



NRC Publications Archive Archives des publications du CNRC

Design of a transmitter-receiver that uses ultra-wideband technology Berthiaume-Dutrisac, Z.

NRC Publications Record / Notice d'Archives des publications de CNRC:

<https://nrc-publications.canada.ca/eng/view/object/?id=0a79a1e0-a43d-42b7-b17c-26514c325fa9>

<https://publications-cnrc.canada.ca/fra/voir/objet/?id=0a79a1e0-a43d-42b7-b17c-26514c325fa9>

Access and use of this website and the material on it are subject to the Terms and Conditions set forth at

<https://nrc-publications.canada.ca/eng/copyright>

READ THESE TERMS AND CONDITIONS CAREFULLY BEFORE USING THIS WEBSITE.

L'accès à ce site Web et l'utilisation de son contenu sont assujettis aux conditions présentées dans le site

<https://publications-cnrc.canada.ca/fra/droits>

LISEZ CES CONDITIONS ATTENTIVEMENT AVANT D'UTILISER CE SITE WEB.

Questions? Contact the NRC Publications Archive team at

PublicationsArchive-ArchivesPublications@nrc-cnrc.gc.ca. If you wish to email the authors directly, please see the first page of the publication for their contact information.

Vous avez des questions? Nous pouvons vous aider. Pour communiquer directement avec un auteur, consultez la première page de la revue dans laquelle son article a été publié afin de trouver ses coordonnées. Si vous n'arrivez pas à les repérer, communiquez avec nous à PublicationsArchive-ArchivesPublications@nrc-cnrc.gc.ca.





National Research
Council Canada

Institute for
Information Technology

Conseil national
de recherches Canada

Institut de technologie
de l'information

NRC - CNRC

Conception d'un émetteur-récepteur utilisant la technologie à bande ultra large *

Berthiaume-Dutrisac, Z.
Août 2000

* published as NRC/ERB-1078. Août 2000. 90 pages. NRC 44144.

Copyright 2000 by
National Research Council of Canada

Permission is granted to quote short excerpts and to reproduce figures and tables from this report, provided that the source of such material is fully acknowledged.



***Conception d'un émetteur-récepteur
utilisant la technologie
à bande ultra large***

Zoé Berthiaume-Dutrisac
Août 2000

École Polytechnique de Montréal
Université de Montréal



Conception d'un émetteur-récepteur utilisant la technologie à bande ultra large

Projet de fin d'études soumis
comme condition partielle à l'obtention du
diplôme de baccalauréat en ingénierie

DÉPARTEMENT DE GÉNIE ÉLECTRIQUE ET DE GÉNIE INFORMATIQUE

Présenté par : Zoé Berthiaume-Dutrisac
Matricule : 61564
Directeur de projet : Michel Lemire
Codirecteur de projet : Larry Korba
Entreprise : Conseil national de recherches
du Canada

Le 21 août 2000

CNRC-NRC

Sommaire

Le projet, réalisé au sein du *Groupe de réseautage* de l'*Institut de technologie de l'information* du *Conseil national de recherches du Canada* à Ottawa, a consisté en l'application de la technologie récente à bande ultra large à un système de télécommunications. Les formes d'ondes à bande ultra large sont des impulsions non sinusoïdales de très courte durée et couvrent donc fondamentalement une très grande partie du spectre des fréquences. Dû au très faible rapport cyclique de ces impulsions, la puissance moyenne du signal résultant est elle aussi très faible.

Le projet s'est inscrit dans le cadre de la réalisation d'un système émetteur-récepteur utilisant la technologie à bande ultra large permettant la communication sans fil entre les ports série de deux stations de travail. Le travail effectué pour ce projet de fin d'études a consisté en une recherche sur cette nouvelle technologie et sur les systèmes à spectre étalé ainsi qu'en la conception théorique d'un système émetteur-récepteur unidirectionnel à bande ultra large à l'aide d'un outil de simulation. Le modèle simplifié du système qui a été simulé fonctionne bien dans le contexte considéré.

Mots clés : Bande ultra large, *ultra-wideband (UWB)*, impulsions ultra brèves, ondes non sinusoïdales, transmissions à spectre étalé, *spread spectrum communications*, simulation, sans fil.

Summary

The project investigated the potential application of the new ultra-wideband (UWB) technology to the design of a secure communication system. This technology is based on the transmission of extremely short pulses, usually less than a nanosecond in duration, which yield broadband signals of very low power densities. During the summer, this technology along with the theory of spread spectrum systems and pseudonoise codes have been studied. Based on this theory, a model of an ultra-wideband transceiver for the transmission of information on a one-way serial link between two workstations has been developed using a simulation tool. This model functions well in the simulated environment.

Keywords : ultra-wideband (UWB), nonsinusoidal, spread spectrum communications, simulation, wireless.

Remerciements

J'aimerais d'abord remercier Monsieur Michel Lemire qui a accepté de prendre la direction de ce projet de fin d'études. Je désire également exprimer ma gratitude à Monsieur Larry Korba qui m'a accueillie dans son groupe, a consenti à devenir le codirecteur de mon projet et m'a proposé un sujet passionnant pour ce travail.

Je désire remercier les autres membres du *Groupe de réseautage* du CNRC, en particulier Messieurs Roger Impey et Randy Born, pour leur agréable compagnie durant mon été au sein du groupe. Ma reconnaissance va aussi à Mélanie LaBarre, Christian Dannewitz et พรรณณมอล เต็มดี (Punnarumol Temdee), mes amis du CNRC qui ont permis de rendre mon séjour à Ottawa plus agréable.

Enfin, je tiens à exprimer ma reconnaissance à Monsieur Michael Moher du *Centre de recherches sur les communications du Canada* pour son aide technique en particulier pour la section du rapport concernant le cas d'un système à plusieurs usagers.

Table des matières

LISTE DES TABLEAUX.....	VI
LISTE DES FIGURES.....	VII
LISTE DES SYMBOLES ET DES ACRONYMES.....	VIII
LEXIQUE FRANÇAIS-ANGLAIS.....	XII
INTRODUCTION	1
1. PROBLÉMATIQUE	3
2. ÉTUDE THÉORIQUE ET REVUE DE LA DOCUMENTATION	4
2.1 TECHNOLOGIE À BANDE ULTRA LARGE.....	4
2.1.1 <i>Bref historique de l'ULB</i>	4
2.1.2 <i>Définition générale de l'ULB et principes</i>	5
2.1.3 <i>Vue d'ensemble des travaux de deux compagnies</i>	7
2.1.4 <i>Avantages et inconvénients de l'ULB</i>	15
2.1.5 <i>Applications de l'ULB</i>	17
2.2 PRINCIPES DES SYSTÈMES DE TRANSMISSION À SPECTRE ÉTALÉ.....	18
2.2.1 <i>Description générale du principe d'étalement du spectre en séquence directe</i>	18
2.2.2 <i>Génération et choix des codes pseudo-aléatoires</i>	21
2.2.3 <i>Réception de signaux à étalement du spectre en séquence directe</i>	29
2.3 COMPARAISON ENTRE LES SIGNAUX À BANDE ULTRA LARGE ET LES SIGNAUX À SPECTRE ÉTALÉ	35
3. MÉTHODOLOGIE DE TRAVAIL ET RÉSULTATS	37
3.1 CHOIX DE L'OUTIL DE CONCEPTION	37
3.2 HYPOTHÈSES SIMPLIFICATRICES	39
3.3 INTERFACE AVEC LES PORTS SÉRIE	39
3.4 SIMULATION DE L'ÉMETTEUR-RÉCEPTEUR THÉORIQUE SUR <i>EXTEND</i>	42
3.4.1 <i>Principe général</i>	43
3.4.2 <i>Émetteur</i>	45

3.4.3 Modélisation du canal.....	47
3.4.4 Récepteur.....	48
3.4.5 Transmission de monocycles.....	54
4. DISCUSSION.....	57
4.1 LIMITES ET AMÉLIORATIONS DU SYSTÈME	57
4.2 CAS DE PLUSIEURS USAGERS.....	60
4.3 RECOMMANDATIONS	64
CONCLUSION	65
BIBLIOGRAPHIE	66
ANNEXE	69

Liste des tableaux

TABLEAU 2.1 : TABLE DE VÉRITÉ POUR LA SOMME MODULO-2	20
TABLEAU 2.2 : TABLE DE VÉRITÉ POUR LE PRODUIT.....	20

Liste des figures

FIGURE 2.1 : MONOCYCLE	9
FIGURE 2.2 : SÉRIE DE MONOCYCLES	9
FIGURE 2.3 : PRINCIPE DE MODULATION DE POSITION POUR L'INFORMATION.....	10
FIGURE 2.4 : SCHÉMA SIMPLIFIÉ D'UN ÉMETTEUR-RÉCEPTEUR UTILISANT LA MODULATION <i>TM-UWB</i>	13
FIGURE 2.5 : IMPULSION GAUSSIENNE	13
FIGURE 2.6 : DEUX IMPULSIONS FORMANT UN DOUBLET.....	14
FIGURE 2.7 : SÉQUENCE DE DOUBLETS REPRÉSENTANT <i>011</i>	14
FIGURE 2.8 : SCHÉMA GÉNÉRAL D'UN SYSTÈME À ÉTALEMENT DE SPECTRE EN SÉQUENCE DIRECTE.....	19
FIGURE 2.9 : EXEMPLE DE SIGNAUX $m(t)$, $c(t)$ ET $s(t)$	21
FIGURE 2.10 : EXEMPLE DE REGISTRE À DÉCALAGE COMPOSÉ DE 4 BASCULES D.....	22
FIGURE 2.11 : FORME D'ONDE DE LA SÉQUENCE <i>001101011110001</i> GÉNÉRÉE PAR LE REGISTRE À DÉCALAGE DE LA FIGURE 2.10.....	23
FIGURE 2.12 : FONCTION D'AUTOCORRÉLATION PÉRIODIQUE D'UNE SÉQUENCE MAXIMALE.....	26
FIGURE 2.13 : CONFIGURATION D'UN GÉNÉRATEUR DE CODES GOLD	27
FIGURE 2.14 : DÉMODULATEUR THÉORIQUE POUR TRANSMISSION EN BANDE DE BASE	30
FIGURE 2.15 : EXEMPLE DE FILTRE ADAPTÉ À LIGNE À RETARD	31
FIGURE 2.16 : PRINCIPE DU CORRÉLATEUR DYNAMIQUE	32
FIGURE 2.17 : BOUCLE À RETARD DE PHASE.....	33
FIGURE 2.18 : BOUCLE <i>TAU-DITHER LOOP</i>	34
FIGURE 2.19 : SCHÉMA D'UN RÉCEPTEUR POUR UN SYSTÈME À SPECTRE ÉTALÉ CONVENTIONNEL.....	35
FIGURE 3.1 : INTERFACE DU LOGICIEL DE SIMULATION <i>EXTEND</i>	38
FIGURE 3.2 : SCHÉMA GÉNÉRAL DU SYSTÈME ÉMETTEUR-RÉCEPTEUR	40
FIGURE 3.3 : FORMAT DU SIGNAL À L'ENTRÉE DE L'ÉMETTEUR-RÉCEPTEUR	41
FIGURE 3.4 : PRINCIPE DE TRANSMISSION DE MONOCYCLES POUR CHAQUE CHIP DE CODE <i>PN</i>	44
FIGURE 3.5 : SIGNAUX À L'ÉMETTEUR	47
FIGURE 3.6 : SIGNAUX PERMETTANT DE MIEUX COMPRENDRE LE MÉCANISME DE DÉMODULATION	50
FIGURE 3.7 : SIGNAUX DU MODULE ACQUISITION PERMETTANT DE MIEUX SAISIR LE PROCESSUS D'INTÉGRATION	51
FIGURE 3.8 : SIGNAUX DU MODULE ACQUISITION PERMETTANT DE COMPRENDRE LE PRINCIPE DU CORRÉLATEUR DYNAMIQUE SIMULÉ.....	53
FIGURE 3.9 : SIGNAUX D'INTÉRÊT POUR LE SYSTÈME SIMULÉ ENTIER	54
FIGURE 4.1 : GÉNÉRATION D'UNE FAMILLE DE CODES GOLD ÉQUILIBRÉS DE PÉRIODE $2^{13} - 1$	61
FIGURE 4.2 : DENSITÉS DE PROBABILITÉ ASSOCIÉES AUX PROBABILITÉS DE FAUSSE ALARME ET DE DÉTECTION.....	63

Liste des symboles et des acronymes

δ :	décalage dans le temps pour la modulation de position de l'information dans le système de <i>Time Domain</i> (s)
γ :	décalage dans le temps dans la boucle à retard de phase (s)
$\varphi(\tau)$:	fonction d'autocorrélation continue pour un décalage dans le temps τ
φ_τ :	fonction d'autocorrélation discrète (dans le cas de codes pseudo-aléatoires) pour un décalage de τ chips
κ :	décalage dans le temps entre les signaux $r_1(t)$ et $r_2(t)$ dans l'explication du cas de 2 usagers dans le système (s)
τ :	décalage dans le temps entre deux signaux dans l'expression de la corrélation continue (s) ou décalage entre deux codes pseudo-aléatoires dans l'expression de la corrélation discrète (chips)
τ_d :	temps appelé en anglais <i>dwell time</i> pendant lequel s'effectue la corrélation (intégration) dans le module acquisition d'un récepteur (s)
τ_p :	délai du signal reçu $r(t)$ par rapport au code pseudo-aléatoire généré localement au récepteur (s)
τ_r :	décalage dans le temps permettant la mise à jour de la phase du code pseudo-aléatoire dans le module acquisition d'un récepteur (s)
$a(t)$:	autocorrélation du signal de l'utilisateur 1 dans l'explication du cas de 2 usagers dans le système
A :	amplitude de crête d'un monocycle (V)
AMRC :	accès multiple par répartition de code
BER :	<i>bit error rate</i>
c_i :	valeur de la $i^{\text{ème}}$ chip d'un code pseudo-aléatoire
c_j :	élément du code pseudo-aléatoire propre à un récepteur particulier correspondant au $j^{\text{ème}}$ monocycle dans le système de <i>Time Domain</i>
$\{c_n\}$:	ensemble des chips d'un code pseudo-aléatoire dans une période (une séquence)
$c(t)$:	signal représentant le code pseudo-aléatoire
CDMA :	<i>code division multiple access</i>

CNRC :	Conseil national de recherches du Canada
d :	information à transmettre (I ou 0) dans le système de <i>Time Domain</i>
<i>DSP</i> :	<i>digital signal processor</i>
f_c :	fréquence centrale du signal (Hz)
f_H :	fréquence supérieure de la bande de fréquences du signal (Hz)
f_{if} :	fréquence intermédiaire d'un récepteur superhétérodyne conventionnel (Hz)
f_L :	fréquence inférieure de la bande de fréquences du signal (Hz)
$f(t)$:	signal à corrélérer servant d'exemple pour l'explication de la corrélation continue
<i>FCC</i> :	<i>Federal Communications Commission</i>
$g(t)$:	deuxième signal à corrélérer servant d'exemple pour l'explication de la corrélation croisée continue
<i>GPS</i> :	<i>global positioning system</i>
i :	indice d'une chip de code pseudo-aléatoire
$i(t)$:	intercorrélation du signal de l'utilisateur 1 avec le signal de l'utilisateur 2 décalé de κ dans l'explication du cas de 2 usagers dans le système
<i>IEEE</i> :	<i>Institute of Electrical and Electronics Engineers, Inc.</i>
j :	indice d'un monocycle
k :	indice d'un émetteur d'un système d'émetteurs-récepteurs
L :	période d'un code pseudo-aléatoire (chips)
LB_R :	largeur de bande relative
L_c :	nombre de chips par bit d'information et gain résultant du traitement
$m(t)$:	signal représentant l'information à transmettre
$m_o(t)$:	expression d'un monocycle définie par l'équation 2.3
$\max\{a(t)\}$:	valeur du pic principal d'autocorrélation de $r_1(t)$ dans l'explication du cas de 2 usagers dans le système
$\max_2\{a(t)\}$:	valeur du pic secondaire d'autocorrélation de $r_1(t)$ dans l'explication du cas de 2 usagers dans le système
$\max\{i(t)\}$:	valeur maximale, en valeur absolue, de la corrélation croisée entre $r_1(t)$ et $r_2(t)$ pour tous les décalages τ dans l'explication du cas de 2 usagers dans le système
n :	nombre de bascules d'un registre à décalage
$n(t)$:	signal bruit
N_h :	valeur maximale que peut prendre l'élément de code pseudo-aléatoire c_j dans le système de <i>Time Domain</i>
N_s :	nombre de monocycles par bit d'information

p :	longueur d'une plage de code pseudo-aléatoire (chips)
P :	niveau de puissance de crête d'un signal à bande ultra large permis au-dessus de la limite moyenne d'émission admise par la partie 15 des règles de la commission (dB)
$P()$:	probabilité
P_D :	probabilité de détection
P_{FA} :	probabilité de fausse alarme
PN :	<i>pseudonoise</i> ou pseudo-aléatoire
$r(t)$:	signal reçu au récepteur
$r_1(t)$:	signal reçu au récepteur correspondant à l'utilisateur 1 dans l'explication du cas de 2 usagers dans le système
$r_2(t)$:	signal reçu au récepteur correspondant à l'utilisateur 2 dans l'explication du cas de 2 usagers dans le système
R_b :	débit binaire de l'information (bits/s ou bps)
R_c :	débit du code pseudo-aléatoire (chips/s ou cps)
RF :	radiofréquence
$s(t)$:	signal transmis dans le canal
S :	seuil du comparateur du module acquisition
t_0 :	temps séparant deux impulsions formant un doublet dans le système de <i>Æther Wire</i> (s)
t_1 :	temps entre deux doublets dans le système de <i>Æther Wire</i> (s)
$t(n)$:	expression permettant le calcul des 3 valeurs possibles de l'intercorrélation des codes Gold
T_1 à T_7 :	éléments d'une ligne à retard correspondant à un retard d'une chip chacun
T_b :	durée d'un bit d'information (s)
T_c :	durée d'une chip du code pseudo-aléatoire (s)
T_f :	période de répétition des monocycles (s)
T_m :	durée d'un monocycle (s)
T_p :	constante multipliant l'élément de code pseudo-aléatoire c_j pour obtenir le décalage dans le temps pour la deuxième modulation de position dans le système de <i>Time Domain</i> (s)

u :	deuxième code pseudo-aléatoire servant d'exemple pour l'explication de la corrélation croisée discrète
<i>UART :</i>	<i>universal asynchronous receiver transmitter</i>
UIT :	Union internationale des télécommunications
ULB :	technologie à bande ultra large
<i>UWB :</i>	<i>ultra-wideband</i>
<i>VCO :</i>	<i>voltage controlled oscillator</i>

Lexique français-anglais

Accès multiple par répartition de code :	<i>code division multiple access (CDMA)</i>
Avis d'information publique :	<i>Notice of Inquiry</i>
Avis de proposition de réglementation :	<i>Notice of Proposed Rule Making</i>
Boucle à retard de phase :	<i>delay-locked loop</i>
Brouillage intentionnel :	<i>jamming</i>
Carnet :	<i>notebook</i>
Circuit de pont d'échantillonnage :	<i>sampling bridge circuit</i>
Code de pseudo-bruit :	<i>pseudonoise code ou pseudonoise sequence</i>
Code pseudo-aléatoire :	<i>pseudorandom code ou pseudorandom sequence ou PN code</i>
Corrélateur dynamique (corrélation dynamique) :	<i>sliding correlator (sliding correlation)</i>
Détection de fronts :	<i>leading edge detection</i>
Distorsion en fréquence :	<i>frequency distortion</i>
Émetteur-récepteur asynchrone universel :	<i>universal asynchronous receiver transmitter (UART)</i>
Entrelacement :	<i>interleaving</i>
Équilibré (séquence équilibrée) :	<i>balanced (balanced sequence)</i>
Gain résultant du traitement :	<i>processing gain</i>
Largeur de bande relative :	<i>fractional bandwidth</i>
Ligne à retard à prises :	<i>tapped delay line</i>
Oscillateur contrôlé en tension :	<i>voltage controlled oscillator (VCO)</i>
Outil de création de graphiques :	<i>plotter</i>
Plage :	<i>run</i>
Poursuite :	<i>tracking</i>
Prise de contact :	<i>handshaking</i>
Prise :	<i>tap</i>
Probabilité de détection :	<i>detection probability</i>
Probabilité de fausse alarme :	<i>false alarm probability</i>
Processeur de signaux numériques :	<i>digital signal processor (DSP)</i>

Récepteur en râteau :	<i>rake receiver</i>
Registre à décalage :	<i>shift register</i>
Sauts de fréquence :	<i>frequency hopping</i>
Séquence directe :	<i>direct sequence</i>
Séquence maximale :	<i>maximal length sequence ou maximal sequence ou m-sequence</i>
Système de transmission à spectre étalé :	<i>spread spectrum system</i>
Système mondial de radiorepérage :	<i>global positioning system</i>
Taux d'erreur binaire :	<i>bit error rate (BER)</i>
Technologie à bande ultra large :	<i>ultra-wideband (UWB) technology</i>

Introduction

Les télécommunications sont présentement en plein essor. De nouveaux protocoles et standards de transmission, algorithmes de codage de données, puces, etc., sont mis au point à un rythme soutenu. Ces développements en matière de communication se basent, pour la plupart, sur la transmission d'ondes continues souvent sinusoïdales. Le projet présenté ici porte plutôt sur l'émission d'impulsions de très courte durée, impulsions qui permettent également de transmettre de l'information. Cette technologie, appelée *technologie à bande ultra large*, offre des avantages intéressants par rapport aux techniques actuelles de transmission d'information et pourrait révolutionner le domaine des télécommunications d'ici quelques années.

Le travail présenté dans ce rapport a été effectué pour le *Groupe de réseautage du Conseil national de recherches du Canada*. Il s'agit d'amorcer la conception d'un émetteur-récepteur utilisant la technologie à bande ultra large pour transmettre des informations entre les ports série de deux ordinateurs. Le présent document décrit le processus qui a mené à la conception de ce système.

Le premier chapitre énonce simplement le problème. Il permet de cerner le but du travail et les moyens entrepris pour y arriver. Dans le chapitre deux, on présente l'étude théorique de la documentation consultée sur la technologie à bande ultra large et sur les systèmes à spectre étalé. Le lecteur qui est déjà familier avec la théorie associée à ces technologies peut laisser tomber le deuxième chapitre et s'y référer au besoin. On expose ensuite, dans le troisième chapitre, la méthode employée pour concevoir l'émetteur-récepteur ainsi que certains des résultats obtenus. Finalement, dans le dernier chapitre, on discute des limites du modèle et des améliorations possibles pouvant y être apportées, on traite du cas d'un système à plusieurs usagers et on formule des recommandations destinées aux membres du *Groupe de réseautage* pour la poursuite des travaux.

Il est à noter que la plupart des traductions de mots techniques de ce rapport ont été obtenues grâce à la base de données terminologique de l'Union internationale des télécommunications (UIT) [1]. Dans tout le document, les mots anglais correspondant sont souvent inscrits, lors de leur première apparition, entre parenthèses après la traduction parce que ces mots anglais sont généralement mieux connus dans le domaine considéré. Par la suite, le lecteur pourra se référer au *Lexique français-anglais* pour trouver l'équivalent anglais d'un mot français traduit.

1. Problématique

Le projet de fin d'études présenté a été réalisé au sein du *Groupe de réseautage* de l'*Institut de technologie de l'information* du *Conseil national de recherches du Canada* (CNRC) à Ottawa. Les recherches menées par ce groupe de travail portent entre autres sur les communications personnelles et sur la sécurité des transmissions. Dans cette perspective, l'utilisation de la « nouvelle » technologie à bande ultra large (*ultra-wideband* ou *UWB*), qui repose sur la transmission d'impulsions ultra brèves et non sinusoïdales, semble une avenue intéressante pouvant permettre des communications sécurisées sur de courtes distances et pourrait avantageusement concurrencer la technologie *Bluetooth*.

Plus spécifiquement, le *Groupe de réseautage* était intéressé à évaluer la possibilité de réaliser un système émetteur-récepteur utilisant la technologie à bande ultra large et permettant la communication sans fil entre les ports série de deux stations de travail. Cette technologie pourrait éventuellement être implantée au sein du groupe et permettre le remplacement de certains câbles reliant les ports série de différents ordinateurs par des liens radiofréquences (RF). Dans ce contexte, le travail demandé consistait d'abord en une recherche sur la technologie à bande ultra large et en l'étude des différents documents disponibles s'y rapportant. Afin d'être capable de concevoir un système émetteur-récepteur utilisant cette technologie et de permettre la communication simultanée de plusieurs paires de stations de travail, les principes des systèmes de transmission à spectre étalé (*spread spectrum systems*) incluant l'accès multiple par répartition de code (AMRC et en anglais *code division multiple access* ou *CDMA*) devaient ensuite être étudiés. Enfin, une conception théorique et générale du système émetteur-récepteur en question allait devoir être entreprise.

2. Étude théorique et revue de la documentation

Dans ce chapitre, la recherche réalisée sur la nouvelle technologie à bande ultra large et sur les systèmes de transmission à spectre étalé est présentée dans ses grandes lignes. La dernière section établit une comparaison entre ces deux modes de transmission de l'information.

2.1 Technologie à bande ultra large

La présente section donne un aperçu des connaissances acquises durant la recherche effectuée sur la technologie à bande ultra large. On expose d'abord un court historique de cette technologie et on en donne ensuite une description. On décrit également de façon brève les travaux de deux compagnies. On énumère enfin les principaux avantages, inconvénients et applications de la technologie à bande ultra large.

La technologie à bande ultra large est surtout connue sous le nom de *ultra-wideband technology*, mais les termes *impulse radio*, *impulse*, *nonsinusoidal*, *baseband*, *carrierless*, *carrier-free*, *time domain* et *super wideband* peuvent, dans la littérature, se substituer à *ultra-wideband* [2] [3] [4]. Étant donné que cette technologie est relativement récente, nous n'avons pu trouver de documents rédigés en français sur ce sujet. Les termes employés dans ce rapport concernant cette technologie sont donc des traductions personnelles des termes anglais mieux connus. De plus, à partir d'ici et dans tout le document, l'acronyme *ULB* remplacera l'expression *technologie à bande ultra large* ou à *ultra large bande* afin d'éviter d'alourdir le texte inutilement.

2.1.1 Bref historique de l'ULB

Quoique plusieurs scientifiques aient travaillé sur la génération d'impulsions ultra brèves depuis les années 1960, le mot *ultra-wideband* ou *UWB* n'est apparu que vers 1989 et l'utilisation de l'ULB par le gouvernement et l'armée a commencé vers 1970. Les pionniers de cette technologie sont Gerald F. Ross, Kenneth W. Robbins, Henning F. Harmuth, Rexford M. Morey, et Paul van Etten. Plus de 200 articles ont été publiés sur le sujet dans divers journaux de l'*Institute of Electrical and Electronics Engineers, Inc. (IEEE)* dont l'article fondamental de cette technologie [5] et une centaine

de brevets concernant l'ULB ont été déposés [2]. Les références de la plupart de ces documents sont disponibles sur le site [6].

La *Federal Communications Commission (FCC)* américaine n'a pas encore approuvé l'utilisation commerciale de l'ULB. Elle a d'abord fait paraître un avis d'information publique (*Notice of Inquiry*) [7] en août 1998 afin d'évaluer la possibilité de permettre l'utilisation de systèmes employant l'ULB. À la suite de cette publication, elle a reçu près de 150 réponses et commentaires provenant de divers organismes et entreprises impliqués de près ou de loin dans l'utilisation de l'ULB. En mai 2000, à partir de ces réponses, la *FCC* a adopté un avis de proposition de réglementation (*Notice of Proposed Rule Making*) [8] dans lequel elle reconnaît les avantages que pourraient apporter les systèmes utilisant l'ULB dans de nombreux domaines. Elle y propose de commencer à identifier des amendements possibles à la partie 15 des règles de la commission afin de permettre l'utilisation de dispositifs utilisant l'ULB. Elle entend toutefois s'assurer, grâce à de nombreux tests et analyses, que l'utilisation de systèmes recourant à l'ULB n'interférera pas avec des dispositifs déjà existants, tels le système mondial de radiorepérage (*global positioning system* ou *GPS*) et les systèmes à bord des avions. En 2001, les amendements finaux à la partie 15 des règles de la commission concernant l'ULB devraient être apportés.

2.1.2 Définition générale de l'ULB et principes

Comme l'ULB n'est pas encore approuvée, il n'existe pas de définition unique de cette technologie ni de standard. Différentes entreprises travaillent sur le développement de dispositifs utilisant l'ULB et ont chacune leur propre façon de l'exploiter selon l'application à laquelle le dispositif fabriqué est destiné. Une définition générale de l'ULB est donnée ici. On présentera, dans la prochaine section, une explication des particularités techniques propres à deux compagnies.

La technologie à bande ultra large se base sur la transmission et la réception d'ondes non sinusoïdales. Les formes de ces ondes sont des impulsions ultra brèves (d'une durée de l'ordre de la nanoseconde) et couvrent donc fondamentalement une très grande partie du spectre des fréquences. On peut croire que les dispositifs utilisant l'ULB standardisée auront une fréquence centrale entre 20 MHz et 60 GHz et une largeur de

bande se situant entre 150 MHz et 30 GHz, et pourront fonctionner sur une distance de 1 cm à 30 m [8]. La forme des ondes émises par les appareils utilisant l'ULB n'est pas toujours la même d'une organisation à l'autre même si les ondes sont toutes ultra brèves. Il est donc impossible d'en donner une équation générale ici. De plus, la forme de ces ondes et son spectre sont affectés par l'antenne qui émet le signal, car celle-ci agit comme un filtre passe-bas [8].

Selon Taylor [3], le terme *bande ultra large* désigne les systèmes qui transmettent et reçoivent des ondes dont la largeur de bande relative LB_R (*fractional bandwidth*) est supérieure ou égale à 0,25. La largeur de bande relative LB_R est définie de la façon suivante :

$$LB_R = \frac{f_H - f_L}{f_c} \text{ où } f_c = \frac{f_H + f_L}{2}, \quad (2.1)$$

avec f_c qui est la fréquence centrale du signal, f_H , la fréquence supérieure de la bande de fréquences du signal et f_L , la fréquence inférieure de cette bande. La définition donnée par Taylor revient donc à dire que la largeur de bande $f_H - f_L$ doit être supérieure ou égale à 25 % de la fréquence centrale pour qu'on puisse parler de bande ultra large. La FCC abonde dans ce sens dans son avis de proposition de réglementation [8] en proposant, pour approbation, cette même définition, mais en y ajoutant deux éléments :

- la largeur de bande doit être celle mesurée à -10 dB;
- un signal peut aussi être considéré à bande ultra large s'il occupe plus de 1,5 GHz du spectre des fréquences même s'il ne satisfait pas à la première définition.

Ainsi, par exemple, pour une fréquence centrale f_c de 2 GHz, la largeur de bande minimale à -10 dB requise pour qu'on puisse parler de bande ultra large, selon cette définition qui risque d'être adoptée, est 500 MHz.

Les puissances moyennes associées aux signaux à bande ultra large sont en général très faibles parce que le rapport cyclique, qui est la largeur de l'impulsion en unités de temps sur la période de répétition des impulsions, est aussi très faible. Dans [8], on indique qu'un signal à bande ultra large a une puissance moyenne ne dépassant généralement pas 10 mW et une puissance de crête de 10 W au maximum. La FCC

propose, toujours dans [8], de limiter la puissance de crête émise par un signal à bande ultra large. En effet, selon elle, cette puissance ne devrait jamais dépasser un certain niveau P au-dessus de la limite moyenne d'émission permise par la partie 15 des règles de la commission. La formule soumise pour calculer ce niveau P en dB sur toute la largeur de bande est la suivante :

$$P = 20 + 20\log_{10}\left[\frac{(\text{largeur de bande à } -10\text{dB du signal en hertz})}{50 \times 10^6 \text{ Hz}}\right]. \quad (2.2)$$

Pour une largeur de bande de 1 GHz à -10 dB, par exemple, la puissance de crête émise ne devrait pas excéder la limite moyenne d'émission permise par la partie 15 des règles par plus de 46 dB. De plus, la *FCC* propose que la valeur maximale de P soit fixée à 60 dB.

Les techniques de génération des impulsions de l'ordre de la nanoseconde caractérisant l'ULB sont nombreuses. Ces impulsions ultra brèves peuvent être produites par des commutateurs ultra rapides, tels des transistors avalanche, des commutateurs activés par la lumière ou des diodes tunnels [2] [3] [9]. Elles sont ensuite transmises et reçues grâce à des antennes qui peuvent émettre et capter des signaux à très large bande. Plusieurs documents tels [3], [10], [11], [12] et [13] traitent de ce sujet. Enfin, la détection des impulsions envoyées et la réception peuvent être réalisées grâce à différentes méthodes, telles la détection de fronts (*leading edge detection*), l'utilisation d'un multivibrateur monostable, l'utilisation d'un circuit de pont d'échantillonnage (*sampling bridge circuit*), l'intégration et l'échantillonnage, l'utilisation d'un filtre adapté et l'utilisation d'un corrélateur [9]. Certaines méthodes de réception seront abordées à la section 2.2 *Principes des systèmes de transmission à spectre étalé*.

2.1.3 Vue d'ensemble des travaux de deux compagnies

Les principales entreprises développant des dispositifs utilisant l'ULB sont *Æther Wire & Locations, Inc.*, *ANRO Engineering, Inc.*, *Fantasma Networks, Inc.*, *Lawrence Livermore National Laboratory*, *McEwan Technologies*, *Multispectral Solutions, Inc.*, *Time Domain Corporation* et *XtremeSpectrum, Inc.* Nous présentons ici, dans ses grandes lignes, la théorie associée aux dispositifs conçus par *Time Domain Corporation* et *Æther Wire & Location, Inc.* En effet, notre étude a porté sur leurs

travaux en particulier étant donné que de nombreux articles techniques et brevets sont disponibles sur la façon dont ils implémentent l'ULB, surtout pour *Time Domain*. La source principale des informations sur les travaux de ces deux compagnies consignées dans les prochains paragraphes sont les documents [4], [14], [15], [16], [17] et [18] pour *Time Domain* et le document [19] pour *Æther Wire & Location, Inc.*

Time Domain Corporation est une entreprise américaine qui a développé une puce baptisée *PulsON* basée sur ce qu'ils appellent *Time Modulated Ultra Wideband (TM-UWB) architecture*. Leur puce, qui peut être utilisée dans des applications radar, de localisation et en communications sans fil, transmet un train d'impulsions dans lequel chaque impulsion est émise à un moment précis. Ces impulsions sont appelées *monocycles*. La *figure 2.1* montre un monocycle ayant une largeur de 0,5 ns. Cette impulsion centrée à $t = 0$ s a comme équation l'équivalent de la dérivée d'une gaussienne :

$$m_o(t) = -6A \sqrt{\frac{e\pi}{3}} \frac{t}{T_m} e^{-6\pi \left(\frac{t}{T_m} \right)^2}, \quad (2.3)$$

où A est l'amplitude de crête du monocycle en volts, t , le temps et T_m , la durée de l'impulsion en secondes. La valeur de T_m pour le monocycle de la *figure 2.1* est 0,5 ns. On détermine la fréquence centrale f_c du monocycle en faisant l'inverse de sa durée :

$$f_c = \frac{1}{T_m}. \quad (2.4)$$

La largeur de bande à mi-puissance vaut autour de 116 % de cette fréquence centrale. Ainsi, pour le monocycle de la *figure 2.1*, la largeur de bande à mi-puissance est d'environ 2,3 GHz. Les impulsions générées par *Time Domain* ont des durées qui se situent entre 0,20 ns et 1,5 ns. Évidemment, l'émission d'un monocycle gaussien parfait est impossible en pratique. L'équation 2.3 est donc théorique.

Afin de permettre la communication, *Time Domain* génère une série d'impulsions, plutôt qu'un seul monocycle, d'abord espacées de T_f secondes. La *figure 2.2* montre une séquence de monocycles. Cette figure n'est pas à l'échelle. En effet, la période T_f est en réalité beaucoup plus grande que la durée T_m du monocycle. Ainsi, comme mentionné

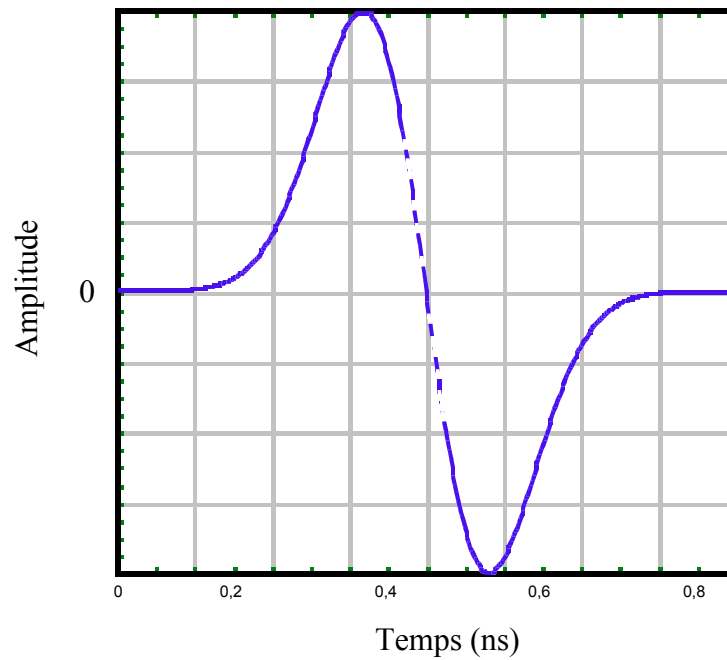


Figure 2.1 : Monocycle
(généré à partir du logiciel de simulation *Extend*)

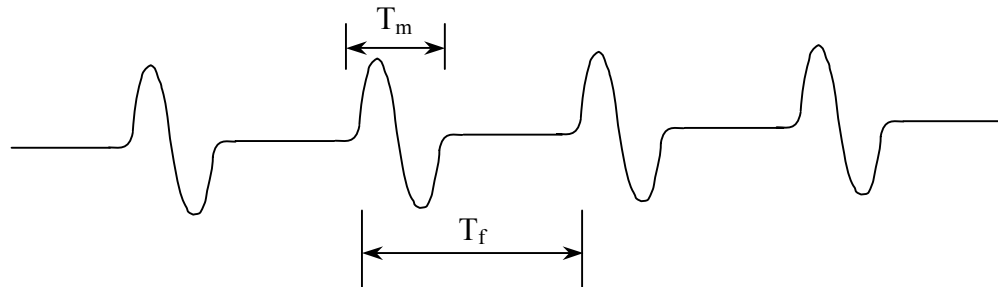


Figure 2.2 : Série de monocycles

précédemment, on obtient un rapport cyclique $\frac{T_m}{T_f}$ très faible, ce qui permet d'avoir un

signal dont la puissance moyenne est également très faible. Par exemple, pour un monocycle dont la puissance est de 14,5 W, la durée T_m , 0,5 ns et le taux de répétition, $1,5 \times 10^6$ monocycles par secondes, on obtient une période T_f de 666,7 ns

($\frac{1}{1,5 \times 10^6 \text{ monocycles/s}} = 666,7 \text{ ns}$), un rapport cyclique de 0,075 %, une puissance

moyenne de 10,9 mW ($\frac{T_m}{T_f} \times \text{puissance du monocycle} = \frac{0,5 \text{ ns}}{666,7 \text{ ns}} \times 14,5 \text{ W} = 10,9 \text{ mW}$) et

une densité de puissance moyenne très faible de 4,7 pW/Hz

($\frac{\text{puissance moyenne}}{\text{largeur de bande}} = \frac{10,9 \text{ mW}}{2,3 \text{ GHz}} = 4,7 \text{ pW/Hz}$) [20]. *Time Domain* produit des séquences

de monocycles dont le rapport cyclique est moins de 1 %. Selon *Time Domain*, la transmission d'impulsions uniformément espacées possède deux inconvénients : la création de raies spectrales dans le domaine des fréquences qui peuvent interférer avec les systèmes radio existants à courte distance et la possibilité de collisions entre des signaux provenant de différents usagers. Ces problèmes sont vraisemblablement éliminés avec la modulation de position employée par *Time Domain*.

Time Domain utilise deux modulations de position : pour l'information et pour le code pseudo-aléatoire. Les deux modulations sont abordées ici. La modulation de position est premièrement utilisée afin de transmettre de l'information. Cette technique de modulation permet l'utilisation d'un récepteur optimal à filtre adapté. Chaque élément binaire¹ d'information est transmis avec un décalage dans le temps selon que le bit est un 0 ou un 1. Par exemple, pour un 0, on peut appliquer un décalage négatif dans le temps de δ (une centaine de picosecondes) par rapport à une référence (située à T_f), alors que, pour un 1, on appliquerait un décalage positif du même δ par rapport à la référence. La *figure 2.3* illustre ce principe. La modulation de position pour l'information permet

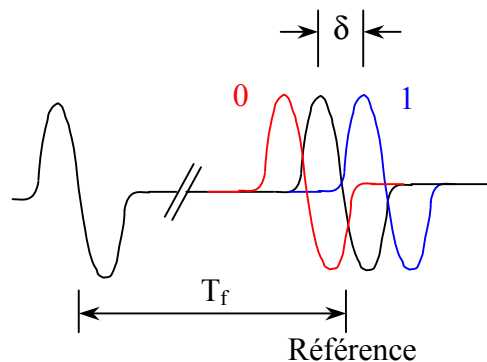


Figure 2.3 : Principe de modulation de position pour l'information

¹ On utilisera, à partir de maintenant, le mot *bit* pour désigner un *élément binaire*.

d'obtenir une densité spectrale de puissance un peu plus uniforme sur toute la bande de fréquences. Contrairement à ce que peut laisser penser la *figure 2.3*, chaque bit d'information est représenté par plus d'un monocycle (N_s monocycles). *Time Domain* semble utiliser plus de 200 monocycles par bit à transmettre. Ainsi, ce sont plusieurs monocycles qui sont décalés dans un sens ou dans l'autre pour chaque bit d'information. Cette technique permet de récupérer l'information envoyée même si certaines impulsions ne sont pas reçues.

Deuxièmement, en plus de la modulation de position pour l'information, le signal émis est codé grâce à un code pseudo-aléatoire de manière à séparer les différents usagers du système², à uniformiser l'énergie dans le domaine des fréquences et à diminuer l'influence du brouillage intentionnel (*jamming*). Chaque usager (système émetteur-récepteur) possède un code pseudo-aléatoire qui lui est propre. Ce code permet une deuxième modulation de position des monocycles. Pour un émetteur particulier, chaque impulsion d'indice j est ainsi décalée dans le temps de $c_j T_p$, où c_j représente l'élément du code pseudo-aléatoire périodique propre à ce récepteur correspondant au $j^{\text{ième}}$ monocycle et T_p , une constante en secondes. Si on a $0 \leq c_j < N_h$, où N_h est la valeur maximale que peut prendre l'élément de code c_j , avec T_f séparant chaque monocycle, on doit s'assurer que $N_h T_p \leq T_f$. Le décalage dans le temps dû au code pseudo-aléatoire est beaucoup plus grand (de l'ordre de plusieurs nanosecondes) que le décalage dans le temps dû à l'information, ce qui permet d'utiliser un corrélateur au récepteur.

Ainsi, le signal $s(t)$ transmis par le $k^{\text{ième}}$ émetteur d'un système d'émetteurs-récepteurs serait de la forme :

$$s^{(k)}(t) = \sum_j m_o(t - jT_f - c_j^{(k)}T_p - \delta d_{\lfloor j/N_s \rfloor}^{(k)}) \quad (2.5)$$

Dans cette équation, $m_o(t)$ représente le monocycle décrit par l'équation 2.3, j est l'indice du monocycle (le $j^{\text{ième}}$ monocycle), T_f est la période de répétition des monocycles, $c_j^{(k)}T_p$ est le temps de décalage selon le code pseudo-aléatoire du $k^{\text{ième}}$ usager et pour

² La notion de code pseudo-aléatoire et de systèmes à plusieurs usagers sera présentée plus en détail à la section 2.2 *Principes des systèmes de transmission à spectre étalé*.

le $j^{\text{ième}}$ monocycle et δ est le décalage dû à l'information. Enfin, d est l'information du $k^{\text{ième}}$ usager et vaut 1 si le bit envoyé est un 1 et -1 si le bit envoyé est 0 . Comme N_s est le nombre de monocycles par bit d'information, $[j/N_s]$ représente l'indice du bit d'information ($[j/N_s]$ implique qu'on prend la partie entière de j/N_s). Ainsi, le dernier terme de l'équation 2.5 tient bien compte du fait que la donnée ne change qu'après un certain nombre (multiple de N_s) de monocycles. Chaque bit d'information a donc une durée $T_b = N_s T_f$. Le débit binaire R_b de l'information résultant est donné par :

$$R_b = \frac{1}{T_b} = \frac{1}{N_s T_f} \text{ bits/s.} \quad (2.6)$$

À partir de cette équation, pour une période de répétition des monocycles T_f et un débit binaire R_b donnés, on peut trouver le nombre de monocycles N_s requis par bit d'information et obtenir le rapport cyclique correspondant.

Une fois le signal décrit ci-dessus généré et envoyé, le récepteur doit retrouver l'information transmise. Pour ce faire, *Time Domain* génère localement au récepteur un signal gabarit qui ressemble au signal transmis. Après avoir corrélé ce signal avec le signal reçu et s'être assuré de la synchronisation de ces deux signaux, on peut démoduler l'information contenue dans le signal reçu. On trouve de plus amples détails sur le fonctionnement d'un récepteur à corrélation dans la section 2.2 *Principes des systèmes de transmission à spectre étalé*. Un schéma extrêmement simplifié d'un émetteur-récepteur utilisant la modulation *TM-UWB* de *Time Domain* est présenté à la figure 2.4.

La compagnie ***Æther Wire & Location, Inc.*** travaille sur un projet financé par la *Defense Advanced Research Projects Agency* dont le but est le développement de dispositifs de la grosseur d'une pièce de monnaie appelés *localizers* qui permettent de localiser des objets avec une précision de quelques centimètres sur des distances de plusieurs kilomètres. En effet, comme la précision dans la détermination de distances est proportionnelle à la largeur de bande du signal utilisé, l'ULB permet d'obtenir d'excellentes résolutions. La localisation d'objets s'appuie sur l'utilisation d'un réseau de *localizers* (émetteurs-récepteurs) qui s'échangent de l'information sur les distances qui les séparent afin de trouver la position d'un d'entre eux par rapport à une unité de commande.

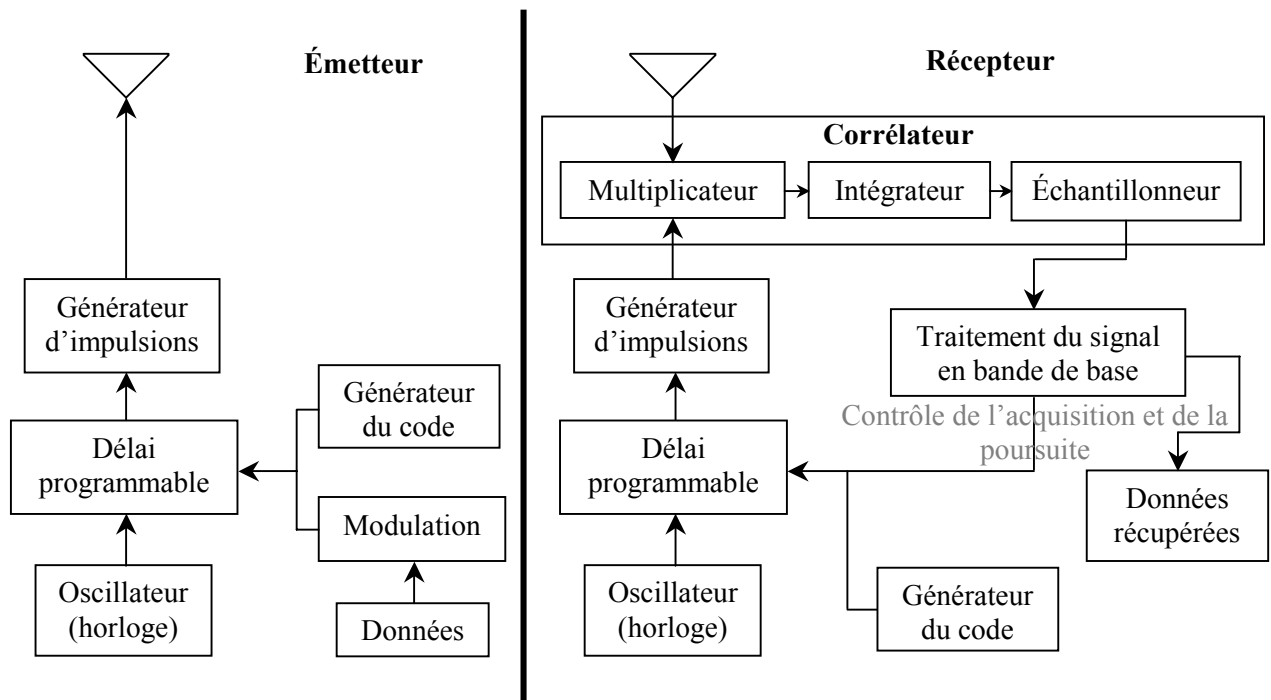


Figure 2.4 : Schéma simplifié d'un émetteur-récepteur utilisant la modulation *TM-UWB* (adaptation d'une figure de [14])

Les signaux utilisés par *Æther Wire* ne sont pas exactement les mêmes que ceux de *Time Domain*. Leur impulsion de base, qu'ils appellent *impulsion gaussienne*, est montrée à la *figure 2.5*. Afin d'avoir une impulsion équilibrée pour diminuer la

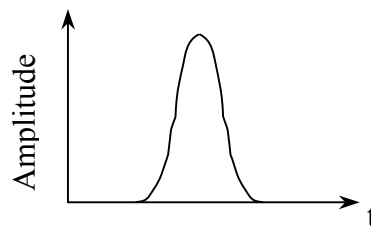


Figure 2.5 : Impulsion gaussienne

consommation d'énergie, ils envoient des *doublets* d'impulsions. Un doublet est présenté à la *figure 2.6* où t_0 correspond au temps séparant les deux impulsions constituant le doublet. L'utilisation de doublets permet également de manipuler le spectre de fréquences correspondant en jouant sur t_0 et sur la distance entre les doublets t_1 .

Contrairement à *Time Domain*, *Æther Wire* utilise la modulation antipodale pour transmettre l'information. Pour envoyer un I , un doublet commençant par un pic positif

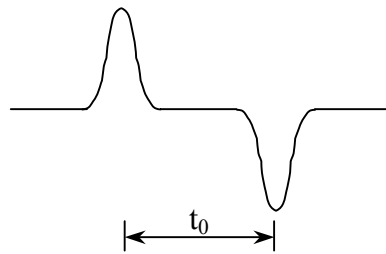


Figure 2.6 : Deux impulsions formant un doublet

est émis et l'inverse pour un 0. La *figure 2.7* illustre ce principe. Évidemment, en plus de la modulation de l'information, un code pseudo-aléatoire est appliqué sur le signal afin de séparer les différents *localizers* fonctionnant au même moment. Ce code module les

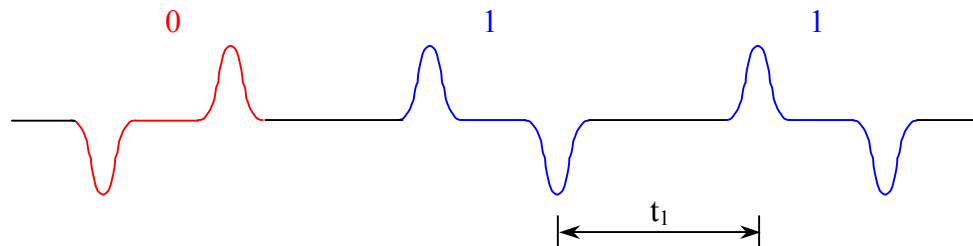


Figure 2.7 : Séquence de doublets représentant 011

doublets de façon antipodale comme pour l'information. Pour récupérer les données transmises, *Æther Wire* utilise un corrélateur dynamique³ (*sliding correlator*) appelé *Time-Integrator Correlator* basé sur l'article [21]. Le récepteur détectera un pic de corrélation positif ou négatif à la sortie du corrélateur selon la donnée envoyée. Dans le système de cette compagnie, le processus de synchronisation est assez complexe. L'explication de cette procédure sort du cadre du présent projet. Le lecteur qui désire obtenir de plus amples renseignements à ce sujet peut consulter le document [19].

³ La section 2.2 *Principes des systèmes de transmission à spectre étalé* traite de la notion de corrélateur dynamique de façon plus détaillée.

2.1.4 Avantages et inconvénients de l'ULB

Quoique chaque entreprise traite l'ULB d'une façon différente, toutes les organisations s'entendent sur les nombreux avantages qu'offre cette nouvelle technologie. Selon les documents [2], [3], [14] et [19], l'ULB présente les avantages suivants :

- 1) Les signaux caractéristiques de l'ULB, tels qu'ils ont été décrits précédemment, ont une faible densité spectrale de puissance étant donné que la puissance du signal est répartie sur une grande largeur de bande. Cette particularité confère aux systèmes utilisant l'ULB une faible probabilité de détection et d'interception. De plus, les signaux à bande ultra large interfèrent peu ou pas avec d'autres signaux, tels les signaux à bande étroite, car leur puissance est très faible sur une petite partie du spectre. L'utilisation de ce type de signal permet donc plus de sécurité.
- 2) L'ULB ne souffre pas des effets des trajets multiples, car le chemin direct arrive bien avant que l'annulation avec des chemins secondaires se produise. On peut même tirer avantage des chemins multiples à l'aide d'un type de récepteur appelé *récepteur en râteau* (*rake receiver*) [14].
- 3) Les signaux à ultra large bande possèdent de bonnes capacités de pénétration probablement dues à leur grande largeur de bande. Ils peuvent ainsi traverser des surfaces, tels des murs.
- 4) Les signaux à bande ultra large permettent une grande précision dans la mesure des distances, car la résolution est inversement proportionnelle à la durée de l'impulsion.
- 5) Les systèmes utilisant l'ULB sont assez peu complexes, en partie parce qu'ils ne nécessitent pas l'utilisation d'un étage intermédiaire (fréquence intermédiaire f_{if}) à l'entrée du récepteur [14]. Ces systèmes peuvent également être intégrés en une seule puce, quelques pièces étant extérieures à cette puce, entraînant un faible coût de fabrication.

- 6) Enfin, comme les systèmes utilisant l'ULB ont des puissances moyennes très faibles, ils peuvent potentiellement fonctionner sur des bandes de fréquences déjà occupées par d'autres systèmes (à large bande, à bande étroite, etc.) sans créer d'interférences. L'ULB répond donc au besoin de plus en plus élevé de bandes de fréquences alors que le spectre est déjà encombré. Il pourrait ainsi permettre de régler le problème de disponibilité de fréquences dans le spectre.

Les inconvénients de l'utilisation de l'ULB sont plus difficiles à trouver parce que cette technologie n'a pas encore été éprouvée et parce que les entreprises exploitant l'ULB ne sont pas intéressées à présenter les mauvais côtés de la technologie sur laquelle ils travaillent. De plus, ces inconvénients dépendent en grande partie de l'application pour laquelle le système est conçu. On peut tout de même affirmer que, comme toute technologie RF sans fil, l'ULB n'échappe pas aux lois qui régissent les systèmes de transmission. Il faut donc faire des compromis lors de la conception du système, tel celui entre le rapport signal sur bruit et la largeur de bande [2]. De plus, l'ULB ne pourra probablement jamais arriver à supplanter les systèmes optiques à extrêmement grands débits (de l'ordre de plusieurs gigabits par seconde) quoique les systèmes à ultra large bande soient moins coûteux [2]. La *FCC*, dans son avis de proposition de réglementation [8], a également relevé plusieurs commentaires d'organisations qui manifestaient des inquiétudes au sujet de l'augmentation du niveau global de bruit avec l'utilisation de systèmes à bande ultra large. La *FCC* croit plutôt que l'effet cumulatif de plusieurs dispositifs utilisant l'ULB sera en général négligeable et dépendra des émissions propres à chaque système. Elle croit aussi qu'un système n'utilisant pas l'ULB ne sera affecté que par l'émetteur à bande ultra large le plus près de la zone d'intérêt de ce système et émettant à la même fréquence. Des tests et des analyses devront tout de même vérifier ces conclusions [8]. Enfin, il est certain que les signaux caractéristiques de l'ULB sont sujets aux phénomènes de dispersion et de distorsion en fréquence puisqu'ils s'étendent sur une grande largeur de bande [3]. La dispersion est le phénomène par lequel les différentes composantes fréquentielles d'un signal ne se propagent pas à la même vitesse dans un canal. La distorsion en fréquence (*frequency distortion*) peut être définie de la façon suivante : « *Distortion that is caused by nonuniform attenuation or gain of the*

*different frequencies in a signal*⁴. » Les récepteurs à bande ultra large doivent donc tenir compte de ces effets.

Il existe également des inconvénients liés à la façon dont les systèmes utilisant l'ULB sont conçus. Par exemple, la transmission d'impulsions uniformément espacées possède deux inconvénients selon *Time Domain* qui ont été mentionnés à la section 2.1.3 *Vue d'ensemble des travaux de deux compagnies* dans le paragraphe sur cette entreprise. Un autre exemple est le désavantage lié au fait de transmettre plusieurs impulsions par bit d'information, technique utilisée par *Time Domain*. En émettant plus d'impulsions que de bits de données, l'énergie totale transmise est d'autant augmentée. Le concepteur devra donc diminuer son taux de transmission s'il désire garder la même énergie moyenne qu'en envoyant une impulsion par bit d'information [2]. En fait, les compagnies travaillant sur l'ULB insistent habituellement sur les désavantages reliés à la façon dont leurs concurrents utilisent cette technologie. Il faut donc rester critique face aux différents inconvénients présentés par ces entreprises.

2.1.5 Applications de l'ULB

Les applications possibles de l'ULB sont nombreuses [3] [8] [14]. Il est impensable d'en dresser une liste exhaustive. On peut toutefois les classer grossièrement en deux catégories : les applications radar et les applications en communications. Les dispositifs radar utilisant l'ULB peuvent servir à mesurer des distances ou des positions avec une plus grande résolution que les dispositifs radar existants ou à obtenir des images d'objets enfouis sous la terre et placés derrière des surfaces. Ainsi, les services de police, d'incendie ou d'urgence, par exemple, pourraient utiliser cette nouvelle technologie pour trouver des personnes cachées derrière des murs ou retrouver des corps parmi des décombres. Elle pourrait également permettre de trouver l'emplacement des clous, des fils électriques ou des poutres dans les murs et de localiser des canalisations d'eau ou de gaz naturel enfouies. On pourrait même utiliser l'ULB dans les appareils photographiques à mise au point automatique pour permettre une plus grande précision dans le calcul des distances ou, en médecine, dans les détecteurs de contractions du cœur.

⁴ [22], p. 432.

En télécommunications – et ce sont les applications qui nous intéressent dans le cadre de ce projet –, l’ULB peut permettre les communications sécurisées sans fil de la voix et des données à potentiellement haut débit sur de courtes distances. On pourrait la voir bientôt utilisée pour l’Internet à large bande, la téléphonie, le câble ou les réseaux d’ordinateurs. Dans les systèmes de communications, l’ULB se situe au niveau de la couche physique. *XtremeSpectrum, Inc.* [23] croit qu’on pourra utiliser l’ULB pour transmettre des données jusqu’à 100 Mbps.

2.2 Principes des systèmes de transmission à spectre étalé

Dans cette section, certains éléments des systèmes de transmission à spectre étalé sont présentés. Ces principes, étudiés en vue de les appliquer aux systèmes utilisant l’ULB et afin de mieux comprendre les systèmes décrits dans les travaux sur l’ULB, sont exposés ici en surface seulement. Le lecteur qui désire obtenir de plus amples informations à ce sujet peut consulter les références qui sont fournies dans cette section. On présente d’abord une description générale des systèmes à spectre étalé. On explique ensuite la façon dont les codes pseudo-aléatoires sont choisis et générés. Enfin, la dernière sous-section expose les grands principes de réception des signaux à étalement du spectre en séquence directe.

2.2.1 Description générale du principe d’étalement du spectre en séquence directe

Les systèmes de transmission à spectre étalé se basent sur l’émission et la réception de signaux dont la largeur de bande est beaucoup plus grande que le débit binaire R_b de l’information à transmettre. Cette expansion de largeur de bande peut se faire de deux façons : par séquence directe (*direct sequence*) ou par sauts de fréquence (*frequency hopping*). Nous concentrons notre discussion sur les systèmes à séquence directe puisque ce sont eux qui s’appliquent à l’ULB. Les informations contenues dans cette première sous-section proviennent, pour la plupart, du volume [24].

L'étalement de spectre en séquence directe se fait par la multiplication de l'information à transmettre de débit R_b par un code pseudo-aléatoire, aussi appelé *signature*, ayant un débit R_c . On a :

$$L_c = \frac{R_c}{R_b} = \frac{T_b}{T_c}, \quad (2.7)$$

où $T_b = \frac{1}{R_b}$ est la durée d'un bit d'information et $T_c = \frac{1}{R_c}$ est la durée d'une impulsion

rectangulaire du code, appelée *chip*⁵. L_c est habituellement un entier, est supérieur à 1 puisqu'il mesure l'étalement du spectre et représente le nombre de chips par bit d'information. On appelle également ce rapport *gain résultant du traitement (processing gain)*. En ce sens, il représente une mesure de la résistance à l'interférence et au brouillage intentionnel obtenue en augmentant la largeur de bande du signal transmis. La *figure 2.8* présente un schéma général d'un système à étalement de spectre en séquence directe. La somme modulo-2 (porte OU exclusif) pour des données en

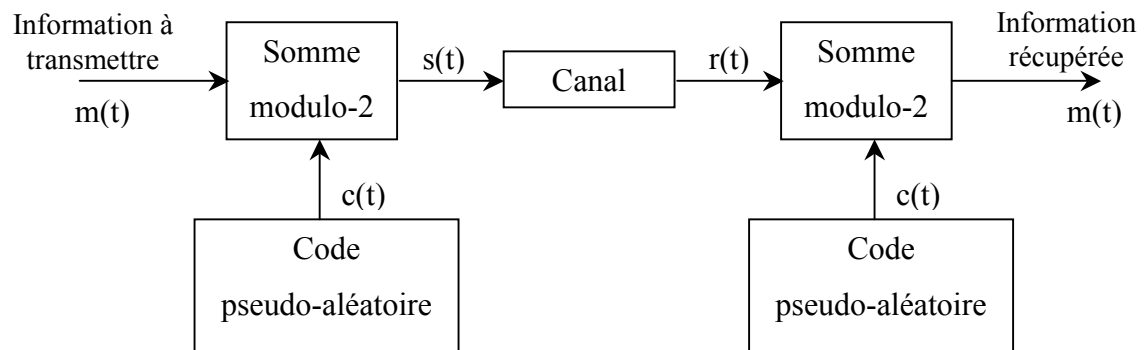


Figure 2.8 : Schéma général d'un système à étalement de spectre en séquence directe

format 0 et 1 est l'équivalent d'une multiplication de signaux en format⁶ -1 et 1. Les *tableaux 2.1* et *2.2* présentent les tables de vérité pour la somme modulo-2 et pour le produit de signaux où l'indice i représente une chip du code.

⁵ Nous utilisons le mot anglais *chip* dans le présent document quoique l'UIT propose la traduction *élément* [1]. En effet, nous croyons que le mot *élément* entraîne de la confusion puisque c'est un terme générique et qu'il est utilisé ailleurs dans le texte pour représenter d'autres notions.

⁶ À partir de maintenant, on appellera le format 0 et 1 *binaires simple* et le format -1 et 1 *cas antipodal*.

Tableau 2.1 : Table de vérité pour la somme modulo-2

m_i	c_i	$s_i = m_i \oplus c_i$
0	0	0
0	1	1
1	0	1
1	1	0

Tableau 2.2 : Table de vérité pour le produit

m_i	c_i	$s_i = m_i \times c_i$
-1	-1	1
-1	1	-1
1	-1	-1
1	1	1

La *figure 2.9* montre un exemple de signaux $m(t)$, $c(t)$ et $s(t)$ en bande de base tels que définis à la *figure 2.8* et permet d'illustrer le principe d'étalement de spectre et de produit des signaux dans le cas antipodal. On peut remarquer, à partir du *tableau 2.2* et de la *figure 2.9*, que le signal transmis $s(t)$ est égal au code $c(t)$ lorsque $m(t)$ vaut 1 alors qu'il est égal à $-c(t)$ lorsque $m(t)$ vaut -1.

Au récepteur, en multipliant le signal reçu $r(t)$ par le même code pseudo-aléatoire $c(t)$ que celui appliqué à l'entrée, on récupère l'information transmise lorsque les deux signaux multipliés sont synchronisés⁷. En effet, l'influence du code pseudo-aléatoire est ainsi éliminée.

Dans un système à plusieurs usagers (AMRC ou CDMA), chaque usagé possède un code pseudo-aléatoire qui lui est propre, ce qui oblige le récepteur à connaître le code associé à l'utilisateur qui l'intéresse permettant ainsi une séparation entre les différents utilisateurs du système.

⁷ On traitera plus en détail de la synchronisation à la sous-section 2.2.3 *Réception de signaux à étalement du spectre en séquence directe*.

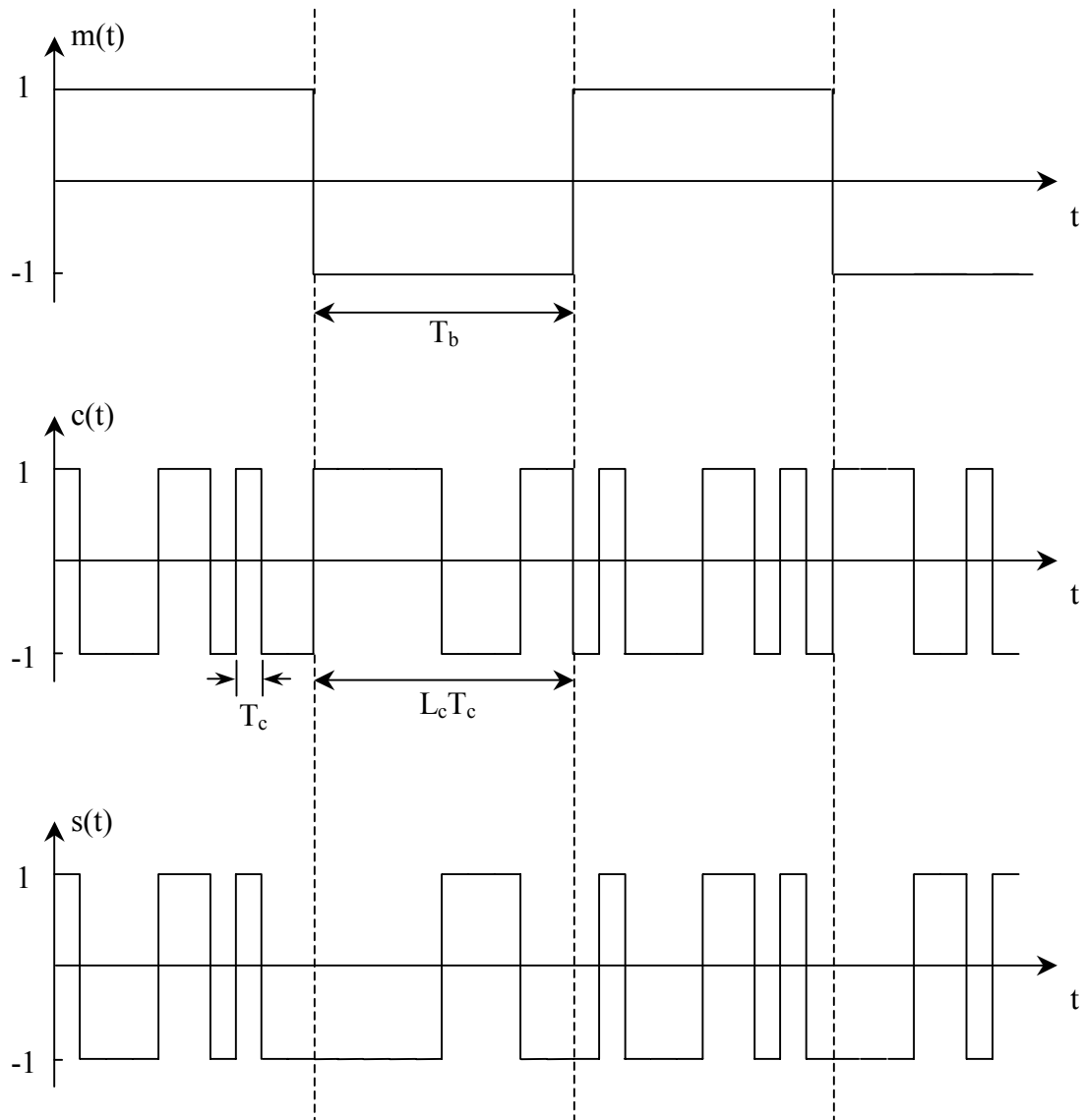


Figure 2.9 : Exemple de signaux $m(t)$, $c(t)$ et $s(t)$

2.2.2 Génération et choix des codes pseudo-aléatoires

Les séquences de code $c(t)$ doivent ressembler le plus possible à des séquences aléatoires afin, entre autres, d'empêcher quelqu'un de trouver le code utilisé et de se servir de cette information pour brouiller intentionnellement le signal envoyé. Comme il est impossible de générer des séquences complètement aléatoires avec une mémoire finie, on doit utiliser des séquences périodiques de période L très longue afin de simuler un phénomène aléatoire [25]. Ces séquences périodiques et déterministes sont appelées *codes pseudo-aléatoires* ou *codes de pseudo-bruit* (*pseudorandom sequences*,

pseudonoise sequences ou *PN code*⁸) puisqu'elles semblent aléatoires. Les informations présentées dans cette sous-section peuvent être retrouvées, exposées de différentes façons et avec plus ou moins de rigueur mathématique, dans n'importe quelle des références suivantes sauf lorsque précisé explicitement : [25], [26], [27] ou [28].

La **génération de codes PN** peut se faire grâce à des registres à décalage (*shift registers*) composés de bascules D. La *figure 2.10* illustre un exemple de générateur pseudo-aléatoire formé de 4 bascules D et de 2 prises ou branchements (*taps*) sans compter la prise de rétroaction (position 0). Ces 2 prises sont combinées en un OU

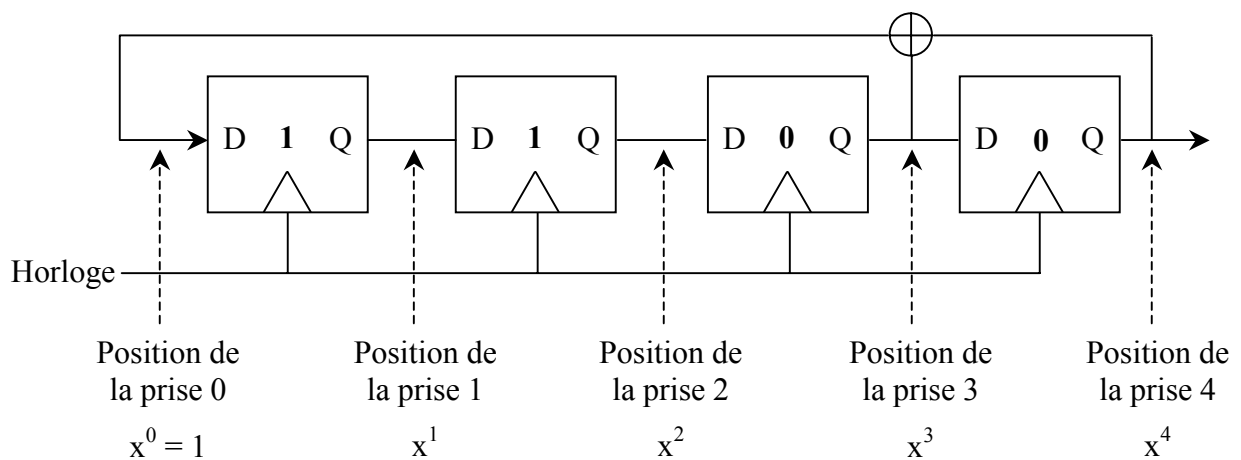


Figure 2.10 : Exemple de registre à décalage composé de 4 bascules D

exclusif, symbolisé par le symbole $\oplus$ sur la *figure 2.10*, ce qui permet de générer des séquences différentes des valeurs initiales placées dans les bascules. Ces valeurs initiales sont 1100 pour l'exemple de la *figure 2.10*. À chaque coup d'horloge, les valeurs contenues dans les bascules sont décalées vers la droite. La séquence obtenue d'un registre à décalage se trouve à la sortie de la dernière bascule.

Les séquences générées par un registre à décalage dépendent de la longueur, des prises de rétroaction et des valeurs initiales de ce registre. La séquence obtenue par la configuration de la *figure 2.10* est périodique de période 15. Cette séquence périodique

⁸ Dans la suite du document, on fera référence aux codes pseudo-aléatoires par l'expression *code PN*, cette abréviation étant courante dans la littérature.

est 001101011110001 [29]. Elle est illustrée à la *figure 2.11* dans le cas antipodal. Tant que l'horloge fonctionne, on obtient cette séquence de façon cyclique à tous les 15 coups d'horloge. La position des prises permet d'identifier le polynôme caractéristique de cette

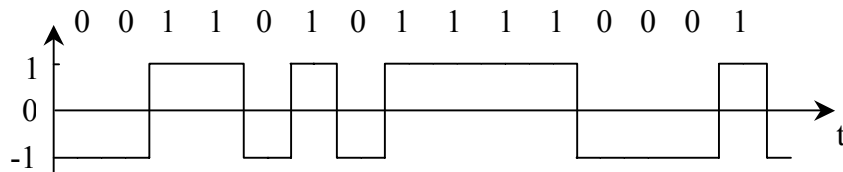


Figure 2.11 : Forme d'onde de la séquence 001101011110001 générée par le registre à décalage de la *figure 2.10*

séquence et de caractériser cette configuration, c'est-à-dire de lui donner un nom ou une notation. Le polynôme caractéristique [27] de la séquence générée par le registre à décalage de la *figure 2.10* est $f(x) = x^0 + x^3 + x^4 = 1 + x^3 + x^4$. Pour caractériser une configuration de registre à décalage, il faut connaître le nombre de bascules et la position des prises de la droite vers la gauche en hexadécimal [29]. Ainsi, par exemple, on peut faire référence à la configuration de la *figure 2.10* en utilisant la notation R4-0xC, où R4 signifie qu'elle contient 4 bascules (R mis pour *register*) et C, l'équivalent en hexadécimal des valeurs binaires 1100, signifie que les prises sont aux positions 4 et 3. La séquence générée à partir de cette configuration dépend ensuite des valeurs initiales des bascules qu'on a mises à 1100 à la *figure 2.10* et qui sont, ici, par hasard égales aux valeurs binaires de la position des prises.

La séquence obtenue à la *figure 2.10* est une séquence qu'on appelle *séquence maximale* (*maximal length sequence*, *maximal sequence* ou *m-sequence*). Une séquence maximale est une séquence périodique pour laquelle la longueur L de la période est maximale pour le nombre n de bascules du registre à décalage et vaut $L = 2^n - 1$. La longueur L représente le nombre de chips dans une période. On peut générer une séquence maximale uniquement lorsque le nombre de prises, excluant la prise de rétroaction, est pair. Les séquences maximales possèdent entre autres les propriétés suivantes :

- 1) Elles sont équilibrées (*balanced*), c'est-à-dire qu'elles possèdent un 1 de plus que de 0 dans une période complète de $2^n - 1$ éléments binaires. Dans l'exemple de

la *figure 2.10*, la période compte 8 uns et 7 zéros. La probabilité que, à un coup d'horloge quelconque, la sortie du registre à décalage soit un 1 ou un 0 s'approche de 0,5 plus la période L de la séquence est longue :

$$\begin{aligned} P(0) &= \frac{1}{2} \left(1 - \frac{1}{L} \right) \\ P(1) &= \frac{1}{2} \left(1 + \frac{1}{L} \right) \end{aligned} \quad (2.8)$$

- 2) La somme modulo-2 d'une séquence maximale $\{c_n\}$ en binaire simple terme à terme avec cette même séquence décalée dans le temps par moins de L cycles d'horloge donne la séquence maximale $\{c_n\}$ décalée dans le temps par rapport aux deux séquences de départ.
- 3) La distribution statistique des 1 et des 0 dans une période d'une séquence maximale est bien définie et toujours la même. Nous utiliserons ici le mot *plage* pour représenter le mot anglais *run* qui désigne une série de chips identiques d'une certaine longueur. Par exemple, 000 est une plage de 0 de longueur 3. Ainsi, il a été montré qu'il existe exactement $2^{[n-(p+2)]}$ plages de longueur $p, p \leq n$, de 0 et de 1 pour chaque période d'une séquence maximale, sauf qu'il n'y a qu'une plage de 1 de longueur n et qu'une plage de 0 de longueur $n-1$. Également, il n'y a pas de plage de 0 de longueur n ou de plage de 1 de longueur $n-1$. Par exemple, pour la séquence de l'exemple de la *figure 2.10*, il y a bel et bien une seule plage de 0 de longueur $p = 2$ ($2^{[4-(2+2)]} = 1$) et une seule plage de 1 de longueur 2. La position de ces différentes plages varie d'une séquence maximale à l'autre, mais le nombre de chacune de ces plages d'une longueur particulière est toujours le même pour toutes les séquences de même période L .

Les trois propriétés précédentes sont celles d'un phénomène aléatoire. Elles permettent donc de démontrer que les séquences maximales sont bel et bien des codes *PN* et que, comme mentionné précédemment et quoique ce soit évident même sans preuve, plus le code est long, plus il ressemble à un phénomène aléatoire.

Le **choix des codes *PN*** à utiliser dans un système à spectre étalé doit être basé sur les propriétés de corrélation de ces codes. Commençons par définir les termes *autocorrélation* et *corrélation croisée* (également appelée *intercorrélation*).

L'autocorrélation est habituellement définie comme la mesure de ressemblance entre un signal $f(t)$ et une copie de ce même signal décalé. La fonction d'autocorrélation ϕ pour un décalage de τ est :

$$\phi(\tau) = \int_0^t f(t) f(t - \tau) dt . \quad (2.9)$$

La fonction d'autocorrélation est habituellement tracée pour tous les décalages τ par rapport au signal de référence $f(t)$. Dans le cas de la corrélation entre codes PN où les chips sont des éléments discrets, on peut remplacer l'intégrale continue par une sommation. On a donc la fonction d'autocorrélation suivante sur une période du code PN pour un décalage τ par rapport au code de référence⁹ :

$$\phi_\tau = \sum_{i=1}^L c_i c_{i-\tau} , \quad (2.10)$$

dans le cas de chips en format antipodal, où $c_{i-\tau}$ est la $i^{\text{ème}}$ chip décalée d'un nombre entier de chips τ par rapport au code de référence dont la valeur de la $i^{\text{ème}}$ chip est c_i et L est, comme précédemment, le nombre de chips par période de la séquence du code PN . De la même façon que pour l'équation 2.9, on parle généralement de la fonction d'autocorrélation lorsqu'on trace ϕ_τ pour tous les décalages τ . Pour illustrer le principe de l'autocorrélation telle que définie par l'équation 2.10, prenons l'exemple de la figure 2.10 où la séquence 001101011110001 devient, dans le cas antipodal, la séquence : -1-111-11-11111-1-1-11. L'autocorrélation de ce code pour $\tau = 0$ se calcule, selon l'équation 2.10, en multipliant terme à terme les valeurs des chips et en additionnant les résultats :

$$\begin{array}{r} -1-111-11-11111-1-1-11 \\ \times \quad -1-111-11-11111-1-1-11 \\ \hline 1+1+1+1+1+1+1+1+1+1+1+1+1 = 15. \end{array} \quad (2.11)$$

La corrélation croisée, quant à elle, est la mesure de la ressemblance entre deux signaux différents, $f(t)$ et $g(t)$, ou, dans le cas des systèmes à spectre étalé, entre deux codes distincts, c et u . On obtient la définition générale de l'intercorrélation en

⁹ Par code de référence, on entend le code pour lequel le déphasage τ vaut 0.

remplaçant $f(t - \tau)$ par $g(t - \tau)$ dans l'équation 2.9 et $c_{i-\tau}$ par $u_{i-\tau}$ dans l'équation 2.10. Ce sont les définitions reliées aux codes des systèmes à spectre étalé qui seront utilisées dans cette sous-section à partir de maintenant.

On ne saurait trop insister sur l'importance des propriétés de l'autocorrélation dans le choix des codes utilisés dans les systèmes à spectre étalé et également des propriétés de l'intercorrélation dans le cas des environnements à plusieurs usagers (AMRC). En effet, afin que le récepteur soit capable de se synchroniser sur le code PN approprié à la bonne phase pour récupérer l'information, la valeur de l'autocorrélation de ce code non décalé ($\tau=0$) doit être de beaucoup supérieure à sa valeur d'autocorrélation pour tous les décalages de phase ainsi qu'aux valeurs de corrélation croisée de ce code avec tous les déphasages des codes des autres usagers du système. Les techniques de réception et de synchronisation présentées dans la prochaine sous-section vont permettre de mieux comprendre le lien entre la corrélation et la réception de l'information.

La fonction d'autocorrélation pour une période complète d'une séquence maximale a toujours l'allure de la courbe présentée à la *figure 2.12*. Cette fonction atteint

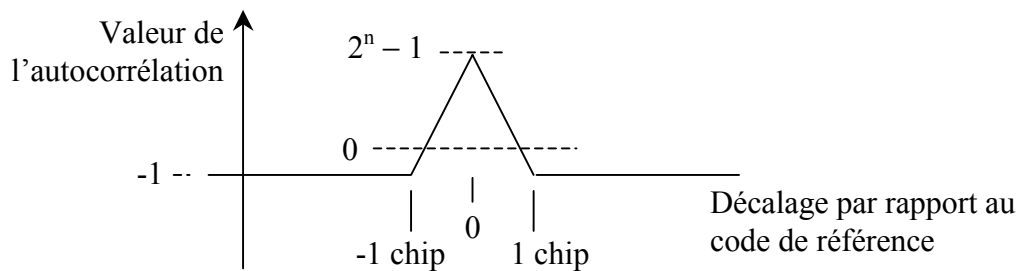


Figure 2.12 : Fonction d'autocorrélation périodique d'une séquence maximale

un maximum égal à $2^n - 1$, avec n égal au nombre de bascules du registre à décalage, lorsque les deux versions de la séquence maximale sont en phase. À l'équation 2.11, on avait bel et bien trouvé une autocorrélation de $2^4 - 1 = 15$ pour un décalage de $\tau=0$. On remarque, d'après la *figure 2.12*, qu'il est facile de savoir à quel moment le code décalé est en phase avec le code de référence. En effet, le pic d'autocorrélation (lorsque les deux versions du code sont en phase) est largement supérieur à la valeur de l'autocorrélation pour des décalages par rapport au code de référence et d'autant plus quand le code est

long. Les propriétés d'autocorrélation pour les séquences maximales sont donc excellentes.

Par contre, les séquences maximales n'ont pas de faibles corrélations croisées et il n'existe pas de formule générale pour déterminer l'intercorrrelation de deux de ces séquences [30]. Le rapport du pic de corrélation croisée d'une séquence maximale avec une autre séquence maximale sur le pic d'autocorrélation de la première séquence peut facilement atteindre 0,6, ce qui est inacceptable dans un système de communication à plusieurs usagers. Il est possible de trouver des séquences maximales qui ont de faibles intercorrélations, mais il y en aura peu. Le système ne pourra donc pas recevoir un grand nombre d'utilisateurs.

Gold (1967) et Kasami (1968) ont montré qu'il existe des codes, appelés respectivement *codes Gold* et *codes Kasami*, qui ne sont pas maximaux et qui possèdent de plus faibles corrélations croisées que les séquences maximales. Les propriétés de corrélation de ces deux codes sont semblables. La plus grande différence entre ces deux familles de codes réside dans le nombre de codes pouvant être générés [30]. Notre brève description se limitera toutefois ici aux codes Gold. Les codes Gold sont générés à partir de deux séquences maximales obtenues à partir de n bascules, tel que montré à la figure 2.13. Quoique ces codes soient dérivés de séquences maximales, ils ont des corrélations croisées bornées contrairement à ces séquences et ces intercorrélations sont

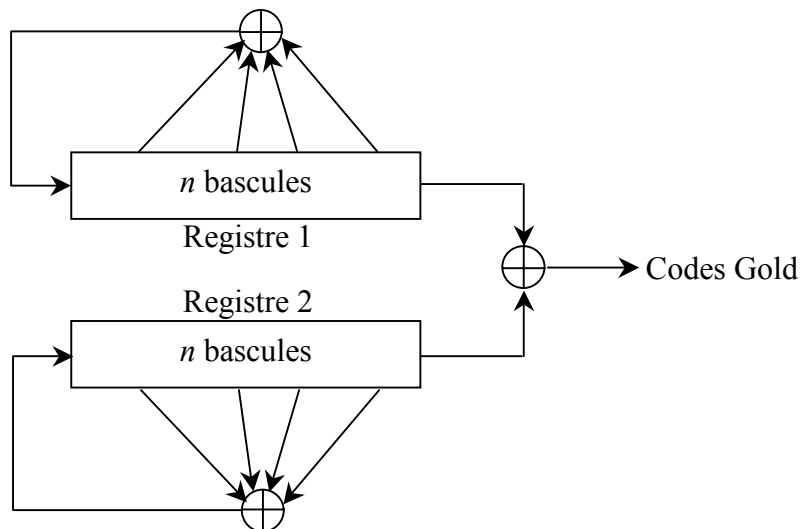


Figure 2.13 : Configuration d'un générateur de codes Gold
(adaptation d'une figure de [27])

connues et uniformes. Elles ne peuvent prendre que 3 valeurs : -1 , $-t(n)$ et $t(n)-2$, où

$$t(n) = \begin{cases} 2^{(n+1)/2} + 1, n \text{ impair} \\ 2^{(n+2)/2} + 1, n \text{ pair} \end{cases}, \quad (2.12)$$

avec n égal au nombre de bascules des registres à décalage des deux séquences maximales d'origine. Contrairement aux séquences maximales, la fonction d'autocorrélation des codes Gold peut prendre plus que deux valeurs et présente donc des pics secondaires en plus du pic principal (sans déphasage) comme celui de la *figure 2.12*. Toutefois, on peut affirmer que les pics secondaires de l'autocorrélation de ces codes sont bornés supérieurement par $t(n)$. Plus les codes utilisés sont longs, plus le rapport de $t(n)$ sur le pic d'autocorrélation principal est faible, donc plus les propriétés des codes Gold sont intéressantes. Ces rapports sont beaucoup plus faibles que leur équivalent pour des séquences maximales ayant la même longueur L de période [24]. De plus, ces codes, qui présentent de très bonnes propriétés de corrélation, sont nombreux, ce qui permet de concevoir des systèmes comportant beaucoup d'utilisateurs. Les codes Gold sont aussi particulièrement efficaces pour l'acquisition dans le processus de synchronisation [28]. La façon de générer des codes Gold et Kasami est exposée dans les références données au début de cette sous-section. On peut trouver des explications plus détaillées sur la sélection de codes Gold à l'appendice 7 de [28] et au chapitre 11 de [27].

Enfin, il est clair que, pour permettre une bonne réception des signaux à étalement du spectre en séquence directe, les codes doivent présenter un grand pic d'autocorrélation principal par rapport aux pics secondaires d'autocorrélation et par rapport aux valeurs d'intercorrélation avec les autres codes du système sur une période complète des séquences. Par contre, lorsque la période L des codes est très longue, les récepteurs peuvent vérifier une partie plus petite des codes que leur période complète afin d'accélérer le processus de synchronisation. On appelle l'autocorrélation et la corrélation croisée de parties de périodes de codes des *corrélations partielles*. Les valeurs de ces corrélations partielles doivent être faibles quelle que soit la longueur de ces parties de codes utilisées et quels que soient les déphasages entre les codes, sauf, évidemment, lors de l'autocorrélation pour un déphasage nul. En effet, si ce n'est pas le cas, le récepteur risque de se verrouiller sur le mauvais code ou sur une mauvaise phase du bon code.

Il faut donc s'assurer qu'aucune partie des codes choisis de la longueur vérifiée par le récepteur ne se retrouve dans le code considéré et dans le code d'aucun autre usager. Plusieurs articles traitent de ce sujet et certains auteurs sont identifiés à la page 729 de [24].

2.2.3 Réception de signaux à étalement du spectre en séquence directe

Les récepteurs de signaux à étalement du spectre en séquence directe comportent trois modules qui effectuent chacun une tâche particulière : la démodulation, l'acquisition et la poursuite (*tracking*) [31]. Dans cette sous-section, nous décrivons brièvement ces trois fonctions du récepteur ainsi que la façon dont elles peuvent être implantées dans un système à spectre étalé.

La **démodulation** permet de récupérer l'information envoyée par l'émetteur, c'est-à-dire de « désétaler » le signal reçu. C'est la fonction de base du récepteur. Son implantation dépend évidemment de la modulation utilisée. Afin d'appliquer les concepts abordés ici à l'ULB, nous concentrons notre étude sur les signaux transmis en bande de base pour les modules démodulation et acquisition. Comme il a déjà été vu à la sous-section 2.2.1 *Description générale du principe d'étalement du spectre en séquence directe*, la démodulation de l'information est obtenue par la multiplication du signal reçu $r(t)$ dans le cas antipodal par une copie $c(t)$ du code PN de l'utilisateur correspondant générée localement au récepteur, qu'on appellera à partir de maintenant *code PN local*. Par cette multiplication, le code PN est retiré du signal à démoduler. On retrouve donc, en théorie, le signal bande de base $m(t)$ tel que montré dans la partie supérieure de la *figure 2.9*. Comme ce signal est également additionné de bruit $n(t)$, le récepteur doit décider, à tous les T_b , durée d'un bit d'information, si c'est un 1 ou un 0 qui a été transmis. On doit donc procéder à une intégration et remise à 0 suivies d'un échantillonnage à tous les T_b . Cette valeur d'échantillon est ensuite comparée à un seuil afin de déterminer quel bit a été transmis [32]. La *figure 2.14* illustre ce principe en présentant un démodulateur théorique pour une transmission en bande de base [24] [32]. Dans les systèmes à spectre étalé conventionnels pour lesquels il y a utilisation d'une porteuse, la multiplication par le code PN local est suivie d'une récupération des données à partir de la porteuse et souvent d'un décodeur.

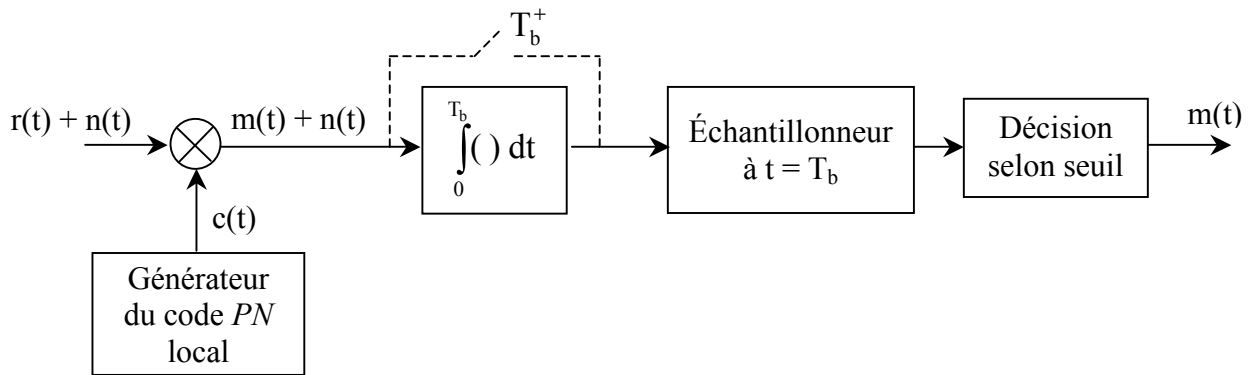


Figure 2.14 : Démodulateur théorique pour transmission en bande de base

La récupération de l'information grâce au système de démodulation qui vient d'être décrit ne peut fonctionner que si le code *PN* local est synchronisé avec le signal reçu. Cette fonction est assurée par la partie synchronisation du récepteur qui est formée des modules acquisition et poursuite. Les explications qui suivent peuvent être trouvées dans les volumes [24], [25], [26], [28] et [33] sauf si précisé explicitement.

L'**acquisition** est le processus par lequel le signal reçu est aligné grossièrement avec le code *PN* local habituellement à plus ou moins une fraction de chip près. Une séquence connue est habituellement envoyée par l'émetteur afin de permettre au récepteur d'acquies le signal avant le début de la démodulation. Plusieurs techniques peuvent être employées pour permettre l'acquisition. Selon les documents consultés, ces méthodes varient. Elles peuvent également porter des noms divers et être classées en différentes catégories selon les auteurs. Nous nous contenterons ici de décrire rapidement deux de ces techniques qui sont assez généralement utilisées et qui, à elles seules, permettent de bien comprendre la suite du présent rapport. Ces deux méthodes sont l'utilisation d'un filtre adapté et l'utilisation d'un corrélateur dynamique.

Par définition, un filtre adapté est « un filtre linéaire réalisable qui fournit à sa sortie un rapport signal à bruit maximal pour le signal auquel il est adapté¹⁰. » Dans le cas des systèmes à spectre étalé, le filtre est adapté à la forme d'onde du code *PN* propre à l'utilisateur considéré. Il peut être implémenté à l'aide d'une ligne à retard à prises (*tapped delay line*), comme l'illustre la figure 2.15. Dans cette figure, le filtre est adapté à la

¹⁰ [32], p. 157.

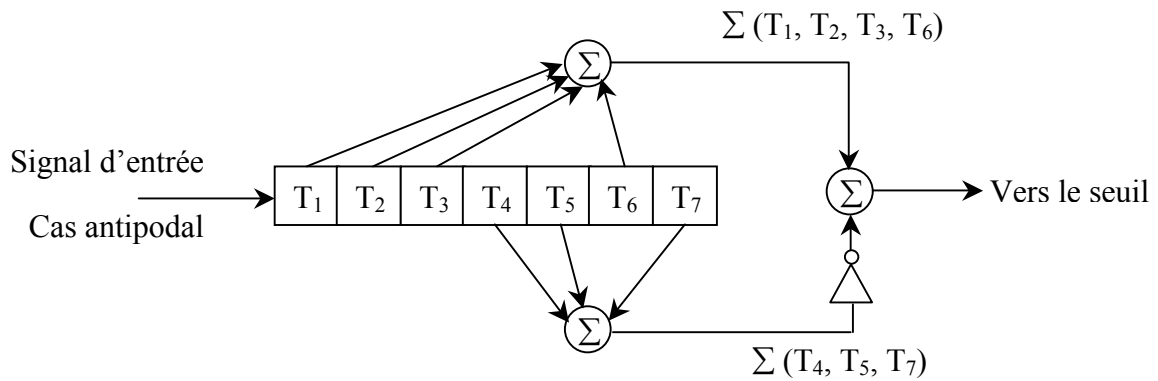


Figure 2.15 : Exemple de filtre adapté à ligne à retard
(adaptation d'une figure de [28])

séquence maximale de 7 chips *0100111* générée par un registre à décalage à 3 bascules dont les valeurs initiales sont *010*. Les blocs T_1 à T_7 sont des éléments de la ligne à retard qui retardent chacun de 1 chip le signal entrant. Le signal reçu circule dans la ligne jusqu'à ce qu'elle soit remplie par la séquence. Les chips présentes dans les blocs T_1 , T_2 , T_3 et T_6 , qui valent toutes 1 lorsque la séquence est dans la phase appropriée par rapport au filtre adapté, et celles dans les blocs T_4 , T_5 et T_7 , qui valent toutes -1 lorsque la séquence est dans la phase appropriée, sont sommées séparément. On obtient la sortie du filtre adapté en additionnant ces deux sommes après avoir inversé celle de T_4 , T_5 et T_7 . Lorsque la séquence correspondant au bon code *PN* est dans la phase appropriée par rapport au filtre adapté, la sortie de la dernière sommation vaut idéalement $2^n - 1 = 2^3 - 1 = 7$, soit la valeur du pic d'autocorrélation de la séquence maximale. Un seuil est alors franchi et la synchronisation initiale est considérée acquise. On peut trouver de l'information supplémentaire sur l'utilisation de filtres adaptés pour la synchronisation en particulier dans les ouvrages [25] et [28].

Contrairement au filtre adapté qui est une technique de corrélation passive, le corrélateur dynamique est une forme de corrélation active. Certains livres dont [25] désignent cette méthode par l'expression *single dwell serial PN acquisition system*. La figure 2.16 montre un exemple de corrélateur dynamique. Le signal reçu est d'abord corrélé avec le code *PN* local. Cette corrélation s'effectue durant un certain temps τ_d , appelé en anglais *dwell time*, qui correspond à un nombre entier de chips. Si le résultat de l'intercorrélation se situe sous un certain seuil S (branche *Non* sur la figure), la phase du

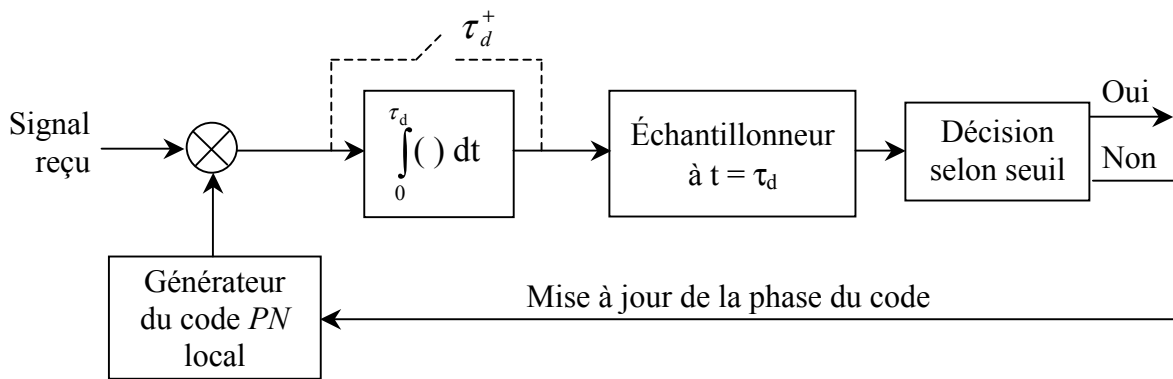


Figure 2.16 : Principe du corrélateur dynamique

code *PN* local est ensuite mise à jour, c'est-à-dire que le code *PN* local est retardé d'un temps τ_r , correspondant à une fraction de chip qui vaut souvent $\frac{1}{2}T_C$ secondes. Ainsi, tant que le seuil n'est pas franchi, à chaque τ_d , le code *PN* généré localement est décalé dans le temps par rapport au signal reçu par incrément de τ_r et une nouvelle corrélation croisée est effectuée entre le signal entrant et le code décalé. Comme le code *PN* local semble « glisser » par rapport au signal reçu, la corrélation dynamique est appelée *sliding correlation* en anglais. Si, par contre, le résultat de la corrélation dépasse le seuil S utilisé (branche *Oui* sur la figure), c'est-à-dire que le signal reçu est approximativement en phase avec le code *PN* local, on peut dire que l'acquisition est terminée et le récepteur peut arrêter la recherche et passer à la phase de la poursuite. À partir des explications précédentes concernant l'acquisition, on peut remarquer l'importance des propriétés de corrélation des codes *PN* choisis. En effet, ce sont ces propriétés qui détermineront si le récepteur peut se synchroniser sur le signal utile ou non, car le module acquisition utilise la corrélation pour déterminer si le signal entrant est le signal désiré.

Une fois l'acquisition terminée, le signal se trouve aligné à une fraction de chip près avec le code *PN* local. Le module **poursuite** prend ensuite le relais afin d'aligner ces deux signaux plus précisément et afin de compenser la dérive des horloges s'il y a lieu. La poursuite est effectuée de façon continue grâce à une boucle de rétroaction et permet ainsi de maintenir une synchronisation fine entre les signaux. Si la synchronisation est perdue, le récepteur peut retourner à la phase acquisition pour revenir ensuite à la poursuite. Il existe principalement deux façons d'effectuer la poursuite : en utilisant une boucle à

retard de phase (*delay-locked loop*) ou en utilisant une boucle appelée *tau-dither loop* en anglais¹¹. Pour synchroniser le signal d'entrée avec le code *PN* local, ces deux méthodes dégradent intentionnellement le résultat de la corrélation de ces deux signaux afin d'utiliser l'information tirée de cette dégradation pour diminuer le décalage qui les sépare. Pour ce faire, elles utilisent le fait que la fonction de corrélation est symétrique par rapport au pic de corrélation comme on peut le voir à la *figure 2.12*.

Le principe de la boucle à retard de phase est présenté à la *figure 2.17* pour un signal à spectre étalé conventionnel. Cette boucle effectue la corrélation croisée entre le

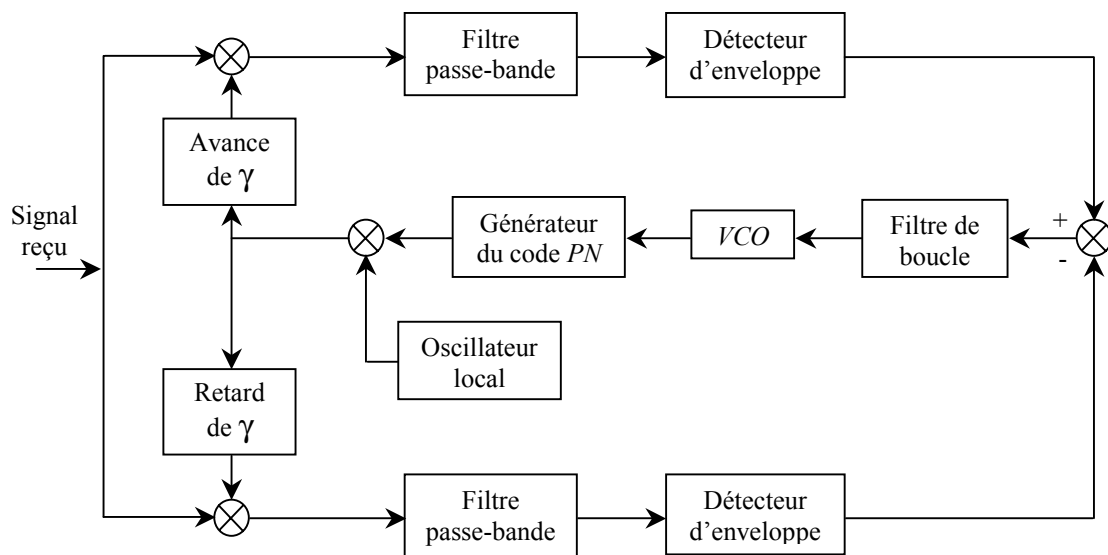


Figure 2.17 : Boucle à retard de phase
(adaptation d'une figure de [24])

signal reçu et deux versions du code *PN* local : une en avance d'un temps γ et l'autre en retard de ce même temps γ . Les deux versions du code sont ainsi séparées de 2γ , où $2\gamma \leq T_c$. Les résultats de ces deux intercorrélations sont ensuite soustraits l'un de l'autre. Puis, le signal différence est filtré et commande un oscillateur contrôlé en tension (*voltage controlled oscillator* ou *VCO*) ou une horloge contrôlée en tension qui permet d'ajuster la fréquence du code *PN* local. Si le signal reçu n'est pas parfaitement synchronisé avec le générateur du code *PN*, le signal différence filtré sera différent de 0

¹¹ Nous n'avons pas trouvé de traduction de cette expression ni été capables de créer nous-mêmes une traduction qui nous semble adéquate.

et sera proportionnel en signe et en amplitude à l'erreur de phase. La sortie du *VCO* sera alors ajustée à la hausse ou à la baisse afin de « rattraper » ou « d'attendre » respectivement le signal entrant. Au point d'équilibre, les deux branches de la boucle sont égales, car la boucle se trouve centrée sur le pic de la courbe d'autocorrélation symétrique des signaux, et le signal différence est nul. Le code *PN* local est alors parfaitement synchronisé avec le signal entrant.

La principale lacune de la boucle à retard de phase est sa sensibilité au déséquilibre des gains des deux branches de la boucle. En effet, si les amplitudes des deux bras ne sont pas équilibrées, le signal différence ne sera pas nul lorsque les signaux seront synchronisés. Une façon de remédier à ce problème est d'utiliser une boucle *tau-dither loop*, car elle ne possède qu'une seule branche, comme le montre la *figure 2.18* pour un signal à spectre étalé conventionnel.

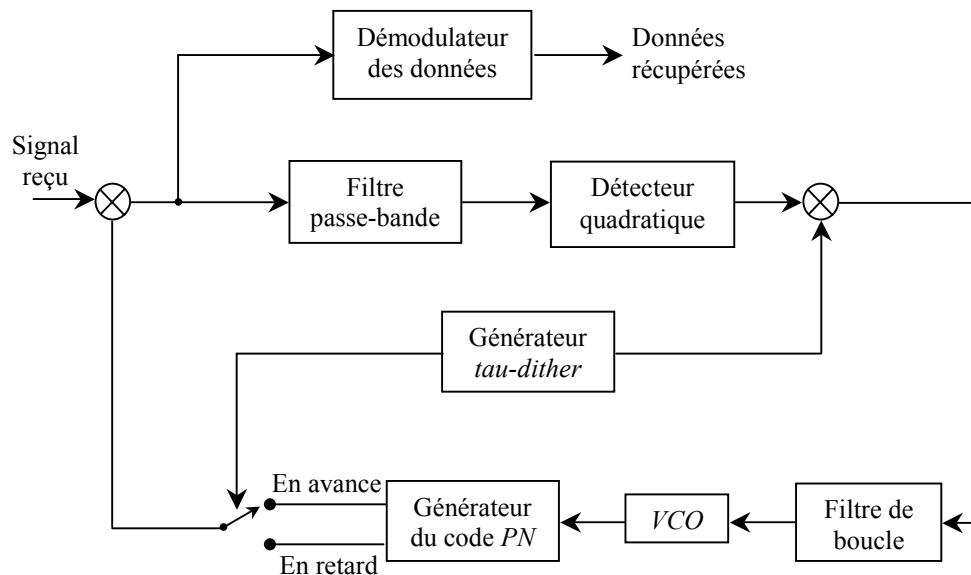


Figure 2.18 : Boucle *tau-dither loop*
(adaptation d'une figure de [33])

Une fois le récepteur en mode poursuite, il peut commencer à démoduler, en parallèle, le signal reçu en utilisant la technique décrite plus haut et illustrée à la *figure 2.14*. Un exemple typique de récepteur pour un système à spectre étalé avec porteuse est montré à la *figure 2.19* où les trois parties d'un récepteur sont clairement indiquées.

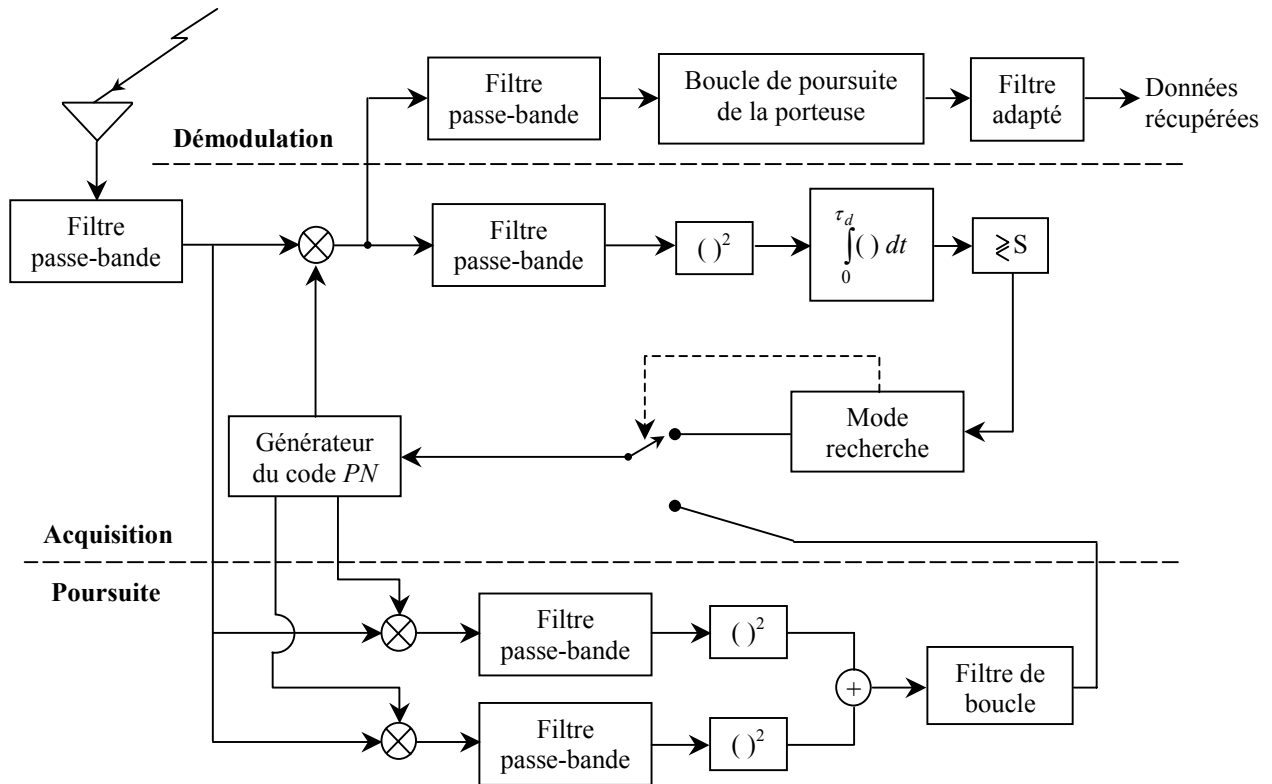


Figure 2.19 : Schéma d'un récepteur pour un système à spectre étalé conventionnel
(adaptation d'une figure de [27])

2.3 Comparaison entre les signaux à bande ultra large et les signaux à spectre étalé

À la lumière des informations contenues dans les deux dernières sections et dans les documents cités, il est possible de faire ressortir les principales différences qui existent entre les signaux à bande ultra large et les signaux à spectre étalé. Ces différences ne sont habituellement pas indépendantes l'une de l'autre :

- 1) La différence fondamentale réside dans le fait que les signaux à spectre étalé sont modulés par une porteuse sinusoïdale, ce qui implique qu'on transmet une onde sinusoïdale, alors que les signaux à bande ultra large sont des trains d'impulsions, des trains de parties de sinusoïdes qui ne sont pas des ondes continues et qui sont sans porteuse [3].
- 2) La largeur de bande des signaux à spectre étalé se situe habituellement entre 1 % et 25 % de la fréquence centrale alors qu'elle est supérieure à 25 % de la fréquence centrale pour les signaux à bande ultra large, tel que défini à la

sous-section 2.1.2 *Définition générale de l'ULB et principes*. Pour fins de comparaison, mentionnons que les signaux à bande étroite ont généralement une largeur de bande inférieure à 1 % de leur fréquence centrale [3].

- 3) Contrairement aux systèmes à spectre étalé, l'utilisation d'un code *PN* dans les systèmes utilisant l'ULB ne sert pas à augmenter la largeur de bande des signaux, car les signaux à bande ultra large ont fondamentalement une largeur de bande très grande due à leur forme [17]. Le code *PN* sert plutôt, comme il a été énoncé à la sous-section 2.1.3 *Vue d'ensemble des travaux de deux compagnies*, à séparer les différents usagers du système, à uniformiser l'énergie dans le domaine des fréquences et à diminuer l'influence du brouillage intentionnel.
- 4) Comme il a été mentionné à la sous-section 2.1.4 *Avantages et inconvénients de l'ULB*, les signaux à bande ultra large peuvent potentiellement être utilisés dans les mêmes bandes de fréquences que les signaux d'autres systèmes de transmission, car leur puissance moyenne est très faible. Au contraire, les fournisseurs de services des systèmes à spectre étalé possèdent leur propre bande de fréquences d'opération.
- 5) Dans les systèmes à étalement du spectre en séquence directe, l'énergie par bit d'information dépend du débit binaire des données pour une puissance de crête d'émission constante, contrairement aux systèmes utilisant l'ULB. Le taux de transmission des bits pour l'ULB influe seulement sur la puissance moyenne. Ainsi, pour un faible taux de transmission, la puissance moyenne du signal à bande ultra large correspondant sera également faible dû au faible rapport cyclique en résultant [20].

De nombreuses autres différences peuvent découler de celles présentées. Les différences qui viennent d'être exposées ici sont, à notre avis, les plus significatives.

Quoique les signaux à spectre étalé et les signaux à bande ultra large soient fondamentalement différents, les systèmes utilisant l'ULB se basent sur les mêmes concepts pour la transmission et la réception que ceux utilisés pour les systèmes à spectre étalé sauf pour la génération des formes d'ondes en soi. Ainsi, grâce à la revue des techniques utilisées dans les systèmes à spectre étalé présentée à la section précédente, il devient plus facile d'appliquer ces méthodes lors de la conception théorique de l'émetteur-récepteur utilisant l'ULB dont la description est donnée au chapitre suivant.

3. Méthodologie de travail et résultats

À la suite de la recherche effectuée sur les systèmes utilisant l'ULB et sur les systèmes à spectre étalé, une conception théorique et générale d'un système émetteur-récepteur utilisant l'ULB a été amorcée. Cette conception théorique a pour objectif de permettre de bien cerner les concepts impliqués dans l'élaboration d'un émetteur-récepteur et la façon dont ils seront appliqués avant de s'engager dans l'implémentation du système. Ce chapitre présente les étapes de conception du système à bande ultra large et certains résultats obtenus. La première section décrit les raisons qui ont motivé le choix de l'outil de conception utilisé. On énumère ensuite les principales hypothèses posées permettant de simplifier le problème. On explique par la suite la façon dont on pourra concevoir l'interface avec les ports série des deux stations de travail. Enfin, on expose la manière dont le système a été conçu et on donne quelques résultats intéressants.

3.1 Choix de l'outil de conception

Le choix d'un outil de conception pour l'élaboration théorique de l'émetteur-récepteur utilisant l'ULB n'est pas chose facile. En effet, le système à concevoir nécessite l'utilisation de composants employés dans différents domaines, tels l'électronique, la physique, les micro-ondes ainsi que la conception logicielle et algorithmique. Sur les recommandations du codirecteur du projet, le logiciel de simulation appelé *Extend* [34] a été choisi afin d'effectuer une conception théorique avant l'implémentation. Cet outil permet de modéliser des systèmes complexes utilisés dans des secteurs aussi variés que les finances, le génie, les chaînes de production, la biologie, etc.

Extend est un logiciel facile à utiliser qui possède une interface attrayante, tel que montré à la *figure 3.1*. Afin de simuler un système, l'utilisateur doit générer un document appelé *modèle* comprenant des blocs et des connexions entre ces blocs. *Extend* inclut différentes librairies de blocs, telles les librairies d'outils de création de graphiques (*plotters*), de composants électroniques usuels, de filtres, d'opérations mathématiques, d'interfaces avec d'autres programmes, etc. Chaque élément d'une librairie peut être inséré sur le modèle, copié, déplacé et modifié. La fonction accomplie par chacun de ces

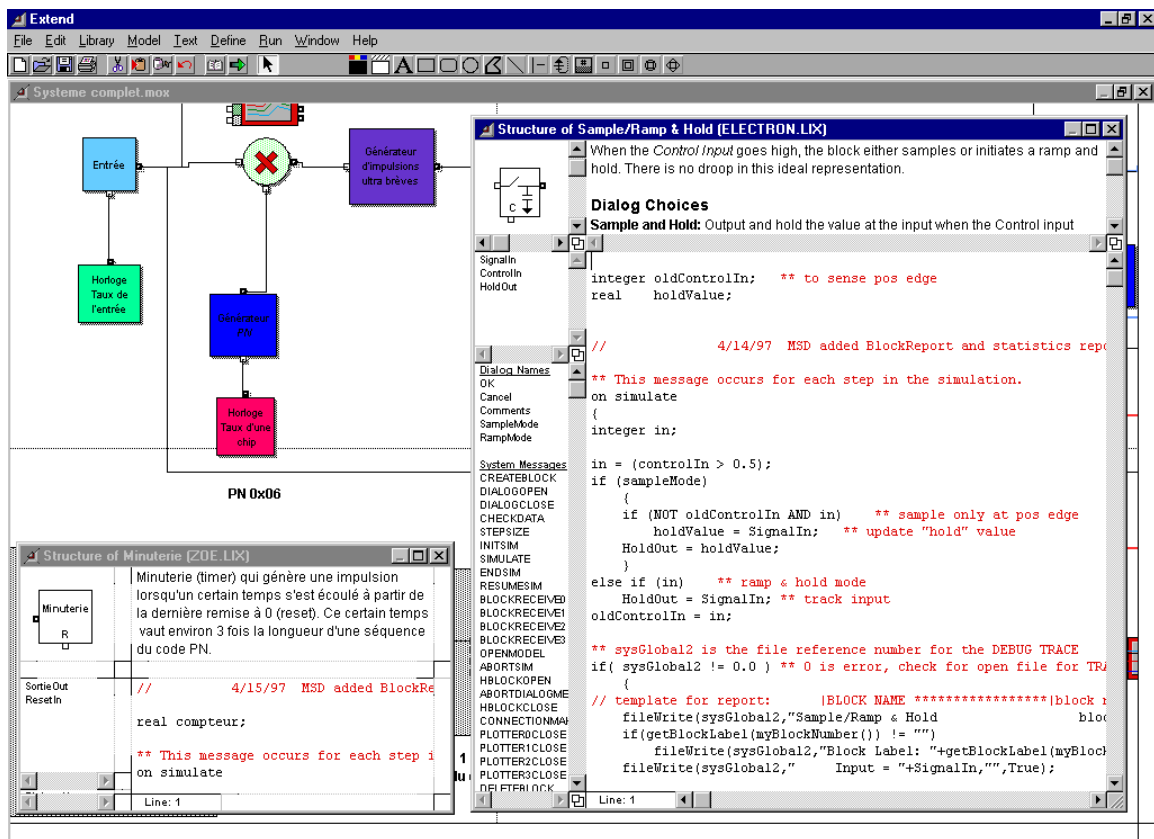


Figure 3.1 : Interface du logiciel de simulation *Extend*

blocs est programmée dans un langage semblable au C appelé *ModL*. Sur la *figure 3.1*, on voit l'intérieur du bloc *Sample/Ramp & Hold* où le code est inscrit. Cette fenêtre apparaît en cliquant deux fois sur l'icône du bloc dans le modèle *Extend*. L'utilisateur peut modifier le code, construire ses propres blocs et bibliothèques (comme le bloc *Minuterie* de la bibliothèque *zoe.lix* qu'on voit sur la *figure 3.1*), programmer lui-même ses blocs, regrouper plusieurs blocs ensemble pour créer de la hiérarchie dans le système, changer les paramètres des blocs et de la simulation, etc. Il peut également se servir d'une fenêtre appelée *notebook*, qu'on traduira *carnet*, pour mettre les paramètres de divers éléments de la simulation qu'il veut changer ainsi que le résultat des simulations en un seul endroit pratique.

Un autre avantage de cet outil est la possibilité pour d'autres utilisateurs de simuler le modèle, grâce à une version gratuite de *Extend*, qui leur permet de changer des paramètres si désiré sans toutefois pouvoir modifier le modèle. *Extend* est donc un logiciel qui convient bien à la conception théorique du présent projet.

3.2 Hypothèses simplificatrices

Afin de simplifier le problème lors d'une première approche, des hypothèses de travail ont été posées. Elles ont permis d'entreprendre la conception de l'émetteur-récepteur telle que présentée dans ce chapitre. Par la suite, le modèle pourra être amélioré et raffiné en laissant tomber certaines de ces hypothèses. Les principales hypothèses sur lesquelles se base le travail effectué sont donc les suivantes :

- 1) La communication entre les ports série est unidirectionnelle. Le port série d'un ordinateur est ainsi relié à l'émetteur alors que le port série du deuxième ordinateur est branché au récepteur.
- 2) Il n'existe qu'une seule paire de stations de travail utilisant l'émetteur-récepteur à concevoir, c'est-à-dire qu'il n'y a qu'un seul usager dans le système.
- 3) Le taux de transmission des données série est fixe et vaut 115 200 bauds.
- 4) La probabilité de transmission d'un 1 par l'ordinateur émetteur est égale à la probabilité d'émettre un 0. Ainsi, $P(1) = P(0) = 0,5$.
- 5) Les stations de travail sont assez près l'une de l'autre et sont situées dans un bâtiment, à l'intérieur.
- 6) Le signal résultant de la superposition des signaux à spectre étalé et des signaux à bande étroite qui peuvent être émis près du système émetteur-récepteur est considéré comme du bruit blanc gaussien théorique.

D'autres simplifications plus spécifiques seront introduites dans les prochains paragraphes au fur et à mesure du développement de la méthodologie de travail.

3.3 Interface avec les ports série

Comme la communication s'effectue entre les ports série de deux stations de travail de façon unidirectionnelle, on devra concevoir une interface entre un port série et l'émetteur d'un côté et entre le deuxième port série et le récepteur de l'autre. La présente section expose la façon dont ces interfaces pourraient être conçues. Une brève explication des fondements de la communication série est d'abord présentée. Les informations peuvent être transmises un bit à la fois d'un ordinateur à un autre sur un câble par l'intermédiaire de leur port série. Chaque ordinateur possède une puce appelée *émetteur-récepteur asynchrone universel* (*universal asynchronous receiver transmitter*

ou *UART*) qui lui permet de contrôler son port série et d'y envoyer de l'information grâce au standard RS-232C. Les connecteurs série possèdent 9 (DB-9) ou 25 (DB-25) broches. Chacun de ces connecteurs possède entre autres les broches suivantes : une broche pour la transmission des données série, une broche pour la réception de l'information et quatre broches pour les signaux de contrôle *Clear To Send*, *Data Terminal Ready*, *Data Set Ready* et *Data Carrier Detect*. Lorsque rien n'est transmis, la broche de transmission de l'information est maintenue au niveau logique 1. Le début de la transmission a lieu lorsqu'un 1 logique est présent sur les broches des quatre signaux de contrôle qui viennent d'être nommés. Le format standard des signaux de sortie d'un port série est entre +5 et +15 V pour un 0 logique et entre -5 et -15 V pour un 1 logique [35].

Les ordinateurs avec lesquels le *Groupe de réseautage* du CNRC travaille sont des stations de travail fonctionnant sous *Windows NT* qui possèdent des connecteurs DB-9. Les tensions des signaux de sortie de leurs ports série sont +12 et -12 V. L'interface à ces ports série pourra être réalisée grâce à deux modules ajoutés à l'émetteur-récepteur simulé sur *Extend*, tel que montré sur la *figure 3.2*. Pour l'instant,

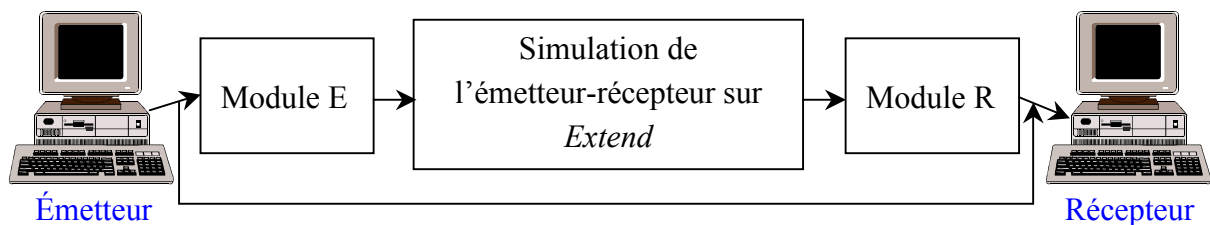


Figure 3.2 : Schéma général du système émetteur-récepteur

l'émetteur-récepteur simulé sur *Extend* ne transmet que les données série, c'est-à-dire que seul le fil reliant la broche de transmission de l'ordinateur émetteur et la broche de réception de l'ordinateur récepteur sera remplacé par un lien RF. Les signaux de contrôle seront toujours transmis par le câble habituel, tel que représenté sur la *figure 3.2* par la flèche qui passe sous les blocs des modules et de la simulation. Par contre, les modules E et R devront avoir accès à ces signaux. Comme les ports série des deux ordinateurs sont identiques et qu'il n'y a pas de modem entre ceux-ci, on devra connecter les broches de l'émetteur et celles du récepteur en utilisant un type de connexion appelé *null modem* [36] en anglais qui permet aux broches appropriées de communiquer ensemble.

Afin de permettre l'acquisition par le récepteur du signal émis, le module E devra transmettre une série de I pendant un temps assez long. Ce temps a été choisi égal à 10 périodes de code PN. Il devra ensuite envoyer une séquence de bits connue par l'émetteur et le récepteur afin de reconnaître le début de la transmission. Pour la simulation, la séquence 1011100 est envoyée. Lors de la conception réelle du système, il est évident qu'une séquence plus longue devrait être employée. Le module E devra également convertir les données du format +12 et -12 V au format binaire simple, possiblement à l'aide d'un circuit formé d'un amplificateur suiveur, d'un redresseur ainsi que d'un limiteur, et les ajouter à la suite de la série de I et de la séquence connue. La *figure 3.3* montre le format du signal à l'entrée de l'émetteur-récepteur simulé sur *Extend*. Le module R devra ensuite récupérer ces informations, reconnaître la séquence connue et transmettre seulement les données série au port série de l'ordinateur récepteur.

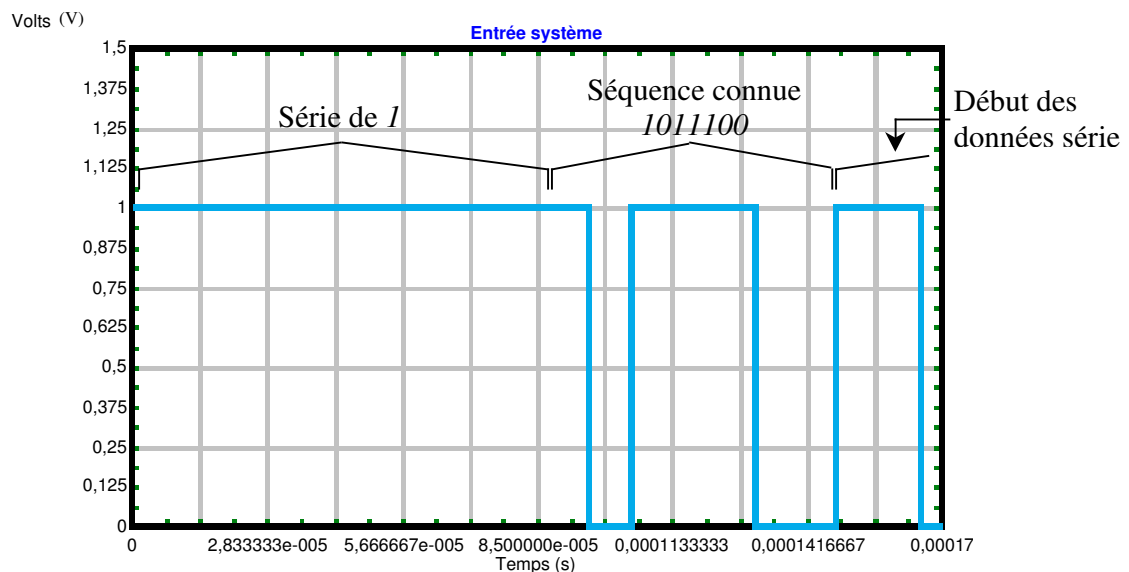


Figure 3.3 : Format du signal à l'entrée de l'émetteur-récepteur
(généré à partir du logiciel *Extend*)

Lors de la conception des modules E et R, il faudra peut-être tenir compte du temps pris par l'émetteur-récepteur utilisant l'ULB pour transmettre les données étant donné l'ajout de la série de I et de la séquence connue avant la transmission des informations réelles. Si ce temps supplémentaire par rapport à la transmission par câble série n'est pas négligeable, les modules devront comprendre des circuits permettant d'interagir avec les ordinateurs grâce aux signaux de contrôle.

3.4 Simulation de l'émetteur-récepteur théorique sur *Extend*

Le système simulé sur le logiciel *Extend* est composé d'un émetteur, d'un récepteur et d'un canal RF modélisé. En annexe, on retrouve l'ensemble des blocs et des sous-blocs qui ont été conçus et simulés sur *Extend*. À la page A1 de cet appendice, l'émetteur-récepteur entier est présenté. Des numéros sont associés à chacun des blocs qui ont été créés afin de mieux comprendre la hiérarchie qui existe entre ceux-ci. Les numéros de blocs utilisés dans le texte de la présente section font référence à cette numérotation. Dans l'annexe, les blocs non numérotés sont, pour la plupart, des blocs de premier niveau, c'est-à-dire qu'ils ne contiennent que du code. Durant la conception du système, nous avons conçu de nombreux blocs dont celui de l'outil de création de graphiques permettant de représenter quatre traces d'ondes carrées une au-dessus de l'autre et modifié le code de la plupart des blocs d'*Extend* utilisés. Nous n'avons toutefois pas cru bon de donner le détail du code *ModL* de chacun de ces blocs afin de ne pas alourdir le présent rapport inutilement. À la page A7 de l'annexe, on présente le carnet associé au modèle du système conçu. Ce carnet contient certains paramètres de la simulation qui peuvent être modifiés directement et des graphiques qui sont mis à jour au fur et à mesure de la progression de chaque simulation permettant de voir rapidement certains résultats.

Dans la présente section, la façon dont l'émetteur-récepteur a été simulé est décrite. On expose d'abord le principe général qui a guidé la conception du système. On présente ensuite les trois parties du modèle : l'émetteur, le canal modélisé et le récepteur. Enfin, on aborde la transmission des monocycles en donnant quelques résultats de calculs propres au système simulé. De la sous-section 3.4.2 *Émetteur* jusqu'à la fin du chapitre, tous les graphiques présentés ont été générés par *Extend* à partir du modèle simulé.

3.4.1 Principe général

Plusieurs choix ont été effectués pour la conception du système présenté à la page *A1* de l'annexe¹² afin de simplifier le problème. D'abord, il a été décidé de concevoir un système utilisant l'ULB basé sur le principe de séquence directe avec modulation antipodale. Ainsi, les bits¹³ de données série, convertis en format antipodal, sont multipliés par un code *PN* antipodal. Pour chaque chip du code, une impulsion sera émise. Si la chip vaut 1 , cette impulsion commencera par un pic positif, alors qu'elle commencera par un pic négatif si la chip correspondante vaut -1 , un peu de la même façon que pour le système présenté par *Æther Wire*. Les impulsions simulées sont des monocycles tels que décrits par *Time Domain* puisqu'une équation simple permet de les modéliser.

La *figure 3.4* montre le principe de transmission de monocycles utilisé dans le système simulé mais pour un code *PN* simplifié ayant un taux de transmission trois fois plus élevé que celui des informations série. Cette figure n'est pas à l'échelle puisque la durée d'une impulsion est de beaucoup inférieure à la période de répétition des monocycles, ce qui implique un faible rapport cyclique. Le troisième signal, celui des impulsions, est obtenu du signal (non montré sur la figure) issu de la multiplication des données série et du code *PN*. La *figure 3.4* montre bien que l'onde carrée représentant l'information est convertie en un train d'impulsions avant la transmission. Ainsi, le bloc *Générateur d'impulsions ultra brèves* du système illustré à la page *A1* aura comme rôle cette conversion et le bloc *Récepteur d'impulsions ultra brèves* devra détecter les monocycles et générer une onde carrée positive ou négative selon le cas pour la suite du processus de réception.

¹² À partir de maintenant, l'expression *de l'annexe* à la suite d'un numéro de page d'annexe commençant par *A* sera sous-entendue.

¹³ Le mot *bit* sera dorénavant utilisé comme synonyme du terme *symbole* pour les éléments de données série même si le taux de transmission de ces informations est donné en bauds. En effet, nous ne sommes pas intéressés à connaître la signification des symboles (s'ils représentent un bit ou plusieurs bits) transmis par l'ordinateur émetteur puisque le système se contente de transmettre l'information. On dira donc que le débit binaire des données série R_b est de 115 200 bits/s (bps).

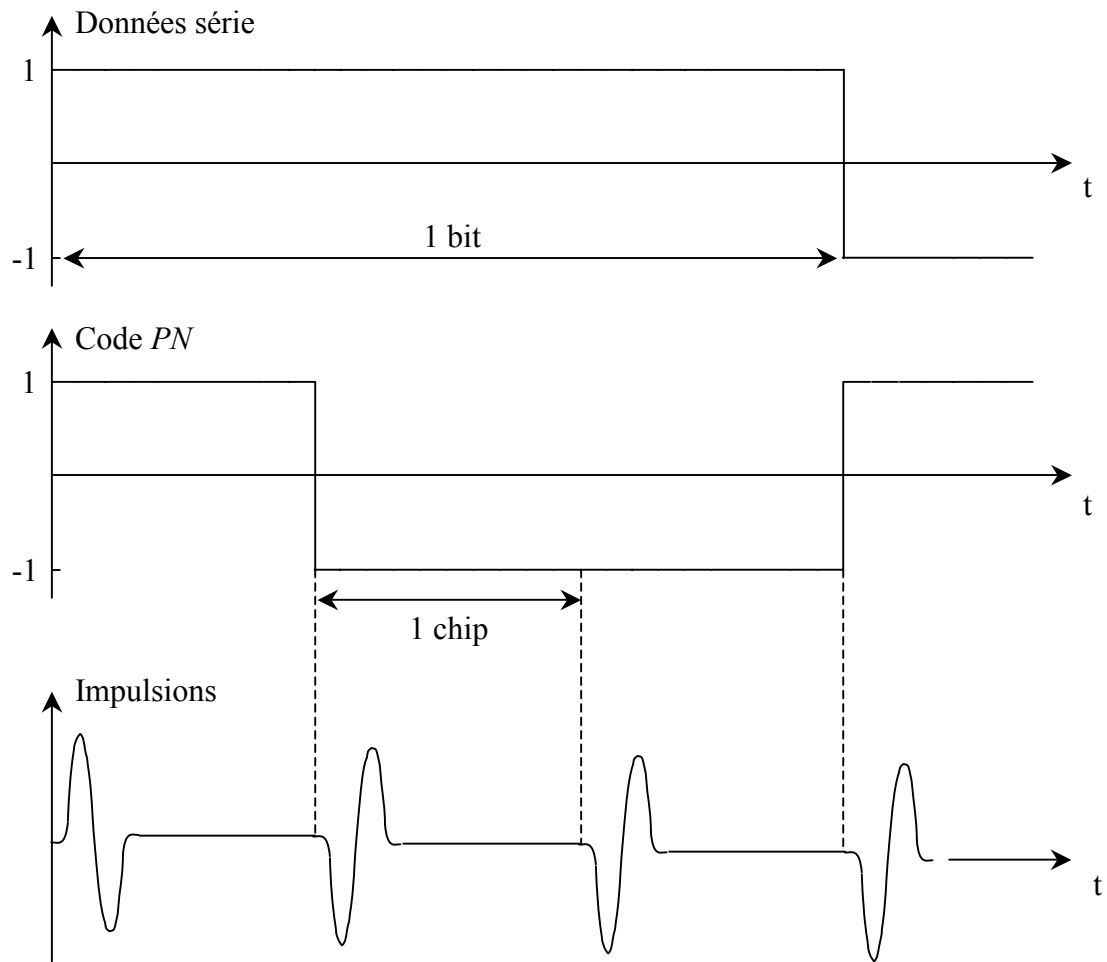


Figure 3.4 : Principe de transmission de monocycles pour chaque chip de code *PN*

On note également que le train d'impulsions ainsi généré est constitué de monocycles espacés régulièrement de T_c , la durée d'une chip du code *PN*. On a vu, à la sous-section 2.1.3 *Vue d'ensemble des travaux de deux compagnies*, que l'utilisation de monocycles uniformément séparés possédait deux désavantages. Par contre, selon [37], dans le cas d'impulsions régulièrement espacées et modulées de façon antipodale par un code pseudo-aléatoire, les raies spectrales sont éliminées. La technique qui vient d'être décrite est donc utilisée ici.

Ensuite, afin de simplifier la conception du récepteur, une autre décision a été prise. Contrairement à ce qui est montré à la *figure 2.9*, dans le système simulé, une période L de code *PN* correspond exactement à la durée d'un seul bit de données série,

c'est-à-dire que $L_c = \frac{T_b}{T_c} = L$. Ce choix, avec le fait qu'on ait choisi de générer une

impulsion par chip, implique que L impulsions sont émises pour chaque bit d'information série ($N_s = L$). Comme il a été mentionné à la sous-section 2.1.4 *Avantages et inconvénients de l'ULB*, le débit d'information sera alors limité si on désire obtenir une puissance moyenne faible, puissance qui est fixée par le rapport cyclique du signal émis. Cette contrainte n'est pas critique puisque le débit des données série, qui vaut 115 200 bps, est relativement faible. En outre, en transmettant plusieurs impulsions par bit, pour un taux d'erreur binaire (*bit error rate* ou *BER*) donné relié à la probabilité d'erreur par bit, on peut diminuer la puissance de chaque monocycle, et ainsi la puissance moyenne du signal, par rapport à la transmission d'une seule impulsion par bit [20]. On peut en effet récupérer l'information même si certains monocycles ne sont pas reçus.

3.4.2 Émetteur

L'émetteur est basé sur celui de la *figure 2.8* où un code *PN* multiplie l'entrée. Le signal d'entrée, montré à la *figure 3.3*, est généré par le bloc 1 présenté à la page A2. Les deux boîtes multicolores sont des outils de création de graphiques qui servent d'entrée à un bloc addition dont la sortie est le signal à transmettre. Les deux boîtes $\pm t$ bleu-gris sont des blocs retardant le signal appliqué à leur entrée d'un délai par rapport au temps $t = 0$ s du début de la simulation, délai pouvant être modifié. Les données série, qui sont les informations provenant de l'ordinateur émetteur et que l'on désire transmettre, sont également simulées sur *Extend* grâce au bloc 12 à l'intérieur du bloc 1. Pour permettre de tester le système plus facilement, nous avons généré des données cycliques. La séquence périodique de l'information série est *11010*. Le bloc 12 présenté à la page A6 montre le registre à décalage qui a été utilisé pour obtenir la séquence désirée. Le bloc rose du générateur de données série émet une impulsion en début de simulation pour donner les valeurs initiales aux bascules.

Une séquence maximale de période L égale à 7 chips a été choisie comme code *PN* afin de faciliter la simulation et de déboguer plus aisément le système conçu. Le logiciel *PN SIM* [38], qui permet de générer, d'afficher et de corrélérer différents codes *PN* (séquences maximales, codes Gold, etc.) entre eux, montre qu'il existe deux

configurations de registres à décalage permettant d'obtenir des séquences maximales de période 7 : R3-0x06 et R3-0x05. La première configuration a été utilisée. Les valeurs initiales données aux trois bascules sont *111* afin de reconnaître facilement le début de la séquence lors de tests. La façon dont la séquence maximale est générée par le bloc 4, le *Générateur PN* du système simulé, est montrée à la page A3. La séquence maximale produite par le registre à décalage utilisé est *1110010*. La fonction d'autocorrélation de cette séquence maximale est la même que celle présentée à la *figure 2.12* en remplaçant $2^n - 1$ par sa valeur qui est $L = 7$.

L'horloge commandant les données série (bloc 2, page A2) a une fréquence égale au taux de transmission des informations série R_b de l'ordinateur émetteur, soit 115 200 Hz. Comme une période du code *PN* doit entrer dans un bit de données, le débit R_c du code *PN*, aussi appelé *taux d'une chip* dans le modèle *Extend*, doit être 7 fois plus élevé que celui des données. Il vaut ainsi 806 400 chips/s (cps). La fréquence de l'horloge commandant les bascules du code *PN* (bloc 5, page A4) a donc été fixée à 806 400 Hz.

Le signal d'entrée et le code *PN* sont ensuite convertis en format antipodal et multipliés ensemble grâce au bloc multiplicateur (bloc 3) présenté à la page A3. Le signal ainsi généré est enfin transmis. La *figure 3.5* montre un graphique présentant les signaux d'intérêt à l'émetteur. On peut remarquer que le signal transmis est antipodal contrairement aux trois autres signaux qui sont en binaire simple.

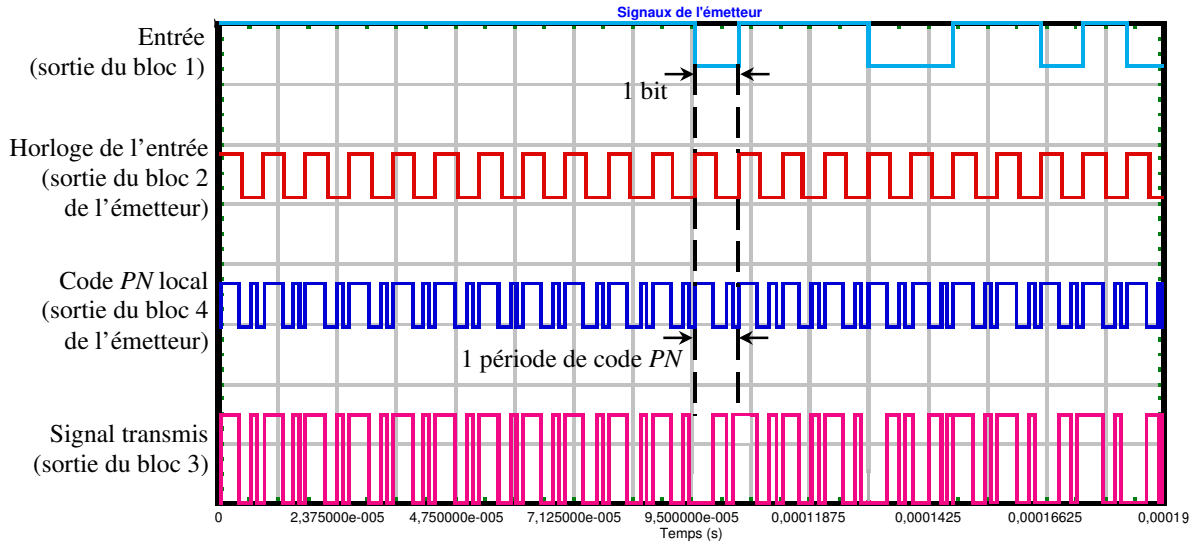


Figure 3.5 : Signaux à l'émetteur

3.4.3 Modélisation du canal

Le canal dans lequel le signal est transmis est un lien RF sans fil. Un modèle simple et classique pour les canaux de communication radio est celui où le signal émis est atténué et ajouté à du bruit [24] [27]. Selon ce modèle, le signal reçu $r(t)$ a la forme suivante :

$$r(t) = (1-\alpha) s(t) + n(t), \quad (3.1)$$

où α est le facteur d'atténuation, $s(t)$ est le signal émis et $n(t)$ est le signal bruit. Selon l'hypothèse 5 de la section 3.2 *Hypothèses simplificatrices*, les ordinateurs sont placés à une faible distance l'un de l'autre. Nous avons donc négligé l'atténuation lors de la modélisation du canal, ce qui implique que α vaut 0. Nous avons également négligé la dispersion en fréquence du signal qui peut faire partie d'une modélisation plus complexe du canal. Le bruit $n(t)$ est composé du bruit thermique, bruit dû aux composantes électroniques, et du bruit causé par les autres signaux de l'environnement. Le bruit thermique peut être considéré comme du bruit gaussien [24]. De plus, d'après l'hypothèse 6, les autres signaux sont aussi équivalents à du bruit gaussien. Ainsi, $n(t)$

peut être modélisé comme un bruit blanc gaussien caractérisé par sa moyenne nulle et par sa densité spectrale de puissance $\frac{N_0}{2}$ ou par sa variance σ^2 :

$$n(t) = N\left(0, \frac{N_0}{2}\right), \quad (3.2)$$

où $\frac{N_0}{2} = \sigma^2$.

Le bloc 6 montré à la page A4 est le canal modélisé par l'équation :

$$r(t) = s(t) + N\left(0, \frac{N_0}{2}\right). \quad (3.3)$$

La boîte sur laquelle est inscrit *Gaussian* est un bloc faisant partie d'une librairie de *Extend* qui génère un bruit gaussien dont l'utilisateur a spécifié la moyenne et l'écart type σ . Dans le carnet associé au modèle sur *Extend* montré en page A7, on peut modifier directement ces paramètres de bruit sans avoir à ouvrir le bloc 6 sur le modèle.

En plus du bruit, une boîte de délai est ajoutée au canal modélisé. Ce bloc permet de simuler le retard du signal reçu ou tout déphasage de ce signal par rapport au code *PN* local du récepteur. On appellera ce délai τ_p . De la même façon que pour les paramètres du bruit, on peut faire varier ce délai directement à partir du carnet.

3.4.4 Récepteur

Le récepteur simulé comporte deux modules principaux : le module démodulation¹⁴ et le module acquisition. Le module poursuite, partie intégrante d'un récepteur à spectre étalé, n'a pas été simulé, car nous prenons pour acquis que les informations transmises (ou les trames si, éventuellement, l'information est séparée en trames) sont de courte durée et que la dérive des horloges est négligeable pendant ce bref laps de temps. De plus, comme on le verra plus loin dans cette sous-section, l'acquisition simulée permet de synchroniser le code *PN* local avec le signal reçu à plus ou moins une

¹⁴ Nous utilisons le mot *démodulation* ici même si le signal transmis ne module pas une porteuse à cause de la similitude de ce module avec celui de la démodulation dans les systèmes à spectre étalé.

chip près, ce qui répond amplement aux besoins du système simulé actuel et ce qui rend inutile, pour l'instant, l'utilisation d'un module poursuite.

Le module démodulation se base sur le principe présenté à la *figure 2.14* à la sous-section 2.2.3 *Réception de signaux à étalement du spectre en séquence directe*. Le multiplicateur du récepteur, bloc 7 montré à la page A4, multiplie le signal reçu antipodal avec le code *PN* local, généré par le bloc 4 du récepteur, converti en signal antipodal. Ces deux blocs sont communs aux deux modules du récepteur. Lorsque l'acquisition est terminée, la démodulation peut commencer. Le bloc 10, présenté en page A6, permet l'intégration et la remise à zéro ainsi que l'échantillonnage du résultat de la multiplication pour le module démodulation. L'intégration se fait sur toute la durée T_b d'un bit d'information, la remise à zéro a lieu à T_b^+ grâce à l'utilisation d'un multivibrateur monostable et les échantillons utilisés pour la décision sont pris à tous les T_b . Le signal est échantillonné grâce à un bloc appelé *Sample & Hold*. Le comparateur (bloc 11, page A6) permet de décider, à partir de ces échantillons, si l'information envoyée est un 1 ou un 0 à chaque bit. Selon l'hypothèse 4 de la section 3.2 *Hypothèses simplificatrices*, $P(1) = P(0) = 0,5$. Ce n'est pas tout à fait le cas avec les données générées à l'entrée du système, mais l'approximation peut tout de même être utilisée. Comme le signal entrant dans le bloc 11 est antipodal avec $P(1) = P(0) = 0,5$, le seuil optimal du comparateur est 0 [32], valeur à laquelle on l'a fixé. Ainsi, si la valeur de l'échantillon à un multiple de T_b secondes vaut plus que 0, alors un 1 sera émis tandis qu'on décidera qu'un 0 a été transmis si la valeur de cet échantillon est inférieure à 0.

La *figure 3.6* présente deux graphiques qui permettent de mieux comprendre le fonctionnement du démodulateur. Ces traces ont été obtenues, comme on peut le voir sur le deuxième graphique, pour une entrée comprenant seulement les données série (11010 cyclique) et sans délai τ_p entre le signal reçu et le code *PN* local. Ce signal d'entrée, *Entrée série du système*, est également le signal qui se trouve à la sortie du multiplicateur du récepteur (sortie du bloc 7) lorsque le signal reçu est parfaitement synchronisé avec le code *PN* local, comme il a été expliqué à la sous-section 2.2.1 *Description générale du principe d'étalement du spectre en séquence directe*. Ce signal est donc, théoriquement, égal à celui qui est intégré. Comme on peut le remarquer sur la figure, la sortie du système, c'est-à-dire le signal récupéré, est retardée

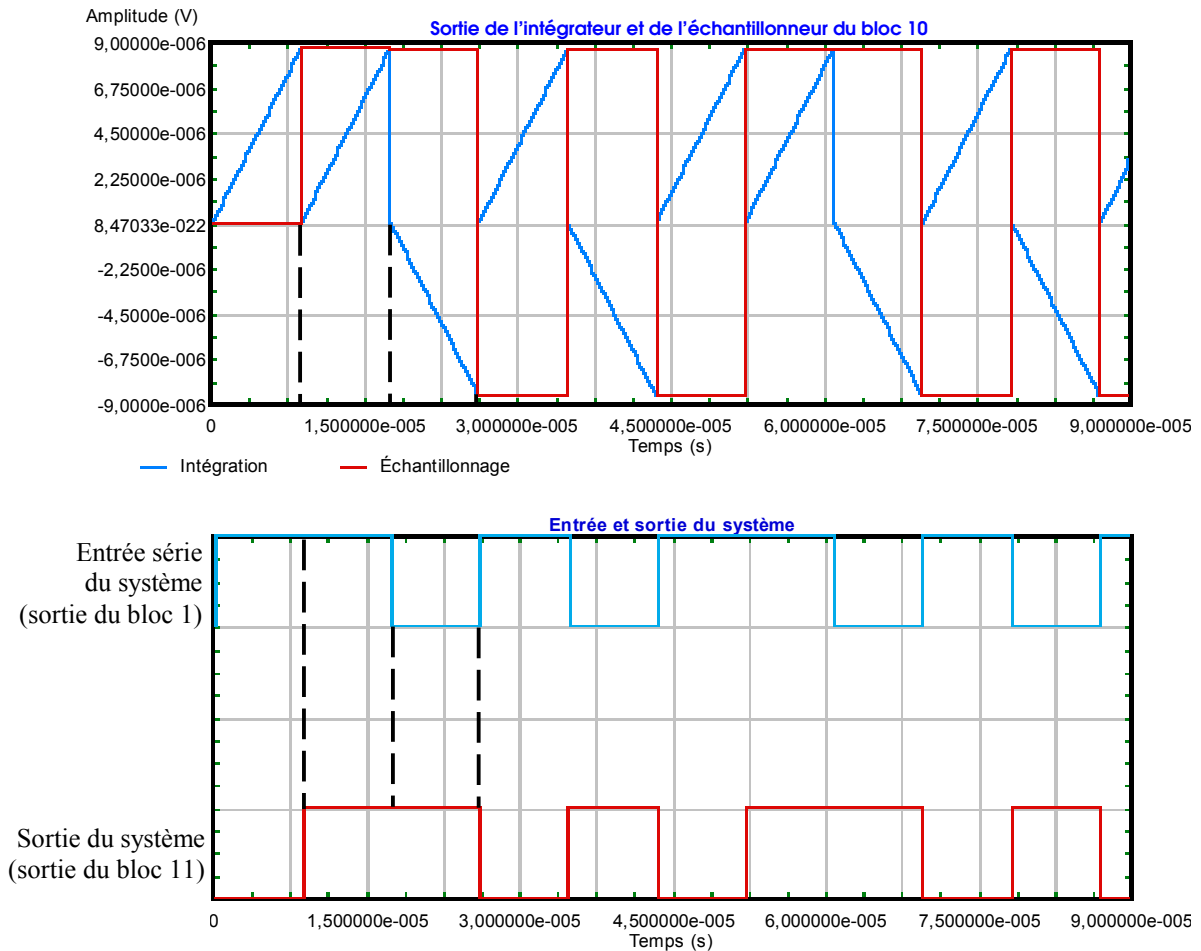


Figure 3.6 : Signaux permettant de mieux comprendre le mécanisme de démodulation

de $1 / 115\,200\text{ Hz} = 8,68\text{ }\mu\text{s}$, soit la durée T_b d'un bit de données. En effet, l'intégration d'un bit prend la durée totale de ce bit et la valeur de l'échantillon n'est disponible qu'à la fin du bit considéré.

L'acquisition du signal reçu est obtenue grâce à un corrélateur dynamique basé sur celui présenté à la *figure 2.16* de la sous-section 2.2.3 *Réception de signaux à étalement du spectre en séquence directe*. Comme le code *PN* utilisé est court, le temps τ_d d'intégration est pris égal à une période de ce code, soit 7 chips ou $\tau_d = T_b = 8,68\text{ }\mu\text{s}$. Le bloc 8 permet de mettre à jour la phase du code *PN* local en le retardant de $\tau_r = T_c$ secondes tant que le seuil du comparateur du bloc 9 n'est pas franchi. Le bloc 9, montré à la page A5, génère un 0 lorsque son signal d'entrée dépasse un certain seuil, c'est-à-dire que le signal reçu et le code *PN* local sont synchronisés à plus ou moins une chip ($1 \times T_c$) près. Il génère par contre un 1 lorsque le signal reçu n'est pas acquis et

qu'il faut alors décaler le code PN local grâce au bloc 8. Le seuil du comparateur du bloc 9 a été fixé à $3,5 \mu V$ par essais et erreurs, ce qui a suffi pour les besoins du modèle actuel. Au chapitre 4. *Discussion*, une méthode formelle de détermination du seuil sera brièvement présentée pour le cas de plusieurs usagers dans le système. À la figure 3.7, les deux graphiques montrent comment l'intégration du signal résultant de la multiplication est réalisée par rapport à la période du code PN local. Les signaux montrés sont ceux obtenus pour un système sans bruit et pour un délai τ_p de $3,72 \mu s$, soit environ 3 chips de

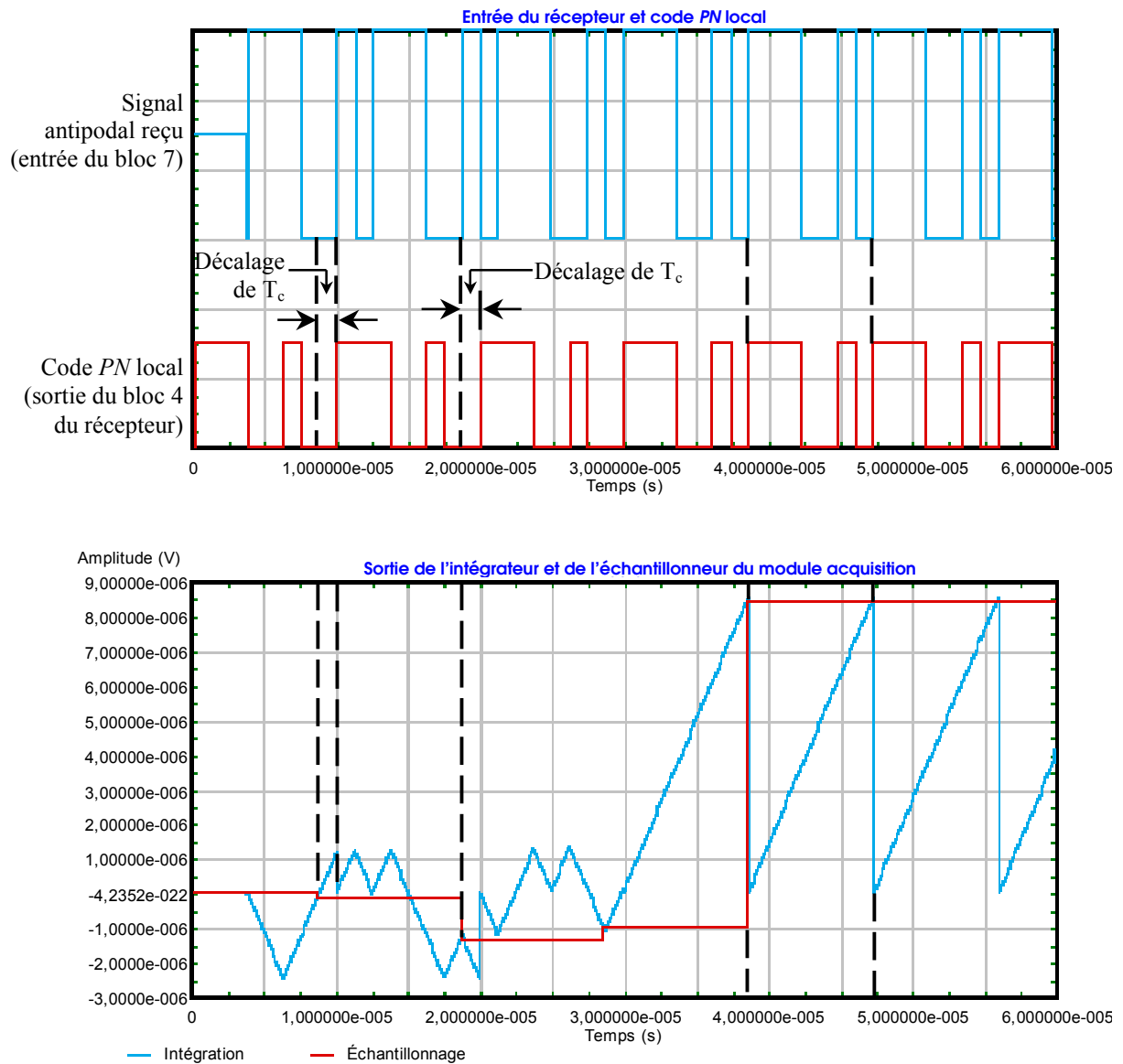


Figure 3.7 : Signaux du module acquisition permettant de mieux saisir le processus d'intégration

code. On remarque, sur cette figure, que le module acquisition intègre une période complète de chaque séquence de code PN local. On note également que la valeur de l'échantillon est prise à la fin de chaque période de code PN local. Pendant le décalage de T_c , il n'y a donc ni intégration ni prise d'échantillon.

Comme la génération des différents signaux de commande par le bloc 8 montré à la page 45 est assez complexe, seule une description du fonctionnement global de ce bloc sera donnée ici. La boîte appelée $\pm t$ spécial bleue est un bloc retardant de T_c le signal d'horloge y entrant lors d'une transition montante de la sortie du commutateur du bloc 8. La sortie de cette boîte sert d'horloge, dont certains créneaux semblent manquer, au générateur PN du récepteur. Lorsque la sortie du bloc 9 reste à 0 pendant plus de 3 périodes du code PN , on peut dire que le signal reçu est acquis à plus ou moins une chip près. Le bloc minuterie passe alors de 0 à 1 pour faire commuter l'interrupteur (ou commutateur) du bloc 8 de façon permanente, ce qui a pour effet de déconnecter la boucle de rétroaction du module acquisition afin de désactiver ce module pour le reste de la transmission (simulation). De fait, une fois l'acquisition terminée, la boîte $\pm t$ spécial ne générera plus de délai et l'horloge du générateur PN sera identique à celle de l'émetteur. La figure 3.8 illustre ce qui vient d'être expliqué en présentant différents signaux associés à l'acquisition du signal pour un canal sans bruit avec un délai τ_p de 7 μ s, soit plus de 5 chips de code PN . Il est intéressant de noter que, dans ce cas, 6 décalages du code PN local d'une durée de 1 chip sont nécessaires avant que le signal reçu soit acquis à plus ou moins 1 chip près, c'est-à-dire avant que le seuil soit dépassé.

Le processus d'acquisition qui vient d'être décrit doit se faire au tout début de la réception du signal. En effet, comme il a été mentionné précédemment, au début de la transmission, une série de dix 1, dont la durée de chacun des 1 est égale à la durée d'un bit de données série, soit 8,68 μ s, est envoyée. Ces 1, multipliés par le code PN local à l'émetteur, donne le code PN local lui-même, ce qui rend possible l'acquisition. Une fois les données commencées (séquence connue et informations série), l'acquisition n'est plus possible avec la configuration actuelle du récepteur et ce module est donc désactivé. Le signal d'entrée doit par conséquent être acquis avant que la série de 1 ne soit terminée.

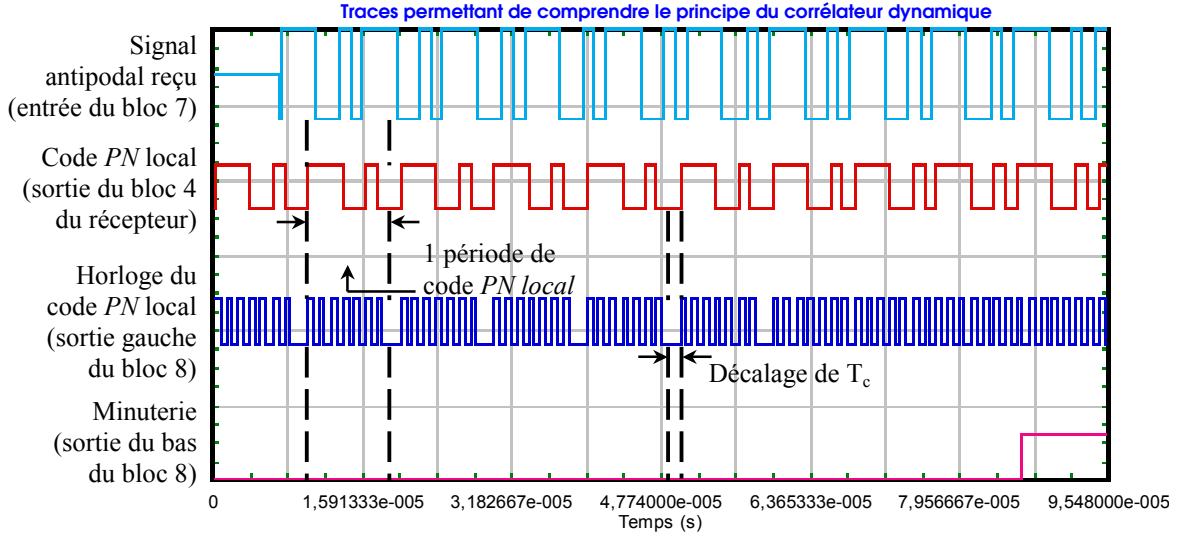


Figure 3.8 : Signaux du module acquisition permettant de comprendre le principe du corrélateur dynamique simulé

Le signal de minuterie, qui passe à 1 lorsque le signal est acquis, en plus de servir de déclencheur au commutateur du bloc 8, donne également le signal de départ, grâce au commutateur montré à la page A1, au module démodulation pour qu'il commence à démoduler le signal reçu. La figure 3.9 présente l'essentiel des signaux représentant le système pour un bruit blanc gaussien dans le canal ayant un écart type de 0,2 et pour un délai τ_p de 1,3 μ s, soit un peu plus d'une chip de code. La simulation est arrêtée à 200 μ s et, pendant ce laps de temps, 5 bits d'information série utile sont reçus. On peut remarquer, sur cette figure, que le signal de minuterie peut changer de valeur lors de l'entrée de la séquence connue et des données série réelles. Par contre, comme il a été mentionné dans un paragraphe précédent, seule la première montée à 1 a une influence sur le système. De plus, on note que l'ajout de bruit au signal n'empêche pas la parfaite démodulation du signal reçu. En effet, l'intégrale du bruit blanc gaussien est nulle, car sa moyenne est nulle [32] :

$$E = \{n(t)\} = \frac{1}{T} \int_0^T n(t) dt = 0. \quad (3.4)$$

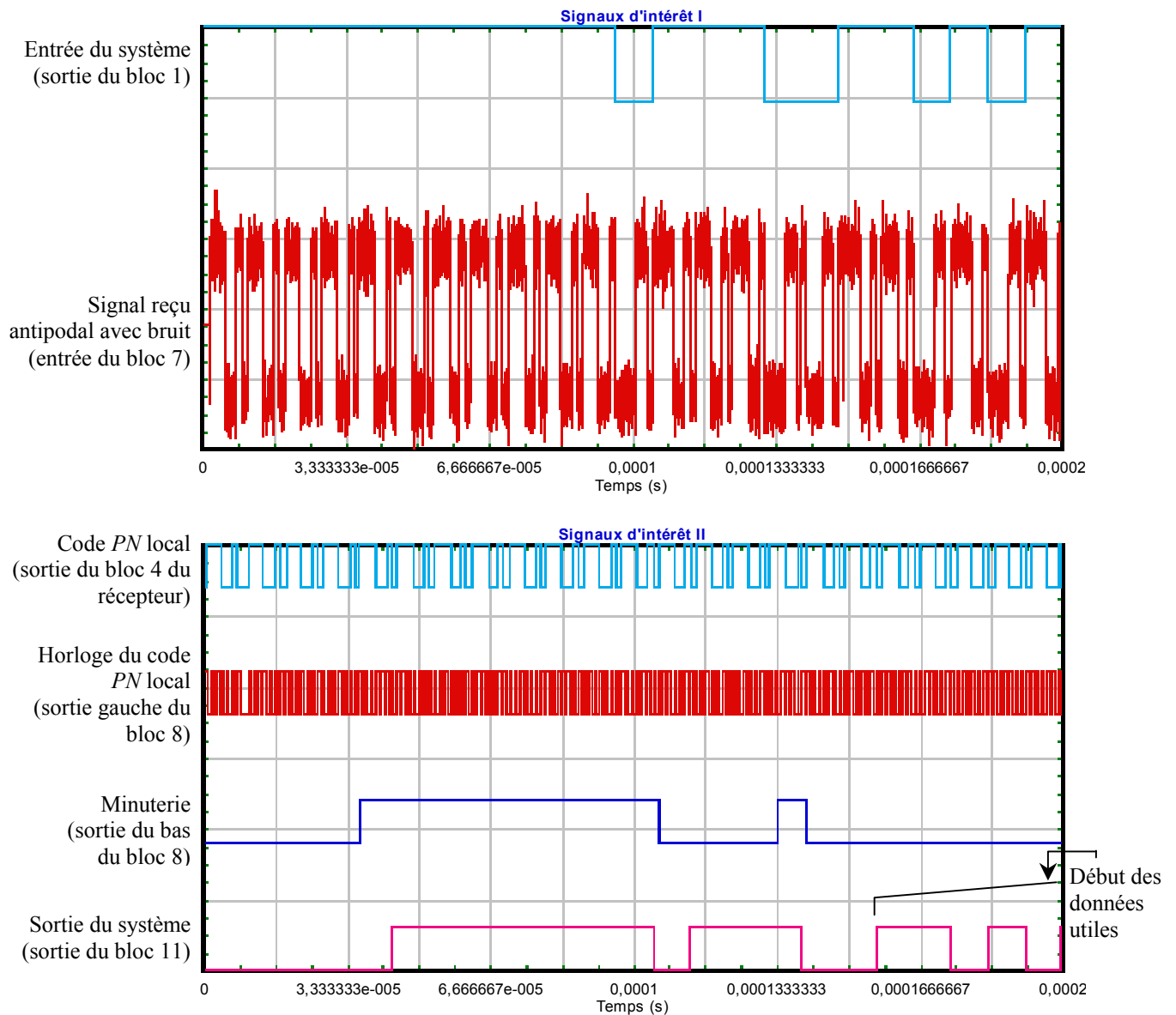


Figure 3.9 : Signaux d'intérêt pour le système simulé entier

3.4.5 Transmission de monocycles

Pour l'instant, les blocs *Générateur d'impulsions ultra brèves* et *Récepteur d'impulsions ultra brèves* n'ont pas été simulés. En effet, comme les monocycles ont des durées de l'ordre de la nanoseconde, le pas de la simulation (ou d'échantillonnage du signal) doit être de l'ordre de la picoseconde pour obtenir une bonne résolution. Puisqu'on utilise un rapport cyclique très faible, la simulation doit être exécutée pendant assez longtemps pour transmettre une certaine quantité d'information, ce qui implique

que le traçage de graphiques dans *Extend* prend énormément de points et donc de mémoire vive. Par exemple, envoyer les 5 premiers bits d'information prend environ 190 μ s en comptant les bits d'acquisition et ceux de la séquence connue. Ainsi, chaque graphique qui doit être tracé contiendra, pour ces seuls bits transmis, environ 190×10^6 points. Le travail avec le signal à bande ultra large a donc été effectué dans un fichier *Extend* différent de celui du modèle principal de la page A1.

La durée des monocycles générés à toutes les chips a été choisie égale à environ 1 ns puisqu'on a mis $T_m = 1$ ns dans l'équation 2.3. D'après l'équation 2.4, la fréquence centrale de chaque monocycle est de 1 GHz et leur largeur de bande vaut environ $116 \% \times 1 \text{ GHz} = 1,16 \text{ GHz}$. Comme la fréquence des chips est de 806 400 cps, la durée d'une chip est de $1 / 806\,400 \text{ Hz} = 1,24 \mu\text{s}$. Le rapport cyclique obtenu dans le cas considéré est donc $\frac{1 \text{ ns}}{1,24 \mu\text{s}} \times 100 = 0,08\%$.

On peut trouver la puissance requise par monocycle selon la probabilité de détection de chacun de ces monocycles par le bloc *Récepteur d'impulsions ultra brèves*. Cette probabilité de détection dépendra en partie de la structure du bloc récepteur de monocycles. Il faudra également tenir compte de l'effet du canal sur l'énergie des impulsions et du niveau de bruit. La puissance d'un monocycle, pour une résistance de charge supposée égale à 1 Ω , est donnée par la relation :

$$P_m = \frac{1}{T_m} \int_0^{T_m} m_o^2(t) dt, \quad (3.5)$$

où $m_o(t)$ est l'équation d'un monocycle (équation 2.3) et où T_m est la durée de l'impulsion qui, dans notre cas et selon une mesure sur *Extend*, vaut 1,32 ns. Cette intégrale, une fois résolue, donne une puissance P_m égale à $0,297 A^2$, où A est l'amplitude de crête du monocycle. Pour l'instant, comme les ordinateurs émetteur et récepteur sont placés à une distance relativement faible l'un par rapport à l'autre à l'intérieur d'un bâtiment selon l'hypothèse 5 de la section 3.2 *Hypothèses simplificatrices* et comme l'atténuation dans le canal est supposée nulle, on peut croire qu'une puissance de 5 W par monocycle transmis est suffisante pour permettre leur détection même dans un environnement bruité. Cette puissance posée, on peut trouver la puissance moyenne du signal émis,

soit $\frac{1\text{ns}}{1,24\mu\text{s}} \times 5\text{W} = 4\text{mW}$ dans notre cas, et la valeur de l'amplitude des monocycles,

soit $A = 4,1\text{ V}$ ($0,297A^2 = 5\text{ W} \Rightarrow A = 4,1\text{ V}$). Puisque chaque bit de données est représenté par 7 impulsions (7 chips par bit), la puissance d'un monocycle obtenue ici est 7 fois inférieure à la puissance qui serait requise si on envoyait une impulsion par bit d'information pour une probabilité d'erreur par bit donnée. La puissance moyenne du signal est elle aussi divisée par 7.

4. Discussion

D'après ce qui a été montré dans le chapitre précédent, on peut dire que le système simulé fonctionne bien dans le contexte considéré. Toutefois, le modèle utilisé est extrêmement simplifié. Il faut donc être prudent quant à l'interprétation de ce bon fonctionnement. Le présent chapitre permet d'apporter de nouveaux éléments au modèle actuel en décrivant, en premier lieu, les limites de ce modèle et les améliorations possibles au système. En second lieu, une des améliorations pouvant être apportées au modèle est traitée brièvement. Cette amélioration est la possibilité d'avoir plusieurs usagers dans le système, c'est-à-dire plusieurs paires d'ordinateurs communiquant entre eux simultanément. Enfin, dans la dernière section de ce chapitre, on émet des recommandations pour la poursuite de la conception du système.

4.1 Limites et améliorations du système

La limite qui semble la plus évidente dans le cas du modèle actuel est une limite d'application. Le modèle simulé sur *Extend* a permis de comprendre le fonctionnement du système et peut servir à obtenir des résultats sans avoir à réaliser l'émetteur-récepteur pratiquement. Il est par contre évident que le système ne serait pas implanté de la façon dont il a été conçu sur *Extend* à l'aide de blocs. On pourrait par exemple implémenter ce système dans un processeur de signaux numériques (*digital signal processor* ou *DSP*) et les composantes utilisées dans *Extend* n'auraient alors pas leur équivalent en pratique. Il est important de saisir que l'implémentation du système est en fait autrement plus complexe que ce que peut laisser paraître la simulation dans *Extend*, car elle met en jeu de nombreux concepts dans différents domaines.

Les autres limites sont reliées au fait que nous avons posé plusieurs hypothèses, qui sont en partie décrites tout au long du présent rapport, afin de simplifier la conception du système. De nombreuses améliorations liées à ces limites peuvent ainsi être apportées au système afin de le rendre plus réaliste ou plus performant. Ces améliorations, qui sortent du cadre de ce projet de fin d'études, ne sont données ici qu'à titre d'information.

On pourrait d'abord améliorer, à court terme, le modèle existant simulé sur *Extend* de plusieurs façons. On peut ajouter un bloc s'assurant que les registres à décalage ne

tombent pas à l'état 0 partout afin d'éviter qu'ils ne restent bloqués dans cet état, générer, au récepteur, les signaux d'horloge au taux des données série et au taux d'une chip à partir de la même horloge afin d'empêcher qu'ils ne soient désynchronisés, générer les données série à partir d'une source aléatoire plutôt qu'utiliser un registre à décalage, utiliser une séquence connue plus longue, comme mentionné précédemment, avant la transmission des données série, diminuer le temps de décalage τ_r de T_c à $\frac{1}{2}T_c$ dans la boucle de rétroaction du module acquisition, utiliser une technique d'acquisition plus complexe et plus efficace, ajouter le module poursuite et permettre que le système puisse reprendre la boucle d'acquisition si la synchronisation est perdue, ajouter plusieurs autres paramètres de simulation au carnet associé au modèle, etc.

En plus de ces améliorations à la simulation actuelle, des modifications plus poussées peuvent être apportées directement au modèle simulé ou lors de l'implantation réelle du système. Certaines de ces améliorations sont présentées dans les paragraphes qui suivent.

On pourrait utiliser un code *PN* de période largement supérieure à 7 afin de le faire ressembler le plus possible à un phénomène aléatoire. On devrait alors probablement faire correspondre une séquence périodique du code avec plus d'un bit de données, ce qui complexifie la conception du récepteur, mais peut permettre de diminuer la puissance moyenne du signal et de sécuriser davantage la communication.

En réalité, le taux de transmission des données série d'un ordinateur n'est pas fixe. Il faudrait donc permettre des débits binaires variables et avoir des horloges accordables à l'émetteur et au récepteur. De plus, en pratique, l'horloge commandant le code *PN* de l'émetteur pourra être générée directement à partir du signal d'horloge de référence pris sur l'*UART* de l'ordinateur considéré afin que le code et les données séries soient parfaitement synchronisés.

On pourrait rendre la communication entre les ports série des deux ordinateurs bidirectionnelle. Cette modification implique que chaque ordinateur serait doté d'un émetteur-récepteur à bande ultra large. Il faudrait alors utiliser une technique permettant d'empêcher les signaux de l'émetteur et du récepteur d'un ordinateur donné d'interférer entre eux. La technique choisie pour y arriver peut être l'entrelacement (*interleaving*) des

impulsions ou l'utilisation d'une période T_f de répétition des monocycles différente pour l'émetteur et pour le récepteur d'une station de travail donnée [18]. La communication bidirectionnelle nécessite également des processus complexes de prise de contact (*handshaking*) entre les ordinateurs s'ajoutant aux actuels signaux de contrôle des ports série.

Dans le modèle actuel, les deux ordinateurs communiquant entre eux ont été supposés être placés près l'un de l'autre (hypothèse 5 de la section 3.2 *Hypothèses simplificatrices*). On a donc pu négliger les effets dus à la grande largeur de bande des signaux comme la dispersion et la distorsion en fréquence. Il faudrait, par contre, en tenir compte dans la modélisation du canal en la raffinant davantage si on décide de séparer les stations de travail par une distance supérieure à quelques mètres. Des mesures devraient alors être prises pour modéliser, entre autres, l'atténuation de propagation à l'intérieur de bâtiments, comme il a été traité dans [39] pour des ondes sinusoïdales à 900 MHz, pour différentes fréquences. Dans le cas de la propagation des signaux à bande ultra large à l'extérieur de bâtiments, comme il pourrait être le cas si les liens série relient des ordinateurs portatifs, on devra tenir compte de l'atténuation de fréquences particulières dans l'atmosphère [3]. On devra également corriger ces effets en concevant des circuits compensateurs dans le système. Par exemple, pour compenser cette distorsion en fréquence, un dispositif de préaccentuation des fréquences atténuées dans le canal pourrait être utilisé à l'émetteur (utilisation d'égalisateurs) [3] [40]. De plus, dans le cas d'ordinateurs portatifs, on devra prendre en considération l'effet Doppler lors de la modélisation ou de la conception du système puisque les deux ordinateurs peuvent être en mouvement l'un par rapport à l'autre [3].

On pourrait aussi permettre à plusieurs paires d'ordinateurs de communiquer entre elles. Le cas d'un système à plusieurs usagers est traité brièvement dans la prochaine section. À la limite, si on désire que toutes les stations de travail puissent transmettre de l'information à toutes les autres, on devra concevoir un système beaucoup plus complexe incluant possiblement une station de base comme dans les systèmes de téléphonie cellulaire et un système de contrôle de la puissance des émetteurs-récepteurs pour contrer le *near-far problem* en AMRC.

4.2 Cas de plusieurs usagers

Le système conçu dans le cadre de ce projet de fin d'études permet la communication d'un ordinateur émetteur à un ordinateur récepteur. Dans un environnement où plusieurs stations de travail sont en opération simultanément comme c'est le cas des locaux du *Groupe de réseautage*, il serait plus commode de permettre la communication entre les ports série de divers ordinateurs. On se trouverait ainsi à avoir un système à plusieurs usagers ou à AMRC. La présente section donne un aperçu des principales modifications qui devraient être apportées au système actuel pour permettre à plusieurs ordinateurs de communiquer entre eux par paires.

Il faudrait d'abord utiliser un code *PN* différent. En effet, quoiqu'il existe une autre façon de générer des séquences maximales de longueur 7 (R3-0x05), l'utilisation de ces configurations à séquence maximale n'est pas appropriée dans le cas d'environnements à plusieurs utilisateurs, car leurs propriétés de corrélation croisée ne sont pas adéquates en plus de permettre peu d'usagers dans le système.

Afin de permettre la communication série de nombreuses stations de travail en même temps, on peut utiliser des codes Gold, dont on a déjà parlé à la section 2.2.2 *Génération et choix des codes pseudo-aléatoires*. Avec la configuration illustrée à la *figure 2.13*, il est possible d'obtenir $2^n + 1$ codes Gold différents de longueur $2^n - 1$ ayant de faibles corrélations croisées et de supporter de nombreux usagers [27]. Sur ces $2^n + 1$ codes, environ la moitié sont des codes équilibrés, la définition de ces codes ayant été donnée à la section 2.2.2. Les codes Gold nous permettent ainsi de contenir beaucoup plus d'usagers qu'en utilisant des séquences maximales seules. Un exemple de code qui peut être implanté dans un système comme le nôtre est celui utilisé et présenté dans [31]. Cette famille de codes Gold, dont la génération est expliquée dans [27], peut être produite grâce à la configuration de registres à décalage de la *figure 4.1*, où les séquences maximales générant cette famille ont été choisies avec soin. On peut remplacer les x de cette figure indifféremment par 1 ou par 0. Les codes ainsi générés ont une période de $2^{13} - 1 = 8191$ chips. Il deviendrait donc difficile, avec l'utilisation de ces codes et de l'ULB, de faire correspondre un bit de données à une période complète de code *PN*, car le taux de transmission de chips

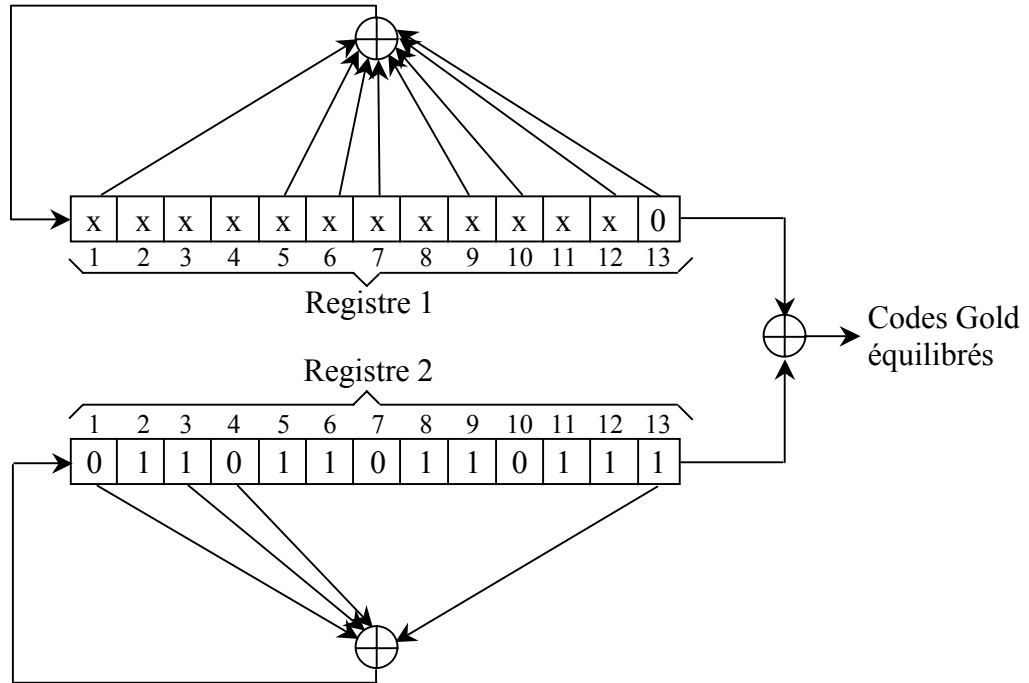


Figure 4.1 : Génération d'une famille de codes Gold équilibrés de période $2^{13} - 1$
(adaptation d'une figure de [31])

serait alors de 944 Mcps pour un débit binaire du port série de 115,2 kbps ($8191 \text{ chips/bit} \times (115,2 \times 10^3 \text{ bits/s}) \approx 944 \text{ Mcps}$). Il faudrait alors probablement coder une période de séquence pseudo-aléatoire sur plus d'un bit d'information pour être capable de générer les impulsions ultra brèves adéquatement.

Ensuite, la conversion des monocycles en ondes carrées au récepteur n'est plus possible dans le cas de plusieurs usagers dans le système. De fait, les impulsions de ces différents usagers sont toutes reçues au récepteur et ce dernier ne peut les distinguer. Il faut alors utiliser, au récepteur, un signal généré localement qui ressemble au signal reçu, c'est-à-dire qui est composé d'impulsions, plutôt que d'utiliser une onde carrée comme signal gabarit. Le récepteur peut ainsi corrélérer dynamiquement les impulsions reçues avec les impulsions locales codées selon le code *PN* correspondant à l'utilisateur propre à ce récepteur. Il doit par la suite vérifier si le résultat de la corrélation dépasse un certain seuil *S* indiquant que l'acquisition est terminée. Ce seuil est le seul outil qui permet de différencier les différents usagers du système : il doit donc être fixé avec soin. On doit, à ce moment, se fier aux mathématiques.

Le seuil S du corrélateur du module acquisition est déterminé par deux mesures de performance :

- la probabilité de fausse alarme P_{FA} (*false alarm probability*) : probabilité que le signal reçu soit déclaré acquis alors qu'il n'est pas dans la bonne phase par rapport au signal généré localement correspondant à l'utilisateur considéré;
- la probabilité de détection P_D (*detection probability*) : probabilité que, lorsque les signaux sont alignés dans la bonne phase, le signal reçu soit déclaré acquis, c'est-à-dire qu'il y ait une détection correcte.

Une ébauche de calcul théorique du seuil S est présentée ici pour le cas de deux usagers dans le système dont les signaux à bande ultra large sont supposés reçus avec une puissance égale sans atténuation. Le signal de l'utilisateur 1, $r_1(t)$, et le signal de l'utilisateur 2, $r_2(t)$, sont les signaux non bruités reçus par les récepteurs et correspondent à des impulsions émises à toutes les T_c secondes selon le code propre à chacun de ces usagers. On considère qu'on envoie, comme pour le système simulé, une série de I au début de la communication afin de permettre l'acquisition. Les récepteurs reçoivent le signal $r(t) = r_1(t) + r_2(t)$ et du bruit $n(t)$ qu'on présume blanc gaussien. Prenons également le cas général pour lequel le signal $r_2(t)$ est retardé d'un temps κ par rapport au signal $r_1(t)$. On a donc $r(t) = r_1(t) + r_2(t - \kappa)$ additionné de bruit. L'explication suivante est valide pour déterminer le seuil S du comparateur du module acquisition du récepteur de l'utilisateur 1. Pour le récepteur de l'utilisateur 2, un raisonnement semblable pourrait être suivi.

Le récepteur de l'utilisateur 1 effectue une corrélation croisée entre le signal reçu et le signal généré localement qui est théoriquement égal à $r_1(t)$:

$$\phi(\tau) = \int_0^{\tau_d} r_1(t) r(t - \tau) dt = \int_0^{\tau_d} r_1(t) (r_1(t - \tau) + r_2(t - \kappa - \tau)) dt. \quad (4.1)$$

On peut séparer cette intégrale en deux selon le signal considéré :

$$\phi(\tau) = \int_0^{\tau_d} r_1(t) r_1(t - \tau) dt + \int_0^{\tau_d} r_1(t) r_2(t - \kappa - \tau) dt = a(t) + i(t), \quad (4.2)$$

où $a(t)$ est l'autocorrélation du signal de l'utilisateur 1 et $i(t)$, l'intercorrélation du signal de l'utilisateur 1 avec le signal de l'utilisateur 2 décalé de κ , qu'on peut considérer comme de l'interférence. À partir de ces équations, on peut donner les expressions des probabilités

de fausse alarme P_{FA} et de détection P_D . Pour la probabilité de fausse alarme, on a :

$$P_{FA} = \Pr[(\max_2\{a(t)\} + \max\{i(t)\} + n(t)) > S], \quad (4.3)$$

où $\max_2\{a(t)\}$ est la valeur du pic secondaire d'autocorrélation de $r_1(t)$ et $\max\{i(t)\}$, la valeur maximale, en valeur absolue, de la corrélation croisée entre $r_1(t)$ et $r_2(t)$ pour tous les décalages τ . Pour la probabilité de détection, on a :

$$P_D \geq \Pr[(\max\{a(t)\} - \max\{i(t)\} + n(t)) > S], \quad (4.4)$$

où $\max\{a(t)\}$ est la valeur du pic principal d'autocorrélation de $r_1(t)$.

Puisqu'on connaît parfaitement les signaux $r_1(t)$ et $r_2(t)$, on peut facilement déterminer les valeurs de $\max\{a(t)\}$, $\max_2\{a(t)\}$ et $\max\{i(t)\}$. La figure 4.2 permet de mieux visualiser les probabilités des équations 4.3 et 4.4. Il est clair qu'on désire avoir

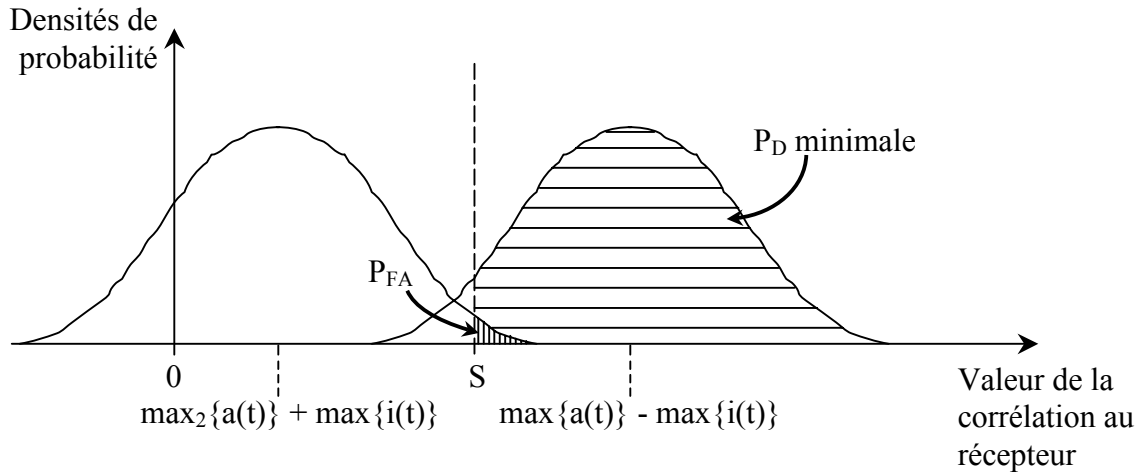


Figure 4.2 : Densités de probabilité associées aux probabilités de fausse alarme et de détection

une probabilité de fausse alarme la plus faible possible et une probabilité de détection la plus élevée possible. On doit donc choisir un seuil S qui permet d'obtenir le meilleur compromis entre ces deux probabilités pour le système considéré. Le même principe que celui qui vient d'être décrit peut être étendu au cas d'un système à plus de deux usagers. Certaines informations supplémentaires ou complémentaires sur le calcul des probabilités peuvent être trouvées dans [3], [25] et [41].

4.3 Recommandations

Maintenant que des améliorations possibles au système actuel viennent d'être proposées, nous pouvons formuler quelques recommandations pour la suite des travaux sur l'émetteur-récepteur à bande ultra large.

Il faudrait d'abord former une équipe composée de scientifiques de différents domaines qui pourraient chacun travailler sur une partie du système à concevoir. Les expertises minimales recherchées sont les suivantes : systèmes de télécommunications, théorie des codes pseudo-aléatoires, physique des impulsions brèves, antennes, traitement de signaux numériques ainsi que conception et implémentation de circuits analogiques et numériques.

Une fois cette équipe créée, les membres devraient convenir des caractéristiques et des propriétés de l'émetteur-récepteur à concevoir. Certaines pistes de réflexion pour y parvenir peuvent être tirées de la section *4.1 Limites et améliorations du système*.

Pour concevoir le système qu'elle aura bien défini, l'équipe aura besoin de matériel et d'équipements spécialisés. Notamment, pour être en mesure de visualiser les signaux à bande ultra large, il leur faudra un oscilloscope d'échantillonnage puissant capable d'échantillonner un signal à toutes les 50 picosecondes ou moins [42]. Les oscilloscopes de la série *TDS8000* de Tektronix [43], à titre d'exemple, permettent d'échantillonner les signaux à toutes les 10 femtosecondes (0,01 ps).

Conclusion

Le présent rapport a permis de présenter notre projet de fin d'études qui a porté sur la conception d'un émetteur-récepteur à bande ultra large. On y a exposé les grandes lignes de la théorie de la technologie à bande ultra large et des systèmes à spectre étalé que nous avons étudiée pour mener à bien ce projet. La façon dont le travail a été réalisé est expliquée et on y retrouve les résultats obtenus à partir du système conçu. On y a également proposé une discussion sur les limites et les améliorations possibles du système ainsi que des recommandations pour les travaux futurs.

Il ressort du présent document qu'un système émetteur-récepteur utilisant la technologie à bande ultra large permettant la communication unidirectionnelle entre les ports série de deux ordinateurs a été simulé sur le logiciel *Extend*. Le modèle est composé d'un émetteur, d'un canal RF modélisé et d'un récepteur séparé en module acquisition et module démodulation. Il permet bien, théoriquement, la transmission d'information série.

Cette simulation simple a permis de mieux comprendre les mécanismes de base mis en jeu dans la conception d'un système à bande ultra large. À partir de maintenant, le système pourra être amélioré afin de le rendre plus réaliste et plus performant. Comme il n'existe pas de standard pour la technologie à bande ultra large utilisée et comme, selon nos informations, aucune entreprise canadienne ne travaille sur ce sujet, l'équipe du CNRC qui pourrait être formée aurait tout intérêt à continuer le développement du système émetteur-récepteur afin de devenir le pionnier canadien dans ce domaine prometteur.

Bibliographie

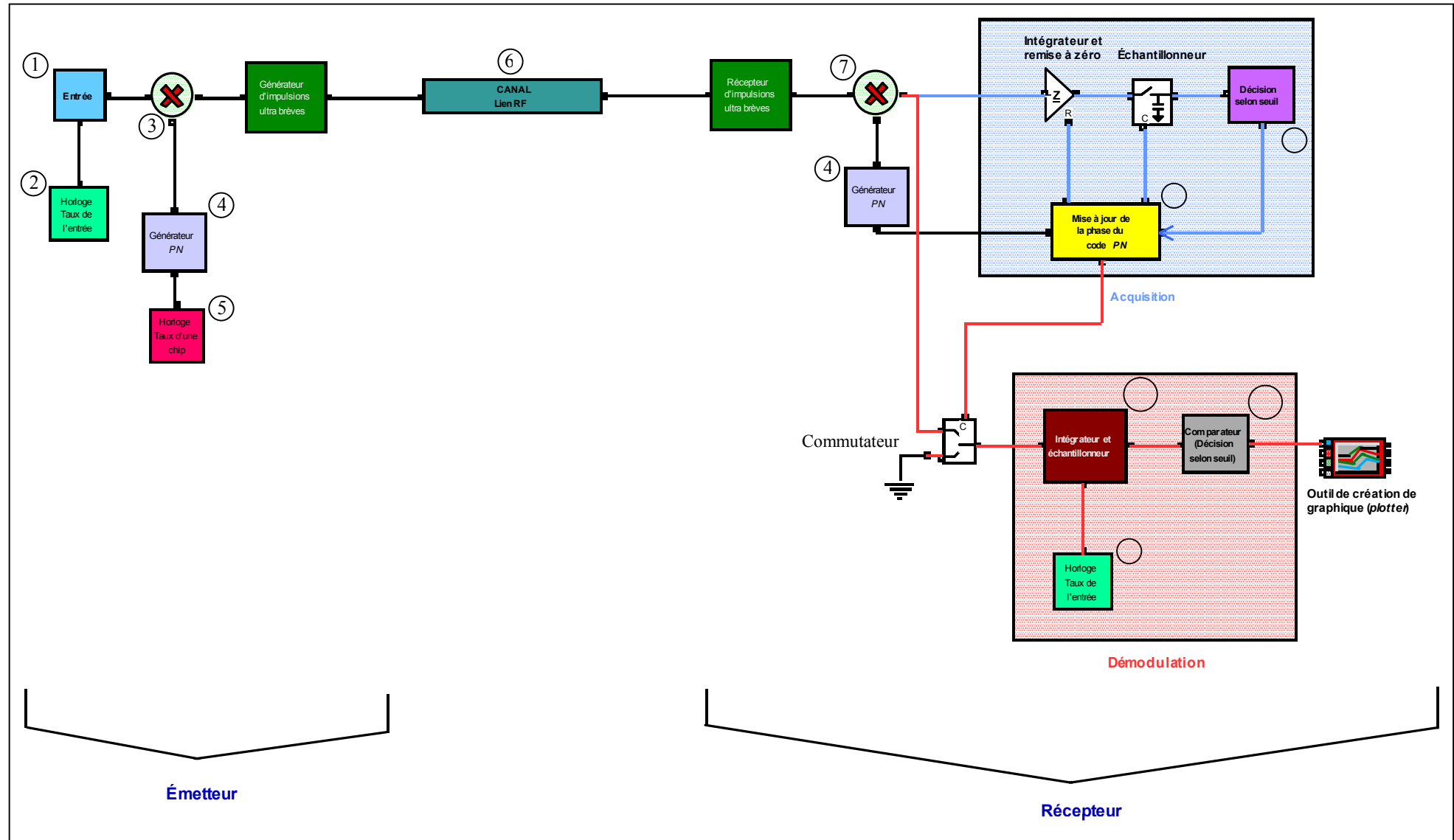
- [1] Union internationale des télécommunications. (pages consultées à partir du 23 juin 2000), « ITU Telecommunication Terminology Database (TERMITE) », [en ligne]. Adresse URL: <http://www.itu.int/search/wais/Termite/index.html>.
- [2] Multispectral Solutions, Inc. (pages consultées à partir du 11 mai 2000), « Welcome to Multispectral Solutions, Inc. », [en ligne]. Adresse URL: <http://www.multispectral.com>.
- [3] TAYLOR, James D., *et al. Introduction to Ultra-wideband Radar Systems*, James D. Taylor Editor, CRC Press, Boca Raton, 1995, 670 p.
- [4] WITHINGTON, Paul. « Impulse Radio Overview », Time Domain Corporation.
- [5] BENNETT, C. Leonard et Gerald F. ROSS. « Time-Domain Electromagnetics and Its Applications » dans *Proceedings of the IEEE*, vol. 66, n° 3, mars 1978, p.299-318.
- [6] Aether Wire & Location, Inc., (pages consultées à partir du 4 mai 2000), « Welcome to Aether Wire », [en ligne]. Adresse URL: <http://www.aetherwire.com/CDROM/Welcome1.html>.
- [7] Federal Communications Commission. « Notice of Inquiry », FCC 98-208, ET Docket 98-153, *In the Matter of Revision of Part 15 of the Commission's Rules Regarding Ultra-Wideband Transmission Systems*, Washington, 1^{er} septembre 1998.
- [8] Federal Communications Commission. « Notice of Proposed Rule Making », FCC 00-163, ET Docket 98-153, *In the Matter of Revision of Part 15 of the Commission's Rules Regarding Ultra-Wideband Transmission Systems*, Washington, 11 mai 2000.
- [9] BARRETT, Terence W. « History of UltraWideBand (UWB) Radar & Communications : Pioneers and Innovators », *Progress In Electromagnetics Symposium 2000 (PIERS2000)*, Cambridge, juillet 2000.
- [10] HARMUTH, Henning F. *Antennas and Waveguides for Nonsinusoidal Waves*, Academic Press, Inc., Orlando, 1984, 276 p.
- [11] HARMUTH, Henning F. *Radiation of Nonsinusoidal Electromagnetic Waves*, Academic Press, Inc., New York, 1990, 338 p.
- [12] HARMUTH, Henning F. « Radiation of Nonsinusoidal Waves by a Large-Current Radiator » dans *IEEE Transactions on Electromagnetic Compatibility*, vol. EMC-27, n° 2, mai 1985, p.77-87.
- [13] FULLERTON, Larry W. « Time Domain Radio Transmission System », brevet 5 363 108, 8 novembre 1994.
- [14] PETROFF, Alan et Paul WITHINGTON. « Time Modulated Ultra-Wideband (TM-UWB) Overview », Time Domain Corporation, Huntsville, janvier 2000.

- [15] SCHOLTZ, Robert A. « Multiple Access with Time-Hopping Impulse Modulation », article invité, *IEEE MILCOM'93*, Boston, 11 au 14 octobre 1993.
- [16] FULLERTON, Larry W. « Time Domain Radio Transmission System », brevet 4 979 186, 18 décembre 1990.
- [17] FULLERTON, Larry W. et Ivan A. COWIE. « Ultrawide-band Communication System And Method », brevet 5 677 927, Pulson Communications Corporation, 14 octobre 1997.
- [18] FULLERTON, Larry W. « Full Duplex Ultrawide-Band Communication System and Method », brevet 5 687 169, Time Domain Systems, Inc., 11 novembre 1997.
- [19] FLEMING, Robert et Cherie KUSHNER. « Low-Power, Miniature, Distributed Position Location and Communication Devices Using Ultra-Wideband, Nonsinusoidal Communication Technology », rapport technique semi-annuel, contrat J-FBI-94-058, Æther Wire & Location, Inc., juillet 1995.
- [20] FONTANA, Robert J. « On “Range-Bandwidth per Joule” for Ultra Wideband and Spread Spectrum Waveforms », Multispectral Solutions, Inc., Gaithersburg, 20 mars 1998.
- [21] BURKE, Barry E. et Dan L. SMYTHE, Jr. « A CCD Time-Integrating Correlator » dans *IEEE Journal of Solid-State Circuits*, vol. SC-18, n° 6, décembre 1983, p.736-744.
- [22] WEIK, Martin H. *Communications Standard Dictionary*, Van Nostrand Reinhold, New York, 1989, 1219 p.
- [23] MCCORKLE, John. « A Tutorial on Ultrawideband Technology », soumission produite pour *IEEE P802.15 Working Group for Wireless Personal Area Networks (WPANS)*, XtremeSpectrum, Inc., Greenbelt, mars 2000.
- [24] PROAKIS, John G. *Digital Communications*, McGraw-Hill, Inc., New York, 1995, 928 p.
- [25] SIMON, Marvin K., *et al.* *Spread Spectrum Communications Handbook*, McGraw-Hill, Inc., New York, 1994, 1227 p.
- [26] VITERBI, Andrew J. *CDMA : Principles of Spread Spectrum Communication*, Addison-Wesley Publishing Company, Reading, 1995, 245 p.
- [27] HOLMES, Jack K. *Coherent Spread Spectrum Systems*, John Wiley and Sons Inc., New York, 1982, 624 p.
- [28] DIXON, Robert C. *Spread Spectrum Systems*, John Wiley & Sons, Inc., New York, 1984, 422 p.
- [29] SCHWARZ, Richard. « An Introduction to Linear Recursive Sequences in Spread Spectrum Systems », Sigtek Inc., Columbia.

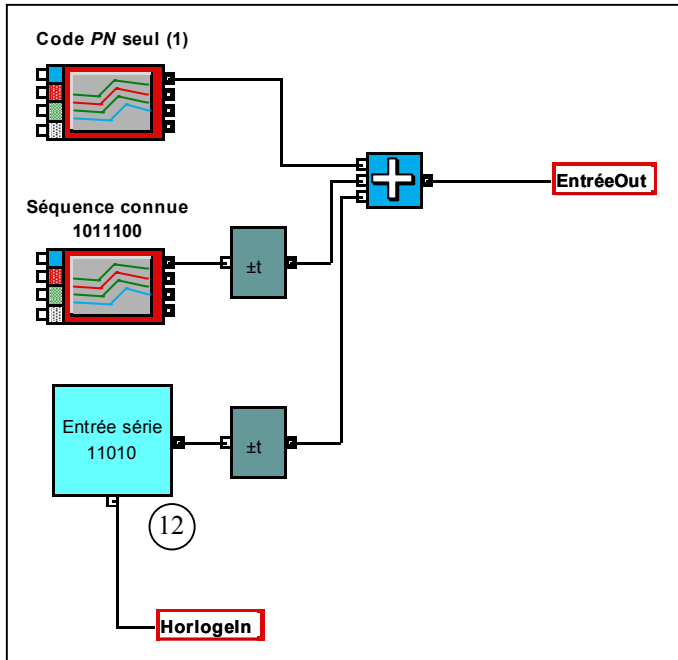
- [30] Delft University of Technology Circuits & Systems Group. (pages consultées à partir du 26 mai 2000), « Spread Spectrum Techniques », [en ligne]. Adresse URL: <http://cobalt.et.tudelft.nl/~wissce/techn/techniques.html>.
- [31] KEMDIRIM, Donald C. et Jim S. WIGHT. « DS SSMA with Some IC Realizations » dans *IEEE Journal on Selected Areas in Communications*, vol. 8, n° 4, mai 1990, p. 663-674.
- [32] HACCOUN, David. *Transmissions numériques : notes de cours*, École Polytechnique de Montréal, 2000, 261 p.
- [33] SKLAR, Bernard. *Digital Communications : Fundamentals and Applications*, Prentice Hall, New Jersey, 1988, 776 p.
- [34] Imagine That, Inc. (pages consultées à partir du 23 mai 2000), « Imagine That, Inc. Simulation-based products and services powered by Extend », [en ligne]. Adresse URL: <http://www.imaginedthatinc.com/>.
- [35] Computer Tips. (page consultée à partir du 12 mai 2000), « ctips.com », [en ligne]. Adresse URL: <http://www.ctips.com/rs232.html>.
- [36] CAMPBELL, Joe. *C Programmer's Guide to Serial Communications*, Howard W. Sams & Company, Carmel, 1989, 655 p.
- [37] FONTANA, Robert J. « A Note on Power Spectral Density Calculations for Jittered Pulse Trains », Multispectral Solutions, Inc., 2000.
- [38] Associated Professional Systems. « PN SIM : Linear Recursive Sequence Analysis Software User's Guide », Maryland.
- [39] LAFORTUNE, Jean-François et Michel LECOURS. « Measurement and Modeling of Propagation Losses in a Building at 900 MHz » dans *IEEE Transactions on Vehicular Technology*, vol. 39, n° 2, mai 1990, p. 101-108.
- [40] LEMIRE, Michel. *Recueil de problèmes : introduction aux communications et Notes de cours*, École Polytechnique de Montréal, 1999.
- [41] HARMUTH, Henning F. « Radar Equation for Nonsinusoidal Waves » dans *IEEE Transactions on Electromagnetic Compatibility*, vol. 31, n° 2, mai 1989, p. 138-147.
- [42] WITHINGTON, Paul. « Re: Ultra-Wideband Testing by NTIA », commentaire en réponse au *Plan for Developing Measurement Methods, Characterizing the Signals and Estimating Their Effects on Existing Systems* du *National Telecommunications and Information Administration*, Time Domain Corporation, Huntsville, juillet 2000.
- [43] Tektronix Inc., (page consultée à partir du 19 mai 2000), « Tektronix Oscilloscopes | DPO, DSO, Sampling and Handheld Oscilloscopes », [en ligne]. Adresse URL: <http://www.tek.com/Masurement/scopes/>.

Annexe

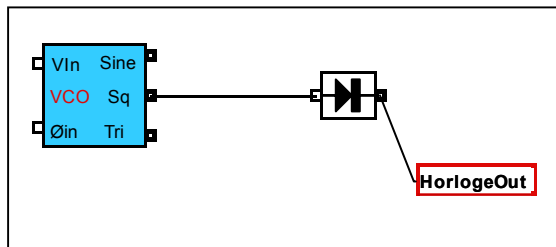
Système émetteur-récepteur complet utilisant la technologie à bande ultra large



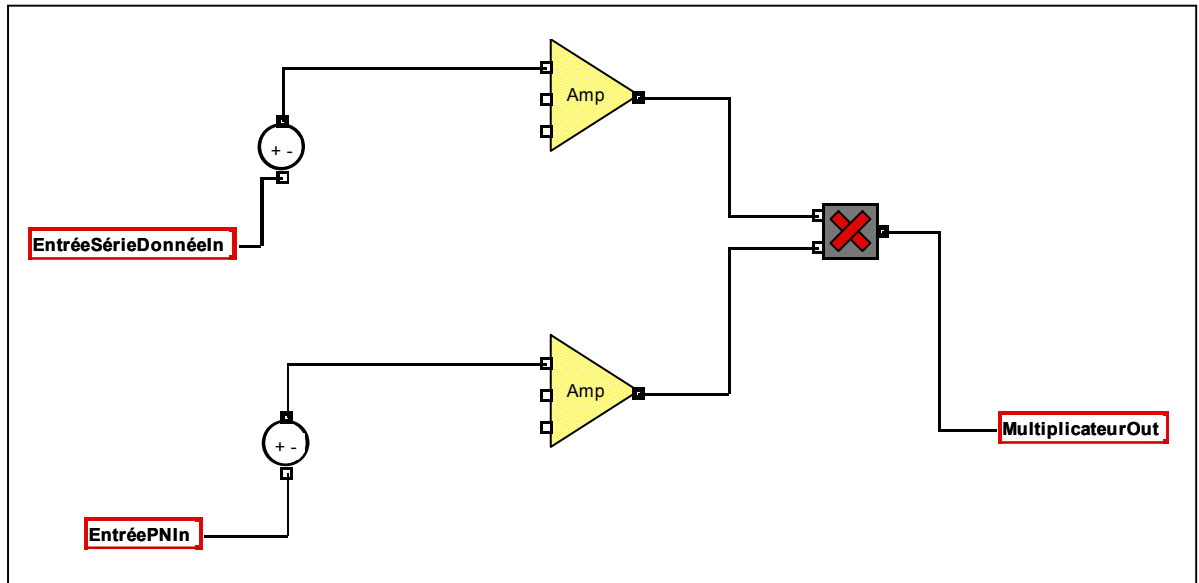
① **Entrée du système**



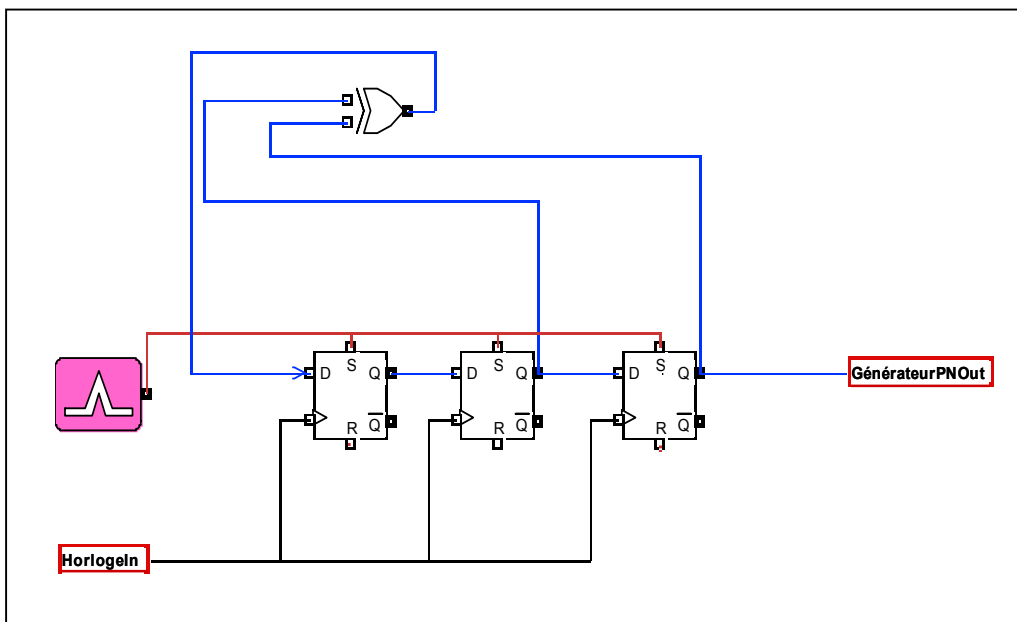
② **Horloge au taux de l'entrée (115 200 Hz)**



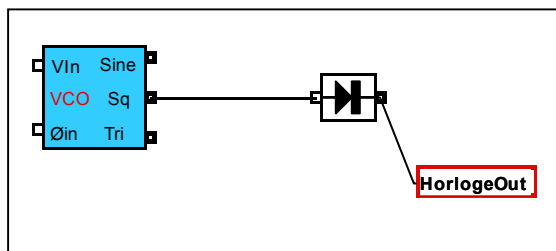
③ Multiplicateur de l'émetteur



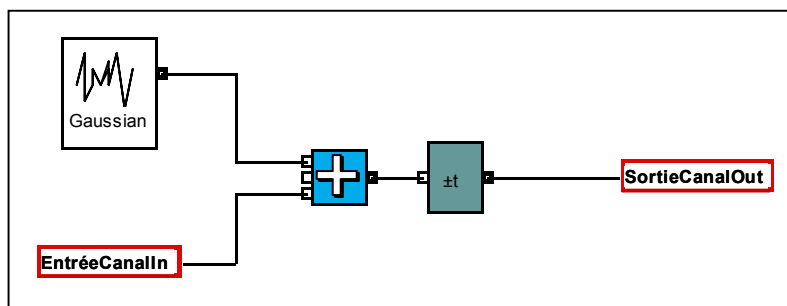
④ Registre à décalage générant le code PN



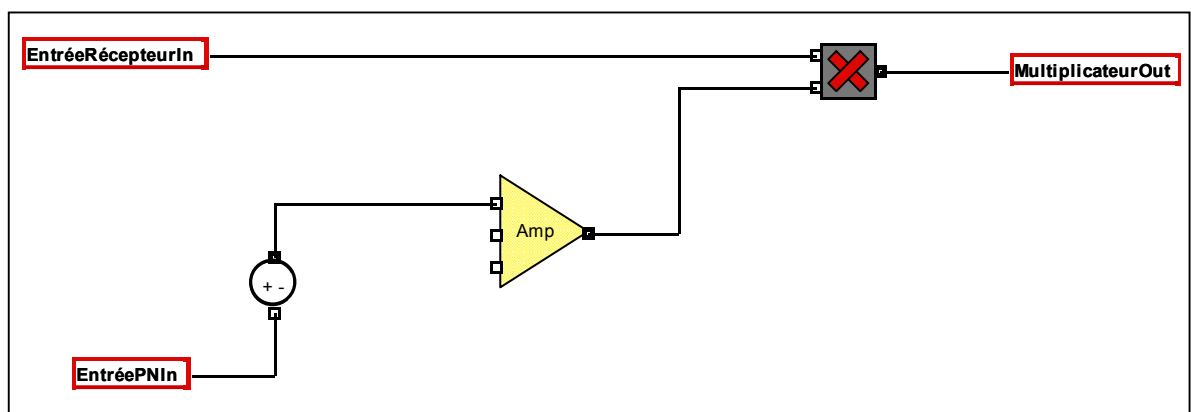
5 Horloge au taux du code *PN* (806 400 Hz)



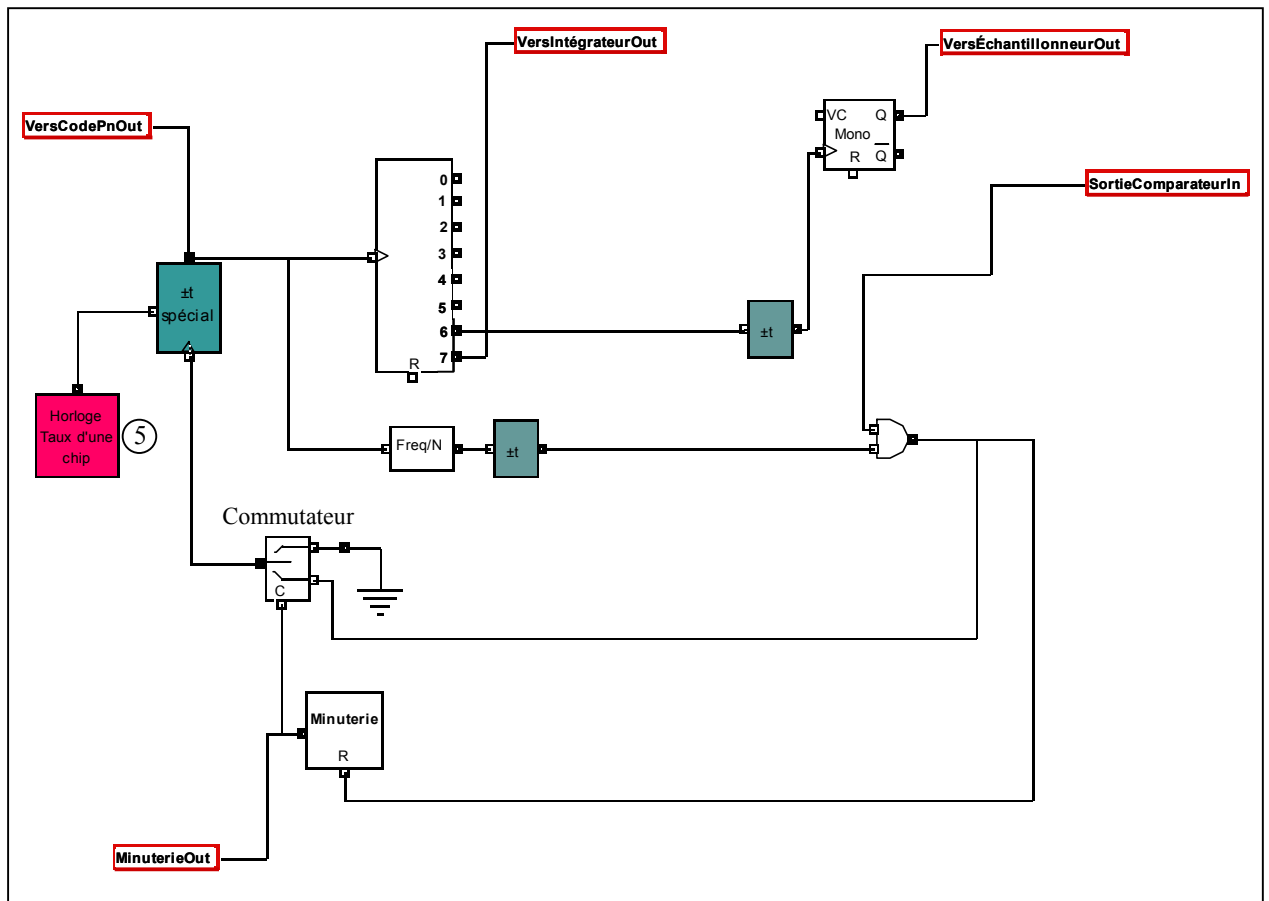
6 Canal modélisé



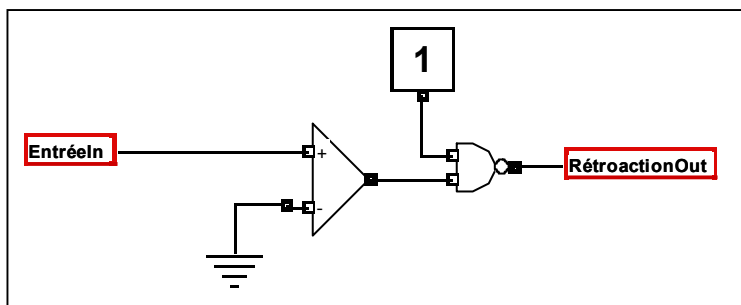
7 Multiplicateur du récepteur



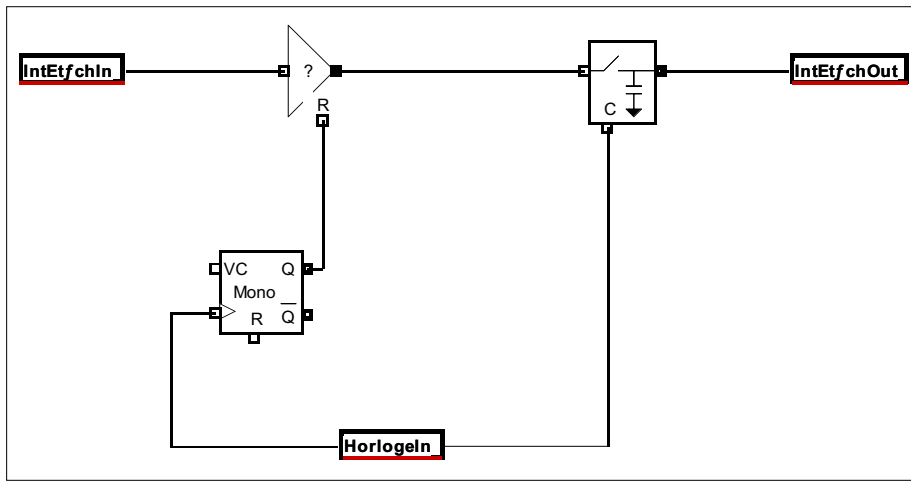
8 Bloc permettant la mise à jour de la phase du code *PN* pour l'acquisition



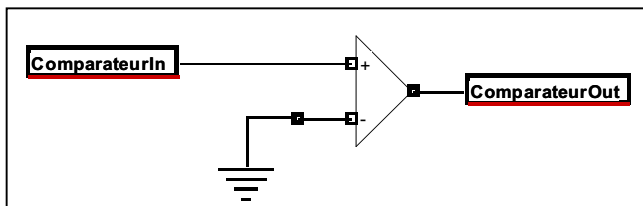
9 Bloc décidant si l'acquisition est terminée



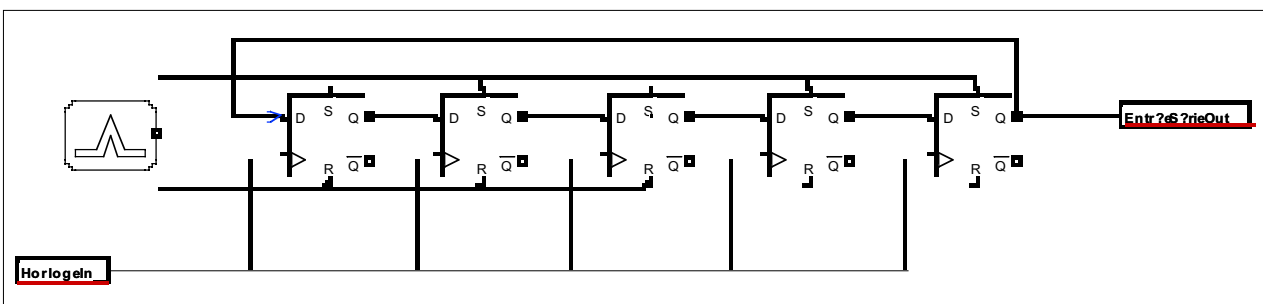
10 Intégrateur et échantillonneur du démodulateur



11 Comparateur permettant de décider si un 0 ou un 1 a été reçu



12 Générateur des données série



Carnet associé au modèle

1. Bruit dans le canal

Moyenne :

Écart-type :

Informations générales

Débit du port série : 115 200 bauds

Durée d'un bit d'information : 8,68e-6 s

Débit du code *PN* : 806 400 chips/s

Durée d'une chip du code *PN* : 1,24e-6 s

2. Délai du signal reçu :

secondes

3. Graphiques

