS $(255, 127)$. Ensuite le code concaténé Gabidulin-convolutif est comparé avec le système constitué d'un code RS associé à un code convolutif et un entrelaceur. Les résultats de simulation indiquent que dans les conditions de canal et de bruit considérés le schéma proposé possède des performances bien meilleures que le code RS utilisé dans la norme CPL-BE G3.

Chapitre 5

Stratégies de relayage de codes fontaines pour les réseaux CPL-BE

Sommaire

5.1	Introduction	80
5.2	Relayage coopératif de codes fontaines : sans codage réseau	80
5.2.1	Codage coopératif	82
5.2.2	Stratégies de relayage	82
5.3	Codes fontaines combinés avec le codage réseau pour les CPL-BE	90
5.3.1	Revue de la littérature du codage réseau pour les CPL-BE	91
5.3.2	Optimisation de la matrice de probabilité jointe . .	101
5.3.3	Génération du degré de sortie	102
5.3.4	Algorithme de ré-allocation	103
5.3.5	Résultats de simulation	108
5.4	Conclusion	113

5.1 Introduction

Dans ce chapitre, le canal CPL-BE est considéré comme étant un canal à effacement de trames. Les schémas de retransmission traditionnels de trames effacées peuvent conduire à un nombre élevé de retransmissions et d'acquittements, avec comme conséquence une congestion du réseau CPL-BE. Avec leurs performances quasi-optimales (sur des canaux à effacement de trames), les codes fontaines [80, 69] qui ont déjà été adoptés dans divers standards (3GPP par exemple [81]) peuvent donc constituer un candidat potentiel en tant que protocole de transfert de données et/ou pour protéger les messages sur les canaux à effacements CPL-BE [82]. Ce chapitre est dédié à la conception et à l'analyse de protocoles de transmission de codes fontaines dans les réseaux CPL-BE pour des scénarios multi-saut. Il comprend deux parties.

La première partie traite du relayage coopératif de codes fontaines ; il s'agit ici d'une seule source qui émet vers une destination et qui est aidée par plusieurs relais intermédiaires lors de la transmission.

Dans la seconde partie, nous présentons des algorithmes pour combiner le codage fontaine et le codage réseau pour la topologie spécifique des réseaux CPL-BE.

Pour ces deux parties, des résultats de simulation confirmant les bonnes performances des méthodes proposées sont présentés.

5.2 Relayage coopératif de codes fontaines : sans codage réseau

Compte tenu des mauvaises conditions du canal CPL-BE, les messages à travers le réseau ont souvent besoin d'être transmis en plusieurs sauts entre un émetteur et un récepteur. Ceci permet de limiter les pertes de trajet, de fournir plus de capacité et de réduire les niveaux d'interférence. Nous étudions un système de communication unicast coopératif, où les transmissions entre la source et les relais et entre les relais et la destination utilisent des codes



FIGURE 5.1: Relayage de codes fontaines sur un réseau mono-chemin

fontaines. La source transmet des mots de code fontaine à une destination via un certain nombre de relais qui emploient un système « decode and forward ».

Le modèle de canal relais est représenté sur la figure 5.1. Il présente le schéma le plus simple de coopération, qui constitue la base pour le routage multi-sauts, dans lequel une trame est acheminée de manière séquentielle à partir d'une source vers la destination à travers une série de bonds (relais).

Les nœuds intermédiaires entre la source et la destination qui peuvent entendre les trames provenant de la source assistent la source dans la transmission vers la destination. Nous fournissons une évaluation analytique de la probabilité de décodage à la destination en fonction du nombre d'utilisation du canal (M) pour un nombre différent de relais. Ces résultats sont ensuite validés par des simulations.

5.2.0.1 Revue de la littérature

Les auteurs dans [83] et [84] présentent et comparent différents schémas de transmission multi-sauts sur le réseau CPL-BE. Ils préconisent l'utilisation de la redondance incrémentale, où différents nœuds envoient différents symboles de parité à la destination. Une manière d'accomplir la redondance incrémentale est l'utilisation de codes fontaines. Ainsi, les auteurs dans [85] fournissent une formulation mathématique de la probabilité de décodage d'un schéma de redondance incrémentale basé sur les codes fontaines quand un relais est utilisé.

Dans cette première partie, nous continuons en ligne avec [85] et [83] et nous étudions les performances des schémas de redondance incrémentale basés sur les codes fontaines pour les réseaux CPL-BE.

5.2.1 Codage coopératif

Étant donné que les CPL-BE ont une structure de bus, on peut supposer que chaque utilisateur connecté au bus est en mesure d'entendre des trames transmises avec un taux de perte de trames plus ou moins important en fonction des conditions du canal. Par exemple, si le premier relais est en mesure de recevoir un nombre suffisant de trames pour décoder (K trames linéairement indépendantes non effacées), alors il est le premier « relais à succès ». Pendant le même temps, les autres relais ($\mathcal{R}_i$) ont également reçu un certain nombre $N_i \leq K$ de trames. Le deuxième relais sera alors assisté par le premier relais ayant décodé avec succès dans le décodage. En fait, le premier relais agit maintenant comme une source et envoie les x_2 trames manquantes pour que le deuxième relais réussisse à être en mesure de décoder. On a alors, pour le deuxième relais, $N_2 + x_2 = K$ trames linéairement indépendantes. Pour plusieurs relais (n), le processus est répété de façon itérative jusqu'à ce que la destination soit capable de décoder [86].

5.2.2 Stratégies de relayage

Nous estimons tout d'abord le taux d'erreurs trames entre deux nœuds $\mathcal{R}_i$ et $\mathcal{R}_j$ du réseau CPL-BE : $P_e^{R_i R_j}$. $P_e^{R_i R_j}$ est ensuite utilisé pour indiquer la probabilité d'effacement sur le canal entre les nœuds $\mathcal{R}_i$ et $\mathcal{R}_j$.

5.2.2.1 Transmission directe

En considérant d'abord une transmission directe de codes fontaines linéaires aléatoires de la source à la destination, le nombre de trames codées reçues non effacées à la destination $\mathcal{D}$ (noté N) est aléatoire et est lié au nombre de trames transmises (M , pour $0 < N \leq M$) par la distribution binomiale suivante :

$$\mathcal{B}_{M, P_e^{SD}}(N) = \binom{M}{N} (1 - P_e^{SD})^N P_e^{SD M-N}. \quad (5.1)$$

La loi de probabilité que le nombre de trames reçues permette un décodage au récepteur peut être écrite comme dans [87]. Il s'agit ici de la probabilité

que toutes les combinaisons linéaires reçues permettent de décoder les K messages envoyés.

$$f(N) = \begin{cases} 0 & \text{si } N < K \\ \prod_0^{K-1} (1 - 2^{i-N}) & \text{si } N = K \\ \frac{2^{-N}(2^K - 1)}{1 - 2^{K-N+1}} \prod_0^{K-1} (1 - 2^{i-N+1}) & \text{si } N > K \end{cases} \quad (5.2)$$

Ici un code fontaine linéaire aléatoire est utilisé pour des besoins de simplification, mais nous pouvons généraliser la démarche à d'autres distributions (Soliton idéale, Soliton Robuste (SR), etc).

Voici l'approche que nous avons suivi : pour le cas d'une transmission point-à-point sur un canal à effacement avec une probabilité d'effacement P_e . Pour que la destination soit capable de décoder au bout de M transmissions, il est évident alors que le dernier paquet (le M -ième paquet) clôt le processus et ne doit pas être effacé ; sinon on tomberait dans le cas d'un décodage réussi avant les M transmissions (qui n'est pas la probabilité recherchée). Donc à $M - 1$ transmissions, on dispose de i paquets reçus avec une probabilité $\mathcal{B}_{M-1, P_e}(i)$. Ajoutés au M -ième paquet qui est non effacé, il existe au total $i + 1$ paquets disponibles à la destination, et donc la probabilité de décodage avec succès est $f(i + 1)$. En sommant le tout (sur i), la probabilité pour que le décodage s'effectue en M transmissions peut s'écrire comme suit :

$$\begin{aligned} P(M, P_e) &= (1 - P_e) \times \sum_{i=K-1}^{M-1} \binom{M-1}{i} \times (1 - P_e)^i \times (P_e)^{M-1-i} \times f(i + 1) \\ P(M, P_e) &= (1 - P_e) \sum_{i=K-1}^{M-1} \mathcal{B}_{M-1, P_e}(i) \times f(i + 1). \end{aligned} \quad (5.3)$$

Où $\mathcal{B}(\cdot)$ et $f(\cdot)$ sont respectivement données dans les équations (5.1) et (5.2).

5.2.2.2 Utilisation d'un seul relais

Tout d'abord, rappelons les hypothèses avec lesquelles nous effectuons notre analyse :

- Les réseaux CPL ont une structure de bus. D’une part chaque utilisateur connecté au bus est en mesure d’entendre les trames transmises et d’autre part deux nœuds situés sur ce bus ne peuvent transmettre en même temps sans collisions.
- Le nœud source a K paquets qu’il encode avec un code fontaine et qu’il transmet vers la destination via un ou deux relais.
- Le relayage des codes fontaines s’effectue selon le principe decode and forward. Un relais reçoit un certain nombre de paquets (qui lui permette de décoder) et par la suite il remplace le nœud source dans la transmission vers la destination.
- Le nombre d’utilisations du canal entre la source et la destination est toujours M . Par exemple, j transmissions entre la source et le relais et $M - j$ transmissions entre le relais et la destination : l’utilisation du canal demeure M . Et pour ces M transmissions, nous évaluons de façon précise la probabilité que la destination soit capable de décoder avec l’utilisation d’un ou de deux relais.

Le premier terme

$$\left(\sum_{j \geq M} P^{(1)}(j, P_e^{SR}) \times P^{(1)}(M, P_e^{SD}) \right),$$

correspond à la probabilité que la destination décode sans avoir besoin du relais : ceci correspond à l’occurrence de deux événements : le relais n’arrive pas à décoder après avoir reçu jusqu’à M paquets (événement A) et la destination est capable de décoder à la réception de M paquets (événement B) ;

$$P(A \cap B) = P(A) \times P(B).$$

$P(\bar{A})$ représente la probabilité que le relais décode avec un nombre de paquets inférieur à M ,

$$\begin{aligned}
 P(A) &= 1 - P(\bar{A}), \\
 P(A) &= 1 - \sum_{j=0}^{M-1} P^{(1)}(j, P_e^{SR}), \\
 P(A) &= \sum_{j \geq M} P^{(1)}(j, P_e^{SR}),
 \end{aligned}$$

et

$$P(B) = P^{(1)}(M, P_e^{SD}).$$

Il existe deux scénarios possibles :

- Soit la destination décode sans avoir besoin du relais : cette probabilité à été détaillée plus haut.
- Soit la destination décode à l'aide d'un relais. Cette probabilité correspond à la somme

$$(1 - P_e^{RD}) \sum_{j < M} P^{(1)}(j, P_e^{SR}) \left(\sum_{s=0}^j \sum_{t=0}^{M-j-1} \mathcal{B}_{j, P_e^{SD}}(s) \mathcal{B}_{M-j-1, P_e^{RD}}(t) f(t+s+1) \right)$$

Toute l'analyse est effectuée sur la base du nombre de transmissions M : M transmissions source + relais. En fait quand un nœud transmet, il est seul à occuper le canal, d'où l'intérêt de minimiser le nombre de transmissions. Pour le scénario considéré : une source ($\mathcal{S}$), un relais ($\mathcal{R}$) et une destination ($\mathcal{D}$), la loi de probabilité qui décrit le succès du décodage au niveau de la destination ($\mathcal{D}$) après M transmissions (de la source et/ou du relais) est donnée par :

$$\begin{aligned}
 P^{(2)}(M, P_e^{SD}, P_e^{SR}, P_e^{RD}) &= \sum_{j \geq M} P^{(1)}(j, P_e^{SR}) \times P^{(1)}(M, P_e^{SD}) \\
 &\quad + (1 - P_e^{RD}) \sum_{j < M} P^{(1)}(j, P_e^{SR}) \\
 &\quad \left(\sum_{s=0}^j \sum_{t=0}^{M-j-1} \mathcal{B}_{j, P_e^{SD}}(s) \mathcal{B}_{M-j-1, P_e^{RD}}(t) f(t+s+1) \right).
 \end{aligned} \tag{5.4}$$

Dans l'équation 5.4, la première somme correspond au cas où la destination décode sans avoir eu besoin du relais, et la seconde somme notée par $P_r^{(2)}$ représente le cas où la destination décode avec l'aide d'un relais. Cette seconde somme est effectuée sur le nombre de trames j envoyées par la source à la fois à la destination et au relais jusqu'à ce que le relais puisse décoder. Et lorsque le relais décode, il complète la transmission en envoyant les $M - j$ trames manquantes (l'utilisation du canal reste toujours M).

5.2.2.3 Utilisation de deux relais

Dans ce qui suit on calcule la probabilité de décodage à la destination dans le cadre de l'utilisation de deux relais.

$$\begin{aligned}
 p^{(3)}(M) = & P^{(1)}(M, P_e^{SD}) \sum_{j \geq M} P^{(1)}(j, P_e^{SR_1}) \sum_{l \geq M} P^{(2)}(l) \\
 & + (1 - P_e^{R_1D}) \sum_{l \geq M} P^{(2)}(l) \sum_{j < M} P^{(1)}(j, P_e^{SR_1}) \left(\sum_{s=0}^j \sum_{t=0}^{M-j-1} \mathcal{B}_{j, P_e^{SD}}(s) \mathcal{B}_{M-j-1, P_e^{R_1D}}(t) \right. \\
 & \left. f(t + s + 1) \right) + \\
 & (1 - P_e^{R_2D}) \sum_{l \geq M} P^{(1)}(l, P_e^{SR_1}) \sum_{j < M} P^{(2)}(j) \left(\sum_{s=0}^j \sum_{t=0}^{M-j-1} \mathcal{B}_{j, P_e^{SD}}(s) \mathcal{B}_{M-j-1, P_e^{R_2D}}(t) \right. \\
 & \left. f(t + s + 1) \right) + \\
 & (1 - P_e^{R_2D}) \sum_{j < M} P_r^{(2)}(j) \times \sum_{p < j} P^{(1)}(p, P_e^{SR_1}) \\
 & \left(\sum_{t=0}^{M-j-1} \sum_{s=0}^{j-p} \sum_{l=0}^p \mathcal{B}_{p, P_e^{SD}}(l) \mathcal{B}_{j-p, P_e^{R_1D}}(s) \mathcal{B}_{M-j-1, P_e^{R_2D}}(t) f(l + t + s + 1) \right)
 \end{aligned} \tag{5.5}$$

Avec

$$\begin{aligned}
 P^{(2)}(j) &= P^{(2)}(j, P_e^{SR_2}, P_e^{SR_1}, P_e^{R_1R_2}), \\
 P_r^{(2)}(j) &= P_r^{(2)}(j, P_e^{SR_2}, P_e^{SR_1}, P_e^{R_1R_2}).
 \end{aligned}$$

Ici la première somme représente le cas où la destination décode avant les deux relais $\mathcal{R}_1$ et $\mathcal{R}_2$, les deux sommes suivantes représentent le cas où la destination décode soit avant $\mathcal{R}_1$ ou $\mathcal{R}_2$ et, enfin, la quatrième somme représente le cas où $\mathcal{R}_1$ décode avant $\mathcal{R}_2$ et $\mathcal{R}_2$ décode avant la destination.

Il est possible que $\mathcal{R}_2$ décode avant $\mathcal{R}_1$ malgré le fait que $\mathcal{R}_1$ soit plus proche de la source. Cela est possible à cause de la présence du bruit impulsif et/ou du bruit à bande étroite qui pourrait affecter uniquement $\mathcal{R}_1$. Cependant ces événements sont plutôt rares et ont été négligés dans le cadre de l'analyse analytique.

Nous négligeons l'overhead qui permet au relais d'informer l'émetteur que les trames transmises sont suffisantes pour le décodage. On ne tient pas compte non plus de la corrélation des probabilités d'effacement sur les différents liens¹.

Dans la suite de ce manuscrit l'overhead est défini pour mesurer les données de codage supplémentaires qui doivent être ajoutées aux données source afin d'assurer la récupération des données source. On suppose que la taille des données source est de K . L'overhead est mesuré comme étant le pourcentage de l'ensemble des données transmises sur celui des données source avant que le décodage ne soit possible. Dans les différentes simulations du taux d'erreurs trames, on tient compte du bruit impulsif. On suppose par ailleurs que les effacements ne proviennent que de l'altération des trames en raison du bruit de transmission. Dans ce cas, le taux d'effacements est égal au taux d'erreurs trames.

Afin d'avoir des valeurs réalistes pour les taux d'effacements e.g, les probabilités P_e^{SD} , P_e^{SR} , $P_e^{R_1D}$, $P_e^{R_2D}$ etc. , nous présentons dans la figure 5.3, le taux d'erreurs trames d'un réseau typique CPL-BE (PRIME) pour différentes distances (distance de la source à la destination et distance inter-nœuds) en fonction du Rapport Signal sur Bruit (RSB). Ces taux d'erreurs trames peuvent être utilisés pour déterminer les probabilités d'effacement $P_e^{R_iR_j}$ du canal entre R_i et R_j si la distance entre ces nœuds est connue.

1. En raison de la corrélation du bruit et de la structure en bus du canal, les probabilités d'effacement sur un même lien sont souvent corrélées.

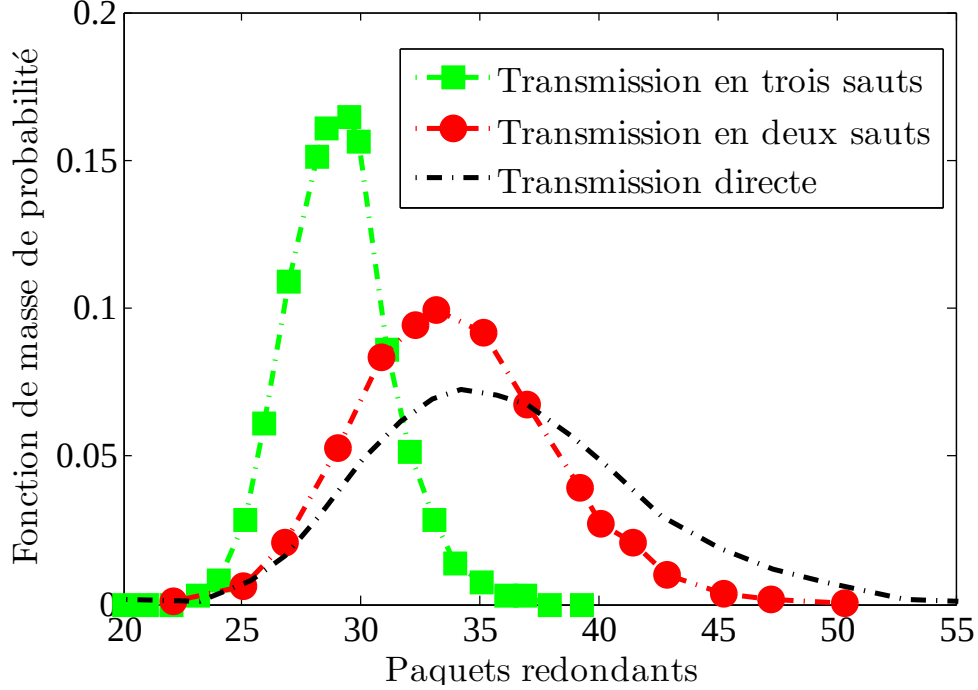


FIGURE 5.2: Loi de probabilité du nombre requis de transmissions (pour décoder) pour une transmission directe, ou avec l'utilisation d'un ou de deux relais ($K = 20$).

Nous fixons, par exemple, la distance inter-nœud à 100 m^2 et la distance de la source à la destination également à 300 m . Dans la figure 5.2, les probabilités d'effacement de trames sur les liens directs entre la source et la destination (lien trois-sauts), les liens à deux sauts, les liens à un saut ont été simulées pour une distance de 300 m (source-destination) avec une topologie type que celle présentée dans le chapitre 1.

La figure 5.2 met en évidence l'utilité du relayage coopératif des codes fontaines. On observe en effet une plus petite moyenne dans les transmissions requises pour les liens à deux sauts par rapport à ceux à un saut ou par rapport à une transmission directe.

2. Cette distance est très réaliste pour un réseau CPL-BE outdoor

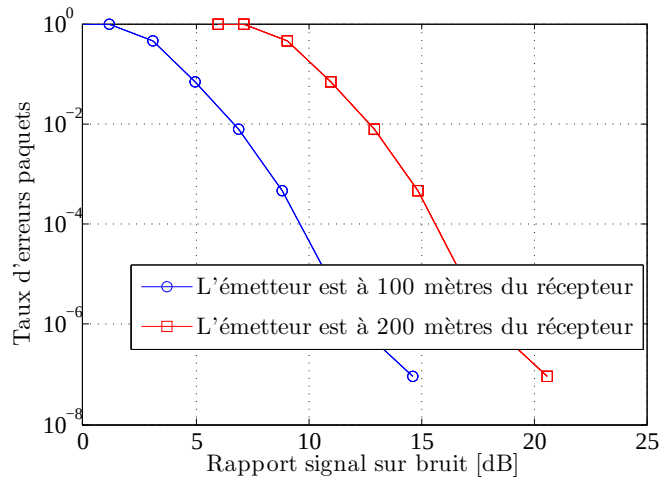


FIGURE 5.3: Taux d'erreurs sur les trames de PRIME avec différentes distances

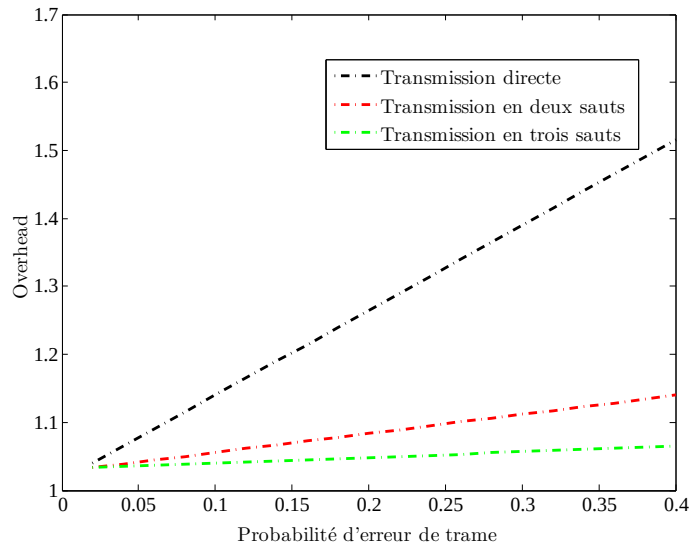


FIGURE 5.4: Probabilité d'échec du décodage en fonction de l'overhead sur un canal CPL pour différents schémas de transmission ($K = 20$)

La figure 5.4 donne la probabilité d'échec du décodage en fonction de l'overhead sur un canal CPL pour différents schémas de transmissions (utilisation d'un nombre différents de relais), et confirme que le schéma avec deux relais permet une meilleure utilisation du canal en particulier quand il y'a une augmentation du taux d'erreur de trame.

Le système ainsi proposé limite la redondance, améliore la fiabilité et réduit la puissance de transmission qui est une contrainte des systèmes CPL-BE.

5.3 Codes fontaines combinés avec le codage réseau pour les CPL-BE multi-sauts

Les codes fontaines sont asymptotiquement optimaux, ce qui signifie que plus la taille de la génération augmente plus le décodage est optimal en terme d'overhead [72]. Il est donc préférable pour une taille de génération équivalente ($2K$) d'utiliser un seul code fontaine de taille $2K$ que deux codes fontaines de tailles K . Par conséquent, avec le codage fontaine distribué, le codage réseau permet de générer ce code fontaine de taille $2K$ à partir de deux codes fontaines de taille K . Dans cette section, nous présentons quelques stratégies pour combiner le codage réseau et les codes fontaines pour certaines applications SG comme par exemple l'agrégation de données de mesure à une station de base.

Les transmissions multi-sauts sur un réseau à topologie radiale sont considérées et des algorithmes pour fusionner les données au niveau des nœuds relais (codage réseau inter-session) sont proposés. Nous concevons les algorithmes au niveau des relais afin d'assurer un bon décodage des codes fontaines. Il faut pour cela respecter certaines propriétés telles que les distributions d'entrée et de sortie des trames envoyées. Cela nous permet de bénéficier des avantages du codage réseau tout en évitant - dans le même temps - les inconvénients de la complexité de décodage quadratique en $\mathcal{O}(K^3)$.

Ainsi le problème des codes fontaines distribués est analysé pour

l'architecture particulière du réseau CPL-BE pour le SG, où un nœud source doit relayer d'autres trames d'autres nœuds source tout en ajoutant ses propres trames. Un algorithme de relayage basé sur la probabilité conjointe sur le degré reçu au relais et le degré sortant du relais est proposé. Deux approches pour le calcul des probabilités conjointes sont données.

5.3.1 Revue de la littérature du codage réseau pour les CPL-BE

Récemment, plusieurs articles ont exploré l'idée du codage réseau pour les applications des réseaux SG. Les auteurs dans [88] montrent que l'utilisation du codage réseau peut accroître la résilience des réseaux SG. Dans l'objectif d'obtenir plus de fiabilité, il est proposé dans [89] d'utiliser le codage réseau dans le transfert de données dans les réseaux SG et un algorithme pour accroître l'efficacité spectrale est proposé. Les auteurs dans [90] suggèrent également l'utilisation du codage réseau pour les CPL-BE en « indoor » dans le contexte des réseaux SG. Ils ont implémenté le codage réseau sur $\mathbb{F}_{2^8}$ et ont évalué sa performance pour différents protocoles ARQ. Les auteurs dans [91] ont proposé d'utiliser le codage réseau avec une stratégie de combinaison clairsemée (pas très dense) au départ pour limiter la complexité du décodage et une recombinaison plus dense à la fin, afin de couvrir toutes les trames source. Ils ont comparé leurs schémas avec les stratégies maître/esclave de collecte de données de mesure actuelles et ont montré qu'au-delà de la fiabilité fournie par le codage réseau, ils pouvaient augmenter la vitesse de collecte de données dans les réseaux SG.

5.3.1.1 Combinaison codage réseau et codage fontaine

La combinaison du codage réseau et du codage fontaine est très souvent un problème non trivial. Dans les réseaux où les nœuds sources effectuent un codage fontaine, et où il existe des nœuds relais intermédiaires qui effectuent un codage réseau, plusieurs approches ont été proposées pour combiner le codage réseau et le codage fontaine. Une motivation pour l'association du codage réseau et du codage fontaine pour un canal à effacement est le fait

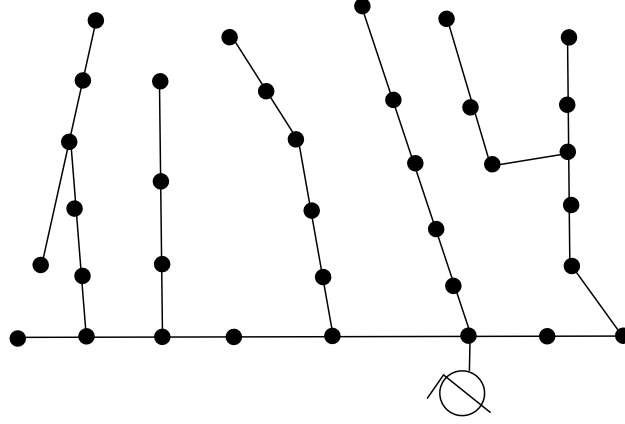


FIGURE 5.5: La distribution de nœuds inspiré du modèle test IEEE 34 nœuds

que celle-ci augmente la taille du code fontaine généré à la sortie du relais. Si *combiné de façon optimale*, le code fontaine résultant est meilleur en terme d'overhead et est appelé code fontaine distribué. Dans cette section, nous considérons le codage conjoint fontaine-réseau dans le contexte des topologies particulières des réseaux CPL-BE pour les SG. Le réseau de distribution IEEE 34-nœuds représenté sur la figure 5.5 [92] révèle que le réseau CPL-BE pour les SG présente une topologie radiale, qui est considéré dans le présent document.

La collecte des informations se fait à partir des extrémités (capteurs) vers les centres de contrôle (concentrateurs), tel que présenté dans la figure 5.6-(a). Lors d'un sondage, chaque nœud esclave envoie ses données au concentrateur soit directement ou par l'intermédiaire d'autres nœuds qui fonctionnent à ce moment là comme répéteurs [93]. Une fonction de répétition est définie comme partie intégrante de la couche MAC des réseaux CPL-BE pour les SG à des fins d'extension de portée. Dans cette section, tous les flux de trames dans le réseau sont encodés avec un code fontaine. Il est clair qu'un volume croissant de données est attendu à l'approche du concentrateur.

En se référant à la figure 5.6-(a), à l'entrée du relais R_1 et en provenance du nœud source S_1 se trouve un flux LT-encodé avec une distribution de sortie de degrés optimale. Dans le reste de ce manuscrit, S_1 est le nœud en amont, qui transmet K_1 trames à une destination via un autre nœud source S_2 qui lui aussi transmet K_2 trames à la destination. Cette section

traite de la reconstruction du code fontaine à la sortie du relais, sans la nécessité de décoder puis ré-encoder le flux de trames total. Nous limitons ainsi l'augmentation du retard. Cette problématique est très proche de celle des codes fontaines distribués abordé dans la littérature.

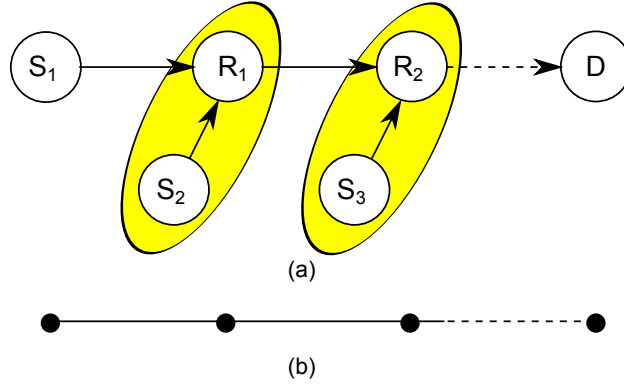


FIGURE 5.6: (a) Modèle de collecte de données CPL-BE pour le SG, (b) Représentation équivalente comme dans le modèle test IEEE-34-nœuds

5.3.1.2 Revue de la littérature sur l'association codage fontaine et codage réseau

Les auteurs dans [94] ont proposé et évalué à l'aide d'approches heuristiques, plusieurs stratégies pour le « xor-age » des trames au relais. Ils ont proposé un algorithme qui effectue les combinaisons « eXclusive OR (XOR) » avec un degré prédéterminé et qui favorise aussi la simple transmission des paquets de faibles degrés. De la même manière, les auteurs dans [95] ont utilisé des buffers pour l'association des codes fontaines et du codage réseau. Bénéficiant du grand nombre de paquets de degré 2, ils proposent des méthodes avec lesquelles les paquets doivent être choisis, combinés puis « raffinés » pour rendre la distribution d'entrée des degrés la plus uniforme possible.

Les auteurs dans [96] ont introduit le codage LT distribué (Distributed LT (DLT)). Dans leur analyse, ils effectuent un codage distribué au niveau des nœuds sources avec une distribution de degré obtenue à partir de la dé-convolution de la Distribution de Soliton Robuste (DSR). Ils ont formé

dans chaque nœud source une distribution de sortie de degrés de telle sorte qu'au niveau du relais après une simple opération « xor », la destination reçoit une distribution totale proche de Soliton Robuste. Les auteurs dans [97] ont étendu le travail de [96] et ont proposé les codes Selective Distributed LT (SDLT) qui prennent en compte un plus grand nombre de nœuds source et offrent une combinaison sélective dans le relais qui consiste à combiner de façon aléatoire les paquets codés reçus de différents nœuds source. Le nombre de paquets « xoré » est fixé sur la base d'une distribution de degré optimisée en suivant l'approche d'analyse en arbre and-or [98]. Plus récemment, les auteurs dans [99] donnent une méthode pour concevoir la distribution de degré au niveau des nœuds source distribués, sur la base de la factorisation polynomiale. Ils introduisent des méthodes pour traiter les nœuds source de façon équitable quand à leur contribution dans les paquets codés.

Les auteurs dans [100] utilisent la relaxe de certaines contraintes de la distribution Soliton Robuste et considèrent uniquement des codes « Soliton Like Rateless Codes (SLRC) » qui préservent les propriétés critiques de la distribution de sortie des codes LT. Il n'y a pas beaucoup de dégradations en termes d'overhead par rapport à la DSR. En fait, chaque nœud source génère une distribution semblable à SR et le relais émet une distribution de degré qui est également semblable à SR.

L'inconvénient majeur de ces méthodes est que la distribution de degrés des nœuds source dépend de la topologie du réseau et doit être connue a priori. Les auteurs dans [101] ont formé toutes les combinaisons « xor » possibles avec les paquets disponibles dans les buffers (un buffer étant dédié à chaque nœud source qui transmet au relais). La distribution de sortie des degrés est mise à jour de sorte que le relais ne tire que parmi les degrés disponibles. Puis, après l'échantillonnage d'un degré de la distribution de sortie des degrés, le relais fait le meilleur choix possible pour la « combinaison » à envoyer en tenant compte de la distribution d'entrée des degrés. Chaque fois qu'un degré n'est pas disponible dans le buffer des paquets « xor-és », ils proposent de mettre à jour la distribution de sortie des degrés avec un algorithme appelé DDU (Degree Distribution Update) afin de promouvoir le degré manquant dans l'échantillonnage du degré de sortie

suivant. De cette façon, ils veillent à ce que la distribution de degrés reçu soit SR en moyenne.

Pour que la distribution d'entrée des degrés soit la plus uniforme possible, l'algorithme proposé dans [101] utilise une métrique notée TN (Transmission Number). La mesure de TN quantifie, pour chaque paquet codé, le nombre moyen de fois que les paquets source qui constituent ce paquet ont été envoyés. Basé sur cette métrique, chaque fois qu'un paquet de degré d doit être envoyé, [101] constitue une liste de tous les paquets disponibles (en tant que combinaison ou pas d'autres packets) et choisit au hasard parmi les paquets ayant la plus petite métrique TN.

La construction de codes fontaines distribués est bien étudiée dans la littérature mais surtout pour les topologies Y ou pour des topologies plus générales (plusieurs sources, plusieurs relais, une destination). Les procédés de combinaison de codes fontaines dans la littérature, s'ils sont utilisés pour les topologies rencontrées dans les réseaux CPL-BE ne donnent pas un résultat optimal.

Ici, grâce à la topologie particulière que l'on voit souvent dans les réseaux CPL-BE pour les SG, la conception du codage distribué est effectuée de telle manière que tous les flux suivent une distribution de sortie de degrés SR et il n'y a pas de mise en mémoire tampon dans les relais.

L'abstraction des deux topologies de base (réseau Y et réseau en ligne) est présentée, ce qui aide à comprendre la logique derrière notre travail. Plus précisément, le fait que le nombre de degrés de liberté associé à la conception de codes fontaines distribués pour les réseaux à topologie en ligne est plus important que celui lié à la conception de codes fontaines distribués pour les réseaux Y. Les degrés de liberté supplémentaires de la topologie en ligne permettent d'optimiser le réseau de manière plus efficace.

La topologie de réseau représentée dans la figure 5.7, dans laquelle deux sources indépendantes, $\mathcal{S}_1$ et $\mathcal{S}_2$ transmettent des informations à une destination via un relais commun $\mathcal{R}$ et où il n'y a pas de lien de communication direct entre les sources et la destination est communément appelée réseau Y.

L'abstraction d'une transmission multi-sauts mono-chemin entre une

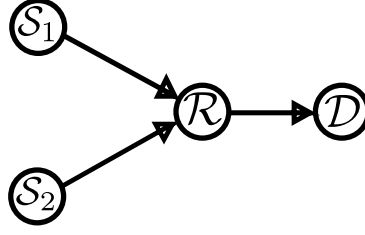


FIGURE 5.7: Topologie d'un réseau en Y

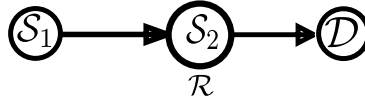


FIGURE 5.8: Topologie d'un réseau linéaire

source et une destination est modélisée par un réseau linéaire comme on peut le voir dans la figure 5.8. Cette topologie peut être considérée comme un réseau Y, où le nœud relais a accès aux paquets de la source locale (S_2).

Nous simplifions ainsi la problématique de fusion de codes fontaines sur des « réseaux Y » en prenant en compte le fait qu'en pratique dans les topologies CPL-BE pour les applications SG, le relais R est co-localisé avec la source S_2 . Le relais a donc accès aux paquets de degré 1 provenant du nœud source S_2 . En outre, notre proposition est adaptative et fonctionne pour les flux de données asymétriques en provenance de deux nœuds source connectés au relais, ce qui se passe dans la pratique et cette asymétrie augmente lorsque nous nous approchons du concentrateur.

5.3.1.3 Stratégie de relayage

Pour un réseau multi-saut composé de deux sources, S_1 appelé nœud en amont, et S_2 (S_2 étant co-localisé avec le relais) avec respectivement K_1 et K_2 paquets d'information, l'objectif est de générer un code LT de taille $K = K_1 + K_2$ en préservant les distributions d'entrée et de sortie des degrés. Pour respecter la distribution d'entrée des degrés au relais, à chaque fois que le relais émet un paquet codé de degré i , la probabilité que ce paquet

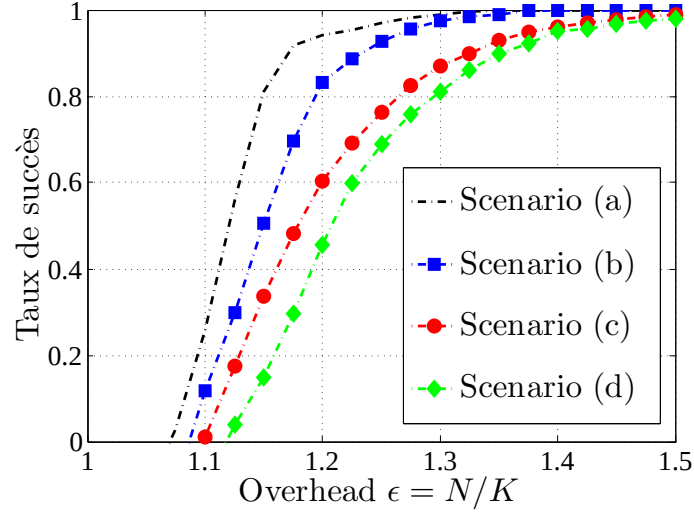
comprenne $(i - j)$ paquets de S_1 est calculée comme suit :

$$\frac{\binom{K_1}{i-j} \binom{K_2}{j}}{\binom{K}{i}}, \quad 1 \leq i \leq K, \quad 0 \leq j \leq i.$$

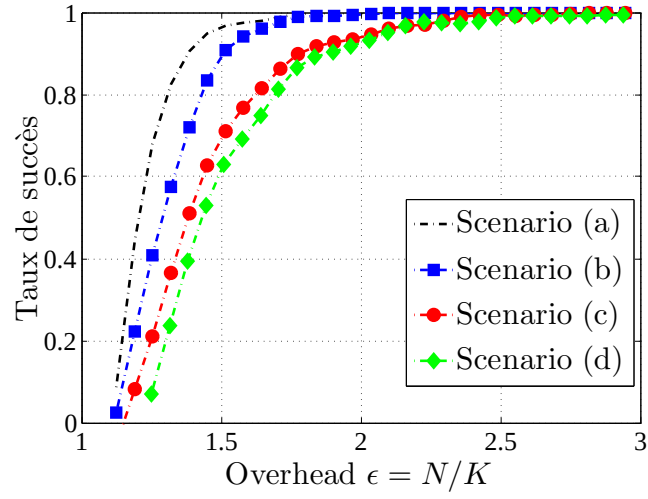
Au relais, nous construisons un code LT de taille $K = K_1 + K_2$ sans avoir à notre disposition tous les $K = K_1 + K_2$ paquets source, ce qui empêche le relais de toujours respecter la contrainte d'un échantillonnage aléatoire et uniforme. La question est alors la suivante : quel sera l'impact d'un échantillonnage non-uniforme des K paquets source. Afin de quantifier son importance, des simulations ont été effectuées. Une transmission multipoint-à-point avec deux sources est considérée avec $K = \{1000, 200\}$, $K_1 = K_2 = \frac{K}{2}$, les performances en termes de taux de succès de décodage sont évaluées pour une redondance donnée, définie par $\epsilon = N/K$, où N est le nombre de paquets reçus à la destination avant que celle-ci ne soit capable de décoder.

Dans un premier temps, tous les paquets de degré 2 ou de degré 3 ou de degré 4 sont sélectionnés soit dans S_1 ou dans S_2 (On appelle ce scénario le scénario (b)). Par conséquent, un paquet de degré 2 ou 3 ou 4 n'est jamais codé à partir des paquets provenant des deux sources. Tous les autres paquets sont générés en utilisant un échantillonnage uniforme à partir de l'ensemble des paquets. Ensuite, les paquets de degrés 5, 6 et 7 sont également ajoutés à ce processus et sont prélevés de la même façon que les paquets de degrés 2, 3 et 4 (scénario (c)). Et puis, tous les paquets (de degré $d \leq \lfloor \frac{K}{2} \rfloor$) sont sélectionnés soit dans S_1 ou dans S_2 (scénario (d)). La figure 5.9 montre les résultats. La courbe de référence en noir correspond au cas idéal où nous avons un code LT de taille K (scénario (a)).

En utilisant un échantillonnage non uniforme pour générer les paquets codés, des transmissions redondantes sont plus susceptibles d'apparaître. Cela entraînera une moins bonne performance du décodage LT comme en témoigne l'écart entre les courbes où l'échantillonnage est non uniforme et la courbe de référence. Cet écart est particulièrement important, lorsque le nombre de degrés choisis de manière non uniforme augmente. À partir de la figure 5.9, nous notons que des améliorations en termes d'overhead peuvent



(a) $K = 1000$
et $K_1 = 500$



(b) $K = 200$
et $K_1 = 100$

FIGURE 5.9: Taux de succès du décodage pour des codes LT en terme d'overhead lorsque les paquets codés ne sont pas choisis de façon aléatoire et uniforme (avec K paquets source, K_1 paquets au relais, $c = 0.05$ et $\delta = 0.5$)

être obtenus en maintenant un échantillonnage uniforme sur l'ensemble des $K_1 + K_2$ paquets au relais.

Afin d'avoir un réseau SG adaptatif, flexible et efficace, les contraintes suivantes sont prises en compte dans le présent manuscrit. Tous les nœuds (relais) dans le réseau SG utilisent le même algorithme indépendamment de la topologie du réseau et de leur emplacement dans le réseau. Tous les flux de paquets dans le réseau sont LT-encodés avec la distribution RS. Le débit d'entrée des paquets au relais est le même que le débit de sortie e.g., le relais envoie un paquet codé à la réception d'un paquet. Les algorithmes proposés ne nécessitent pas la mise en mémoire des paquets reçus et il n'y a pas de décodage LT au relais. Avec ces propriétés, ajouter ou supprimer un nœud n'a pas d'effet sur la gestion du système, ce qui rend la conception du système simple.

L'objectif du relais (ayant à sa disposition K_2 paquets) est d'envoyer des paquets avec une DSR de taille K lors de la réception de paquets avec une DSR de taille K_1 . En outre, le relais doit respecter la sélection uniforme et aléatoire des paquets ce qui résulte en la distribution d'entrée de degrés étant binomiale sur tous les $K = K_1 + K_2$ paquets.

Nous définissons la matrice $\mathbf{P}_{(K+1) \times (K_1+1)}$, avec $K + 1$ lignes et $K_1 + 1$ colonnes, avec l'entrée $p_{i,j}$ étant la probabilité conjointe que le paquet de sortie du relais soit de degré i et comprenne j paquets de S_1 . Dans un cas idéal, la matrice $\mathbf{P}$ est calculée comme suit :

$$\begin{aligned}
 p_{i,j} &= \Pr\{d = i, d_1 = j\} \\
 &= \Pr\{d = i\} \Pr\{d_1 = j \mid d = i\} \\
 &= \begin{cases} \mu_K(i) \frac{\binom{K_1}{j} \binom{K_2}{i-j}}{\binom{K}{i}}, & \text{pour } 0 \leq j \leq i \\ 0, & \text{pour } i + 1 \leq j \leq K_1 \end{cases} \quad (5.6)
 \end{aligned}$$

Bien que la première ligne de $\mathbf{P}$ est toujours tout à zéro, l'indexation commence à partir de zéro pour des fins de notation. Par exemple, les 5 premières lignes et colonnes de la matrice $\mathbf{P}$ pour $c = 0.05$, $\delta = 0,5$ et

$K_1 = K_2 = 50$ sont calculées comme suit :

$$\mathbf{P}_{101 \times 51} = \begin{bmatrix} 0 & 0 & 0 & 0 & 0 & \dots \\ 0.016 & 0.016 & 0 & 0 & 0 & \dots \\ 0.11 & 0.22 & 0.11 & 0 & 0 & \dots \\ 0.018 & 0.058 & 0.058 & 0.018 & 0 & \dots \\ 0.005 & 0.02 & 0.03 & 0.02 & 0.005 & \dots \\ \dots & \dots & \dots & \dots & \dots & \dots \end{bmatrix} \quad (5.7)$$

Le relais doit donc envoyer 22% de paquets de degré 2, formés par « xor-age » d'un paquet de la source en amont et d'un paquet provenant de son propre ensemble de paquets.

Étant donné que d est le degré de sortie, en marginalisant $p(d, d_1)$ par rapport à d_1 , nous obtenons une distribution soliton robuste. Cette probabilité marginale $P_{\text{out}}(d = i)$ est un vecteur de taille $K + 1$ obtenu en additionnant les colonnes de la matrice $\mathbf{P}$:

$$P_{\text{out}}(i) = \sum_{j=0}^{K_1} p_{i,j}, \quad 0 \leq i \leq K.$$

Définissons les probabilités marginales $P_{S_1}(d_1 = j)$ comme étant un vecteur de taille $K_1 + 1$ qui peut être obtenu comme la somme des lignes de la matrice $\mathbf{P}$:

$$P_{S_1}(j) = \sum_{i=0}^K p_{i,j}, \quad 0 \leq j \leq K_1.$$

Étant donné que la distribution reçue à partir de S_1 est une DSR, il n'est pas toujours possible d'avoir à la fois une distribution d'entrée binomiale et une distribution de sortie SR à la sortie du relais. Pour mieux comprendre cela, examinons la matrice donnée dans (5.7) par exemple. Selon cette matrice, le relais doit générer 22% de paquets avec $d_1 = 1$ et $d = 2$. Puisque nous ne voulons pas de décodage au relais, il n'est tout simplement pas possible d'avoir cette quantité de paquets de degré 1 provenant de S_1 en raison de la forme de la DSR au niveau de S_1 .

Pour éviter ce problème, nous mettons à jour $\mathbf{P}$ de sorte que le besoin

total pour les paquets de degré j venant de S_1 ($P_{S_1}(j)$) est inférieur ou égal au pourcentage disponible de paquets de degré j venant de S_1 ($\mu_{K_1}(j)$).

Nous devons donc modifier la matrice $\mathbf{P}$ afin de rendre possible la conservation de la DSR en sortie sur les K paquets, lors de la réception d'une DSR sur K_1 paquets. Comme fonction d'optimisation, nous essayons de minimiser l'écart avec la distribution idéale binomiale donnée dans (5.6). En fait l'optimisation vise à déterminer la distribution conjointe de degrés optimale permettant de minimiser l'overhead. L'overhead dépend de la distribution de degrés des paquets de façon non triviale.

5.3.2 Optimisation de la matrice de probabilité jointe

Dans un premier temps, nous formulons le problème comme étant un problème d'optimisation standard qui peut être résolu par logiciel (nous avons utilisé pour la résolution, le package CVX de Matlab [102]). Ensuite, nous présentons une méthode simple qui donne une solution réalisable pour la matrice $\mathbf{P}$ (voir l'algorithme 2). On note par $\bar{\mathbf{P}}$ la matrice qui satisfait les contraintes et décrit les distributions conjointes possibles d, d_1 des degrés. Les contraintes étant que la somme des colonnes de $\bar{\mathbf{P}}$ (donnant P_{out}) doit correspondre à une DSR et la somme des lignes de $\bar{\mathbf{P}}$ (donnant P_{S_1}) doit toujours être possible de générer à partir de la distribution de sortie des degrés reçue de S_1 .

$$\begin{aligned}
 & \min_{\bar{\mathbf{P}}} && \|\bar{\mathbf{P}} - \mathbf{P}\|^2 \\
 & \text{s.c.} && \\
 & && \bar{\mathbf{P}}^T \cdot \mathbb{1}_{(K+1)} \leq \mu_{K_1} \\
 & && \bar{\mathbf{P}} \cdot \mathbb{1}_{(K_1+1)} = \mu_K \\
 & && \text{triu}(\bar{\mathbf{P}}) = 0 \\
 & && \bar{\mathbf{P}} \geq 0.
 \end{aligned} \tag{5.8}$$

Ces contraintes sont utilisées dans un problème d'optimisation convexe défini dans 5.8, où nous tirons $\bar{\mathbf{P}}^*$ ayant la distance minimale -de Frobenius- par

rapport à $\mathbf{P}$.

Dans l'équation 5.8, $\leq$ et $\geq$ sont évalués élément par élément.

Il est facile de transformer ce problème d'optimisation formulé matriciellement à un problème d'optimisation formulé avec des vecteurs, en réarrangeant les éléments non nuls de $\bar{\mathbf{P}}$ dans un vecteur et en ajustant la formulation des contraintes. Le code MATLAB, utilisant le package CVX pour l'optimisation sous forme de vecteur est donnée ci-dessous.

```
n = size(xo);
cvx_begin
    variable x(n)
    minimize norm(x-xo) %- Obj. function
    subject to          %- Constraints
        A*x'== beq' ;
        B*x'<= bneq';
        x>= 0;
cvx_end
```

où A et B sont deux matrices qui regroupent les contraintes. Afin de montrer que ce problème d'optimisation a toujours une réponse, nous utilisons un algorithme simple, expliqué dans la section 5.3.4.

5.3.3 Génération du degré de sortie

À chaque fois qu'un degré $d_1 = j$ est reçu de S_1 , le relais a deux options afin de générer le degré de sortie $d = i$.

- Soit il consulte la colonne avec l'indice j de la matrice $\bar{\mathbf{P}}^*$ et forme une loi de probabilité en normalisant cette colonne. d est tiré de cette loi de probabilité ainsi formée. Il délivre alors un paquet de degré d qui comprend le paquet reçu et $d - j$ de ses propres paquets échantillonnés de façon aléatoire et uniforme.
- Soit il choisit d'envoyer un paquet codé comportant exclusivement ses propres paquets.

Le relais ne reçoit jamais de paquets de degré $d_1 = 0$. Par conséquent, afin de générer les paquets provenant exclusivement de son propre ensemble

de paquets, ce qui arrive avec une probabilité $P_{\text{excl-S2}} = P_{S1}(0)$, le relais fonctionne comme suit. Lors de la réception d'un paquet à partir de S_1 de degré $d_1 = j$ qui arrive avec la probabilité $\mu_{K_1}(j)$, à $P_{S1}(j)/\mu_{K_1}(j)$ du temps le paquet reçu est mélangé avec $d - d_1$ paquets de S_2 pour former le paquet de sortie comme décrit précédemment. Pour le reste du temps, la colonne zéro de la matrice $\bar{\mathbf{P}}^*$ est utilisée pour produire des paquets provenant exclusivement de S_2 . Une loi de probabilité, décrite dans l'équation 5.9 permet le choix du degré d du paquet qui sera envoyé par le relais :

$$\Pr(D = d) = \begin{cases} \bar{\mathbf{P}}^*(d, j) / \left(\sum_{d=0}^K \bar{\mathbf{P}}^*(d, j) \right) & \text{pour } j \neq 0 \\ \bar{\mathbf{P}}^*(d, 0) / \left(\sum_{d=0}^K \bar{\mathbf{P}}^*(d, 0) \right), & \text{pour } j = 0 \end{cases} \quad (5.9)$$

Le pourcentage des paquets provenant exclusivement de S_2 , sera le même que $P_{S1}(0)$ tandis que la distribution de sortie des degrés reste SR. Le détail de la procédure est donnée dans le pseudo-code de l'algorithme 1.

5.3.4 Algorithme de ré-allocation

Cet algorithme permet de trouver une « bonne » solution $\mathbf{P}_o$ au problème d'optimisation.

À l'étape d'initialisation, l'algorithme commence avec :

$$\begin{aligned} \mathbf{P}_o &= 0, \\ \mu_{K_1}^{\text{residual}} &= \mu_{K_1}. \end{aligned}$$

Ensuite, il scanne la matrice idéale $\mathbf{P}$ ligne par ligne à partir de la deuxième ligne. Pour la i -ème ligne de $\mathbf{P}$, il distribue $\mu_{K_1}^{\text{residual}}(0 : i)$ aux éléments de $\mathbf{P}_o(i, 0 : i)$ proportionnellement aux pourcentages qui sont dans $\mathbf{P}(i, 0 : i)$. Avec la condition que

Algorithm 1: Algorithme de fusion des packets au relais

Entrée: $\bar{\mathbf{P}}^*(d, d_1)$
 Pck_{in} $\triangleright$ Paquet reçu de S_1 .
 μ_{K_1} $\triangleright$ Dist. de sortie des degrés de S_1 .
Sortie: Pck_{out} $\triangleright$ Paquet envoyé à la destination.

- 1 $j = \text{degree}(\text{Pck}_{in})$
- 2 $P_{S1}(j) = \sum_{i=0}^K \bar{\mathbf{P}}^*(i, j)$
- 3 $P_{using}(j) = P_{S1}(j)/\mu_{K_1}(j) \triangleright$ Prob d'utiliser Pck_{in} dans le processus de fusion.
- 4 **si** $\text{rand} \leq P_{using}(j)$ **alors**
 - 5 Former une loi de prob (Eq. 5.9, $j \neq 0$) pour le choix du degré d .
 - 6 Choisir d en échantillonnant la loi de prob.
 - 7 Échantillonner de façon aléatoire $d - j$ packets du relais pour former Pck_r .
 - 8 Envoyer $\text{Pck}_{out} = \text{Pck}_{in} \oplus \text{Pck}_r$.
- 9 **sinon**
 - 10 Choisir un degré d en fonction de la loi de prob décrite dans 5.9 ($j = 0$).
 - 11 Échantillonner de façon aléatoire d packets du relais pour former Pck_{out} .
- 12 **retourner** Pck_{out}

$$\mathbf{P}_o(i, j) \leq \mu_{K_1}^{residual}(j), \quad 0 \leq j \leq i,$$

$$\sum_{j=0}^{\min(i, K_1)} \mathbf{P}_o(i, j) = \mu_K(i).$$

Enfin, $\mu_{K_1}^{residual}(0 : i)$ est mis-à-jour en conséquence, et les pourcentages restants de paquets de degrés $0, 1, \dots, i$ non encore attribués seront utilisés dans les itérations suivantes.

L'algorithme se termine avec :

$$\sum_{j=0}^{\min(i, K_1)} \mathbf{P}_o(i, j) = \mu_K(i),$$

pour tout $i, 1 \leq i \leq K$;

Algorithm 2: Calcul de P_o , la matrice de fusion

Entrée: $\mathbf{P}(i, j)$
 μ_{K_1} ▷ La dist. de sortie des degrés reçue de S_1 .
Sortie: $\mathbf{P}_o(i, j)$

- 1 $\mathbf{P}_o = \text{zeros}(K + 1, K_1 + 1)$
- 2 $\mu_{K_1}^{residual} = \mu_{K_1}$ ▷ Variable temporaire pour μ_{K_1} .
- 3 **pour** $i \leftarrow 1$ **à** K **faire**
- 4 Allouer les éléments de $\mu_{K_1}^{residual}(0 : i)$ aux éléments qui sont dans $\mathbf{P}_o(i, 0 : i)$ de façon proportionnelle aux probabilités qui sont dans $\mathbf{P}(i, 0 : i)$, en maintenant la condition : $\sum_{j=0}^{\min(K_1, i)} \mathbf{P}_o(i, j) = \mu_K(i)$
- 5 Mettre à jour $\mu_{K_1}^{residual}(0 : i)$ ▷ Retirer le pourcentage alloué pour cette i -ème étape.
- 6 **retourner** $\mathbf{P}_o$

5.3.4.1 Exemple de calcul de la matrice $\mathbf{P}_o$

Nous présentons une trace de l'exécution de l'algorithme 2 avec les paramètres suivants $c = 0.05$, $\delta = 0,5$ et $K_1 = K_2 = 50$. Les 4 premières

5.3 Codes fontaines combinés avec le codage réseau pour les CPL-BE

lignes et colonnes de la matrice $\mathbf{P}_o$ sont calculées à chaque itération et exprimés ci-dessus :

$\mathbf{P}(i, j)$	$j = 0$	$j = 1$	$j = 2$	$j = 3$	$\dots$
$i = 0$	0	0	0	0	$\dots$
$i = 1$	0.016	0.016	0	0	$\dots$
$i = 2$	0.1099	0.2243	0.1099	0	$\dots$
$i = 3$	0.0184	0.0575	0.0575	0.0184	$\dots$
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$
$\dots$	$\dots$	$\dots$	$\dots$	$\dots$	$\dots$

$\mathbf{P}_o(i, j)$ est initialisée à la matrice tout zéro et $\mu_{K_1}^{\text{residual}} = \mu_{K_1}$ (obtenu dans l'équation 3.24).

$\mathbf{P}_o(i, j)$	$j = 0$	$j = 1$	$j = 2$	$j = 3$	$\dots$
$i = 0$	0	0	0	0	$\dots$
$i = 1$	0	0	0	0	$\dots$
$i = 2$	0	0	0	0	$\dots$
$i = 3$	0	0	0	0	$\dots$
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$
$\dots$	$\dots$	$\dots$	$\dots$	$\dots$	$\dots$
$\mu_{K_1}^{\text{residual}}$	1	0.0316	0.4442	0.1519	$\dots$

Puisque le relais à un accès direct aux paquets de la source, il est toujours possible de générer les paquets de degré $d_1 = 0$, nous fixons arbitrairement

$$\mu_{K_1}^{\text{residual}}(0) = 1;$$

$\mu_{K_1}^{\text{residual}}(0)$ étant le pourcentage de paquets de degré $d_1 = 0$ provenant de S_1 dans le processus de combinaison.

Dans la première itération, nous allouons $\mu_{K_1}^{\text{residual}}(0 : 1)$ aux éléments dans $\mathbf{P}_o(1, 0 : 1)$. L'allocation est faite de telle manière à maintenir les pourcentages qui se trouvent dans $\mathbf{P}(1, 0 : 1)$. $\mu_{K_1}^{\text{residual}}$ est mis-à-jour sur les indices $(0 : 1)$.

$\mathbf{P}_o(i, j)$	$j = 0$	$j = 1$	$j = 2$	$j = 3$	$\dots$
$i = 0$	0	0	0	0	$\dots$
$i = 1$	0.016	0.016	0	0	$\dots$
$i = 2$	0	0	0	0	$\dots$
$i = 3$	0	0	0	0	$\dots$
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$
$\dots$	$\dots$	$\dots$	$\dots$	$\dots$	$\dots$
$\mu_{K_1}^{\text{residual}}$	0.9842	0.0292	0.4442	0.1519	$\dots$

Dans la seconde itération, nous allouons $\mu_{K_1}^{\text{residual}}(0 : 2)$ aux éléments dans $\mathbf{P}_o(2, 0 : 2)$. Et comme dans l'itération précédente, l'allocation est faite de telle manière à maintenir au mieux les différentes proportions trouvées dans $\mathbf{P}(2, 0 : 2)$. $\mu_{K_1}^{\text{residual}}$ est mis-à-jour sur les indices $(0 : 2)$ par la suite.

$\mathbf{P}_o(i, j)$	$j = 0$	$j = 1$	$j = 2$	$j = 3$	$\dots$
$i = 0$	0	0	0	0	$\dots$
$i = 1$	0.016	0.016	0	0	$\dots$
$i = 2$	0.2075	0.0292	0.2075	0	$\dots$
$i = 3$	0	0	0	0	$\dots$
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$
$\dots$	$\dots$	$\dots$	$\dots$	$\dots$	$\dots$
$\mu_{K_1}^{\text{residual}}$	0.7767	0	0.2343	0.1519	$\dots$

Dans une troisième itération, nous allouons $\mu_{K_1}^{\text{residual}}(0 : 3)$ aux éléments se trouvant dans $\mathbf{P}_o(3, 0 : 3)$ de la même manière que précédemment. $\mu_{K_1}^{\text{residual}}$ est mis-à-jour en conséquence.

$\mathbf{P}_o(i, j)$	$j = 0$	$j = 1$	$j = 2$	$j = 3$	$\dots$
$i = 0$	0	0	0	0	$\dots$
$i = 1$	0.016	0.016	0	0	$\dots$
$i = 2$	0.2075	0.0292	0.2075	0	$\dots$
$i = 3$	0.0296	0	0.0926	0.0296	$\dots$
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$
$\dots$	$\dots$	$\dots$	$\dots$	$\dots$	$\dots$
$\mu_{K_1}^{\text{residual}}$	0.7471	0	0.1417	0.1223	$\dots$

Le même procédé de mis à jour est effectué pour la i -ème itération, jusqu'à $i = K$.

5.3.5 Résultats de simulation

Dans cette section, les résultats de simulation sont présentés pour les algorithmes proposés et sont essentiellement comparés aux scénarios de simulation suivants :

- un *code LT standard* comme base de référence. Nous considérons une transmission point-à-point avec une seule source qui transmet $(K_1 + K_2)$ paquets LT-encodés à la destination.
- un multiplexage temporel de deux codes LT, c'est-à-dire, le relais envoie alternativement les paquets LT-encodés provenant de S_1 et ses propres paquets LT-encodés.
- Les algorithmes DLT proposés dans [96], ceux proposés dans [99], le SLRC proposé dans [103] pour les topologies Y.

Nous évaluons la performance en termes de taux de succès de décodage étant donné une redondance spécifiée définie par $\epsilon = N/K$, où N est le nombre de paquets reçus à la destination avant qu'elle ne soit capable de décoder. La figure 5.10 montre la performance de notre algorithme avec les matrices $\mathbf{P}_o$ et $\bar{\mathbf{P}}^*$ (comme entrées de l'algorithme 1), par rapport à un code LT standard et par rapport au système de transmission à multiplexage temporel. On observe que, d'une part, le processus de fusion proposé surpasse le système de transmission à multiplexage temporel. D'autre part, le processus de fusion proposé, utilisé avec les matrices $\mathbf{P}_o$ et $\bar{\mathbf{P}}^*$ ont des performances similaires, montrant ainsi que la solution donnée dans l'algorithme 2 pour $\mathbf{P}_o$ constitue une excellente solution avec une complexité moindre par rapport à son homologue basée sur la programmation linéaire.

La figure 5.11 montre la performance de notre algorithme comparé aux schémas de fusion de codes LT rencontrés dans la littérature. Cette figure permet de rendre compte des performances supérieures de l'algorithme que nous avons proposé. La conception du processus de fusion est effectuée au plus proche des propriétés des réseaux CPL-BE. La topologie considérée ayant

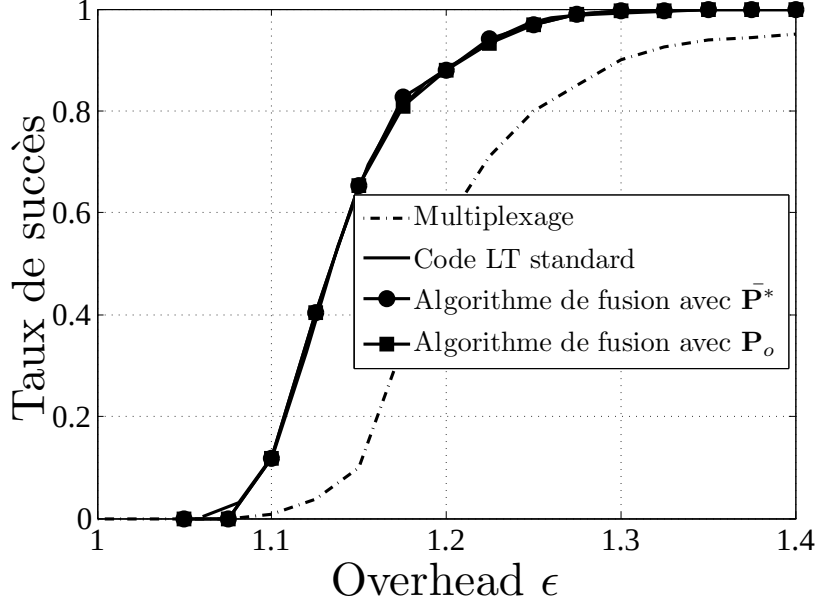


FIGURE 5.10: Taux de succès du processus de décodage pour les matrices de fusion $\mathbf{P}_o$ et $\bar{\mathbf{P}}^*$ en terme d'overhead (avec $K_1 = K_2 = 250$, $c = 0.05$, $\delta = 0.5$)

plus de degrés de liberté par rapport aux réseaux Y traditionnels, notre conception est plus efficace et peut être mieux optimisée.

Les performances en terme de scalabilité de l'algorithme sont présentées dans la figure 5.12, ce sont les performances lorsque le réseau augmente en taille. Dans ce cas, il peut arriver que le relais fusionne deux sources inégales, c'est-à-dire, l'une ayant une taille plus grande que l'autre. En fait, plus on se rapproche du concentrateur, plus la fusion des paquets au niveau du relais apparait inégale en faveur du nœud transmettant au relais. Maintenant, comme nous pouvons le voir dans la figure. 5.12, la source en amont contribue à hauteur de 75% du nombre total de paquets à transmettre à la destination. Le processus de fusion présente de bonnes performances en dépit de l'inégalité des sources.

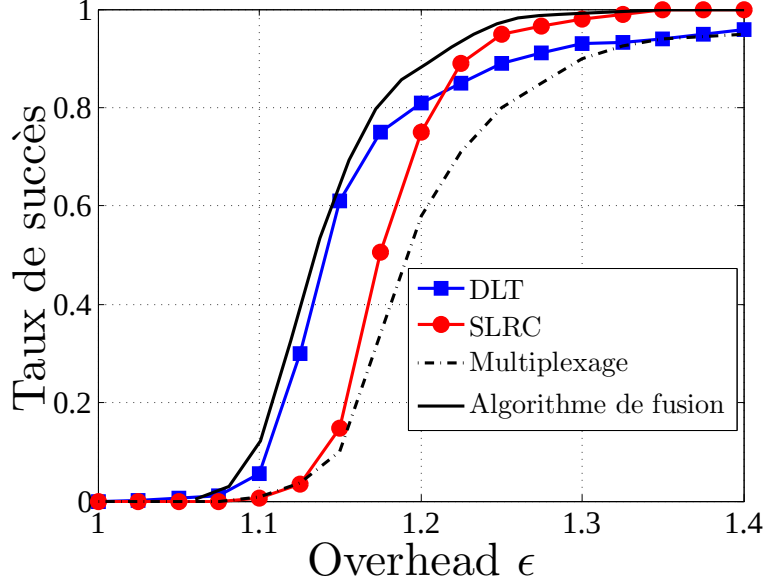


FIGURE 5.11: Taux de succès du processus de fusion, en fonction de l'overhead, comparé à différentes stratégies dans l'état de l'art (avec $K_1 = 250$, $K_2 = 250$, $c = 0.05$, $\delta = 0.5$)

Enfin, nous fournissons des résultats de simulation pour un scénario multi-sauts avec 3 relais afin d'apporter des mesures quantitatives de la performance du système proposé, quand il est mis à l'échelle. Dans les résultats de simulation présentés dans la figure 5.13, relais et nœuds source envoient le même nombre de paquets de sorte que la taille de la génération de code LT double à chaque saut. Le premier nœud source envoie $K_1 = 100$ paquets et après l'application du processus de fusion dans les 3 relais, la destination reçoit $K = 2^3 \times K_1$ paquets. Ces configurations multi-sauts justifient par ailleurs la pertinence de notre travail pour les topologies rencontrées dans les CPL-BE pour les réseaux SG. Comme on peut le voir sur la figure 5.13, notre système est plus efficace que l'alternative d'utiliser plusieurs codes fontaines en parallèle et il n'y a pas de dégradation par

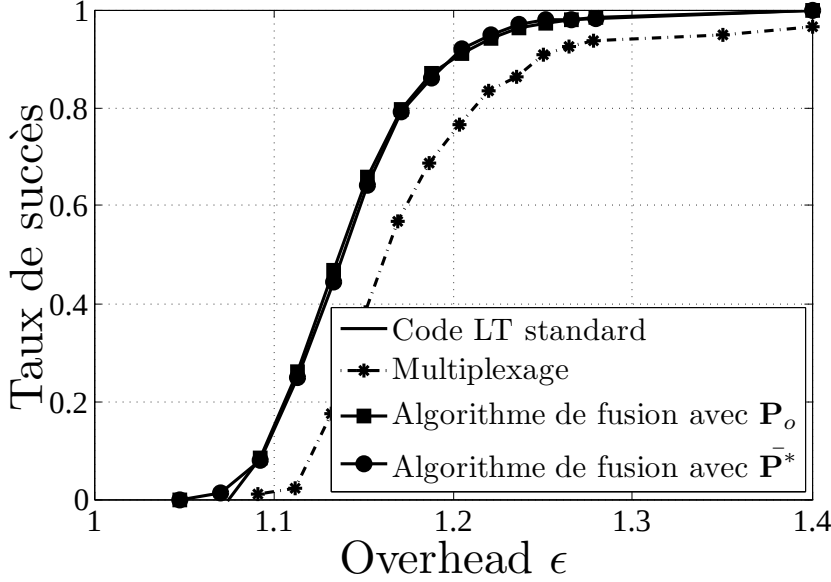


FIGURE 5.12: Taux de succès du processus de fusion proposé en fonction de l'overhead (avec $K_1 = 450$, $K_2 = 150$, $c = 0.05$, $\delta = 0.5$)

rapport à un code LT standard.

Deux points expliquent la bonne performance de notre approche dans un scénario cascadi. D'abord, les performances des codes fontaines (LT) sont très affectées par la distribution de sortie des degrés, et notre approche reconstruit de manière exacte la distribution de sortie idéale des degrés à la sortie du relais. Contrairement à beaucoup de techniques de l'état de l'art, qui se contentent d'approcher la distribution de sortie optimale des degrés requise à la sortie du relais.

Les performances des codes fontaines (LT) sont également affectées par la distribution d'entrée des degrés, mais dans une moindre mesure. En outre, l'approximation que nous avons faite de la distribution d'entrée des degrés ne dévie pas (du cas idéal) suffisamment pour dégrader les performances d'un scénario à sauts multiples.

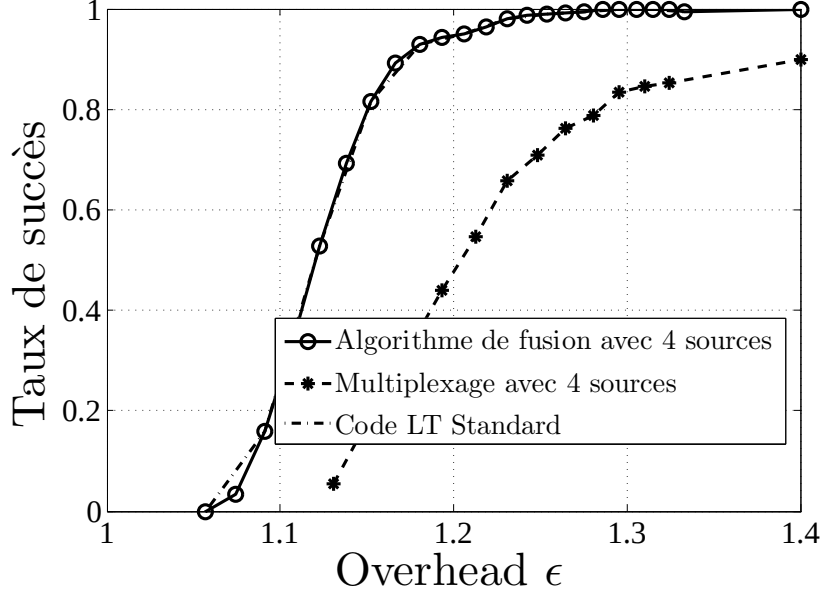


FIGURE 5.13: Taux de succès du processus de fusion proposé en fonction de l'overhead pour un scénario mis en cascade (avec 3 sauts $K_1 = 100$, $c = 0.05$, $\delta = 0.5$)

5.3.5.1 Analyse de la complexité

Le calcul des matrices de fusion est effectué une fois pour toutes et ces matrices sont enregistrées dans des tableaux Look Up Table (LUT). La complexité de calcul des matrices de fusion correspond à la complexité de la résolution d'un problème de programmation linéaire. L'algorithme de réaffectation a une complexité de l'ordre de $\mathcal{O}(K^2)$. La fusion au relais consiste à effectuer une recherche dans la table LUT; la complexité de ces opérations peut être négligée. Nos algorithmes ne sont donc pas plus complexes comparés aux autres approches de codage fontaines distribués.

5.4 Conclusion

L'architecture du réseau CPL-BE et les applications nécessaires pour permettre le SG rendent ce réseau adéquat pour l'application du codage fontaine mais aussi du codage réseau.

Dans la première partie de ce chapitre, nous discutons du relayage de codes fontaines de façon coopérative : ici un seul flux LT doit être relayé à travers un réseau CPL-BE. En nous limitant uniquement à deux relais, nous évaluons de façon analytique les lois de probabilité régissant le succès du décodage au niveau de la destination.

Dans une seconde partie, nous combinons le codage réseau et fontaine appliqué à la topologie des réseaux CPL-BE. Les codes fontaines par leur faible complexité de décodage permettent de bénéficier du codage réseau sans en avoir les inconvénients à savoir, la complexité du décodage. Ainsi, nous considérons le relayage de paquets codés fontaine sur un lien de transmission multi-sauts et nous proposons des algorithmes pour combiner les paquets aux relais, de manière à préserver les propriétés importantes qui permettent le décodage optimal à la destination. Les résultats de simulations confirment la bonne performance des algorithmes proposés pour un réseau réaliste.

Chapitre 6

Conclusion générale et perspectives

Nous avons proposé et implémenté des techniques de codage correcteurs d'erreurs et d'effacements pour fiabiliser la couche PHY et/ou les couches plus hautes des réseaux CPL-BE pour les SG. Ces techniques ont été utilisées seules ou jointes à des techniques de coopération ou de codage réseau pour les communications multi-sauts sur des réseaux CPL-BE à topologie linéaire. Les critères à optimiser étant les complexités, le taux d'erreur bloc des transmissions et le niveau d'overhead. Dans cette dernière partie, nous résumons ci-dessous les principales contributions de cette thèse et nous identifions quelques perspectives.

6.1 Conclusion

Les CPL-BE ont été envisagés comme moyen de communication pour permettre la mise en œuvre du SG dans le réseau d'accès. Aussi, comme les CPL-BE n'ont pas été conçus pour propager des signaux hautes fréquences, les concepteurs de systèmes à base de CPL doivent faire face à de nombreux défis pour les communications par CPL-BE. Un de ces défis consiste en la fiabilisation des communications sur le réseau CPL.

Dans la première partie de ce manuscrit, nous nous sommes intéressés à la

fiabilisation des liens sur les canaux CPL-BE, en présence d'un bruit impulsif dominé par sa composante périodique synchrone, du bruit à bande étroite et du bruit de fond. Nous avons implémenté les codes à métrique rang pour des systèmes CPL-BE, où les transmissions entre émetteurs et récepteurs peuvent être assimilées à des transmissions de matrices avec des erreurs confinées dans les lignes et/ou les colonnes : on parle de motifs d'erreurs criss-cross. Les codes à métrique rang se sont avérés intéressants pour combattre les motifs d'erreurs criss-cross. Des simulations ont été effectuées en considérant des environnements réalistes et des normes et standards déjà existants. Comparés aux normes actuelles, les schémas proposés sont plus robustes face à l'environnement bruyant du canal CPL-BE sans pour autant être plus complexes.

Dans une seconde partie, nous avons implémenté les codes de type fontaine, particulièrement intéressants pour les canaux CPL-BE, car ils s'adaptent à toutes les statistiques d'effacements du canal CPL. Il s'agit ensuite d'effectuer le relayage de ces codes fontaines ; pour cela le relayage coopératif d'un flux unicast de codes LT a été analysé pour une liaison à plusieurs sauts. En considérant le fait que le réseau CPL est un bus, différents relais sont autorisés à relayer de façon coopérative et avec un overhead minimal un flux fontaine sur un réseau multi-sauts et mono-chémin.

Enfin, compte tenu de la topologie du réseau CPL pour les SG, ce réseau peut bénéficier de l'association du codage fontaine et du codage réseau inter-session. Nous avons proposé des algorithmes qui permettent cette association et qui ont de bonnes performances malgré une faible complexité. Nous avons ainsi obtenu des procédés de fusion de codes fontaines adaptés aux réseaux CPL-BE multi-sauts pour des applications SG.

6.2 Perspectives

Je conclus ce travail de recherche par une discussion sur les directions de recherche possibles pour l'amélioration de la fiabilité sur les CPL-BE pour le SG.

- Tout d'abord nous pouvons étendre la discussion de la première section

du chapitre V au cas où il y a plusieurs relais. Et calculer la probabilité que la destination décode en supposant que le décodage se fasse suivant un certain ordre. Cette hypothèse simplificatrice consiste à supposer que les relais plus proches de la source décodent plus rapidement que les relais qui sont plus éloignés. De cette façon nous avons une borne inférieure sur la probabilité de décodage au niveau de la destination.

- Ensuite, nous pouvons élargir l’approche proposée dans la deuxième section du cinquième chapitre pour étudier des politiques de relayage optimales de codes fontaines non-binaires sur des topologies mono-chémin. Nous pourrions également envisager le processus de fusion de codes fontaines quand ils sont décodés de façon soft.
- Enfin, la solution proposée dans la section II du chapitre V peut s’appliquer également à une topologie de bus caractéristique des réseaux CPL-BE. La présence d’un lien de communication direct entre la source et la destination dans un réseau mono-chémin se traduit en ce que la destination reçoit des paquets LT-codés provenant directement de la source. En définitive, la destination reçoit des paquets LT-codés de taille de génération K_1 provenant d’une source et un autre flux de paquets LT-codés de taille de génération $K_1 + K_2$ provenant d’une autre source (qui agit comme relais). Les changements doivent être prévus dans le processus de codage au niveau du relais de sorte que la distribution reçue à la destination ne diffère pas d’une distribution de Soliton Robuste.

Liste des publications acceptées ou soumises

Conférences et journaux

- [A1] Kabore A.W., Meghdadi V., Cances J.-P., “Cooperative Relaying in NarrowBand PLC Networks using Fountain Codes,” in the 18th IEEE International Symposium on Power Line Communications and its Applications (ISPLC), pp. 306-310 avril 2014.
- [A2] —, “Performance of Gabidulin Codes for Narrowband PLC Smart Grid Networks,” in the 19th International Symposium on Power Line Communications and its Applications (ISPLC), pp. 262-267, avril 2015.
- [S1] Kabore A.W., Meghdadi V., Cances J.-P., “Fountain Codes Combined with Network Coding for Tree Topology Multihop Powerline Smart Grid Networks” soumis à Transactions on Networking.
- [S2] —, “LT Codes Combined with Network Coding for Multihop Powerline Smart Grid Networks” article soumis à ICT.
- [S3] —, “On the Optimization of the Physical Layer Frame Lengths of Narrowband Power-Line Communications” soumis à Eusipco.

Rapports Techniques :

- [T1] Kabore A.W., Meghdadi V., Cances J.-P., “LT Codes Combined with Network Coding for Multihop Powerline Smart Grid Networks” Arxiv à <http://arxiv.org/abs/1509.06019> 2015.
- [T2] —, “ “Performance des codes à métrique de rang pour les réseaux CPL à bande étroite” disponible dans Researchgate.

Bibliographie

- [1] IEEE, “IEEE Standard for Low-Frequency (less than 500 kHz) Narrowband Power Line Communications for Smart Grid Applications,” *IEEE Std 1901.2-2013*, pp. 1–269, Dec 2013.
- [2] D. Silva and F. Kschischang, “Fast Encoding and Decoding of Gabidulin Codes,” in *IEEE International Symposium on Information Theory (ISIT)*, June 2009, pp. 2858–2862.
- [3] A. Ali, *Smart Grids : Opportunities, Developments, and Trends*, ser. Green Energy and Technology. Springer, 2013.
- [4] S. Galli, A. Scaglione, and Z. Wang, “For the Grid and Through the Grid : The Role of Power Line Communications in the Smart Grid,” *CoRR*, 2010. [Online]. Available : <http://arxiv.org/abs/1010.1973>
- [5] A. Amarsingh, H. Latchman, and D. Yang, “Narrowband Power Line Communications : Enabling the Smart Grid,” *IEEE Potentials Magazine*, vol. 33, no. 1, pp. 16–21, Jan 2014.
- [6] S. Galli and T. Lys, “Next generation Narrowband (under 500 kHz) Power Line Communications (PLC) standards,” *China Communications*, vol. 12, no. 3, pp. 1–8, March 2015.
- [7] A.J.H.Vinck, “Coding for a Terrible Channel,” in *Telecommunications and Information Science and Technologies COST*, july 2005.
- [8] J. Lin, “Robust Transceivers to Combat Impulsive Noise in Powerline Communications,” Ph.D. dissertation, The university of Texas at Austin, 2014.
- [9] M. Nassar, J. Lin, Y. Mortazavi, A. Dabak, I. H. Kim, and B. Evans, “Local utility Power Line Communications in the 3-500 khz Band :

- Channel Impairments, Noise, and Standards,” *IEEE Signal Processing Magazine*, vol. 29, no. 5, pp. 116–127, Sept 2012.
- [10] D. MacKay, “Fountain Codes,” *IEEE Proceedings-Communications*, vol. 152, no. 6, pp. 1062–1068, Dec 2005.
- [11] L. Lampe and A. H. Vinck, “On Cooperative Coding for Narrow Band PLC Networks,” *AEU-International Journal of Electronics and Communications*, vol. 65, no. 8, pp. 681–687, 2011.
- [12] R. Prior, D. Lucani, Y. Phulpin, M. Nistor, and J. Barros, “Network Coding Protocols for Smart Grid Communications,” *IEEE Transactions on Smart Grid*, vol. 5, no. 3, pp. 1523–1531, May 2014.
- [13] A. W. Kabore, V. Meghdadi, J. P. Cances, P. Gaborit, and O. Ruatta, “Performance of Gabidulin codes for Narrowband PLC Smart Grid Networks,” in *2015 International Symposium on Power Line Communications and its Applications (ISPLC)*, March 2015, pp. 262–267.
- [14] A. W. Kabore, V. Meghdadi, and J. P. Cances, “Cooperative Relaying in Narrow-Band PLC Networks using Fountain Codes,” in *Power Line Communications and its Applications (ISPLC)*, March 2014, pp. 306–310.
- [15] A. Moscatelli, “From Smart Metering to Smart Grids : PLC Technology Evolutions.” Avril 2011, [http ://www.ieee-isplc.org/2011/Moscatelli](http://www.ieee-isplc.org/2011/Moscatelli).
- [16] A. Rossello Busquet, L. Dittmann, and J. Soler, “Intelligent Control of Home Appliances via Network,” Ph.D. dissertation, Technical University of Denmark, 2013.
- [17] G. Guerassimoff and N. Maïzi, *Au-delà du concept, comment rendre les réseaux plus intelligents*. Presses de l’école des mines, 2013.
- [18] G. W. Arnold, “Laying the Foundation for the Electric Grid’s Next 100 Years,” Avril 2011.
- [19] V. Giordano, F. Gangale, G. Fulli, M. S. Jiménez, I. Onyeji, A. Colta, I. Papaioannou, A. Mengolini, C. Alecu, T. Ojala *et al.*, “Smart Grid projects in Europe,” *JRC Reference Reports*, 2011.

BIBLIOGRAPHIE

- [20] P. A. Brown, "Power line communications - past present and future," in *International Symposium on Power Line Communications and its Applications*, ISPLC, March 1999, pp. 1–8.
- [21] P. J. VAN RENSBURG, "Effective Coupling for Power-Line Communications," Ph.D. dissertation, University of Johannesburg, 2008.
- [22] H. C. Ferreira, H. M. Grové, O. Hooijen, and A. Han Vinck, "Power Line Communication," *Wiley Encyclopedia of Electrical and Electronics Engineering*, 2001.
- [23] H. Ferreira, L. Lampe, J. Newbury, and T. Swart, *Power Line Communications : Theory and Applications for Narrowband and Broadband Communications over Power Lines*. Wiley, 2011.
- [24] M. Tlich, A. Zeddami, F. Moulin, and F. Gauthier, "Indoor Power-Line Communications Channel Characterization up to 100 Mhz ; part ii : Time-Frequency Analysis," *IEEE Transactions on Power Delivery*, vol. 23, no. 3, pp. 1402–1409, July 2008.
- [25] T. Banwell and S. Galli, "On the Symmetry of the Power Line Channel," in *International Symposium on Power Line Communications and its Applications*, 2001, pp. 325–330.
- [26] W. Liu, *Emulation of Narrowband Powerline Data Transmission Channels and Evaluation of PLC Systems*, ser. Forschungsberichte aus der Industriellen Informationstechnik / Institut für Industrielle Informationstechnik (IIIT), Karlsruher Institut für Technologie. KIT Scientific Publishing, 2013.
- [27] T. O. meter Consortium, "Report on Final Test Results and Recommendations," 2011.
- [28] H. Philipps, "Modelling of Powerline Communication Channels," in *IEEE Int. Symp. on Power Line Communications and Its Applications*, 1999, pp. 14–21.
- [29] M. Zimmermann and K. Dostert, "A Multipath Model for the Powerline Channel," *IEEE Transactions on Communications*, vol. 50, no. 4, pp. 553–559, 2002.

- [30] O. G. Hooijen, “On the Relation between Network-Topology and Power Line Signal Attenuation,” in *International Symposium on Power Line Communications*, 1998, pp. 45–56.
- [31] S. Galli and T. Banwell, “A Novel Approach to the Modeling of the Indoor Power Line Channel—Part ii : Transfer Function and Its Properties,” *IEEE Transactions on Power Delivery*, vol. 20, no. 3, p. 1869, 2005.
- [32] A. M. Tonello, “Wideband Impulse Modulation and Receiver Algorithms for Multiuser Power Line Communications,” *EURASIP Journal on advances in signal processing*, vol. 2007, no. 17, 2007.
- [33] M. Tlich, A. Zeddami, F. Moulin, and F. Gauthier, “Indoor Power-Line Communications Channel Characterization up to 100 Mhz—part i : one-parameter deterministic model,” *IEEE Transactions on Power Delivery*, vol. 23, no. 3, pp. 1392–1401, 2008.
- [34] M. Katayama, T. Yamazato, and H. Okada, “A Mathematical Model of Noise in Narrowband Power line Communication Systems,” *IEEE Journal on Selected Areas in Communications*, vol. 24, no. 7, pp. 1267–1276, July 2006.
- [35] M. Nassar, A. Dabak, I. H. Kim, T. Pande, and B. Evans, “Cyclostationary Noise Modeling in Narrowband Powerline Communication for Smart Grid Applications,” in *International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, March 2012, pp. 3089–3092.
- [36] W. Liu, “Emulation of Narrowband Powerline Data Transmission Channels and Evaluation of PLC Systems,” Ph.D. dissertation, Karlsruher Institut für Technologie (KIT), 2013.
- [37] D. Middleton, “Statistical-Physical Models of Electromagnetic Interference,” *IEEE Transactions on Electromagnetic Compatibility*, vol. EMC-19, no. 3, pp. 106–127, Aug 1977.
- [38] M. Zimmermann and K. Dostert, “Analysis and Modeling of Impulsive Noise in Broad-band Powerline Communications,” *IEEE Transactions on Electromagnetic Compatibility*, vol. 44, no. 1, pp. 249–258, Feb 2002.

BIBLIOGRAPHIE

- [39] M. Nassar, K. Gulati, Y. Mortazavi, and B. Evans, “Statistical Modeling of Asynchronous Impulsive Noise in Powerline Communication Networks,” in *Global Telecommunications Conference*, Dec 2011, pp. 1–6.
- [40] J. Bausch, T. Kistner, M. Babic, and K. Dostert, “Characteristics of Indoor Power line Channels in the Frequency Range 50-500 khz,” in *IEEE International Symposium on Power Line Communications and its Applications*. IEEE, 2006, pp. 86–91.
- [41] V. Oksman and J. Zhang, “G.HNEM : the new ITU-T Standard on NarrowBand PLC Technology,” *IEEE Communications Magazine*, vol. 49, no. 12, pp. 36–44, December 2011.
- [42] C. E. Shannon, “A Mathematical Theory of Communication,” *ACM Mobile Computing and Communications Review*, vol. 5, no. 1, pp. 3–55, 2001.
- [43] N. Sendrier, “Introduction à la théorie de l’information,” Cours théorie de l’information, août 2007.
- [44] D. J. MacKay, *Information theory, inference and learning algorithms*. Cambridge university press, 2003.
- [45] N. Shlezinger and R. Dabora, “On the Capacity of Narrowband PLC channels,” *IEEE Transactions on Communications*, vol. 63, no. 4, pp. 1191–1201, 2015.
- [46] R. Lidl and H. Niederreiter, *Introduction to finite fields and their applications*. Cambridge university press, 1994.
- [47] E. Berlekamp, *Algebraic Coding Theory*, ser. McGraw-Hill series in systems science. Aegean Park Press, 1984.
- [48] T. K. Moon, *Error Correction Coding : Mathematical Methods and Algorithms*. Wiley-Interscience, 2005.
- [49] I. S. Reed and G. Solomon, “Polynomial Codes over certain Finite Fields,” *Journal of the society for industrial and applied mathematics*, vol. 8, no. 2, pp. 300–304, 1960.

- [50] W. W. Peterson, "Encoding and Error-Correction Procedures for the Bose-Chaudhuri Codes," *IRE Transactions on Information Theory*, vol. 6, no. 4, pp. 459–470, 1960.
- [51] J. L. Massey, "Shift-Register Synthesis and BCH Decoding," *IEEE Transactions on Information Theory*, vol. 15, no. 1, pp. 122–127, 1969.
- [52] L. Welch and E. Berlekamp, "Error Correction for Algebraic Block Codes," Dec. 30 1986, uS Patent 4,633,470.
- [53] P. Elias, "Coding for Noisy Channels," in *IRE International Convention Record (part 4)*, 1965, pp. 37–46.
- [54] B. Claude and G. Alain, "Turbo Codes," *Encyclopedia of Telecommunications*, 2003.
- [55] A. Viterbi, "Error Bounds for Convolutional Codes and an Asymptotically Optimum Decoding Algorithm," *IEEE Transactions on Information Theory*, vol. 13, no. 2, pp. 260–269, April 1967.
- [56] L. Bahl, J. Cocke, F. Jelinek, and J. Raviv, "Optimal Decoding of Linear Codes for Minimizing Symbol Error Rate (corresp.)," *IEEE Transactions on Information Theory*, vol. 20, no. 2, pp. 284–287, Mar 1974.
- [57] G. D. Forney Jr, "Concatenated Codes." Ph.D. dissertation, Massachusetts Institute of Technology, 1965.
- [58] D. J. MacKay, "Good Error-Correcting Codes based on Very Sparse Matrices," *IEEE Transactions on Information Theory*, vol. 45, no. 2, pp. 399–431, 1999.
- [59] M. Sipser and D. A. Spielman, "Expander Codes," *IEEE Transactions on Information Theory*, vol. 42, no. 6, 1996.
- [60] P. Delsarte, "Bilinear Forms over a Finite Field, with Applications to Coding Theory," *Journal of Combinatorial Theory, Series A*, vol. 25, no. 3, pp. 226–241, 1978.
- [61] E. M. Gabidulin, "Theory of Codes with Maximum Rank Distance," *Problems on Information Transmission*, vol. 21, pp. 1–12, Jan 1985.

BIBLIOGRAPHIE

- [62] R. M. Roth, “Maximum-Rank Array Codes and their Application to Crisscross Error Correction,” *IEEE Transactions on Information Theory*, vol. 37, no. 2, pp. 328–336, 1991.
- [63] P. Loidreau, “A Welch–Berlekamp Like Algorithm for Decoding Gabidulin Codes,” in *Coding and cryptography*. Springer, 2006, pp. 36–45.
- [64] G. Richter and S. Plass, “Fast Decoding of Rank-Codes with Rank Errors and Column Erasures,” in *Proceedings of the International Symposium on Information Theory ISIT*. IEEE, 2004, pp. 398–398.
- [65] E. M. Gabidulin, “Theory of Codes with Maximum Rank Distance,” *Problemy Peredachi Informatsii*, vol. 21, no. 1, pp. 3–16, 1985.
- [66] G. Bumiller, “Powerline Channel Adopted Layer-design and Link-layer for Reliable Data Transmission,” in *International Symposium on Power Line Communications and its Applications (ISPLC)*, March 2015, pp. 256–261.
- [67] S. Liu, F. Yang, and S. J., “An Optimal Interleaving Scheme with Maximum Time-Frequency Diversity for PLC Systems,” *IEEE Transactions on Power Delivery*, no. 99, 2014.
- [68] D. J. MacKay, “Fountain Codes,” *IEE Proceedings Communications*, vol. 152, no. 6, pp. 1062–1068, 2005.
- [69] J. Byers, M. Luby, and M. Mitzenmacher, “A Digital Fountain Approach to Asynchronous Reliable Multicast,” *IEEE Journal on Selected Areas in Communications*, vol. 20, no. 8, pp. 1528–1540, Oct 2002.
- [70] M. Luby, “LT Codes,” in *The 43rd Annual IEEE Symposium on Foundations of Computer Science*, 2002, pp. 271–280.
- [71] A. Shokrollahi, “Raptor Codes,” *IEEE Transactions on Information Theory*, vol. 52, no. 6, pp. 2551–2567, June 2006.
- [72] A. Shokrollahi and M. Luby, *Raptor Codes*, ser. Foundations and trends in communications and information theory. Now Publishers, 2011.

- [73] R. Ahlswede, N. Cai, S.-Y. R. Li, and R. W. Yeung, "Network Information Flow," *IEEE Transactions on Information Theory*, vol. 46, no. 4, pp. 1204–1216, 2000.
- [74] T. Ho, M. Médard, R. Koetter, D. R. Karger, M. Effros, J. Shi, and B. Leong, "A Random Linear Network Coding Approach to Multicast," *IEEE Transactions on Information Theory*, vol. 52, no. 10, pp. 4413–4430, 2006.
- [75] S. C. Liew, S. Zhang, and L. Lu, "Physical-layer Network Coding : Tutorial, Survey, and beyond," *Physical Communication*, vol. 6, pp. 4–42, 2013.
- [76] R. Koetter and F. R. Kschischang, "Coding for Errors and Erasures in Random Network Coding," *IEEE Transactions on Information Theory*, vol. 54, no. 8, pp. 3579–3591, 2008.
- [77] W. Chu, C. Colbourn, and P. Dukes, "Constructions for Permutation Codes in Powerline Communications," *Designs, Codes and Cryptography*, vol. 32, no. 1-3, pp. 51–64, 2004.
- [78] D. Galda and H. Rohling, "Narrow Band Interference Reduction in OFDM-based Power Line Communication Systems."
- [79] T. Shongwe, V. Papilaya, and A. Vinck, "Narrow-band Interference model for OFDM systems for Powerline Communications," in *the 17th IEEE International Symposium on Power Line Communications and Its Applications (ISPLC)*, March 2013, pp. 268–272.
- [80] J. W. Byers, M. Luby, M. Mitzenmacher, and A. Rege, "A digital fountain approach to reliable distribution of bulk data," in *Proceedings of the ACM SIGCOMM*, ser. SIGCOMM, 1998, pp. 56–67.
- [81] G. T. . V6.1.0, "Technical specification group services and system aspects; multimedia broadcast/multicast service; protocols and codecs," June 2005.
- [82] P. Amirshahi, S. Navidpour, and M. Kavehrad, "Fountain Codes for Impulsive Noise Correction in Low-voltage Indoor Power-Line Broadband Communications," in *Consumer Communications and Networking Conference*, vol. 1. IEEE, 2006, pp. 473–477.

BIBLIOGRAPHIE

- [83] A. V. V. Balakirsky, "Potential Performance of PLC Systems composed of several Communication Links," *International symposium on Power Line Communications and Its applications (ISPLC)*, vol. 65, pp. 12–16, Apr 2005.
- [84] L. Lampe and A. Vinck, "Cooperative Multihop Power Line Communications," *IEEE International Symposium on Power Line Communications and Its Applications*, 2012.
- [85] K. Y. E. Kurniawan, S. Sun and K. F. E. Chong, "Network Coded Transmission of Fountain Codes over Cooperative Relay Networks," *information theory*, vol. 65, 2010.
- [86] L. Lampe and A. Vinck, "On Cooperative Coding for Narrow Band PLC Networks," *International Journal of Electronics and Communications*, vol. 65, pp. 681–687, jan 2010.
- [87] J. E. Gentle, *Random Number Generation and Monte Carlo Methods (Statistics and Computing)*. springer, 2003.
- [88] Y. Phulpin, J. Barros, and D. Lucani, "Network Coding in Smart Grids," in *IEEE International Conference on Smart Grid Communications (SmartGridComm)*, 2011, pp. 49–54.
- [89] M. Karthick and K. Sivalingam, "Network Coding based Reliable and Efficient Data Transfer for Smart Grid Monitoring," in *IEEE International Conference on Advanced Networks and Telecommunications Systems (ANTS)*, Dec 2013, pp. 1–6.
- [90] J. Bilbao, A. Calvo, I. Armendariz, P. Crespo, and M. Medard, "Reliable Communications with Network Coding in Narrowband Powerline Channel," in *18th IEEE International Symposium on Power Line Communications and its Applications (ISPLC)*, March 2014, pp. 316–321.
- [91] R. Prior, D. Lucani, Y. Phulpin, M. Nistor, and J. Barros, "Network Coding Protocols for Smart Grid Communications," *IEEE Transactions on Smart Grid*, vol. 5, no. 3, pp. 1523–1531, May 2014.
- [92] D. system subcommittee, "IEEE 123 Node Test Feeder." [Online]. Available : <http://ewh.ieee.org/soc/pes/dsacom/testfeeders/>

- [93] J. Selga, A. Zaballos, J. Abella, and G. Corral, “Model for Polling in Noisy Multihop Systems with Application to PLC and AMR,” in *IEEE Symposium on Computers and Communications ISCC*, July 2008, pp. 664–669.
- [94] A. Apavatjrut, C. Goursaud, K. Jaffrès-Runser, C. Comaniciu, and J.-M. Gorce, “Toward Increasing Packet Diversity for Relaying LT Fountain Codes in Wireless Sensor Networks,” *IEEE Communications Letters*, vol. 15, no. 1, pp. 52–54, January 2011.
- [95] M.-L. Champel, K. Huguenin, A.-M. Kermarrec, and N. Le Scouarnec, “LT Network Codes,” in *30th International Conference on Distributed Computing Systems (ICDCS)*, June 2010, pp. 536–546.
- [96] S. Puducheri, J. Kliever, and T. Fuja, “The Design and Performance of Distributed LT Codes,” *IEEE Transactions on Information Theory*, vol. 53, no. 10, pp. 3740–3754, Oct 2007.
- [97] D. Sejdinovic, R. Piechocki, and A. Doufexi, “And-Or tree analysis of Distributed LT codes,” in *IEEE Information Theory Workshop on Networking and Information Theory*, June 2009, pp. 261–265.
- [98] M. Luby, M. Mitzenmacher, and M. A. Shokrollahi, “Analysis of Random Processes via And-Or Tree Evaluation,” in *Proceedings of the 9th Annual ACM-SIAM Symposium on Discrete Algorithms, SODA 1998*, no. ALGO-CONF-1998-001, 1998.
- [99] R. Rafie Borujeny and M. Ardakani, “Fountain Code Design for the Y-Network,” *IEEE Communications Letters*, vol. 19, no. 5, pp. 703–706, May 2015.
- [100] A. Liao, S. Yousefi, and I.-M. Kim, “Soliton-Like Network Coding for a Single Relay,” in *12th Canadian Workshop on Information Theory (CWIT)*, May 2011, pp. 86–89.
- [101] S. Jafarizadeh, “Distributed Coding and Algorithm Optimization for Large-Scale Networked Systems,” Ph.D. dissertation, The University of SYDNEY, 2014.
- [102] M. Grant and S. Boyd, “CVX : Matlab software for disciplined convex programming, version 2.1,” <http://cvxr.com/cvx>, Mar. 2014.

BIBLIOGRAPHIE

- [103] A. Liao, S. Yousefi, and I.-M. Kim, “Binary Soliton-Like Rateless Coding for the Y-network,” *IEEE Transactions on Communications*, vol. 59, no. 12, pp. 3217–3222, 2011.

Étude des codes correcteurs d'erreurs pour les couches PHY/ MAC des réseaux smart grid

Ce manuscrit traite de la fiabilisation des CPL-BE dans le contexte SG avec l'application des techniques de codage correcteur d'erreurs et d'effacements. Après une introduction sur le concept de smart grid, le canal CPL-BE est caractérisé précisément et les modèles qui le décrivent sont présentés. Les performances des codes à métrique rang, simples ou concaténés avec des codes convolutifs, particulièrement intéressants pour combattre le bruit criss-cross sur les réseaux CPL-BE sont simulées et comparées aux performances des codes Reed-Solomon déjà présents dans plusieurs standards. Les codes fontaines sont utilisés et les performances de schémas coopératifs basés sur ces codes fontaines sur des réseaux CPL-BE linéaires multi-sauts sont étudiées. Enfin des algorithmes permettant de combiner le codage réseau et le codage fontaine pour la topologie particulière des réseaux CPL pour les SG sont proposés et évalués.

Mots clés : CPL-BE, smart grid, codage de canal, codes fontaines, codes à métrique rang, codes Gabidulin, codage réseau.

Study of the error correction codes for the PHY/MAC layers of the smart grid networks

This PhD dissertation deals with the mitigation of the impact of the Narrowband PowerLine communication (NB-PLC) channel impairments e.g., periodic impulsive noise and narrowband noise, by applying the error/erasure correction coding techniques. After an introduction to the concept of smart grid, the NB-PLC channels are characterized precisely and models that describe these channels are presented. The performance of rank metric codes, simple or concatenated with convolutional codes, that are particularly interesting to combat criss-cross errors on the NB-PLC networks are simulated and compared with Reed-Solomon codes (already present in several NB-PLC standards). Fountain codes are used for the NB-PLC networks and the performance of cooperative schemes based on these fountain codes on linear multi-hop networks are studied. Finally, algorithms that combine the network coding and fountain codes for the particular topology of PLC networks for the smart grid are proposed and evaluated.

Keywords : Smart grid, NB-PLC, fountain codes, rank metric codes, Gabidulin codes, network coding, channel coding.