



Nouvelle technique de communication pour les réseaux longue portée basseconsommation : optimisation et comparaison

Yoann Roth, Jean-Baptiste Doré, Laurent Ros, Vincent Berg

► To cite this version:

Yoann Roth, Jean-Baptiste Doré, Laurent Ros, Vincent Berg. Nouvelle technique de communication pour les réseaux longue portée basseconsommation : optimisation et comparaison. Journées scientifiques 2016 de l'URSI-France Energie et radiosciences, Union Radio-Scientifique Internationale, Mar 2016, Cesson-Sévigné, France. hal-01418439

HAL Id: hal-01418439

<https://hal.science/hal-01418439>

Submitted on 16 Dec 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.

Nouvelle technique de communication pour les réseaux longue portée basse consommation : optimisation et comparaison

Yoann ROTH^{1, 2}, Jean-Baptiste DORÉ¹, Laurent ROS², Vincent BERG¹

¹ CEA, LETI, MINATEC Campus, F-38054 Grenoble, France, {yoann.roth, jean-baptiste.dore, vincent.berg}@cea.fr,

² Univ. Grenoble Alpes, GIPSA-Lab, F-38000 Grenoble, France, laurent.ros@gipsa-lab.grenoble-inp.fr

Mots-clés : Turbo, Modulations Orthogonales, EXIT Chart, IdO

Keywords : Turbo, Orthogonal Modulations, EXIT Chart, IoT

Résumé — L'internet des objets a pour vocation de connecter des milliards de terminaux entre eux. Ces derniers doivent être bon marché, économes en énergie, et capables de communiquer même à très longue portée. Alors que les technologies actuelles de communications machine-to-machine ont tendance à utiliser les techniques de répétition ou d'étalement afin de répondre aux contraintes notamment en termes de sensibilité, ce travail propose une utilisation plus sophistiquée de la répétition dans une technique dénommée Turbo-FSK. Cette dernière implique l'utilisation d'une modulation orthogonale de fréquence (Frequency-Shift-Keying (FSK)) et d'un turbo décodage spécifique aux formes d'onde FSK au niveau du récepteur. On montre alors qu'une communication robuste est possible même avec un émetteur très simple, la complexité étant déportée au niveau du récepteur. Les résultats de simulations sont comparés à des techniques de références, dont celles utilisant la répétition, démontrant qu'un gain significatif est obtenu même avec des petites tailles de paquets.

Abstract — The Internet of Things aims to connect several billions of devices. Terminals are expected to be low cost, low power, and to be able to achieve successful communication at long range. While current Machine-to-Machine technologies tend to use spreading factors to meet the requirements, we propose a more sophisticated use of redundancy in a scheme called Turbo-FSK. This scheme involves Frequency-Shift-Keying (FSK) modulation at the transmission, and a turbo-decoding dedicated to the FSK waveforms at the receiver. Highly robust communication is achieved with a mere transmitter, as complexity is deported on the receiver side. Results are compared to common modulations using spreading factors, showing a significant gain in performance is achieved even with small packet sizes.

Introduction

Les communications dites Machine-to-Machine (M2M) se développent de façon exponentielle, et on attend plusieurs milliards d'objets connectés à l'Internet des Objets (IdO) dans les années à venir [1]. Dans ce contexte, la mise en place de réseaux longue portée basse consommation est une solution envisagée afin de respecter les contraintes au niveau du terminal : faible consommation, faible coût au niveau de l'infrastructure, dépenses d'investissement de capital (CAPEX) limitées. Le choix de la couche physique est alors critique et cette dernière doit être sélectionnée de telle sorte que la communication soit robuste et fonctionnelle, même à de très faibles niveaux de sensibilités. Par ailleurs, les cellules contiendront un très grand nombre de nœuds émettant chacun de façon sporadique ; on s'attend donc à un faible débit par terminal. C'est le scénario déjà envisagé par les premières solutions industrielles pour l'IdO [2, 3]. Les modulations orthogonales d'ordre M sont un choix naturel pour les communications bas débit et contraintes en énergie. Efficace énergétiquement, elles atteignent la limite théorique de la capacité, énoncée par Claude Shannon, pour une taille d'alphabet infini [4], et donc des durées de forme d'onde infinies. Ceci n'étant pas réalisable, une alternative pour améliorer l'efficacité énergétique de la transmission est le codage canal, et particulièrement l'application du principe turbo [5]. L'idée d'utiliser des codes orthogonaux dans un schéma de concaténation parallèle avec turbo décodage a été proposée en 2001 dans [6], en utilisant la famille de codes binaires de Hadamard. Le schéma présenté montre d'excellentes performances pour de grandes tailles de bloc et d'alphabet. Par ailleurs, la modulation de fréquence (Frequency-Shift-Keying, FSK) est un choix intéressant, pour notre contexte, de modulation orthogonale du fait de son enveloppe constante, de la simplicité du modulateur et de sa robustesse aux canaux multi-trajets sélectifs en fréquence. Le schéma proposé par [6] a d'ailleurs été adapté aux formes d'onde FSK [7], la Turbo-FSK, démontrant la possibilité d'un tel système et son intérêt. Cependant, l'analyse dans [7] manque de comparaisons réalistes au vu des applications visées, et seul quelques cas d'utilisation sont présentés.

Dans ce travail, nous proposons d'approfondir le travail de [7], en effectuant d'abord une analyse théorique plus poussée (à base des technique d'analyses des systèmes itératifs) de la technique proposée dans [7], tout en l'étendant à la notion de modulation Turbo-Orthogonale. L'idée est d'utiliser l'outil d'analyse EXIT Chart (EXtrinsic Information Transfer Chart en anglais), permettant d'observer les échanges d'information à l'intérieur du décodeur, et de l'adapter à la modulation étudiée. L'optimisation des paramètres du code devient alors possible, et la solution optimisée sera alors ensuite confrontée aux techniques et standard actuellement en concurrence. À notre connaissance, ni l'analyse EXIT du schéma proposé dans [7], ni la comparaison avec l'état de l'art des modulations utilisées pour l'IdO n'ont été réalisées. Dans cette étude, on insistera sur les motivations justifiant l'utilisant d'un schéma tel que [7]. Après une

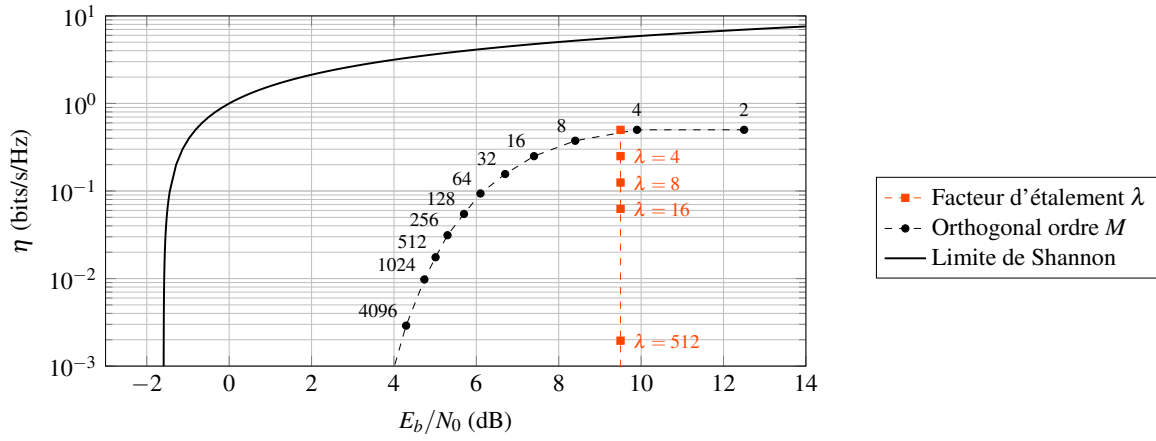


FIGURE 1 – L'efficacité spectrale en fonction de E_b/N_0 , pour différentes modulations et un TEB de 10^{-5} .

présentation du schéma, on effectuera l'analyse théorique et l'optimisation des paramètres, tout en considérant la nécessité de travailler avec des petites tailles de bloc. La comparaison du schéma avec l'état de l'art sera ensuite réalisée.

1 Motivations

Une technique largement utilisée afin d'améliorer la sensibilité d'un système est la répétition par un facteur λ , appelé facteur d'étalement. À la réception, en présence d'un Bruit Blanc Additif Gaussien (BBAG), le Rapport Signal à Bruit (RSB) en utilisant le facteur d'étalement est augmenté de $10\log_{10}(\lambda)$ dB. Cependant, le débit de la transmission est lui aussi réduit d'un facteur λ . On peut remarquer que l'influence du facteur d'étalement est différente si l'on s'intéresse, pour une fiabilité de transmission donnée, au rapport de l'énergie d'un bit "utile" sur la densité spectrale de puissance du bruit, noté E_b/N_0 , avec

$$\frac{E_b}{N_0} = \frac{\text{RSB}}{\eta}. \quad (1)$$

La Figure 1 représente l'efficacité spectrale normalisée (notée η , en bits/s/Hz) en fonction de l'efficacité énergétique par bit normalisée E_b/N_0 . Sont placées dans ce plan les performances de deux types de modulation : les modulations orthogonales d'ordre M , et la modulation de phase à deux états (BPSK) à laquelle on applique un facteur d'étalement λ . On y représente également la limite de Shannon [8], définie par

$$\eta \leq \log_2(1 + \text{RSB}) \Leftrightarrow \frac{E_b}{N_0} \geq \frac{2^\eta - 1}{\eta}. \quad (2)$$

Cette limite correspond à une efficacité énergétique maximale. On constate bien sur la Figure 1 que le facteur d'étalement ne permet pas de se rapprocher de la limite, mais seulement de diminuer l'efficacité spectrale (et donc la sensibilité). En revanche, augmenter la taille M de l'alphabet d'une modulation orthogonale permet de se rapprocher de cette limite, et de l'atteindre pour une taille d'alphabet infinie [4]. L'idée proposée dans [7] est de combiner une modulation orthogonale avec le principe turbo, afin de se rapprocher au maximum de la limite de Shannon tout en ayant une taille d'alphabet raisonnable.

2 Modèle du système proposé

On propose dans cette section de décrire les structures de l'émetteur et du récepteur de la technique, ainsi que les étapes de calcul réalisées lors du décodage souple.

2.1 L'émetteur

Le schéma de l'émetteur Turbo-Orthogonal est présenté Figure 2. Il est constitué de K branches, qui vont chacune encoder une version entrelacée des bits d'information. On définit la taille du bloc de bits d'information N fourni à chaque branche comme un multiple de l'entier r , tel que $N = P \times r$. L'encodeur convolutif orthogonal est représenté Figure 3. À partir d'un ensemble de r bits, l'encodeur effectue la somme binaire (*i.e* la parité), puis accumule cette valeur dans une mémoire. On utilise ensuite les r bits ainsi que la parité accumulée pour choisir quel mot de code orthogonal transmettre. La taille de l'alphabet orthogonal A dépend de l'entier r choisi : pour pouvoir associer chacune des possibilités de mot binaire, l'alphabet doit contenir 2^{r+1} mots de code. Chaque branche encode alors P mots de code. L'accumulation de la parité agit comme un moyen de partager de l'information entre deux mots de code consécutifs, et induit la présence d'un treillis (celui de l'accumulateur). Afin de fermer ce treillis, il est nécessaire d'émettre un symbole supplémentaire. Le rendement total peut alors s'exprimer

$$\eta = \frac{P \times r}{K(P+1)2^{r+1}}. \quad (3)$$

Une étape de conversion parallèle/série est ensuite effectuée afin de transmettre les mots de code orthogonaux sur le canal.

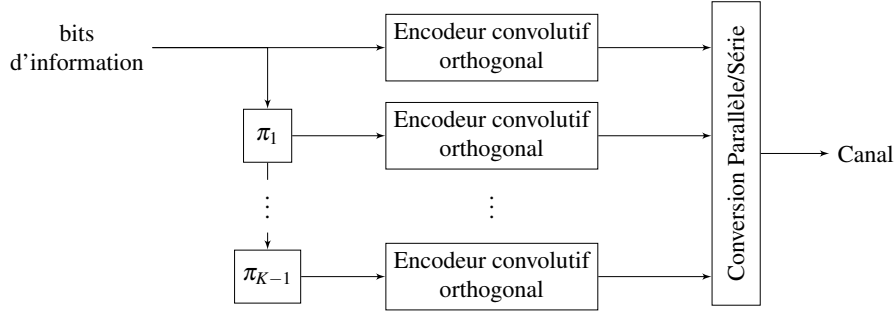


FIGURE 2 – L'émetteur Turbo-Orthogonal.

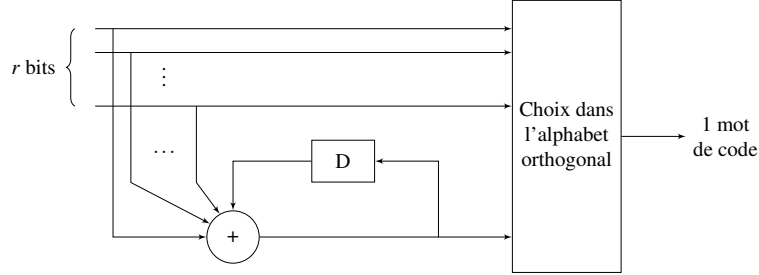


FIGURE 3 – L'encodeur convolutif orthogonal.

2.2 Canal

Le modèle de canal envisagé est le canal BBAG, dont la sortie est définie par

$$y = x + n \quad (4)$$

où n est un bruit blanc gaussien complexe centré de variance σ^2 .

2.3 Le récepteur

La structure du récepteur est donnée Figure 4. Une conversion Série/Parallèle est d'abord réalisée sur les mots de code orthogonaux qui ont transité par le canal. On obtient K vecteurs $y^{(k)}$, $k \in [0, \dots, K-1]$, chacun composés de $P+1$ mots de code ($y^{(k)} = [y^0, y^1, \dots, y^P]$). Le bloc "Projection Alphabet" consiste ensuite à réaliser le produit de $y^{(k)}$ avec la matrice transposée Hermitienne $\bar{A}$, *i.e* la projection du mot de code sur l'alphabet. On obtient la matrice $Y^{(k)}$ de taille $(P+1) \times 2^{r+1}$, qui sera utilisée comme observation du canal par les décodeurs. Le principe turbo consiste à échanger de façon itérative l'information entre les décodeurs. Après le désentrelacement qui convient, un décodeur utilisera donc l'observation du canal qui lui est associée ainsi que l'information extrinsèque des autres décodeurs (l'information *a priori*) afin de réaliser une estimation des bits d'information. L'information propagée est la matrice L des rapports de vraisemblance logarithmiques (RVL), de taille $P \times r$ et définie par

$$L = \left\{ L(b_{p,n}) = \log \frac{p(b_{p,n} = 1)}{p(b_{p,n} = 0)} \right\}_{\substack{p \in [0, \dots, P-1] \\ n \in [0, \dots, r-1]}}, \quad (5)$$

où $p(b_{p,n} = 1)$ (resp. $p(b_{p,n} = 0)$) est la probabilité que le bit $b_{p,n}$ soit égal à 1 (resp. 0). L est initialisée à 0 et ne contient pas d'information sur la terminaison du treillis. Le signe de $L(b_{p,n})$ est directement relié à l'information binaire, et le module correspond à sa vraisemblance. À la fin de chaque itération, L contient toute l'information de tous les décodeurs et une décision peut alors être prise.

2.4 Décodage souple

Afin de réaliser une estimation des bits d'information, chaque décodeur calcule les Probabilités *A Posteriori* (PAP) des bits, en utilisant l'observation du canal et l'information *a priori* provenant des autres décodeurs. On considère le p -ième mot de code y^p du k -ième étage. Le mot d'information décodé est noté $d = \{d_n\}_{n \in [0, \dots, r-1]}$. c^i est un mot de code de l'alphabet A et $b^i = \{b_n^i\}_{n \in [0, \dots, r-1]}$ est le mot d'information associé. Afin de généraliser le calcul, on considère que l'alphabet orthogonal est composé de $M = 2^{r+1}$ mots de code à valeurs complexes.

La PAP d'un mot de code est la probabilité d'avoir ce mot de code connaissant l'observation, avec

$$p(c^i | y^p) = \frac{p(y^p | c^i) p(c^i)}{p(y^p)} \quad (6)$$

d'après la formule de Bayes. $p(c^i)$ est la probabilité d'avoir le mot de code c^i , ou que le mot d'information décodé d soit b^i . En considérant que y^p et c^i sont des vecteurs de taille M composés de valeurs complexes, on a alors

$$p(c^i | y^p) = \frac{1}{p(y^p)} \prod_{m=0}^{M-1} p(y_m^p | c_m^i) \prod_{n=0}^{r-1} p(d_n = b_n^i). \quad (7)$$

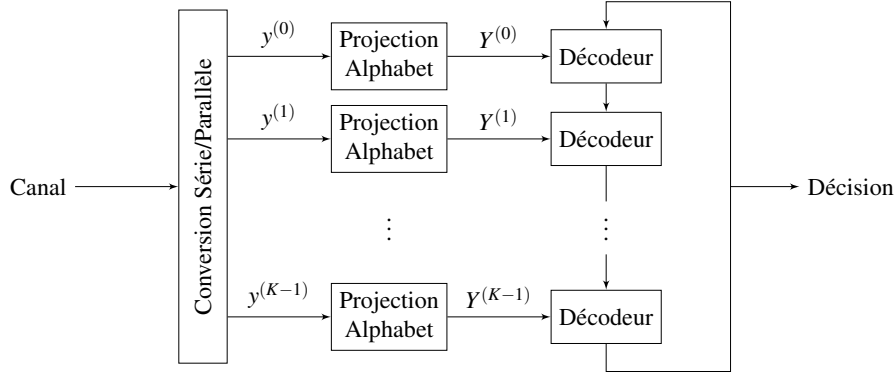


FIGURE 4 – Le récepteur Turbo-Orthogonal.

Dans le cas du canal BBAG complexe, on a

$$p(y_m^p | c_m^i) = \frac{1}{2\pi\sigma^2} \exp \left\{ -\frac{1}{2\sigma^2} \|y_m^p - c_m^i\|^2 \right\}, \quad (8)$$

le premier produit de (7) devient donc

$$\prod_{m=0}^{M-1} p(y_m^p | c_m^i) = B \exp \left\{ \frac{1}{\sigma^2} \sum_{m=0}^{M-1} \text{Re} \left(y_m^p \cdot \overline{c_m^i} \right) \right\}, \quad (9)$$

où $\overline{c_m^i}$ est le complexe conjugué de c_m^i , Re la partie réelle et B une constante.

En utilisant la définition des RVL de (5) et $p(d_n = 1) = 1 - p(d_n = 0)$, on a

$$p(d_n = b_n^i) = C \exp \left\{ L(d_n) (1 - 2b_n^i) / 2 \right\} \quad (10)$$

avec C une constante. Au final, en utilisant (9) et (10), (7) devient

$$p(c^i | y^p) = D \exp \left\{ \frac{1}{\sigma^2} \sum_{m=0}^{M-1} \text{Re} \left(y_m^p \cdot \overline{c_m^i} \right) + \sum_{n=0}^{r-1} \frac{1 - 2b_n^i}{2} L(d_n) \right\}, \quad (11)$$

où D est une constante. Cette expression fait apparaître le résultat de la projection du mot de code reçu y^p sur l'alphabet (la partie observation du canal, qui n'est calculée qu'une seule fois) et une combinaison particulière des RVL des bits d'information (la partie *a priori*, mise à jour à chaque itération). Ainsi, en fournissant le résultat de la projection des mots de code sur l'alphabet ainsi que la matrice L , les PAP de tous les mots de code c^i de l'alphabet peuvent être calculées avec (11).

Comme mentionné précédemment, l'utilisation du code convolutif au niveau de l'émetteur implique un partage d'information entre deux mots de code consécutifs et induit la construction d'un treillis. Ceci permet l'utilisation de l'algorithme Bahl, Cocke, Jelinek, et Raviv (BCJR) [9], tel que proposé dans [6], afin de mettre à jour les probabilités avec la connaissance des probabilités des autres mots de code. On dénote c^i , $i \in [0, \dots, M-1]$ les mots de code de l'alphabet. On peut alors construire le treillis pour le cas particulier $r = 2$, qui est représenté sur la Figure 5. On constate que chaque transition est représentée par $M/4$ mots de code (donc 2 dans le cas $r = 2$). Chaque branche est également orthogonale aux autres branches, ce qui est une des caractéristiques fondamentales de cette technique. L'algorithme BCJR sert alors à recalculer les probabilités en prenant en compte la connaissance que l'on a du treillis. On note les probabilités mises à jour $P(c^i | y^p)$.

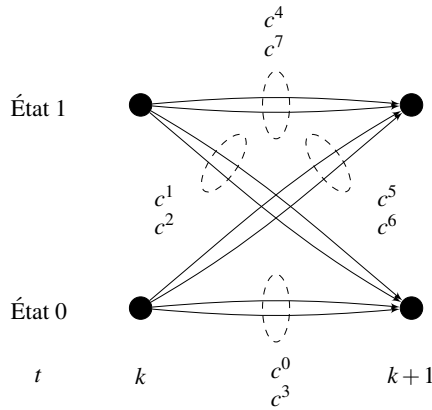


FIGURE 5 – Treillis pour le cas $r = 2$.

Les RVL des bits d'information peuvent alors être calculés grâce à

$$L(d_n|y^p) = \log \sum_{i, b_n^i=1}^{2^{r+1}} P(c^i|y^p) - \log \sum_{i, b_n^i=0}^{2^{r+1}} P(c^i|y^p), \quad (12)$$

où les indices des sommes indiquent que l'on doit sommer les probabilités des mots de code d'indice i correspondants à un mot d'information pour lequel le bit b_n^i vaut 1 (somme de gauche) ou 0 (somme de droite). Ce décodage souple s'avère être équivalent au Maximum *A Posteriori* (MAP).

3 Analyse théorique des performances et confrontations aux simulation

La technique proposée dans la partie précédente dispose de deux paramètres K et r qui vont influencer sur ses performances. On cherche alors à effectuer une analyse théorique afin de déterminer le jeu de paramètres maximisant l'efficacité énergétique. Un outil d'analyse efficace pour les processus itératifs est le suivi de l'échange d'information extrinsèque entre les décodeurs, également appelé EXtrinsic Information Transfer (EXIT) chart en anglais. Cet outil a été introduit par ten Brink [10], dans le cas particulier des codes convolutifs concaténés en parallèle (PCCC en anglais, ou aussi Turbo Code).

3.1 Calcul de la courbe EXIT

Chaque décodeur est décrit par une courbe EXIT qui lui est caractéristique. L'idée est de mesurer l'information extrinsèque de sortie du décodeur lorsqu'on lui fournit une certaine information *a priori*, qui est simulée par un modèle.

3.1.1 Information mutuelle

La métrique utilisée est l'Information Mutuelle (IM). L'IM entre deux sources X et Y peut être définie par la divergence de Kullbach-Leiber [11]

$$I(X, Y) = \int_x \int_y f_{X,Y}(x, y) \log \frac{f_{X,Y}(x, y)}{f_X(x) f_Y(y)} dx dy, \quad (13)$$

avec f la fonction de densité des distributions. On considère la source X comme étant les bits d'information, prenant donc comme valeurs $+1$ et -1 (avec le mapping $0 \rightarrow 1, 1 \rightarrow -1$). L'IM peut alors s'exprimer

$$I(X, Y) = \frac{1}{2} \sum_{x=\pm 1} \int_y f_{Y|X}(y | X=x) \log \frac{2f_{Y|X}(y | X=x)}{f(y|x=+1) + f(y|x=-1)} dy. \quad (14)$$

3.1.2 Modèle *a priori*

Le modèle d'information *a priori* choisi est un modèle gaussien $\mathcal{N}(\mu_A, \sigma_A)$ tel que suggéré dans [10]. Bien qu'étant une hypothèse très simplificatrice, une observation de l'allure des RVL échangés observés dans le cas Turbo-Orthogonal permet de la valider. On s'affranchit du paramètre μ_A en supposant la consistance des RVL *a priori*. On a alors $\mu_A = \sigma_A^2/2$, et les RVL *a priori* s'expriment donc

$$L_A = \mu_A \cdot x + n_A \quad (15)$$

où x est l'information ($x = \pm 1$) et n_A un bruit blanc gaussien centré d'écart type σ_A . On peut alors simplifier (14) en prenant pour Y les RVL *a priori* et définir la fonction J telle que [10]

$$I_A = J(\sigma_A) = 1 - \int_{-\infty}^{+\infty} \frac{1}{\sqrt{2\pi}\sigma_A} \exp \left\{ -\frac{1}{2\sigma_A^2} \left(z - \frac{\sigma_A^2}{2} x \right)^2 \right\} \log(1 + e^{-z}) dz \quad (16)$$

3.1.3 L'information extrinsèque dans le cas Turbo-Orthogonal

D'une façon générale, l'information contenue dans les RVL de sortie d'un décodeur itératif se compose de trois éléments : l'information *a priori* L_A , provenant des autres décodeurs, l'information provenant du canal L_{ch} , et l'information extrinsèque générée par le décodeur L_E . Le RVL L s'exprime alors

$$L = L_A + L_{ch} + L_E \quad (17)$$

On peut donc extraire L_E en soustrayant simplement les autres termes au RVL mesuré L . Dans le cas du Turbo Code classique, les mots de code reçus contiennent une partie systématique (directement les bits transmis), et une partie parité (les bits de redondance ajoutés par l'encodage). On peut alors directement retrouver le RVL provenant du canal en ne gardant que la partie systématique des mots de code provenant du canal. Dans le cas Turbo-Orthogonal, ceci n'est pas toujours possible. Si on prend un alphabet orthogonal composé de mots de Hadamard, on se retrouve dans le cas précédent, car le code de Hadamard est systématique. En revanche, si on utilise un alphabet FSK, composé d'éléments de la matrice de Fourier discrète, les mots de code ne sont plus systématiques. Il n'est alors plus possible d'extraire l'information provenant du canal pour les bits d'information seuls. Tout décodage souple de l'alphabet (sans faire complètement le décodage turbo) entraîne l'ajout d'une information extrinsèque due à l'orthogonalité du code. Il est alors nécessaire de prendre en compte cette information résiduelle provenant du canal.

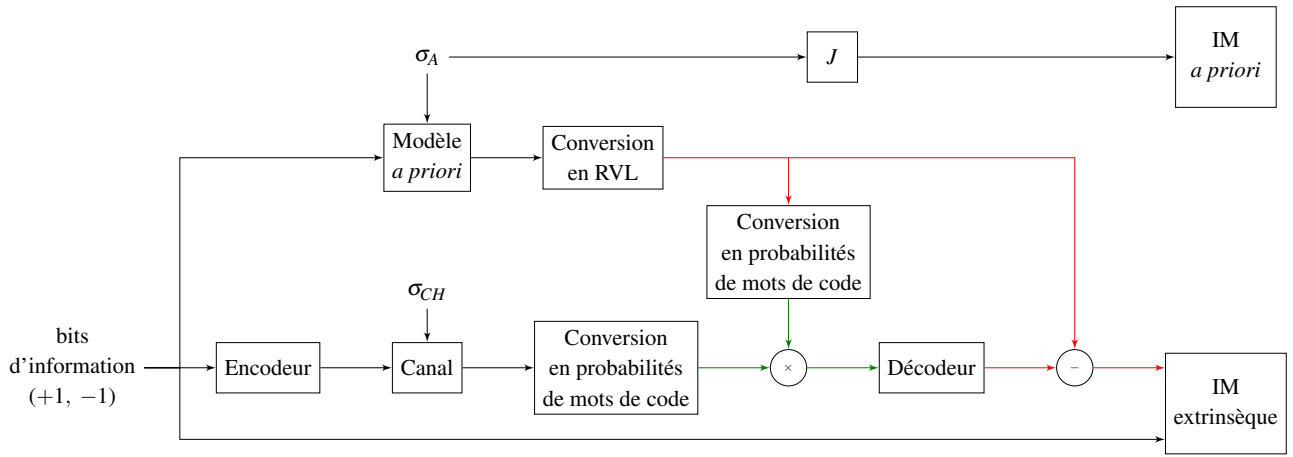


FIGURE 6 – Schéma bloc du calcul de l'EXIT chart.

3.1.4 Schéma bloc

Sur la Figure 6, on représente le schéma bloc permettant de mesurer l'EXIT chart. En vert sont représentées les probabilités de mots de code de l'alphabet. La multiplication réalisée est celle de (7) entre le terme de gauche (provenant du canal) et le terme de droite (provenant des autres décodeurs). En rouge, on représente les RVL. On constate que comme précisé précédemment, on ne peut que soustraire l'information *a priori*, n'ayant pas accès directement à l'information des bits d'informations provenant du canal. Le bloc de calcul de l'IM extrinsèque réalise la somme d'intégrales de (14). L'IM *a priori* est mesurée avec la fonction J de (16). Pour une valeur spécifique de σ_{ch} , on fait varier σ_A afin d'obtenir toutes les valeurs d'IM *a priori* entre 0 et 1. On mesure pour chacun de ces valeurs l'IM extrinsèque de sortie, et si le processus converge vers le point (1,1), il n'y a pas d'erreurs. Il est nécessaire d'utiliser des très grandes tailles de bloc de bits d'information, afin d'avoir une bonne statistique pour les mesures de fonction de densité de probabilités. Il est à noter également que le calcul de la courbe EXIT ne prend pas en compte la fonction d'entrelacement.

3.1.5 Cas multi-dimensionnel

Dans le cas où le récepteur est composé de deux décodeurs, on peut alors représenter dans un plan en deux dimensions les EXIT chart des deux décodeurs et suivre l'échange d'information qui se produit lors des itérations. En revanche, si le nombre de décodeurs est supérieur à deux, comme dans le cas du système étudié ici (K décodeurs identiques), un travail de généralisation doit être réalisé. L'échange d'information s'observe alors dans un espace à $K - 1$ dimensions [12]. Cependant, si tous les décodeurs sont identiques, la convergence vers le point $(1, \dots, 1)$ peut s'estimer en calculant une "ligne de surface". Pour cette mesure, l'IM *a priori* s'exprime alors [12]

$$I_A = J\left(\sqrt{K-1} \cdot \sigma_A\right). \quad (18)$$

Ceci traduit le fait que chaque décodeur reçoit l'information des $K - 1$ autres décodeurs. Le calcul de l'IM extrinsèque reste inchangé.

3.1.6 L'EXIT chart

On réalise une mesure d'EXIT pour un cas particulier du schéma Turbo-Orthogonal, celui $r = 6$, $K = 4$ et où un alphabet FSK de taille $M = 2^{r+1} = 128$ est utilisé. On utilise le processus décrit sur la Figure 6, pour finalement obtenir les courbes EXIT présentées sur la Figure 7. Le nombre de bits d'information par bloc a été fixé à $N = 100000$. L'EXIT chart a été mesuré pour trois valeurs de E_b/N_0 . On a également représenté la ligne médiane allant de $(0,0)$ à $(1,1)$. On constate que l'intersection entre cette ligne médiane et la courbe noire ($E_b/N_0 = -1.6\text{dB}$) apparaît pour $I_A \simeq 0.1$. On peut alors conclure qu'à ce niveau de E_b/N_0 , le système ne pourra atteindre le point $(1,1)$, et que donc peu importe le nombre d'itérations, un certain nombre d'erreurs persisteront. En revanche, pour des valeurs de E_b/N_0 de 0 et 1dB, il n'y a pas intersection avec la ligne médiane, et donc le point $(1,1)$ peut être atteint après un certain nombre d'itérations.

Une recherche exhaustive nous permet de déterminer la plus petite valeur de E_b/N_0 pour laquelle la courbe EXIT est "ouverte", c'est-à-dire qu'il n'y a plus intersection avec la ligne médiane. On appellera dans la suite cette valeur le "point d'ouverture" de la courbe EXIT. Pour le cas choisi précédemment, on trouve que le point d'ouverture est situé à $E_b/N_0 = -1.12\text{dB}$, et la courbe EXIT pour cette valeur est représentée sur la Figure 8 (a). On constate bien qu'il n'y a pas intersection avec la ligne médiane. Le système permettrait donc de corriger toutes les erreurs induites par le canal. Pour vérifier cela, on réalise un calcul de Taux d'Erreur Binaire (TEB) pour le même cas $r = 6$, $K = 4$, où une très grande taille de bloc est utilisée ($N = 100000$). Pour le décodeur, on utilise la règle du MAP sans aucune approximation, et on réalise 100 itérations de décodage. On obtient alors la Figure 8 (b). On constate que le TEB commence à diminuer de façon significative après $E_b/N_0 = -1.12\text{dB}$, pour atteindre une valeur de 10^{-5} à $E_b/N_0 = -1.03\text{dB}$. La courbe EXIT permet donc bien de prédire le point à partir duquel le TEB va chuter (le "waterfall" en anglais).

3.2 Optimisation des paramètres

Puisque la courbe EXIT permet de prédire de façon relativement précise le moment où va se déclencher le waterfall, on peut réaliser une recherche exhaustive du point d'ouverture pour chacun des deux paramètres r et K . On obtient alors les valeurs présentées dans le Tableau 1. On a noté avec "*" le point d'ouverture minimum, et ce pour chaque valeur de r . On constate qu'il n'y a pas de logique

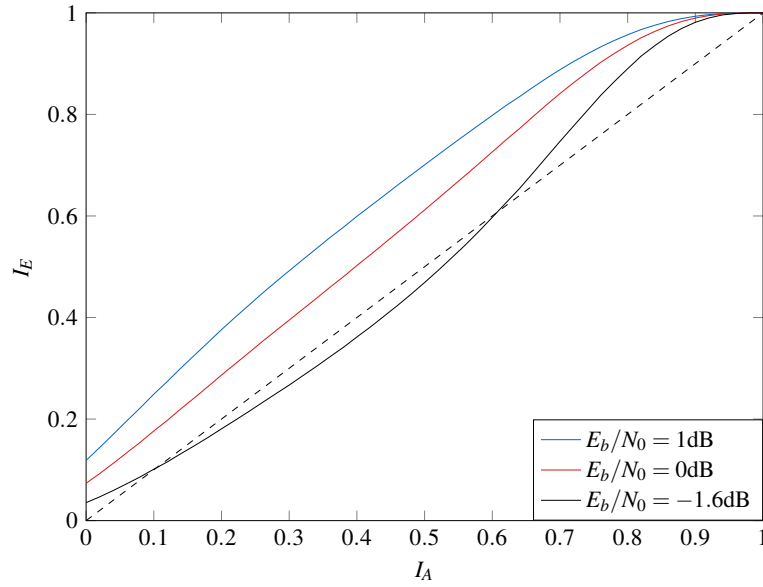


FIGURE 7 – L'EXIT chart du décodeur pour le cas $r = 6$, $K = 4$, pour différentes valeurs de E_b/N_0

$K \backslash r$	3	4	5	6	7	8	9	10
3	0.33	-0.17	-0.59	-0.94	-1.23*	-1.30*	-1.19*	-1.07*
4	0.02	-0.46	-0.84	-1.12*	-1.20	-1.04	-0.83	-0.59
5	-0.06*	-0.51*	-0.86*	-1.09	-1.03	-0.77	-0.47	-0.18
6	-0.06*	-0.51*	-0.84	-1.02	-0.84	-0.52	-0.14	0.19

TABLE 1 – Valeurs de E_b/N_0 de point d'ouverture de la courbe EXIT, pour tous les jeux de paramètres. Les courbes EXIT sont mesurées pour une taille de bloc $N = 100000$.

précise sur la valeur K à choisir pour obtenir une faible valeur de E_b/N_0 pour le point d'ouverture, d'où la nécessité notamment de réaliser cette recherche de manière exhaustive. Le point d'ouverture le plus bas est pour le cas $r = 8$, $K = 3$, qui se trouve alors à seulement 0.3dB de la limite de Shannon (environ égale à -1.59 dB pour les faibles efficacités spectrales). Dans le contexte des réseaux longue portée basse consommation, le besoin est plutôt orienté vers l'utilisation de petites tailles de paquet (de l'ordre d'une centaine d'octets). L'analyse EXIT est intéressante pour déterminer le point de waterfall, mais elle implique cependant l'utilisation de taille de bloc très importante. Réduire la taille va donc modifier la position du waterfall, et il est alors intéressant de voir dans quelle mesure cela modifier l'optimisation des paramètres.

On effectue alors des simulations de TEB pour chacun des couples (r, K) , en fixant la taille de bloc à la valeur $N = 1000$. Le décodeur réalise 10 itérations, et on note la valeur de E_b/N_0 pour laquelle le TEB vaut 10^{-4} . On obtient les valeurs du Tableau 2. On note ici encore les valeurs de E_b/N_0 les plus basses par "*", et ce pour chaque valeur de r . On constate que les positions de ces minimums sont en accord avec celles déterminées par l'analyse EXIT. L'écart en E_b/N_0 avec les points d'ouverture dépend cependant des valeurs de paramètres, variant entre $+0.80$ et $+1.25$ dB.

L'efficacité spectrale du système étant fonction de r et de K , chaque jeu de paramètre possède son efficacité spectrale propre,

$K \backslash r$	3	4	5	6	7	8	9	10
3	1.40	0.89	0.48	0.14	-0.13	-0.22*	-0.24*	0.00*
4	0.95	0.49	0.13	-0.10*	-0.17*	0.02	0.35	0.52
5	0.84	0.40	0.08*	-0.04	0.08	0.35	0.78	1.07
6	0.74*	0.37*	0.13	0.04	0.35	0.73	1.10	1.46

TABLE 2 – Valeurs de E_b/N_0 à partir desquelles le TEB est inférieur à 10^{-4} , pour tous les jeux de paramètres. Le TEB est mesuré pour une taille de bloc $N = 1000$. L'algorithme du MAP est utilisé, 10 itérations sont réalisées.

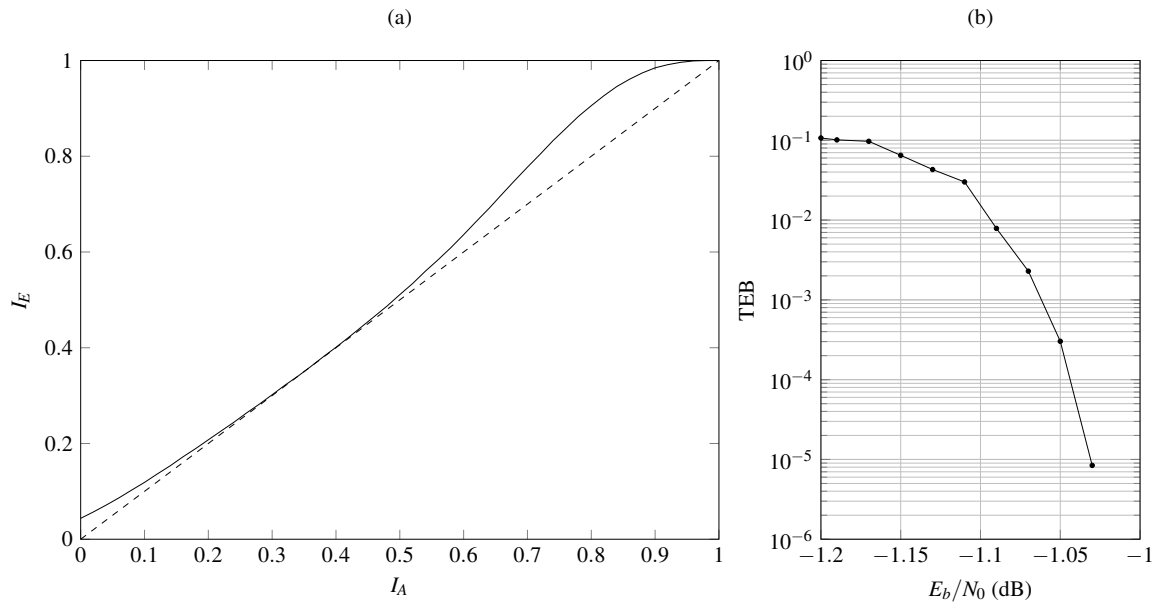


FIGURE 8 – (a) L’EXIT chart du cas $r = 6$, $K = 4$ au point d’ouverture $E_b/N_0 = -1.12$ dB. (b) Performances en TEB pour le même cas et pour une taille de bloc $N = 100000$. L’algorithme MAP est utilisé, 100 itérations sont réalisées.

calculée avec l’équation (3). On peut alors représenter, pour chaque cas, l’efficacité spectrale en fonction du E_b/N_0 soit d’ouverture de l’EXIT, soit de $TEB=10^{-4}$ (avec $N = 1000$). On obtient la Figure 9. Chaque r est distingué par une couleur, et les quatre valeurs possibles pour K sont représentées. Puisque augmenter K réduit l’efficacité spectrale, le point possédant la plus haute efficacité spectrale correspond au point $K = 3$, le deuxième à $K = 4$ et ainsi de suite. On constate ici encore qu’augmenter la valeur de K n’est pas toujours synonyme d’une amélioration de performances, notamment pour les fortes valeurs de r . On constate également qu’un fort gain est envisageable comparé aux modulations orthogonales non codées, et ce même avec une petite taille de bloc.

4 Comparaison des solutions

Afin de mieux situer le système Turbo-Orthogonal dans le contexte des techniques de communication M2M, on propose de confronter ses performances à trois autres techniques : le standard IEEE 802.15.4k [13], la technique définie dans le brevet [14], qui semble être la base de la technique LoRa [3], et enfin une autre technique qui consisterait à utiliser une modulation orthogonale, avec un Turbo Code (TC) comme codage canal.

4.1 Les solutions envisagées

On propose ici une courte description des solutions envisagées. Pour chacune des techniques, le paramètre λ va être réglé pour atteindre une efficacité spectrale donnée, afin de réaliser une comparaison équitable entre toutes les techniques.

4.1.1 Le standard 802.15.4k

Ce standard est dédié aux réseaux basse consommation pour le suivi des infrastructures [13]. La couche physique la plus robuste définie par le standard (et choisie pour notre comparaison) consiste en une modulation BPSK différentielle (DBPSK), avec une Code Convolutif (CC). Le standard permet également l’utilisation d’un facteur d’étalement λ allant de 16 à 32768. L’efficacité spectrale normalisée en bande du standard est définie par

$$\eta_1 = \frac{1}{2\lambda}. \quad (19)$$

4.1.2 La technique de type LoRa

La modulation LoRa est destinée aux réseaux longue portée basse consommation et à l’IdO [3]. Cette technique consiste en la concaténation d’un alphabet orthogonal avec un code bloc de Hamming. L’alphabet orthogonal choisi est de taille λ , où $\log_2 \lambda$ est compris entre 7 et 12. Le code correcteur génère des bits de redondance pour chaque bloc de 4 bits d’informations. Le nombre de bits de redondance p peut être choisi entre 0 (pas de codage) et 4, offrant ainsi un certain nombre de rendement de code possible, qu’on dénote $R = 4/(4 + p)$. L’efficacité spectrale normalisée en bande de la technique est définie par

$$\eta_2 = \frac{\log_2(\lambda)}{\lambda R}. \quad (20)$$

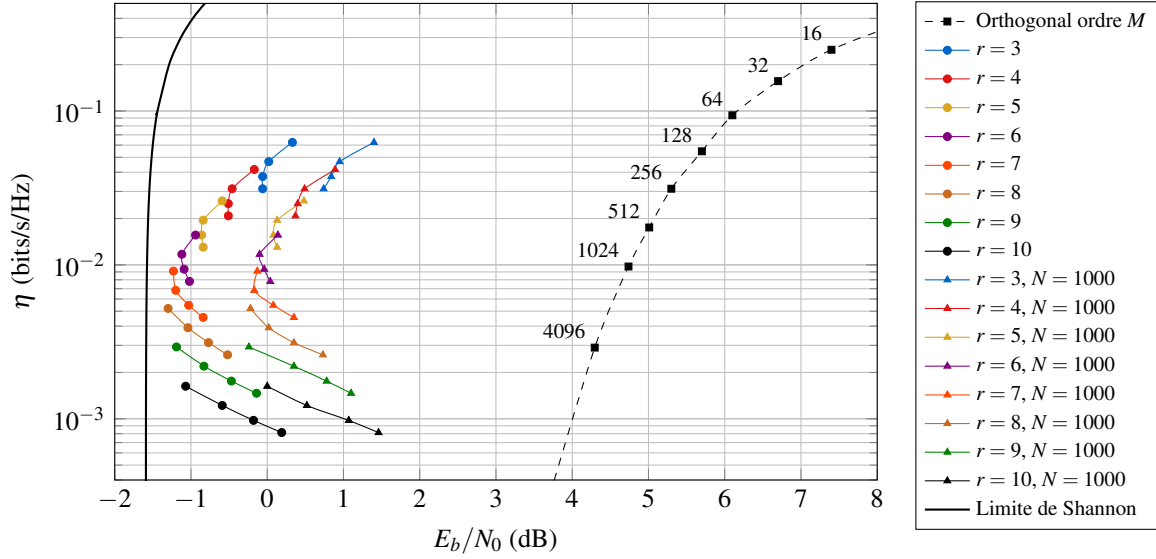


FIGURE 9 – Présentation des points de d'ouverture de l'EXIT dans le plan η vs E_b/N_0 .

4.1.3 Modulation orthogonale et Turbo Code

Bien que non suggéré dans la littérature, il est possible d'envisager comme couche physique une modulation orthogonale, bas débit et robuste, avec un codage canal de type Turbo Code. L'émetteur resterait ainsi relativement peu complexe. On utilise le TC de rendement $1/3$ spécifié par la norme 3G [15, 16], et une modulation orthogonale de type FSK où l'alphabet est de taille M . On suppose également la présence d'un facteur d'étalement λ . L'efficacité spectrale normalisée en bande est alors définie par

$$\eta_3 = \frac{\log_2(M)}{3\lambda M}. \quad (21)$$

4.1.4 La Turbo-FSK

On choisit, pour le Turbo-Orthogonal, un alphabet orthogonal de type FSK [7]. Afin d'harmoniser les notations, on note par λ ce qui était précédemment le paramètre K . La taille de l'alphabet utilisé vaut $M = 2^{r+1}$. L'efficacité spectrale normalisée en bande est maintenant définie par

$$\eta_4 = \frac{P(\log_2(M) - 1)}{\lambda(P+1)M}. \quad (22)$$

où P est tel que défini dans (3).

4.2 Comparaison des performances

Pour comparer les performances, on sélectionne un cas particulier du Turbo-Orthogonal, le cas $r = 6$, $K = 4$. Ce cas présente en effet de bonnes performances avec une taille d'alphabet relativement faible. Avec les paramètres définis pour la Turbo-FSK paragraphe 4.1.4, ce cas correspond à une taille d'alphabet de 128 et à $\lambda = 4$. On prend une taille de bloc $N = 1000$. L'équation (22) nous donne un rendement égal à $\eta_4 = 1.17 \cdot 10^{-2}$. On peut alors, pour chaque technique, jouer sur le paramètre λ afin d'approcher à chaque fois la même efficacité spectrale. Ceci donne les paramètres indiqués Tableau 3. Afin de réaliser les simulations, on suppose un modèle de canal à bruit blanc additif gaussien, avec réception cohérente. La taille des paquets transmis est fixée à $N = 1000$ (parfois arrondi pour être un multiple du nombre de bits par symbole), taille appropriée dans le contexte des réseaux longue portée basse consommation. Au décodage, on utilise un décodage souple pour le CC et le code de Hamming, et pour les processus turbo, la règle du MAP est choisie et 10 itérations sont réalisées. Un entrelacement aléatoire est réalisé, sauf pour le cas du TC 3G où on utilise l'entrelaceur spécifié par la norme.

On réalise pour les paramètres donnés Tableau 3 une mesure de TEB. Le résultat est reporté sur la Figure 10. On a représenté le TEB en fonction du RSB, où la conversion depuis le E_b/N_0 est faite grâce à l'équation (1). On représente également le Taux d'Erreur Paquet (TEP) en fonction du RSB sur la Figure 11 (un paquet est déclaré erroné si il existe au moins une erreur). Ces deux figures permettent de constater le réel intérêt à utiliser un décodage turbo à la réception. En effet, à un TEB de 10^{-5} , la 128-FSK avec le TC de la 3G offre un gain d'environ 3dB par rapport à la modulation de type LoRa. La Turbo FSK avec les paramètres sélectionnés permet elle un gain de 1.7dB par rapport à la FSK avec un TC, et donc de 4.7dB par rapport à la modulation de type LoRa.

Les niveaux de RSB observés sont directement reliés au niveau de sensibilité de telle ou telle technique. Les quatre modulations permettent donc bien d'atteindre des sensibilités très faibles, tout en ayant une faible efficacité spectrale. Les gains peuvent eux s'interpréter également comme les gains en sensibilité. Ainsi, utiliser la Turbo-FSK plutôt que le standard 802.15.4k avec la même bande permettra un fonctionnement avec une sensibilité réduite de 5.8dB pour un même TEP de 10^{-3} . Ou, pour les mêmes performances de sensibilité et de TEP, la puissance du signal émis peut être divisée par 3.8, ou encore la distance entre l'émetteur et le récepteur augmenté d'un facteur 1.9 (en considérant une perte en espace libre).

Technique	802.15.4k	Technique LoRa	FSK+TC	Turbo-FSK
Modulation	DBPSK	512-Orthog	128-FSK	128-FSK
Encodage	CC [171 133]	Hamming	TC [13 15]	Turbo-FSK
Débit binaire	1/2	4/6	1/3	-
λ	43	9	2	4
η ($\cdot 10^{-2}$)	1.163	1.172	0.907	1.170

TABLE 3 – Paramètres de chacune des modulations sélectionnées pour la comparaison. La taille de bloc est fixée à $N = 1000$.

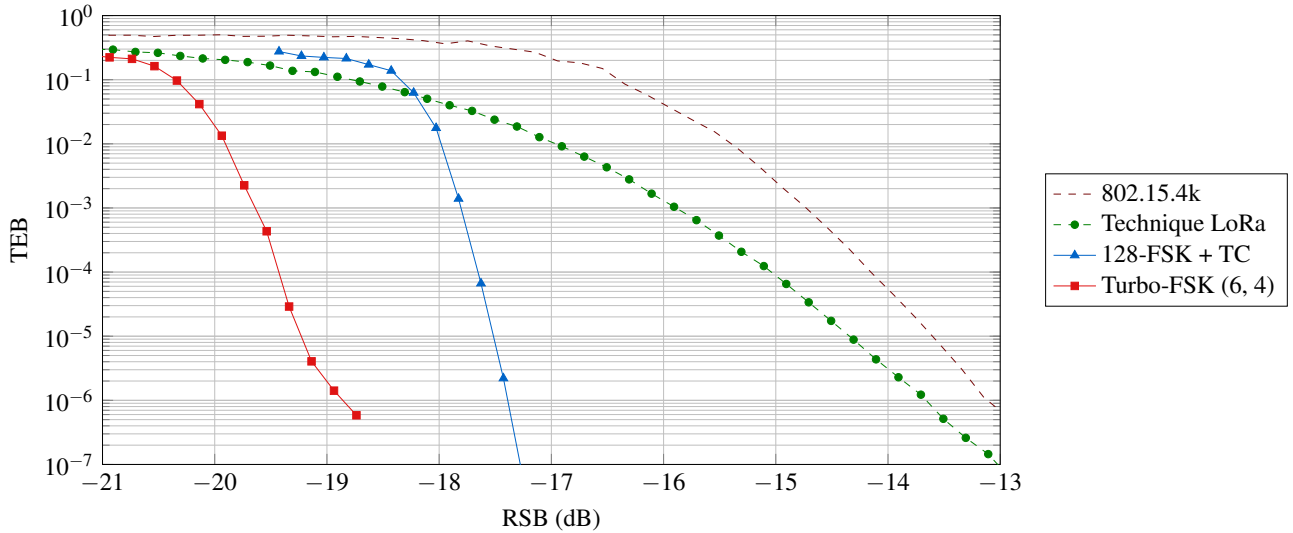


FIGURE 10 – TEB des couches physiques comparées. Les paramètres utilisés sont ceux du Tableau 3.

5 Conclusions

La technique Turbo-Orthogonale se présente comme un candidat sérieux pour la couche physique des réseaux longue portée et basse consommation. Après la description de la technique et de la méthode d'optimisation des paramètres, basée sur l'analyse EXIT, force est de constater que d'importants gains en sensibilités peuvent être envisagés en utilisant cette couche physique par rapport aux techniques actuellement utilisées. De même, les multiples points de fonctionnement proposés permettent un certain nombre de degrés de liberté pour atteindre des contraintes en efficacité spectrale ou en efficacité énergétique. La structure même de l'encodeur est très simple, impliquant une faible consommation au niveau de l'émetteur, un des critères fondamentaux pour la conception de la liaison montante pour le type de réseau ciblé. Le gain en performance se fait bien entendu au détriment d'une complexité augmentée, mais cette dernière est reportée sur la station de base, pour laquelle on suppose une ressource illimitée. Un des points clés restant à étudier est l'étude de la synchronisation et de la détection du paquet, qui peuvent en effet s'avérer complexe pour les niveaux de RSB envisagés.

Références

- [1] T. Rebeck, M. Mackenzie, and N. Afonso, "Low-powered wireless solutions have the potential to increase the M2M market by over 3 billion connections," *Analysys Mason*, Sept 2014.
- [2] "SigFox website," <http://www.sigfox.com/>, accessed : 3 mars 2016.
- [3] "LoRa Alliance," <https://www.lora-alliance.org/>, accessed : 3 mars 2016.
- [4] J. Proakis, *Digital Communications 3rd Edition*, ser. Communications and signal processing. McGraw-Hill, 1995.
- [5] C. Berrou, A. Glavieux, and P. Thitimajshima, "Near Shannon limit error-correcting coding and decoding : Turbo-codes. 1," in *IEEE International Conference on Communications (ICC)*. Geneva., vol. 2, May 1993, pp. 1064–1070.
- [6] L. Ping, W. Leung, and K. Y. Wu, "Low-rate turbo-Hadamard codes," *IEEE Transactions on Information Theory*, vol. 49, no. 12, pp. 3213–3224, Dec 2003.
- [7] Y. Roth, J.-B. Dore, L. Ros, and V. Berg, "Turbo-FSK : A New Uplink Scheme for Low Power Wide Area Networks," in *2015 IEEE 16th International Workshop on Signal Processing Advances in Wireless Communications (SPAWC)*, June 2015, pp. 81–85.
- [8] C. Shannon, "A mathematical theory of communication," *The Bell System Technical Journal*, vol. 27, no. 3, pp. 379–423, July 1948.

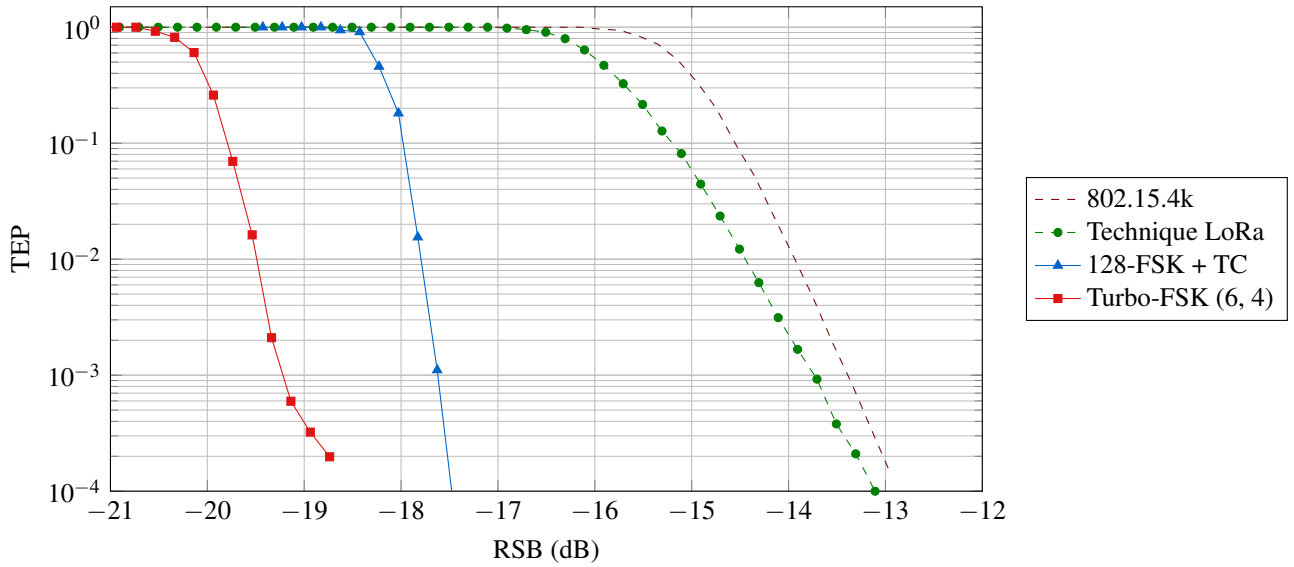


FIGURE 11 – TEP des couches physiques comparées. Les paramètres utilisés sont ceux du Tableau 3.

- [9] 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.
- [10] S. ten Brink, "Convergence behavior of iteratively decoded parallel concatenated codes," *IEEE Transactions on Communications*, vol. 49, no. 10, pp. 1727–1737, Oct 2001.
- [11] S. Kullback and R. A. Leibler, "On information and sufficiency," *Ann. Math. Statist.*, vol. 22, no. 1, pp. 79–86, 03 1951. [Online]. Available : <http://dx.doi.org/10.1214/aoms/1177729694>
- [12] S. ten Brink, "Convergence of multidimensional iterative decoding schemes," in *Conference Record of the Thirty-Fifth Asilomar Conference on Signals, Systems and Computers*, vol. 1, Nov 2001, pp. 270–274 vol.1.
- [13] "802.15.4k : Low-Rate Wireless Personal Area Networks (LR-WPANs) Amendment 5 : Physical Layer Specifications for Low Energy, Critical Infrastructure Monitoring Networks." *IEEE Standard for Local and metropolitan area networks*, pp. 1–149, Aug 2013.
- [14] O. Seller and N. Sornin, "Low power long range transmitter," *US Patent 20140219329 A1*, Aug 2014.
- [15] "LTE Evolved Universal Terrestrial Radio Access (E-UTRA) : Multiplexing and Channel Coding," *3GPP TS 36.212, V12.6.0, Release 12*, pp. 12–15, Oct 2015.
- [16] M. C. Valenti and J. Sun, "The UMTS turbo code and an efficient decoder implementation suitable for software defined radios," *International Journal of Wireless Information Networks*, vol. 8, pp. 203–216, 2001.