

UNIVERSITE D'ANTANANARIVO

ECOLE SUPERIEURE POLYTECHNIQUE

DEPARTEMENT TELECOMMUNICATION

MEMOIRE DE FIN D'ETUDES

en vue de l'obtention

du DIPLOME d'INGENIEUR

Spécialité : Télécommunication

Option : Radiocommunications

par : ANDRY Jean Onésime

***LE DS-CDMA ET SES APPLICATIONS
DANS LE SYSTEME DE TELEPHONIE
MOBILE UMTS***

Soutenu le mercredi 9 Octobre 2002 devant la Commission d'Examen composée de :

Président :

M. RANDRIAMITANTSOA Paul Auguste

Examineurs :

1. M. RATSIHOARANA Constant

2. M BOTO ANDRIANANDRASANA Jean Espérant

Directeur de mémoire :

M. RAKOTOMALALA Mamy Alain

REMERCIEMENTS

J'aimerais adresser mes vifs remerciements à :

- Monsieur RANDRIAMITANTSOA Paul Auguste, Chef du Département Télécommunications, qui m'a fait l'honneur de présider le jury de ce mémoire
- Monsieur RAKOTOMALALA Mamy Alain qui m'a témoigné son entière confiance et qui m'a donné de précieux conseils en tant que Directeur de ce mémoire.
- Aux membres du jury

Monsieur RATSIHOARANA Constant

Monsieur RANDRIATSIRESY Ernest

Monsieur BOTO ANDRIANANDRASANA Jean Espérant

- Au corps enseignant au sein du Département Télécommunications à l'ESPA.
- A toute ma famille qui m'a toujours soutenu au cours de mes études.
- A mes collègues et amis qui, de près ou de loin, ont contribué à la réalisation de ce travail.

TABLES DES MATIERES

TABLES DES MATIERES	i
ABREVIATIONS.....	viii
GLOSSAIRE.....	xi
INTRODUCTION.....	1

PREMIERE PARTIE

FONDEMENTS THEORIQUES DU DS-CDMA

CHAPITRE I GENERALITES SUR L'ACCES MULTIPLE

I.1 INTRODUCTION.....	4
I.2 COMPARAISON DES DIFFERENTES METHODES D'ACCES MULTIPLE.....	5
<i>I.2.1 L'accès multiple à répartition fréquentielle</i>	<i>6</i>
I.2.1.1 Les caractéristiques les plus importantes du FDMA.....	6
I.2.1.2 Avantages du FDMA.....	7
I.2.1.3 Inconvénients du FDMA.....	7
<i>I.2.2 L'accès multiple à répartition dans le temps.....</i>	<i>8</i>
I.2.2.1 Avantages du TDMA.....	8
I.2.2.2 Inconvénients du TDMA.....	9
<i>I.2.3 L'accès multiple à répartition dans les codes</i>	<i>9</i>
<i>I.2.4 Le mode de duplexage TDD.....</i>	<i>11</i>
<i>I.2.5 Le mode de duplexage FDD.....</i>	<i>12</i>
I.3 CAS PARTICULIERS.....	12

CHAPITRE II HISTORIQUE DE DEVELOPPEMENT DES TECHNIQUES D'ETALEMENT DE SPECTRE

II.1 HISTORIQUE.....	13
----------------------	----

CHAPITRE III LES PRINCIPES DE BASE DE L'ETALEMENT DE SPECTRE

III.1 LES NOTIONS DE BASE.....	15
--------------------------------	----

III.1.1 Les notions de base.....	15
III.1.2 Le chip.....	15
III.1.3 La longueur d'une séquence.....	15
III.1.4 Le Gain de traitement (Processing Gain).....	15
III.2 LES DIFFERENTES TECHNIQUES D'ETALEMENT DE SPECTRE.....	16
III.2.1 L'étalement de spectre par séquence directe (DS-SS).....	16
III.2.2 L'étalement de spectre par saut de fréquence(FH-SS).....	19
III.3 COMPARAISON ENTRE L'ETALEMENT PAR SEQUENCE DIRECTE ET CELUI PAR SAUT DE FREQUENCE.....	22
III.4 DETAIL DE FONCTIONNEMENT D'UN SYSTEME A ETALEMENT DE SPECTRE PAR SEQUENCE DIRECTE.....	23
III.4.1 Introduction.....	23
III.4.2 Suppositions.....	24
III.4.3 Système de base.....	24
 CHAPITRE IV LES CODES D'ETALEMENT UTILISES DANS UN SYSTEME DS-CDMA	
IV.1 RAPPEL SUR LE BRUIT BLANC GAUSSIEN.....	27
IV.2 METHODES DE SEPARATION DES UTILISATEURS DANS UN SYSTEME DS-CDMA..	28
IV.2.1 Notion d'orthogonalité des codes.....	28
IV.2.2 Propriétés d'auto-corrélation.....	28
IV.2.3 Propriétés d'inter-corrélation.....	29
IV.3 LES CODES UTILISES POUR L'ACCES MULTIPLE ORTHOGONAL.....	29
IV.3.1 Les codes de Walsh-Hadamard.....	29
IV.3.1.1 Méthodes de génération des codes de Walsh-Hadamard.....	29
IV.3.1.2 Propriétés d'auto-corrélation d'un code de Walsh-Hadamard.....	31
IV.3.1.3 Propriétés d'inter-corrélation d'un code de Walsh-Hadamard.....	32
IV.3.1.4 La notion de facteur d'étalement variable.....	34
IV.4 LES CODES UTILISES POUR L'ACCES MULTIPLE NON ORTHOGONAL.....	36
IV.4.1 Principes de génération des séquences PN.....	37

<i>IV.4.2 Notion de longueur d'un code PN.....</i>	<i>38</i>
<i>IV.4.3 Propriétés communes des codes PN.....</i>	<i>38</i>
<i>IV.4.4 Les types de codes PN.....</i>	<i>39</i>
IV.4.4.1 Les M-séquences.....	39
<i>IV.4.4.1.1 Principe de génération des M-séquences.....</i>	<i>39</i>
<i>IV.4.4.1.2 Propriétés des M-séquences.....</i>	<i>40</i>
<i>IV.4.4.1.3 Conclusion.....</i>	<i>42</i>
IV.4.4.2 Les codes de Gold.....	43
<i>IV.4.4.2.1 Principe de génération des codes de Gold.....</i>	<i>43</i>
<i>IV.4.4.2.2 Propriétés des codes de Gold.....</i>	<i>45</i>
<i>IV.4.4.2.3 Conclusion.....</i>	<i>47</i>
IV.4.4.3 Les codes de Kasami.....	47
<i>IV.4.4.3.1 Principe de génération des codes de Kasami.....</i>	<i>47</i>
<i>IV.4.4.3.2 Propriétés des codes de Kasami.....</i>	<i>48</i>
<i>IV.4.4.3.3 Conclusion.....</i>	<i>51</i>
IV.5 L'ENCHAINEMENT DES CODES ORTHOGONAUX AVEC DES CODES PN.....	52
<i>IV.5.1 Propriétés des codes enchainés.....</i>	<i>53</i>
<i>IV.5.1.1 Propriétés d'inter-corrélation.....</i>	<i>54</i>
<i>IV.5.1.2 Propriétés d'auto-corrélation.....</i>	<i>55</i>
<i>IV.5.1.3 Résumé des résultats précédents.....</i>	<i>55</i>
<i>IV.5.2 Conséquences de l'utilisation des codes enchainés sur la capacité d'un système.....</i>	<i>56</i>
<i>IV.5.2.1 Commentaires et interprétation de ces résultats</i>	<i>57</i>
 CHAPITRE V EMISSION, RECEPTION et SYNCHRONISATION	
V.1 ARCHITECTURE GENERALE D'UN EMETTEUR DS-SS.....	58
V.2 ARCHITECTURE GENERALE D'UN RECEPTEUR DS-SS.....	58
V.3 L'OPERATION DE DE-ETALEMENT.....	59
<i>V.3.1 Le "Matched Filter".....</i>	<i>60</i>
<i>V.3.2 Le corrélateur actif.....</i>	<i>61</i>
V.4 L'OPERATION DE SYNCHRONISATION.....	62
<i>V.4.1 Phase d'acquisition (Coarse Synchronisation).....</i>	<i>62</i>
<i>V.4.1.1 L'approche entièrement parallèle.....</i>	<i>63</i>
<i>V.4.1.2 L'approche entièrement sérielle.....</i>	<i>64</i>
<i>V.4.1.3 L'approche hybride.....</i>	<i>65</i>
<i>V.4.1.4 La phase d'acquisition dans la pratique.....</i>	<i>67</i>

V.4.2 Phase de poursuite (<i>Tracking Phase</i>).....	68
V.4.2.1 La boucle de verrouillage de délai DLL (<i>Delay Locked Loop</i>).....	68
 CHAPITRE VI LES PERFORMANCES D'UN SYSTEME DS-CDMA	
VI.1 DEFINITION D'UN SYSTEME DS-CDMA.....	71
VI.2 CARACTERISTIQUES GLOBALES D'UN SYSTEME DS-CDMA.....	71
VI.3 PERFORMANCES D'UN SYSTEME DS-CDMA.....	72
VI.3.1 Les puissances mises en jeu.....	72
VI.3.2 Résistance aux effets des propagations par trajets multiples et des interférences.....	73
VI.3.2.1 Introduction.....	73
VI.3.2.2 Résistance aux interférences.....	74
VI.3.2.3 Les propagations par trajets multiples.....	76
VI.3.3 Capacité d'un système DS-CDMA.....	77
VI.3.3.1 Les avantages inhérents à un système DS-CDMA.....	77
VI.3.3.1.1 La prise en compte de l'activité vocale.....	77
VI.3.3.1.2 Limite floue de la capacité (<i>soft limit capacity</i>).....	77
VI.3.3.1.3 Effet de sectorisation.....	77
VI.3.3.1.4 La réutilisation des fréquences.....	78
VI.3.3.2 Capacité d'un système DS-CDMA à cellule unique.....	78
VI.3.3.2.1 La sectorisation.....	80
VI.3.3.2.2 La prise en compte de l'activité vocale.....	80
VI.3.3.3 Capacité de la liaison montante d'un système DS-CDMA dans une configuration pluricellulaire.....	81
VI.3.3.4 Capacité de la liaison descendante d'un système DS-CDMA dans une configuration pluricellulaire.....	83
VI.3.3.5 Exemples et comparaisons.....	85
VI.3.3.5.1 Paramètres.....	85
VI.3.3.5.2 Calculs.....	85
VI.3.3.5.3 Comparaisons.....	85
VI.4 PROBLEMES INHERENTS A UN SYSTEME DS-CDMA.....	86
VI.4.1 L'interférence causée par l'accès multiple (<i>MAI : Multiple Access Interference</i>).....	86
VI.4.2 L'effet proche-lointain (<i>near- far effect</i>).....	86

DEUXIEME PARTIE	
APPLICATIONS DU DS-CDMA A UN SYSTEME	
DE TELEPHONIE MOBILE TEL QUE L'UMTS	
CHAPITRE VII UTILISATION DU DS-CDMA DANS UN RESEAU TELEPHONIQUE	
MOBILE	
VII.1 INTRODUCTION.....	89
VII.2 LES PARAMETRES D'OPTIMISATION.....	90
<i>VII.2.1 La puissance d'émission du mobile.....</i>	<i>90</i>
<i>VII.2.2 La puissance d'émission de la station de base.....</i>	<i>91</i>
<i>VII.2.3 La sectorisation.....</i>	<i>92</i>
<i>VII.2.4 La mise en place de plusieurs couches de cellules.....</i>	<i>93</i>
VII.3 LES NOUVELLES NOTIONS INTRODUITES PAR LE DS-CDMA.....	94
<i>VII.3.1 Les récepteurs RAKE.....</i>	<i>94</i>
<i>VII.3.2 Le Soft Handover.....</i>	<i>96</i>
<i>VII.3.3 Le Softer Handover.....</i>	<i>97</i>
<i>VII.3.4 L'Interfrequency Handover.....</i>	<i>97</i>
<i>VII.3.5 Les récepteurs à détecteur Multi-Utilisateurs.....</i>	<i>98</i>
 CHAPITRE VIII LE WCDMA : INTERFACE RADIO DE L'UMTS	
VIII.1 INTRODUCTION.....	100
VIII.2 LES SYSTEMES DE COMMUNICATIONS MOBILES DE LA TROISIEME	
GENERATION (3G).....	100
VIII.3 LES BANDES DE FREQUENCE UTILISEES PAR L'UMTS.....	100
VIII.4 LES SERVICES FOURNIS PAR L'UMTS.....	101
VIII.5 L'ARCHITECTURE D'UN SYSTEME UMTS.....	101
VIII.6 LA TECHNOLOGIE WCDMA.....	102
VIII.7 LES CANAUX DE TRANSPORT.....	102

VIII.8 LES CANAUX PHYSIQUES.....	104
<i>VIII.8.1 La liaison montante.....</i>	<i>104</i>
<i>VIII.8.2 La liaison descendante.....</i>	<i>105</i>
VIII.9 ETALEMENT DE SPECTRE ET MODULATION SUR LES CANAUX DEDIES.....	106
VIII.10 ETALEMENT DE SPECTRE SUR LES CANAUX COMMUNS.....	108
<i>VIII.10.1 Les CCPCH primaires et secondaires.....</i>	<i>108</i>
<i>VIII.10.2 Le SCH.....</i>	<i>109</i>
<i>VIII.10.3 Le PRACH.....</i>	<i>110</i>
VIII.11 LES COUCHES DE CELLULES UTILISEES DANS L'UMTS.....	111
VIII.12 LA REALISATION DU SOFT HANDOVER DANS UN RESEAU UMTS.....	112
VIII.13 LA REALISATION DE L'INTERFREQUENCY HANDOVER DANS UN RESEAU UMTS.....	113

TROISIEME PARTIE

SIMULATIONS

CHAPITRE IX SIMULATIONS SOUS MATLAB

IX.1 INTRODUCTION.....	114
IX.2 SIMULATION DU COMPORTEMENT D'UN SYSTEME DS-CDMA	114
<i>IX.2.1 Présentation et interprétation des résultats.....</i>	<i>117</i>
IX.2.1.1 Les données émises et les codes utilisés par les utilisateurs.....	117
IX.2.1.2 Cas du canal idéal, 3 utilisateurs, le récepteur connaît le code utilisé à l'émission.....	118
IX.2.1.3 Cas du canal idéal, 3 utilisateurs, le récepteur ne connaît pas le code utilisé à l'émission.....	118
IX.2.1.4 Cas du canal idéal, 3+20 utilisateurs, Gold(8,2).....	118
IX.2.1.5 Cas du canal idéal, 3+20 utilisateurs, Gold(11,2).....	118
IX.2.1.6 Cas du canal bruité (SNR=10dB), 3 utilisateurs, Gold(8,2).....	119
IX.2.1.7 Cas du canal bruité (SNR=0dB), 3 utilisateurs, Gold(8,2).....	120
IX.2.1.8 Cas du canal bruité (SNR=-4dB), 3 utilisateurs, Gold(8,2).....	120
IX.2.1.9 Cas du canal bruité (SNR=-4dB), 3 utilisateurs, Gold(11,2).....	120
IX.2.1.10 Cas du canal avec interférence à bande étroite (SIR=10db), 3 utilisateurs, Gold(8,2).....	121
IX.2.1.11 Cas du canal avec interférence à bande étroite (SIR=0db), 3 utilisateurs, Gold(8,2).....	122
IX.2.1.12 Cas du canal avec interférence à bande étroite (SIR=0db), 3 utilisateurs, Gold(11,2).....	122

IX.2.1.13 Cas du canal avec trajets multiples, 3 utilisateurs, Gold(11,2).....	123
IX.2.2 Conclusion.....	124
IX.3 UN EXEMPLE DE PROGRAMME D'AIDE A L'OPTIMISATION DE LA CAPACITE D'UN RESEAU DS-CDMA PLURICELLULAIRE CHARGE.....	124
IX.3.1 Définition des capacités égales.....	126
IX.3.2 Définition de la capacité optimale d'un réseau DS-CDMA.....	126
IX.3.3 Le facteur de compensation de puissance (PCF: Power Compensation Factor).....	127
IX.3.4 Les paramètres utilisés par le programme de simulation.....	127
IX.3.5 Calculs.....	129
IX.3.5.1 Capacités égales, distribution uniforme des utilisateurs.....	129
IX.3.5.2 Répartition optimale, distribution uniforme des utilisateurs.....	130
IX.3.5.3 Capacités égales, distribution non uniforme des utilisateurs.....	131
IX.3.5.4 Répartition optimale, distribution non uniforme des utilisateurs.....	131
IX.3.5.5 Répartition optimale, distribution non uniforme des utilisateurs, avec une contrainte de 8 utilisateurs par cellule au minimum.....	133
IX.3.5.6 Répartition optimale, distribution non uniforme des utilisateurs, utilisation de la compensation de puissance.....	133
IX.3.5.7 Répartition optimale, distribution non uniforme des utilisateurs, utilisation de la compensation de puissance, avec une contrainte de 8 utilisateurs par cellule au minimum.....	134
IX.3.6 Conclusion.....	135
CONCLUSION.....	136
ANNEXE 1.....	138
ANNEXE 2.....	140
ANNEXE 3.....	144
ANNEXE 4.....	148
ANNEXE 5.....	152
BIBLIOGRAPHIE.....	155

ABREVIATIONS

3G	Third Generation Mobile Network
3GPP	Third Generation Partnership Project
AMPS	Advanced Mobile Phone System
AMRC	Accès Multiple à Répartition dans les Codes
AMRF	Accès Multiple à Répartition Fréquentielles
AMRT	Accès Multiple à Répartition dans le Temps
AuC	Authentication Center
BCCH	Broadcast Control CHannel
BPSK	Binary Phase Shift Keying
BSC	Base Station Controller
BTS	Base Transceiver Station
CCPCH	Common Control Physical CHannel
CDMA	Code Division Multiple Access
CN	Core Network
D-AMPS	Digital - Advanced Mobile Phone System
D-BPSK	Differential Binary Phase Shift Keying
DCCH	Dedicated Control CHannel
DLL	Delay-Locked Loop
DPCCH	Dedicated Physical Control CHannel
DPCH	Downlink Physical Channel
DPDCH	Dedicated Physical Data CHannel
DS	Direct Sequence
DS-CDMA	Direct Sequence - Code Division Multiple Access
DS-SS	Direct Sequence – Spread Spectrum
DTCH	Dedicated Traffic CHannel
EDGE	Enhanced Data rate for GSM Evolution
ETSI	European Telecommunications Standards Institute
FACH	Forward Access CHannel
FDD	Frequency Division Duplex
FDMA	Frequency Division Multiple Access

FH	Frequency Hopping
FH-SS	Frequency Hopping – Spread Spectrum
FSK	Frequency Shift Keying
FW	Frequency Word
GPRS	Global Packet Radio Service
GPS	Global Positioning System
GSM	Global System for Mobile communication
IMT-2000	International Mobile Telecommunications 2000
IT	Intervalle Temporel
LA	Location Area
LMMSE	Linear Minimum Mean Square Error
MAI	Multiple Access Interference
MSC	Mobile Switching Center
MSK	Minimum Shift Keying
MUD	Multi-User Detection
NCO	Numerical Controlled Oscillator
OVSF	Orthogonal Variable Spreading Factor
PCH	Paging CHannel
PIC	Parallel Interference Cancellation
PN	Pseudo-Noise
PPGC	Preferred Pair Gold Codes
PRACH	Physical Random Access CHannel
PS	Packet Switched
PSK	Phase Shift Keying
QAM	Quadrature Amplitude Modulation
QPSK	Quadrature Phase Shift Keying
RACH	Random Access CHannel
RF	Radiofréquence
RIF	Réponse Impulsionnelle Finie
RNC	Radio Network Controller
RRC	Radio Resource Control
RTCP	Réseau Téléphonique Commuté Public

SCH	Synchronization CHannel
SIC	Successive Interference Cancellation
SMS	Short Message Service
SNR	Signal to Noise Ratio
SS	Spread Spectrum
SSRG	Simple Shift Register Generator
TD-CDMA	Time Division - Code Division Multiple Access
TDD	Time Division Duplex
TDMA	Frequency Division Multiple Access
TEB	Taux d'Erreurs Binaires
UE	User Equipment
UIT	Union Internationale des Télécommunications
UMTS	Universal Mobile Telecommunications System
UTRAN	UMTS Terrestrial Radio Access Network
WCDMA	Wideband – Code Division Multiple Access

GLOSSAIRE

Diversité	<i>C'est la transmission du même message d'information via plusieurs trajets ou supports distincts dont les statistiques d'évanouissement sont indépendantes</i>
Diversité fréquentielle	<i>C'est un cas particulier de diversité où le message est transmis par deux ou plusieurs canaux ayant des fréquences porteuses différentes, suffisamment éloignées entre elles pour que leurs statistiques d'évanouissement soient indépendantes</i>
Diversité spatiale	<i>C'est un cas particulier de diversité où le message est transmis sur une seule fréquence porteuse mais on utilise plusieurs points de réception (stations de base ou antennes) suffisamment espacés entre eux pour que leurs statistiques d'évanouissement soient indépendantes</i>
Diversité temporelle	<i>C'est un type de diversité spécifique à un système DS-CDMA, dans la mesure où plusieurs copies du même signal originel sont reçues par le récepteur avec des retards aléatoires et peuvent être combinées pour combattre l'évanouissement sur un des trajets empruntés par ces signaux.</i>
DS-CDMA	<i>C'est une méthode d'accès multiple à répartition dans les codes utilisant la séquence directe comme méthode d'étalement de spectre.</i>
DS-SS	<i>Etalement de spectre par séquence directe. L'étalement de spectre est obtenu par multiplication du signal de données par une séquence (code) spécifique à chaque utilisateur dont la fréquence (débit) est largement supérieure à celle des données.</i>
Effet Proche-Lointain ou effet Near-Far	<i>Masquage des utilisateurs reçus avec de faibles puissances, par d'autres reçus avec de fortes puissances (se trouvant près de la station de base par exemple). Ce phénomène est spécifique à un système DS-CDMA.</i>
Hard Handover	<i>C'est un cas particulier de transfert de communication intercellulaire, où la communication avec la cellule courante est d'abord interrompue puis une nouvelle est établie avec une autre cellule.</i>
Hard limit capacity	<i>Terme utilisé pour indiquer que la capacité d'un système donné est strictement limité par la quantité de ressources disponibles (fréquences, temps).</i>
Hot Spot	<i>Une zone englobant une forte densité d'utilisateurs entraînant une distribution non uniforme de ces derniers à travers le réseau.</i>
Interférence due à l'accès multiple	<i>Type d'interférence spécifique à un système DS-CDMA dans lequel tous les utilisateurs émettent dans la même bande de fréquence. Dans un tel système, l'interférence augmente donc avec le nombre d'utilisateurs dans ce système.</i>
Interfrequency Handover	<i>C'est un transfert de communications entre plusieurs fréquences porteuses (dans un réseau DS-CDMA) correspondant soit à différentes couches de cellules ou à des systèmes différents coexistant à l'intérieur d'un réseau.</i>

Liaison descendante	<i>C'est la liaison unidirectionnelle dans le sens station de base vers un terminal mobile.</i>
Liaison montante	<i>C'est la liaison unidirectionnelle dans le sens terminal mobile vers la station de base</i>
Récepteur Multi-Utilisateurs	<i>Type de récepteur particulier utilisé dans un système DS-CDMA effectuant les démodulations conjointes de l'ensemble des utilisateurs afin de choisir l'utilisateur le plus vraisemblable. Il permet d'éliminer l'effet Near Far ainsi que d'atténuer les effets de l'interférence intracellulaire.</i>
Récepteur RAKE	<i>Type de récepteur particulier utilisé dans un système DS-CDMA ayant la capacité de combiner les signaux issus de différents trajets pour offrir une diversité temporelle à la structure de réception utilisée.</i>
Soft Handover	<i>C'est un cas particulier de communication intercellulaire où le mobile établit d'abord plusieurs communications avec plusieurs cellules, et c'est seulement ensuite que se fait le choix de la cellule qui prendra le contrôle du mobile.</i>
Soft limit capacity	<i>Terme utilisé pour indiquer que la capacité d'un système donné n'est pas limitée par la quantité de ressources disponibles (fréquences, temps). C'est le cas d'un système DS-CDMA, où la capacité n'est limitée que par le niveau d'interférence dans le système.</i>
Softer Handover	<i>C'est un transfert de communication entre les secteurs d'une même cellule. La procédure suivie est analogue à celle utilisée dans le cas du Soft Handover.</i>
UMTS	<i>Standard européen implémentant un système de téléphonie mobile de la troisième génération conforme à la recommandation de l'UIT</i>
WCDMA	<i>Version à large bande du DS-CDMA. Il occupe une bande passante allant de 5 à 15 MHz contre 1.25 MHz pour un système DS-CDMA standard</i>

INTRODUCTION

Depuis l'invention du téléphone en 1876, on s'est déjà rendu compte qu'il est irraisonnable (et pratiquement impossible) de relier chaque paire d'abonnés du réseau entre eux par une liaison propriétaire. Si c'est le cas, on doit mettre en place un nombre excessivement élevé de connexions. La solution consistait donc à mettre en place des centres de commutation qui ont pour rôle d'aiguiller les appels vers les destinations correspondantes.

Cette technique permettait donc de relier un grand nombre d'utilisateurs sans que l'infrastructure nécessaire ne soit trop complexe. L'astuce consistait alors à faire partager le même support de transmission (le câble ou la paire torsadée) par plusieurs utilisateurs. Ce support est alloué pour une communication pendant toute sa durée et devient indisponible pour les autres utilisateurs. C'était la forme la plus rudimentaire de l'accès multiple dans l'histoire des télécommunications modernes.

L'utilisation des moyens radioélectriques a permis de révolutionner le domaine des télécommunications, dans la mesure où elle a permis d'ouvrir la porte de la mobilité aux matériels de télécommunications. Depuis quelques dizaines d'années, des systèmes de téléphonie mobile ont vu le jour. Les premiers systèmes utilisaient des fréquences différentes pour chaque communication. Dans de tels systèmes, un canal est alloué pour toute la durée d'une communication et devient par la suite indisponible pour les autres utilisateurs. Cette forme d'accès multiple est appelée FDMA (Frequency Division Multiple Access) ou AMRF (Accès Multiple à Répartition Fréquentielle).

Cependant avec l'essor incroyable que connaît maintenant le monde de la téléphonie mobile, les ressources fréquentielles deviennent vite insuffisantes lorsque le FDMA est utilisé comme méthode d'accès multiple. Des méthodes ont été proposées alors pour rendre efficace l'utilisation du spectre des fréquences qui est une ressource limitée en quantité.

Avec l'avènement des technologies numériques, une nouvelle méthode d'accès multiple a vu le jour : le TDMA (Time Division Multiple Access) ou AMRT (Accès Multiple à Répartition dans le Temps) qui consiste à multiplexer un certain nombre de communications sur une seule porteuse en allouant à chacune d'entre elles un intervalle temporel de largeur fixe. Mais le TDMA ne permet pas encore une utilisation optimale de la ressource fréquentielle, et des saturations ont été vite atteintes dans les réseaux de téléphonie mobile utilisant le TDMA comme le GSM.

Le CDMA (Code Division Multiple Access) ou AMRC (Accès Multiple à Répartition dans les Codes) a été proposée comme alternative au TDMA pour permettre une utilisation plus efficace du spectre des fréquences, et fera l'objet de ce mémoire. Ce type d'accès multiple est caractérisé par une utilisation de toute une bande de fréquence (largement supérieure à celle requise pour transmettre l'information voulue) par tous les utilisateurs à tout moment. On parle alors d'étalement de spectre. Chaque utilisateur se voit ensuite attribué un code particulier pour effectuer l'étalement de spectre, et par lequel on pourra l'identifier.

Les fondements théoriques de cette technologie ont été déjà trouvés depuis la fin de la deuxième guerre mondiale. Seulement, à cette époque, la réalisation d'un tel système était tellement difficile (donc aussi très coûteuse, du moins à cette époque), que l'idée d'utilisation de cette nouvelle technique dans le cadre d'une application commerciale (pour le grand public) fut rapidement abandonnée. Les recherches militaires ont été les seules applications pour lesquelles ces coûts étaient justifiés. Cependant avec les progrès accomplis dans le domaine de la microélectronique et de la microinformatique, l'implémentation d'un système d'accès multiple à répartition de code est devenue abordable. Rien ne s'oppose donc à ce que cette nouvelle technique soit mise à profit pour des applications commerciales, vu ses atouts par rapport aux autres alternatives.

Ce mémoire s'intitule « Le DS-CDMA et ses applications dans un système de téléphonie mobile UMTS ». Il est consacré à l'étude de l'étalement de spectre par séquence directe, du CDMA et leurs applications dans un système de téléphonie mobile.

De nombreux ouvrages ont déjà traité ce sujet. Ce mémoire n'a pas la prétention de tout couvrir ni sur l'étalement de spectre, ni sur le CDMA. Si c'était le cas, on aboutirait facilement à un document ayant plus de mille pages. Par contre, il vise à donner une image facilement compréhensible des avantages fournis par le CDMA, tout en se limitant à un domaine très largement prisé aujourd'hui: l'étalement de spectre par séquence directe. On parle alors de DS-CDMA (Direct Sequence – CDMA).

Ce mémoire se divise en trois parties. La première partie sera consacrée à l'étude des fondements théoriques du DS-CDMA. On y trouvera notamment les principes de bases de l'étalement de spectre par séquence directe ainsi que les études des propriétés des codes utilisés pour l'étalement de spectre. Des structures de base seront proposées pour effectuer l'émission et la réception de signaux à spectre étalé.

On analysera ensuite les performances d'un système DS-CDMA, et des comparaisons seront faites pour montrer l'efficacité d'un système DS-CDMA par rapport aux systèmes conventionnels utilisés actuellement.

La deuxième partie mettra en relief les applications du DS-CDMA comme méthode d'accès multiple au sein d'un système de téléphonie mobile, en particulier dans le système de la troisième génération UMTS (Universal Mobile Telecommunications System). Cette partie servira aussi à introduire de nouvelles notions propres à un système de téléphonie mobile utilisant le DS-CDMA.

Enfin des résultats de simulation seront présentés dans la troisième partie. A travers ces simulations, on pourra discerner le comportement d'un système DS-CDMA en présence d'un canal bruité, d'un canal à trajets multiples et d'une interférence à bande étroite dans le canal de transmission. Un exemple de programme d'aide à l'optimisation d'un réseau DS-CDMA pluricellulaire, basé sur la simulation d'un tel système sera aussi proposé. Cette dernière partie sert donc à décrire les comportements et les performances d'un système DS-CDMA dans un environnement plus réaliste (qui prend en compte certains paramètres traduisant une partie de la réalité physique de l'environnement considéré).

PREMIERE PARTIE :
FONDEMENTS THEORIQUES DU DS-CDMA

Chapitre I

GENERALITES SUR L'ACCES MULTIPLE

I.1 Introduction

Dans plusieurs applications en télécommunications, un support de transmission commun doit être partagé par plusieurs utilisateurs. Des exemples de systèmes utilisant l'accès multiple peuvent inclure les réseaux locaux, les communications sans fil et/ou mobile, liaison montante des communications par satellites ainsi que les réseaux de télémétrie.

Traditionnellement les ressources étaient attribuées aux différents utilisateurs en partitionnant le canal en plusieurs sous – canaux, et en allouant chaque sous – canal à un usager différent. Ce partage du canal peut se faire dans le domaine temporel (TDMA) ainsi que dans le domaine fréquentiel (FDMA).

Plus tard, en remarquant que la ressource fréquentielle devient de plus en plus rare, les travaux de recherche se sont orientés vers des nouveaux concepts permettant une utilisation plus rationnelle du spectre des fréquences. Ces objectifs peuvent être atteints même partiellement en utilisant un nouveau mode d'accès multiple qui consiste à différencier chaque utilisateur par un code approprié, tout en utilisant une même bande de fréquence que les autres. Cette révolution apporte aussi bien des avantages totalement inconnus des autres méthodes d'accès multiple tel que l'accès aléatoire, et un niveau de confidentialité élevé, une meilleure résistance aux interférences et aux brouillages, de faibles puissances mises en jeu.....

Cette nouvelle méthode qui se nomme CDMA a été déjà utilisée pendant plusieurs décennies par les militaires pour les avantages cités ci-dessus.

Actuellement cette technique est entrain d'être vulgarisée dans les applications commerciales et grand public. Mais avant d'entrer dans les détails passons en revue les différentes techniques utilisables actuellement pour implémenter l'accès multiple, ainsi que leurs défauts et avantages respectifs.

I.2 Comparaison des différentes méthodes d'accès multiple [1] [2] [3] [4]

Le site cellulaire gère le trafic des communications avec un grand nombre de mobiles opérant plus ou moins indépendamment les uns des autres. Deux grandes questions doivent être abordées à ce niveau :

- 1- *La gestion du spectre* c'est à dire la manière dont on gère l'allocation totale du spectre en circuits téléphoniques ;
- 2- *L'accès* c'est à dire comment ces circuits seront ensuite affectés suivant les demandes des utilisateurs particuliers (les possibilités d'accès multiple et aléatoire).

Même si la première question est d'une importance évidente, la seconde n'en est pas moindre. Ceci est induit du fait que les systèmes modernes seront désormais à accès multiple, ce qui signifie que chaque circuit téléphonique sera partagé entre plusieurs utilisateurs et n'importe quel utilisateur peut y accéder pour n'importe quel appel téléphonique (*circuit téléphonique banalisé*). La banalisation complète est donc une nécessité pour les systèmes de téléphonie modernes.

. En général, il y a trois architectures possibles basées sur trois différents concepts d'accès multiple (par ordre d'apparition chronologique) :

- ◆ Le FDMA (Frequency Division Multiple Access) ou AMRF (Accès Multiple à Répartition Fréquentielle)
- ◆ Le TDMA (Time Division Multiple Access) ou AMRT (Accès Multiple à Répartition dans le Temps)
- ◆ Le CDMA (Code Division Multiple Access) ou AMRC (Accès Multiple à Répartition dans les Codes)

Pour mettre en œuvre une communication en « duplex », on peut aussi choisir entre plusieurs méthodes (appelées *méthodes de duplexage*) dont les plus connus sont sans doute :

- ◆ le TDD (Time Division Duplex)
- ◆ le FDD (Frequency Division Duplex)

Ces méthodes sont utilisées conjointement avec les méthodes d'accès multiples citées plus haut.

I.2.1 L'accès multiple à répartition fréquentielle

Dans une architecture FDMA, chaque utilisateur se voit attribué une bande de fréquence donnée pour toute la durée de la communication. Les canaux (fréquences) disponibles sont assignés à la demande, sur la base du premier arrivé, premier servi, à des utilisateurs demandeurs pour écouler un appel, ou à ceux pour qui un appel arrivant a été reçu. Chaque fréquence ne transmet qu'une seule communication à un moment donné.

Dans les systèmes FDMA modernes, un ou plusieurs des canaux disponibles est utilisé comme *canal de contrôle*. Dans le cas d'un appel lancé par un mobile, la demande d'appel comprenant la transmission des bits de numérotation et quelquefois d'autres opérations d'initialisation, est transmise sur ce canal du mobile vers la base.

La base transmet ensuite sur ce canal des instructions au mobile pour le commuter vers une fréquence de conversation pour transmettre. L'assignation de fréquence est menée à bonne fin rapidement et automatiquement ; elle est transparente pour l'utilisation mobile.

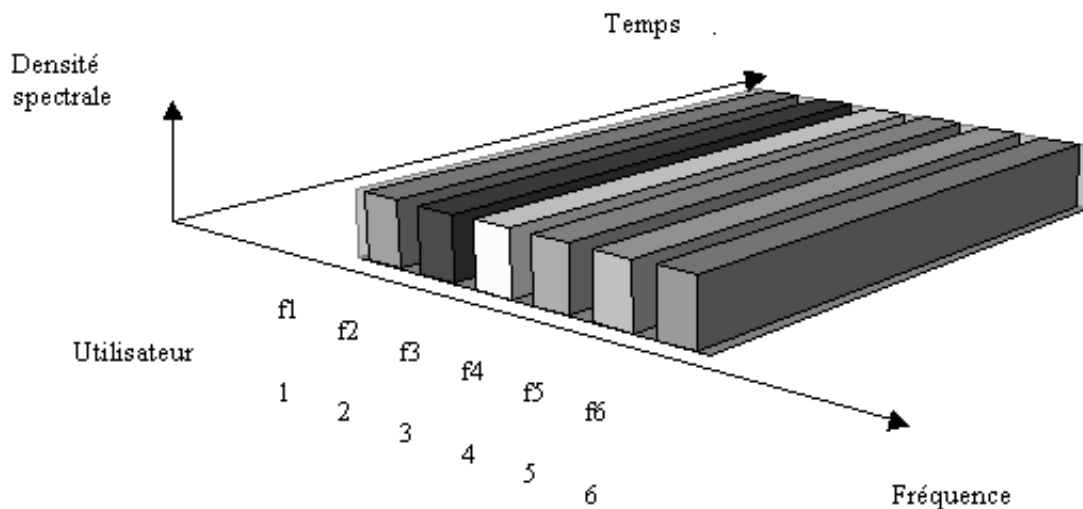


Figure 1.1 : partage des ressources dans le cas du FDMA

I.2.1.1 Les caractéristiques les plus importantes du FDMA

- **Un seul circuit par porteuse** : chaque canal FDMA est conçu pour ne transporter qu'un seul circuit téléphonique.

- **Transmission continue** : une fois que le canal de conversation a été attribué, l'unité mobile et la base transmettent toutes les deux continûment et simultanément.
- **Largeur de bande étroite** : car les canaux FDMA ne transmettent qu'un circuit par porteuse
- **Nécessité d'un espace de garde entre les canaux**: car il est impossible de concevoir un émetteur qui n'émet strictement que dans la bande de fréquence allouée
- **C'est une technique entièrement analogique.**

I.2.1.2 Avantages du FDMA

- **Une certaine immunité face à l'étalement de délai** : car généralement le débit est faible, d'où la durée des symboles est assez importante, ce qui procure une certaine immunité face à *l'étalement de délai*.
- **Complexité moindre de l'unité mobile** : à priori les unités en question n'ont pas besoin de dispositifs d'égalisations, de synchronisations et des formatages complexes. Les bits supplémentaires nécessaires pour assurer la synchronisation est très faible car la transmission du ou vers le mobile est continue (l'usage du canal n'est pas multiplexé).

I.2.1.3 Inconvénients du FDMA

- **Nécessité d'un grand nombre d'équipement pour desservir un nombre donné d'abonné** : un émetteur, un récepteur, deux codecs, et deux modems pour chaque circuit.
- **Le coût partagé par abonné est élevé**
- **La nécessité d'un duplexeur** : par le fait que l'émetteur et le récepteur doivent opérer en même temps.
- **Complexité du transfert automatique de cellule** : Du fait que la transmission FDMA est continue, la réalisation du transfert vers un autre canal dans une autre cellule est plus difficile que dans le système TDMA, où l'on puisse profiter d'un intervalle de temps libre pour réaliser le transfert.

I.2.2 L'accès multiple à répartition dans le temps

Dans un système TDMA, chaque canal radio transporte un certain nombre de circuits banalisés qui sont multiplexés dans le temps. Chaque canal est ainsi divisé en plusieurs intervalles de temps. La réalisation de l'accès multiple en TDMA se fait par une correspondance « temps – fréquence ». L'attribution des circuits aux utilisateurs se fait appel par appel, sous le contrôle de la station de base. L'émission d'un appel est réalisée sur un intervalle de temps de contrôle séparé.

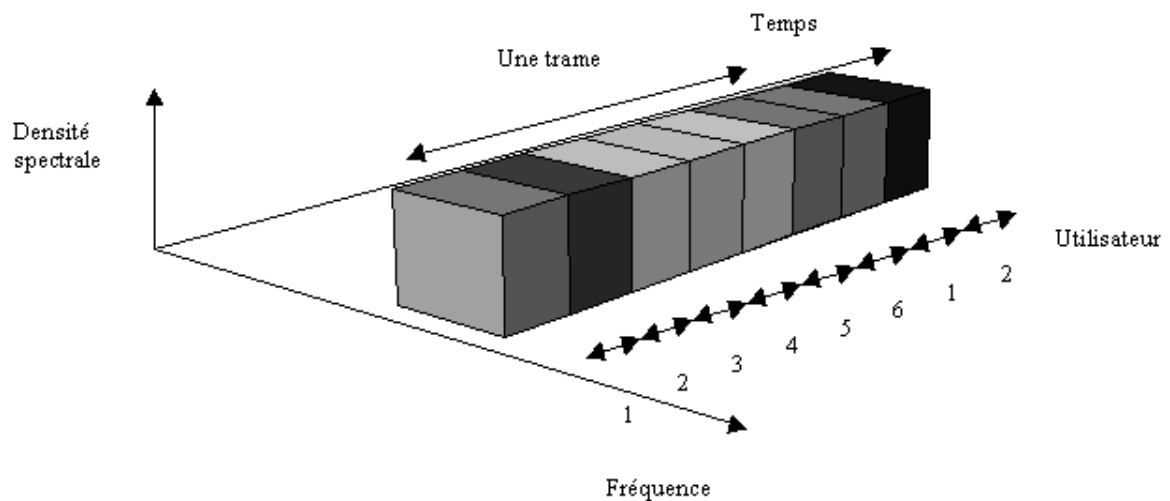


Figure 1.2 : partage des ressources dans le cas du TDMA

I.2.2.1 Avantages du TDMA

- **Transmission par paquets** : ceci influe sur l'effet des interférences, car à un moment donné il n'y a qu'une partie des mobiles en opération qui sont réellement en émission.
- **Coût de système partagé plus bas** : car chaque canal radio est effectivement partagé par un grand nombre d'abonné.
- **Non nécessité d'un duplexeur** : il devient possible d'émettre et de recevoir sur des intervalles de temps différents, ce qui permet de remplacer les circuits duplexeurs par des circuits de commutation rapide, et qui de ce fait permet de réduire le coût de l'unité mobile.

- **Facilité de transfert automatique de cellule** : par le fait que l'émetteur TDMA est inactif pendant des intervalles de temps libres, ce qui permet de mettre une procédure de transfert automatique de cellule plus efficace.

I.2.2.2 Inconvénients du TDMA

- **Une grande sensibilité aux «délais d'étalement»** : Le fait de multiplexer un certain nombre de circuits sur un même canal, provoque inévitablement l'accroissement du débit de ce dernier. Ceci entraîne la diminution de la durée d'un symbole jusqu'à frôler l'étalement moyen des données, d'où la nécessité d'une égalisation plus poussée.
- **Une plus grande complexité du mobile** : Le traitement numérique nécessaire augmente notablement la complexité du mobile qui devient un dispositif de communication complexe et non plus un simple émetteur – récepteur.
- **Une plus grande quantité de bits supplémentaires** : Le récepteur doit récupérer la synchronisation à chaque paquet. Des bandes de garde sont nécessaires pour séparer les intervalles de temps les uns des autres au cas où des délais de propagation non égalisés feraient qu'un utilisateur lointain se glisse dans l'intervalle de temps adjacent de celui d'un utilisateur proche. On a besoin pour cela dans un système TDMA, de plus de bits supplémentaires pour implémenter ces temps de garde. Ces bits peuvent atteindre la proportion de 20 à 30 % du total des bits transmis, ce qui est vraiment très pénalisant.
- **Une limitation de la zone de couverture** : Le temps de garde entre les différentes communications multiplexées sur un seul canal ne peut pas dépasser une valeur donnée. Si l'émetteur est très éloigné du récepteur, le temps de garde peut dépasser la valeur limite et il y aura chevauchement entre deux communications adjacentes. Ceci va se traduire par la perte de la liaison : la communication ne pourra pas aboutir.

I.2.3 L'accès multiple à répartition dans les codes

Chez les deux systèmes précédemment cités, chaque utilisateur émet durant son intervalle de temps (TDMA) ou à travers la fréquence allouée par le système (FDMA). Ceci

a pour principal avantage de supprimer théoriquement les interférences entre les différents usagers du système.

Cependant cette allocation statique de la ressource fréquentielle ne fait pas une utilisation efficace du canal lorsque la transmission effectuée par chaque utilisateur n'est pas fréquente ou se fait en rafale. Par exemple en téléphonie mobile, les études menées ont permis de déterminer que les usagers de ces téléphones parlent moins de 50% du temps durant lequel dure la conversation, laissant donc leur canal oisif pendant les 50% restant. On voit donc aisément que pour de tels cas, une technique de multiplexage plus efficace est souhaitable, qui sera capable d'allouer dynamiquement les ressources en canal en fonction des besoins des utilisateurs actifs en un moment donné.

Les techniques d'étalement du spectre sont à l'origine de l'architecture CDMA. Dans ce type d'accès multiple, une bande de fréquence très large est allouée pour tous les utilisateurs et la distinction entre les diverses communications se fait au moyen de techniques de codage appropriées. Ici on affecte à un mobile un code particulier qui possède une certaine orthogonalité (voir Chapitre IV) vis à vis des autres codes utilisés par les autres usagers. Cependant en pratique, (et comme nous le verrons plus tard) l'attribution de codes parfaitement orthogonaux à tous les utilisateurs peuvent introduire une limitation sensible de la capacité du système. La plupart du temps on doit se contenter d'utiliser des codes « presque orthogonaux », ce qui va se traduire par une augmentation progressive de la probabilité de « collision » de 2 communications utilisant 2 codes se « chevauchant » partiellement, avec le nombre de communications en cours.

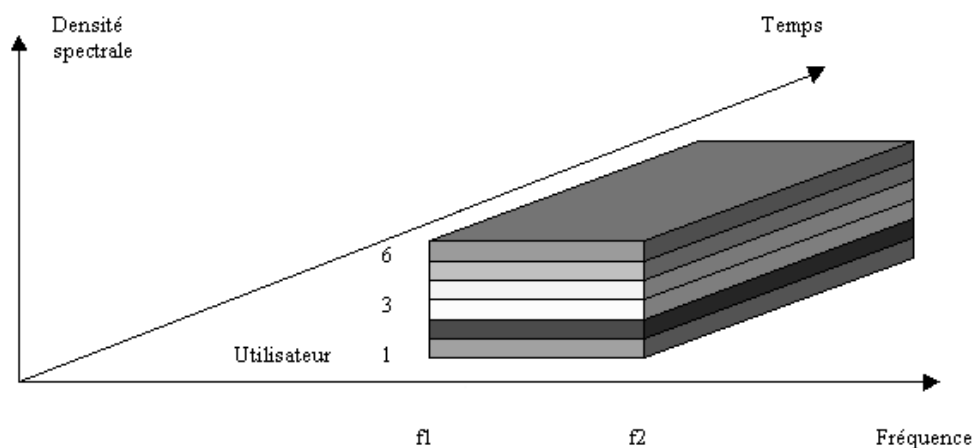


Figure 1.3 : partage des ressources dans le cas du CDMA

A la différence des systèmes FDMA et TDMA, la surcharge du système ne se traduit plus par l'impossibilité de communiquer, ou par une attente longue, mais par une qualité de transmissions de plus en plus dégradée.

Mais CDMA possède plusieurs avantages, qui peuvent faire pencher la balance de son côté lors du choix pour un nouveau système (surtout pour les mobiles):

- plus grande capacité en terme de nombre d'utilisateur par cellule à cause d'une meilleure efficacité spectrale et une meilleure utilisation de la ressource « fréquence »
- meilleure confidentialité des transmissions
- meilleure résistance contre les brouillages et les interférences
- une bonne protection contre les « fadings » dus aux propagations par trajets multiples
- une meilleure répartition de la gestion du système
- une meilleure flexibilité en vue d'une maintenance dans le futur
- une bonne adaptation à la transmission de données numériques par « rafales » puisqu'il est facile de créer un circuit virtuel permanent qui n'est utilisé que lorsque les données sont à transmettre.

1.2.4 Le mode de duplexage TDD

Dans cette méthode, un canal est utilisé mais à un instant particulier, seule une transmission dans un sens est effectuée. Une station émet tandis que l'autre reçoit. C'est seulement après que les rôles s'inversent, permettant d'effectuer des transferts de données dans les deux sens («*duplex*») en n'utilisant qu'un seul canal.

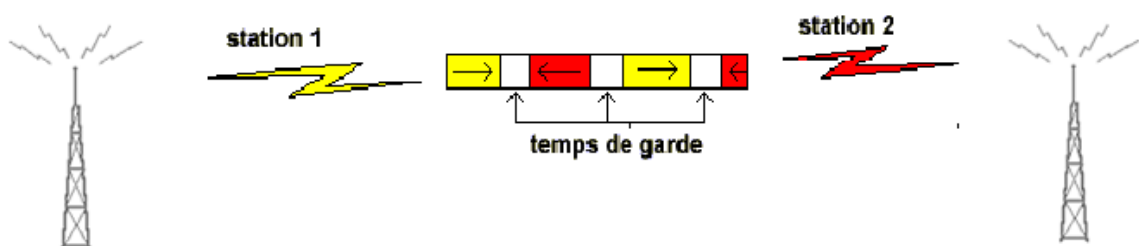


Figure 1.4 : la méthode de duplexage TDD

I.2.5 Le mode de duplexage FDD

Cette méthode consiste à allouer deux chemins différents pour les communications. La station émettrice se voit attribuer une fréquence alors que la station réceptrice se voit en affecter une autre. Les deux stations peuvent émettre et recevoir en même temps. Par contre, ceci a pour inconvénient de réduire le nombre de canaux disponibles dans le système. C'est un système sous - optimal car les stations peuvent ne pas avoir besoin d'un flux ininterrompu de données. Sur la figure suivante montre la communication entre deux stations utilisant deux bandes de fréquence.

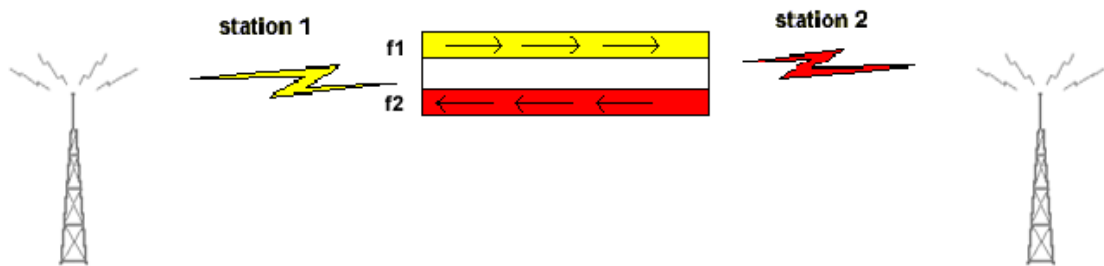


Figure 1.5 : la méthode de duplexage FDD

I.3 Cas particuliers:

Dans le cadre de la norme UMTS (*Universal Mobile Telecommunications System*), des combinaisons particulières ont été retenues :

- ◆ le WCDMA (*Wideband CDMA*) qui utilise conjointement le FDD et CDMA
- ◆ le TD-CDMA (*Time Division CDMA*) qui utilise conjointement le TDD et le CDMA.

Chapitre II

HISTORIQUE DE DEVELOPPEMENT DES TECHNIQUES D'ETALEMENT DE SPECTRE

II.1 Historique

Les techniques d'étalement de spectre ont leurs origines dans les applications militaires. Les recherches ont été accélérées par les besoins militaires de la deuxième guerre mondiale. Cependant, des idées ont déjà fait leurs chemins bien avant le commencement de cette guerre.

- En 1935, des ingénieurs allemands travaillant chez Telefunken, Paul Kotowski et Kurt Dannehl ont postulé un brevet sur un dispositif employé pour masquer des signaux vocaux en les combinant avec un bruit à large bande produit par un générateur rotatif. Le récepteur dispose aussi d'un autre générateur qui est synchronisé et qui est utilisé dans le but de récupérer le signal vocal bruité. Cette invention fut le point de départ du développement des systèmes d'étalement du spectre à Séquence Directe (DS-SS : Direct Sequence - Spread Spectrum).
- En 1949, John Pierce a écrit un mémorandum où il décrivait un système de multiplexage dans lequel un support commun transporte des signaux codés n'ayant pas besoin d'être synchronisés. Claude Shannon et Robert Pierce ont introduit, toujours en 1949 les idées de base de l'accès multiple à répartition dans les codes (AMRC) en décrivant les effets de l'interférence moyenne et la dégradation qui peut s'ensuivre dans un système AMRC.
- En 1950, De Rosa-Rogoff a proposé un système à étalement de spectre par séquence directe (DS-CDMA) et introduit la notion de Gain de Traitement (Processing Gain) ainsi que les équations correspondantes.
- En 1956, Price et Green déposent un brevet sur le premier récepteur RAKE qui a la particularité de pouvoir combiner les signaux issus des propagations par trajets multiples.
- En 1961, l'effet de masquage des autres utilisateurs par un utilisateur puissant dans un système AMRC a été identifié par Magnuski sous le nom de l'effet Near-Far.

- En 1978, l'utilisation de l'étalement de spectre dans un réseau de téléphonie cellulaire a été suggérée par Cooper et Nettleton.
- Dans les années 80, les recherches effectuées par Qualcomm dans le domaine des techniques d'étalement de spectre par séquence directe (DS-SS), ont abouti à la finalisation du standard IS-95 en 1993.
- En 1985, commencement de l'étude du projet IMT-2000
- En 1986, le concept de la détection multi-utilisateur optimale a été formulé par Verdu.
- En 1987, commencement de l'étude de la déclinaison européenne de l'IMT-2000 : l'UMTS
- En Février 1992, les bandes de fréquences pour les systèmes de téléphonie mobile de la troisième génération satisfaisant la recommandation de l'IMT-2000 ont été définies.
- En 1995, les déclinaisons régionales de l'IMT-2000 ont été définies et retenues:
 - UMTS en Europe
 - Cdma2000 aux Etats-Unis
 - ARIB au Japon
 - TTA I et TTA II en Corée du Sud.
- En 1996, le système IS-95 est maintenant opérationnel et les opérations commerciales ont pu démarrer.
- En Septembre 1998, le premier appel effectué avec un terminal mobile WCDMA a été couronné de succès sur un réseau expérimental de Nokia à Tokyo, Japon
- En Mars 1999, l'UIT a approuvé les interfaces radio utilisées par les systèmes de la troisième génération.
- En Juillet 2000, la responsabilité de la norme GSM a été transférée de l'ETSI au 3GPP
- En Décembre 2001, lancement commercial du premier système commercial UMTS en Norvège par l'opérateur Telenor.

Chapitre III

LES PRINCIPES DE BASE DE L'ÉTALEMENT DE SPECTRE

III.1 Les notions de base [2] [5] [9]

L'approche traditionnelle dans les communications numériques consiste à transmettre le plus d'information possible dans une bande de fréquence la plus étroite possible. Par contre, au sein du concept de l'étalement de spectre (SS : Spread Spectrum), l'information à transmettre est étalée à travers une bande de fréquence beaucoup plus large que celle nécessaire avec l'approche traditionnelle.

III.1.1 Les séquences d'étalement

Une séquence d'étalement est une suite d'éléments binaire, cadencé à un rythme largement supérieur au débit des données d'information à traiter. La multiplication de ces dernières avec une séquence d'étalement fait que les données obtenues après cette opération sont étalées à travers une large bande largement supérieure à celle requise pour la transmission des données d'information.

III.1.2 Le chip

L'étalement de spectre est effectué par une fonction qui est indépendante de l'information à transmettre. Cette fonction est matérialisée par une séquence d'étalement et chaque élément binaire de cette séquence d'étalement est appelé "**chip**".

III.1.3 La longueur d'une séquence

Une séquence d'étalement est caractérisée par sa longueur qui n'est autre que le nombre de "chip" contenu dans une période toute entière de la séquence.

III.1.4 Le Gain de traitement (Processing Gain) ou le Facteur d'étalement

C'est le rapport entre la bande de fréquence occupée par le signal étalé par rapport à celle occupée par les bits d'informations. C'est encore le rapport entre les débits de la séquence d'étalement et des bits d'informations.

$$G_p = \frac{R_c}{R_s} = \frac{T_s}{T_c} \quad (3.1)$$

Ce rapport reflète la robustesse du système face aux interférences et les brouillages à bande étroite. En pratique, ce facteur est choisi comme un entier.

III.2 Les différentes techniques d'étalement de spectre [1] [5]

L'étalement de spectre peut être effectué de plusieurs manières, dont les plus utilisées sont :

- ♦ L'étalement de spectre par séquence directe
- ♦ L'étalement de spectre par saut de fréquence

Actuellement, l'étalement par séquence directe est sans doute la méthode la plus populaire à cause de ses qualités. Cependant, nous allons exposer ici les fondements de ces deux méthodes.

III.2.1 L'étalement de spectre par Séquence Directe (DS-SS : Direct Sequence Spread Spectrum)

Généralement les données sont étalées en bande de base, et le signal résultant sert alors à moduler une porteuse dans un second étage. L'opération inverse est réalisée dans le récepteur pour retrouver les données émises.

Dans cette technique, une séquence d'étalement est utilisée conjointement avec un modulateur PSK M-aire (modulateur PSK à M états, où $M = 2^n$ et n est le nombre de bits correspondant à chaque symbole) pour déplacer la phase du signal PSK suivant un rythme R_c dit "débit de chip" ($R_c = \frac{1}{T_c}$, où T_c est la durée d'un symbole élémentaire d'étalement). Ce "débit de chip" est choisi comme multiple entier du "débit de symbole" (symbole d'information) avec $R_s = \frac{1}{T_s}$

La largeur de la bande utilisée pour la transmission est déterminée par le "débit de chip" et par le filtrage effectué en bande de base. Comme nous le verrons plus loin, plus cette largeur de bande est grande, plus le système sera performant. Cependant cette largeur

est limitée par la fréquence de chip R_c maximale autorisée par l'implémentation utilisée (rapidité des circuits et matériels utilisés, puissance de calcul disponible,...).

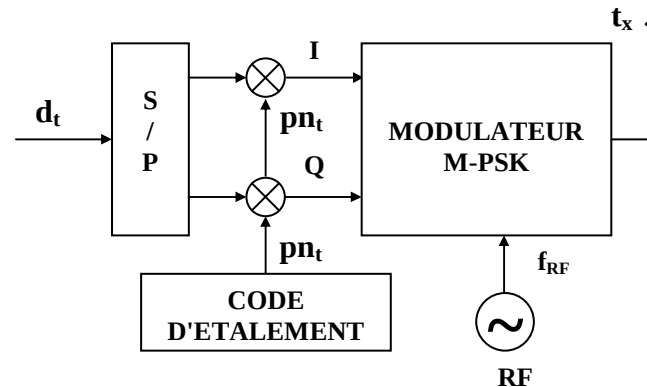


Figure 3.1 : étalement de spectre par séquence directe

Sur cette figure d_t représente les données à transmettre, pn_t la séquence d'étalement utilisée pour l'étalement (la même séquence est utilisée pour étaler les 2 voies I et Q), f_{RF} la porteuse radiofréquence et enfin t_x le signal modulé en M-PSK et à spectre étalé.

La figure 3.2 illustre le processus d'étalement de spectre par séquence directe.

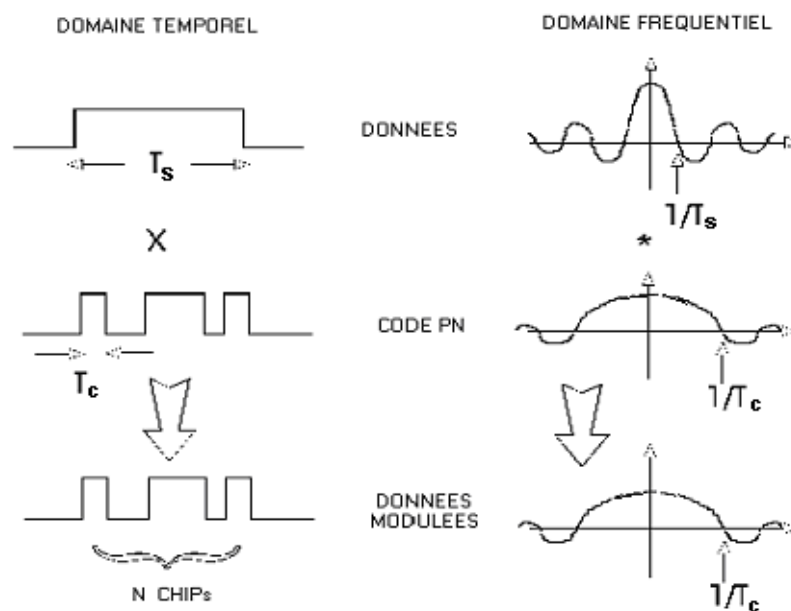


Figure 3.2 : le processus d'étalement par séquence directe

Des signaux binaires codés en NRZ sont utilisés pour représenter les données d'information et les séquences d'étalement. On sait que leurs transformées de Fourier sont des fonction **sinc** (sinus cardinal) avec respectivement des zéros aux points $\frac{1}{T_s}$ et $\frac{1}{T_c}$ (voir paragraphe III.4.2.1). Les zéros au point $R_c = \frac{1}{T_c}$ forment donc un signal à spectre étalé. On peut voir donc que le signal SS possède en effet une largeur de bande beaucoup plus grande que celle du message à transmettre.

Le spectre du signal t_x peut donc être déduit, et est donné sur la figure 3.3.

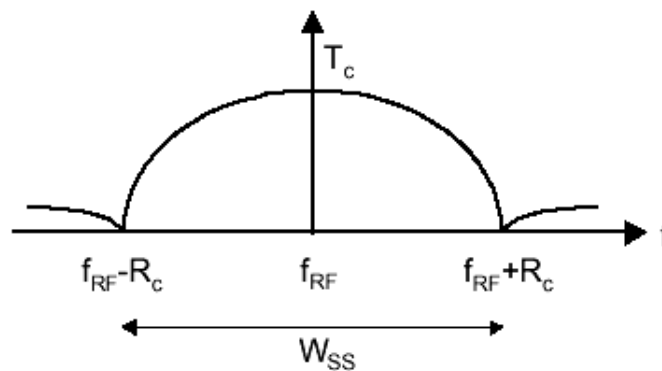


Figure 3.3 : Spectre d'une transmission DS-SS

A ce stade, on peut distinguer deux types de système à étalement de spectre par séquence directe:

- les systèmes à codes courts: la longueur du code est égale à la durée d'un symbole de données. Une nouvelle période du code commence à chaque début d'un symbole.
- les systèmes à codes longs: la longueur du code correspond à la durée de plusieurs symboles de données. Dans ce cas, une réalisation différente du code d'étalement est associée à chaque symbole.

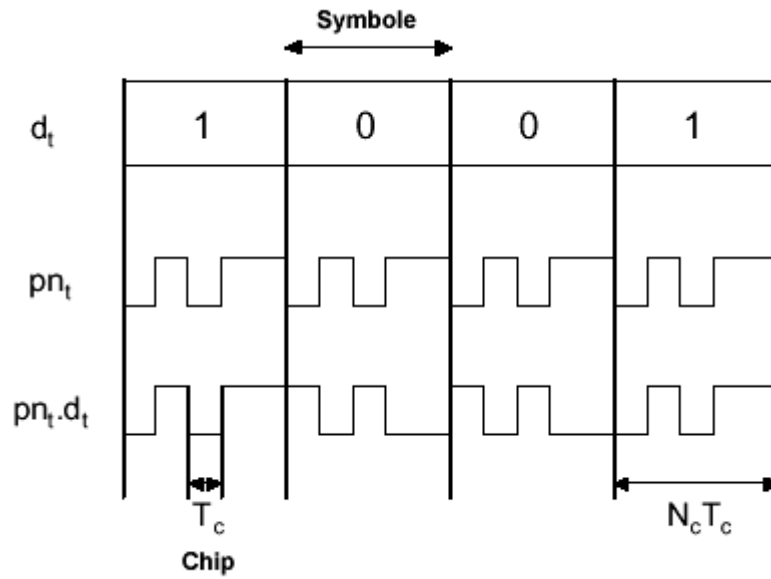


Figure 3.4 : étalement par séquence directe à code court

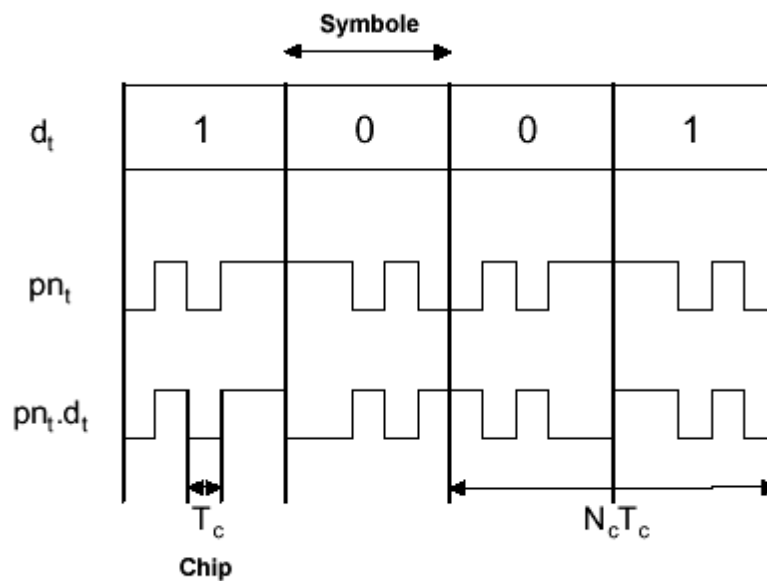


Figure 3.5 : étalement par séquence directe à code long

III.2.2 L'étalement de spectre par saut de fréquence (FH-SS : Frequency Hopping Spread Spectrum)

Dans le cadre de cette technique, une séquence d'étalement est utilisée conjointement avec un modulateur FSK M-aire (modulateur FSK à M états avec $M = 2^n$, où n est le nombre de bits correspondant à un symbole) pour déplacer la fréquence porteuse du signal

FSK suivant un "rythme" R_h . Ce terme sous entend que la porteuse occupe une seule fréquence (comme dans le cas d'une transmission à bande étroite) mais seulement pendant une durée $T_h = \frac{1}{R_h}$. A chaque saut de fréquence, un générateur de séquence d'étalement fournit une séquence de n chips appelée *Frequency Word (FW)* à un synthétiseur de fréquence pour lui dicter une fréquence parmi les 2^n possibles. La fréquence change donc à un rythme égal à R_h par seconde, et on a l'impression que le signal occupe toute une bande de fréquence pour une durée d'observation assez grande. La bande de fréquence occupée par la transmission est donc déterminée par les positions extrêmes des sauts ainsi que par la bande passante ΔF_{ch} occupée par chaque saut.

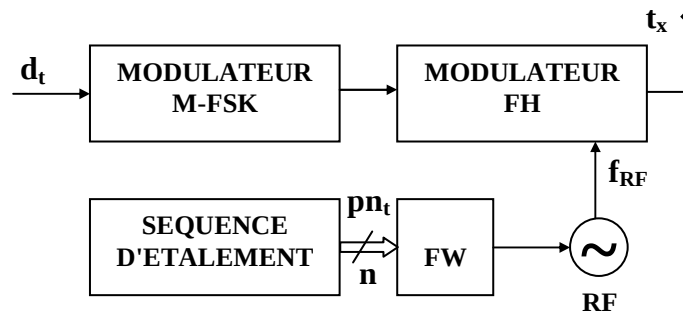


Figure 3.6 : étalement de spectre par saut de fréquence

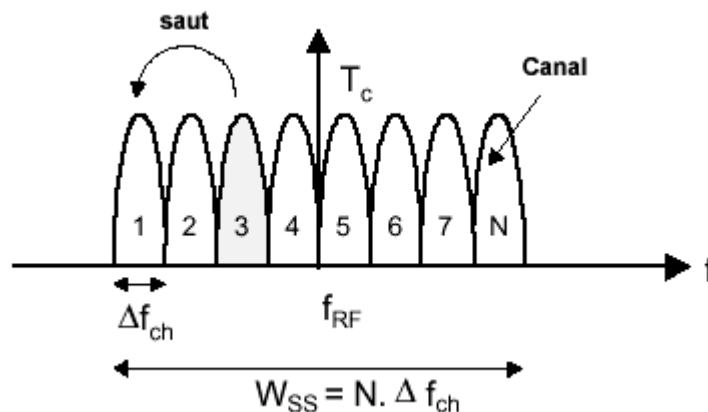


Figure 3.7 : spectre d'une transmission FH-SS

Comme dans le cas de l'étalement par séquence directe, on distingue deux types de systèmes à saut de fréquence:

- les systèmes à saut de fréquence lent : à chaque saut de fréquence correspond plusieurs symboles de données.
- les systèmes à saut de fréquence rapide : à chaque symbole de donnée correspond plusieurs sauts de fréquences

Ceci peut être illustré sur la figure ci-dessous :

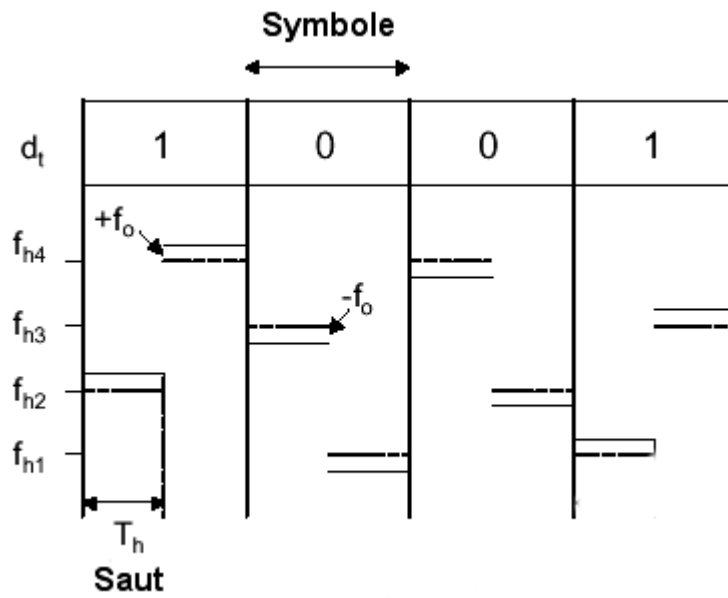


Figure 3.8 : saut de fréquence rapide

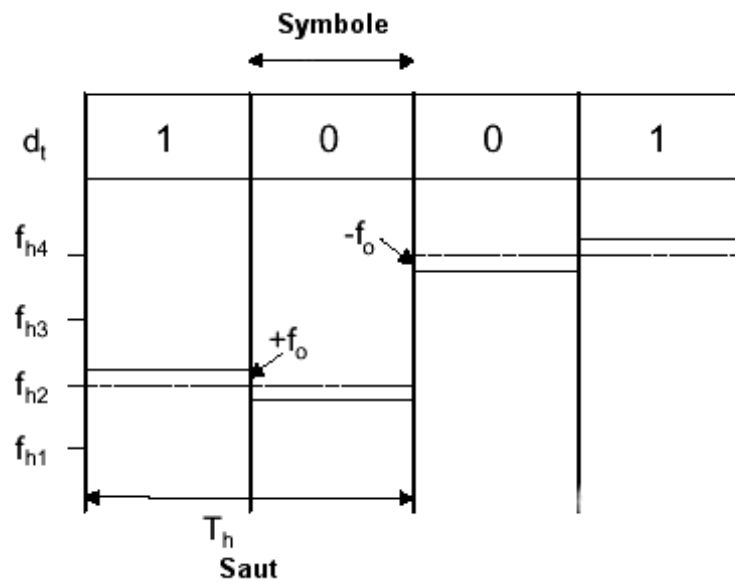


Figure 3.9 : saut de fréquence lent

Remarquons aussi que pour un système à saut de fréquence, la bande utilisée pour la transmission n'est limitée que par les limites en fréquence des divers circuits (oscillateurs, modulateur,.....).

D'autre part, les sauts de fréquences provoquent généralement des discontinuités de phases, et c'est pour cela qu'une démodulation non cohérente est adoptée chez le récepteur.

III.3 Comparaison entre l'étalement par séquence directe et celui par saut de fréquence [6] [7]

La technique d'étalement de spectre par saut de fréquence n'étaie pas le signal. Ce qui fait qu'il n'y a pas de « gain de traitement », qui devrait améliorer le rapport signal sur bruit au niveau du récepteur après l'opération de dé-étalement. En d'autres mots, on a besoin d'émettre plus de puissance pour avoir le même rapport signal sur bruit que l'on aurait obtenu si on avait utilisé l'étalement par séquence directe.

On a aussi plus de difficulté pour synchroniser le récepteur avec l'émetteur car les fréquences et le temps doivent à tout moment correspondre, alors que dans le cas d'une séquence directe, seulement le timing des chips est à prendre en considération. L'utilisation du saut de fréquence conduit donc à une recherche du signal (jusqu'à son verrouillage) qui peut prendre beaucoup de temps. Ceci a pour effet d'allonger inutilement les temps de latence, alors qu'un système à séquence directe peut se verrouiller sur la séquence de chip sur quelques bits seulement.

Généralement, pour rendre la synchronisation initiale possible, le système FH doit rester sur une fréquence fixe avant d'effectuer les sauts (c'est à dire avant toute communication proprement dite). Si le brouilleur se trouve sur cette fréquence (« Parking Frequency »), le système FH ne pourra jamais effectuer les sauts suivants. Et si jamais on a pu effectuer ces sauts, c'est pratiquement impossible de ré-synchroniser lorsque le récepteur perdra la synchronisation.

D'autre part, les systèmes FH coûtent plus chers que les DS parce que les premiers ont besoin de circuits spéciaux (synthétiseur de fréquence très rapide) pour effectuer les sauts de fréquences rapides et la synchronisation y afférente.

Le système FH est par contre préférable pour traiter des problèmes liés aux propagations par trajets multiples. C'est le cas parce que ce système ne reste jamais sur une

seule fréquence. Et il est improbable qu'on obtienne un zéro de transmission pour plusieurs fréquences pendant une période très courte. Ainsi, les systèmes FH combattent mieux les propagations par trajets multiples que les systèmes DS.

Les systèmes FH peuvent transporter plus de données que les systèmes DS pour une même bande de fréquence occupée parce que le signal émis occupe la même bande passante que celle de l'information à transmettre (rappelons-le : il n'y a pas d'étalement du signal)

Cependant, les systèmes FH ne possèdent pas plus d'avantages particuliers que les systèmes DS face aux interférences. Si le premier réduit l'impact des interférences en éliminant tout simplement les brouillages, le second le fait en étalant (donc en réduisant la densité de puissance) les brouillages.

Enfin les systèmes DS sont très sensibles aux effets « Proche-Lointain » (Near-Far effects). Ce problème sera abordé dans le Chapitre VI.

III.4 Détail de fonctionnement d'un système à étalement de spectre par séquence directe

III.4.1 Introduction

Une communication utilisant l'étalement de spectre fait intervenir plusieurs paramètres, quelques-uns sont internes au système, tandis que les autres sont dictés par l'environnement utilisé pour cette communication. Les paramètres internes sont essentiellement constitués par la modulation utilisée (généralement l'une des modulations suivantes: BPSK (Binary Phase Shift Keying), QPSK (Quadrature Phase Shift Keying), D-BPSK (Differential Binary Phase Shift Keying) ou MSK (Minimum Shift Keying)) et la longueur du code utilisé (code long ou code court, voir Chapitre IV pour plus d'explications). Les propagations par trajets multiples induites par le milieu de transmission viennent affecter aussi le système. D'autre part, on a aussi à tenir compte des interférences provoquées par les autres utilisateurs dans le système. Ainsi, un système de communication réel, utilisant l'étalement de spectre ne peut être décrit complètement que par un ensemble complexe de paramètre.

Dans ce chapitre, nous essaierons de mettre en valeurs les propriétés de l'étalement de spectre par séquence directe. Ainsi nous adopterons des hypothèses simplificatrices, pour ne garder que le strict minimum pour pouvoir montrer clairement le fonctionnement d'un système simple à étalement de spectre par séquence directe.

III.4.2 Suppositions

Pour simplifier, nous adoptons les suppositions suivantes :

- La modulation utilisée est la BPSK (Binary Phase Shift Keying)
- Il n'y a pas de propagation par trajets multiples
- Il n'y a que les deux entités (l'émetteur et le récepteur) dans le système, ce qui fait qu'il n'y a aucune interférence.
- Le récepteur est synchronisé sur l'émetteur (du point de vue code et fréquence)
- On utilise un code d'étalement c court (un code long pourra être considéré comme une série de codes courts)
- On utilise un code d'étalement c de longueur N doté des propriétés suivantes :

Auto-corrélation :

$$Ra_c(l) = \sum_{n=0}^{N-1} c(n)c(n+l) = \begin{cases} N & l = 0 \\ 0 & l \neq 0 \end{cases} \quad (3.1)$$

Inter-corrélation :

$$Rc_c(l) = \sum_{n=0}^{N-1} c_1(n)c_2(n+l) = \begin{cases} 0 & c_1 \neq c_2 \quad \forall l \\ 0 & c_1 = c_2 \quad l \neq 0 \\ N & c_1 = c_2 \quad l = 0 \end{cases} \quad (3.2)$$

III.4.3 Système de base [5]

Soit le système montré sur la figure 3.10. Sur cette figure, on voit que l'étalement est effectué en bande de base. C'est alors que le signal résultant est utilisé pour moduler une porteuse. L'opération inverse est effectuée chez le récepteur : le signal reçu est d'abord démodulé, le signal en bande de base obtenu fera ensuite l'objet d'un dé-étalement.

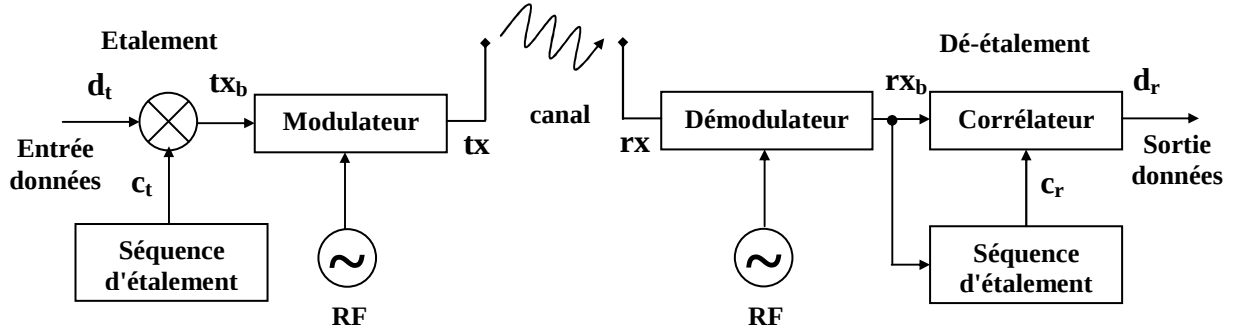


Figure 3.10 : système DS-SS de base

d_t représente les bits de données dont le débit est égal à $R_s = \frac{1}{T_s}$ (débit de symbole, égal au débit binaire R_b pour la modulation BPSK). c_t est la séquence d'étalement dont le débit est égal à $R_c = \frac{1}{T_c}$ (débit de chip, multiple entier du débit de symbole). Ainsi, on a donc $T_s = NT_c$ (étant donné qu'on utilise un code court de longueur N). Les bits de données d_t sont directement multipliés avec le code d'étalement c_t pour produire le signal en bande de base tx_b :

$$tx_b = d_t \cdot c_t \quad (3.3)$$

soit

$$tx_b(t) = d_t \cdot \text{rep}_{T_c} \left[\text{rect} \left(\frac{2t}{T_c} \right) \right] \quad (3.4)$$

où $\text{rep}_{T_c} \left[\text{rect} \left(\frac{2t}{T_c} \right) \right]$ représente la fonction porte de durée $\frac{T_c}{2}$ et de période T_c

Le spectre d'une telle fonction est donnée par

$$TX_B(f) = \frac{1}{2} \sum_{n=-\infty}^{\infty} \text{sinc} \left(\frac{n}{T_c} \frac{T_c}{2} \right) \delta \left(f - \frac{n}{T_c} \right) \quad (3.5)$$

Le signal $tx_b(t)$ possède donc un spectre dont la largeur de la première lobe de

l'enveloppe est égale à $\frac{1}{T_c}$. Donc l'effet de la multiplication de $\mathbf{d}_t$ avec le code d'étalement

$\mathbf{c}_t$ est que la bande de fréquence occupée par le signal en bande de base passe de $\mathbf{R}_s$ à $\mathbf{R}_c$.

Dans le récepteur, l'opération inverse est réalisée. Soit $\mathbf{d}_t(\mathbf{k})$ le k-ième bit émis, et $\mathbf{d}_r(\mathbf{k})$ le bit correspondant attendu à la réception. En supposant que $\mathbf{r}\mathbf{x}_b = \mathbf{A}\mathbf{t}\mathbf{x}_b$ où A est une constante réelle symbolisant l'atténuation, on a :

$$\begin{aligned}\mathbf{d}_r(\mathbf{k}) &= \sum_{n=0}^{N-1} \mathbf{r}\mathbf{x}_b(n) \mathbf{c}_r(n) \\ &= A \sum_{n=0}^{N-1} \mathbf{d}_t(\mathbf{k}) \mathbf{c}_t(n) \mathbf{c}_r(n)\end{aligned}\tag{3.6}$$

Deux cas peuvent alors se présenter :

- $\mathbf{c}_r = \mathbf{c}_t$ et les 2 séquences sont synchronisées (voir les suppositions faites plus haut). Or à l'intérieur de l'intervalle $[0, NT_c]$, $\mathbf{d}_t(\mathbf{k})$ est constante alors on a

$$\begin{aligned}\mathbf{d}_r(\mathbf{k}) &= A \mathbf{d}_t(\mathbf{k}) \sum_{n=0}^{N-1} \mathbf{c}_t(n) \mathbf{c}_t(n) \\ &= A N \mathbf{d}_t(\mathbf{k})\end{aligned}\tag{3.7}$$

On voit donc que le signal émis est reproduit sur $\mathbf{d}_r$. La multiplication du signal à spectre étalé $\mathbf{r}\mathbf{x}_b$ avec le même code d'étalement que celui employé pour l'émission a pour effet de dé-étaler la bande passante occupée par $\mathbf{r}\mathbf{x}_b$ vers une bande égale au débit de symbole $\mathbf{R}_s$.

- $\mathbf{c}_r \neq \mathbf{c}_t$, alors il n'y a aucun dé-étalement. Le signal $\mathbf{d}_r$ reste un signal à spectre étalé. D'où un récepteur qui ne connaît pas la séquence utilisée par l'émetteur ne pourra pas reproduire les données émises.

Chapitre IV

LES CODES D'ÉTALEMENT UTILISÉS DANS UN SYSTÈME DS-CDMA

IV.1 Rappel sur le bruit blanc Gaussien [5]

Le bruit blanc Gaussien centré (de moyenne zéro) possède la même densité spectrale pour toutes les fréquences. L'adjectif "blanc" est utilisé dans le sens où la lumière blanche contient toutes les fréquences de la bande visible des radiations électromagnétiques.

La fonction d'auto-corrélation d'un bruit blanc Gaussien est donnée par la transformée inverse de Fourier de la densité spectrale de puissance du bruit $G_v(f)$:

$$R_{vv}(\tau) = \int_{-\infty}^{+\infty} v(t) \cdot v(t + \tau) dt = F^{-1}\{G_v(f)\} = \frac{N_0}{2} \delta(\tau) \quad (4.1)$$

La fonction d'auto-corrélation $R_{vv}(\tau)$ est égale à zéro pour $\tau \neq 0$. Ceci signifie que deux échantillons différents de bruit blanc Gaussien, aussi semblables soient-ils au moment où on les prend, sont dé-corrélés. Le bruit blanc Gaussien $v(t)$ est totalement dé-corrélé de sa version décalée dans le temps pour tout $\tau \neq 0$.

L'amplitude d'un bruit blanc Gaussien est définie suivant une densité de probabilité $p(v)$:

$$p(v) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\left[\frac{1}{2}\left(\frac{v}{\sigma}\right)^2\right]} \quad (4.2)$$

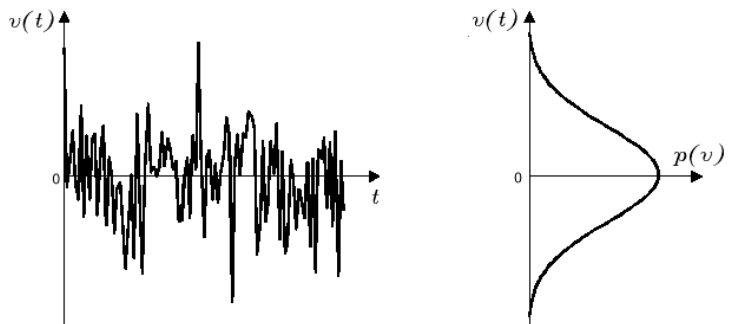


Figure 4.1 : représentation temporelle et densité de probabilité de l'amplitude du bruit blanc

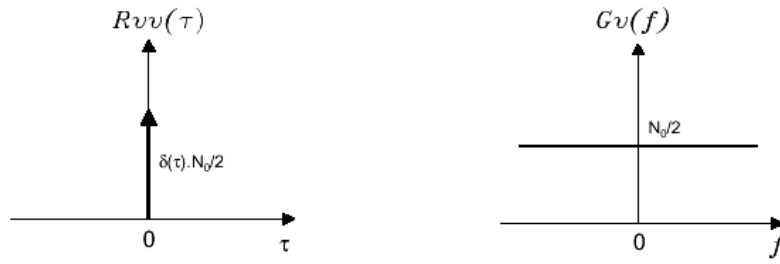


Figure 4.2 : auto-corrélation et spectre du bruit blanc

IV.2 Méthodes de séparation des utilisateurs dans un système DS-CDMA

Il y a deux méthodes pour séparer les utilisateurs dans un système DS-CDMA :

- ◆ Accès multiple orthogonal (utilisé dans les systèmes DS-CDMA synchrones à l'aide de codes orthogonaux)
- ◆ Accès multiple non orthogonal (utilisé dans les systèmes DS-CDMA asynchrones à l'aide des codes non orthogonaux)

IV.2.1 Notion d'orthogonalité des codes [3]

Dans le cas général, on dit que deux signaux périodiques de période T sont orthogonaux quand la valeur de leur fonction d'inter-corrélation est nulle pour un décalage de temps nul:

$$\int_0^T S_1(t)S_2(t)dt = 0 \quad (4.3)$$

Donc deux codes $c_1(n)$ et $c_2(n)$ de longueur N (c'est à dire de période NT_c) sont dits orthogonaux si on a

$$\sum_{n=0}^{N-1} c_1(n)c_2(n) = 0 \quad (4.4)$$

Si un ensemble de codes ne satisfait pas à l'égalité précédente, ces codes sont dits non orthogonaux.

IV.2.2 Propriétés d'auto-corrélation

La fonction d'auto-corrélation d'un code c quelconque est définie comme l'ensemble de nombre de conformité moins le nombre de différence dans une comparaison terme à

terme pour une période entière du code avec une version décalée du code lui-même, en fonction du retard τ .

$$R_a(\tau) = \int_{-N_c T_c / 2}^{N_c T_c / 2} c(t).c(t + \tau) dt \quad (4.5)$$

Idéalement, cette fonction d'auto-corrélation devrait être semblable à celle du bruit blanc : maximale quand il n'y a pas de décalage, et nulle autrement. Elle traduit donc, l'aptitude du code à combattre efficacement les effets des propagations par trajets multiple, étant donné que ceux-ci introduisent des images du signal original mais décalées dans le temps (retard).

IV.2.3 Propriétés d'inter-corrélation

C'est la mesure de la conformité entre 2 codes différents c_i et c_j . Quand cette fonction $R_c(\tau)$ prend la valeur zéro pour tout τ , les codes sont dits orthogonaux. La fonction d'inter-corrélation décrit donc l'interférence entre les codes c_i et c_j :

$$R_c(\tau) = \int_{-N_c T_c / 2}^{N_c T_c / 2} c_i(t).c_j(t + \tau) dt \quad (4.6)$$

Idéalement, cette fonction ne devrait avoir une valeur non nulle que pour $i = j$. Cette fonction traduit donc l'aptitude du code considéré à combattre les interférences introduites par les autres codes employés dans le système.

IV.3 Les codes utilisés pour l'accès multiple orthogonal [3] [5] [13]

Dans ce cas de figure, chaque utilisateur se voit attribuer un ou plusieurs codes orthogonaux. Cependant, ceci requiert une synchronisation entre les utilisateurs parce que ces codes ne sont pas orthogonaux que s'ils sont alignés dans le temps.

IV.3.1 Les codes de Walsh-Hadamard

IV.3.1.1 Méthodes de génération des codes de Walsh-Hadamard

Les codes de Walsh-Hadamard sont générés dans un ensemble de $N = 2^n$ codes de longueur $N = 2^n$. L'algorithme de génération de ces codes est donnée par la relation suivante:

$$\mathbf{H}_{2N} = \begin{bmatrix} \mathbf{H}_N & \mathbf{H}_N \\ \mathbf{H}_N & -\mathbf{H}_N \end{bmatrix} \text{ avec } \mathbf{H}_1 = [1] \quad (4.7)$$

Les lignes (ou les colonnes) de la matrice $\mathbf{H}_N$ constituent les codes de Walsh-Hadamard. Voici quelques exemples de matrices de Walsh-Hadamard.

$$\mathbf{H}_2 = \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix} \quad \mathbf{H}_4 = \begin{bmatrix} 1 & 1 & 1 & 1 \\ 1 & -1 & 1 & -1 \\ 1 & 1 & -1 & -1 \\ 1 & -1 & -1 & 1 \end{bmatrix} \quad \mathbf{H}_8 = \begin{bmatrix} 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 \\ 1 & -1 & 1 & -1 & 1 & -1 & 1 & -1 \\ 1 & 1 & -1 & -1 & 1 & 1 & -1 & -1 \\ 1 & -1 & -1 & 1 & 1 & -1 & -1 & 1 \\ 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 \\ 1 & -1 & 1 & -1 & 1 & -1 & 1 & -1 \\ 1 & 1 & -1 & -1 & 1 & 1 & -1 & -1 \\ 1 & -1 & -1 & 1 & -1 & 1 & 1 & -1 \end{bmatrix}$$

Chaque matrice est donc en fait un ensemble de codes de même longueur (nombre de colonnes), chaque ligne représentant un code. Dans chaque cas, la première ligne (ligne 0) de la matrice est constituée entièrement de "1" et chacune des autres lignes contient $\frac{N}{2}$ "-1" et $\frac{N}{2}$ "1". La ligne $\frac{N}{2}$ débute avec $\frac{N}{2}$ "1" et se termine par $\frac{N}{2}$ "-1".

Notons aussi que les lignes sont mutuellement orthogonales :

$$\sum_{k=0}^{N-1} \mathbf{h}_{ik} \cdot \mathbf{h}_{jk} = 0 \quad (4.8)$$

pour toutes les lignes i et j .

Dans la littérature, on emploie les notations suivantes pour référencer aisément les fonctions de Walsh-Hadamard qui constituent les matrices vues plus haut :

- La longueur des codes ayant été fixée (ou encore le gain de traitement) à N , les fonctions constituant la matrice peuvent être écrites comme suit : **had(n,t)** où $n=0....N$, et n est le numéro de la ligne considérée dans la matrice.
- Cette notation est équivalente à $C_{N,n}$ dans la mesure où cette dernière est plus pratique pour manipuler une structure en arbre (voir plus loin)

IV.3.1.2 Propriétés d'auto-corrélation d'un code de Walsh-Hadamard

La fonction d'auto-corrélation d'un code de Walsh-Hadamard $C_{N,n}$ peut s'écrire en fonction du décalage k

$$R_a(k) = \sum_{j=0}^{N-1} C_{N,n}(j) \cdot C_{N,n}(j+k) \quad (4.9)$$

$$\text{où } k \in \left[-\frac{N}{2}, \frac{N}{2} \right]$$

Cette fonction n'est cependant pas très satisfaisante car il comporte plusieurs pics pour $k \neq 0$. De plus, elle n'est pas homogène pour des codes différents. Pour les figures suivantes, on a pris des codes de longueur $N = 64$ (voir figure 4.3).

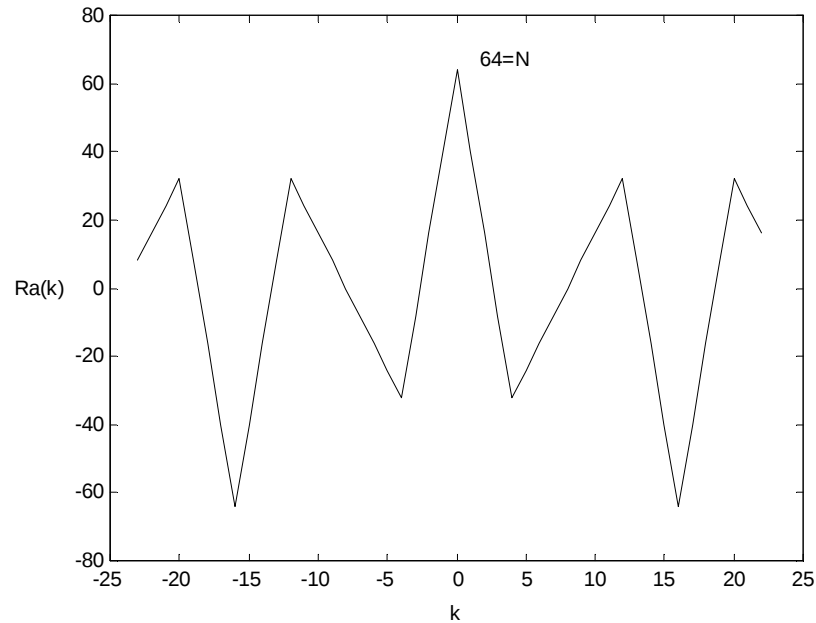


Figure 4.3 : auto-corrélation de **had (29, t)**

Les valeurs d'auto-corrélation des codes de Walsh-Hadamard ne se ressemblent pas pour des codes différents. De plus, ces fonctions d'auto-corrélation comporte de nombreux pics (le nombre varie selon le code considéré) pour des retards non nuls. Sachant que le dispositif de synchronisation se base toujours sur des propriétés d'auto-corrélation du code utilisé, les codes de Walsh-Hadamard ne peuvent être utilisés seuls que dans les environnements synchrones.

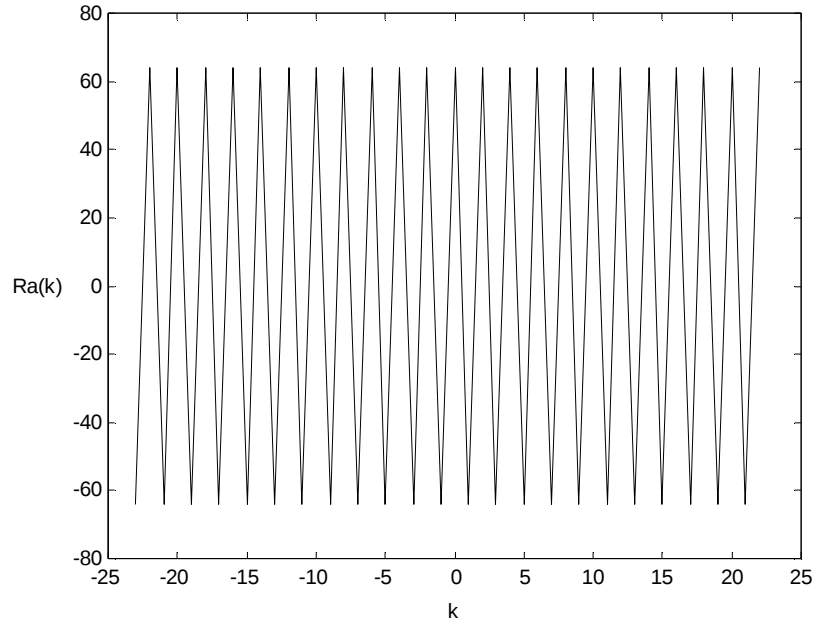


Figure 4.4 : auto-corrélation de **had(2,t)**

IV.3.1.3 Propriétés d'inter-corrélation d'un code de Walsh-Hadamard

La fonction d'inter-corrélation de deux codes $C_{N,n}$ et $C_{N,m}$ peut s'écrire :

$$R_c(k) = \sum_{j=0}^{N-1} C_{N,n}(j) \cdot C_{N,m}(j+k) \quad (4.10)$$

$$\text{où } k \in \left[-\frac{N}{2}, \frac{N}{2} \right]$$

Ici cette fonction prend une valeur nulle lorsque $m \neq n$ et que $k = 0$. C'est le seul point commun de toutes les fonctions d'inter-corrélation possibles, car pour les codes de Walsh-Hadamard ces dernières ne se ressemblent pas. Donc le seul avantage obtenu avec l'utilisation de ces codes, est qu'ils sont orthogonaux lorsqu'ils sont parfaitement alignés dans le temps.

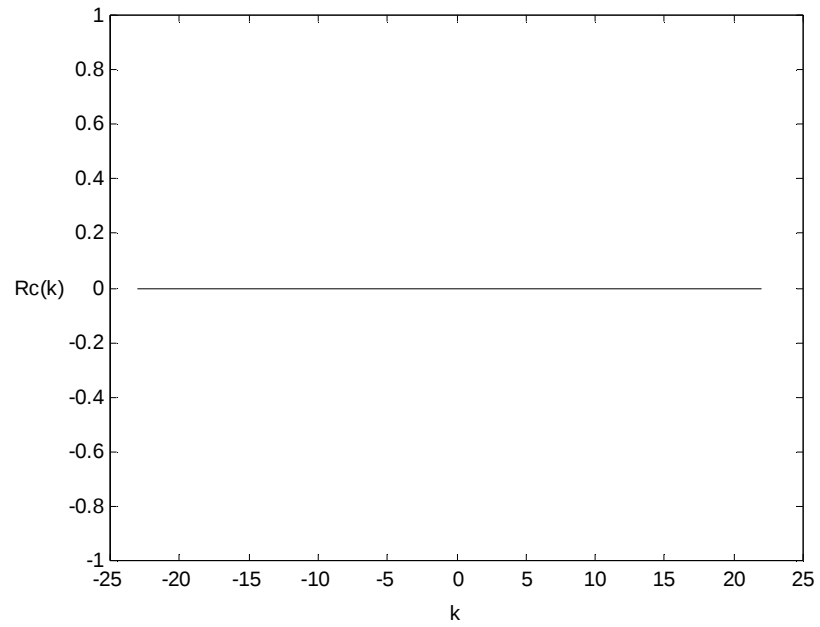


Figure 4.5 : inter-corrélation de **had(21,t)** et **had(59,t)**

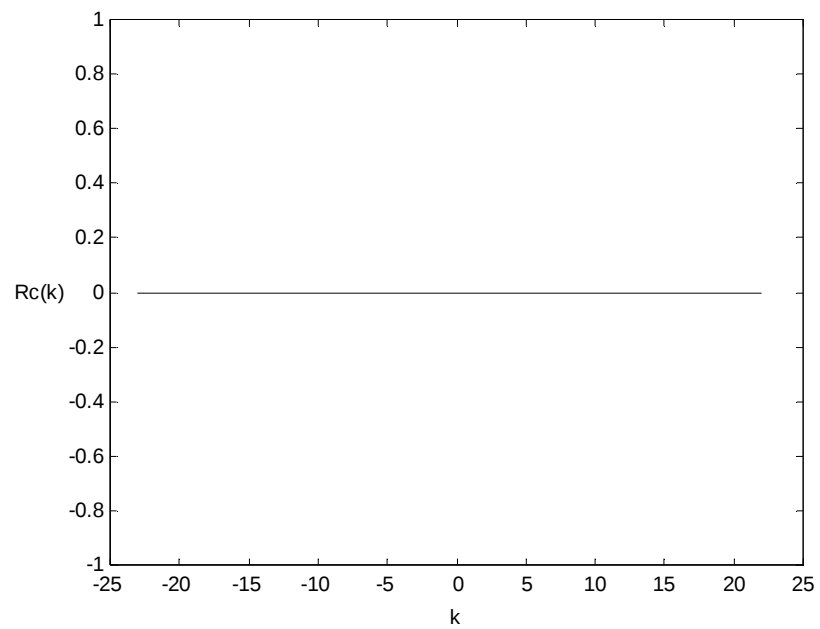


Figure 4.6 : inter-corrélation de **had(59,t)** et **had(29,t)**

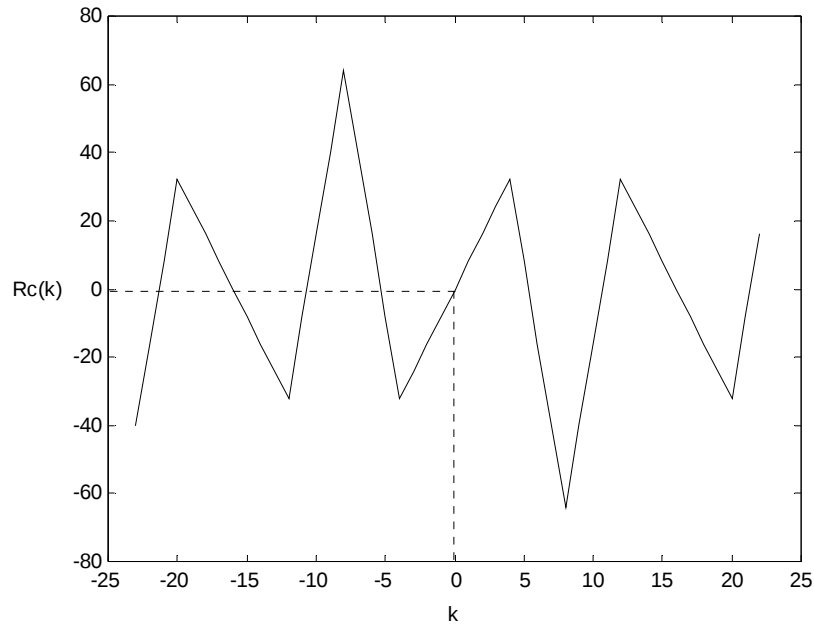


Figure 4.7 : inter-corrélation de **had(21,t)** et **had(29,t)**

Les figures 4.5, 4.6 et 4.7 illustrent ce propos pour $N=64$. On voit sur ces exemples l'irrégularité de la fonction d'inter-corrélation des codes de Walsh-Hadamard. De plus, sur la figure 4.7, on voit que cette fonction peut prendre des valeurs élevées pour $k \neq 0$, ce qui ne nous permet pas de les utiliser dans un environnement asynchrone.

IV.3.1.4 La notion de facteur d'étalement variable

En dépit de ses défauts énumérés précédemment, les codes de Walsh-Hadamard sont importants car ils se trouvent à la base des codes orthogonaux utilisés pour différents facteurs d'étalement. Cette propriété est vraiment intéressante quand on veut partager la même bande de fréquence avec des signaux ayant différents Facteurs d'Etalement (c'est à dire utilisant des codes de longueurs différentes). Les codes qui possèdent cette propriété sont appelés codes OVSF (Orthogonal Variable Spreading Factor). Il est plus aisé de construire ces codes en utilisant une autre approche que la manipulation de matrice. Une structure en arbre permet une meilleure visualisation des relations entre les longueurs des différents codes ainsi que l'orthogonalité entre eux.

Les codes se trouvant sur des niveaux différents sur l'arbre sont de longueurs différentes. Et étant donné que tous les codes l'arbre peuvent être obtenus par une concaténation successive du code actuel avec son inverse, les codes issus d'un code parent ne sont pas orthogonaux à ce code parent.

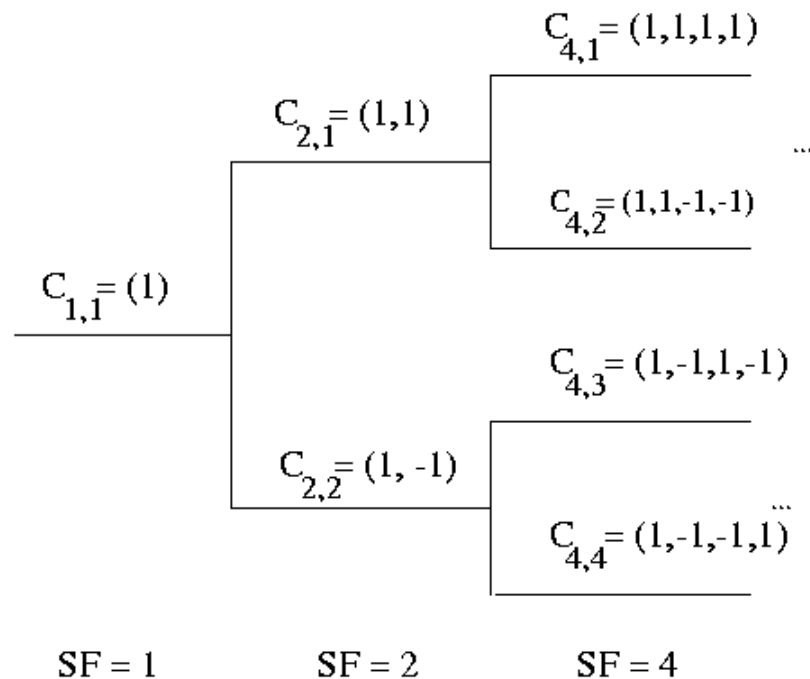


Figure 4.8 : structure en arbre des codes OVSF

L'OVSF consiste à associer des codes de longueurs différentes pour des débits de données différents afin de garder un débit unique lors du multiplexage des divers signaux sur le canal de transmission. Ainsi, les données à haut débit (comme celles des séquences vidéo par exemple) seront associées à des codes de courtes longueurs et les données de faible débit (comme celles de la parole) seront associées à des codes plus longs. De ces constatations, on peut faire quelques remarques :

- Le gain de traitement qu'on peut avoir varie inversement en fonction des débits de données
- Le nombre de codes disponibles pour les transferts diminue au fur et à mesure qu'on monte en débit
- Les codes qui se trouvent au dessous d'un code déjà utilisé sont inutilisables.

Exemple : Le standard UMTS autorise des transferts de données jusqu'à 2 Mbits/s. Cependant dans ce système, le débit chip à atteindre est de 4.096 Mc/s, ce qui nous donne un gain de traitement égal à 2 et un nombre de codes disponibles égal à 2. D'où le transfert à 2 Mbits/s n'est possible que pour un utilisateur seulement à un moment donné (dans une cellule), l'autre code de longueur 2 ne pouvant pas être alloué sous peine de bloquer les appels vocaux et autres services à faible débit.

L'inter-corrélation entre deux codes Walsh-Hadamard appartenant à une même matrice (à une même colonne sur l'arbre) est égale à zéro, dans le cas d'une synchronisation parfaite. De plus, deux codes appartenant à deux branches disjointes sont orthogonaux même s'ils sont de longueurs différentes.

Exemple : prenons les codes $C_{4,2}$ et $C_{2,2}$. Etant donné qu'ils sont de longueurs différentes, on répétera $C_{2,2}$ pour obtenir la même longueur que $C_{4,2}$. Alors on obtient les deux mots de codes :

$$C_{2,2} \Rightarrow (+1 - 1 \mid +1 - 1) \text{ et } C_{4,2} \Rightarrow (+1 + 1 - 1 - 1).$$

Ainsi, on voit que les deux codes sont orthogonaux (après avoir effectué une multiplication élément par élément).

Cependant, la fonction d'auto-corrélation des codes de Walsh-Hadamard ne possède pas de bonnes caractéristiques. Elle comporte plus d'un pic et donc, il n'est pas possible pour le récepteur de détecter le début des codes sans un dispositif de synchronisation externe. L'inter-corrélation est aussi différente de zéro pour un non-alignement temporel entre les deux codes. Ainsi les utilisateurs non synchronisés peuvent interférer avec les autres. C'est pourquoi les codes de Walsh-Hadamard ne peuvent être utilisés seuls que dans le cadre d'un système DS-CDMA synchrone.

IV.4 Les codes utilisés pour l'accès multiple non orthogonal [3] [5] [8] [9] [24] [25]

L'idée derrière ce concept consiste à trouver un ensemble de codes ayant de bonnes caractéristiques de corrélation sans qu'ils soient parfaitement orthogonaux. Les séquences PN répondent à ce critère.

Une séquence PN agit comme une porteuse ayant des propriétés du bruit (mais déterministe). Cette porteuse est employée pour étaler l'énergie du signal à transmettre à travers une bande passante beaucoup plus grande. Le choix d'un bon code PN est important, parce que le type et la longueur du code déterminent les limites en termes de capacité du système.

Une séquence PN (ou code PN) est tout simplement une séquence pseudo-aléatoire composée de 1 et de 0. La fonction d'auto-corrélation des codes PN possède des propriétés similaires à celles du bruit blanc (c'est à dire un seul pic d'auto-corrélation au point 0).

Le terme pseudo-aléatoire vient en fait des considérations suivantes :

- Pas aléatoire, mais la séquence PN apparaît aléatoire pour des usagers qui ne connaissent pas le code
- Déterministe, parce que la séquence est périodique et est connu à l'avance par l'émetteur et le récepteur. Plus la longueur de la période est grande, plus le signal transmis possède une apparence aléatoire, donc difficile à détecter.

IV.4.1 Principe de génération des séquences PN

Un registre à décalage peut être utilisé pour générer une séquence PN, en effectuant une combinaison linéaire appropriée des sorties et en réinjectant ce résultat à l'entrée du registre. Un tel dispositif est appelé "Simple Shift Register Generator" (SSRG). Le SSRG est dit linéaire si la fonction de retour (feed-back) peut s'exprimer par une somme modulo - 2 (XOR)

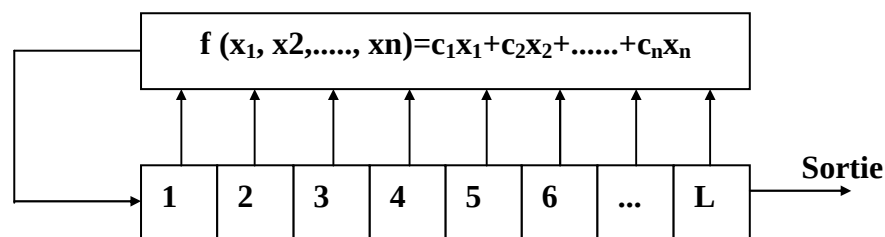


Figure 4.9 : principe de génération des séquences PN

La fonction de retour $f(x_1, x_2, \dots, x_n)$ est une somme modulo-2 des contenus x_i du registre à décalage multipliés par les coefficients c_i ($c_i = 0$ = circuit ouvert, $c_i = 1$ = connecté).

Un SSRG de L étages produit des séquences qui dépendent de la longueur L elle même, des positions des prises de retour (*feed-back tap*) et des conditions initiales.

Le code PN ainsi généré peut être aussi caractérisé par le polynôme

$$P(x) = 1 + P_1 X^1 + P_2 X^2 + P_3 X^3 + \dots + P_n X^n \quad (4.11)$$

où le degré du polynôme est égal au nombre d'étages dans le registre utilisé. Les termes du polynôme $(X^1, X^2, \dots, X^n)$ représente les étages du registre, et P_0 à P_n sont des coefficients binaires qui déterminent les étages impliqués dans la boucle de retour. Par définition, l'étage X^0 correspond toujours à 1 avec un coefficient égal à 1.

Exemple :

Un SSRG de 7 étages, ayant ses prises de retour sur les étages 1, 2,4 et 7 peut être représenté par le polynôme $1 + X + X^2 + X^4 + X^7$

IV.4.2 Notion de longueur d'un code PN

- On parle de code court si la même séquence PN est appliquée à chaque symbole $N_c \cdot T_c = T_s$
- On parle de code long si la période d'une séquence PN entière est beaucoup plus longue que la durée d'un symbole de données. D'où une réalisation toujours différente du code PN est associée à chaque symbole de données $N_c \cdot T_c \gg T_s$.

IV.4.3 Propriétés communes des codes PN

Les codes PN possèdent une certaine ressemblance sur les propriétés d'auto-corrélation. La fonction d'auto-corrélation d'un code PN quelconque s'écrit :

$$R_a(k) = \sum_{j=0}^{N_c-1} pn(j) \cdot pn(j+k) \quad (4.12)$$

$$\text{où } k \in \left[-\frac{N}{2}, \frac{N}{2} \right]$$

et

$$R_a(k) = N_c \quad \text{pour } k = 0 \quad (4.13)$$

$$R_a(k) \ll N_c \quad \text{pour } k \neq 0 \quad (4.14)$$

IV.4.4 Les types de codes PN

IV.4.4.1 Les M-séquences

IV.4.4.1.1 Principe de génération des M-séquences

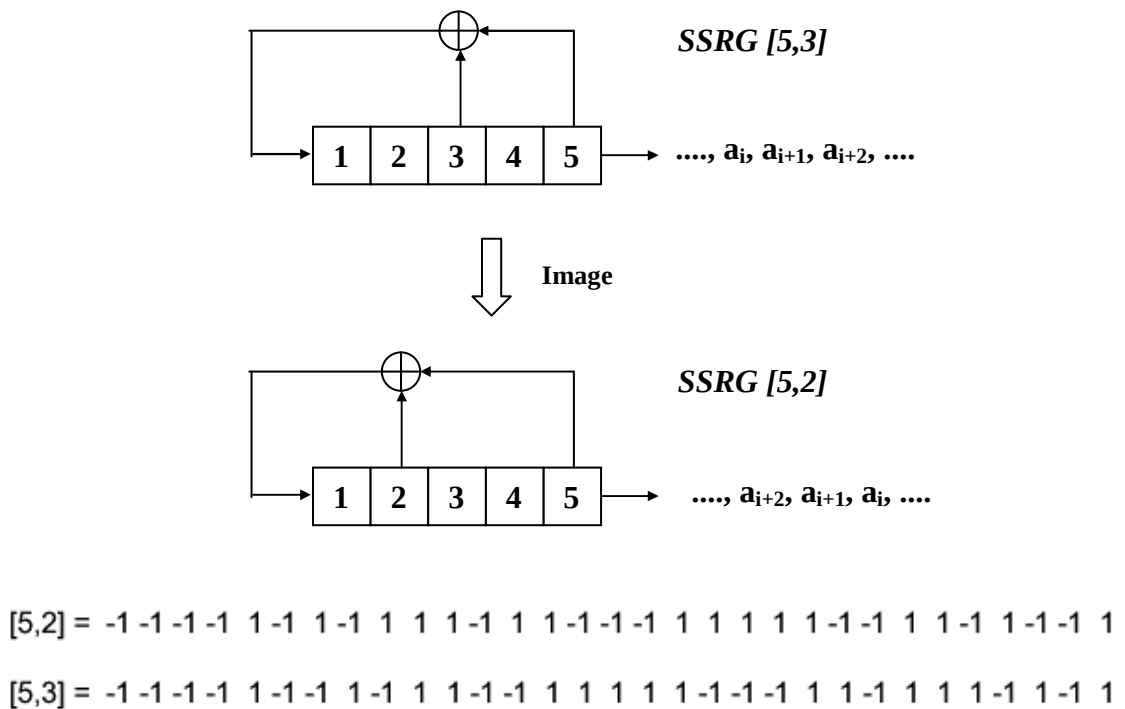


Figure 4.10 : génération de M-séquence avec un SSRG

Quand la période (la longueur) de la séquence PN est exactement égale à $N_c = 2^L - 1$, la séquence PN est dite "**séquence de longueur maximale**" ou encore "**M-séquence**". Une M-séquence générée avec un SSRG linéaire possède un nombre pair de prises de retour. Si un SSRG à L étages comporte des prises de retour sur les étages L, k, m et génère une séquence. ..., a_i , a_{i+1} , a_{i+2} , alors le SSRG inverse doit posséder des prises de retour sur les étages L, L - k, L - m, et la séquence générée devient, a_{i+2} , a_{i+1} , a_i , ...

Dans l'Annexe 1, on peut trouver un extrait de séquences générées avec un SSRG linéaire (sans les ensembles images correspondant à un registre inverse) avec les prises de retour (*feed-back taps*) correspondantes.

IV.4.4.1.2 Propriétés des M-séquences

➤ **Auto-corrélation**

La fonction d'auto-corrélation d'une M-séquence prend la valeur -1 pour toute les valeurs du décalage de la séquence à l'exception du zone qui se trouve entre un décalage à gauche et un décalage à droite de la séquence initiale, à l'intérieur duquel la fonction d'auto-corrélation varie d'une façon linéaire de -1 jusqu'à $2^{L-1} = N_c$ (voir figure 4.11).

Et comme la hauteur du pic d'auto-corrélation augmente avec la longueur N_c de la M-séquence, cette fonction devient alors de plus en plus une approximation de la fonction d'auto-corrélation du bruit blanc quand la longueur de la M-séquence augmente.

➤ **Inter-corrélation**

La fonction d'inter-corrélation de deux M-séquences $\mathbf{pn}_m$ et $\mathbf{pn}_n$ peut s'écrire :

$$R_c(\mathbf{k}) = \sum_{j=0}^{N_c-1} \mathbf{pn}_n(j) \cdot \mathbf{pn}_m(j+\mathbf{k}) \quad (4.15)$$

$$\text{où } \mathbf{k} \in \left[-\frac{N}{2}, \frac{N}{2} \right]$$

Celle-ci n'est pas aussi régulière que la fonction d'auto-corrélation. De plus, cette fonction d'inter-corrélation peut avoir des valeurs élevées pour certaines valeurs de $\mathbf{k}$. Ainsi, dans un système où on utilise les M-séquences comme code d'étalement, l'interférence augmente rapidement avec le nombre d'utilisateurs dans ce système.

La propriété d'inter-corrélation des M-séquences n'est pas utilisée dans la pratique pour la séparation des utilisateurs d'un système. On lui préfère généralement la propriété d'auto-corrélation (qui est très bonne), ce qui limite son utilisation dans le domaine des environnements synchrones.

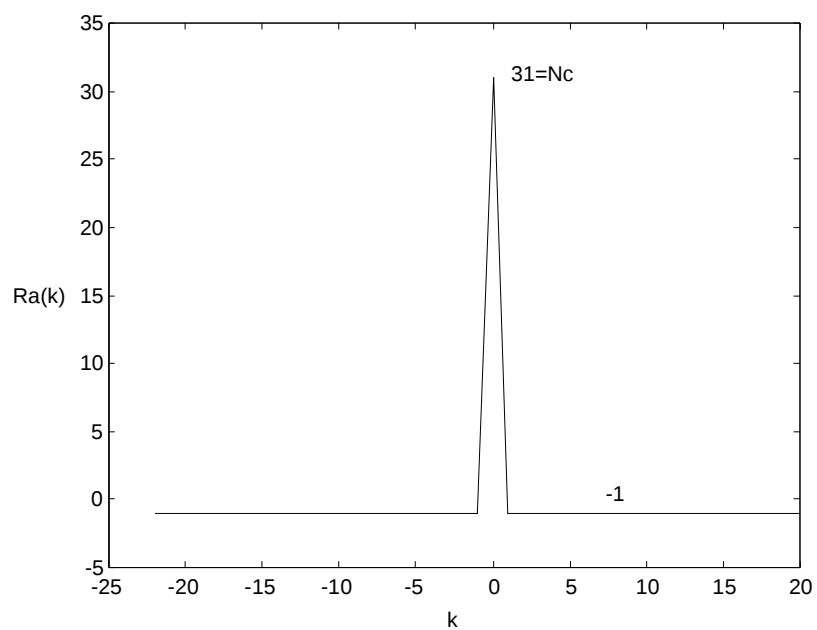


Figure 4.11 : auto-corrélation d'une M-séquence (ici le code SSRG [5,3])

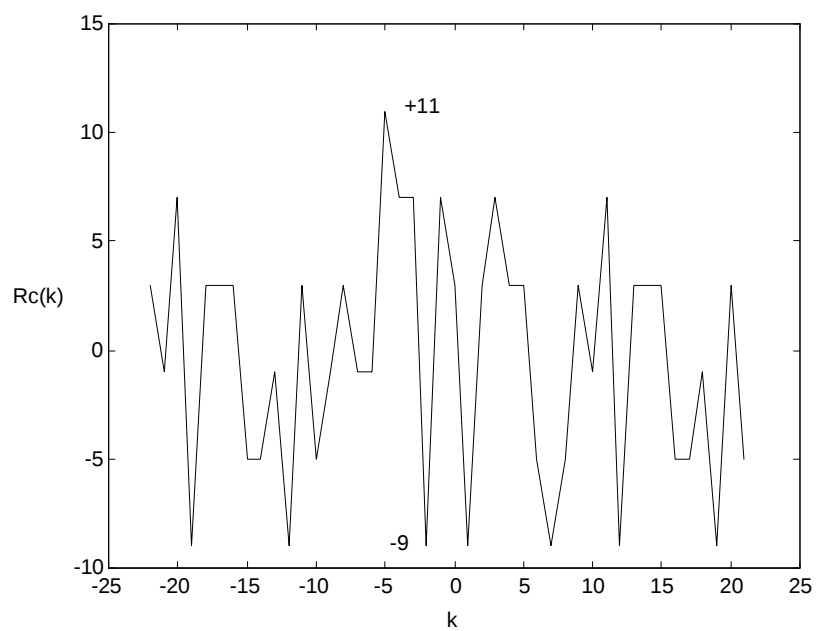


Figure 4.12 : inter-corrélation entre les codes SSRG[5,2] et SSRG[5,3]

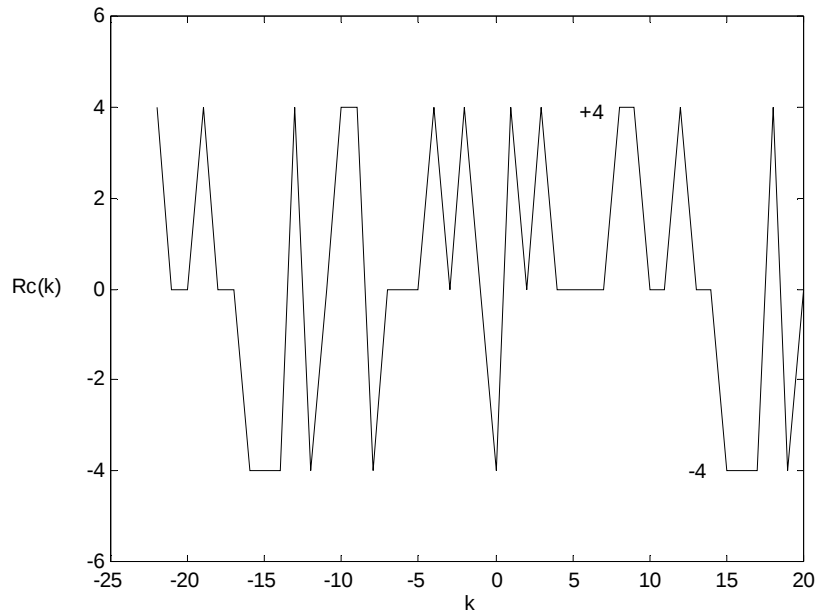


Figure 4.13 : inter-corrélation entre les codes SSRG[5,3] et SSRG[5,4,3,2]

IV.4.4.1.3 Conclusion

Les M-séquences ne possèdent pas de bonnes propriétés d'inter-corrélation. Par contre leurs fonctions d'auto-corrélation peuvent approcher celle du bruit blanc pour une longueur suffisamment grande de la M-séquence.

Cette propriété remarquable est mise à profit dans la pratique, en utilisant la même séquence pour tous les utilisateurs d'un système mais affectée d'un décalage différent pour chaque utilisateur. Mais pour éviter tout faux verrouillage, il faut que l'origine des temps soit la même pour tous les utilisateurs. Ce qui implique une référence de temps commune pour tout le système; autrement dit, le système doit être synchrone.

C'est le cas du système américain IS-95, qui est l'un des ancêtres des systèmes DS-CDMA commerciaux. Dans ce système, les stations de base utilisent la même M-séquence mais décalée différemment dans le temps pour la liaison descendante. La référence de temps commune utilisée est donnée par des récepteurs GPS (Global Positioning System) placés dans chaque station de base.

IV.4.4.2 Les codes de Gold

IV.4.4.2.1 Principe de génération des codes de Gold

Les propriétés d'auto-corrélation des M-séquences ne peuvent pas être améliorées. Or un environnement multi-utilisateur (CDMA) a besoin d'un ensemble de codes qui possèdent la même longueur et avec de bonnes propriétés d'inter-corrélation.

Les codes de Gold conviennent pour cela, car un grand nombre de codes (de même longueur et avec des bonnes propriétés d'inter-corrélation) peut être généré. Ces codes peuvent être générés en utilisant une addition modulo-2 de deux M-séquences de même longueur. Les séquences de codes sont alors additionnées chip par chip. Chaque changement de phases entre les deux codes M-séquences provoque la génération d'une nouvelle séquence.

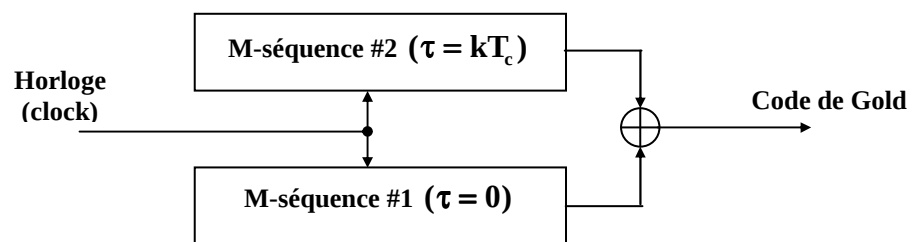


Figure 4.14 : le principe de génération des codes de Gold

Un générateur de code de Gold à base de deux registres de longueur L peut générer $2^L - 1$ séquences (de longueur $N_c = 2^L - 1$) plus les deux séquences de base M-séquences, donnant un total de $2^L - 1$ séquences. En plus de l'avantage de pouvoir générer un grand nombre de codes, les codes de Gold peuvent être choisis de telle façon que parmi un ensemble de codes disponibles à partir d'un générateur donné l'auto-corrélation et l'inter-corrélation entre les codes sont uniformes et limitées. Quand des M-séquences spécialement choisies appelées " M-séquences préférées " (Preferred M-sequences) sont utilisées, les codes de Gold générés possèdent des fonctions d'inter-corrélation qui ne prennent que 3 valeurs. L'Annexe1 fournit un extrait de tels codes avec leurs valeurs d'inter-corrélation normalisées.

Cet important sous-ensemble des codes de Gold est appelé "Preferred Pair Gold Codes (PPGC)". Dans la littérature, le plus souvent, on fait allusion aux codes PPGC quand on

parle de codes de Gold. Dans la suite, nous utiliserons indifféremment les termes « code de Gold » et « code PPGC ».

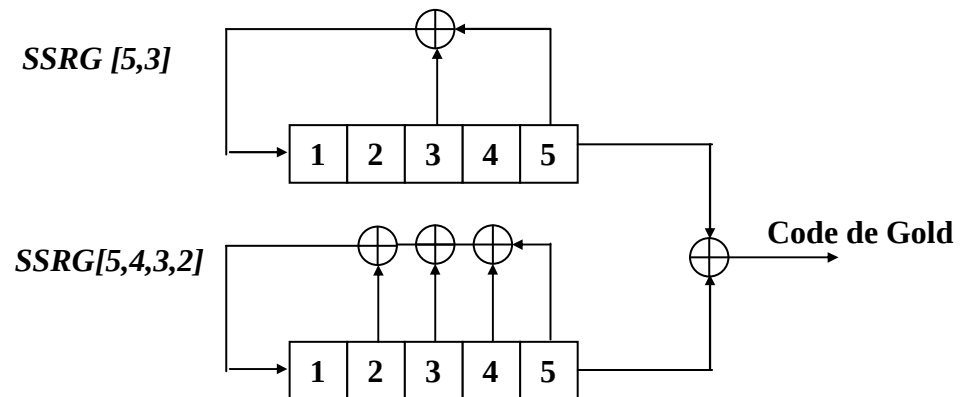


Figure 4.15 : exemple de génération de codes PPGC

L	N_c	Les 3 valeurs d'inter-corrélation normalisées	Fréquence d'occurrence
Impair	$2^L - 1$	$-\frac{1}{N_c}$	0.50
		$-\frac{2^{\frac{L+1}{2}} + 1}{N_c}$	0.25
		$\frac{2^{\frac{L+1}{2}} - 1}{N_c}$	0.25
Pair (différent de $4k$)	$2^L - 1$	$-\frac{1}{N_c}$	0.75
		$-\frac{2^{\frac{L+2}{2}} + 1}{N_c}$	0.125
		$\frac{2^{\frac{L+2}{2}} - 1}{N_c}$	0.125

Tableau 4.1 : propriétés des codes PPGC

Nous procédons comme suit pour la numérotation des codes de Gold de longueur $2^n - 1$:

- Gold (n, 0) désigne le code obtenu à partir des 2 M-séquences sans décalage
- Gold (n, 1) désigne le code obtenu à partir de la première M-séquence et la version décalée (par rapport à la première) de 1 chip de la seconde M-séquence
- Et ainsi de suite.....

IV.4.4.2.2 Propriétés des codes de Gold

➤ **Auto-corrélation**

La fonction d'auto-corrélation d'un code de Gold n'est pas aussi bonne que pour les M-séquences. Cependant, on peut toujours distinguer le pic caractéristique des séquences PN pour $k = 0$. Un exemple est donné à la figure 4.16

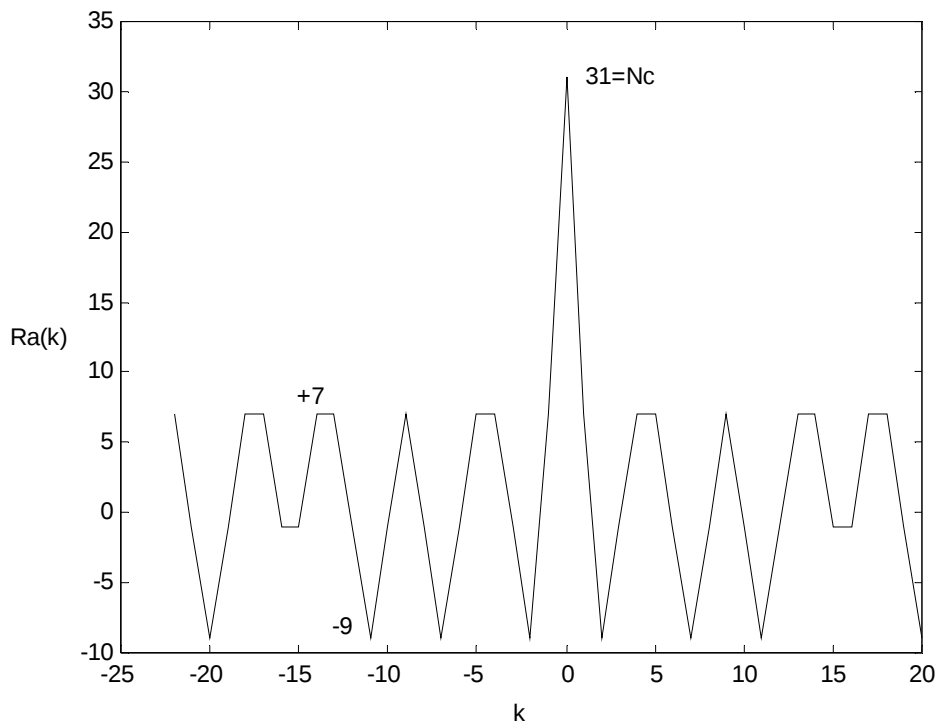


Figure 4.16 : auto-corrélation d'un code Gold(5,0)

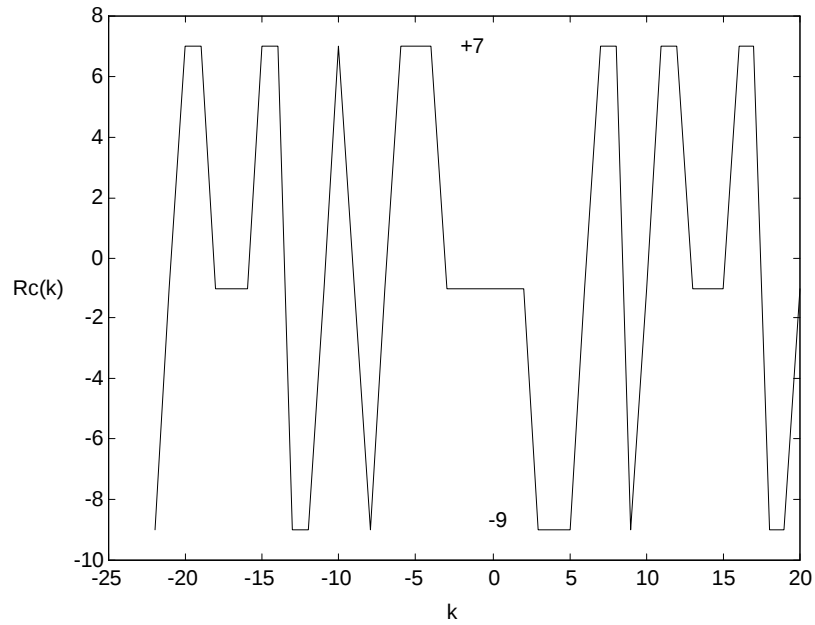


Figure 4.17 : inter-corrélation entre *Gold(5,0)* et *Gold(5,1)*

➤ Inter-corrélation

Comme vu précédemment, les fonctions d'inter-corrélation des codes de Gold ne possèdent que 3 valeurs. Les figures 4.17, 4.18 et 4.19 montrent trois exemples de fonctions d'inter-corrélation de codes de Gold.

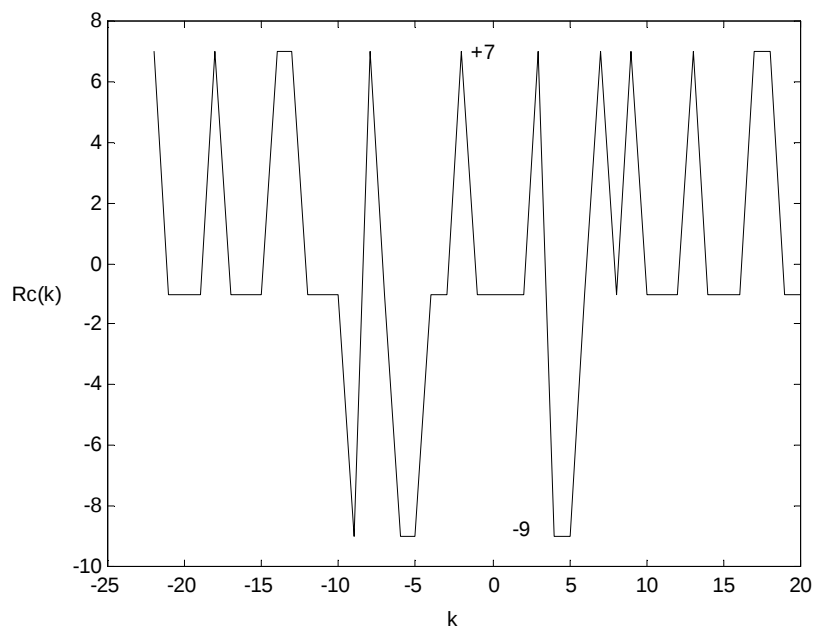


Figure 4.18 : inter-corrélation entre *Gold(5,1)* et *Gold(5,2)*

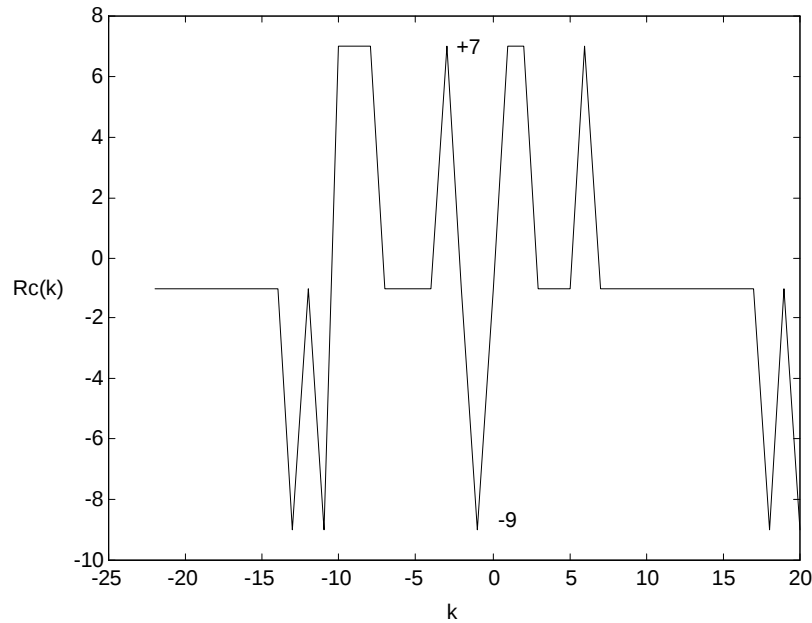


Figure 4.19 : inter-corrélation entre *Gold(5,1)* et *Gold(5,3)*

IV.4.4.2.3 Conclusion

Etant donné que les codes de Gold possèdent de bonnes propriétés d'inter-corrélation et que le nombre de codes disponibles pour une longueur donnée est assez élevé, ils peuvent servir dans des systèmes asynchrones. Cette propriété remarquable offre l'avantage de ne plus requérir une référence commune pour tout le système. Ceci est utilisé par les Européens pour mettre en œuvre l'UMTS sans recourir au système GPS (pour fournir la base de temps commune, s'ils avaient opté pour un système synchrone). C'est une façon pour eux de marquer leur indépendance vis-à-vis des Américains.

IV.4.4.3 Les codes de Kasami

IV.4.4.3.1 Principe de génération des codes de Kasami

Pour une M-séquence **a**, **a'** désigne une séquence obtenue en ne prenant qu'un chip parmi **q** chips (décimation) dans la séquence **a**. En choisissant $\mathbf{q} = 2^{n/2} + 1$, où **n** est le degré du polynôme générateur de la séquence **a**, on constate que **a'** est une séquence périodique de période $2^{n/2} - 1$. En répétant **a'** **q** fois, on obtient une nouvelle séquence **b** est obtenue (**b** est maintenant de période $\mathbf{N}_c = 2^n - 1$). Avec **a** et **b**, on peut former un nouvel ensemble de

séquence en additionnant (modulo-2) la séquence **a** et les $2^{n/2} - 2$ versions décalées de **b**. Incluant **a** et **b**, on aboutit à un ensemble de $2^{n/2}$ codes. C'est "*l'ensemble restreint*" de codes de Kasami (small set of Kasami codes).

Maintenant, soient deux M-séquences **a'** et **a''** obtenues par une décimation de **a** d'ordre $2^{n/2} + 1$ et $2^{(n+2)/2} + 1$ respectivement, et en additionnant **a**, **a'** et **a''** pour différentes valeurs de **a'** et **a''** on aboutit à "*l'ensemble élargi*" de codes de Kasami (large set of Kasami codes).

Le nombre de codes devient $2^{3n/2}$ si **n** = 0 (modulo 4), et $2^{3n/2} + 2^{n/2}$ pour **n** = 2(modulo 4). Cet ensemble "élargi" de codes de Kasami présente donc l'avantage de fournir un très grand nombre de codes, ce qui donne la possibilité de choisir les codes ayant de bonnes propriétés de corrélation. Cependant, comme nous le verrons dans le prochain paragraphe, les codes appartenant à l'ensemble restreint de codes de Kasami possèdent de plus bonnes propriétés de corrélation que ceux de l'ensemble élargi.

Remarquons que l'ensemble restreint est un sous ensemble de l'ensemble élargi (parce qu'une M-séquence peut être considérée comme un cas particulier de codes de Gold). Remarquons aussi que les codes de Kasami n'existent que pour **m** pair.

Dans la suite, nous procéderons comme suit pour la numérotation des codes de Kasami:

- kas (n, m) désigne un code de Kasami de longueur $N_c = 2^n - 1$ appartenant à l'ensemble restreint; m désigne le décalage appliqué à la séquence **a'**.
- kas(n,k,m) désigne un code de Kasami de longueur $N_c = 2^n - 1$ mais appartenant à l'ensemble élargi; m et k désignent respectivement les décalages appliqués aux séquences **a'** et **a''** (par rapport à la séquence **a**).

IV.4.4.3.2 Propriétés des codes de Kasami

➤ **Auto-corrélation**

Les propriétés d'auto-corrélation des codes de Kasami diffèrent selon qu'ils appartiennent à l'ensemble restreint ou à l'ensemble élargi.

- Codes appartenant à l'ensemble restreint

Cet ensemble de codes possède de meilleures propriétés d'auto-corrélation que celles des codes de Gold, mais encore bien plus modestes que celles des M-séquences.

Ici, les fonctions d'auto-corrélation ne prennent que trois valeurs : -1 , $-(2^{n/2} + 1)$, $2^{n/2} + 1$ en dehors de synchronisation. Un exemple est donné à la figure 4.20

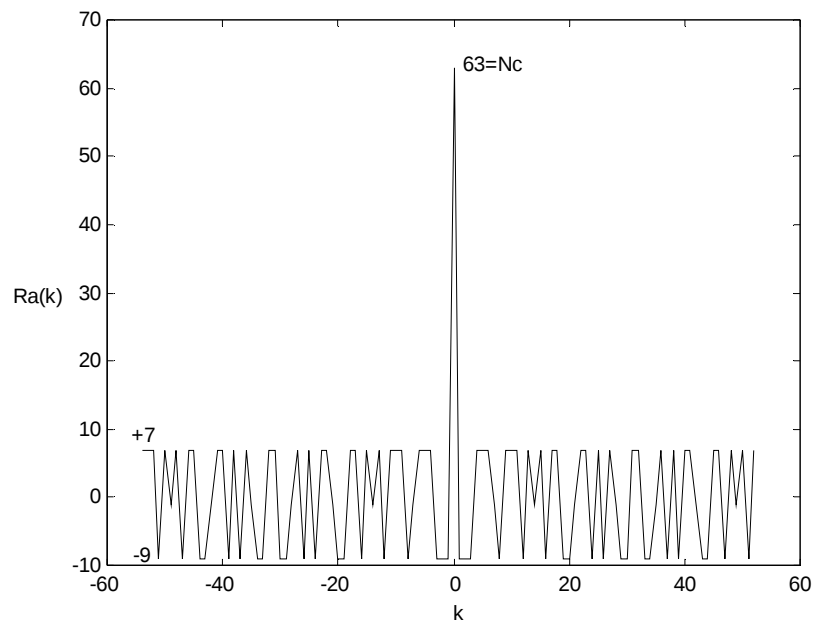


Figure 4.20 : auto-corrélation du code kas (6,1)

- Codes appartenant à l'ensemble élargi

Cet ensemble de codes possède des propriétés d'auto-corrélation similaires à celles des codes de Gold. Mais ici, on a l'avantage de pouvoir choisir parmi un plus grand nombre de codes, ceux ayant les meilleurs comportements.

Les fonctions d'auto-corrélation ne prennent ici que 5 valeurs : -1 , $-1 \pm 2^{n/2}$, $-1 \pm 2^{n/2+1}$.

Un exemple est donné à la figure 4.21

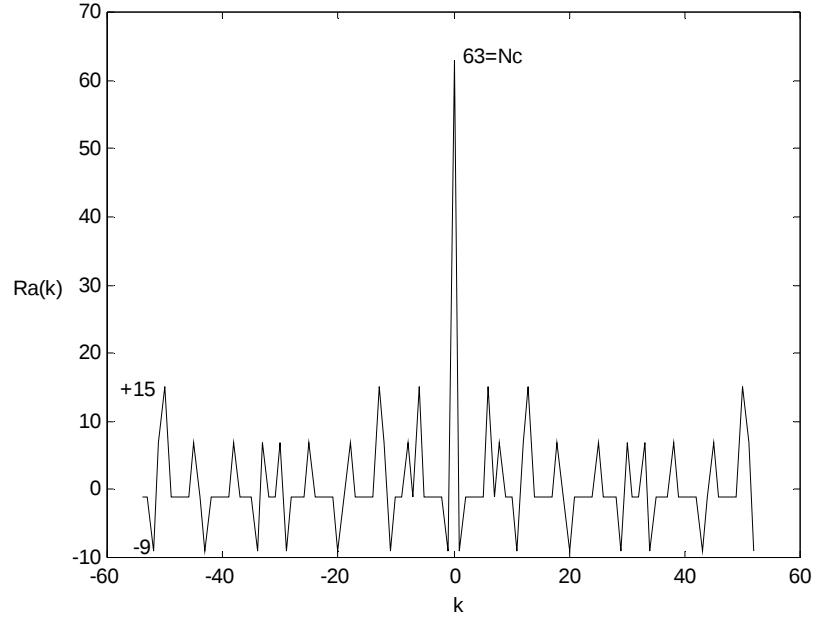


Figure 4.21 : auto-corrélation du code $kas(6,0,1)$

➤ Inter-corrélation

- Codes appartenant à l'ensemble rétreint

Encore une fois, les fonctions d'inter-corrélation ne prennent que 3 valeurs : $-1, -(2^{n/2} + 1), 2^{n/2} + 1$. Cet ensemble possède des propriétés d'inter-corrélation optimale, approchant la limite inférieure définie par Welch. Une limite inférieure sur l'inter-corrélation entre toute paire de séquences binaires de période N_c dans un ensemble de N séquences est donné par:

$$\Phi \geq N_c \sqrt{\frac{N-1}{N_c N - 1}} \quad (4.16)$$

Pour $N_c = 63$ et $N = 8$ on a $\Phi = 7.43$. De l'autre coté, on a $-(2^{n/2} + 1) = -9$ et $2^{n/2} + 1 = 7$, qui se rapprochent manifestement de Φ .

- Codes appartenant à l'ensemble élargi

Cet ensemble possède des valeurs d'inter-corrélation similaires à celles des codes de Gold mais pouvant avoir 5 valeurs : $-1, -1 \pm 2^{n/2}, -1 \pm 2^{n/2+1}$. Un exemple est donné sur la figure 4.23.

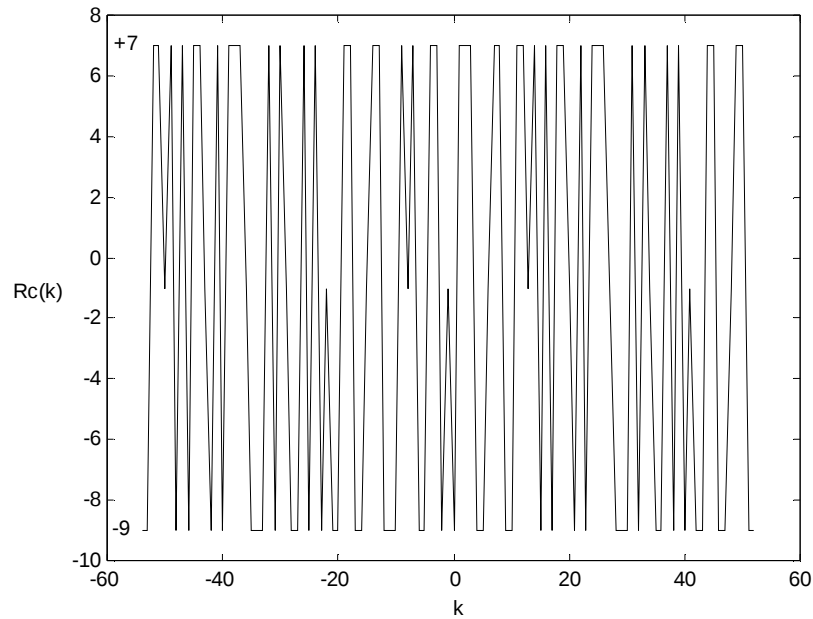


Figure 4.22 : inter-corrélation des codes *kas* (6,1) et *kas* (6,2)

IV.4.4.3.3 Conclusion

Les codes appartenant à l'ensemble restreint possèdent de bonnes propriétés d'inter-corrélation. Cependant, le nombre de codes disponibles est faible par rapport aux codes de Gold. De plus, la génération des codes Kasami est difficile à réaliser du point de vue matériel (l'opération de décimation et de répétition nécessitant des horloges de fréquences largement supérieures à celle du débit de chip à atteindre).

De ce fait, dans la pratique les codes de Kasami appartenant à l'ensemble restreint ne sont utilisés que pour des systèmes dont le nombre d'utilisateurs est faible.

D'un autre côté, les codes appartenant à l'ensemble élargi possèdent le grand avantage de fournir un grand nombre de codes pour une longueur donnée.

Au sein du système UMTS, des codes de Kasami appartenant à l'ensemble élargi de longueur 256 sont utilisées comme alternative aux codes Gold pour la liaison montante dans le cas où la station de base est dotée d'un récepteur multi-utilisateurs.

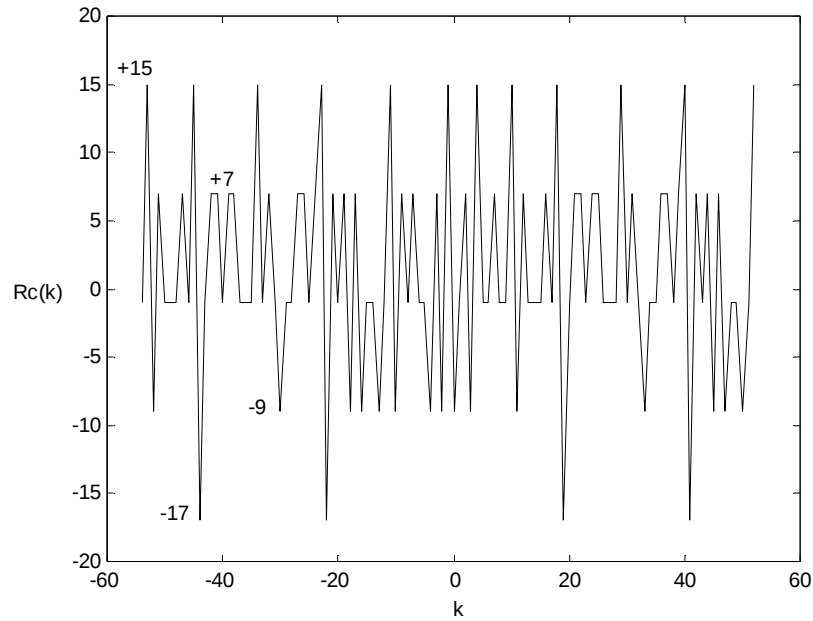


Figure 4.23 : inter-corrélation des codes $kas(6, 0, 1)$ et $kas(6, 0, 2)$

IV.5 L'enchaînement des codes orthogonaux avec des codes PN [12]

Comme nous l'avons vu plus haut, l'auto-corrélation des codes de Walsh-Hadamard pour des retards non nuls laisse vraiment à désirer. Ceci rendra la synchronisation des séquences très difficile. De plus, les fonctions d'inter-corrélation de ces codes sont très irrégulière et ne se ressemblent pas pour tous les codes. Il est impossible de choisir un sous-ensemble de taille raisonnable de codes de Walsh-Hadamard ayant de bonnes propriétés d'inter-corrélation entre chaque paire de codes appartenant à ce sous-ensemble.

D'autre part, les codes PN bien qu'ayant de bonnes propriétés d'auto-corrélation ont le fâcheux défaut de ne pas être orthogonaux (même si les 2 codes considérés sont parfaitement synchronisés)

Pour tirer profit des avantages fournis séparément par les deux types de codes, l'opération d'étalement de spectre se fera successivement avec les deux types de codes : d'abord avec les codes de Walsh-Hadamard, puis avec les codes PN. Du point de vue pratique, les codes de Walsh-Hadamard servent à séparer les canaux dans un système, un code PN sert ensuite à masquer les défauts des codes orthogonaux tout en conservant leurs orthogonalités. C'est pourquoi, dans la littérature les codes orthogonaux sont dits "**codes de canalisation**" et les codes PN des "**codes de brouillage**".

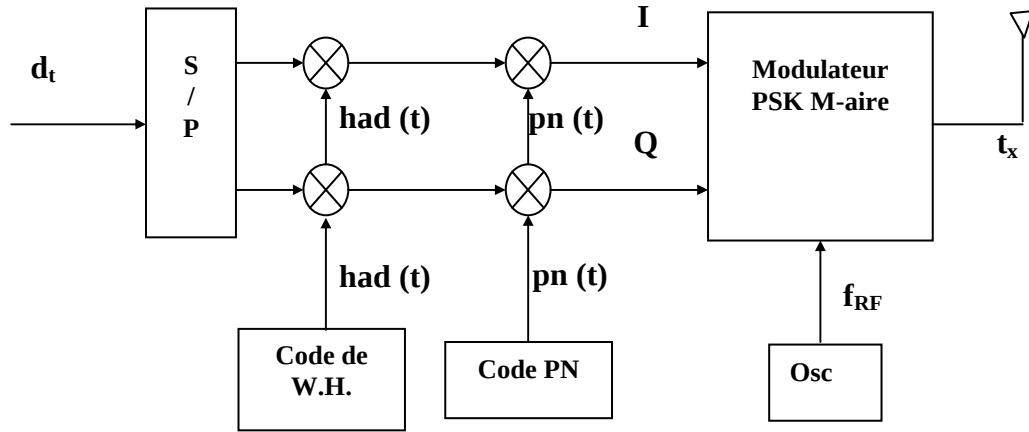


Figure 4.24 : étalement de spectre avec des codes enchaînés

IV.5.1 Propriétés des codes enchaînés

Prenons donc comme code PN, une M-séquence de longueur P (code long). Supposons aussi que la longueur code de Walsh-Hadamard est égal à M, avec $P \gg M$. Le signal à spectre étalé obtenu peut être écrit comme suit :

$$\begin{aligned}
 a_k(t) &= \sum_{j=-\infty}^{\infty} a_j^{(k)} \psi(t - jT_c) \\
 &= \sum_{j=-\infty}^{\infty} d_j w_j^{(k)} \psi(t - jT_c)
 \end{aligned} \tag{4.17}$$

où $a_j^{(k)} \in \{1, -1\}$ est la séquence d'étalement formée par l'enchaînement des codes utilisée pour le canal k , $d_j \in \{1, -1\}$ représente la M-séquence et $w_j^{(k)} \in \{1, -1\}$ le code de Walsh-Hadamard pour le canal k . $\psi(t)$ représente la forme d'onde du chip, c'est à dire :

$$\psi(t) = \begin{cases} 1, & 0 \leq t < T_c \\ 0, & \text{autrement} \end{cases} \tag{4.18}$$

où T_c est la durée d'un chip.

Dans les paragraphes qui vont suivre, nous essaierons de faire une comparaison entre l'utilisation des codes enchaînés et une M-séquence utilisée seule. La comparaison se fera en termes d'aptitude à traiter les interférences causées par les autres codes utilisés. Nous emploierons pour cela des fonctions de corrélations partielles.

IV.5.1.1 Propriétés d'inter-corrélation

Comme nous employons un code PN long, on doit tenir compte du fait que les bits de données successifs ne sont pas multipliés avec la même suite de chips. Chaque chip possible $\mathbf{d}$ de la M-séquence peut ainsi se présenter au début de chaque bit de donnée, avec $0 \leq \mathbf{d} \leq \mathbf{P} - 1$.

Comme dans [12] nous définissons la fonction d'inter-corrélation partielle (ou d'auto-corrélation partielle si $\mathbf{k} = \mathbf{i}$) comme suit:

$$\mathbf{R}_{\mathbf{k},\mathbf{i}}(\mathbf{d},\tau) = \int_0^\tau \mathbf{a}_{\mathbf{k}}(\mathbf{t} + \mathbf{dT}_c - \tau) \mathbf{a}_{\mathbf{i}}(\mathbf{t} + \mathbf{dT}_c) d\mathbf{t} \quad (4.19)$$

$$\hat{\mathbf{R}}_{\mathbf{k},\mathbf{i}}(\mathbf{d},\tau) = \int_\tau^T \mathbf{a}_{\mathbf{k}}(\mathbf{t} + \mathbf{dT}_c - \tau) \mathbf{a}_{\mathbf{i}}(\mathbf{t} + \mathbf{dT}_c) d\mathbf{t} \quad (4.20)$$

Notre but est de calculer la variance du terme d'interférence $\mathbf{b}_{-1}^{(\mathbf{k})} \mathbf{R}_{\mathbf{k},\mathbf{i}}(\mathbf{d},\tau) + \mathbf{b}_0^{(\mathbf{k})} \hat{\mathbf{R}}_{\mathbf{k},\mathbf{i}}(\mathbf{d},\tau)$ où $\mathbf{b}_0^{(\mathbf{k})}$ et $\mathbf{b}_{-1}^{(\mathbf{k})}$ sont respectivement le bit considéré et le bit précédent de l'utilisateur $\mathbf{k}$. Ces bits de données sont supposés être des variables indépendants et identiquement distribués avec des probabilités égales de prendre des valeurs 1 ou -1.

La variance en question est donc donnée par l'expression :

$$\mathbf{E} \left\{ \left[\mathbf{b}_{-1}^{(\mathbf{k})} \mathbf{R}_{\mathbf{k},\mathbf{i}}(\mathbf{d},\tau) + \mathbf{b}_0^{(\mathbf{k})} \hat{\mathbf{R}}_{\mathbf{k},\mathbf{i}}(\mathbf{d},\tau) \right]^2 \right\} = \mathbf{E} \left[\mathbf{R}_{\mathbf{k},\mathbf{i}}^2(\mathbf{d},\tau) \right] + \mathbf{E} \left[\hat{\mathbf{R}}_{\mathbf{k},\mathbf{i}}^2(\mathbf{d},\tau) \right] \quad (4.21)$$

Après des calculs [12], on aboutit aux résultats suivants:

➤ Si $\tau > T_c$, alors on a :

$$\mathbf{E} \left\{ \left[\mathbf{b}_{-1}^{(\mathbf{k})} \mathbf{R}_{\mathbf{k},\mathbf{i}}(\mathbf{d},\tau) + \mathbf{b}_0^{(\mathbf{k})} \hat{\mathbf{R}}_{\mathbf{k},\mathbf{i}}(\mathbf{d},\tau) \right]^2 \right\} = \begin{cases} \frac{2}{3} \mathbf{MT}_c^2, & \tau = \mathbf{k}' T_c \text{ où } \mathbf{k}' \in \mathbf{N} \\ \mathbf{MT}_c^2, & \text{autrement} \end{cases} \quad (4.22)$$

Cette variance est identique à celle obtenue lors de l'utilisation des M-séquences seules.

➤ Si $\tau = 0$ on a donc à calculer une fonction d'auto-corrélation périodique:

$$\mathfrak{R}_{k,i}(\mathbf{d}, \mathbf{t}) = \int_0^T \mathbf{a}_k(\mathbf{t} - \tau) \mathbf{a}_i(\mathbf{t}) d\mathbf{t} \quad (4.23)$$

ce qui donne pour les codes enchaînés :

$$\mathfrak{R}_{k,i}(\mathbf{d}, 0) = \int_0^T \mathbf{d}_i(\mathbf{t}) \mathbf{d}_k(\mathbf{t}) \mathbf{w}_i(\mathbf{t}) \mathbf{w}_k(\mathbf{t}) d\mathbf{t} = 0 \quad (4.24)$$

à cause de l'orthogonalité des codes de Walsh-Hadamard. Pour une M-séquence seule, on aurait eu:

$$\mathbb{E}\{\mathfrak{R}_{k,i}^2(\mathbf{d}, 0)\} = \mathbf{M}\mathbf{T}_c^2 \quad (4.25)$$

➤ Si $0 \leq \tau \leq T_c$ on a :

$$\mathbb{E}\left\{\left[\mathbf{b}_{-1}^{(k)} \mathbf{R}_{k,i}(\mathbf{d}, \tau) + \mathbf{b}_0^{(k)} \hat{\mathbf{R}}_{k,i}(\mathbf{d}, \tau)\right]^2\right\} = \begin{cases} \frac{1}{3} \mathbf{M}\mathbf{T}_c^2, & \text{pour le code enchaîné} \\ \frac{2}{3} \mathbf{M}\mathbf{T}_c^2, & \text{pour la M séquence} \end{cases} \quad (4.26)$$

IV.5.1.2 Propriétés d'auto-corrélation

Etant donné que l'auto-corrélation reflète l'influence des propagations par trajets multiple (retards) sur le signal considéré (c'est à dire lorsque $\mathbf{k} = \mathbf{i}$), nous ignorons le cas où $\tau = 0$. Pour les codes enchaînés lorsque $\tau > T_c$, la variance de la fonction d'auto-corrélation est la même que celle de l'inter-corrélation [12]. Ceci signifie que le comportement insatisfaisant des codes de Walsh-Hadamard (ayant une variance beaucoup plus grande que $\frac{2}{3} \mathbf{M}\mathbf{T}_c^2$) disparaît.

Pour la M-séquence, l'auto-corrélation possède la même variance que celle obtenue pour l'inter-corrélation.

IV.5.1.3 Résumé des résultats précédents

Les codes enchaînés possèdent un meilleur comportement par rapport aux M-séquences seules car ceux-ci possèdent une valeur d'inter-corrélation nulle (c'est à dire que

ceux-ci sont orthogonaux) pour $\tau = 0$. Ils peuvent être donc plus performants que les M-séquences.

De plus, le comportement insatisfaisant des codes de Walsh-Hadamard est masqué par l'ajout de la M-séquence. L'emploi des codes enchaînés se révèle donc bénéfique, dans la mesure où ils tirent profit des atouts de chacun des codes qui les composent.

IV.5.2 Conséquences de l'utilisation de codes enchaînés sur la capacité d'un système

Une étude sur l'influence de l'utilisation de codes enchaînés sur la capacité d'un système multicellulaires a été effectuée dans [12]. Le système étudié est constitué d'un réseau à 18 cellules, soumis à des fadings dus aux propagations par trajets multiples. Pour cela, on a utilisé les modèles de Rayleigh et de Rice. Le modèle de Rayleigh est adopté pour modéliser le cas du "mobile encaissé" (le mobile ne reçoit pas d'onde directe, mais seulement les ondes réfléchies). Ce modèle est donc utilisé dans le cas où la visibilité directe n'existe pas. Le modèle de Rice permet de tenir compte de l'existence d'une visibilité directe. Selon [12], et d'après des résultats de mesures récentes, dans un environnement urbain l'enveloppe des signaux suivent plutôt une distribution de Rice.

L'exposant de perte de propagation est noté par γ . Dans le cas du Fading de Rayleigh, on a adopté 2 cas ayant chacun trois trajets de puissances relatives (0, -6, -8 dB) et (0, -2, -3 dB), le premier trajet étant pris comme référence. Pour le fading de Rice, le facteur de Rice utilisé est **R = 7 dB** (c'est le rapport entre la puissance de l'onde directe et celle résultante des ondes réfléchies).

Notre but consiste à comparer les capacités du réseau, lorsqu'on utilise des codes enchaînés et lorsqu'on n'utilise que des M-séquences. Les résultats peuvent être donnés sous forme d'amélioration par rapport au cas où on n'utilise que des M-séquences seules.

	$\gamma=2$	$\gamma=4$
Rayleigh (0, -6, -8 dB)	7.7%	19.2%
Rayleigh (0, -2, -3 dB)	13.3%	26.7%
Rice (R=7 dB)	83.3%	255.9%

Tableau 4.2 : amélioration de la capacité du réseau pour différents paramètres

IV.5.2.1 Commentaires et interprétation de ces résultats

L'amélioration de la capacité est beaucoup plus significative quand les puissances des ondes réfléchies augmentent. Ceci est dû au fait qu'une grande partie des interférences provoquées par les données destinées aux autres utilisateurs a été éliminée par l'orthogonalité des codes enchaînés. D'une manière similaire, on assiste à une amélioration remarquable de la capacité dans le cas du Fading de Rice. Ceci peut être expliquée par le fait que l'existence de la visibilité directe introduit une hausse des interférences provoquées par les données destinées aux autres utilisateurs.

On voit aussi que l'amélioration de la capacité augmente lorsque l'exposant de perte de propagation augmente, étant donné que l'interférence provoquée par les autres cellules (qui ne peut être éliminée par l'utilisation des codes enchaînés) devient moins significative.

En guise de conclusion, l'augmentation remarquable de la capacité, plus particulièrement dans le cas du fading de Rice (caractéristique des milieux urbains à visibilité directe), justifie l'utilisation des codes enchaînés dans les réseaux cellulaires DS-CDMA. La plupart des systèmes DS-CDMA actuels (IS-95, UMTS,...) utilisent des codes enchaînés à cause de leurs performances.

Chapitre V

EMISSION, RECEPTION et SYNCHRONISATION

V.1 Architecture générale d'un émetteur DS-SS [5]

Considérons ici une architecture typique d'un émetteur utilisant l'étalement par séquence directe (DS-SS) :

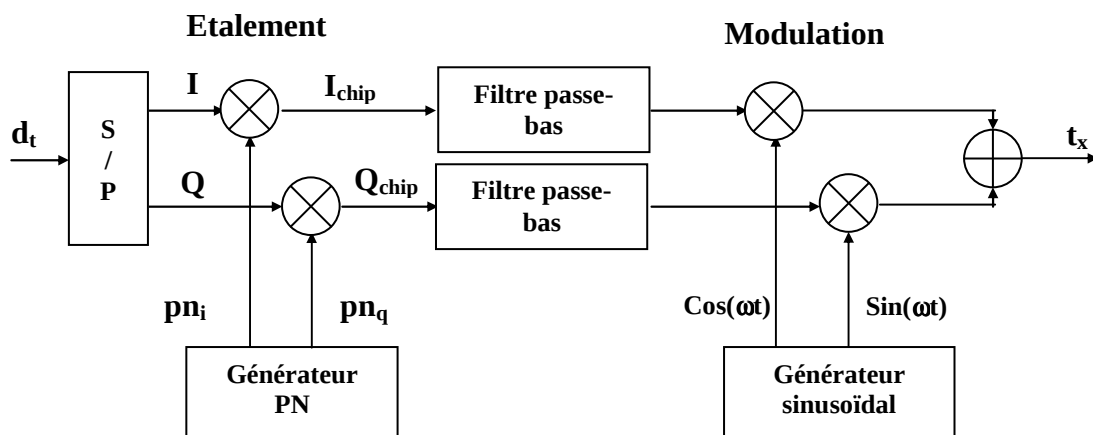


Figure 5.1 : synoptique d'un émetteur DS-SS

L'étalement de spectre est effectué en bande de base, puis le signal obtenu sert à moduler une porteuse pour pouvoir être émis. Sur la figure précédente, d_t représente les données à transmettre.

V.2 Architecture générale d'un récepteur DS-SS [5]

Une architecture typique d'un récepteur DS-SS est montrée sur la figure 5.2. Sur cette figure, on a représenté l'intérieur du bloc de dé-étalement pour une raison de clarté. Ici, l'opération de dé-étalement est effectuée par de simples corrélateurs (voir plus loin). Mais dans le cas général, on peut le faire avec d'autres techniques comme nous la verrons dans le prochain paragraphe. Cette figure sert donc à décrire d'une manière plus générale et plus claire la réception des signaux à spectre étalé.

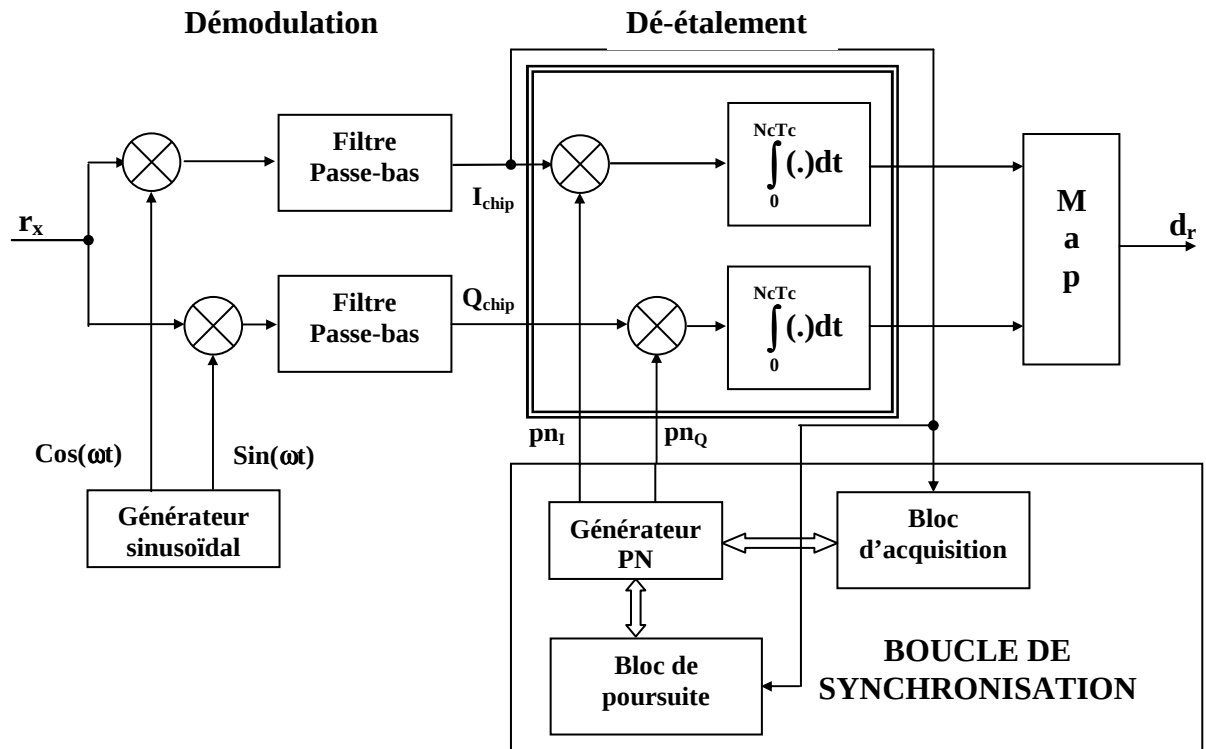


Figure 5.2 : synoptique d'un récepteur DS-SS

Les blocs de base constituant le récepteur DS-SS sont :

- Un démodulateur IQ cohérent avec un synthétiseur de forme d'onde (synthèse numérique directe) calé sur la fréquence intermédiaire (f_{IF})
 - Des filtres destinés à filtrer les séquences de « chip » (ce sont en général des filtres à cosinus surélevé)
 - Le bloc de dé-étalement (corrélation entre les symboles reçus et les séquences PN générées localement pn_I et pn_Q)
- Des boucles de synchronisation pour la porteuse fréquence intermédiaire (f_{IF} , la mesure de l'erreur de phase $\Delta\phi_{IF}$ après l'opération de dé-étalement permet de réduire l'influence du bruit sur la boucle) et pour la fréquence de « chip » (f_{chip})

V.3 L'opération de dé-étalement [5]

Il y a au moins 2 architectures utilisées pour dé-étaler le signal :

- le « matched filter »

- le corrélateur actif (« *active correlator* »)

V.3.1 Le « *Matched Filter* »

Généralement, le “matched filter” effectue une convolution en utilisant un filtre à réponse impulsionnelle finie (RIF) dont les coefficients ne sont rien d’autre que la séquence PN attendue prise dans l’ordre inversé dans le temps.

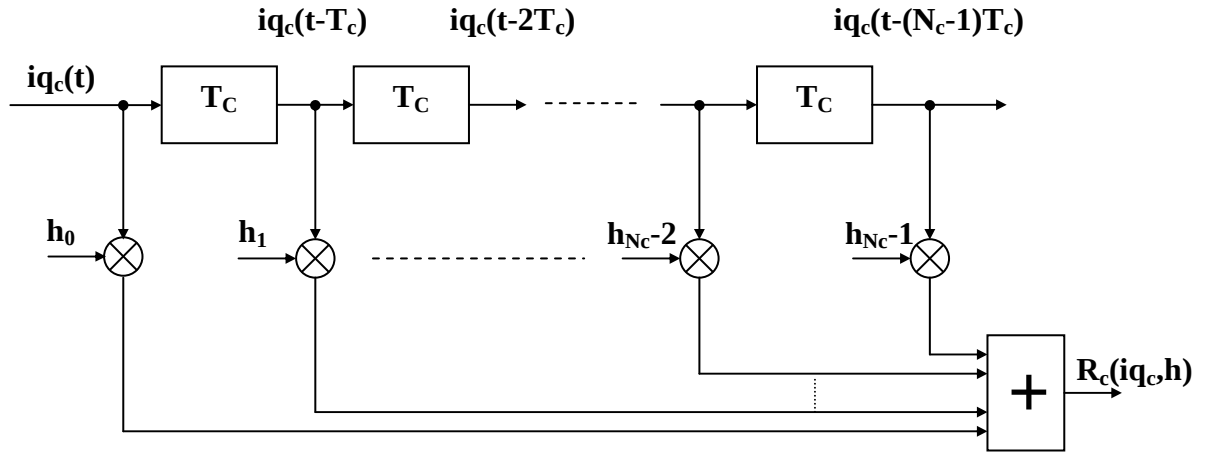


Figure 5.3 : dé-étalement avec un Matched Filter

Sur la figure précédente, $\mathbf{iq_c(t)}$ représente le signal reçu (dont le spectre est étalé). Les bits de données émis peuvent être retrouvés par le signe de la sortie du Matched Filter (voir plus loin).

Exemple :

$$\mathbf{pn_t} = [\mathbf{pn_0 \ pn_1 \ pn_2 \ pn_3 \ pn_4 \ pn_5 \ pn_6}] = [+1 \ +1 \ +1 \ -1 \ +1 \ -1 \ -1]$$

$$\mathbf{h} = [\mathbf{h_0 \ h_1 \ h_2 \ h_3 \ h_4 \ h_5 \ h_6}] = [-1 \ -1 \ +1 \ -1 \ +1 \ +1 \ +1]$$

A la sortie du filtre RIF on retrouve la convolution des signaux I-Q, $\mathbf{iq_c}$ (démodulés et filtrés, $\mathbf{I_{chip}}$ ou $\mathbf{Q_{chip}}$ dans le schéma - bloc du récepteur) avec la réponse impulsionnelle du filtre RIF $\mathbf{h} = [\mathbf{h_0 \ h_1 \ \dots \ h_{N_c-1}}]$. A cause de l’inversion dans le temps, la sortie du filtre est la corrélation du signal à spectre étalé en bande de base $\mathbf{rx_b}$ avec la séquence PN locale.

$$\mathbf{R_c(\tau)} = \sum_{i=0}^{N_c-1} \mathbf{iq_c(t-i.T_c).h_i} = \sum_{i=0}^{N_c-1} \mathbf{iq_c(t-i.T_c).p_{N_c-1-i}} \quad (5.1)$$

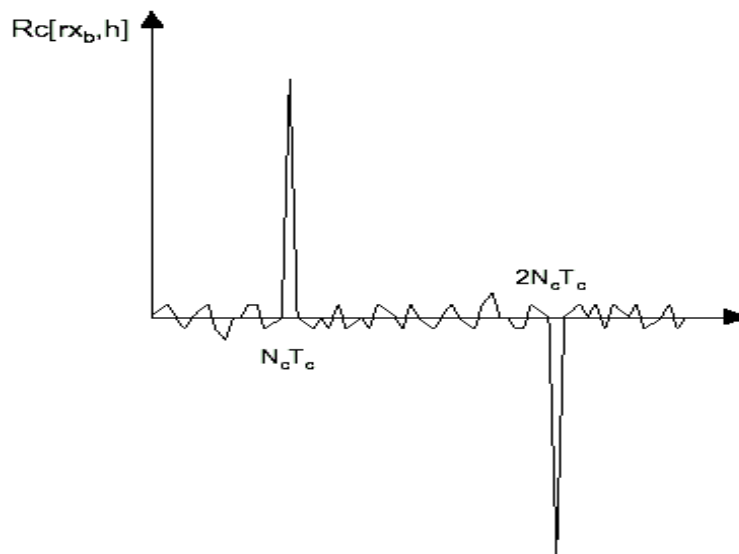


Figure 5.4 : allure du signal à la sortie du Matched Filter de la figure 5.3

Si le récepteur n'est pas synchronisé, le signal reçu se propagera tout simplement à travers le filtre. La sortie de ce dernier donnera la valeur de la fonction de corrélation entre les deux séquences. La présence d'un pic confirmera que le bon code a été reçu, et fournira un « timing » assez précis à propos de la synchronisation du signal reçu.

La sortie R_c du filtre RIF donne immédiatement les données décorrélées : la polarité des grands pics de corrélations indique la valeur des données.

V.3.2 Le corrélateur actif

Quand l'information sur le timing est déjà disponible, alors un simple récepteur à corrélateur actif peut être utilisé. Ce récepteur fonctionne correctement lorsque la séquence PN est précisément identifiée et correctement synchronisée. La synchronisation peut être obtenue en déplaçant le signal de référence par rapport au signal reçu, jusqu'à ce que la sortie de l'intégrateur dépasse un certain seuil. Cependant ceci est un processus très lent, surtout pour les codes d'étalement assez longs.

Le corrélateur actif est utilisé lorsque la séquence utilisée pour l'opération de dé-étalement est déjà synchronisée avec celle utilisée par le signal reçu.

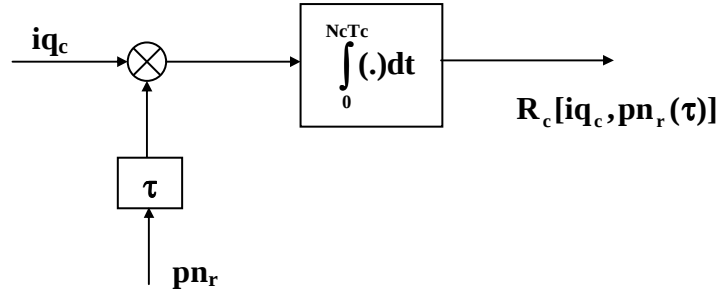


Figure 5.5 : le corrélateur actif

V.4 L'opération de synchronisation [5] [15] [16]

Pour pouvoir fonctionner correctement, un système de communication à spectre étalé nécessite une synchronisation entre les séquences PN générées localement $\mathbf{pn}_r$ (dans le récepteur) et celles utilisées par l'émetteur $\mathbf{pn}_t$. Cette synchronisation doit être satisfaite aussi bien en terme de débits qu'en terme de phase ou de position (ou encore de décalage).

Du fait que le pic d'auto- corrélation de la fonction des séquences PN est très étroit et décroît d'une manière linéaire, un non-alignement de $T_c / 2$ dans la séquence PN provoque une perte d'un facteur de 2 dans le gain de traitement G_p .

Le processus de synchronisation des séquences PN générées localement est habituellement exécuté en deux phases. La première appelée "acquisition " consiste à amener les 2 séquences d'étalement vers un alignement approximatif , l'une avec l'autre. Une fois la séquence PN "acquise", la seconde phase appelée "poursuite" (*tracking*) est amorcée. Elle consiste en un alignement fin, et au maintien de ce dernier quelle que ce soient les circonstances grâce à une boucle de retour. C'est essentiel pour obtenir la puissance de corrélation la plus élevée, et donc le plus grand gain de traitement G_p chez le récepteur.

V.4.1 Phase d'acquisition (Coarse Synchronisation)

Dans le cas général, le problème d'acquisition consiste en une recherche à travers une région à deux dimension : temps - fréquence (chip, porteuse), dans le but de synchroniser le signal à spectre étalé reçu avec la séquence PN générée localement. Mais étant donné que la fréquence varie peu (effet Doppler), la synchronisation dans le domaine temporelle

détermine la plupart du temps la performance du système. C'est ce que nous allons aborder dans la suite.

Une caractéristique commune de toutes les méthodes d'acquisition est que le signal reçu ainsi que la séquence PN générée localement sont appliqués à un corrélateur avec un temps de décalage grossier (généralement $T_c / 2$) pour produire une mesure de la similarité entre les deux. Cette mesure est ensuite comparée à un seuil pour pouvoir décider si les deux signaux sont en synchronisme.

Il y a trois approches pour traiter le problème d'acquisition :

- une approche entièrement parallèle
- une approche entièrement sérielle
- une approche combinant les deux approches précédentes

V.4.1.1 L'approche entièrement parallèle

C'est l'approche optimale du point de vue "temps d'acquisition". Cependant, ce dispositif est le plus complexe à réaliser. Un Matched Filter calcule la fonction de corrélation pour chaque code décalé d'un temps $k \frac{T_c}{2}$ ($k = 0, \dots, 2N_c$, où N_c est la longueur du code utilisé). Il suffit ensuite de choisir le code (avec décalage) qui a produit le maximum de produit d'inter-corrélation avec la séquence incidente. Cette méthode fournit le temps d'acquisition le plus court ($T_{acq} = N_c T_c$), mais une implémentation entièrement parallèle requiert un grand lot de composants et de matériel. Dans cette optique, on a besoin de $2N_c$ Matched Filter(s). La complexité du matériel nécessaire augmente donc avec la longueur du code PN (période de la séquence PN). Ainsi cette méthode est la plupart du temps employée pour les systèmes à codes courts. Un exemple d'architecture entièrement parallèle est donné sur la figure 5.6.

V.4.1.2 L'approche entièrement sérielle

Pour l'approche sérielle on utilise un simple corrélateur à intégrateur (voir figure 5.7), ce qui nécessite une intégration à travers toute une période $N_c \cdot T_c$ de la séquence PN pour calculer un point de la fonction de corrélation. Un long temps d'acquisition est nécessaire mais ici, on peut se contenter d'un matériel dérisoire.

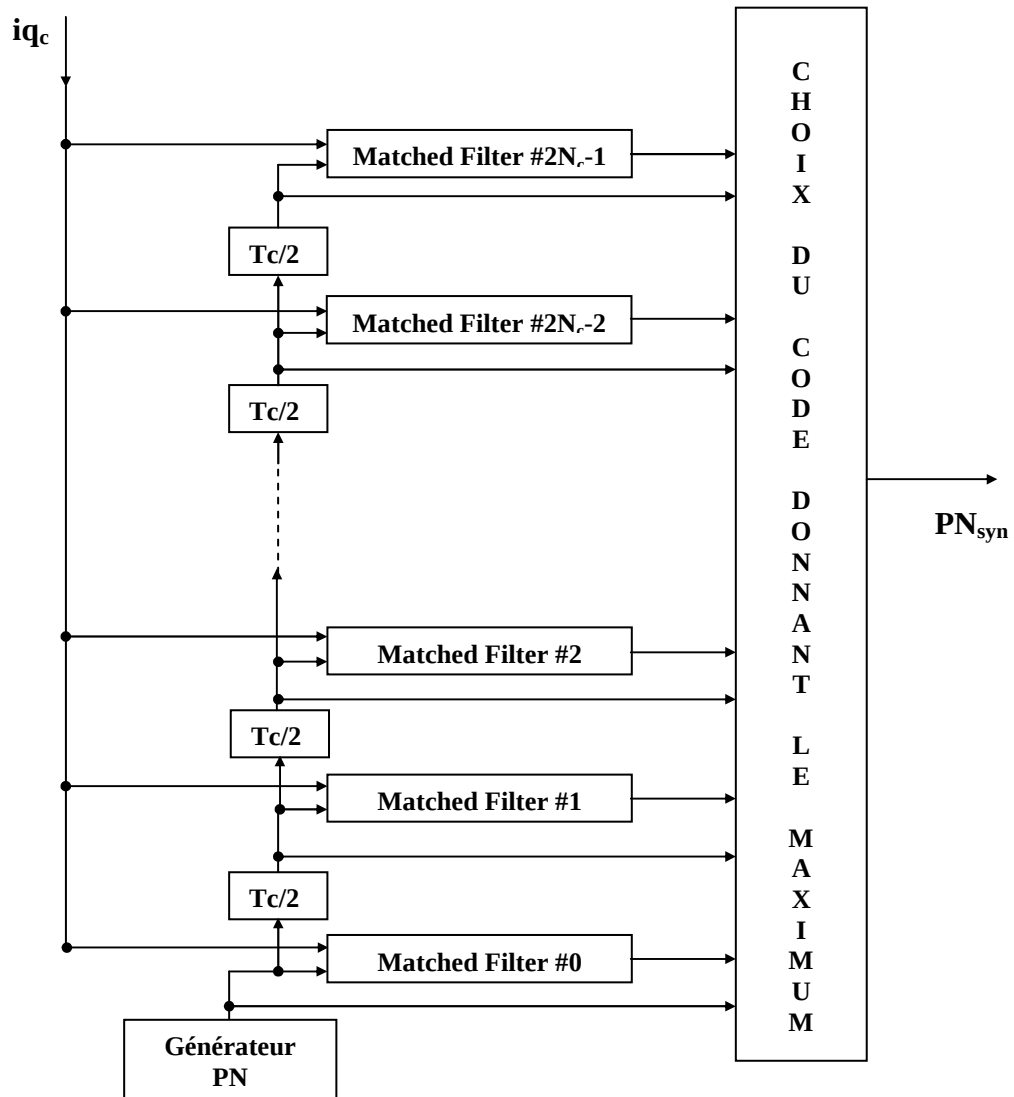


Figure 5.6 : acquisition optimale à l'aide d'un Matched Filter

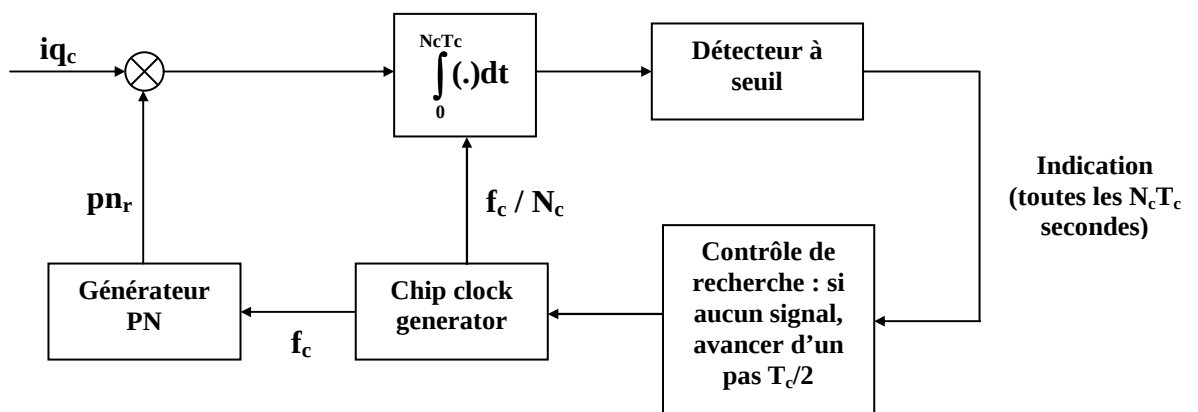


Figure 5.7 : approche entièrement sérielle de l'acquisition

La corrélation est effectuée pendant une période de la séquence PN. Après chaque intervalle d'intégration la sortie du corrélateur est comparée avec un seuil pour déterminer si la séquence PN appropriée (celle que le récepteur connaît) est présente à l'entrée. Si le seuil n'est pas dépassé, la séquence PN générée localement $\mathbf{pn}_r$ est avancée de $\frac{T_c}{2}$ secondes et le processus de corrélation se répète.

Ces opérations sont effectuées jusqu'à ce qu'un signal soit détecté ou que le temps d'acquisition T_{acq} ait été dépassé. Pour un pas de $\frac{T_c}{2}$, et dans le pire des cas, le temps d'acquisition peut s'écrire :

$$T_{acq} = \left\lceil \frac{N_c T_c}{\frac{T_c}{2}} \right\rceil N_c T_c = 2N_c^2 T_c \quad (5.2)$$

Ceci représente beaucoup de temps pour le cas des codes longs (N_c grand), ce qui devient vite inacceptable.

Exemple : Pour le système UMTS, on a $N_c = 2^{18} - 1$ (liaison descendante) pour un débit de chip R_c égal à $4.096 \cdot 10^6$ chips / s, ce représente une période de $\frac{N_c}{R_c} = 6.4 \cdot 10^{-2}$ s. Le temps d'acquisition est donc ici égal à : $T_{acq} = 3.35 \cdot 10^4$ s = 9.32 h, ce qui n'est pas acceptable.

Remarquons qu'ici on suppose que le récepteur connaît déjà le code utilisé à l'émission. Si ce n'est pas le cas, ce temps sera encore allongé (pour tenir compte du temps utilisé pour communiquer le code vers le récepteur). Nous verrons plus loin, comment surmonter ce problème dans la pratique.

V.4.1.3 L'approche hybride (combinant les deux approches précédentes)

L'idée derrière l'approche hybride consiste à trouver un compromis entre le temps d'acquisition et la complexité du dispositif à réaliser. Le temps d'acquisition est amélioré au détriment d'une complexité supérieure à celle obtenue pour l'approche sérielle. Un exemple de réalisation visant à améliorer le temps d'acquisition dans un facteur de 3 est donné par la

figure 5.8.

Plusieurs corrélateurs à intégrateurs sont placés en parallèle (3 dans cet exemple) avec des séquences PN espacées d'un demi - chip ($T_c / 2$). Après un temps d'intégration $N_c T_c$, les résultats fournis par les sorties du corrélateur sont comparés. La fonction de corrélation est donc calculé en 3 points successifs (espacés de $T_c / 2$). Quand aucune des sorties n'excède le seuil, les séquences seront avancées de $3T_c / 2$ secondes. Quand le seuil est enfin dépassé, le corrélateur dont la sortie donne la plus grande valeur est choisi.

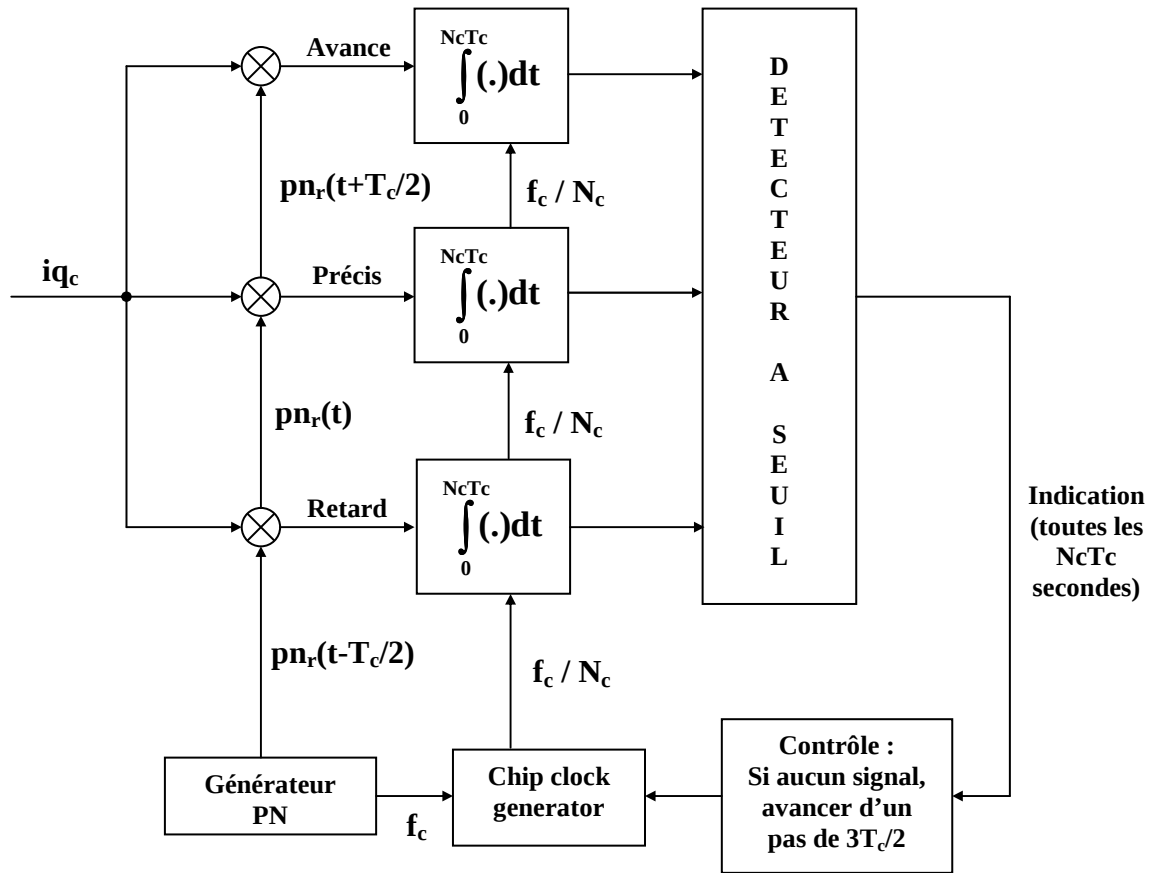


Figure 5.8 : acquisition série/parallèle

Comme précédemment, dans le plus pire des cas, on a

$$T_{acq} = \left(\frac{N_c T_c}{3T_c / 2} \right) \cdot N_c T_c = \frac{2}{3} N_c^2 T_c \quad (5.3)$$

Le temps de recherche a été diminué pour implémentation plus complexe et plus coûteuse.

V.4.1.4 La phase d'acquisition dans la pratique [6] [14] [20]

Dans la pratique, on n'utilise pas un code long dans toute sa longueur. Ceci veut dire qu'on utilise seulement une portion de ce code. En procédant de la sorte, on dispose d'un ensemble de segments de codes connu à la fois par l'émetteur et le récepteur.

Dans le cas de la liaison descendante utilisée par l'UMTS, un code de Gold de longueur $2^{18} - 1$ est divisé en des portions de 10 ms (ce qui représente 40960 chips pour un débit de 4.096 Mc/s). Et il y a 512 segments de codes (répartis en 16 groupes de 32 codes) différents disponibles pour cette liaison descendante.

L'astuce consiste à fournir un moyen permettant au récepteur de se synchroniser rapidement sur un canal non modulé (dit canal de synchronisation primaire). Or, à chaque groupe correspond un code sur un autre canal (dit canal de synchronisation secondaire); on a donc 16 codes disponibles sur ce canal. Ce canal est modulé par une séquence fixe de 16 bits pour tout le réseau (déjà connue par le récepteur). La synchronisation sur ce canal permet donc de connaître le groupe du code utilisé (après avoir essayé les 16 codes) ainsi que le timing du code de Gold utilisé (en analysant les phases des données reçues).

Ayant obtenu ces informations (timing et groupe de code), il ne reste plus qu'à essayer les 32 codes appartenant au groupe pour savoir le bon code avec le bon timing. On remarque donc que le premier canal ne sert qu'à obtenir le timing du second canal (pour ne plus avoir à essayer les versions décalées de chacun des 16 codes du second canal).

Les codes employés pour les canaux de synchronisation possèdent une longueur de 256. Sachant qu'on utilise un "matched filter" (approche parallèle), le temps d'acquisition nécessaire à la synchronisation sur le canal de synchronisation primaire est donc

$$\text{de } T_{\text{acq1}} = \frac{256}{4.096 \cdot 10^6} = 6.25 \cdot 10^{-5} \text{ s}.$$

Maintenant que le timing du second canal est acquis, le temps nécessaire à la connaissance du groupe de code et du timing à propos du code de Gold utilisé par la station

$$\text{de base (en utilisant un corrélateur actif) est de } T_{\text{acq2}} = 256 * \frac{256}{4.096 \cdot 10^6} = 16 \text{ ms (étant donné}$$

qu'on a à calculer 256 corrélations : combinaison des 16 codes et des 16 versions décalées de

la séquence de 16 bits).

Et il ne reste qu'à essayer chacun des 32 codes appartenant au groupe. Cette fois aussi, on utilise un corrélateur actif (approche entièrement sérielle), étant donné la longueur du code à traiter (40960 chips). Le temps nécessaire pour cette opération est de $T_{acq3} = 32 * 2 * \frac{40960}{4.096 \cdot 10^6} = 640 \text{ ms}$.

Le temps d'acquisition total est donc de $T_{acq} = T_{acq1} + T_{acq2} + T_{acq3} = 656 \text{ ms}$ ce qui est acceptable.

V.4.2 Phase de poursuite (Tracking Phase) [5] [15]

L'acquisition est destinée à synchroniser le récepteur sur le signal reçu à une fraction de période de chip près (généralement $\frac{T_c}{2}$). C'est là donc qu'intervient la phase de poursuite qui vise à réduire à zéro l'erreur de synchronisation. Cette correction doit se faire continuellement, étant donné que les horloges utilisées ne sont pas infiniment précises. Elle maintient donc le générateur de code PN chez le récepteur en synchronisme avec le signal reçu. Ceci est nécessaire pour obtenir un "gain de traitement" G_p (processing gain) maximal. Et une erreur de phase se traduit directement par une perte de performance.

Exemple : Une erreur de phase de $T_c / 2$ sur les séquences PN se traduirait par une division par 2 du gain de traitement G_p (voir figure 5.10).

V.4.2.1 La boucle de verrouillage de délai DLL (Delay-Locked Loop)

Le code PN généré localement $pn_r(t)$ est décalé en phase par rapport au code PN reçu $pn(t)$ d'un temps égal à τ , où $\tau < T_c / 2$. Dans le DLL, deux séquences PN $pn_r(t + T_c/2 + \tau)$ et $pn_r(t - T_c/2 + \tau)$ sont séparées l'une de l'autre de par un chip (T_c). Les sorties Avance et Retard (Early et Late) sont des estimations de la fonction d'auto - corrélation sur deux points temporels :

$$c_e = R_a(\tau - T_c / 2) \quad (5.4)$$

$$c_l = R_a(\tau - T_c / 2) \quad (5.5)$$

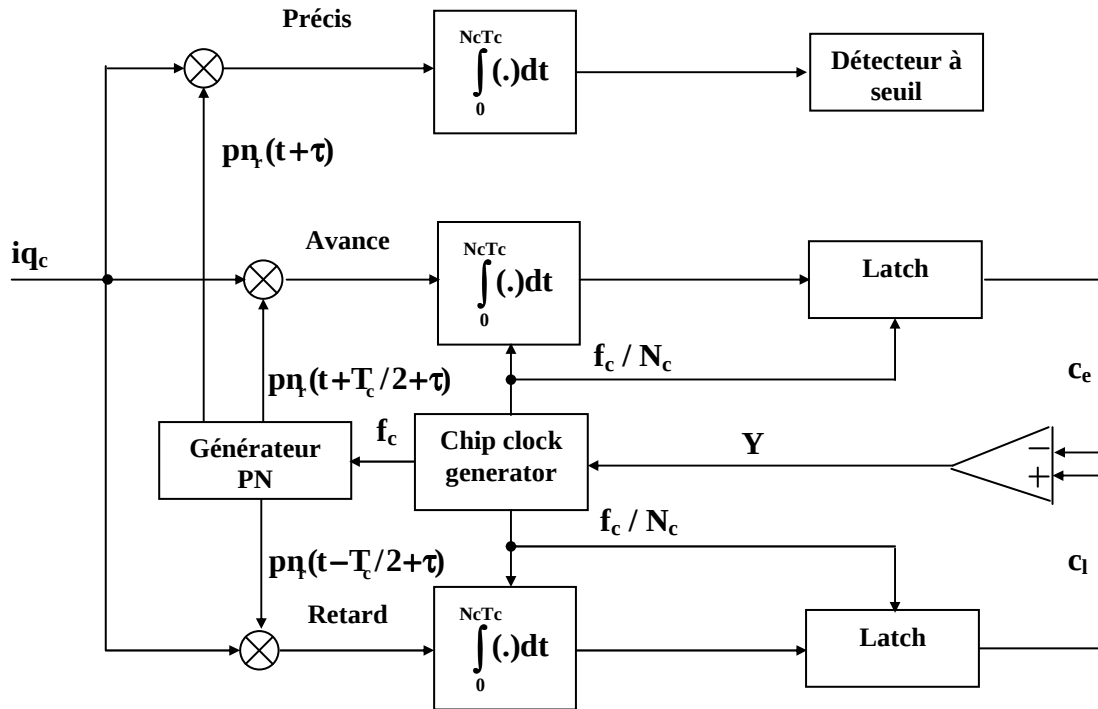


Figure 5.9 : boucle de verrouillage de délai

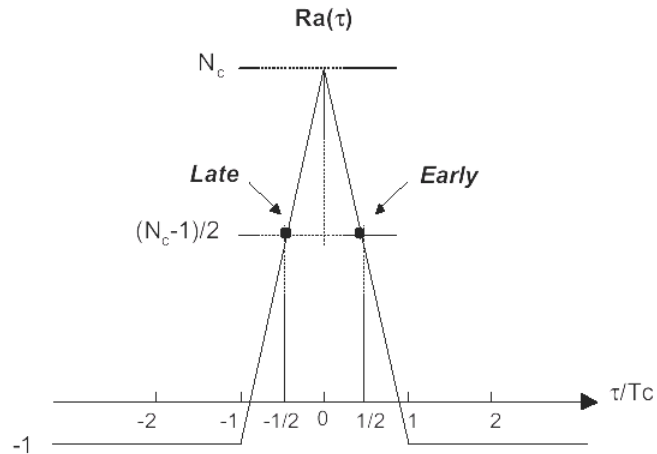


Figure 5.10 : estimation de l'auto-corrélation d'une séquence PN

Lorsque τ est positif, le signal de retour $\mathbf{Y}(\mathbf{t})$ suggère au circuit d'horloge (NCO : Numerical Controlled Oscillator) du générateur de séquence PN d'augmenter sa fréquence, ce qui force τ à décroître. De la même façon, lorsque τ décroît, $\mathbf{Y}(\mathbf{t})$ suggère au NCO

précèdent de réduire sa fréquence, ce qui force τ à croître.

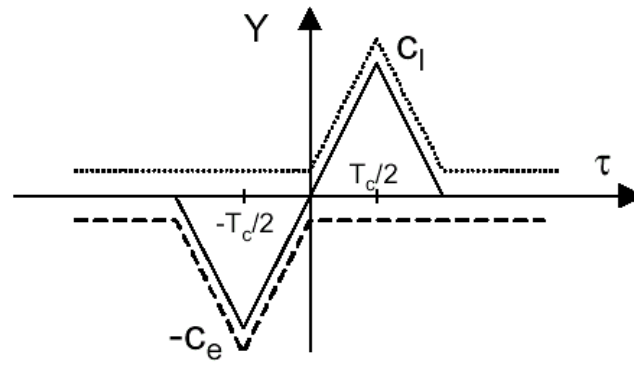


Figure 5.11 : allure de Y en fonction de l'erreur d'alignement τ

Chapitre VI

LES PERFORMANCES D'UN SYSTEME DS-CDMA

VI.1 Définition d'un système DS-CDMA

D'abord le CDMA (Code Division Multiple Access) ou l'AMRC est une méthode utilisée pour multiplexer plusieurs communications sur un seul support de transmission (en utilisant une même bande de fréquence) par le biais de plusieurs codes (ayant une certaine orthogonalité). Ceci n'est possible qu'en effectuant un étalement de spectre. On parle alors de système DS-CDMA lorsque celui-ci utilise la méthode de séquence directe comme méthode d'étalement de spectre.

VI.2 Caractéristiques globales d'un système DS-CDMA [5]

Un système DS-CDMA ne nécessite pas l'allocation de fréquence (ce qui est le cas pour l'AMRF.), ni la synchronisation entre les utilisateurs (ce qui est le cas pour l'AMRT). Chaque usager dans un système DS-CDMA possède un accès complet à toute la bande de fréquence, à n'importe quel moment. Cependant ceci se fait au détriment de la qualité de communication qui se dégrade avec le nombre croissant des usagers du système (augmentation du taux d'erreur binaire).

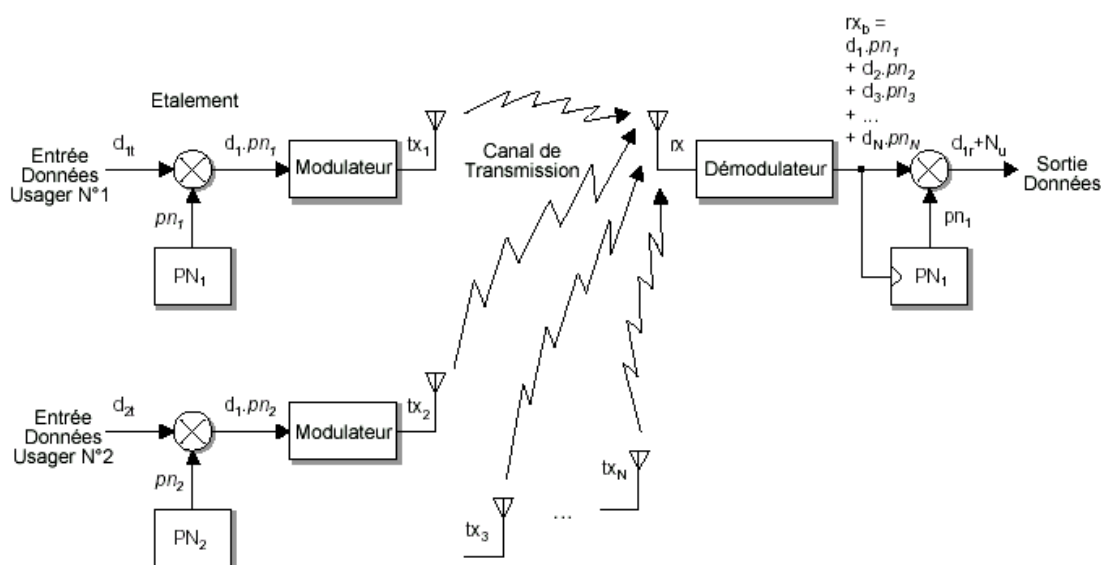


Figure 6.1 : principe d'un système DS-CDMA

Au sein d'un système DS-CDMA, chaque utilisateur :

- ♦ possède son propre code PN
- ♦ utilise la même bande de fréquence radio
- ♦ peut émettre indépendamment des autres utilisateurs (d'une manière asynchrone ou synchrone vis-à-vis des autres utilisateurs).

Les codes les plus utilisés (à cause de leurs performances) sont :

- ♦ les codes de Gold, Kasami (pour les systèmes DS-CDMA asynchrones)
- ♦ les codes Walsh-Hadamard (pour les systèmes DS-CDMA synchrones)

VI.3 Performances d'un système DS-CDMA [5] [10] [11]

VI.3.1 Les puissances mises en jeu

On sait que les performances d'un récepteur numérique dépendent du rapport Energie par bit sur la Densité du bruit $\frac{E_b}{N_0}$. Dans un système conventionnel avec une modulation à deux états, ce rapport prend une valeur proche du rapport signal sur bruit **SNR** étant donné que le rapport $\frac{E_b}{N_0}$ peut s'exprimer par

$$\frac{E_b}{N_0} = \frac{S/R}{N/W} \quad (6.1)$$

où **W** est la bande passante occupée par le signal, **R** est le débit et le rapport **S/N** représente le rapport signal sur bruit.

Pour atteindre un taux d'erreur donné, il faut avoir un rapport $\frac{E_b}{N_0}$ donné qui dépend des caractéristiques du récepteur.

Dans un système DS-CDMA, le rapport **W/R** représente le gain de traitement du système. A la sortie de l'émetteur, on a donc un rapport $\frac{E_b}{N_0}$ très faible équivalent à un faible rapport Signal sur Bruit si la transmission est vue du point de vue classique. Le rapport

Signal sur Bruit peut ainsi prendre des valeurs inférieures à 1 (rapport négatif en dB). Dans [11] on trouve par exemple un rapport **SNR** de -1 dB.

Après l'opération de dé-étalement, la bande passante occupée par les données passe de **W** à **R**, ce qui équivaut à une multiplication du rapport $\frac{E_b}{N_0}$ par le gain de traitement $\frac{W}{R}$ pour assurer une démodulation correcte des données.

En procédant de la sorte, la transmission apparaît donc comme du bruit pour un récepteur classique ou un récepteur qui ne connaît pas le code d'étalement utilisé à l'émission.

VI.3.2 Résistance aux effets des propagations par trajets multiples et des interférences

VI.3.2.1 Introduction

On considère une communication à temps discrets où le signal reçu est une séquence définie par

$$\mathbf{r}_m = \epsilon \mathbf{d}_m + \mathbf{v}_m \quad (6.2)$$

Où $\mathbf{d}_m$ est une séquence de symboles d'information (binaires bipolaires, $d_m \in \{-1, 1\}$) et $\mathbf{v}_m$ est un bruit blanc Gaussien de moyenne zéro, c'est à dire que

$$E[\mathbf{v}_n] = 0 \quad (6.3)$$

$$E[\mathbf{v}_m \mathbf{v}_{m+1}] = \sigma^2 \delta(l) \quad (6.4)$$

Ici $\mathbf{r}_m$ est connu sous la dénomination de « variable de décision », dont les statistiques déterminent les performances du récepteur.

Nous utilisons un récepteur à corrélation pour déterminer lequel de +1 ou -1 a été émis à l'instant m . Nous supposons que l'émetteur est connecté à un générateur dont la sortie fournit des 1 et des -1 avec des probabilités égales. Dans ce cas, un tel récepteur est un simple détecteur de niveau :

- si $\mathbf{r}_m \geq 0$ on décide que le bit émis est un 1
- si $\mathbf{r}_m \leq 0$ on décide que le bit émis est un 0

Il est facile de voir que c'est une variable aléatoire normale (Gaussienne) dont la moyenne est ϵd_m et dont la variance est σ^2

VI.3.2.2 Résistance aux interférences

Maintenant considérons que chaque symbole modulé par une autre séquence ayant des valeurs +1 ou -1 (séquence d'étalement) dont la longueur est N .

Etant donné que chaque bit de durée T est codé avec une séquence de N « chips » de durée $T_c = \frac{T}{N}$.

Le signal reçu peut donc s'écrire :

$$r_n = \epsilon_c d c_n + v_n, \quad n = 0, 1, \dots, N-1 \quad (6.5)$$

$$\text{où } \epsilon_c = \frac{\epsilon}{N} \text{ et } E[v_n^2] = \frac{\sigma^2}{N}$$

Maintenant en spécifiant 2 propriétés de la séquence d'étalement de spectre c_n :

- celle-ci a une valeur moyenne approximativement égale à zéro car

$$\sum_{n=0}^{N-1} c_n \approx 0, \quad (6.6)$$

- l'auto-corrélation possède 2 valeurs :

$$\sum c_n c_{n+i} \approx \begin{cases} N, & i = 0 \\ 0, & \text{autrement} \end{cases} \quad (6.7)$$

Ces conditions sont idéales mais peuvent être approchées en pratique.

Alors, le récepteur à corrélation effectue les opérations suivantes pour obtenir la variable de décision y :

$$y = \sum_{n=0}^{N-1} r_n c_n \quad (6.8)$$

ou encore

$$y = \sum_{n=0}^{N-1} (\epsilon_c d c_n + v_n) c_n \quad (6.9)$$

qui peut encore s'écrire en se basant sur les propriétés de la séquence d'étalement,

$$\mathbf{y} = \mathbf{N}\boldsymbol{\varepsilon}_c \mathbf{d} + \sum_{n=0}^{N-1} \mathbf{v}_n c_n \quad (6.10)$$

Il s'ensuit que notre variable de décision est normale de valeur moyenne $\mathbf{N}\boldsymbol{\varepsilon}_c \mathbf{d} = \boldsymbol{\varepsilon} \mathbf{b}$ et de variance $\boldsymbol{\sigma}^2$.

En comparant ce résultat, avec celui obtenu avec un système qui n'utilise pas l'étalement de spectre, on peut voir que l'étalement de spectre n'introduit pas d'amélioration sur un canal idéal à bruit blanc Gaussien.

Mais malgré cela, les avantages de l'étalement de spectre se matérialisent lorsqu'on cherche à éviter des signaux à bande étroite ou des signaux corrélés (interférences, propagation par trajets multiples, brouillages à bande étroite, ou d'autres signaux provenant des autres émetteurs du réseau).

Maintenant supposons qu'un interférant est présent dans le canal ; une constante inconnue est alors additionnée au signal reçu. Alors nous avons

$$\mathbf{r}_n = \boldsymbol{\varepsilon}_c \mathbf{d} c_n + \mathbf{i}_n + \mathbf{v}_n, \quad \mathbf{n} = 0, 1, \dots, N-1 \quad (6.11)$$

où $\mathbf{i}_n = \mathbf{I}$, est une constante réelle (car nous avons supposé que l'interférence est à bande étroite).

Alors notre récepteur à corrélation produit la variable de décision

$$\mathbf{y} = \sum_{n=0}^{N-1} (\boldsymbol{\varepsilon}_c \mathbf{d} c_n + \mathbf{i}_n + \mathbf{v}_n) c_n \quad (6.12)$$

qui devient

$$\mathbf{y} = \mathbf{N}\boldsymbol{\varepsilon}_c \mathbf{d} + \mathbf{I} \sum_{n=0}^{N-1} c_n + \sum_{n=0}^{N-1} \mathbf{v}_n c_n \quad (6.13)$$

ou encore

$$\mathbf{y} \approx \mathbf{N}\boldsymbol{\varepsilon}_c \mathbf{d} + \mathbf{0} + \sum_{n=0}^{N-1} \mathbf{v}_n c_n \quad (6.14)$$

qui est une variable de décision de valeur moyenne $\mathbf{N}\boldsymbol{\varepsilon}_c \mathbf{d} = \boldsymbol{\varepsilon} \mathbf{d}$ et de variance $\boldsymbol{\sigma}^2$.

Ainsi l'interférence a été supprimée par l'opération de dé-étalement.

La variable de décision aurait une valeur moyenne égale à $\epsilon d + I$ dans un système qui n'utilise pas l'étalement de spectre, un système qui sera donc inutilisable pour $|I|$ suffisamment grand.

Cependant, l'étalement de spectre n'apporte aucune protection contre les interférences à large bande. Ceci n'est pas grave, parce qu'il est très difficile de réaliser un générateur haute fréquence à large bande ayant une puissance suffisante pour brouiller la communication.

Remarque : Un système DS-CDMA ne peut annihiler complètement l'interférence à bande étroite que dans le cas d'une parfaite synchronisation entre l'émetteur et le récepteur (comme c'est le cas ici).

VI.3.2.3 Les propagations par trajets multiples

A partir de maintenant, nous allons normaliser, pour simplifier posons $\epsilon_c = 1$ alors l'énergie par bit est ainsi égale à N .

Considérons maintenant un canal à trajets multiples, (avec un chemin direct et un chemin via réflexion qui génère une copie du signal qui arrive après un délai égal à l fois la durée d'un *chip* avec une atténuation inconnue β) :

$$r_n = \begin{cases} d_m c_n + \beta d_{m-1} c_{N-l+n} & n = 0, 1, \dots, l-1 \\ d_m c_n + \beta d_{m-1} c_{n-l} & n = l, l+1, \dots, N-1 \end{cases} \quad (6.15)$$

où nous avons supposé que $l < N$, c'est à dire que le retard est inférieur à la durée d'un symbole.

On peut donc écrire :

$$y_m = Nd_m + \beta d_{m-1} \sum_{n=0}^{l-1} c_{N-l+n} c_n + \beta d_m \sum_{n=l}^{N-1} c_{n-l} c_n + \sum_{n=0}^{N-1} v_n c_n, \quad (6.16)$$

qui devient

$$y_m \approx Nd_m + 0 + 0 + \sum_{n=0}^{N-1} v_n c_n \quad (6.17)$$

Les effets des propagations par trajets multiples ont été donc annulés par l'opération de dé-étalement. Dans le cas d'un système n'utilisant pas un étalement de spectre, cette situation aboutira à une interférence inter-symbole sévère provoquant une dégradation de la performance.

Remarque : Un système DS-CDMA ne peut pas éliminer les effets des propagations par trajets multiples que dans le cas d'une synchronisation parfaite entre l'émetteur et le récepteur (comme c'est le cas ici).

VI.3.3 Capacité d'un système DS-CDMA [11] [18] [19]

VI.3.3.1 Les avantages inhérents à un système DS-CDMA

VI.3.3.1.1 La prise en compte de l'activité vocale

L'activité de la voix humaine est de l'ordre de 35%. Quand les utilisateurs au sein d'une cellule ne parlent pas, la prise en compte de l'activité vocale permet de réduire l'interférence mutuelle. Ainsi l'interférence est réduite dans un facteur de 65%. Le CDMA est la seule technologie d'accès multiple qui peut profiter naturellement de ce phénomène. On peut montrer que la capacité d'un système CDMA est multipliée par 3 environ à cause de la prise en compte de l'activité vocale.

VI.3.3.1.2 Limite floue de la capacité (soft limit capacity)

La capacité des systèmes DS-CDMA est limitée par les interférences, tandis que l'AMRT et l'AMRF sont limitées par la bande passante disponible (hard limit capacity). Ainsi la limite est floue car l'addition d'un utilisateur résulte en une légère dégradation des performances. Quand le système est très chargé, il n'y a pas de blocage (théoriquement) contrairement aux autres types d'accès multiple. Une autre conclusion qu'on peut tirer de ce fait est que toute réduction sur l'interférence due à l'accès multiple (MAI : Multiple Access Interference) est récompensée directement par une augmentation de la capacité.

VI.3.3.1.3 Effet de sectorisation

Dans les systèmes AMRT et AMRF, la sectorisation est effectuée pour atténuer l'interférence co-canal. Cependant, l'efficacité globale de ces systèmes décroît (car on ne peut pas utiliser les mêmes fréquences pour les secteurs d'une même cellule). D'un autre

coté, la sectorisation augmente la capacité des systèmes DS-CDMA. Cette sectorisation peut être faite tout simplement par l'introduction de 3 équipements radio similaires dans les 3 secteurs. La réduction de l'interférence mutuelle ainsi effectuée se traduit par une augmentation de la capacité (dans un facteur de 3 environ).

En général, toute isolation spatiale à travers l'utilisation des antennes multi-secteurs produit une augmentation de la capacité d'un système DS-CDMA.

VI.3.3.1.4 La réutilisation des fréquences.

Le plus grand avantage du DS-CDMA par rapport aux systèmes conventionnels est que celui ci peut réutiliser le spectre tout entier à travers toutes les cellules étant donné qu'il n'y a pas de concept d'allocation de fréquence dans ce système. Ceci augmente la capacité du système DS-CDMA dans une large proportion (à cause de la réduction du facteur de réutilisation des fréquences).

VI.3.3.2 Capacité d'un système DS-CDMA à cellule unique

Considérons un système DS-CDMA à cellule unique avec N usagers. On suppose qu'un contrôle de puissance approprié est utilisé de telle façon que les signaux sur la liaison montante sont reçus avec la même puissance. Chaque station de base traite le signal utile avec une puissance S et $N-1$ signaux interférants, chacun d'eux ayant une puissance S . Le rapport Energie par bit / Densité du bruit (interférence + bruit) peut être écrite comme suit :

$$\frac{E_b}{N_0} = \frac{S/R}{(N-1)S/W} = \frac{W/R}{N-1} \quad (6.18)$$

où R est le débit binaire et W est la bande passante totale occupée par le signal étalé. Le terme W/R n'est autre que le gain de traitement G_p . Si on tient compte du bruit de fond η dû au bruit thermique, l'équation précédente devient :

$$\frac{E_b}{N_0} = \frac{W/R}{(N-1) + (\eta/S)} \quad (6.19)$$

d'où la capacité en termes de nombre d'usagers est donné par :

$$N = 1 + \frac{W/R}{E_b/N_0} - \frac{\eta}{S} \quad (6.20)$$

où E_b/N_0 est la valeur requise pour un fonctionnement correct du démodulateur. Pour une transmission numérique de la parole ceci requiert un taux d'erreur binaire (TEB) supérieur ou égal à 10^{-3} . En utilisant l'équation précédente, on peut faire une simple comparaison du DS-CDMA avec les autres systèmes d'accès multiple.

Considérons une bande passante de 1.25 MHz et un débit de 8 kbps en utilisant des vocodeurs. Supposons qu'un E_b/N_0 minimal de 5 (7 dB) est requis pour un fonctionnement correct (TEB de 10^{-3}). En ne tenant pas compte du bruit thermique, le nombre d'utilisateurs dans ce système DS-CDMA (dans une bande de 1.25 MHz) est :

$$N = 1 + \frac{1.25 \text{ MHz} / 8 \text{ kHz}}{5} = 32 \text{ utilisateurs}$$

D'un autre côté, un système AMPS (Advanced Mobile Phone System) à cellule unique (AMRF à espacement de canaux de 30kHz) travaillant dans une même largeur de bande peut accommoder jusqu'à

$$\frac{1.25 \text{ MHz}}{30 \text{ kHz}} = 42 \text{ utilisateurs}$$

Pour un système D-AMPS (Digital-AMPS, TDMA avec multiplexage de 3 canaux) ce nombre est de 126 utilisateurs.

Jusqu'à présent, la capacité du système DS-CDMA est beaucoup plus petite que celles des systèmes conventionnels (étant donné que le nombre d'utilisateurs est très inférieur au gain de traitement (W/R) du système). Cependant, il est important de noter le fait que nous n'avons pas encore pris en compte les domaines dans lesquels le DS-CDMA excelle (détection de l'activité vocale, sectorisation, réutilisation des fréquences,...).

Notons aussi que dans un système AMPS à plusieurs cellules (avec un facteur de réutilisation de fréquences égal à 7), le nombre d'utilisateurs par cellules passe de 42 à 6 (et une réduction de 126 à 18 dans le cas du D-AMPS). Ainsi le DS-CDMA montrera toute sa puissance face aux systèmes conventionnels.

Dans les paragraphes qui vont suivre, on va discuter les effets de la détection de l'activité vocale et de la sectorisation sur la capacité, qui sont des méthodes utilisées pour diminuer les interférences mutuelles au sein du système DS-CDMA.

VI.3.3.2.1 La sectorisation

Toute isolation spatiale des usagers dans un système DS-CDMA est traduite directement en une augmentation de la capacité du système. Considérons un exemple où 3 antennes directionnelles ayant chacune un angle d'ouverture de 120°. Maintenant, les sources d'interférences vues par chaque antenne sont approximativement le tiers de celles vue par une antenne omnidirectionnelle. Ceci réduit le terme d'interférence (N-1) dans le dénominateur de l'équation (6.18) par un facteur de 3. Le nombre d'usagers par cellule est ainsi (pour $N \gg 1$) de $N = 3N_s$ où N_s est le nombre d'usagers par secteur.

VI.3.3.2.2 La prise en compte de l'activité vocale

La gestion de l'activité vocale est une fonctionnalité présente dans la plupart des vocodeurs numériques où la transmission est supprimée lors de l'absence de voix. Prenons pour le *facteur d'activité vocale* une valeur de 3/8 (correspondant à l'activité de la voix humaine de 35 à 40 %). Le terme d'interférence dans le dénominateur de l'équation (3) est donc réduit de $(N - 1)$ à $(N - 1)\alpha$.

Ainsi, avec la détection de l'activité vocale et la sectorisation, on a :

$$\frac{E_b}{N_0} = \frac{W/R}{(N_s - 1)\alpha + (\eta/S)} \quad (6.21)$$

Le nombre d'utilisateurs par cellule est donc de :

$$N = 3N_s = 3 \left[1 + \frac{1}{\alpha} \left\{ \frac{W/R}{E_b/N_0} - \frac{\eta}{S} \right\} \right] \quad (6.22)$$

Dans les mêmes conditions que précédemment, la capacité du système DS-CDMA est maintenant de

$$N = 3 \times \left[1 + \frac{8}{3} \left\{ \frac{1.25 \text{ MHz} / 8 \text{ kbps}}{5} \right\} \right] = 253 \text{ utilisateurs}$$

On voit que la capacité d'un système DS-CDMA est maintenant comparable à celles des systèmes conventionnels. Cependant, on n'a pas encore tenu compte des considérations de réutilisation de fréquences. Le plus grand avantage du DS-CDMA vient du fait qu'on peut utiliser les mêmes fréquences dans toutes les cellules (ce qui n'est pas le cas de

l'AMRF et de l'AMRT). Pour tenir compte de ceci, la capacité d'un système DS-CDMA doit être calculée dans le cas d'une configuration pluricellulaire, où il y a des interférences causées par les autres utilisateurs appartenant aux cellules adjacentes.

VI.3.3.3 Capacité de la liaison montante d'un système DS-CDMA dans une configuration pluricellulaire.

A cause du contrôle automatique de la puissance d'émission des mobiles, le niveau des interférences reçues à par les usagers des autres cellules dépend de deux facteurs :

- ◆ L'atténuation dans le trajet vers la station de base de l'utilisateur concerné
- ◆ L'atténuation dans le trajet vers la station de base de l'utilisateur interférant.

En supposant que l'effet de masque combiné à l'imperfection de la boucle de contrôle de puissance puissent être modélisés par une loi *log-normale* [11], la perte de propagation entre l'utilisateur sa station de base est proportionnelle à $10^{\left(\frac{\xi}{10}\right)} r^{-4}$, où ξ est une variable gaussienne log-normale de moyenne zéro et d'écart-type $\sigma = 8 \text{ dB}$ et r est la distance entre la station de base et l'utilisateur. Etant donné que nous ne considérerons que les niveaux moyens de puissance, l'effet des fadings rapides sera ignoré.

Soit un utilisateur interférant dans une cellule à une distance r_m de sa station de base et à r_0 de la station de base de l'utilisateur qui nous intéresse. L'interférant produira dans la cellule de l'utilisateur voulu une interférence égale à

$$\frac{I(r_0, r_m)}{S} = \left[\frac{10^{\left(\frac{\xi_0}{10}\right)}}{r_0^4} \right] \left[\frac{r_m^4}{10^{\left(\frac{\xi_m}{10}\right)}} \right] \leq 1 \quad (6.23)$$

Cette inéquation traduit le fait que l'utilisateur se commutera toujours vers la station de base qui le recevra avec le moins de perte possible. En supposant une densité uniforme des utilisateurs, en normalisant le rayon des cellules à 1 et étant donné que $N = 3N_s$, la densité des utilisateurs est égale à

$$\rho = \frac{2N}{3\sqrt{3}} = \frac{2N_s}{\sqrt{3}} \quad \text{par unité de surface.} \quad (6.24)$$

Donc, le rapport interférence total sur le signal utile est [11] :

$$\frac{I}{S} = \iint_A \psi \left(\frac{r_m}{r_0} \right)^4 \left\{ 10^{(\xi_0 - \xi_m)/10} \right\} \Phi(\xi_0 - \xi_m, r_0 / r_m) p \, dA \quad (6.25)$$

où m est l'index de la station de base pour laquelle on a

$$r_m^4 10^{-\xi_m} = \min_{k \neq 0} r_k^4 10^{-\xi_k} \quad (6.26)$$

et

$$\Phi(\xi_0 - \xi_m, r_0 / r_m) = \begin{cases} 1, & \text{si } (r_m / r_0)^4 10^{(\xi_0 - \xi_m)/10} \leq 1 \\ 0, & \text{autrement} \end{cases} \quad (6.27)$$

et ψ est la variable qui représente l'activité vocale, qui est égale à 1 avec une probabilité α et 0 avec une probabilité $(1 - \alpha)$, et enfin A désigne l'aire du réseau vu sous l'angle central du secteur considéré.

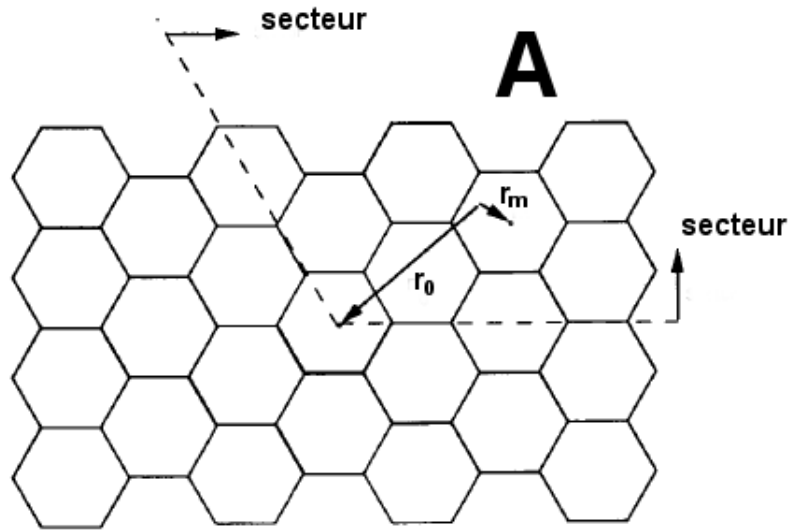


Figure 6.2 : explication géométrique de l'équation (6.24)

On peut montrer que [11], ce rapport (I/S) est une variable aléatoire gaussienne, dont la moyenne et la variance sont majorées par (pour une valeur de $\sigma = 8 \text{ dB}$)

$$E(I/S) \leq 0.247N_s \text{ et } \text{Var}(I/S) \leq 0.078N_s \quad (6.28)$$

Et en incluant ces résultats dans le calcul de la capacité, le rapport $\frac{E_b}{N_0}$ sur la liaison montante pour chacun des usagers voulus devient une variable aléatoire :

$$\frac{E_b}{N_0} = \frac{W/R}{\sum_{i=1}^{N_s-1} X_i + (I/S) + (\eta/S)} \quad (6.29)$$

Ici, le terme $(N_s - 1)\alpha$ dans l'équation (6.20) est remplacé pour chaque usager par les variables aléatoires $\{X_i\}$, où X_i prend la valeur 1 avec une probabilité α et 0 avec une probabilité $(1 - \alpha)$.

$\frac{E_b}{N_0}$ représente un seuil en dessous duquel, la réception correcte des données n'est plus possible (pour un taux d'erreur binaire donné).

En utilisant cette équation, on peut obtenir la probabilité d'obtenir un taux d'erreur binaire supérieur à 10^{-3} :

$$\Pr(\text{TEB} > 10^{-3}) = \sum_{k=0}^{N_s-1} C_k^{N_s-1} \alpha^k (1 - \alpha)^{N_s-1-k} Q\left(\frac{\delta - k - 0.247N_s}{\sqrt{0.078N_s}}\right) \quad (6.30)$$

$$\text{où} \quad \delta = \frac{W/R}{E_b/N_0} - \frac{\eta}{S}$$

Ainsi, pour calculer la capacité du système, on doit se fixer le pourcentage du temps pour lequel ce taux d'erreur sera dépassé (1 % dans notre exemple).

VI.3.3.4 Capacité de la liaison descendante d'un système DS-CDMA dans une configuration pluricellulaire

Dans la liaison descendante, le contrôle de puissance prend la forme d'une allocation de puissance au niveau de l'émetteur de la station de base suivant le besoin de chaque usager individuel dans une cellule donnée. Ceci requiert que le mobile mesure le rapport de la puissance reçue à partir de sa station de base par la puissance totale reçue. En supposant que la station de base possède une estimation précise de S_{T_i} (la puissance reçue par le mobile à

partir de sa station de base) et $\sum_{i=1}^K S_{T_i}$ (la puissance totale reçue par le mobile), on peut minorer le rapport

$$\frac{E_b}{N_0} \geq \frac{W}{R} \frac{\beta \Phi_i S_{T_i}}{\left[\left(\sum_{j=1}^K S_{T_j} \right) + \eta \right]} \quad (6.31)$$

où β est la fraction de la puissance totale de la station de base consacrée aux usagers (le reste est consacré aux signaux pilotes),

et Φ_i est la fraction de la puissance consacrée à l'utilisateur i , donnée par

$$\Phi_i \leq \frac{E_b / N_0}{\beta W / R} \left[1 + \left(\frac{\sum_{j=2}^K S_{T_j}}{S_{T_1}} \right)_i + \frac{\eta}{(S_{T_1})_i} \right] \quad (6.32)$$

et en posant

$$f_i = \left(1 + \frac{1}{S_{T_1}} \sum_{j=2}^K S_{T_j} \right) \quad i = 1, 2, \dots, N_s \quad (6.33)$$

une limite pour le taux d'erreur binaire peut être obtenue comme [11]:

$$\Pr(\text{TEB} > 10^{-3}) = \Pr\left(\sum_{i=1}^{N_s} f_i > \delta'\right) \quad (6.34)$$

$$\text{où} \quad \delta' = \frac{\beta W}{R} \frac{E_b}{N_0}$$

L'expression précédente qui contrairement à celle correspondante à la liaison montante ne peut pas être évaluée analytiquement. Ainsi une simulation s'impose pour déterminer la capacité de la liaison descendante.

Dans [11], on trouve un extrait de simulation dont une partie des résultats est reportée dans le paragraphe suivant pour illustrer les performances d'un système DS-CDMA en termes de capacité dans une configuration pluricellulaire.

VI.3.3.5 Exemples et comparaisons

VI.3.3.5.1 Paramètres

- ♦ La bande utilisée pour l'étalement de spectre est de 1.25 MHz
- ♦ Le débit binaire est de 8 kbps pour un vocodeur de moyenne qualité
- ♦ Le facteur d'activité vocale est de $\alpha = \frac{3}{8}$
- ♦ le nombre de secteur par cellule est de 3
- ♦ dans la liaison descendante, on a $\beta = 0.8$
- ♦ un TEB meilleur que 10^{-3} est atteint pendant 99% du temps

Ces paramètres impliquent les choix [11] de $\delta = 30$ et $\delta' = 38$

VI.3.3.5.2 Calculs

Avec ces paramètres, la liaison montante peut supporter 36 usagers par secteur ou encore 108 usagers par cellule. Ce nombre passe à 132 usagers par cellule si les cellules adjacentes ne fonctionnent qu'à la moitié de cette charge.

Dans les mêmes conditions, la liaison descendante peut traiter 38 usagers par secteur ou encore 114 usagers par cellule.

Remarque : De là, on peut énoncer une règle générale qui veut que la capacité d'un système DS-CDMA soit limitée par celle de la liaison montante. De ce fait, dans la suite on ne considérera que la capacité de la liaison montante. Ceci est toujours valable pour tout système DS-CDMA.

VI.3.3.5.3 Comparaisons

Supposons un système analogique AMPS (canal de 30 kHz, 3 secteurs par cellule et un facteur de réutilisation de fréquence égale à 7). On peut calculer la capacité d'un tel système est de $(1.25 \text{ MHz} / 30 \text{ kHz}) / 7 = 6$ utilisateurs par cellule. Ainsi le DS-CDMA fournit une amélioration de la capacité dans un facteur de 18 comparé à l'AMRF utilisé par le système AMPS.

De l'autre côté, considérons un système numérique utilisant l'AMRT comme l'IS-54. Chaque canal de 30 kHz multiplexe 3 communications en utilisant l'AMRT. La capacité de ce système est donc de 3 fois celle de l'AMPS analogique. Mais cette capacité est encore de 6 fois moins importante que celle fournie par un système DS-CDMA sous les mêmes conditions.

Dans le cas du GSM, chaque canal possède une largeur de 200 kHz et utilise l'AMRT pour multiplexer 8 communications. En utilisant un facteur de réutilisation de fréquence égal à 7, une bande de 1.25 MHz peut donc contenir 6 porteuses correspondants à $(8 \times 6) / 7 = 7$ usagers, ce qui est vraiment loin de ce qu'un système DS-CDMA aurait accepté dans les mêmes conditions.

VI.4 Problèmes inhérents à un système DS-CDMA [5] [6]

VI.4.1 L'interférence causée par l'accès multiple (MAI: Multiple Access Interference)

Le détecteur reçoit un signal composé essentiellement de la somme de tous les signaux envoyés par tous les usagers, qui se superposent dans le temps et dans le domaine fréquentiel. Le MAI fait allusion à l'interférence entre les usagers d'un système DS-CDMA utilisant l'étalement par Séquence Directe. C'est un facteur limitant la capacité et la performance d'un système DS-CDMA.

Dans un système DS-CDMA conventionnel, un signal provenant d'un utilisateur particulier est détecté en utilisant la corrélation entre le signal reçu avec le code correspondant régénéré. Ce type de détecteur ne prend donc pas en compte l'existence d'une quelconque interférence MAI. D'où une meilleure stratégie consiste à utiliser une détection d'un usager parmi plusieurs d'autres. L'information à propos des utilisateurs multiples est utilisée pour pouvoir détecter chaque utilisateur individuel, et choisir le plus vraisemblable. C'est la détection multi-utilisateur (voir chapitre VII).

VI.4.2 L'effet proche-lointain (near-far effect)

C'est un problème spécifique à un système DS-SS, qui se présente lorsqu'une station (généralement la station de base) reçoit en même temps des signaux provenant de plusieurs utilisateurs. Si pendant la communication entre la station et un utilisateur i (voir figure 6.3),

un autre utilisateur j plus proche de la station émet aussi, l'utilisateur i risque d'être masqué (même si les codes utilisés sont orthogonaux).

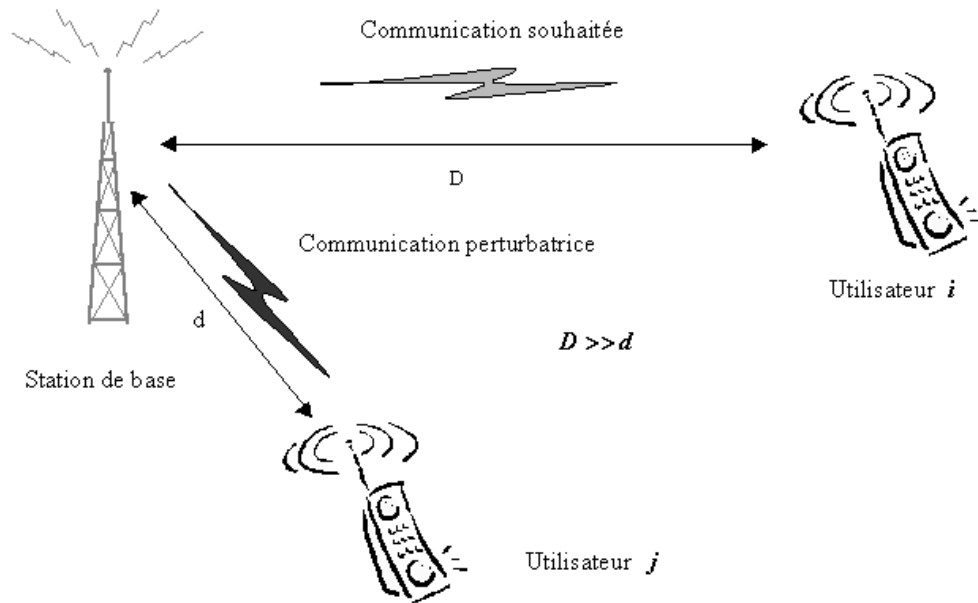


Figure 6.3 : l'effet Near-Far

En effet, comme n'importe quel récepteur, un récepteur DS-CDMA ne fonctionne pas correctement lorsque le niveau du signal qu'il reçoit se trouve au-dessous d'un certain seuil. Or ici, en plus du bruit blanc habituel, on retrouve aussi les signaux émis par les autres utilisateurs indésirables.

Ainsi le rapport signal sur bruit (lorsqu'on considère la communication entre la station et l'utilisateur i peut s'écrire :

$$\text{SNR}_i = \frac{E_b}{N_0 + I_{0i}} \quad (6.35)$$

où E_b est l'énergie par bit avec $E_b = \frac{P_i}{R}$

et où P_i est la puissance du signal provenant de l'utilisateur i , R le débit de données (d'informations), N_0 le bruit thermique, et I_{0i} l'interférence résultant (vu par la station lorsqu'elle s'adresse à l'utilisateur i) provoqué par les autres utilisateurs

$$I_{0i} = \frac{\sum_{j \neq i} P_j}{B} \quad (6.36)$$

où B représente la bande passante utilisée
d'où

$$SNR_i = \frac{P_i}{R \left(N_0 + \frac{\sum_{j \neq i} P_j}{B} \right)} \quad (6.37)$$

Avec l'expression (6.37), on peut voir aisément que le rapport signal sur bruit se dégrade avec l'augmentation du nombre des utilisateurs. Mais pour illustrer le problème de l'effet « proche – lointain », prenons un exemple où il n'y a que deux utilisateurs dans une cellule.

Supposons aussi que le bruit thermique est négligeable, et que les pertes dues à la propagation entre l'utilisateur le plus proche et celui le plus loin diffèrent de 30 dB. Si on utilise une bande passante de 10Mhz pour la transmission d'une information ayant un débit de 10000 bps.

Ainsi :

$$SNR_i = 10 \log \left(\frac{B}{R} \right) + (P_i - P_j)_{dB} = 0 \text{ dB} \quad (6.38)$$

Ici, on voit que si $P_i = P_j$, on aurait un rapport signal sur bruit de 30 dB. On peut donc remédier à ce problème en contrôlant les puissances d'émissions de chaque utilisateur de telle manière que la puissance reçue par la station de base à partir de tous les utilisateurs soit approximativement la même. Cette puissance sera donc celle qui est le minimum nécessaire pour garder une qualité de liaison donnée pour chaque mobile.

Chapitre VII

Utilisation du DS-CDMA dans un réseau téléphonique mobile

VII.1 Introduction

Bien que la théorie de l'étalement de spectre et de l'accès multiple à répartition dans les codes aient été connus depuis les années 50, le concept d'un système cellulaire basé sur le CDMA n'a été suggéré officiellement que depuis 1978 par Cooper et Nettleton. Durant les années 80, Qualcomm avait commencé des études approfondies sur la faisabilité d'un système de téléphonie mobile basé sur l'étalement par séquence directe (DS-CDMA). Ceci a abouti à l'annonce du premier système cellulaire grand public basé sur le CDMA qu'est le standard IS-95 en Juillet 1993. Cependant les opérations commerciales sur des systèmes IS-95 ne commençaient qu'en 1996.

Après cela, on assistait à toute une multitude d'évolution et d'extension qui visent à améliorer les systèmes existant et à préparer la transition vers une nouvelle génération de mobiles qui pourraient s'accommoder plus facilement des transferts de données à très haut débit ainsi que des services multimédia réservés auparavant à des stations de travail fixes.

Actuellement, les systèmes dits " de la troisième génération " sont tous basés sur le DS-CDMA. Ceci traduit l'intérêt croissant que suscite le DS-CDMA auprès des concepteurs et des opérateurs de système de téléphonie mobile à cause de sa flexibilité et de son efficacité en matière d'utilisation de la ressource spectrale. Cependant, l'utilisation du DS-CDMA au sein d'un système de téléphonie mobile bouleverse quelques règles en matière de conception du réseau, et introduit de nouvelles notions pour parvenir à un système plus performant que ceux déjà existant.

Les paramètres les plus courants sur lesquels on peut jouer pour optimiser le fonctionnement du réseau sont :

- ◆ La puissance d'émission du mobile
- ◆ La puissance d'émission de la station de base
- ◆ La sectorisation
- ◆ La mise en place de plusieurs couches de cellules

Les nouvelles notions introduites par le DS-CDMA sont :

- ◆ Les récepteurs RAKE
- ◆ Le Soft Handover
- ◆ Le Softer Handover
- ◆ L'Interfrequency Handover
- ◆ Les récepteurs à détecteurs multi-utilisateurs

VII.2 Les paramètres d'optimisation

Il faut se souvenir que toute réduction de l'interférence mutuelle dans un système DS-CDMA se traduit directement en une augmentation de la capacité du système. Les paramètres d'optimisation seront manipulés donc en vue de l'amélioration du rapport Signal sur Interférence.

VII.2.1 La puissance d'émission du mobile [5] [6] [17]

Comme nous l'avons vu dans le Chapitre VI, un contrôle sur la puissance d'émission s'impose pour contrer l'effet Proche-Lointain (Near-Far effects). De ce point de vue, le contrôle de la puissance n'est pas un moyen d'optimisation. Il s'agit d'un moyen conditionnant le bon fonctionnement du réseau tout entier. D'un autre point de vue, ce paramètre peut jouer un rôle important dans les parties du réseau soumises à des charges importantes.

Souvent les utilisateurs des services d'un réseau ne sont répartis uniformément à travers de celui-ci. Ceci est dû d'abord à une répartition non uniforme de la population sur une région donnée. Ensuite, il y a une forte concentration de demandes dans les zones d'activités (bureaux, sièges d'entreprises,). Dans ces régions (dites "Hot Spots" ou "point chaud"), on constate donc de très fortes interférences intracellulaires (du fait du nombre élevé d'utilisateur dans cette cellule). Mais le comble est que ces interférences font que la capacité de la cellule qui se trouve au centre du "Hot Spot" diminue considérablement (voir Simulation, dans la troisième partie). Un moyen d'amener cette capacité à une valeur raisonnable consiste à faire augmenter les puissances d'émission des mobiles afin que le rapport Signal sur Interférence conserve une valeur supérieure à un seuil donné.

Il y a deux principes adoptés pour le contrôle de puissance:

- Le contrôle en boucle ouverte (Open Loop Power Control), où le terminal mobile mesure l'interférence sur le canal sur la base de la puissance reçue à partir du signal pilote diffusé par la station de base et ajuste sa puissance d'émission en conséquence. Cette technique permet de combattre les variations relativement lentes des paramètres du canal de transmission, mais reste impuissante face aux fading rapides (fading de Rayleigh, et fading de Rice) étant donné que ceux-ci entraînent des fluctuations rapides sur la liaison montante, qui n'est pas corrélée à la liaison descendante. Cependant, ce type de contrôle de puissance est essentiel pour l'établissement d'une connexion par le terminal mobile dans la mesure où une émission trop puissante provoquerait une interférence inacceptable dans la cellule concernée.
- Le contrôle en boucle fermée, où la station de base mesure le rapport Signal sur Interférence, et ensuite envoie vers le terminal mobile une commande lui donnant l'ordre d'ajuster sa puissance d'émission. Cette technique permet (si elle est exécutée à une fréquence assez élevée) de maintenir la puissance reçue à la station de base à un niveau constant pour tous les terminaux mobiles de la cellule correspondante. Le contrôle en boucle fermée permet aussi entre autre de mettre en œuvre le "Soft Handover" (voir plus loin dans ce chapitre) qui doit garantir que le mobile sera toujours connecté à la station de base qui le recevra le mieux.

VII.2.2 La puissance d'émission de la station de base [21] [22]

La puissance d'émission de la station de base détermine l'étendue de la zone de couverture de la cellule correspondante. Si cette étendue augmente, le nombre d'utilisateurs liés à cette cellule augmente aussi. Cependant, les puissances d'émission des utilisateurs se trouvant près du bord de la cellule agrandie seront tellement élevées qu'elles provoquent des interférences inadmissibles dans les cellules adjacentes, ce qui se traduira par une réduction de la capacité totale du réseau. D'où pour limiter les interférences vers les cellules voisines, on ne doit pas dépasser un certain nombre d'utilisateur par unité de surface dans une cellule. Ce qui nous conduit à utiliser des cellules de faibles étendues pour les zones à fortes concentrations d'utilisateurs (Hot Spots). De là, on voit donc que la puissance d'émission de

la station de base doit être ajustée pour modifier la taille de la cellule qui influera à son tour sur la capacité globale du réseau.

VII.2.3 La sectorisation [23] [32]

La sectorisation est grandement facilitée dans un système DS-CDMA parce qu'il suffit de déployer un certain nombre d'émetteur identique (travaillant dans la même bande de fréquence, et ayant des antennes de même caractéristiques) que de secteurs à couvrir.

Au sein d'un réseau de téléphonie mobile, la sectorisation vise à isoler plusieurs parties d'une même cellule. Cette isolation se traduit par un taux d'interférence moindre pour chaque secteur ce qui à son tour engendrera une augmentation de la capacité de la cellule dans le même ordre de grandeur. Donc si on utilise 3 secteurs par cellule, on devrait (dans un réseau DS-CDMA) aboutir à une multiplication par 3 de la capacité totale de la cellule. Cependant ceci suppose qu'il n'y a pas de chevauchement entre les secteurs, ce qui n'est pas faisable en pratique. Des simulations ont été faites dans [32] pour arriver à la conclusion que pour un angle de chevauchement de 5° entre les secteurs, la capacité d'une cellule à 3 secteurs se trouve multipliée par 2.88.

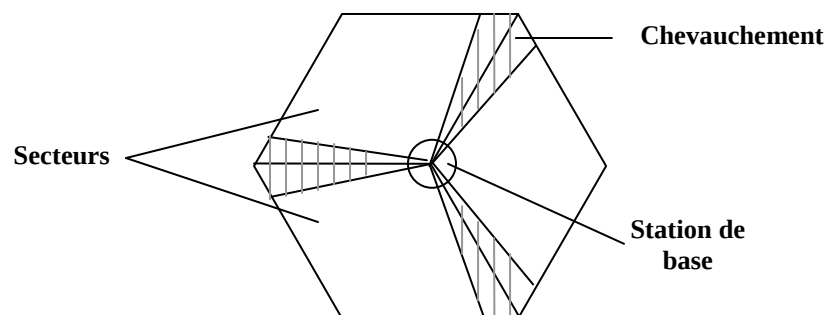


Figure 7.1 : la sectorisation d'une cellule

De plus, avec les antennes actuelles, on peut modifier dynamiquement l'étendue de chaque secteur, pour tenir compte de la variation des charges sur les secteurs et cellules adjacents. On parle alors de sectorisation adaptative. Lorsqu'un secteur est très chargé, on réduit son angle de couverture et en même temps augmenter l'étendue d'un des secteurs adjacents dont la charge est moindre. En fait on effectue des transferts d'appel d'un secteur vers un autre moins chargé (c'est le Softer Handover, voir plus loin dans ce chapitre).

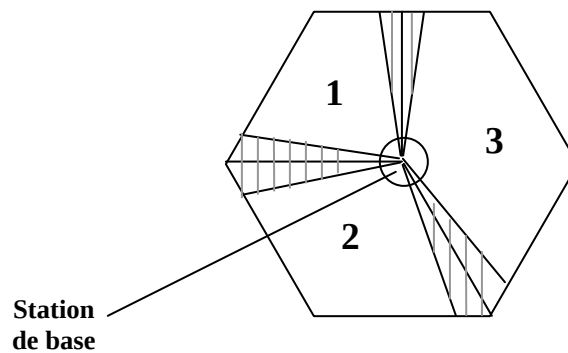


Figure 7.2 : la sectorisation adaptative

Sur la figure précédente, on voit que le secteur 1 est le plus petit de tous, c'est aussi le secteur le plus actif. Lorsque le niveau de l'interférence dans le secteur 1 devient gênant, les usagers se trouvant près des bords se verront affectés sur les autres secteurs (2 ou 3). Ensuite l'angle de couverture du secteur 1 sera réduit d'autant, tandis que le nouveau secteur chargé de l'appel sera agrandi.

VII.2.4 La mise en place de plusieurs couches de cellules [22]

Dans les zones fortement peuplées comme le cas des quartiers administratifs, le trafic augmente rapidement, particulièrement à l'intérieur des bâtiments. Or on sait que lorsque la densité d'appel est très élevée, on doit réduire la couverture de la cellule. A l'intérieur de ces bâtiments, on doit installer de petites stations de base pour servir les usagers qui sont à l'intérieur. Dans un réseau de téléphonie utilisant le DS-CDMA, on utilise plusieurs couches de cellules ayant chacune des rôles spécifiques (voir figure 7.3) :

- Les macrocellules : pour assurer une couverture globale (toute l'étendue de la cellule avec une grande mobilité des usagers).
- Les microcellules : pour assurer la couverture des zones fortement peuplées (à l'extérieur des bâtiments)
- Les picocellules : pour assurer la couverture de l'intérieur des bâtiments des zones fortement peuplées avec une faible mobilité des usagers.

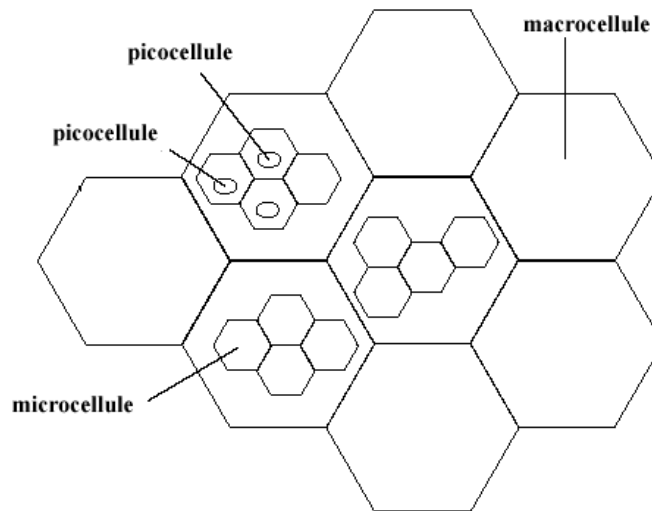


Figure 7.3 : les couches de cellules dans un réseau DS-CDMA

Etant donné que les systèmes cellulaires utilisant le DS-CDMA peuvent partager la même bande de fréquence entre les cellules adjacents aussi bien qu'avec les différentes couches de cellules qui se superposent, la mise en place de ces dernières est un moyen pour augmenter la capacité totale du réseau sans avoir recours à une bande de fréquence supplémentaire. Ceci n'exclut pourtant l'utilisation (si possible) d'une bande de fréquence différente pour chaque couche de cellule (ce qui conduirait à une meilleure performance).

VII.3 Les nouvelles notions introduites par le DS-CDMA

Ces éléments sont propres à un système DS-CDMA et servent à augmenter encore les performances de celui-ci.

VII.3.1 Les récepteurs RAKE [6] [16] [28]

Dans un canal de transmission à trajets multiple, le signal original émis se réfléchit sur des obstacles et le récepteur obtient plusieurs copies de ce signal avec différents retards. Un récepteur DS-CDMA peut les éliminer efficacement si les retards sont supérieurs à la durée d'un chip. Dans ce cas, les signaux retardés sont éliminés par la propriété d'auto-corrélation de la séquence d'étalement utilisée. Cependant, on peut obtenir une diversité temporelle en combinant les signaux issus des trajets multiples dans un récepteur RAKE. Ceci permet de

surmonter les effets des fadings car les différents trajets suivis par les signaux parvenant au récepteur possèdent des statistiques d'évanouissement qui sont indépendantes.

Le récepteur RAKE est composé de corrélateurs, chacun recevant un signal correspondant à un trajet. La figure 7.4 montre le principe de base d'un récepteur RAKE.

Après l'étalement de spectre et la modulation le signal est émis et passe à travers un canal à trajets multiples, qui peut être modélisé par des retards et des atténuations. Sur la figure 7.4 on a trois composantes avec trois différents retards (τ_1 , τ_2 et τ_3) et trois différentes atténuations (a_1 , a_2 et a_3), chacune correspondante à un trajet dans le canal de transmission. Dans le récepteur RAKE, à chaque composante correspond un corrélateur (appelé communément "*finger*"). Dans chaque corrélateur, le signal reçu est dé-étalé avec le code d'étalement requis mais retardé de τ_1 , τ_2 ou τ_3 . Ensuite, les signaux sont pondérés et combinés.

Remarquons que pour fonctionner, le récepteur RAKE doit avoir des estimations sur les retards et les atténuations. Ceci est indispensable pour reconstituer convenablement le signal à la sortie. On a volontairement omis les blocs effectuant les estimations sur la figure 7.4 pour des raisons de simplicité.

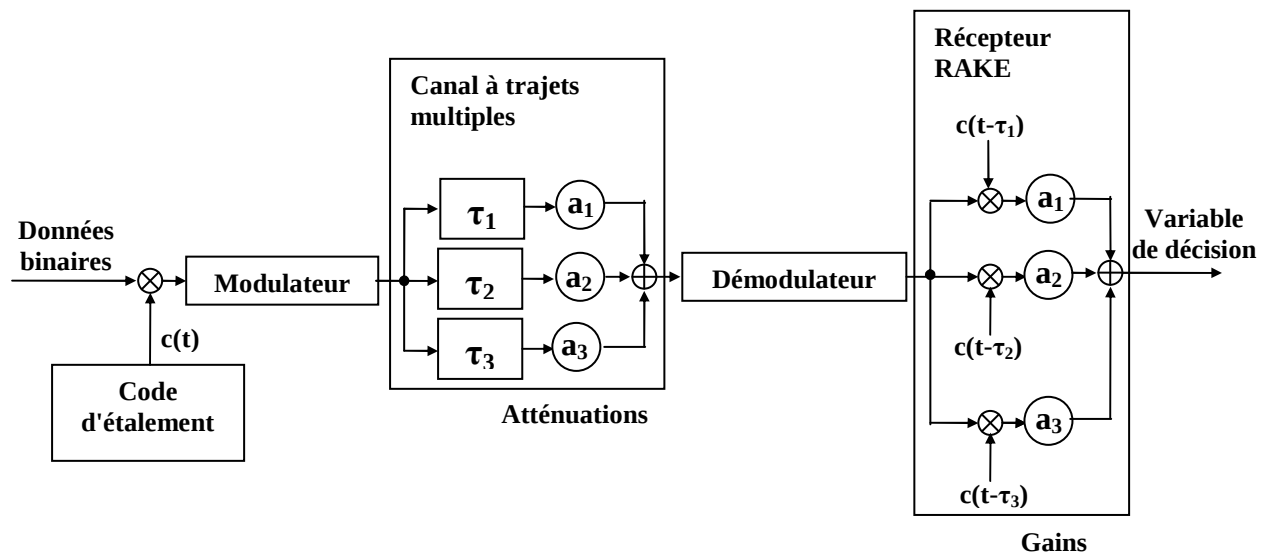


Figure 7.4 : principe de base d'un récepteur RAKE

VII.3.2 *Le Soft Handover* [6] [14] [20] [32]

Habituellement, un terminal mobile effectue un "handover" lorsque le signal provenant de la station de base d'une des cellules voisines est plus fort que celui de la cellule courante avec un seuil donné. C'est le "Hard Handover". Ceci n'est pas utilisable dans un réseau DS-CDMA car il provoquerait une augmentation excessive de l'interférence dans les cellules avoisinantes (étant donné qu'on utilise la même bande de fréquence dans tout le réseau). Ce phénomène est dû au fait que le mobile va émettre avec une puissance supérieure à celle vraiment nécessaire pour être sûre d'être entendu par la nouvelle station de base (ce qui provoquerait un effet proche-lointain non négligeable).

Dans le cas du Soft Handover, le mobile est relié à plusieurs stations de base simultanément. La puissance d'émission du mobile est contrôlée par la station de base qui le reçoit avec la plus forte puissance. De cette manière, la puissance d'émission du mobile est toujours sous contrôle. Le mobile entre dans l'état de Soft Handover lorsque la puissance du signal de la station de base de la cellule voisine dépasse un certain seuil, mais le mobile reste encore sous le contrôle de la station de base actuelle.

Le DS-CDMA se prête bien à l'utilisation du Soft Handover parce que :

- dans la liaison montante, deux ou plusieurs stations de base peuvent recevoir le même signal émis par un mobile ce qui permettrait de choisir en permanence la liaison de meilleure qualité (diversité spatiale). Et ce sans avoir recours à une bande de fréquence supplémentaire.
- dans la liaison descendante, le récepteur RAKE contenu dans le mobile peut combiner (diversité temporelle) les signaux provenant de plusieurs stations de base comme s'il s'agissait des signaux issus d'une propagation par trajets multiples.

Un exemple est montré sur la figure 7.5 avec un contrôleur de station de base (BSC) et deux stations de base (BTS).

A part ces avantages déjà cités, le Soft Handover permet aussi d'avoir un flux de données ininterrompu lors du passage d'une cellule vers une autre (ce que le Hard Handover ne permettrait pas). Ceci peut être assuré car le mobile est toujours connecté à plus d'une station de base à la fois.

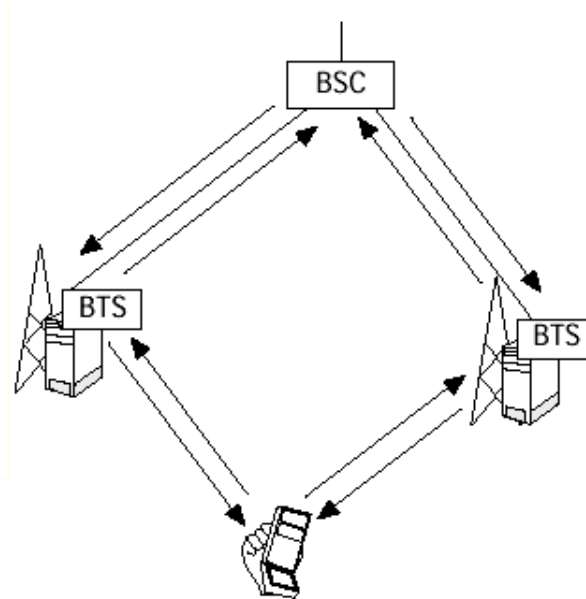


Figure 7.5 : principe du Soft Handover avec deux stations de base

Remarquons que le Soft Handover introduit des communications redondantes, ce qui entraîne des pertes sur la capacité de la liaison montante. Dans [32], on peut voir une simulation qui démontre que ces pertes varient de 0.2 % à 1.1 % pour différentes activités vocales ($3/8$ à $1/2$). Mais comme la capacité de la liaison montante est supérieure à celle de la liaison descendante, ces pertes n'ont aucune influence sur la capacité du système.

VII.3.3 Le Softer Handover [6] [14] [20] [32]

Le Softer Handover est un cas particulier du Soft Handover, car ce n'est que le Soft Handover appliqué à l'intérieur d'une cellule pour transférer les usagers entre les secteurs de celle-ci. Ceci étant, toutes les règles utilisées pour le Soft Handover sont toutes valables pour le Softer Handover.

VII.3.4 L'Interfrequency Handover [6] [14] [20]

On parle de Interfrequency Handover lorsqu'on utilise plus d'une bande de fréquence dans un réseau DS-CDMA. C'est l'action de transfert d'un usager d'une fréquence (celle utilisée par la cellule en cours) vers une autre (celle d'une autre couche de cellule, d'un autre opérateur ou d'un autre système). C'est primordial dans un système pour supporter :

- les structures en couches des cellules, lorsque les couches macro, micro et pico utilisent des fréquences différentes.
- effectuer des handovers entre différents opérateurs : ceci est prévu par l'UMTS.
- effectuer des handovers entre différents systèmes : c'est le cas du GSM qui doit être supporté au sein d'un réseau UMTS.

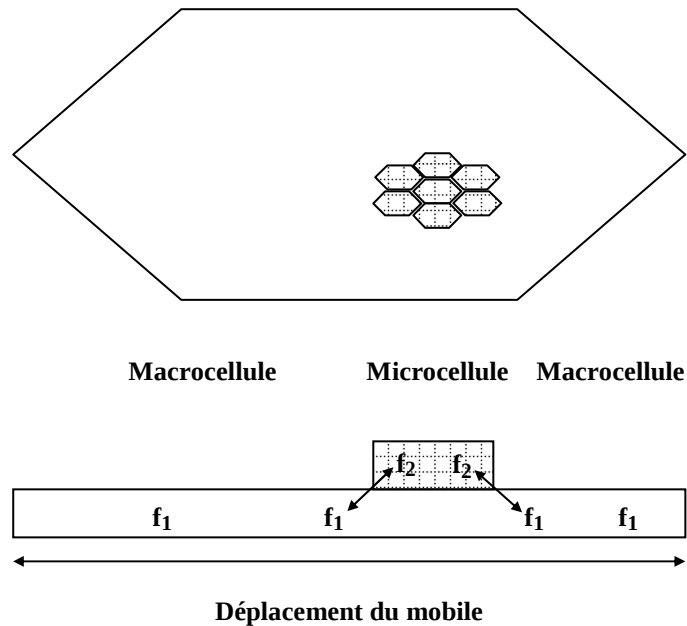


Figure 7.6 : illustration de l'Interfrequency Handover entre deux couches de cellule

Pour assurer la transparence de l'Interfrequency Handover, des mesures sur d'autres fréquences doivent être possibles sans perturber le flux normal des données.

VII.3.5 Les récepteurs à détecteur Multi-Utilisateurs [4] [6] [14] [20] [29] [30]

Les récepteurs DS-CDMA actuels sont basés sur le récepteur RAKE, qui considère les signaux des autres utilisateurs comme des interférences. Cependant dans un récepteur optimal, tous les signaux devraient être détectés ensemble. Dans ce cas, les interférences provenant des autres utilisateurs seront seulement soustraites pour avoir le signal utile. Ceci est possible parce que les interférences sont ici déterministes (étant donné que les codes d'étalement utilisés par les utilisateurs sont connus du système) mais non pas aléatoire.

La capacité d'un système DS-CDMA utilisant les récepteurs RAKE est limitée par l'interférence. La détection multi-utilisateurs (MUD : Multi-User Detection), fournit un moyen pour réduire les effets de l'interférence due à l'accès multiple, et de ce fait augmente la capacité du système. En premier lieu, la MUD est destinée à éliminer l'interférence intracellulaire, c'est-à-dire que dans ce cas la capacité du système est limitée par l'interférence intercellulaire et l'efficacité de l'algorithme employé. En plus de l'augmentation de la capacité du système, la MUD permet aussi de combattre les effets proche-lointain inhérents à tout système DS-CDMA.

Il existe deux grandes catégories de récepteurs à MUD :

- les détecteurs linéaires
- les éliminateurs d'interférence (interference cancellation).

Les détecteurs linéaires recourent à des transformations linéaires pour éliminer les interférences entre les utilisateurs. Les décorrélateurs et les détecteurs LMMSE (Linear Minimum Mean Square Error) sont des exemples de détecteurs linéaires.

Quant aux éliminateurs d'interférence, l'interférence due à l'accès multiple est d'abord estimée puis soustrait du signal reçu. Les PIC (Parallel Interference Cancellation) et SIC (Successive Interference Cancellation) en sont des exemples.

Ces détecteurs MUD sont assez complexes à réaliser, et ils sont pour le moment destinés à l'équipement des stations de base (en remplaçant les récepteurs RAKE), afin de pouvoir supporter un plus grand nombre d'abonnés. Ils peuvent avoir plusieurs structures suivant l'algorithme utilisé, mais leurs études dépassent largement le cadre de cet ouvrage. Pour plus de détails à ce sujet, voir [4] [29] [30].

Chapitre VIII

LE WCDMA : INTERFACE RADIO DE L'UMTS

VIII.1 Introduction

Le WCDMA (Wideband CDMA) est une technologie d'interface radio basée sur le DS-SS-SS, à la différence près que celui-ci utilise une plus grande bande de fréquence que celle utilisée par le DS-SS-SS (5 MHz pour l'UMTS contre 1.25 MHz pour l'IS 95). Cette extension de la bande passante a été sollicitée par une demande toujours croissante en termes de débit de données tout en gardant un facteur d'étalement honorable.

Le WCDMA a été retenue comme interface radio de la déclinaison européenne de l'IMT 2000 (International Mobile Telecommunications 2000), qu'est l'UMTS.

VIII.2 Les systèmes de communications mobiles de la troisième génération (3G) [6] [14] [20]

Les systèmes de communications mobiles de la troisième génération (3G) sont destinés à fournir une mobilité globale avec toute une gamme de services incluant la téléphonie, la messagerie, l'Internet et les transferts de données à haut débit. Ils doivent satisfaire à la norme IMT-2000 élaborée par l'UIT (Union Internationale des Télécommunications). Il appartient ensuite à chaque organisation régionale de proposer leurs propres versions de cette norme.

En Europe, cette organisation est l'ETSI (European Telecommunications Standards Institute). C'est elle qui est donc chargée de la standardisation de la norme UMTS (Universal Mobile Telecommunication System) qui sera aussi l'ultime stade de l'évolution du système GSM.

VIII.3 Les bandes de fréquence utilisées par l'UMTS [6] [14] [20] [26]

En Février 1992, le World Radio Conference avait alloué pour l'usage de l'UMTS les bandes de 1885-2025 et 2110-2200 MHz. Le partage de ces fréquences se fait comme suit :

□ 1920-1980 et 2110-2170 MHz : utilisées en mode FDD (Frequency Division Duplex). On parle alors de WCDMA. Les liaisons montantes et descendantes sont appariées ; l'espacement des canaux est de 5 MHz. Un opérateur aura besoin de 3 ou 4 canaux (2 x 15 ou 2 x 20 MHz) pour mettre en place un réseau à haut débit et à capacité élevée.

□ 1900-1920 et 2010-2025 MHz : utilisées en mode TDD (Time Division Duplex). On parle alors de TD-CDMA. Les canaux ne sont pas appariés, c'est à dire que l'émission et la réception ne sont pas séparées en fréquence. L'espacement des canaux est toujours de 5 MHz.

□ 1980-2010 et 2170-2200 MHz : réservées pour les liaisons par satellite (montante et descendante).

VIII.4 Les services fournis par l'UMTS [6] [14] [20] [26]

Il y a plusieurs classes de services que l'UMTS peut supporter :

- ◆ Conversations (voix, visiophonie)
- ◆ Streaming (multimédia, vidéo à la demande, ...)
- ◆ Interactivité (navigation Web, jeux en réseau, accès à des bases de données)
- ◆ Services d'arrière-plans (E-mail, SMS (Short Message Service), téléchargement).

D'autre part, on ne peut pas oublier non plus que l'UMTS possède une sécurisation réseau plus poussée. Pour mettre en œuvre ces services, les débits maximaux que l'on s'est fixé d'atteindre sont :

- ◆ 144 kbits/s pour les liaisons par satellite et pour les zones rurales.
- ◆ 384 kbits/s pour les zones urbaines (à l'extérieur des bâtiments).
- ◆ 2048 kbits/s pour les zones urbaines (à l'intérieur des bâtiments).

VIII.5 L'architecture d'un système UMTS [14] [20]

Un réseau UMTS est constitué de trois domaines :

- ◆ Le Core Network (CN) ou le réseau cœur
- ◆ L'UMTS Terrestrial Radio Access Network (UTRAN) ou le réseau d'accès

◆ L'Équipement d'abonné (UE : User Equipment)

La principale fonction du cœur du réseau consiste à fournir la commutation, le routage et le transit des trafic des usagers. Il contient aussi les bases de données sur les utilisateurs ainsi que les fonctions de gestion du réseau. L'architecture de base du CN a été fortement inspiré du réseau GSM avec l'extension GPRS (Global Packet Radio Service). Cependant tous les équipements doivent être modifiés et adaptés pour les services et opérations UMTS.

L'UTRAN fournit la méthode d'accès sur l'interface constitué par l'espace libre pour les Equipements d'abonné (UE). Il définit l'interface radio (méthode d'accès multiple et de duplexage ainsi que les paramètres correspondants) utilisée par les UE pour accéder aux services offerts par le réseau UMTS. La station de base est dénommée NodeB et l'équipement qui contrôle les NodeB est appelé Radio Network Controller (RNC).

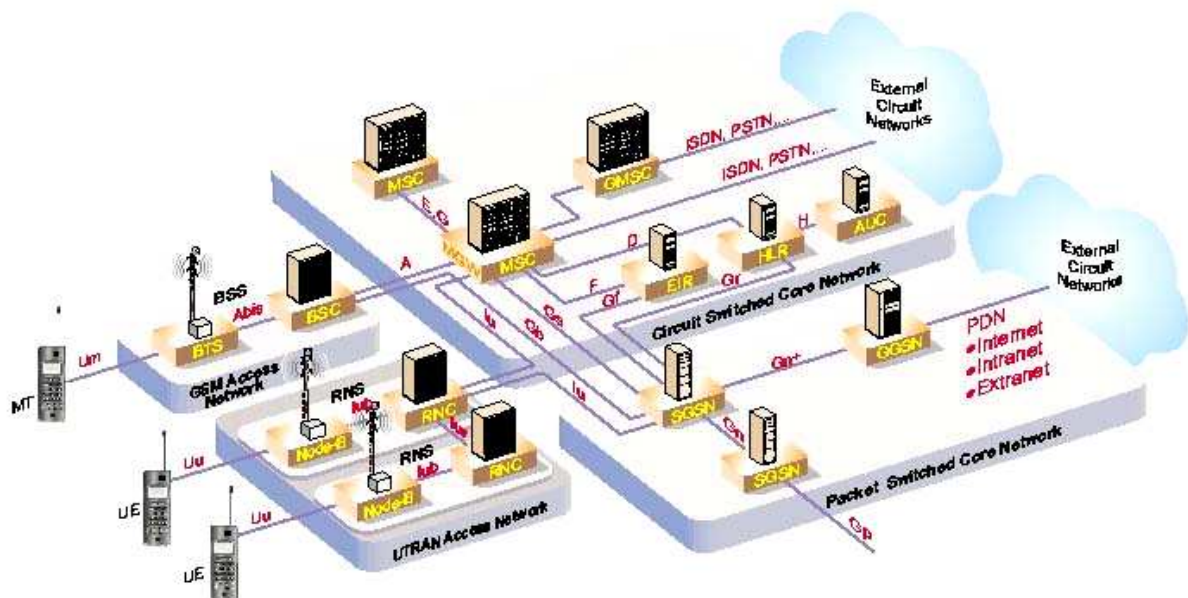


Figure 8.1 : architecture d'un réseau UMTS

VIII.6 La technologie WCDMA [6] [14] [20]

La technologie WCDMA (Wideband CDMA) est l'une des principales technologies utilisées pour l'implémentation des systèmes mobiles de la troisième génération. Elle est basée sur le DS-CDMA, mais utilise une plus grande largeur de bande pour la transmission.

La WCDMA utilise le FDD comme méthode de duplexage. Les paires de bandes de fréquences utilisées sont 1920-1980 et 2110-2170 MHz. La largeur de bande minimale est de 2 x 5MHz. Pour le codage de canal, on utilise le codage convolutif (pour des transferts de données à haut débit, on utilise plutôt les Turbo Codes).

Ici on utilise l'étalement par séquence directe. Le débit de la séquence d'étalement est de 4.096 Mchips/s. Le débit maximal binaire est de 2.048 Mbps, mais en pratique ceci n'est disponible que pour un nombre limité d'utilisateurs par cellule. Ce système utilise un étalement dont le facteur d'étalement peut varier de 4 à 256. Sur la liaison montante, on emploie des codes OVSF pour la séparation des canaux (*channelization codes*, ou encore **codes de canalisation**), et des codes de Gold de longueur 2^{41} pour la séparation des utilisateurs (*scrambling codes*, ou encore **codes de brouillage**). Sur la liaison descendante on utilise des codes OVSF pour la séparation des utilisateurs, et des codes de Gold de longueur 2^{18} pour la séparation des cellules.

Notons aussi que ce système utilise des boucles rapides (1.6 kHz) de contrôles de puissances (ouvertes et fermées).

VIII.7 Les canaux de transport

Les canaux de transport suivants ont été définis pour le WCDMA.

Les trois canaux de contrôle communs sont :

- Le BCCH (Broadcast Control CHannel) : canal de diffusion des informations spécifiques à une cellule (liaison descendante).
- Le PCH (Paging CHannel) : canal d'appel du mobile, utilisé lorsque le NodeB ne connaît pas la localisation du terminal mobile. Dans ce cas, un "message d'appel" est diffusé sur tout le réseau (liaison descendante).
- Le FACH (Forward Access CHannel) : canal destiné à transporter des informations spécifiques à un mobile ou à un groupe de mobiles, lorsque la localisation du (ou des) mobile(s) est connu (liaison descendante).
- Le RACH (Random Access CHannel) : canal destiné à véhiculer des demandes de ressources par le mobile, ainsi que de courts paquets de données (liaison montante).

Il y a aussi deux autres canaux de transport dédiés :

- Le DCCH (Dedicated Control CHannel) : canal transportant les informations de contrôle spécifiques à un mobile (dans les deux sens).
- Le DTCH (Dedicated Traffic CHannel) : canal utilisé pour des transmissions point à point dans les deux sens.

VIII.8 Les canaux physiques [6] [14] [20]

Les canaux logiques sont organisés et traduits en canaux physiques pour pouvoir être émis. Il y a donc des correspondances entre les canaux logiques et les canaux physiques.

VIII.8.1 La liaison montante

- Le DPDCH (Dedicated Physical Data CHannel) : canal physique dédié destiné à porter le DTCH. Il peut y avoir zéro, un ou plusieurs DPDCH à chaque connexion (ceci dépend du débit utilisé). Ce canal est segmenté en plusieurs trames de 10 ms, chaque trame comportant 160×2^k bits (16×2^k kbits/s) où $k=0, 1, \dots, 6$, correspondant à un facteur d'étalement égal à $\frac{256}{2^k}$ avec un débit chip de 4096 Mc/s. Chaque trame de 10 ms est encore subdivisée en 16 IT (slots) de durée 0.625 ms (une période de la boucle de contrôle de puissance). A chaque slot correspondent donc 10×2^k bits. Ici, on voit donc que le contrôle de la puissance d'émission du mobile se fait à une fréquence de 1.6 kHz.
- Le DPCCH (Dedicated Physical Control CHannel) : canal physique dédié destiné à porter DCCH. Il n'y a qu'un DPCCH à chaque connexion. Ce canal porte les bits pilotes, les commandes de contrôle de puissance (TPC : Transmit Power Control), ainsi que les bits indicateurs de format (TFI: Transport Format Indicator). Les bits TFI servent à indiquer au récepteur les formats utilisés pour la transmission des données en cours, en particulier le facteur d'étalement et le codage utilisés.
- Le PRACH (Physical Random Access CHannel) : canal physique destiné à porter RACH.

Remarque : Les canaux DPCCH et DPDCH sont multiplexés dans le code pour la liaison montante (et non pas dans le temps), c'est-à-dire qu'ils utilisent des codes de canalisation différents.

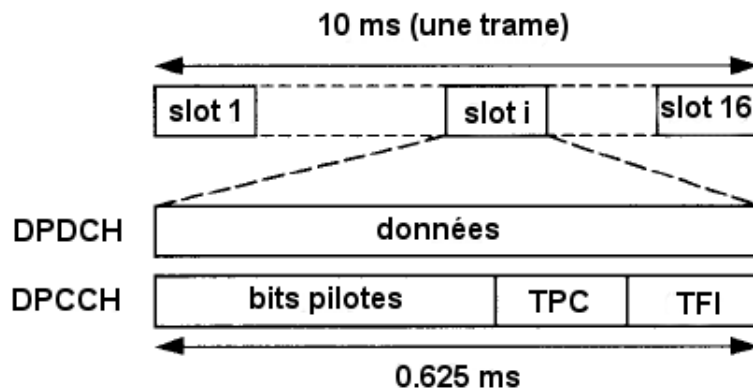


Figure 8.2 : structure de trame pour la liaison montante

VIII.8.2 La liaison descendante.

- Le DPCH (Downlink Physical Channel) : canal physique dédié équivalent à l'ensemble DPDCH et DPCCH multiplexés dans le temps (et transmis avec le même code). DPDCH contient toujours les bits pilotes, TPC et TFI. Ici encore, chaque trame de 10 ms est divisé en 16 slots de durée 0.625 ms.

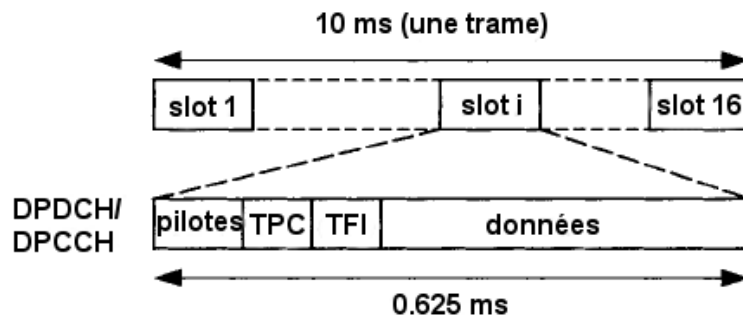


Figure 8.3 : structure de trame pour la liaison descendante

- Le CCPCH primaire (Primary Common Control Physical CHannel) : canal physique commun destiné à porter BCCH. Le débit sur ce canal est fixe et égal à 32 kbps avec un facteur d'étalement égal à 256. Un même code de canalisation est affecté à CCPCH pour toutes les cellules. Ainsi, un terminal mobile peut toujours

trouver le BCCH, sachant que le code de brouillage unique utilisé par la cellule pour la liaison descendante est déjà connu pendant la recherche d'un NodeB initialement effectuée par le mobile.

- Le CCPCH secondaire (Secondary Common Control Physical Channel) : canal physique commun destiné à porter FACH et PCH. Le débit de ce canal peut être différent pour chaque cellule. Le code de canalisation du CCPCH secondaire est transmis sur le CCPCH primaire.
- Le SCH (Synchronization CHannel) : canal physique destiné à assurer la synchronisation entre la séquence PN utilisée à l'émission et celle générée localement dans le récepteur. Le SCH est constitué de deux sous-canaux : le SCH primaire et le SCH secondaire. Le SCH primaire n'est pas modulé et ne sert qu'à acquérir le timing du SCH secondaire. Le SCH secondaire est par contre modulé et porte des informations (numéro du groupe de code) sur le code de brouillage (code long) utilisé par la cellule. Cette information supplémentaire va donc réduire le temps de recherche et d'acquisition du code de brouillage utilisé. Le SCH primaire consiste en un code non modulé de 256 chips, qui est émis une fois par trame. Ce code est le même pour chaque NodeB dans le système et est émis aligné dans le temps avec la limite de chaque trame de 10 ms. Le SCH sert aussi à l'estimation de la puissance reçue à partir d'un NodeB particulier pour savoir si celui-ci pourra prendre en compte un appel parallèlement au NodeB actuel (Soft Handover).

VIII.9 Étalement de spectre et modulation sur les canaux dédiés [6] [14] [20]

La figure 8.4 illustre le processus de modulation et d'étalement de spectre utilisé pour la liaison descendante. Pour la modulation, on a le QPSK, ce qui implique que chaque paire de bits sera étalée avec le même code de canalisation, puis brouillé avec le même code de brouillage.

Différents canaux physiques doivent utiliser différents codes de canalisations. Les codes de canalisation utilisés sont des codes OVSF car ces derniers peuvent maintenir l'orthogonalité mutuelle entre les différents canaux physiques de la liaison descendante même pour des débits différents (donc des facteurs d'étalement différents). L'utilisation des codes OVSF est donc l'une des clés pour avoir une flexibilité de services dans un système WCDMA.

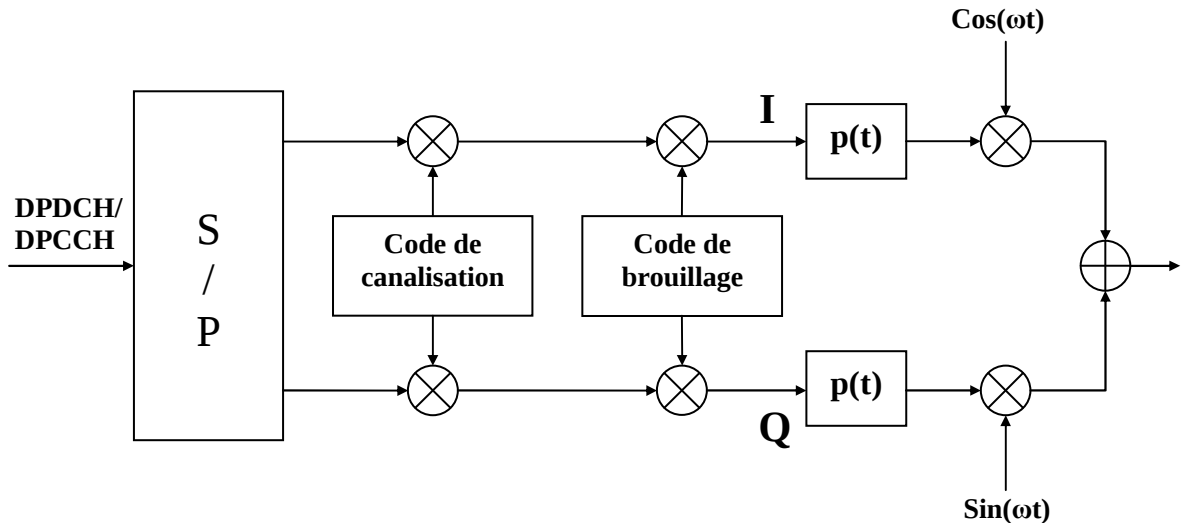


Figure 8.4 : étalement et modulation utilisés pour la liaison descendante

Sur la figure 8.4, $p(t)$ représente un filtre en cosinus surélevé dont le facteur de retombée est $\alpha = 0.22$. La bande occupée par la transmission est donc de $4.096 \cdot (1.22) = 4.997$ MHz, ce qui peut tenir dans les 5 MHz prévus par la norme.

Le code de brouillage utilisé pour la liaison descendante consiste en un segment de 40960 chips d'un code de Gold de longueur 2^{18} (correspondant à une durée de 10 ms). Il y a un total de 512 codes différents disponibles dans le système, divisés en 32 groupes, chacun ayant 16 codes. Cette subdivision a été faite pour permettre une synchronisation rapide sur un NodeB.

Le processus de modulation et d'étalement utilisé pour la liaison montante est donné sur la figure 8.5. Sur cette liaison, on utilise l'étalement complexe de spectre (voir Annexe 3 pour le détail de l'étalement complexe de spectre). Sur la figure 8.5 C_d et C_c représente des codes de canalisations différents; $C_{scrambl}$ est un code de brouillage complexe et $p(t)$ représente un filtre en cosinus surélevé de facteur de retombée $\alpha = 0.22$. La modulation utilisée est le QPSK à deux canaux, parce qu'ici les voies I et Q sont utilisés comme deux canaux BPSK indépendants. Dans le cas d'un unique DPDCH, le DPDCH et le DPCCH sont étalés par deux codes de canalisation différents et sont transmis sur les voies I et Q respectivement. Si on utilise plus d'un DPDCH (débit plus élevé), les DPDCH supplémentaires peuvent être transmis sur I ou sur Q en utilisant des codes de canalisations

supplémentaires. Le signal obtenu est ensuite brouillé par un code de brouillage complexe qui a été alloué pour le mobile.

Les codes de canalisations sont toujours des codes OVSF pour assurer l'orthogonalité entre le DPDCH et le DPCCH.

Pour la liaison montante, on utilise comme codes de brouillage des portions de 10 ms de codes de Gold de longueur 2^{41} , correspondant à 40 960 chips. Cependant, on peut aussi utiliser des codes de Kasami (appartenant à l'ensemble élargi) de longueur 256 (codes courts) comme codes de brouillage. Ces codes courts sont utilisés pour les cellules dont le NodeB est doté d'un récepteur MUD. L'utilisation de ces codes courts autorise l'emploi d'un récepteur MUD de moindre complexité. Cependant, on doit se rappeler le fait que les propriétés d'inter-corrélation des codes courts sont plus mauvaises que celles des codes longs. Ainsi, dans les cellules où un récepteur RAKE ordinaire est employé, on doit utiliser des codes longs.

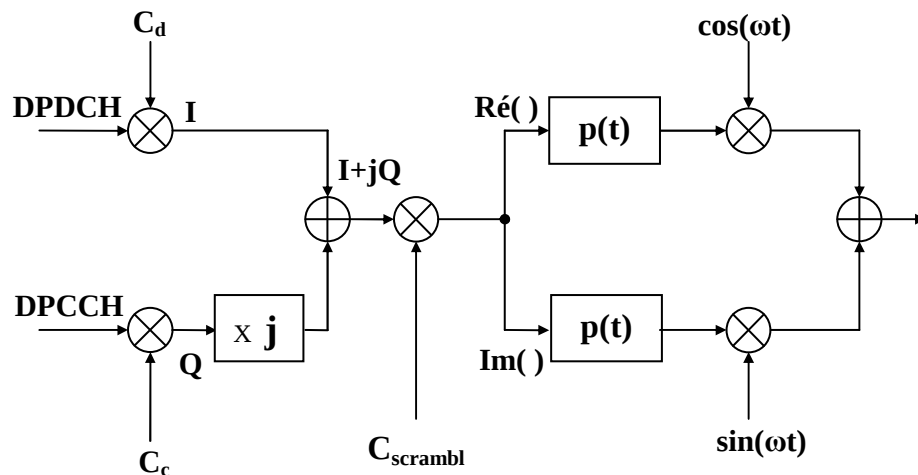


Figure 8.5 : étalement et modulation utilisés pour la liaison montante

VIII.10 Étalement de spectre sur les canaux communs [6] [14] [20]

VIII.10.1 Les CCPCH primaires et secondaires

Comme nous l'avons vu plus haut, ces canaux communs de contrôle sont utilisés pour porter les canaux de transport communs tels que le BCCH, PCH, et le FACH.

La structure de trame utilisée par les CCPCH est similaire à celle utilisée par les DPDCH/DPCCH dans la liaison descendante (multiplexés dans le temps), à la différence

près qu'il n'y a pas de TPC ni de TFI. Etant donné que les CCPCH sont des canaux communs, aucun contrôle de puissance en boucle fermée n'est possible. D'autre part, les CCPCH ont une structure et un débit fixes, ce qui rend l'utilisation du champ TFI superflue.

La canalisation et le brouillage est effectué de la même manière qu'on les a fait pour les DPDCH/DPCCH pour la liaison descendante.

Le CCPCH primaire utilise un code de canalisation prédéfini de longueur 256, spécifique au système, émis continuellement à travers tout le réseau. En procédant de la sorte, le terminal mobile peut donc trouver le CCPCH primaire après la recherche initiale de station de base. Par contre, le CCPCH secondaire emploie un code de canalisation spécifique à une cellule et peut être émis sur toute une cellule ou sur une partie de cette cellule seulement. Ce code de canalisation est diffusé sur le BCCH (qui rappelons-le, est porté par le CCPCH primaire).

VIII.10.2 Le SCH

La recherche initiale d'un NodeB est effectuée avec l'aide du SCH. La structure du SCH est donnée sur la figure 8.6.

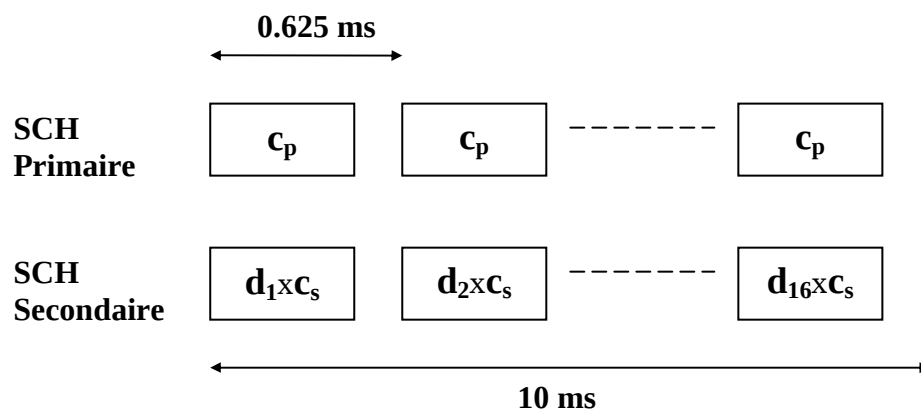


Figure 8.6 : structure du SCH

Le SCH contient deux sous canaux : le SCH primaire et SCH secondaire. Chaque trame de 10 ms est divisée en 16 slots de longueur 0.625 ms. Le SCH primaire est constitué d'un code de Gold (c_p) de longueur 256 chips, non modulé et identique pour chaque slot. Le SCH secondaire est constitué d'un code de Gold (c_s) de longueur 256 (code court), modulé

et émis parallèlement au SCH primaire. Il y a 16 codes différents disponibles pour le SCH secondaire, pour indiquer le groupe du code de brouillage employé par la cellule pour la liaison descendante. La séquence ($d_1 d_2 d_3 \dots d_{16}$) de 16 bits (spécifique au système) utilisée pour moduler le code de Gold sur le SCH secondaire se répète dans chaque trame de 10 ms.

Dans un système WCDMA, la séparation des cellules est obtenue par l'utilisation de différents codes de brouillage. Durant la recherche initiale de la station de base, le mobile doit déterminer le code de brouillage utilisé pour la liaison descendante. Il doit aussi assurer la synchronisation au niveau des trames émises par la station de base.

La procédure utilisée par le mobile pour rechercher un NodeB est donnée dans le Chapitre V, Sous paragraphe V.4.1.4.

VIII.10.3 Le PRACH

Le PRACH sert à transporter les demandes de ressources effectuées par le mobile. Une demande se fait par l'émission d'un paquet dit "d'accès aléatoire". Cette dénomination vient du fait que chaque mobile émet (indépendamment des autres mobiles) des requêtes sur ce canal avant de pouvoir accéder aux services offerts par le réseau.

Un paquet d'accès aléatoire est composé de deux parties : une partie d'en-tête (preamble part) et une partie de message.

L'en-tête est constituée de 16 symboles étalés par un code de Gold de longueur 256 (code d'en-tête). Chaque NodeB dans le réseau dispose de son propre code d'en-tête, qui est diffusé sur BCCH à l'intention des terminaux mobiles.

La partie de message du paquet d'accès aléatoire est constituée par une trame de données (champ d'identification du mobile (MS ID), un champ décrivant le service demandé, un paquet de données (optionnel), et un champ de contrôle de CRC) et une trame de contrôle (TFI) (voir Figure 8.7). Les trames de données et de contrôle sont étalées et modulées comme le sont DPDCH et DPCCH mais en utilisant des codes d'étalement différents (diffusés sur BCCH).

Dans un système WCDMA, on utilise comme méthode d'accès aléatoire l'ALOHA discrétisé ou Slotted-ALOHA, les demandes de connexion étant émises dans des intervalles espacés de 1.25 ms (voir Figure 8.8).

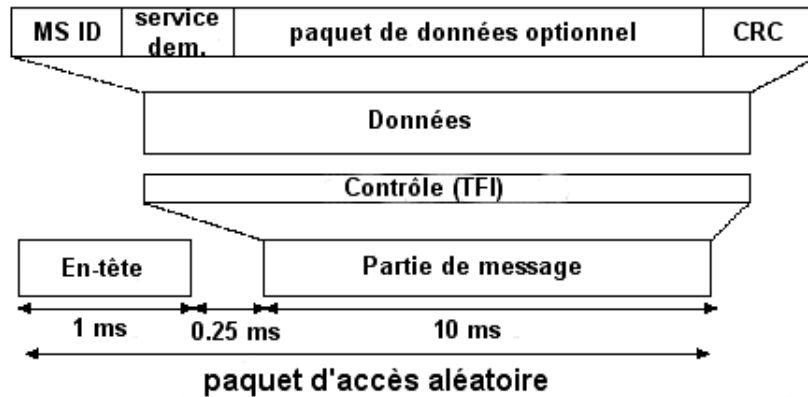


Figure 8.7: structure du paquet d'accès aléatoire

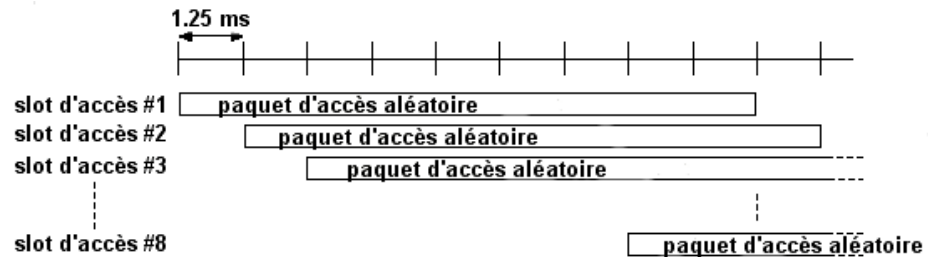


Figure 8.8 : l'accès aléatoire dans un système WCDMA

Entre l'en-tête et la partie de message, il y a une période inutilisée de durée 0.25 ms. Cette période permet la prise en compte de l'en-tête du paquet, et le prochain traitement de la partie de message du paquet d'accès aléatoire.

VIII.11 Les couches de cellules utilisées dans l'UMTS [22] [26] [30]

La norme UMTS prévoit l'utilisation de trois couches de cellules :

- les macrocellules, ayant un rayon de l'ordre de 15 km, pour assurer la couverture globale (campagne, routes, ...) avec une possibilité de déplacement rapide du mobile (voiture, train,...). Le débit maximal est de 384 kbits/s.
- les microcellules, ayant un rayon d'ordre de 500 m pour assurer la couverture dans le milieu urbain (extérieur des bâtiments, rues, ...) avec une vitesse de déplacement limitée (voiture en ville, piéton, ...). Le débit maximal peut varier entre 384 kbits/s et 2 Mbits/s selon le trafic et la vitesse du mobile.

- les picocellules, ayant un rayon de l'ordre de 100 m, pour assurer la couverture à l'intérieur des bâtiments à très forte concentration d'utilisateurs se déplaçant le plus faiblement possible. Le débit maximal est ici de 2 Mbits/s.

La norme UMTS a prévu des fréquences différentes pour chaque couche de cellule. Ainsi une bande appariée de 15 MHz (3 x 5MHz) est nécessaire pour utiliser toute la hiérarchie de couches de cellule.

VIII.12 La réalisation du Soft Handover dans un réseau UMTS

L'UMTS utilise le Soft Handover où le mobile peut se connecter à deux ou plusieurs NodeB dans la même bande de fréquence. Les signaux de la liaison montante sont combinés dans le réseau, tandis que ceux de la liaison descendante le sont dans le récepteur RAKE du mobile. Le mobile cherche continuellement de nouveaux NodeB(s), mais cette recherche est limitée à une liste de NodeB(s) diffusée sur BCCH. Cette liste dicte au terminal mobile l'ordre dans lequel les codes de brouillages utilisés par la liaison descendante seront cherchés. Elle permet aussi de limiter la recherche à un sous-ensemble des codes disponibles.

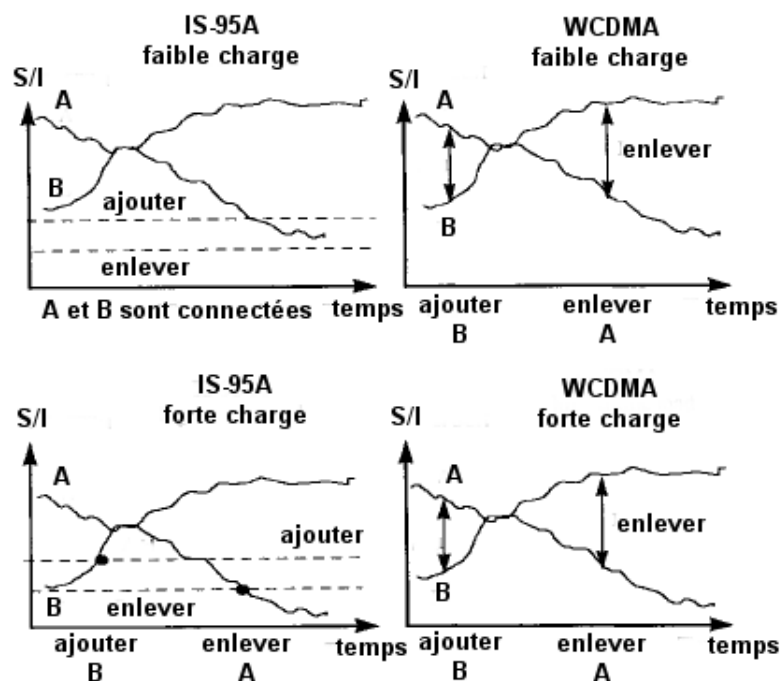


Figure 8.9 : le soft handover dans un système WCDMA

VIII.13 La réalisation de l'Interfrequency Handover dans un réseau UMTS [6] [14] [20]

Etant donné que la transmission se fait d'une manière continue sur la liaison descendante, il n'y a pas de temps pour effectuer une quelconque mesure sur une autre fréquence. Un second récepteur est alors nécessaire pour effectuer cette mesure. Cependant, tous les terminaux UMTS ne disposent pas tous de deux récepteurs.

Donc pour permettre à ces terminaux d'effectuer un Interfrequency Handover, on a recours à une astuce dénommée "slotted mode". Cette astuce consiste à réduire le facteur d'étalement pour une trame dans un facteur de 2.

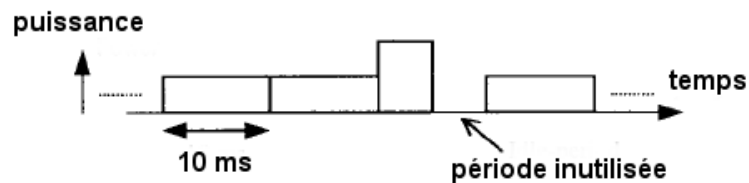


Figure 8.10 : principe du "slotted mode" pour permettre la mesure sur une autre fréquence

Ainsi, une trame de 10 ms peut être transmise en 5 ms. Cette émission est faite avec une puissance plus grande pour compenser la réduction du gain de traitement. En utilisant cette technique une période inutilisée de 5ms est obtenue, pendant laquelle aucune donnée n'est reçue par le terminal et une mesure sur une autre fréquence est possible (voir Figure 8.10).

CHAPITRE IX

SIMULATIONS SOUS MATLAB

IX.1 Introduction

Ce chapitre est consacré à la simulation des systèmes DS-CDMA sous MATLAB. Il se divise en deux parties :

- Simulation du comportement d'un système DS-CDMA à plusieurs utilisateurs. Le comportement d'un tel système est analysé sous l'influence du bruit blanc, d'une interférence à bande étroite, et des propagations à trajets multiples. Cette partie démontre aussi entre autres la faisabilité de la séparation des différents utilisateurs d'un système par le biais d'un ensemble de codes.
- Un exemple de programme d'aide à l'optimisation de la capacité d'un réseau DS-CDMA pluricellulaire chargé. Une méthode de calcul numérique de la capacité optimale d'un réseau DS-CDMA, ainsi qu'une méthode de répartition optimale des utilisateurs en vue de la maximisation de la capacité du réseau sont proposées. Ensuite on analyse l'effet d'un "Hot Spot" sur la capacité de chaque cellule et sur la capacité du réseau tout entier. Et finalement, on propose une méthode d'ajustement de la puissance des mobiles pour optimiser la capacité du réseau tout entier dans le cas de l'existence d'un "Hot Spot".

IX.2 Simulation du comportement d'un système DS-CDMA.

Un système DS-CDMA à plusieurs utilisateurs est simulé dans cette partie. Le nombre d'utilisateurs n'est pas limité. Cependant, seuls les paramètres de 3 utilisateurs sont accessibles. Ceux des autres utilisateurs sont générés par le programme en suivant des lois aléatoires. Remarquons aussi que le programme effectue une simulation en bande de base (c'est-à-dire qu'il n'y a pas de modulation ou de démodulation passe-bande).

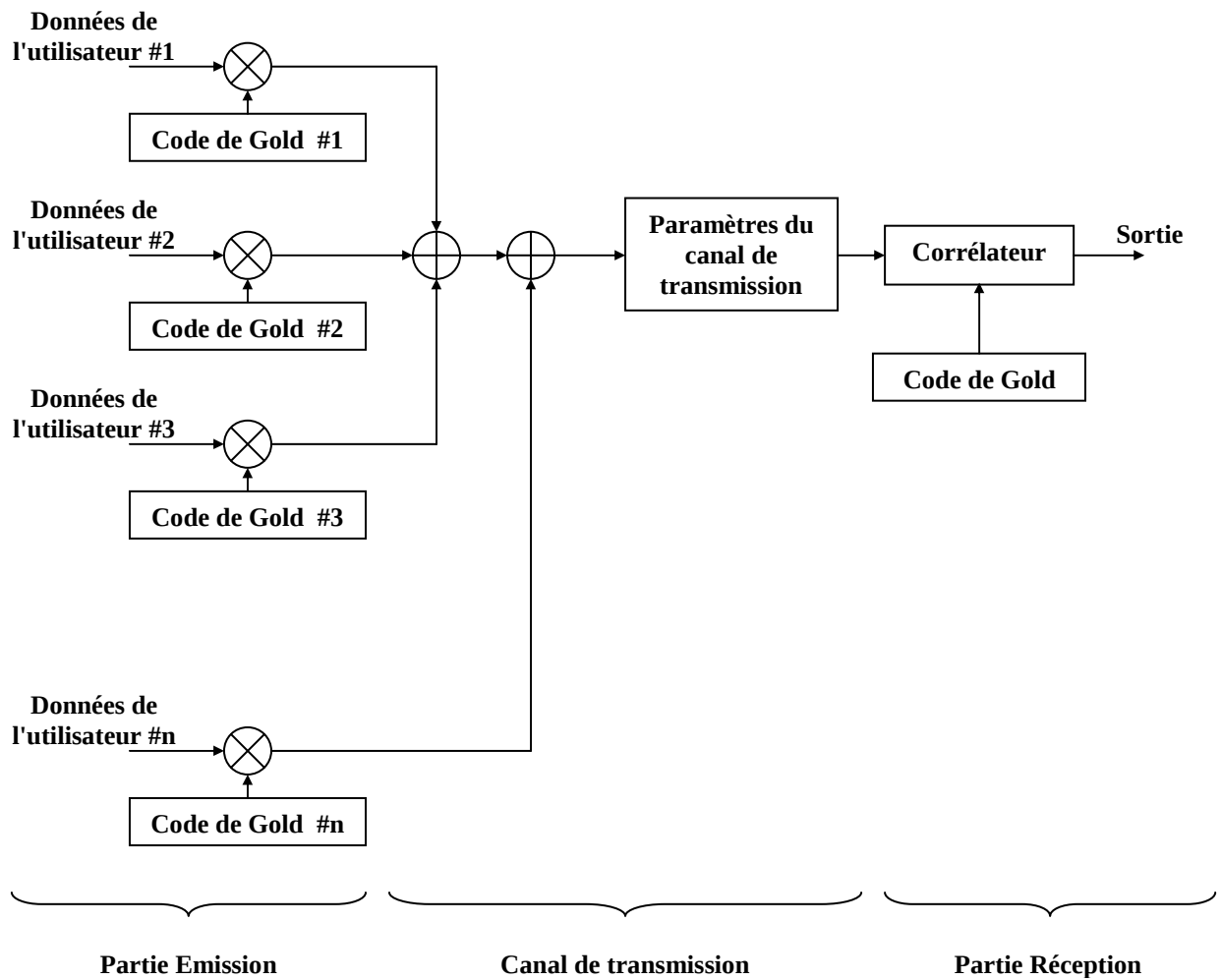


Figure 9.1: structure du système simulé

Le système considéré ici, peut être illustré par la figure 9.1. Chaque utilisateur dans le système peut être caractérisé par les paramètres suivants :

- Les données spécifiques à chaque utilisateur. Une séquence binaire de 16 bits codée en NRZ matérialise alors les données de chaque utilisateur.
- Le code (spécifique à chaque utilisateur) employé pour effectuer l'étalement de spectre. Le programme utilise des codes de Gold pour la séparation des utilisateurs. L'ensemble des codes de Gold (d'une même longueur) utilisé peut être décrit par une paire de polynômes dont les coefficients seront choisis parmi deux listes proposées par le programme. Chaque code dans cet ensemble est

désigné par son numéro, qui représente entre autres le décalage entre les deux polynômes utilisés.

Mais comme nous l'avons signalé plus haut, seuls les paramètres (données et code) des trois premiers utilisateurs seront accessibles, les autres étant générés automatiquement par le programme d'une manière aléatoire. Seul le nombre d'utilisateurs supplémentaires à ajouter dans le système peut être paramétré.

Ensuite les données étalées de chaque utilisateur sont additionnées dans le canal de transmission. Le signal ainsi obtenu fera alors l'objet d'un ou plusieurs traitements supplémentaires selon le profil du canal de transmission choisi depuis l'interface graphique du programme de simulation.

Les paramètres du canal de transmission qu'on peut utiliser dans notre programme sont

- Du bruit blanc additif, Gaussien de moyenne nulle, qui peut être paramétré suivant le rapport Signal sur Bruit (SNR) voulu dans le canal de transmission.
- Des trajets multiples pour la propagation du signal, caractérisés par les amplitudes relatives (par rapport à celle du signal du trajet direct) des signaux correspondants et leurs retards respectifs par rapport au signal du trajet direct (donné en nombre de chips).
- Une interférence à bande étroite, caractérisée par le rapport Signal sur Interférence (SIR) (interférence à bande étroite) en dB. La bande passante occupée par cette interférence est fixée dans le programme et est égal à celle occupée par chaque signal de données non étalé. On utilise pour cela un générateur binaire aléatoire.

La partie réception est effectuée par un corrélateur constitué par un "Matched Filter" dont les coefficients sont fixés par un générateur de codes de Gold. A travers la simulation, on montrera que la sortie donnera une image des données émises par l'utilisateur dont le code utilisé à l'émission est égal à celui utilisé à la réception. L'emploi d'un Matched Filter est justifié par le fait que la sortie de celui-ci donne déjà la polarité des données émises sans l'aide d'un dispositif de synchronisation supplémentaire (voir Chapitre V). D'autre part, le fonctionnement correct d'un système DS-CDMA est conditionné par la bonne marche du processus de synchronisation. Or dans un système réel, un dispositif de synchronisation (qui

peut être implémenté par un Matched Filter) doit précéder le processus de dé-étalement proprement dit. Le corrélateur de la figure 9.1 peut donc être vu comme un dispositif de synchronisation pour un corrélateur actif (qui n'a pas été simulé), mais qui peut déjà donner une image des données reçues, ce qui fait de lui le dispositif idéal pour analyser le comportement du système DS-CDMA considéré.

L'aspect de l'interface graphique utilisée par le programme pour récolter les différents paramètres ainsi que pour afficher les résultats est donné par la figure 9.2.

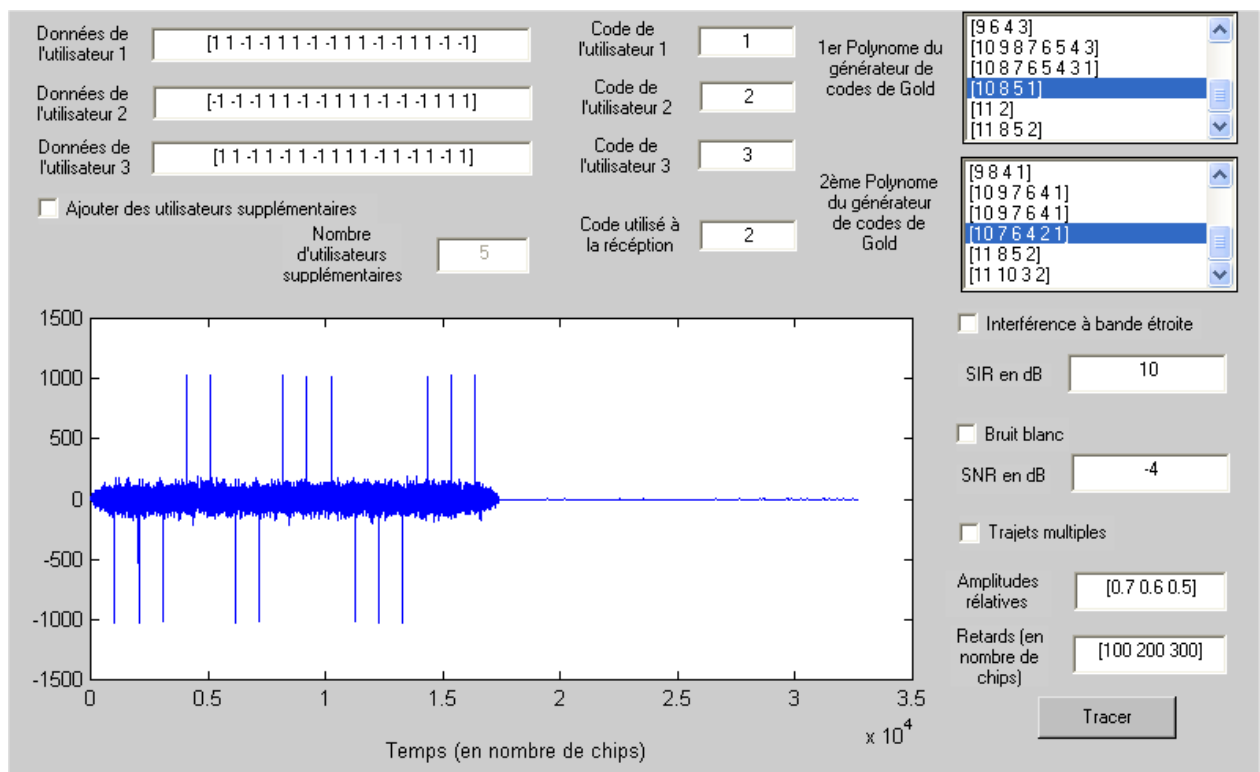


Figure 9.2 : aspect de l'interface graphique utilisée

IX.2.1 Présentation et interprétation des résultats

IX.2.1.1 Les données émises et les codes utilisés par les utilisateurs

Les données de chaque utilisateur sont :

- +1 +1 -1 -1 +1 +1 -1 -1 +1 +1 -1 -1 +1 -1 -1 pour l'utilisateur 1
- -1 -1 -1 +1 +1 -1 -1 +1 +1 +1 -1 -1 -1 +1 +1 +1 pour l'utilisateur 2
- +1 +1 -1 +1 -1 +1 -1 +1 +1 +1 -1 +1 -1 +1 -1 +1 pour l'utilisateur 3

Quant aux codes, l'utilisateur 1 utilise le code **Gold(n,1)**, l'utilisateur 2 **Gold(n,2)** et l'utilisateur 3 le code **Gold(n,3)** où **n** représente le degré des polynômes utilisés (voir Chapitre IV) pour la génération des codes de Gold.

Remarque : Dans toute la suite, l'axe des abscisses désigne le temps (donné en nombre de chips) tandis que l'ordonnée donne la valeur de la sortie de la figure 9.1.

IX.2.1.2 Cas du canal idéal (sans bruit, sans interférence et sans trajets multiples), 3 utilisateurs, le récepteur connaît le code utilisé à l'émission.

Un code **Gold (8,2)** est utilisé à la réception. On voit ici (par la polarité des pics), que les données émises par l'utilisateur 2 sont bien reçues.

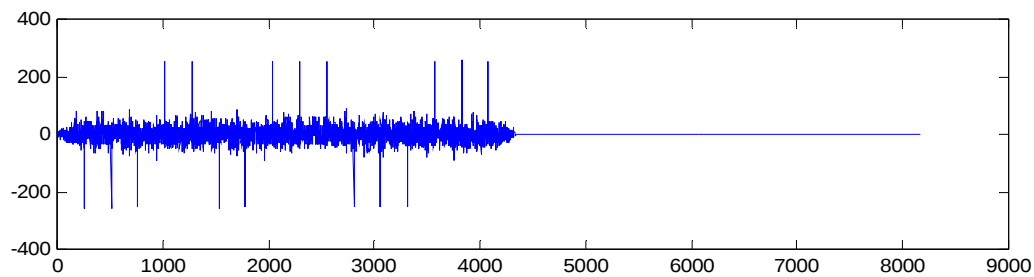


Figure 9.3 : canal idéal, 3 utilisateurs, le récepteur connaît le code

IX.2.1.3 Cas du canal idéal, 3 utilisateurs, le récepteur ne connaît pas le code utilisé à l'émission.

Un code **Gold (8,4)** est utilisé à la réception. Un récepteur qui ne connaît pas le code utilisé à l'émission ne peut pas reproduire les données émises par les utilisateurs.

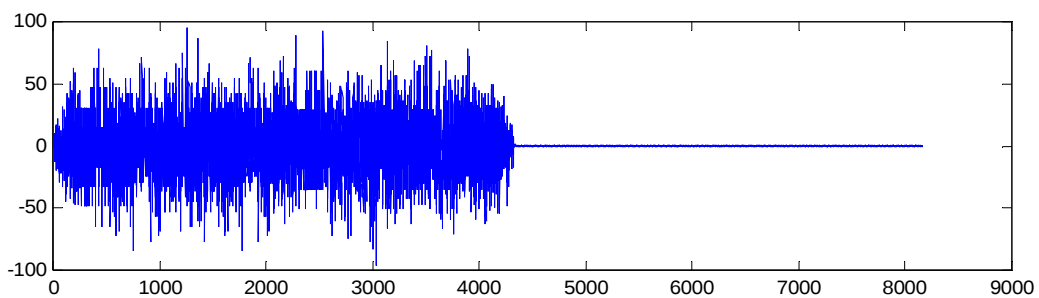


Figure 9.4 : canal idéal, 3 utilisateurs, le récepteur ne connaît pas le code

Dans toute la suite, et sauf indication contraire, le récepteur connaît toujours le code utilisé pour l'émission.

IX.2.1.4 Cas du canal idéal, 3+20 utilisateurs, **Gold (8,2)**.

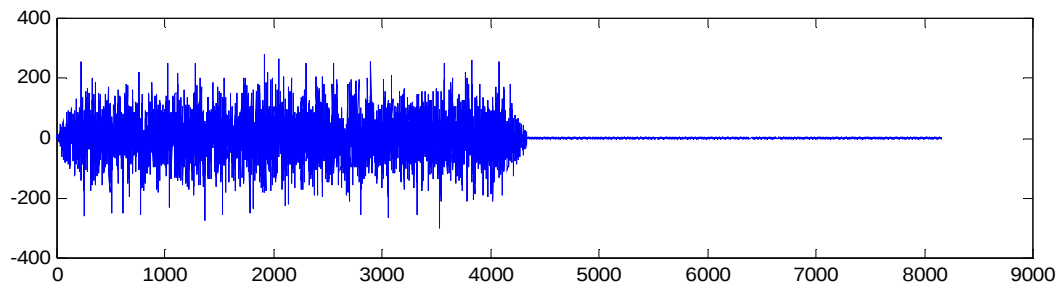


Figure 9.5 : canal idéal, 3+20 utilisateurs, **Gold (8,2)**

Lorsque le nombre d'utilisateurs augmente, l'interférence due aux autres utilisateurs augmente aussi, et peut empêcher toutes les communications lorsque le Gain de Traitement est faible (codes de longueurs plus courtes).

IX.2.1.5 Cas du canal idéal, 3+20 utilisateurs, **Gold (11,2)**.

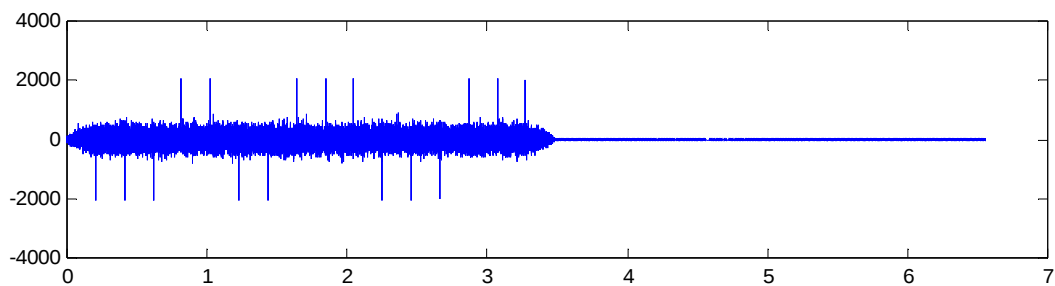


Figure 9.6 : canal idéal, 3+20 utilisateurs, **Gold (11,2)**

Sur cette figure, l'axe de temps est à multiplier par 10^4 . On voit ici que lorsque le Gain de traitement augmente (utilisation de codes plus longs), plusieurs utilisateurs supplémentaires peuvent être admis dans le système.

IX.2.1.6 Cas du canal bruité (SNR= 10dB), 3 utilisateurs, **Gold (8,2)**.

On voit dans ce cas, une légère dégradation de la qualité du signal à la sortie. Cette dégradation est presque invisible sur notre figure quand on la compare à la figure 9.3.

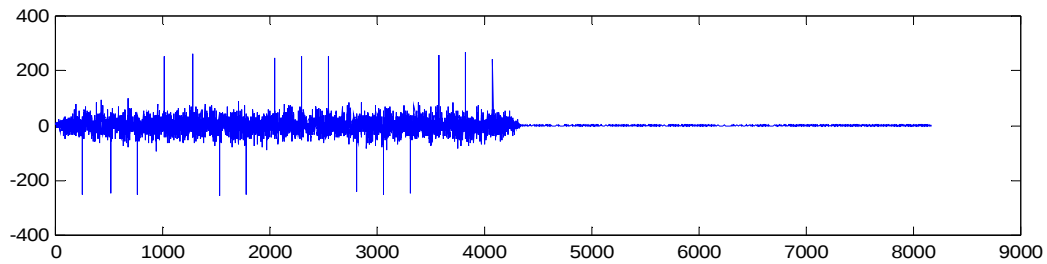


Figure 9.7: canal bruité (SNR=10dB), 3 utilisateurs, **Gold (8,2)**

IX.2.1.7 Cas du canal bruité (SNR= 0dB), 3 utilisateurs, **Gold (8,2)**.

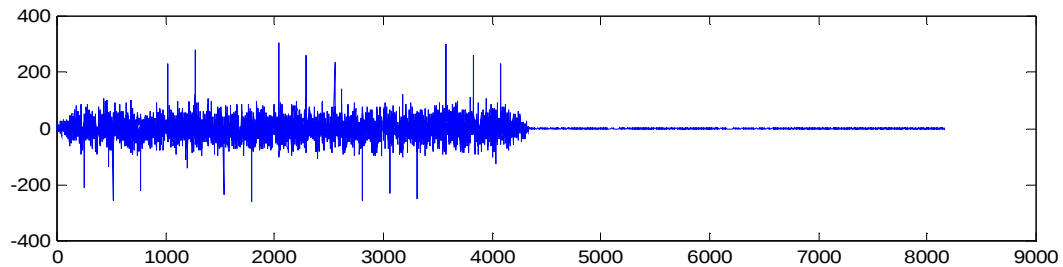


Figure 9.8 : canal bruité (SNR=0dB), 3 utilisateurs, **Gold (8,2)**

La puissance du bruit a été augmentée. On assiste ici à un fait nouveau : les données continuent d'être reçues correctement alors que le signal et le bruit sont de même puissances.

Lorsque le signal est étalé à l'émission, sa puissance se trouve étalée sur une bande de fréquence plus grande, ce qui implique une réduction du rapport signal sur bruit dans le canal. A la réception, lors de l'opération de dé-étalement, l'opération inverse est effectuée ce qui fait que le signal sur bruit se trouve amélioré artificiellement lors du passage vers une bande étroite.

IX.2.1.8 Cas du canal bruité (SNR= -4dB), 3 utilisateurs, **Gold (8,2)**.

Si on envisage de faire augmenter la puissance du bruit au delà de celle du signal utile, on assiste à une dégradation du signal reçu telle qu'on ne pourra plus reconstituer correctement les données émises par l'utilisateur 2. Le gain de traitement n'est plus d'aucune aide si celui-ci est assez faible.

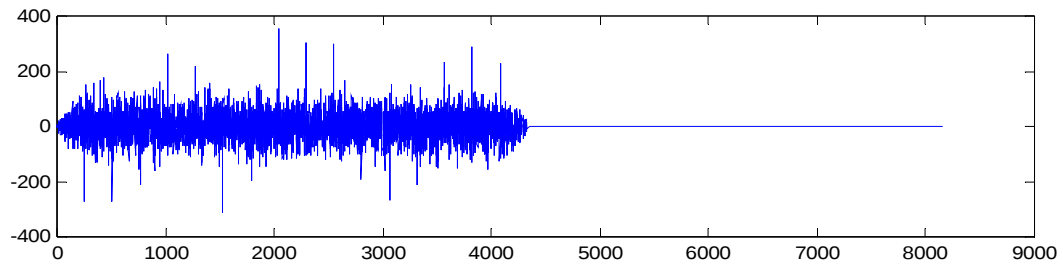


Figure 9.9 : canal bruité (SNR= -4dB), 3 utilisateurs, **Gold (8,2)**

IX.2.1.9 Cas du canal bruité (SNR= -4dB), 3 utilisateurs, **Gold (11,2)**.

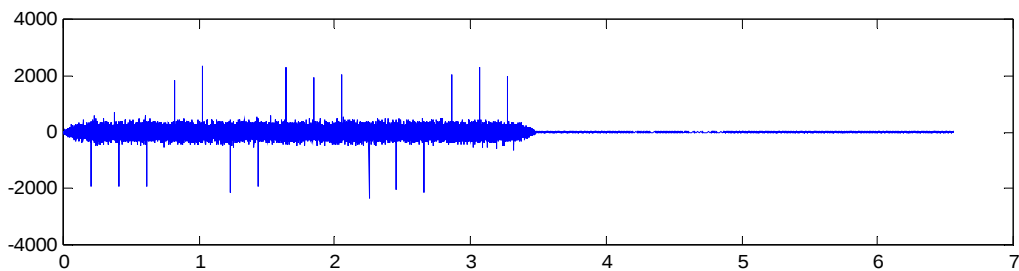


Figure 9.10 : canal bruité (SNR= -4dB), 3 utilisateurs, **Gold (11,2)**

Sur cette figure, l'axe de temps est à multiplier par 10^4 . Contrairement au cas précédent, on voit ici que la communication peut encore se faire même si le bruit est de 4 dB au dessus du signal utile lorsque le gain de traitement est plus élevé. Ce qui fournit une faculté de résistance plus grande au système face au bruit.

IX.2.1.10 Cas du canal avec interférence à bande étroite (SIR= 10dB), 3 utilisateurs, **Gold (8,2)**.

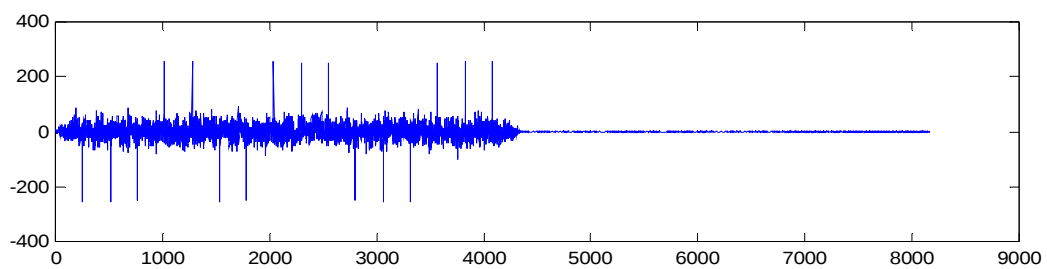


Figure 9.11 : interférence (SIR= 10dB), 3 utilisateurs, **Gold (8,2)**

Une interférence à bande étroite n'a presque aucune influence sur un système lorsque la puissance de celle-ci est faible, car dans ce cas l'interférence ne modifie pas la propriété de corrélation du signal émis.

IX.2.1.11 Cas du canal avec interférence à bande étroite (SIR= 0dB), 3 utilisateurs, **Gold (8,2)**.

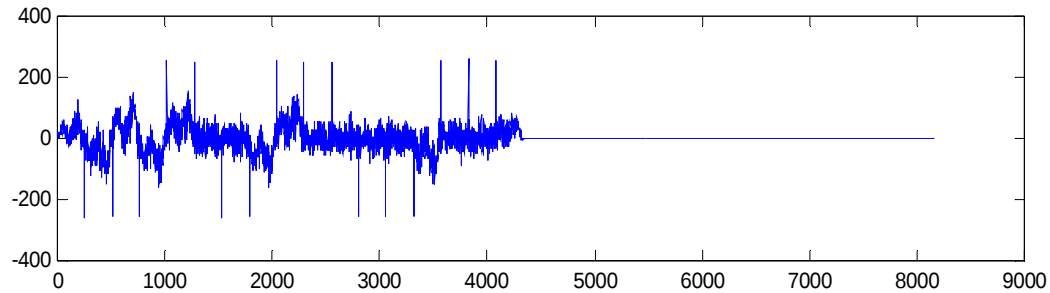


Figure 9.12 : interférence (SIR= 0dB), 3 utilisateurs, **Gold (8,2)**

L'augmentation de la puissance de l'interférence entraîne une déformation du signal reçu qui peut être gênant lorsque ajoutée aux effets du bruit. Cette déformation significative est due au masquage du signal utile par l'interférence lorsque la longueur du code utilisé n'est pas assez longue. Dans le cas de l'interférence à bande étroite, une série toute entière de chips se trouve alors superposée à une même valeur de l'interférence.

IX.2.1.12 Cas du canal avec interférence à bande étroite (SIR= 0dB), 3 utilisateurs, **Gold (11,2)**.

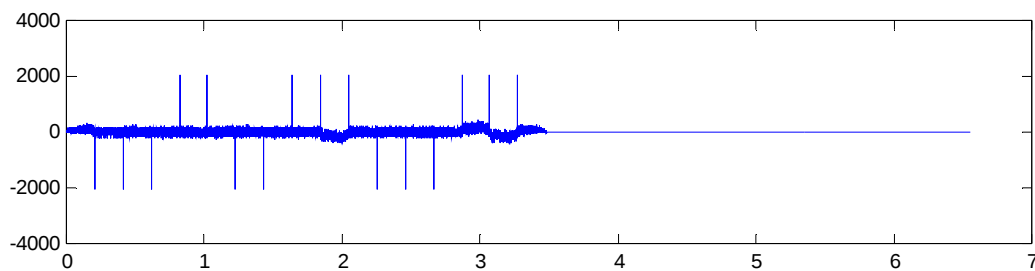


Figure 9.13 : interférence (SIR= 0dB), 3 utilisateurs, **Gold (11,2)**

Sur cette figure, l'axe de temps est à multiplier par 10^4 . On voit ici que la déformation du signal introduite par l'interférence est atténuée par l'augmentation du gain de traitement lors de l'utilisation d'un code plus long.

IX.2.1.13 Cas du canal avec trajets multiples (Amplitudes relatives = 0.7 0.6 0.5), 3 utilisateurs, **Gold (11,2)**.

Les retards des composantes ayant emprunté les chemins via réflexions sont respectivement de 500 chips, 1000 chips et de 1500 chips.

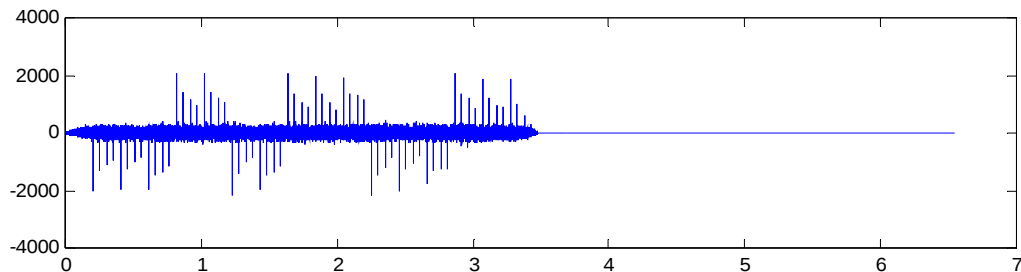


Figure 9.14 : trajets multiples, 4 chemins, **Gold (11,2)**

Ici on voit chaque composante du signal reçu par le récepteur correspond à un pic à la sortie du Matched Filter. Ceci représente un avantage considérable, car on maintenant le choix de combiner les signaux provenant de différents chemins (récepteur RAKE) pour avoir une diversité temporelle, ou tout simplement de ne garder que le signal correspondant au trajet direct en ne sélectionnant que les pics de longueur plus grande qu'un seuil donné.

Remarquons aussi que contrairement aux systèmes conventionnels, les propagations par trajets multiples n'aboutissent pas à une interférence inter-symboles mais à des pics discrets qu'on peut séparer élégamment ou combiner selon le principe d'un récepteur RAKE.

IX.2.2 Conclusion

De cette première partie de la simulation, on peut conclure que le DS-CDMA peut être utilisé comme méthode d'accès multiple pour partager l'accès à un support commun de communication par le biais des codes. De plus, un système DS-CDMA est pourvu de capacités de résistance contre le bruit blanc, l'interférence à bande étroite lorsque la longueur des codes utilisés est assez longue. Cette longueur de code détermine aussi l'aptitude du système à admettre en son sein un plus grand nombre d'utilisateurs.

Comme règle générale donc, on peut dire que plus les codes utilisés sont longs, plus le système est apte à accepter une valeur élevée de l'interférence, et de la puissance du bruit.

De cette première partie, on peut aussi voir qu'il est possible (dans certaines conditions) de réutiliser les bandes de fréquence déjà utilisées pour un système à bande étroites par un système DS-CDMA sans que les deux systèmes ne se perturbent mutuellement.

IX.3 Un exemple de programme d'aide à l'optimisation de la capacité d'un réseau DS-CDMA pluricellulaire chargé.

Ce programme est basé sur une simulation d'un réseau DS-CDMA multicellulaire. Le comportement d'un tel système est régi par les interactions entre les cellules qui le composent. Dans le système simulé, les utilisateurs sont rattachés à la station de base de leurs cellules respectives et ils ne sont pas mobiles (c'est-à-dire qu'il n'y a pas de Handover). De plus, on n'utilise pas de sectorisation, ce qui fait les éventuelles sommations ou intégrations se feront sur toute l'étendue du réseau.

La capacité N d'une cellule dans un réseau DS-CDMA doit satisfaire à l'expression déduite de la formule (6.29) :

$$\frac{E_b}{N_0} = \frac{W/R}{\sum_{i=1}^{N-1} X_i + (I/S) + (\eta/S)} \quad (9.1)$$

où N_0 représente l'ensemble (interférence + bruit).

Pour toute la suite, adoptons les modifications de notation suivantes (pour des raisons de convenance) :

- le rapport Energie par bit sur la densité de l'interférence totale (y compris le bruit) sera noté par $\frac{E_b}{I_0}$
- le rapport Energie par bit sur la densité du bruit (seulement le bruit) sera noté par $\frac{E_b}{N_0}$
- le nombre d'utilisateurs dans la cellule n° i est noté par n_i

Les autres notations utilisées sont conformes à celles utilisées dans le Chapitre VI.

La réécriture de l'équation 9.1 suivant la nouvelle notation donne :

$$\frac{E_b}{I_0} = \frac{W/R}{\sum_{i=1}^{n_i-1} X_i + (I/S) + (\eta/S)} \quad (9.2)$$

Pour avoir un taux d'erreur binaire donné, on doit avoir $\frac{E_b}{I_0} \geq \Gamma$ où Γ représente un seuil, en dessous duquel, le taux d'erreur binaire que l'on s'est fixé ne pourra être atteint (ce nombre dépend des caractéristiques du récepteur utilisé).

Par passage à la moyenne, le terme $\sum_{i=1}^{n_i-1} X_i$ de l'équation (9.2) sera remplacé par $(n_i - 1)\alpha$ (voir Chapitre VI). Ce qui donne

$$\frac{E_b}{I_0} = \frac{W/R}{(n_i - 1)\alpha + E[I/S] + (\eta/S)} \quad (9.3)$$

où I/S est donné par la formule (6.25) pour lequel la fonction $\Phi = 1$ car le mobile est toujours rattaché à la station de base de la cellule où il se trouve.

η représente l'énergie du bruit thermique

α représente l'activité vocale.

Après quelques manipulations, et en supposant qu'il y a M cellules dans le réseau, le rapport Energie par bit sur la densité d'interférence dans la cellule i est donné par :

$$\left(\frac{E_b}{I_0} \right)_i = \frac{E_b}{\alpha R E_b \left(n_i - 1 + \sum_{j=1, j \neq i}^M n_j K_{ji} \right) / W + N_0} \quad (9.4)$$

pour $i=1, \dots, M$.

où

$$K_{ji} = \frac{e^{(\gamma\sigma)^2}}{n_j} \iint_{C_j} \left(\frac{r_j(x, y)}{r_i(x, y)} \right)^4 \rho(x, y) dA(x, y) \quad (9.5)$$

avec $\gamma = \ln(10)/10$

σ est l'écart type de la variable log-normale utilisée pour modéliser l'effet de masque combiné à l'imperfection de la boucle de contrôle de puissance (voir Chapitre VI).

ρ est la densité relative d'utilisateurs au point (x, y)

r_j représente la distance entre l'utilisateur se trouvant aux coordonnées (x, y) et la station de base j .

En réécrivant l'équation 9.4 et en tenant compte de Γ , le nombre d'utilisateurs dans chaque cellule doit satisfaire

$$n_i + \sum_{j=1, j \neq i}^M n_j K_{ji} \leq \frac{W/R}{\alpha} \left(\frac{1}{\Gamma} - \frac{N_0}{E_b} \right) + 1 \equiv c_{\text{eff}} \quad (9.6)$$

pour $i = 1, \dots, M$.

IX.3.1 Définition des capacités égales [33]

C'est la capacité d'une cellule lorsqu'on suppose que les autres cellules aient la même capacité ($n_1 = n_2 = n_3 = \dots = n_i = n$). La valeur prise par cette capacité sera la plus petite des capacités obtenues après le calcul. Dans ce cas, la capacité du réseau est égale à Mn , où

$$n = \min_{1 \leq i \leq M} \left[\frac{c_{\text{eff}}}{1 + \sum_{j=1, j \neq i}^M K_{ji}} \right] \quad (9.7)$$

C'est une hypothèse simplificatrice qui permet de calculer simplement et rapidement la capacité d'une cellule dans un réseau. Cependant, ceci ne permet pas le calcul de la capacité réelle du réseau.

IX.3.2 Définition de la capacité optimale d'un réseau DS-CDMA [33]

La capacité optimale d'un réseau est donnée par la solution du problème d'optimisation suivant :

$$\begin{aligned} & \max_n \sum_{i=1}^M n_i \\ & \text{avec } n_i + \sum_{j=1, j \neq i}^M n_j K_{ji} \leq c_{\text{eff}} \text{ pour } i = 1, \dots, M. \end{aligned} \quad (9.8)$$

IX.3.3 Le facteur de compensation de puissance (PCF : Power Compensation Factor) [33]

C'est le facteur par lequel la puissance d'émission du mobile sera multipliée pour effectuer la correction de celle-ci. A chaque cellule j correspond un facteur spécifique β_j .

On a alors

$$K_{ji} = \frac{e^{(\alpha)^2}}{n_j} \iint_{C_j} \left(\frac{r_j(x,y)}{r_i(x,y)} \right)^4 \beta_j \rho(x,y) dA(x,y) \quad (9.9)$$

Le nombre d'utilisateurs dans chaque cellule doit alors satisfaire :

$$n_i + \sum_{j=1, j \neq i}^M n_j K_{ji} \frac{\beta_j}{\beta_i} \leq \frac{W/R}{\alpha} \left(\frac{1}{\Gamma} - \frac{N_0}{\beta_i E_b} \right) + 1 \equiv c_{\text{eff}}^{(i)} \quad (9.10)$$

Le cas des capacités égales devient alors :

$$n = \min_{1 \leq i \leq M} \left[\frac{c_{\text{eff}}^{(i)}}{\beta_i + \sum_{j=1, j \neq i}^M \beta_j K_{ji}} \right] \quad (9.11)$$

tandis que la capacité optimale du réseau est donnée par la solution du problème d'optimisation :

$$\max_{n, \beta} \sum_{i=1}^M n_i$$

avec $1 \leq \beta \leq \beta^{\max}$

$$\text{et } n_i + \sum_{j=1, j \neq i}^M n_j K_{ji} \frac{\beta_j}{\beta_i} \leq c_{\text{eff}}^{(i)} \text{ pour } i = 1, \dots, M. \quad (9.12)$$

où $\beta^{\max}$ doit être choisie de telle façon que la puissance d'émission du mobile requise n'excède pas la puissance maximale que le mobile peut fournir.

IX.3.4 Les paramètres utilisés par le programme de simulation

L'aspect de l'interface graphique du programme de simulation est donné par 9.17.

Le programme de simulation nécessite l'acquisition des paramètres suivants pour fonctionner, les valeurs entre parenthèses représentent celles utilisées dans notre exemple :

RESULTATS PAR CELLULE										PARAMETRES			
CELLULE N°										RAYON D'UNE CELLULE en m		GAMMA en dB	
1	17	10	20	19	6	1732		9.2					
2	16	11	18	20	24	ACTIVITE VOCALE		0.375					
3	10	12	18	21	22	CENTRE DU HOT SPOT en m		[-4500 2598]					
4	13	13	11	22	21	RAYON DU HOT SPOT en m		3000					
5	4	14	4	23	25	CAPACITE MINIMALE D'UNE CELLULE		4					
6	4	15	5	24	4	PROFILS DE CALCUL							
7	11	16	4	25	4	☐ Capacités égales							
8	14	17	24	26	23	☐ Repartition Optimale							
9	6	18	18	27	24	☐ Hot Spot, Capacités égales							
AFFICHAGE DES RESULTATS						CAPACITE TOTALE DU RESEAU		☐ Hot Spot, Repartition Optimale, sans contraintes					
☒ Nombre d'utilisateurs						370		☐ Hot Spot, Repartition Optimale, avec contraintes					
☐ Coefficients de Correction de Puissance (PCF)						CALCULER		☒ Hot Spot, Repartition Optimale, sans contraintes, avec PCF					
ANDRY JEAN ONESIME Copyright © Août 2002								☐ Hot Spot, Repartition Optimale, avec contraintes, avec PCF					

Figure 9.15 : l'interface graphique utilisée par le programme de simulation

- le rayon d'une cellule (1732 m)
- le seuil Γ (9.2 dB [33])
- le rapport $\frac{I_0}{N_0}$ qui représente entre la densité d'interférence et la densité de bruit (10dB [33])
- l'écart-type σ de la variable log-normale utilisée pour modéliser l'imperfection de la boucle de contrôle de puissance, combinée à l'effet de masque (6 dB [33]).
- le gain de traitement utilisé (c'est-à-dire le facteur d'étalement utilisé) (256, cas de l'UMTS)
- les coordonnées du centre du Hot Spot simulé (-4500m 2598m)
- le rayon du Hot Spot (car celui est de forme circulaire dans notre programme de simulation) (3000m)

- la densité relative d'utilisateurs dans le Hot Spot (par rapport au reste du réseau)
- une contrainte sur la capacité minimale d'une cellule

La numérotation des cellules est faite en spirale en partant du centre.

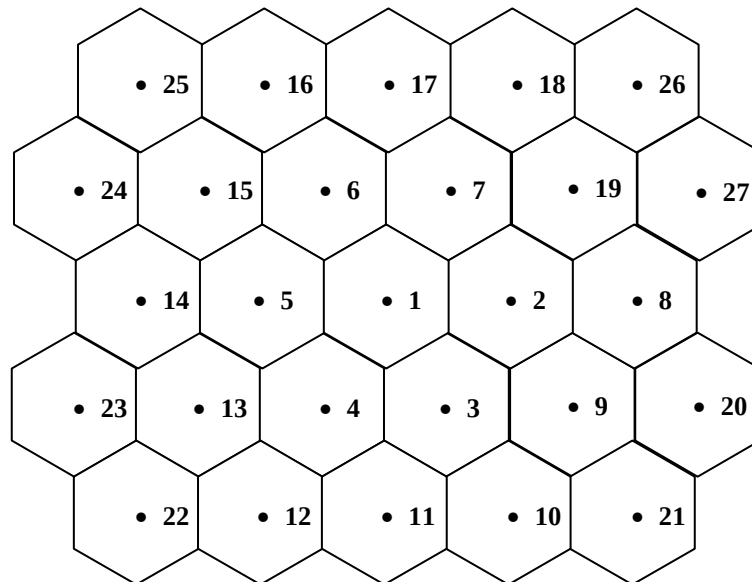


Figure 9.16 : plan de numérotation des cellules

Le centre de la cellule 1 est pris comme origine des coordonnées, ainsi le centre du Hot Spot coïncide alors avec le centre de la cellule 15.

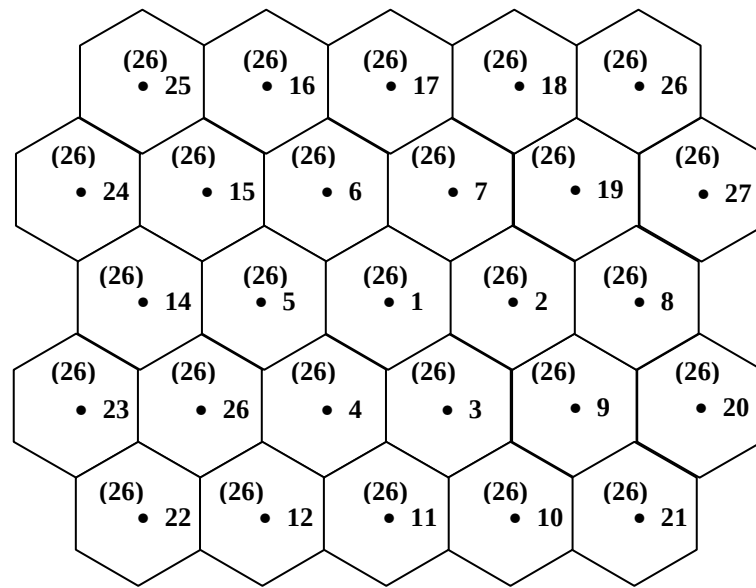
En retour, le programme donnera pour chaque profil de calculs la capacité de chaque cellule, la capacité du réseau tout entier ainsi que le facteur de compensation de puissance utilisé dans chaque cellule lorsque cette option est choisie.

IX.3.5 Calculs

IX.3.5.1 Capacités égales, distribution uniforme des utilisateurs

Le résultat de calcul pour ce cas est donné par la figure 9.17.

On remarque que les capacités de toutes les cellules du réseau sont égales. La capacité totale du réseau est donc égale à 702, la capacité de chaque cellule étant égale à 26.



- numéro de la station de base
- () : capacité de la cellule

Figure 9.17 : cas des capacités égales, sans Hot Spot

IX.3.5.2 Répartition optimale, distribution uniforme des utilisateurs

Les capacités des cellules qui sont sur les bords de la couverture du réseau ont significativement augmenté. Ceci est dû au fait que l'interférence intercellulaire pour ces cellules est plus petite que pour celles qui se trouvent à l'intérieur du réseau.

Cependant, cette augmentation de capacité provoque l'augmentation de l'interférence dans les cellules adjacentes qui ont vu leurs capacités diminuer. Mais quand même, la capacité totale du réseau a pu passer de 702 à 885.

La connaissance en avance de la capacité optimale ainsi que la répartition optimale peut aider le programme de contrôle d'admission d'appel utilisé par le système, dans la mesure où celui-ci ne doit plus accepter d'appels supplémentaires lorsque le quota prévu pour la cellule a été atteint pour le blocage de toute la cellule.

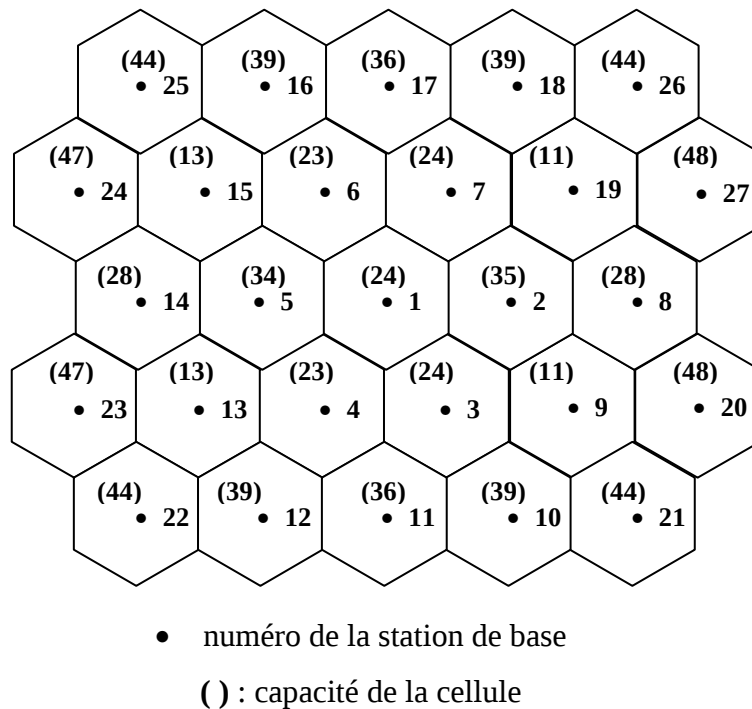


Figure 9.18 : répartition optimale, sans Hot Spot

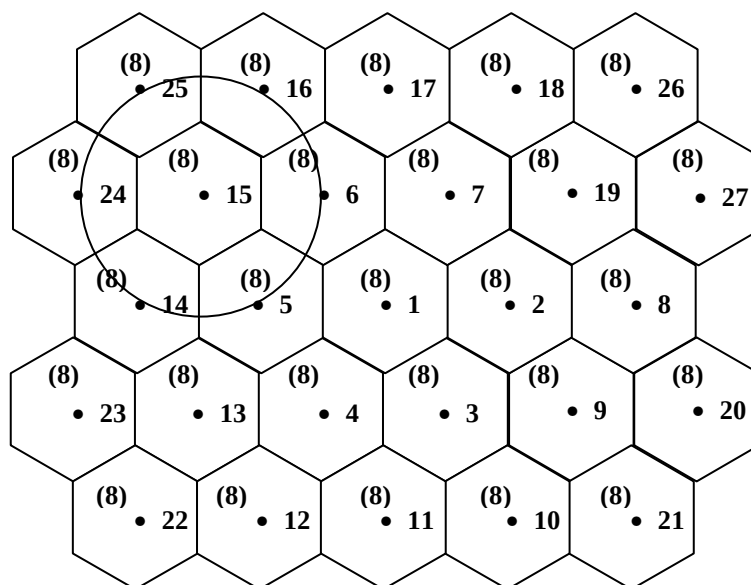
IX.3.5.3 Capacités égales, distribution non uniforme des utilisateurs (figure 9.19).

Cette section sert à simuler l'effet de la densité élevée d'utilisateurs à l'intérieur d'une zone appelée Hot Spot sur les capacités de chaque cellule du réseau. Le Hot Spot en question est une zone circulaire centrée sur la station de base de la cellule 15. Dans cette zone, la densité des utilisateurs est de 5 fois la densité à l'extérieur du Hot Spot. La capacité d'une cellule est de 8. La capacité du réseau est alors égale à 216

IX.3.5.4 Répartition optimale, distribution non uniforme des utilisateurs (figure 9.20).

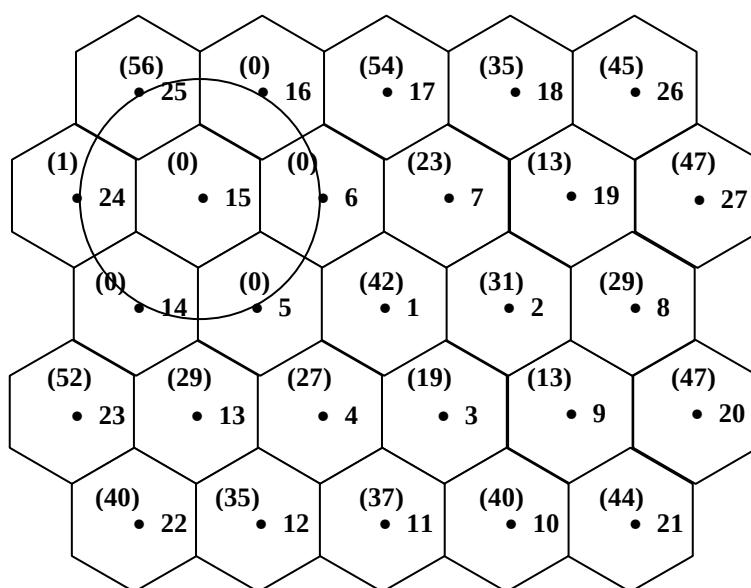
On remarque que les cellules proches du Hot Spot ont vu leurs capacités se réduire d'une manière drastique. Ceci est dû au fait que l'interférence à l'intérieur du Hot Spot est tellement élevée qu'aucune communication ne puisse avoir lieu. D'autre part, le programme cherche à maximiser la capacité du réseau entier, et lorsque les capacités des autres cellules assez proches du Hot Spot augmentent, elles créent des interférences supplémentaires dans le Hot Spot ce qui réduit encore la capacité de ce dernier.

La capacité totale du réseau est maintenant égale à 759



- numéro de la station de base
- () : capacité de la cellule

Figure 9.19: capacités égales, avec un Hot Spot



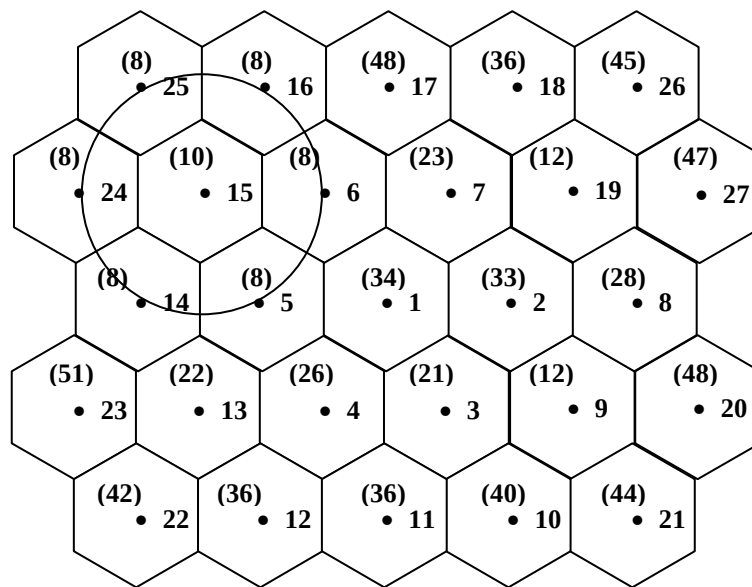
- numéro de la station de base
- () : capacité de la cellule

Figure 9.20 : répartition optimale, avec un Hot Spot

IX.3.5.5 Répartition optimale, distribution non uniforme des utilisateurs, avec une contrainte de 8 utilisateurs par cellule au minimum.

Pour éviter des cellules vides, une astuce consiste à imposer des contraintes sur la capacité minimale d'une cellule, c'est-à-dire à imposer des restrictions sur les cellules adjacentes au Hot Spot pour que les cellules à l'intérieur de celui-ci ont une capacité au moins égale à celle obtenue avec la méthode des capacités égales (8 pour notre cas).

La capacité totale du réseau passe alors de 759 à 742.



- numéro de la station de base
- () : capacité de la cellule

Figure 9.21 : répartition optimale, avec un Hot Spot,
la capacité minimale est égale à 8

IX.3.5.6 Répartition optimale, distribution non uniforme des utilisateurs, utilisation de la compensation de puissance

La méthode employée précédemment entraîne une réduction de la capacité du réseau. Un autre moyen d'augmenter la capacité du réseau consiste à ajuster la puissance d'émission des mobiles surtout ceux qui se trouvent à l'intérieur du Hot Spot. Les résultats obtenus sont présentés sur la figure 9.22.

La capacité totale du réseau passe alors à 802. Cependant, on remarque que les cellules 5, 6, 14, 15 sont toujours vides. Ceci veut dire que les techniques que nous avons déployées jusqu'à présent ne permettent pas de rehausser les capacités de ces cellules tout en optimisant la capacité du réseau tout entier. Il faut alors envisager des solutions sous optimales comme l'utilisation des contraintes sur la capacité d'une cellule.

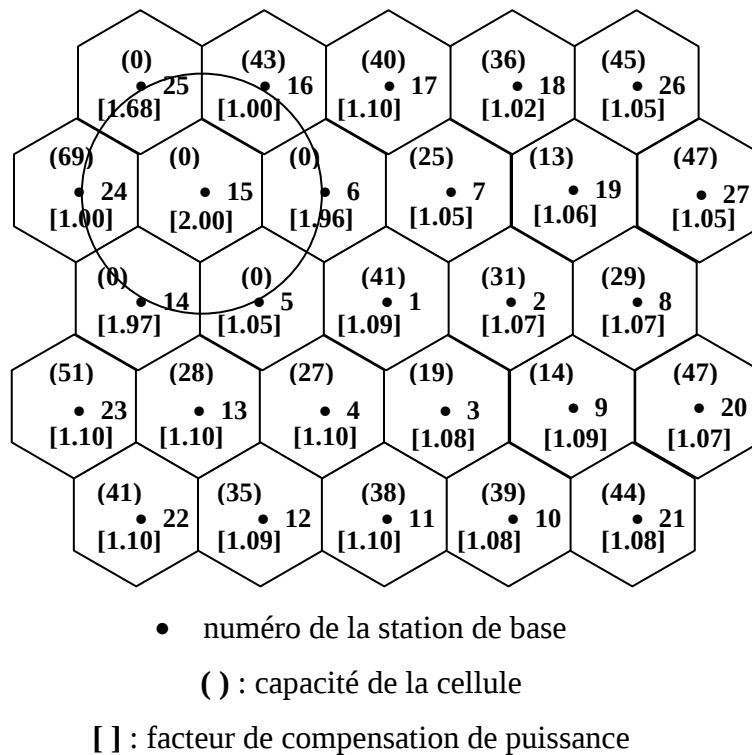
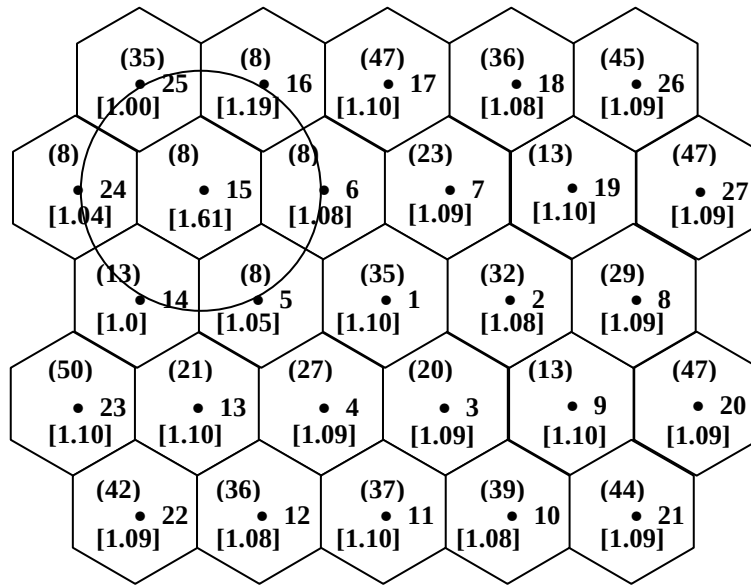


Figure 9.22 : répartition optimale, avec un Hot Spot, utilisation des compensations de puissance

IX.3.5.7 Répartition optimale, distribution non uniforme des utilisateurs, utilisation de la compensation de puissance, avec une contrainte de 8 utilisateurs par cellule au minimum.

Cette méthode consiste à imposer des restrictions sur les capacités des cellules voisines du Hot Spot, pour que les cellules à l'intérieur de celui-ci aient une capacité minimale de 8 utilisateurs par cellule. Cependant, ici on rehausse la puissance des mobiles à l'intérieur du Hot Spot, ce qui a pour effet d'augmenter la capacité totale du réseau par rapport au cas où aucune compensation n'a été faite. La capacité passe alors de 742 (c'est le cas de la figure 9.21) à 771 (figure 9.23).



• numéro de la station de base

() : capacité de la cellule

[] : facteur de compensation de puissance

Figure 9.23 : répartition optimale, avec un Hot Spot, utilisation des compensations de puissance, la capacité minimale est égale à 8.

IX.3.6 Conclusion

Les programmes de simulation peuvent aider lors de la conception et de l'optimisation de la capacité d'un réseau. Ils peuvent éviter les désagréments provoqués par l'inefficacité du système lors de son implantation. Le cas de la figure 9.22 est un exemple type. Bien avant l'implantation réelle du système, on voit déjà à travers la simulation que la compensation des puissances des mobiles ne permet pas de rehausser les capacités des cellules 5, 6, 14, 15. D'autres alternatives doivent alors être explorées comme la modification des rayons des cellules, le déplacement des stations de base ou encore la mise en place de plusieurs couches de cellules avant d'envisager l'installation du système.

Remarquons aussi le programme tourne pendant plusieurs minutes avant de retourner le résultat. Ceci est dû au fait qu'on a à calculer (27×27) intégrales doubles régissant les interactions entre les cellules, avant de faire le calcul d'optimisation proprement dit.

CONCLUSION

Les fondements théoriques d'un système DS-CDMA ont été analysés. Un système DS-CDMA utilise des codes pour séparer les utilisateurs qui partagent la même bande de fréquence. Les performances d'un tel système dépend donc du choix du ou des codes utilisés. Les codes d'étalement les plus utilisés dans un système DS-CDMA ont été étudiés, et classés selon leurs types. Les propriétés de chaque code permettent donc de les suggérer pour des usages spécifiques dans un système DS-CDMA. De plus, on a montré que l'enchaînement des codes orthogonaux avec des codes PN permettent d'avoir de meilleures performances en termes de capacité du réseau dans un environnement soumis à des fading de Rayleigh et de Rice.

Ensuite, des architectures de base utilisées pour l'émission et plus particulièrement pour la réception des signaux à spectre étalé ont été proposées. Pour pouvoir reproduire correctement les données émises, le récepteur doit comporter un dispositif qui assure en permanence la synchronisation entre le code utilisé à l'émission et celui utilisé pour la réception. Etant donné son importance dans un système DS-CDMA, le processus de synchronisation a été abordé ainsi que les structures les plus couramment utilisées pour effectuer cette tâche.

Les performances d'un système DS-CDMA en matière de résistance aux effets des propagations par trajets multiples ainsi qu'aux interférences à bande étroite ont été analysées. On a montré que dans le cas d'une synchronisation parfaite, les effets précédemment décrits peuvent être annihilés. Ensuite la capacité d'un système DS-CDMA monocellulaire et pluricellulaire a été étudiée, et des comparaisons avec les systèmes conventionnels ont été effectuées pour mettre en relief la supériorité d'un système DS-CDMA par rapport aux systèmes conventionnels utilisés actuellement. On a constaté qu'un système DS-CDMA peut accepter un plus grand nombre d'utilisateurs (108 utilisateurs par cellule contre 7 dans un système GSM pour la configuration pluricellulaire considérée). Ceci peut être atteint à cause de la réutilisation de la même bande de fréquence dans tout le réseau, qui élimine aussi entre autres le besoin de planification des fréquences (incontournable pour les systèmes FDMA et TDMA actuels). De plus, on a montré que la capacité d'un système DS-CDMA n'est limitée que par le niveau de l'interférence dans la cellule. Ainsi, toute réduction faite sur l'interférence (ou sur son effet) se traduit directement par une augmentation de la capacité.

Cependant, la réutilisation de la même bande de fréquence n'apporte que des avantages. Des problèmes inhérents à un système DS-CDMA (effets Near-Far, interférence due à l'accès multiple) doivent être alors surmontés. L'utilisation d'une boucle fermée de contrôle de la puissance du mobile peut combattre l'effet Near-Far, tandis que l'emploi de nouvelles structures de réception (détection multi-utilisateur) peut atténuer considérablement l'effet de l'interférence due à l'accès multiple.

La seconde partie de ce mémoire a été consacrée à l'utilisation du DS-CDMA comme méthode d'accès multiple au sein d'un système de téléphonie mobile. De nouvelles notions ont été introduites ainsi que les paramètres sur lesquels on peut jouer pour optimiser la capacité d'un réseau de téléphonie mobile. Ensuite, l'utilisation du DS-CDMA dans sa version à large bande WCDMA comme interface radio du système de téléphonie mobile de la troisième génération UMTS a été explorée. Les principes adoptés dans ce système pour offrir aux abonnés un accès radio efficace, ainsi que la possibilité des transferts de données à haut débit ont été passées en revue.

Les résultats de simulations portant sur un système DS-CDMA à plusieurs utilisateurs ainsi que sur un réseau DS-CDMA pluricellulaire ont été présentés dans la troisième partie de ce mémoire.

De nos jours, l'étalement de spectre par séquence directe est utilisé dans de nombreux domaines (à part les systèmes de téléphonie mobile) comme les systèmes de positionnement par satellite ou GPS (Global Positioning System), les systèmes de réseaux locaux sans fil (IEEE 802.11 et IEEE 802.11bis) et les systèmes de guidage des missiles pour n'en citer que quelques-uns.

Remarquons que les codes BARKER n'ont pas été abordés car de par leurs longueurs et nombres réduits, ils ne peuvent pas être utilisés dans un système DS-CDMA.

Annexe 1

A1.1. Table donnant les prises de retour à utiliser en vue de générer des M-Séquences

Dans la table suivante, on peut trouver les prises de retour (feed-back taps) pour les séquences générées avec un SSRG linéaire (sans les ensembles images correspondant à un registre inverse) ainsi que les nombres de séquences générées pour chaque longueur L du registre.

L	$N_c=2^L-1$	Prises de retour pour les M-séquences	Nombre de M-séquences
2	3	[2,1]	2
3	7	[3,1]	2
4	15	[4,1]	2
5	31	[5,3] [5,4,3,2] [5,4,2,1]	6
6	63	[6,1] [6,5,2,1] [6,5,3,2]	6
7	127	[7,1] [7,3] [7,3,2,1] [7,4,3,2] [7,6,4,2] [7,6,3,1] [7,6,5,2] [7,6,5,4,2,1] [7,5,4,3,2,1]s	18
8	255	[8,4,3,2] [8,6,5,3] [8,6,5,2] [8,5,3,1] [8,6,5,1] [8,7,6,1] [8,7,6,5,2,1] [8,6,4,3,2,1]	16
9	511	[9,4] [9,6,4,3] [9,8,5,4] [9,8,4,1] [9,5,3,2] [9,8,6,5] [9,8,7,2] [9,6,5,4,2,1] [9,7,6,4,3,1] [9,8,7,6,5,3]	48
10	1023	[10,3] [10,8,3,2] [10,4,3,1] [10,8,5,1] [10,8,5,4] [10,9,4,1] [10,8,4,3] [10,5,3,2] [10,5,2,1] [10,9,4,2] [10,6,5,3,2,1] [10,9,8,6,3,2] [10,9,7,6,4,1] [10,7,6,4,2,1] [10,9,8,7,6,5,4,3] [10,8,7,6,5,4,3,1]	60
11	2047	[11,2] [11,8,5,2] [11,7,3,2] [11,5,3,2] [11,10,3,2] [11,6,5,1] [11,5,3,1] [11,9,4,1] [11,8,6,2] [11,9,8,3] [11,10,9,8,3,1]	176

Tableau A1.1 : les prises de retour à utiliser pour générer des M-séquences

A1.2 Table donnant les paires préférées de M-Séquences en vue de générer des codes GOLD

Cette table donne les paires préférées de M-Séquences en vue d'une génération de codes Gold. Elle donne également les valeurs normalisées des valeurs d'intercorrélation (qui ne prennent que trois valeurs dans ce cas).

L	$N_c=2^L-1$	Paires préférées de M-séquences	Valeurs normalisées de l'inter-corrélation			Limite
5	31	[5,3] [5,4,3,2]	7	-1	-9	-29%
6	63	[6,1] [6,5,2,1]	15	-1	-17	-27%
7	127	[7,3] [7,3,2,1] [7,3,2,1] [7,5,4,3,2,1]	15	-1	-17	-13%
8*	255	[8,7,6,5,2,1] [8,7,6,1]	31	-1	-17	+12%
9	511	[9,4] [9,6,4,3] [9,6,4,3] [9,8,4,1]	31	-1	-33	-6%
10	1023	[10,9,8,7,6,5,4,3] [10,9,7,6,4,1] [10,8,7,6,5,4,3,1] [10,9,7,6,4,1] [10,8,5,1] [10,7,6,4,2,1]	63	-1	-65	-6%
11	2047	[11,2] [11,8,5,2] [11,8,5,2] [11,10,3,2]	63	-1	-65	-3%

Tableau A1.2 : les paires préférées de M-séquences pour générer des codes de Gold

Annexe 2

Les systèmes de téléphonie mobile de la première et de la seconde génération

A2.1 Les systèmes de la première génération

Les premiers systèmes de téléphonie mobile étaient analogiques. Même si des embryons avaient déjà vu le jour dès 1946 [2], ce n'est qu'à partir des années 1960 que les premiers véritables systèmes de téléphonie mobile ont été mis en service aux Etats-Unis avec l'avènement du système IMTS (Improved Mobile Telephone System). Ce système comporte 23 canaux étalés de 150 à 450 MHz. Outre le faible nombre des canaux disponibles, la puissance d'émission est assez puissante (200 W) empêchant la mise en place d'autres systèmes similaires à proximité pour améliorer l'efficacité (essentiellement en termes de nombre de canaux disponibles).

En 1982, on assiste à la naissance du système AMPS (Advanced Mobile Phone System) aux Etats-Unis. Un nouveau concept est aussi apparu : c'est la division du territoire à couvrir en une suite de zones contiguës de forme circulaires appelées **cellules**. Dans le système AMPS les cellules ont une taille de l'ordre de 10 à 20 km de diamètre. Chaque cellule dispose de nombreux canaux de communication s'étalant sur une certaine bande de fréquences (5 à 10 appels simultanément par cellule). Les points clés de ce système pour améliorer sa capacité consistent à réduire la dimension des cellules et à ne pas réutiliser les mêmes fréquences dans les cellules adjacentes pour éviter les interférences. Les cellules de petites dimensions supposent des émetteurs de faible puissance (de l'ordre de 0.6 W) donc des équipements moins coûteux et de faible encombrement.

Au centre de chaque cellule on retrouve une station de base dénommée BSC (Base Station Controller) qui comprend l'émetteur radio et l'antenne d'émission. Une ou plusieurs BSC sont reliées à un centre de commutation du service mobile MSC (Mobile Switching Center).

Lors du passage dans une nouvelle cellule, le téléphone mobile est informé de la nouvelle BSC et, si une communication est en cours, les changements de canal de

communication et de bandes de fréquences doivent être effectués. C'est la fonction de *transfert de communication* ou *handover*. Le Handover se traduit par une coupure de la communication pendant environ 300 ms sans que celle-ci s'interrompe. L'assignation des canaux aux mobiles est traitée par le MSC. Les stations BSC, elles, ne sont en réalité dans le cadre de l'AMPS que de simples relais radio.

Le système AMPS dispose de 832 canaux duplex, chacun consistant en une paire de canaux simplex (un pour chaque sens de transmission). La bande de 824 à 849 MHz est utilisée pour l'émission alors que la bande de 869 à 894 MHz est employée pour la réception. Chaque canal simplex a une bande passante d'environ 30 kHz. Le système AMPS utilise le CDMA comme méthode d'accès multiple.

On n'a abordé ici que le système AMPS qui a été mis en place initialement aux Etats-Unis. Cependant, il y a encore plusieurs systèmes similaires et incompatibles entre eux qui ont vu le jour. On note par exemple le système R150 ouvert en 1956 en France, ainsi que ses successeurs R450 et R200. En 1985, France Telecom a mis en place le système **Radiocom 2000** qui fonctionne dans les fréquences 200,400 et 900 MHz.

Le système AMPS a permis de démocratiser le concept de "téléphonie mobile cellulaire". Cependant, il avait encore beaucoup de faiblesses en matière de capacité et surtout en matière de sécurisation des appels (l'AMPS ne fournit aucun support de chiffrement, le système étant aussi analogique). On note aussi l'existence de toute une multitude de norme incompatible entre eux, qui obligent les utilisateurs à changer de postes chaque fois qu'ils changent de pays ou de région.

A2.2 Les systèmes de la seconde génération

Comme nous l'avons dit auparavant, les systèmes de la première génération étaient analogiques. Ceux de la seconde génération sont numériques. Aux Etats-Unis, l'AMPS était le seul système analogique réellement utilisé. Lorsque le numérique est apparu, trois ou quatre systèmes différents sont entrés en compétition sur le marché pour tenter de s'imposer et de survivre. Le premier d'entre eux est compatible du point de vue de mécanisme d'allocation de fréquences et de canaux, avec l'ancien système AMPS. Les standards associés sont nommés IS-54 et IS-135. Un autre système est basé sur l'étalement de spectre à séquence directe (DS-SS) et correspond au standard IS-95.

Le système IS-54 est en double mode : analogique et numérique, il utilise les mêmes canaux de 30 kHz de largeur de bande que ceux du système AMPS. Chaque canal est partagé entre 3 utilisateurs et la transmission s'y effectue à 48 kbit/s. Chaque utilisateur dispose d'un canal à 13 kbit/s ; le reste (soit 9 kbit/s) est utilisé pour la signalisation. Les cellules, comme les stations BSC et MSC ont les mêmes spécifications que celles de l'ancien système AMPS. Seuls l'encodage et la numérisation des signaux sont différents. Les systèmes IS-54 et IS-135 utilisent donc des techniques de multiplexage mixte FDMA-TDMA.

Quant au système IS-95, il utilise le CDMA comme technique d'accès multiple. Il permet des transmissions en mode paquet qui tire profit naturellement du CDMA. Il constitue donc une étape dans la transition vers les systèmes de la troisième génération, tout en gardant une certaine compatibilité avec les systèmes existants (IS-54 et IS-135). L'IS-95 permet des transferts de données à 14,4 kbps et jusqu'à 115 kbps pour la norme IS-95B.

En Europe, c'est le processus inverse qui s'est produit. On s'acheminait vers une norme commune qu'est le GSM (Global System for Mobile communication) qui est largement adoptée en Europe aujourd'hui. Le système GSM travaille dans la bande des 900 MHz ainsi que dans la bande des 1800 MHz (comme bande d'extension). Il utilise aussi des techniques de multiplexage mixte (FDMA-TDMA), chaque porteuse pouvant transporter 8 communications multiplexées dans le temps. Le spectre de fréquence disponible est divisé en 50 bandes de 200 kHz. Pour le GSM, le transfert de données est limité à un débit maximal de 9.6 kbps.

Contrairement aux systèmes de la première génération, des moyens de chiffrement sont utilisés dans tous les systèmes de la seconde génération. Les communications deviennent ainsi très difficilement piratables. En plus de cela, des moyens d'authentifications cryptées permettent de lutter efficacement contre les usurpations d'identité (pratique très courante avec les systèmes de la première génération).

A l'exception du système IS-95, tous les systèmes de la seconde génération énumérés ici ne peuvent traiter des communications en mode paquet, ce qui le rendent inadéquats pour la transmission sporadiques de données (navigation Web, consultation d'E-Mail...). Cet aspect des choses est en grande partie imposé par le choix du FDMA et du TDMA comme techniques de multiplexage. De plus, à cause de l'utilisation de la commutation des circuits,

les débits obtenus pour les transmissions de données sont faibles, car ceux-ci sont alloués d'une manière statique et égale pour tous les utilisateurs même si ces derniers ne l'utilisent pas à ce moment précis.

A l'instar de l'IS-95, le monde GSM a adopté aussi des extensions pour faciliter la migration vers les systèmes de la troisième génération. Ce sont le GPRS (General Packet Radio Service) et l'EDGE (Enhanced Data rates for GSM Evolution). Ces extensions permettent d'ajouter de nouvelles fonctionnalités au réseau GSM existant telles les communications à commutation de paquet ainsi que des débits de données plus élevés. Le GPRS permet d'atteindre un débit de 115 kbps en mode paquet, en allouant dynamiquement les 8 IT (Intervalle de Temps ou Time Slot) de la norme GSM selon le besoin de l'abonné. L'allocation peut être différente pour la liaison montante et celle descendante. Ceci permet d'optimiser l'utilisation de la ressource radio pour les applications telles que la navigation Web, ou navigation Wap. Le GPRS utilise toujours la modulation GMSK du monde GSM. Cependant, les stations de base doivent être changées pour pouvoir manipuler les accès en mode paquet.

L'EDGE utilise une autre approche qui consiste à utiliser une modulation à plusieurs états (QAM-16) pour obtenir un débit élevé. Comme le GPRS, il offre des accès en mode paquet (allocation dynamique des IT). L'EDGE utilise les mêmes structures de multiplexage temporel que le GSM, mais actuellement on ne peut atteindre que 48 kbps pour chaque IT. Ce qui offre tout de même un débit de 384 kbps pour les 8 IT disponibles. Cependant, l'utilisation d'une modulation à plusieurs états nécessite une puissance de réception assez élevée pour atteindre un débit maximal. Ce qui implique l'installation de stations de base supplémentaires pour tirer profit pleinement de l'extension EDGE.

Annexe 3

Etalement complexe de spectre

A3.1 Introduction

Au sein d'un système à étalement de spectre, les signaux I et Q sont étalés avec une séquence PN avant la modulation. Dans les systèmes de téléphonie mobile de la troisième génération, à la place de l'étalement traditionnel de spectre, on a adopté l'étalement complexe de spectre, car en plus de la séparation des utilisateurs, il permet de travailler avec des déséquilibre de puissances sur les voies I et Q, étant donné que chaque voie est utilisée à des fins différents. Cette opération consiste à distribuer d'une manière équitable la puissance entre les axes. Ainsi, le récepteur n'a pas à tenir compte des différences de puissance entre les signaux I et Q.

A3.2 L'étalement complexe de spectre

A3.2.1 Principe de fonctionnement

La figure suivante montre le principe utilisé pour l'étalement complexe de spectre.

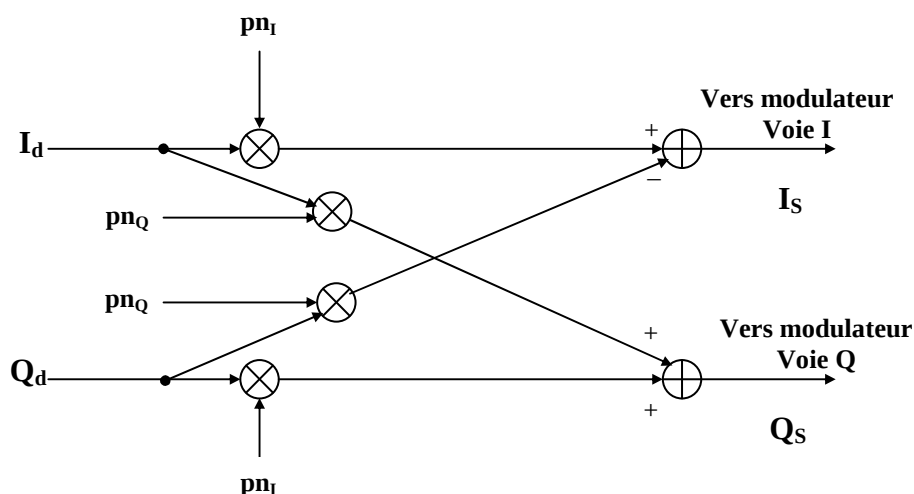


Figure A3.1: principe de l'étalement complexe de spectre

Mathématiquement, l'étalement complexe de spectre effectue la multiplication de deux signaux complexes:

- Un signal complexe de données, représenté par les signaux I_d et Q_d de la figure précédente.
- Un signal complexe de brouillage (c'est à dire une séquence PN, représentée par pn_I et pn_Q sur la figure précédente.

Les signaux I_s et Q_s peuvent s'écrire :

$$I_s = I_d \cdot pn_I - Q_d \cdot pn_Q \quad (A3.1)$$

$$Q_s = I_d \cdot pn_Q + Q_d \cdot pn_I \quad (A3.2)$$

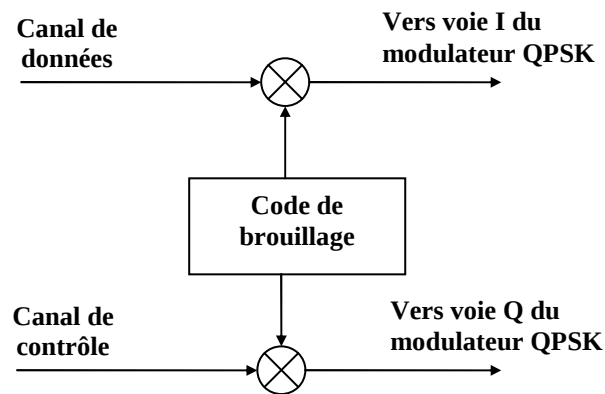
Le signal résultant, aussi bien que le signal de données et le signal de brouillage ne sont pas vraiment complexes. Cependant, ils peuvent être exprimés comme des signaux complexes :

$$\begin{aligned} I_s + jQ_s &= (I_d \cdot pn_I - Q_d \cdot pn_Q) + j(I_d \cdot pn_Q + Q_d \cdot pn_I) \\ &= (I_d + jQ_d)(pn_I + jpn_Q) \\ &= A_d \cdot A_{pn} e^{j(\Phi_d + \Phi_{pn})} \\ &= Ae^{j\Phi} \end{aligned} \quad (A3.3)$$

où A_d et $e^{j\Phi_d}$ représente l'amplitude et la phase de $I_d + jQ_d$; A_{pn} et $e^{j\Phi_{pn}}$ sont l'amplitude et la phase de $pn_I + jpn_Q$. Donc l'amplitude A du signal résultant est le produit des amplitudes des deux signaux. Sa phase Φ est la somme des phases des signaux originaux.

A3.2.1 Utilité de l'étalement complexe du spectre

Comme nous l'avons dit plus haut, les systèmes de la troisième génération utilisent les voies I et Q à des fins différents. La voie I est généralement utilisée pour transmettre les données (canal de données), tandis que la voie Q est employée pour véhiculer des informations de contrôles (canal de contrôle) nécessaires pour la bonne marche du système. Supposons qu'on utilise le QPSK comme modulation, et étudions d'abord le cas où l'on n'utilise pas d'étalement complexe de spectre. Pour un tel cas, on a la figure suivante :



:Figure A3.2 : structure n'utilisant pas l'étalement complexe de spectre

S'il y a un déséquilibre de puissance (moyenne) sur les deux canaux (données et contrôle), la constellation du signal obtenu à la sortie du modulateur QPSK est modifiée. En fait, on arrive à une constellation qui s'apparente à celle obtenue avec une modulation QAM (voir figure suivante).

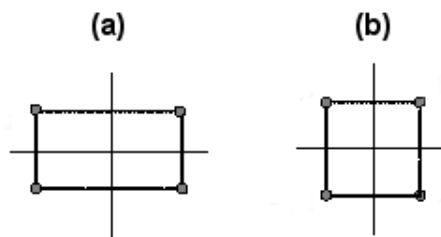


Figure A3.3 : allure de la constellation du signal obtenu à la sortie du modulateur QPSK

(a) canaux déséquilibrés; (b) canaux équilibrés

Ceci est vraiment gênant pour le récepteur, car la modification de la constellation du signal (de QPSK vers QAM) entraîne la modification simultanément de l'amplitude et de la phase du signal reçu, ce qui rend particulièrement difficile l'estimation du symbole transmis. Or on sait que les deux canaux ne peuvent être, en aucun cas, assujettis à être équilibrés (étant donné qu'ils sont utilisés à des fins différents). C'est là donc qu'intervient l'étalement complexe de spectre.

La figure suivante montre les constellations des signaux complexes $I_d + jQ_d$ et $p n_I + j p n_Q$ dans le cas d'une déséquilibre sur I_d et Q_d .

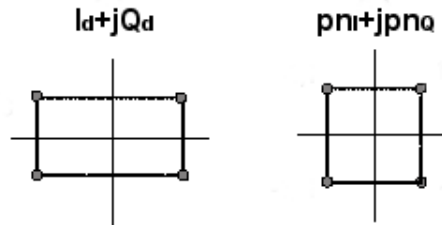


Figure A3.4 : allure des constellations des signaux $I_d + jQ_d$ et $p n_I + j p n_Q$

Si on utilise l'étalement complexe de spectre, on obtient (en appliquant (A3.3)) à la sortie du modulateur QPSK, un signal dont la constellation est de la forme

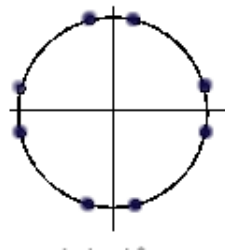


Figure A3.5 : allure de la constellation du signal à la sortie du modulateur QPSK, en utilisant l'étalement complexe de spectre

Sur cette figure on voit donc que l'étalement complexe de spectre a eu pour effet d'uniformiser l'amplitude du signal obtenu. Le récepteur pourra alors estimer efficacement le symbole émis car l'amplitude est restée inchangée.

Annexe 4

Démonstrations des équations (6.18) jusqu'à (6.34)

- **Equation (6.18)**

E_b est égal à $\frac{S}{R}$ où S est la puissance du signal de données et R est le débit de données. Le rapport puissance sur débit donne l'énergie par bit.

N_0 représente la densité du bruit (bruit + interférence). Si on ignore pour le moment l'effet du bruit, chaque utilisateur introduit une densité d'interférence égale à $\frac{S}{W}$ où W est la bande passante utilisée par la transmission. Et comme il y a $N-1$ utilisateurs introduisant des interférences dans un système à N utilisateurs, $N_0 = \frac{(N-1)S}{W}$.

- **Equation (6.19)**

Ici on tient compte de l'énergie ρ du bruit qui s'ajoute au terme d'interférence de l'équation (6.18). N_0 est alors égal à $\frac{(N-1)S}{W} + \frac{\rho}{W}$. En simplifiant on a (6.19).

- **Equation (6.20)**

Elle découle de (6.19).

- **Equation (6.21)**

Comme les utilisateurs ne parlent que pendant une fraction α du temps, d'où le terme d'interférence $\frac{(N-1)S}{W}$ se trouve multiplié par α . En simplifiant et tenant compte qu'on calcule maintenant la capacité d'un secteur (remplacement de N par N_s), on a l'équation (6.21).

- **Equation (6.22)**

Elle découle de (6.22)

- **Equation (6.23)**

La perte de propagation est proportionnelle à $10^{\left(\frac{\xi}{10}\right)} r^{-4}$. Le mobile interférent produira dans l'utilisateur voulu une interférence égale à $I(r_0, r_m) = S \left[\frac{10^{\left(\frac{\xi_0}{10}\right)}}{r_0^4} \right] \left[\frac{r_m^4}{10^{\left(\frac{\xi_m}{10}\right)}} \right]$ étant donné que chaque utilisateur émet avec une puissance S . Le rapport est inférieur à 1 car le mobile est toujours connecté à la station de base correspondant au minimum de perte de propagation.

- **Equation (6.24)**

La surface d'un hexagone de rayon unité est égale à $\frac{3\sqrt{3}}{2}$, d'où la densité $\rho = \frac{2N}{3\sqrt{3}}$ et l'équation (6.24).

- **Equation (6.25)**

C'est le rapport interférence totale sur le signal utile, c'est donc l'intégrale de l'équation (6.23) que multiplie une variable représentant l'activité vocale ψ , la densité ρ ainsi que la fonction Φ (voir équation (6.27)).

- **Equation (6.26)**

Cette équation traduit le fait que le mobile soit toujours connecté à la station de base correspondant à la plus faible perte de propagation.

- **Equation (6.27)**

Cette équation garantit que l'atténuation due à la propagation entre le mobile voulu et sa station de base est inférieure à celle du mobile interférent.

- **Equation (6.28)**

Ces valeurs sont tirées de la simulation effectuée dans [11]

- **Equation (6.29)**

C'est une réécriture de l'équation (6.21) mais en remplaçant le terme $(N_s - 1)\alpha$ par la variable aléatoire X_i prenant la valeur 1 avec une probabilité α et la valeur 0 avec une probabilité $(1 - \alpha)$. Cette réécriture est justifiée par le fait que cette façon de faire n'introduit

pas de modification du résultat final étant donné que l'on ne s'intéresse qu'aux valeurs moyennes.

- **Equation (6.30)**

Comme on le sait, une performance adéquate (TEB inférieur à 10^{-3}) peut être atteinte avec $E_b / N_0 \geq 5$ (7 dB). Alors la performance requise peut être atteinte avec une probabilité $P = \Pr(\text{TEB} > 10^{-3}) = \Pr(E_b / N_0 \geq 5)$. Nous pouvons minorer cette probabilité en majorant sa probabilité complémentaire pour P donnée ($P = 0.99$ par exemple) et en tenant compte de l'équation (6.29) :

$$1 - P = \Pr(\text{TEB} > 10^{-3}) = \Pr\left(\sum_{i=1}^{N_s} X_i + I/S > \delta\right)$$

où

$$\delta = \frac{W/R}{E_b / N_0} - \frac{\eta}{S}, \quad E_b / N_0 = 5$$

Sachant que la variable aléatoire X_i possède une distribution binomiale (équation (6.29)), que I/S est une variable Gaussienne dont le moyenne et la variance sont données par l'équation (6.28) et que toutes les variables sont mutuellement indépendantes on a

$$\begin{aligned} 1 - P &= \sum_{k=0}^{N_s-1} \Pr(I/S > \delta - k | \sum X_i = k) \Pr(\sum X_i = k) \\ &= \sum_{k=0}^{N_s-1} C_k^{N_s-1} \alpha^k (1-\alpha)^{N_s-1-k} Q\left(\frac{\delta - k - 0.247N_s}{\sqrt{0.078N_s}}\right) \end{aligned}$$

$$\text{où } Q(x) = \frac{1}{\sqrt{2\pi}} \int_x^\infty e^{-y^2/2} dy$$

- **Equation (6.31)**

$$E_b = \beta \Phi_i S_{T_i} / R$$

et

$$N_0 = \left[\left(\sum_{j=1}^K S_{T_j} \right) + \eta \right] / W$$

où K est le nombre de stations de base dans le système.

D'où l'équation (6.31).

- **Equation (6.32)**

Elle découle de l'équation (6.31)

- **Equation (6.34)**

Il est clair que $\sum_{i=1}^{N_s} \Phi_i \leq 1$ étant donné que βS_{T_i} est la puissance maximale allouée pour

l'utilisateur considéré. On a alors

$$\sum_{i=1}^{N_s} f_i \leq \frac{\beta W / R}{E_b / N_0} - \sum_{i=1}^{N_s} \frac{\eta}{S_{T_i}} \equiv \delta'$$

Cette expression constitue une contrainte sur l'allocation de la puissance utilisée par la station de base.

Généralement, le bruit de fond est bien plus faible que le total des signaux issus des stations de base, d'où la seconde somme est négligeable. Ainsi on a

$$\delta' = \frac{\beta W}{R} \frac{E_b}{N_0}$$

Donc la performance voulue ne peut être atteinte lorsque le nombre d'utilisateurs par secteur ne satisfait plus la contrainte sur l'allocation de puissance précédente.

D'où

$$1 - P = \Pr(\text{TEB} > 10^{-3}) = \Pr\left(\sum_{i=1}^{N_s} f_i > \delta'\right)$$

et l'équation (6.34).

Annexe 5

Les algorithmes utilisés pour résoudre les problèmes d'optimisation

Le problème général d'optimisation peut s'exprimer par

$$\begin{aligned} \text{minimiser } & f(\mathbf{x}) \\ & \mathbf{x} \in \mathcal{X}^n \end{aligned}$$

étant donné que

$$\mathbf{G}_i(\mathbf{x}) = 0, \quad i = 1, \dots, m_e$$

$$\mathbf{G}_i(\mathbf{x}) \leq 0, \quad i = m_e + 1, \dots, m$$

$$\mathbf{x}_l \leq \mathbf{x} \leq \mathbf{x}_u$$

où $\mathbf{x}$ est le vecteur de variables, $f(\mathbf{x})$ est la fonction objective qui retourne une valeur scalaire, et la fonction vectorielle $\mathbf{G}(\mathbf{x})$ retourne les valeurs des contraintes (égalités et inégalités) évaluées au point $\mathbf{x}$.

Dans le logiciel Matlab, on peut utiliser deux fonctions pour résoudre les problèmes d'optimisation :

- *fminunc* pour les optimisations sans contraintes
- *fmin* pour les optimisations avec contraintes.

A5.1 Les optimisations sans contraintes

On utilise pour cela les méthodes dites « *Quasi-Newton* ». Ces méthodes utilisent les gradients pour déterminer la direction de recherche des extremums. Dans toute la suite, supposons que le problème consiste à trouver le minimum d'une fonction. Les méthodes Quasi-Newton obtiennent les informations à propos de la courbure de la fonction considérée en formulant un problème de la forme :

$$\min_{\mathbf{x}} \frac{1}{2} \mathbf{x}^T \mathbf{H} \mathbf{x} + \mathbf{c}^T \mathbf{x} + \mathbf{b}$$

où la matrice de Hessian $\mathbf{H}$, est une matrice symétrique définie positive, $\mathbf{c}$ est un vecteur constant, et $\mathbf{b}$ est une constante. La solution optimale à ce problème apparaît quand les dérivées partielles par rapport à $\mathbf{x}$ sont égales à 0 c'est à dire :

$$\nabla f(\mathbf{x}^*) = \mathbf{H}\mathbf{x}^* + \mathbf{c} = \mathbf{0}$$

La solution optimale, $\mathbf{x}^*$ peut s'écrire

$$\mathbf{x}^* = -\mathbf{H}^{-1}\mathbf{c}$$

La matrice $\mathbf{H}$ est construite en suivant une itération :

$$\mathbf{H}_{k+1} = \mathbf{H}_k + \frac{\mathbf{q}_k \mathbf{q}_k^T}{\mathbf{q}_k^T \mathbf{s}_k} - \frac{\mathbf{H}_k^T \mathbf{s}_k \mathbf{s}_k^T \mathbf{H}_k}{\mathbf{s}_k^T \mathbf{H}_k \mathbf{s}_k}$$

où

$$\mathbf{s}_k = \mathbf{x}_{k+1} - \mathbf{x}_k$$

$$\mathbf{q}_k = \nabla f(\mathbf{x}_{k+1}) - \nabla f(\mathbf{x}_k)$$

Au démarrage, $\mathbf{H}_0$ peut être choisie comme une quelconque matrice symétrique définie positive, par exemple, la matrice unité. Pour éviter l'inversion de la matrice $\mathbf{H}$, on utilise la même équation que celle pour calculer $\mathbf{H}$ mais en remplaçant $\mathbf{s}_k$ par $\mathbf{q}_k$. L'information sur les gradients peut être obtenue par une fonction analytique ou par les dérivées partielles en utilisant une méthode de différentiation numérique via les différences finies.

Après chaque itération $\mathbf{k}$, une recherche du minimum est effectuée dans la direction

$$\mathbf{d} = -\mathbf{H}_k^{-1} \cdot \nabla f(\mathbf{x}_k)$$

Cette recherche se fait par la modification de la valeur de $\mathbf{x}_k$ en suivant l'algorithme :

$$\mathbf{x}_{k+1} = \frac{1}{2} \frac{\beta_{23} f(\mathbf{x}_1) + \beta_{31} f(\mathbf{x}_2) + \beta_{12} f(\mathbf{x}_3)}{\gamma_{23} f(\mathbf{x}_1) + \gamma_{31} f(\mathbf{x}_2) + \gamma_{12} f(\mathbf{x}_3)}$$

où

$$\beta_{ij} = \mathbf{x}_i^2 - \mathbf{x}_j^2$$

$$\gamma_{ij} = \mathbf{x}_i - \mathbf{x}_j$$

On parle alors *d'interpolation quadratique*. Le minimum de la fonction est obtenue (ici pour $\mathbf{x}_{k+1}$) lorsqu'on a quelques choses du genre :

$$f(x_{k+1}) < f(x_k) \text{ et } f(x_{k+1}) < f(x_{k+2})$$

A5.2 Les optimisations avec contraintes

On utilise pour cela la méthode dite « *programmation quadratique séquentielle* ». A chaque itération, une approximation de la matrice de Hessian de la fonction de Lagrange est faite en utilisant la méthode « *Quasi-Newton* ». Ceci est alors utilisé pour générer un sous-problème de Programmation Quadratique dont la solution est utilisée pour obtenir la direction de recherche.

L'idée principale consiste donc à formuler un sous-problème basé sur une approximation quadratique de la fonction de Lagrange.

$$L(x, \lambda) = f(x) + \sum_{j=1}^m \lambda_j \cdot g_j(x)$$

La matrice de Hessian de la fonction de Lagrange peut être obtenue par

$$H_{k+1} = H_k + \frac{q_k q_k^T}{q_k^T s_k} - \frac{H_k^T H_k}{s_k^T H_k s_k}$$

où

$$s_k = x_{k+1} - x_k$$

$$q_k = \nabla f(x_{k+1}) + \sum_{i=1}^n \lambda_i \cdot \nabla g_i(x_{k+1}) - \left(\nabla f(x_k) + \sum_{i=1}^n \lambda_i \cdot \nabla g_i(x_k) \right)$$

et où λ_i ($i = 1, \dots, m$) est une estimation des multiplicateurs de Lagrange.

Ensuite une méthode similaire à celle adoptée pour l'optimisation sans contraintes est utilisée pour parvenir à la solution finale.

BIBLIOGRAPHIE

- [1] J.G. Remy, J. Cueugnet et C. Siben, *"Systèmes de communications avec les mobiles"*, Eyrolles, Paris 1988.
- [2] A. Tanenbaum, *"Réseaux"*, Dunod / Prentice Hall, Londres 1997.
- [3] H.Leib, *"PCS Third generation CDMA systems, Study of the physical layer"*, McGill University – Department of Electrical and Computer Engineering, August 2001, disponible en téléchargement sur <http://www.tsp.ece.mcgill.ca>
- [4] M.K. Tsatsanis and G.B. Giannakis, *"Optimal decorrelating receivers for DS-CDMA systems: A signal framework"*, IEEE Trans. Signal Proc., vol. 42 n°12, pp. 3044-3055, December 1996
- [5] J.Meel, *"Spread Spectrum"*, De Nayer Instituut, Belgique, Octobre 1999, disponible en téléchargement sur <http://www.sss-mag.com>
- [6] R.Prasad, T. Ojanpera, *"An overview of CDMA evolution toward wideband CDMA"*, IEEE Commun. Surveys, vol. 1 n°1, pp. 1-29, Fourth quarter 1998
- [7] *"Direct Sequence vs. Frequency Hopping"*, Omnispread Communications, Inc., 1999, disponible en téléchargement sur <http://www.omnispread.com>
- [8] N. Doriswamy, *"Reconfigurable modules to support multiple CDMA standards"*, Berkley Varitronics Systems, Inc., disponible en téléchargement sur <http://www.bvsystems.com>
- [9] J. Fakatselis, *"Processing gain for direct sequence spread spectrum communication systems and PRISM"*, Intersil, August 1996
- [10] P. Flikkema, *"Direct sequence spread spectrum"*, disponible en téléchargement sur <http://www.sss-mag.com>
- [11] K.S. Gilhousen, I.M. Jacobs, R. Padovani, A.I. Viterbi, L.A. Weaver, Jr., and C.E. Wheatley, *"On the capacity of a cellular CDMA system"*, IEEE Trans. Veh. Technol., vol. 40 n°2, pp. 303-312, May 1991
- [12] M.H. Fong, V.K. Bhargava, and Q. Wang, *"Concatenated Orthogonal / PN spreading sequences and their application to cellular DS-CDMA systems with integrated traffic"*, IEEE J. Sel. Areas Commun., vol. 14 n°3, pp.547-558, April 1996

- [13] J. Lifterman, *"Les méthodes rapides de transformation du signal : Fourier, Walsh, Hadamard, Haar"*, Masson, Paris, 1980.
- [14] E. Dahlman, B. Gudmundson, M. Nilsson, and J. Sköld, *"UMTS / IMT-2000 Based on Wideband CDMA"*, IEEE Commun. Mag., pp. 70-80, September 1998.
- [15] L. Lindh, *"Synchronization in CDMA systems"*, Licentiate Course on Signal Processing in Communications, Nokia Research Center, FALL-97, November 1997, disponible en téléchargement sur http://keskus.hut.fi/opetus/s38220/reports_97
- [16] Y.S. Chiou, *"Algorithm and performance evaluation of space-time RAKE receiver for W-CDMA systems"*, Institute of Communication Engineering National Chiao Tung University, China, 2000.
- [17] A. Abrardo, G. Giambene, and D. Sennati, *"Optimization of power control parameters for DS-CDMA cellular systems"*, IEEE Trans. Commun. , vol. 49 n°8, pp. 1415-1424, August 2001.
- [18] K.T. Kasengulu, *"About the capacity of a DS Spread Spectrum CDMA cellular system"*, Berocan inc, Montréal, disponible en téléchargement sur <http://www.sss-mag.com>
- [19] D.B. Levey, *"Interference and capacity calculations for CDMA systems"*, Electrobit, September 2001, disponible en téléchargement sur <http://www.electrobit.co.uk>
- [20] E.Dahlman, P.Beming, J. Knutsson, F. Ovesjö, M. Persson, and C. Roobol, *"WCDMA-The radio interface for future mobile multimedia communications"*, IEEE Trans. Veh. Technol., vol 47 n° 4, pp. 1105-1118, November 1998.
- [21] V.V. Veeravalli, and A. Sendonaris, *"The Coverage-Capacity tradeoff in cellular CDMA systems"*, IEEE Trans. Veh. Technol., vol.48 n°5, pp. 1443-1450, September 1999.
- [22] D. Hong, S. Choi, J. Cho, *"Coverage and Capacity analysis for the multi-layer CDMA Macro/Indoor-Pico cells"*, Dept. of Electronic Engineering, Sogang University, South Korea, disponible en téléchargement sur <http://www.sogang.ac.kr>
- [23] A. Ahmad, *"A CDMA network architecture using optimized sectoring"*, IEEE Trans. Veh. Technol., Under revision for publication, disponible en téléchargement sur <http://facweb.cs.depaul.edu/aahmad/papers>
- [24] E.H. Dinan and B. Jabbari, *"Spreading codes for direct sequence CDMA and wideband CDMA cellular networks"*, IEEE Commun. Mag., pp. 48-54, September 1998.

- [25] B.J. Choi, "*Spreading Sequences*", disponible en téléchargement sur <http://www-mobile.ecs.soton.ac.uk/bjc97r/pnseq-1.1>
- [26] "*Minimum spectrum demand per public terrestrial UMTS operator in the initial phase*", Report from the UMTS, n°5, 1998, disponible en téléchargement sur <http://www.umts-forum.org>
- [27] J.E. Padgett, C.G. Günther, and J. Hattori, "*Overview of wireless communications*", *IEEE Commun. Mag.*, pp. 28-41, January 1995.
- [28] M. Valkama, "*An introduction to spread spectrum techniques and WCDMA*", Institute of Communications Engineering, Tampere University of Technology, Finland, 1999.
- [29] L.S. Fai, "*Adaptive multiuser interference cancellation*", Department of Electronic and Information Engineering, The Hong Kong Polytechnic University, Hong Kong, June 2000
- [30] R.A. Cameron, "*Fixed-point implementation of a multistage receiver*", Faculty of the Virginia Polytechnic Institute and State University, Virginia, January 1997.
- [31] P. Herhold, W. Rave, and G. Fettweis, "*Joint deployment of macro- and microcells in UTRAN FDD networks*", Dresden University of Technology, Germany, 2001.
- [32] C.C. Lee, and R. Steele, "*Effect of soft and softer handoffs on CDMA system capacity*", *IEEE Trans. Veh. Technol.*, vol. 47 n° 3, pp. 830-841, August 1998.
- [33] R.G. Akl, M.V. Hedge, M.Naraghi-Pour, and P.S. Min, "*Multicell CDMA Network design*", *IEEE Trans. Veh. Technol.*, vol 50 n°3, pp. 711-722, May 2001.

<i>Nom</i>	ANDRY
<i>Prénom</i>	Jean Onésime
<i>Adresse de l'auteur</i>	Logt 39 Cité Ambohipo, Antananarivo 101, MADAGASIKARA
<i>Titre du mémoire</i>	LE DS-CDMA ET SES APPLICATIONS DANS UN SYSTEME DE TELEPHONIE MOBILE TEL QUE L'UMTS
<i>Nombre de pages</i>	152
<i>Nombre de tableaux</i>	4
<i>Nombre de figures</i>	97
<i>Mots clés</i>	Etalement de spectre, Codes d'étalement, Accès multiple, DS- CDMA, Gain de traitement, Auto-corrélation, Inter-corrélation,
<i>Directeur de mémoire</i>	RAKOTOMALALA Mamy Alain

RESUME :

Le CDMA (Code Division Multiple Access) ou AMRC (Accès Multiple à Répartition dans les Codes) peut utiliser les ressources fréquentielles d'une manière plus efficace que les autres méthodes comme le TDMA ou le FDMA. Cet exposé montre les bénéfices qu'on peut tirer de l'utilisation d'un système DS-CDMA : protection contre les écoutes frauduleuses, élimination des effets des propagations par trajets multiple, suppression des interférences à bande étroite, augmentation à grande échelle de la capacité par rapport à celles proposées par les techniques analogiques et même par des techniques numériques compétitives. Des programmes de simulation sont aussi proposés pour permettre la compréhension des comportements des systèmes DS-CDMA dans le cas d'un canal à trajets multiples, d'un canal bruité, d'une interférence à bande étroite ainsi que dans une configuration multicellulaire.

ABSTRACT:

The use of CDMA (Code Division Multiple Access) can make a more efficient use of the spectrum frequency than other schemes like TDMA or FDMA. This paper shows the benefits that can be obtained with a DS-CDMA system: low probability of intercept, multipath mitigation, narrow-band interference suppression, many-fold increase in capacity over analog and even over competing digital techniques. Simulation programs are also proposed, which are helpful for understanding the behaviour of DS-CDMA systems in the case of multipath channel, noisy channel, narrow-band interference and a multicell configuration.