

REPUBLIQUE ALGERIENNE DEMOCRATIQUE ET POPULAIRE
Ministère De L'enseignement Supérieur Et De La Recherche Scientifique
Université De Batna
Faculté De Technologie



MEMOIRE

Présenté au Département de génie électrique

En vue de l'obtention du diplôme de

Magister en Electronique

Option : **Micro-ondes Pour Télécommunications**

Par

Fateh Bouguerra

Thème

***Contribution à l'optimisation des
télécommunications
dans les réseaux mobiles***

Devant le jury constitué de :

Pr. Nabil Benoudjit	Prof.	Univ. Batna	Président,
Dr. Lamir Saidi	M.C. (A)	Univ. Batna	Rapporteur,
Pr. Abdelmadjid Bensalama	Prof.	Univ. Constantine	Examineur,
Dr. Rédha Benzid	M.C. (A)	Univ. Batna	Examineur,
Dr. Salim Aissi	M.C. (B)	Univ. Batna	Invité.

Avril 2011

***Contribution à l'optimisation des
télécommunications
dans les réseaux mobiles***

*This dissertation is dedicated to my Family; specially my mother,
and to all my friends.*

Special Thanks to Mr. SAIDI Lamir for his efforts and assistance; driving me to the right way and his patient on my errors, he was like a missed colleague.

I would to thank Mr. BENOUDJIT Nabil to have agreed to chair the Jury of this dissertation.

Also, I address my cordial thanks to Mr. BENSALAMA Abdelmadjid, Mr. BENZID Redha, and Mr. AISSI Salim for the honour they made me to judge this modest work.

Thank you all for being with me.

Sommaire

Acronymes et Abréviations	X
Listes des tableaux	XIV
Liste des figures	XV
Introduction générale	XVIII

Chapitre I : Généralités..... 1

1. Introduction	2
2. Evolution et concept des systèmes cellulaires	2
2.1. Systèmes de première génération (analogiques)	2
2.2. Systèmes de seconde génération (digitals)	2
2.3. Systèmes de troisième génération (WCDMA dans l'UMTS)	3
2.4. Réseaux de quatrième génération (All-IP)	3
2.5. Concept cellulaires	4
2.5.1. Réutilisation des fréquences	4
2.5.2. Handover	4
3. 2G (GSM)	4
3.1. Sous-système de station de base (BSS)	5
3.1.1. Base Transceiver Station (BTS)	5
3.1.2. Contrôleur de station de base (BSC)	5
3.2. Sous-système réseau (NSS)	6
3.2.1. Mobil Switching Center (MSC)	6
3.2.2. Home Location Register (HLR)	6
3.2.3. Visitor Location Register (VLR)	6
3.2.4. Authentication Center (AUC)	6
3.2.5. Equipment Identity Registers (EIR)	7
3.3. Système de gestion du réseau (NMS)	7
3.4. Interfaces et signalisations en GSM	7
3.4.1. Interfaces	7
3.4.2. Signalisations	8

4. 3G (UMTS)	8
4.1. Radio Access Network (RAN).....	8
4.1.1. Base Station (BS)	9
4.1.2. Radio Network Controller (RNC).....	9
4.2. Core Network (CN).....	9
4.2.1. WCDMA Mobile Switching Centre (WMSC) and VLR	10
4.2.2. Gateway Mobile Switching Centre (GMSC).....	10
4.2.3. Home Location Register (HLR).....	10
4.2.4. Serving GPRS Support Node (SGSN).....	10
4.2.5. Gateway GPRS Support Node (GGSN)	10
4.3. Network Management System in 3G Networks	10
4.4. Interfaces et signalisation en 3G.....	11
4.4.1. Interfaces	11
4.4.2. Signalisations.....	12
5. Modulations	12
5.1. Binary Phase Shift Keying (BPSK)	13
5.2. Gaussian Minimum Phase Shift Keying (GMSK).....	13
5.3. Octagonal Phase Shift Keying (8-PSK)	13
5.4. Quarter Phase Shift Keying (QPSK)	13
5.5. Modulation par code (CDMA et étalement de spectre).....	14
5.6. Modulation par fréquences orthogonales (OFDM).....	14
6. Techniques de duplexage.....	14
6.1. FDD	14
6.2. TDD	15
7. Techniques d'accès	16
7.1. FDMA	16
7.2. TDMA	17
7.3. CDMA.....	18
7.4. Les performances comparées de TDMA/FDMA/CDMA	18
8. Conclusion.....	19

Chapitre II : DS-CDMA : Codage & détection 20

1. Introduction	21
2. Etalement de spectre.....	21

2.1. Pourquoi étaler un spectre ?	23
2.2. Avantages & inconvénients	25
3. CDMA	25
3.1. FH-CDMA (Frequency Hopping CDMA)	27
3.2. TH-CDMA (Time Hopping CDMA)	28
3.3. DS-CDMA (Direct Sequence CDMA)	28
3.3.1. Avantages	29
3.3.2. Inconvénients	29
4. Codage	30
4.1. Codes Orthogonaux	31
4.2. Codes pseudo aléatoires (PN codes)	32
5. Détection	33
5.1. Détection mono-utilisateur	33
5.2. Détection multi-utilisateurs	34
6. Conclusion	34

Chapitre III : Problèmes & égalisation du canal 35

1. Introduction	36
1.1. Interférence d'accès multiples (MAI)	36
1.2. Etats par trajets multiples de canal	36
1.3. Interférence d'inter symboles (ISI)	36
2. Egalisation	38
2.1. Egalisation avec apprentissage	39
2.1.1. Egalisation adaptative	39
2.1.2. Filtre de Wiener	40
2.1.3. Types d'algorithmes adaptatifs	43
2.1.4. Choix de l'algorithme	45
2.1.5. Algorithme LMS	46
2.1.6. Algorithme NLMS	48
2.1.7. Algorithme RLS	49
2.2. Egalisation aveugle	50
2.2.1. L'algorithme CMA (Constant Modulo Algorithme)	51
2.2.2. Equation de mise à jour de CMA	52
3. Simulations et résultats	53

3.1. Principe	53
3.2. Algorithme LMS.....	54
3.3. Algorithme NLMS	59
3.4. Algorithme RLS	63
3.5. Algorithme CMA	67
4. Conclusion.....	72

Chapitre IV : Optimisation neuronale..... 74

1. Introduction.....	75
1.1. Egalisation linéaire	75
1.2. Egalisation non linéaire	75
2. L'optimisation.....	76
2.1. Qu'est ce que l'optimisation ?	76
2.2. Catégories d'optimisation.....	76
3. Les réseaux de neurones artificiels (RNA)	78
3.1. Architecture et fonctionnement	79
3.2. Fonctions de transfert	80
3.3. Entraînement des réseaux de neurones.....	81
4. Multilayer Perceptrons (MLP)	81
4.1. Algorithme d'apprentissage : Rétropropagation du gradient (BP)	82
4.1.1. Etapes de la BP	82
4.1.2. Principe	82
4.2. Equations du réseau	83
4.2.1. Adaptation des poids	84
4.2.2. Problèmes de la BP	85
4.3. Simulations et résultats.....	86
4.3.1. Valeurs complexes	86
4.3.2. Principe	87
4.3.3. Etapes de simulation.....	88
4.3.4. Résultats et discussions	88
4.4. Conclusion	93
5. Radial Basis Function (RBF)	94
5.1. Types d'apprentissage.....	94
5.1.1. Apprentissage supervisé	94

5.1.2. Apprentissage non supervisé.....	94
5.2. Apprentissage des réseaux à fonctions radiales de bases (RBF)	95
5.3. Equations du réseau	96
5.4. Simulations et résultats	97
5.4.1. Principe.....	97
5.4.2. Etapes de mis en œuvre du réseau RBF	98
5.4.3. Résultats et discussions	98
5.5. Conclusion	102
6. Comparaison entre les égaliseurs RBF & MLR	102
7. Conclusion	104

Conclusion générale	105
----------------------------------	------------

Références	107
------------------	-----

Acronymes et Abréviations

3GPP	The 3 rd Generation Partner Project.
8-PSK	Octagonal Phase Shift Keying.
A-CDMA	Asynchronous CDMA.
ANN	Artificial Neural Networks.
AuC	Authentication Center.
BER	Bit Error Rate.
BP	Back Propagation.
BPSK	Binary Phase Shift Keying.
BS	Base Station.
BSC	Base Station Controller.
BSS	Base Station Sub-system.
BTS	Base Transceiver Station.
C-NBAP	Common NBAP.
CCITT	Comité Consultatif International Téléphonique et Télégraphique.
CCK	Complementary Code Keying.
CDMA	Code Division Multiple Access.
CLMS	Complex LMS.
CM	Constant Modulo.
CMA	CM Algorithm.
CN	Core Network.
CS	Circuit Switched.
D-NBAP	Dedicated NBAP.
DECT	Digital Enhanced Cordless Telecommunications.
DPDCH	Dedicated Physical Data Channel.
DS-CDMA	Direct-Sequence CDMA.
DSSS	Direct Sequence Spread Spectrum.
EDGE	Enhanced Data for Global Evolution.
EIR	Equipment Identity Register.
ETSI	European Telecommunications Standards Institute.
FAP	Fast Affine Projection Algorithm.
FDD	Frequency Division Duplex.
FDMA	Frequency Division Multiple Access.
FH-CDMA	Frequency Hopping CDMA.
FIR	Finite Impulse Response.

FM	Frequency modulation.
FTF	Fast Transversal Filter.
GA	Genetic Algorithms.
GGSN	Gateway GPRS Support Node.
GMSK	Gaussian Minimum Phase shift Keying.
GPRS	General Packet Radio Service.
GPS	Global position system.
GSM	Global System for Mobile communication.
HLR	Home Location Register.
IC	Interference Cancellation.
IIR	Infinite Impulse Response.
IMEI	International Mobile Equipment Identity.
IMT-2000	International Mobile Telecommunications 2000.
IP	Internet Protocol.
IS-95	Interim Standard 95.
ISDN	Integrated Services Digital Network Design.
ISI	InterSymbole Interference.
ITU	International Telecommunication Union.
ITU-T	ITU Telecommunication Standardization Sector.
JD	Joint Detection.
LAPD	Link Access Protocol For The D Channel.
LAPD_m	Link Access Protocol For The D Channel Modified.
LTF	Linear Transversal Filter.
LMS	Least Mean Square.
MAI	Multiple Access Interference.
MBER	Minimum Bit-Error Rate.
MBOK	M-ary Bi-Orthogonal Keying.
MGW	Media Gateway.
MLP	Multi Layer Perceptrons.
MMSE	Minimum Mean Square Error.
MOE	Minimal Output Energy.
MSC	Mobile Switching Center.
MSE	Mean Square Error.
MSK	Minimum Phase shift Keying.
NBAP	Node B Application Part.
NLMS	Normalised Least Mean Squares.
NMS	Network Management System.

NMTs	Nordic Mobile Telephones.
NSS	Network Switching Service.
OFDM	Orthogonal Frequency Division Multiplexing.
OMC	Operation and Maintenance Center.
OSI	Open System Interconnection.
OVSF	Orthogonal Variable Spreading Factor Codes.
PN	Pseudo-Noise sequences.
PS	Packet Switched.
PSK	Phase shift Keying.
PSTN	Public Switch Telephone Network.
QoS	Quality of Service.
QPSK	Quarter Phase Shift Keying.
RAN	Radio Access Network.
RANAP	RAN Application Part.
RBF	Radial Basis Function.
RF	Radio Frequency.
RLS	Recursive Least Square.
RNA	Réseaux de neurones artificiels.
RNC	Radio Network Controller.
RNS	Radio Network Subsystem.
RNSAP	RNS Application Part.
S-CDMA	Synchronous CDMA.
SF	Spreading Factor.
SGSN	Serving GPRS Support Node.
SMSC	Short Message Service Center.
SNR	Signal to Noise Ratio.
SOS	Second Order Statistics.
SS7	Signalling System Number 7.
TACS	Total Access Communication System.
TCSM	Transcoder Sub-Multiplexer.
TDD	Time Division Duplex.
TDMA	Time Division Multiple Access.
TH-CDMA	Time-Hopping CDMA.
TV	Television.
UE	User Equipment.
UMTS	Universal Mobile Telecommunications System.
UTRAN	Universal Terrestrial Radio Access Network.

VLR	Visitor Location Register.
WCDMA	Wideband CDMA.
WLAN	Wireless Local Area Network.
WMSC	WCDMA Mobile Switching Center.
xG	x th Generation Of Wireless Communication Technology (2 nd , 2.5, 3 rd , 3.5, 4 th).

Listes des tableaux

2.1	Caractéristiques de quelques standards de télécommunication.....	28
3.1	Algorithmes adaptatifs à travers les filtres adaptatifs.....	45
4.1	Fonctions de transferts d'un neurone.....	80
4.2	MSE en Fonction du SNR pour les égaliseurs RBF et MLP sur l'ensemble des validation, test et le total.....	102

Listes des figures

1.1	Cellules radio & principes de réutilisation des fréquences.....	4
1.2	Architecture 2G (GSM).....	5
1.3	Architecture 3G (UMTS).....	9
1.4	Le modèle OSI en 3G	12
1.5	Bandes GSM et UMTS (FDD).....	15
1.6	Multiplexage temporel -GSM- (TDD).....	15
1.7	L'accès multiple par répartition de fréquence (FDMA).....	17
1.8	L'accès multiple par répartition de temps (TDMA)	17
1.9	L'accès multiple par répartition de code (CDMA).....	18
2.1	Technique d'étalement de spectre	21
2.2	Etalement de spectre par séquence directe (DSSS)	22
2.3	Comparaison signal à bande étroite / signal à bande étalée	23
2.4	Principe de la technique CDMA	26
2.5	Différents types de CDMA.....	27
2.6	Comparaison des techniques DS, FH et TH-CDMA	27
2.7	Technique DS-CDMA (Uplink)	29
2.8	Types de codes utilisés dans l'interface air.....	30
2.9	Canalisation et brouillage.....	31
2.10	Orthogonal Variable Spreading Factor Codes (OVSF).....	32
2.11	Classification des détecteurs DS-CDMA	33
3.1	Multi-trajets	37
3.2	Interférences inter-symboles.....	37
3.3	Egalisation du canal	38
3.4	Egalisation avec apprentissage	39
3.5	Schéma d'un système de filtrage adaptatif	40
3.6	Problème d'estimation linéaire.....	40
3.7	Catégorisation des filtres adaptatifs	44
3.8	Egalisation aveugle	51
3.9	Principe des simulations	53
3.10	Symboles à égaliser.....	54
3.11	L'erreur MSE du LMS sur 500 itérations	55
3.12	Convergence des coefficients de l'égaliseur LMS sur 500 itérations.....	55

3.13	La capacité de poursuite et l'erreur résiduelle du LMS sur 500 itérations	56
3.14	Constellation des signaux émis, reçus & égalisés avec le LMS	56
3.15	Convolution du canal $C(z)$ avec l'égaliseur LMS.....	57
3.16	Performances du LMS avec un pas $\mu = 0.09$	58
3.17	Performances du LMS avec un pas $\mu = 0.0003$	58
3.18	Performances & robustesse du LMS avec des SNR=30, 20 & 15 dB	59
3.19	L'erreur MSE du NLMS sur 500 itérations	60
3.20	Convergence des coefficients de l'égaliseur NLMS sur 500 itérations	60
3.21	Capacité de poursuite & l'erreur résiduelle du NLMS sur 500 itérations.....	61
3.22	Constellation des signaux émis, reçus (bleu) & égalisés (vert) avec le NLMS	62
3.23	Convolution du canal $C(z)$ avec l'égaliseur NLMS	62
3.24	Performances & robustesse du NLMS avec des SNR = 30, 20 & 15 dB.....	63
3.25	L'erreur MSE du RLS sur 500 itérations	64
3.26	Convergence des coefficients de l'égaliseur RLS sur 500 itérations	64
3.27	Capacité de poursuite & l'erreur résiduelle du RLS sur 500 itérations.....	65
3.28	Constellation des signaux émis, reçus (bleu) & égalisés (vert) avec le RLS	66
3.29	Convolution du canal $C(z)$ avec l'égaliseur RLS	66
3.30	Performances & robustesse du RLS avec des SNR = 30, 20 & 15 dB.....	67
3.31	L'erreur MSE du CMA sur 500 itérations	68
3.32	Convergence des coefficients de l'égaliseur CMA sur 500 itérations	68
3.33	Capacité de poursuite & l'erreur résiduelle du CMA sur 500 itérations.....	69
3.34	Constellation des signaux émis, reçus (bleu) & égalisés (vert) avec le CMA	69
3.35	Convolution du canal $C(z)$ avec l'égaliseur CMA.....	70
3.36	Performances du CMA avec un pas $\mu = 0.023$	71
3.37	Performances du CMA avec un pas $\mu = 0.0001$	71
3.38	Performances & robustesse du CMA avec des SNR = 30, 20 & 15 dB.....	72
4.1	Un groupement possible des techniques d'égalisations	75
4.2	Diagramme d'une fonction ou processus à optimiser	76
4.3	Catégories d'optimisation.....	77
4.4	Techniques de recherche et d'optimisation.....	78
4.5	Modèle de neurone artificiel.....	79
4.6	Structure d'un réseau de neurones non récurrent (statique)	79
4.7	Structure d'un filtre adaptatif split-complex (réel & dual univariable)	86
4.8	Structure d'un égaliseur neuronal MLP	87
4.9	MSE en fonction du nombre de neurones dans la couche cachée du MLP	89
4.10	MSE en fonction du pas d'apprentissage μ du MLP	90

4.11	MSE de validation du MLP en fonction de 50 itérations.....	90
4.12	MSE de validation du MLP en fonction de 100 itérations.....	91
4.13	Poursuite & erreur du MLP sur l'ensemble de tests sur 500 itérations	92
4.14	Constellation des signaux émis, reçus (bleu) & égalisés (vert) avec le MLP	92
4.15	Performances & robustesse de l'égaliseur MLP avec des SNR=30, 20 & 15 dB	93
4.16	Architecture des réseaux RBF	95
4.17	Structure d'un égaliseur neuronal RBF	97
4.18	MSE en fonction de la largeur des gaussiens σ du RBF	99
4.19	MSE en fonction du nombre de neurones dans la couche cachée du RBF.....	99
4.20	Poursuite & erreur du RBF sur l'ensemble de tests sur 500 itérations	100
4.21	Constellation des signaux émis, reçus (bleu) & égalisés (vert) avec le RBF	101
4.22	Performances & robustesse de l'égaliseur RBF avec des SNR=30, 20 & 15 dB	101
4.23	Performances & robustesse des égaliseurs RBF (rouge) & MLP (bleu) avec des SNR=30, 20 et 15 dB sur 500 itérations	103

Introduction générale

La demande croissante en capacité dans les réseaux mobiles est la force motrice derrière l'amélioration des réseaux établis et le déploiement de nouveaux standards de communication mobile dans le monde entier. L'interface air est la plus importante des interfaces dans la plupart des systèmes mobiles. L'importance de cette interface résulte du fait que c'est la seule interface à qui l'abonné mobile est exposé, et la qualité de cette interface est cruciale pour le succès du réseau mobile. La qualité de cette interface dépend principalement de l'utilisation efficace du spectre de fréquence qui lui est assigné, et les techniques d'accès mises en œuvre.

Le choix d'une technique d'accès (FDMA, TDMA, CDMA, ou une hybridation de deux d'entre elles) peut avoir un impact important sur les performances, la QoS (Quality of Service) et la capacité du système. Ce choix est tellement prépondérant qu'on dénomme souvent les systèmes en fonction de l'accès multiple.

Le calcul de capacité montre que la technique séquence directe d'accès multiple par répartition de code (DS-CDMA) a une grande capacité par rapport aux techniques d'accès multiple par répartition de temps/fréquence (TDMA/FDMA). Par conséquent, la plupart des systèmes mobiles de troisième génération vont incorporer une certaine sorte de DS-CDMA.

L'étalement de spectre est l'un des avantages mis en avant pour l'utilisation du DS-CDMA. En effet, la puissance d'un signal, après codage, est étalée sur toute la largeur de la bande de fréquence disponible. De ce fait des caractéristiques importantes apparaissent : robustesse vis-à-vis des brouillages, détection difficile par les utilisateurs non concernés, facilité de mélanger des canaux voix et données, ... etc.

Cependant, en parcourant un trajet entre l'émetteur et le récepteur dans l'interface air, le signal transmis est sujet à de nombreux phénomènes dont la plupart ont souvent un effet dégradant sur la qualité du signal. Cette dégradation se traduit en pratique par des erreurs dans les messages reçus qui entraînent des pertes d'informations pour l'utilisateur ou le système. Les sources principales de dégradation dans les systèmes DS-CDMA dues à des causes internes et externes sont :

- Les atténuations de puissance du signal dues aux effets induits par le phénomène des trajets multiples,
- L'interférence d'inter symboles (ISI) où l'effet de chaque symbole ou chip transmis au dessus d'un canal dispersif se prolonge au delà de l'intervalle de temps employé pour représenter ce symbole,

- L'interférence d'accès multiples (MAI) qui se produit quand un certain nombre d'utilisateurs partage un canal commun simultanément, où les signaux d'autres utilisateurs apparaissent comme une interférence pour un utilisateur donné.

Les récepteurs classiques (matched filter, récepteur RAKE) ne sont pas efficaces pour combattre ces diverses sources d'interférence. Pour cela la suppression d'interférence a été l'objet de beaucoup d'efforts de recherche, plusieurs techniques sont proposées, dont l'égalisation de canal ; une procédure qui est utilisée par les récepteurs afin de réduire l'effet d'interférence entre symboles (ISI) et l'effet des multiples trajets, due à la propagation du signal modulé à travers le canal. Pour empêcher une dégradation sévère des performances, il est nécessaire de compenser la distorsion des canaux. Cette compensation est effectuée par un égaliseur adaptatif linéaire ou non linéaire.

Récemment, les réseaux de neurones artificiels (ANN) et les algorithmes génétiques (GAs) et autres algorithmes évolutionnaires ont reçu une grande attention concernant leurs potentiels comme techniques d'optimisation pour de vastes classes de problèmes ; ils sont utilisés la plupart du temps pour optimiser les problèmes des réseaux de communications. Ces problèmes sont souvent formulés comme une sorte de problèmes d'optimisation combinatoire.

Dans ce mémoire nous développons les deux techniques ; égalisation linéaire et non linéaire ; Nous ne nous intéresserons aux quatre algorithmes les plus utilisés dans l'égalisation adaptative linéaire qui sont LMS, NLMS et RLS pour l'égalisation adaptative avec apprentissage, et le CMA pour l'égalisation adaptative aveugle. Puis nous nous proposons de concevoir deux égaliseurs non linéaires basés sur deux modèles de réseaux de neurones MLP et RBF dans le but d'optimiser les performances du traitement, de poursuite, et de minimiser l'erreur vis-à-vis des effets du canal et du bruit ascendant.

Cette dissertation est divisée en quatre chapitres outre une introduction et une conclusion générales : le premier chapitre introduit des généralités concernant les réseaux mobiles ; les différentes générations, leurs conceptions cellulaires, les modulations et les duplexages utilisés, ainsi que les techniques d'accès et leurs avantages et inconvénients. Le deuxième chapitre décrit la technique DS-CDMA, comment étaler un spectre, les codages mis à profit, et les diverses techniques de détections d'utilisateurs. Le troisième chapitre prend en compte notre problématique, définit les problèmes du canal et dénombre les différentes techniques d'égalisations classiques, une discussion est faite après chaque simulation, une comparaison est également établie entre ces techniques.

Le quatrième chapitre présente notre contribution à ce problème, la conception d'un égaliseur non linéaire grâce aux réseaux de neurones MLP et RBF, les études théoriques

sont également présentées ; après chaque simulation une comparaison est faite avec les méthodes classiques. Une comparaison entre les performances des deux techniques est également exposée selon le critère MSE pour les différentes étapes d'apprentissages.

Généralités

CHAPITRE

1

DANS CE CHAPITRE

- ☒ Introduction
- ☒ Evolution & concept des systèmes cellulaires
- ☒ 2G (GSM)
- ☒ 3G (UMTS)
- ☒ Modulations
- ☒ Techniques de duplexage
- ☒ Techniques d'accès
- ☒ Conclusion

1. Introduction

Les premiers systèmes mobiles fonctionnaient en mode analogique. Les terminaux étaient de tailles importantes, seulement utilisables dans quelques applications spécifiques. Les systèmes mobiles actuels fonctionnent en mode numérique (la voie est échantillonnée, numérisée et transmise sous forme de bits), et, les progrès de la microélectronique ont permis de réduire la taille des téléphones mobiles à un format de poche.

On a vu une grande évolution au cours des vingt années passées à cause des demandes de différents services et le besoin croissant de la communication moderne, soit en miniaturisation, soit en transfert de données avec de haut débits pour plusieurs applications, ce qui a donné des sauts de générations entre plusieurs technologies.

2. Evolution et concept des systèmes cellulaires

Les réseaux mobiles sont différenciés entre eux par le mot « génération », comme première génération, seconde génération...etc. C'est tout à fait approprié parce qu'il y a un grand décalage entre les technologies de chaque génération.

2.1 Systèmes de première génération (analogiques)

Les systèmes mobiles de première génération qui ont commencé dans les années 80, ont été basés sur des techniques de transmission analogique. À ce moment-là, il n'y avait aucune coordination mondiale (ou même au niveau européen) pour le développement des techniques de norme pour les systèmes. Les pays nordiques ont déployé les téléphones mobiles nordiques ou les NMTs, alors que le royaume d'Angleterre et l'Irlande utilisaient le système de communication d'accès total ou le TACS, et ainsi de suite. Le Roaming n'était pas possible et l'utilisation efficace du spectre de fréquence n'était pas présente.

2.2 Systèmes de seconde génération (digitals)

Au milieu des années 80, la commission européenne a commencé une série d'activités pour libéraliser le secteur de communications, y compris celui des communications mobiles. Ceci a eu comme conséquence la création de l'ETSI, qui a hérité de toutes les activités d'étalonnage en Europe. Cela a vu la naissance des premières caractéristiques, le réseau est basé sur la technologie numérique ; ceci est appelé Global System for Mobile communication ou le GSM. Depuis que les premiers réseaux sont apparus au début de 1991, le GSM a graduellement évolué pour répondre aux exigences du trafic de données et notamment l'intégration de beaucoup de services comparativement aux réseaux originaux.

Les éléments principaux de ce système sont les BSS (sous-système de station de base), dans lesquels il y a la BTS (station de base de transmission) et le BSC (contrôleurs de station de base) ; et le NSS (sous-système de commutation de réseau), dans lequel il y a le MSC (centre de commutation mobile) ; VLR (registre d'endroit de visiteur) ; HLR (registre d'endroit maison) ; CA (Centre d'authentification), et EIR (registre d'identité d'équipement) (voir figure 1.2). Ce réseau est capable de fournir tous les services de base tels que les services de la parole et de données jusqu'à 9.6 Kbps, fax, ...etc. Le réseau GSM a également une prolongation aux réseaux de téléphonie fixe. Le prochain avancement dans le système GSM était l'addition de deux plates-formes, appelées respectivement le système de messagerie audio (VMS) et le centre de service messagerie courte (SMSC).

Après viendra le GPRS et le EDGE (2.5 G) pour satisfaire l'augmentation du trafic des données.

2.3 Systèmes de troisième génération (WCDMA dans l'UMTS)

En EDGE, le transfert à haut débit des données était possible, mais ce transfert de paquet sur l'interface air se comporte toujours comme un appel sur un réseau de commutation. Ainsi une partie de cette efficacité de raccordement de paquet est perdue dans l'environnement de commutateur de circuit. D'ailleurs, les normes pour développer les réseaux étaient différentes pour diverses régions du monde. Par conséquent, on a décidé d'avoir un réseau qui fournit des services indépendants de la technologie de la plate-forme et dont les normes de conception sont les mêmes globalement. Ainsi, 3G était née. En Europe ce système est appelé l'UMTS (système mobile terrestre universel). IMT-2000 est le nom d'ITU-T pour le système de troisième génération, alors que CDMA2000 est le nom de la variante 3G américaine. WCDMA est la technologie d'interface pour l'UMTS. Les composants principaux incluent les BS (station de base) ou le noeud B, RNC (contrôleur de réseau de radio) indépendamment du WMSC (centre de commutation mobile à large bande de CDMA) et le SGSN/GGSN. Cette plateforme offre beaucoup de services tels que Internet, vidéophonie, imagerie...etc.

2.4 Réseaux de quatrième génération (All-IP)

La raison fondamentale de la transition au Tout-IP est d'avoir une plateforme commune pour toutes les technologies qui ont été développées jusqu'ici, et d'harmoniser avec les espérances d'utilisateur concernant les nombreux services à fournir. La différence fondamentale entre le GSM/3G et le Tout-IP est que la fonctionnalité du RNC et la BSC est maintenant distribuée au BTS et à un ensemble de serveurs et de gateways (passages).

Ceci signifie que ce réseau sera moins cher et les transferts de données seront beaucoup plus rapides [1].

2.5 Concept cellulaires

Les bases de transmission sont réparties sur le territoire selon un schéma de cellules, Un réseau cellulaire comprend un ensemble de cellules dont la taille dépend de la puissance d'émission des émetteurs et surtout de la nature de l'environnement.

2.5.1 Réutilisation des fréquences

Chaque base utilise un groupe de fréquences différent de ses voisines ; les mêmes fréquences ne sont réutilisées qu'à une distance suffisante (Fig. 1.1) afin de ne pas créer d'interférences.

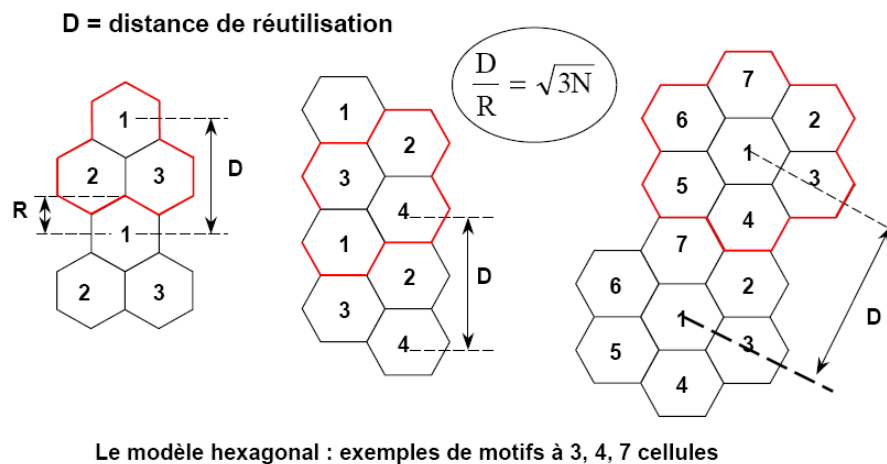


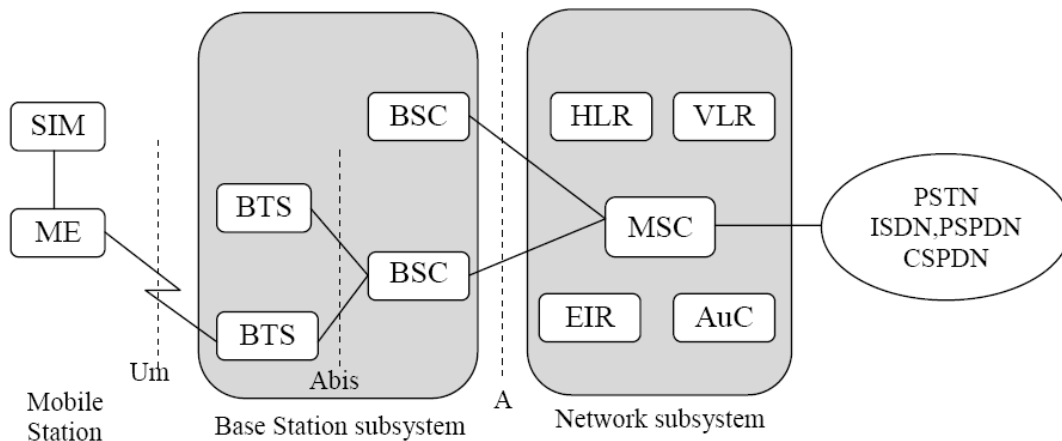
Fig.1.1 — Cellules radio & principes de réutilisation des fréquences

2.5.2 Handover

Pendant une communication, le terminal est en liaison radio avec une station de base bien déterminée. Le Handover est indispensable pour assurer la continuité de services alors que l'utilisateur se déplace. Il est donc nécessaire de changer la station de base avec laquelle le terminal est relié tout en maintenant la communication [2].

3. 2G (GSM)

De tous les systèmes mobiles de seconde génération, le GSM est le plus employé. Dans cette section nous passerons brièvement par ses constituants importants. Ce système peut être divisé en trois parties (sous-système de station de base, sous-système de réseau, et système de gestion du réseau). Voir figure 1.2.

**Fig.1.2 — Architecture GSM**

3.1 Sous-système de station de base (BSS)

Le BSS comprend la station de base (BTS), le contrôleur de station de base (BSC) et le multiplexeur de transcodeur secondaire (TCSM). Ce dernier est parfois physiquement placé au MSC. Par conséquent la BSC a également trois interfaces normalisées au réseau fixe, à savoir Abis, A et X.25.

3.1.1 Base Transceiver Station (BTS)

Elle contrôle l'interface entre le réseau et la station mobile. Par conséquent, elle remplit la fonction importante d'agir en tant que Hub pour la totalité de l'infrastructure en réseau. Des bornes mobiles sont liées au BTS par l'interface air. La transmission et la réception au BTS avec le mobile sont faites par l'intermédiaire d'antennes omnidirectionnelles ou directrices (habituellement ayant des secteurs de 120°). Les fonctions principales de la station de base sont la transmission des signaux dans le format désiré, codage et décodage des signaux, parant les effets de la propagation par trajets multiples en employant des algorithmes d'égalisation, le chiffage des trains de données, mesures de qualité et de la puissance du signal reçu, et le fonctionnement et la gestion de l'équipement de station de base elle-même.

3.1.2 Contrôleur de station de base (BSC)

Ce dispositif commande le sous-système radio, particulièrement les stations de base. Les fonctions principales du contrôleur de station de base incluent la gestion des ressources par radio et de la passation (Handover). Il est également responsable de la commande de la

puissance transmise, et il contrôle le OMC et sa signalisation, et ses configurations de sécurité et alarmes.

3.2 Sous-système réseau (NSS)

Le sous-système réseau agit en tant qu'interface entre le réseau GSM et les réseaux publics, PSTN et ISDN... (Fig. 1.2). Les composants principaux du NSS sont MSC, HLR, VLR, AUC, et EIR.

3.2.1 Mobil Switching Center (MSC)

Le MSC (ou commutateur comme il est généralement appelé) est l'élément le plus important du NSS car il est responsable des fonctions de commutation qui sont nécessaires pour des interconnexions entre les utilisateurs mobiles et d'autres utilisateurs mobiles et fixes du réseau. A cette fin, le MSC se sert de trois composants principaux du NSS : HLR, VLR et AUC.

3.2.2 Home Location Register (HLR)

Le HLR contient les informations inhérentes à chaque abonné mobile, tel que le type d'abonnement, les services que l'utilisateur peut employer, l'endroit courant de l'abonné, et l'état d'équipement mobile. La base de données dans le HLR demeure intacte et sans changement jusqu'à l'arrêt de l'abonnement.

3.2.3 Visitor Location Register (VLR)

Le VLR entre en action une fois que l'abonné entre dans la zone de couverture. A la différence du HLR, le VLR est d'une nature dynamique. Quand l'abonné se déplace à une autre région, la base de données de l'abonné est également décalée au VLR de la nouvelle région.

3.2.4 Authentication Center (AUC)

L'AUC est responsable des actions de maintien de l'ordre dans le réseau. Il contient tous les outils nécessaires pour protéger le réseau contre les faux abonnés et pour protéger également les appels des abonnés réguliers. Il y a deux clefs importantes dans les normes GSM : le chiffage des communications entre les utilisateurs mobiles, et l'authentification des utilisateurs. Les clefs de chiffage sont tenues dans l'équipement mobile et l'AUC et l'information est protégée contre l'accès non autorisé.

3.2.5 Equipment Identity Registers (EIR)

Chaque équipement mobile a sa propre identification personnelle, qui est dénotée par un nombre codifiée sous l'identité internationale d'équipement mobile (IMEI). Le nombre est installé pendant la fabrication de l'équipement et énonce sa conformation aux normes du GSM. Toutes les fois qu'un appel est fait, le réseau vérifie le nombre d'identité ; si ce nombre n'est pas trouvé sur la liste homologuée d'équipement autorisé, l'accès est nié. L'EIR contient cette liste de nombres autorisés et permet à l'IMEI d'être vérifié.

3.3 Système de gestion du réseau (NMS)

La tâche principale du NMS est d'assurer au réseau un fonctionnement flexible. A cette fin, elle a quatre tâches importantes à exécuter : surveillance du réseau, développement du réseau, mesures de réseau, et gestion de défauts.

Une fois que le réseau est en service, le NMS commence à surveiller sa performance. S'il voit un défaut il produit l'alarme appropriée. Quelques défauts peuvent être corrigés par le NMS lui-même (la plupart du temps orienté logiciel), alors que pour d'autres les visites sur sites sont exigées.

Le NMS est également responsable de la collecte des données et l'analyse des performances, menant de ce fait aux décisions précises liées à l'optimisation du réseau. La capacité et la configuration du NMS dépendent de la taille (en termes de capacité et secteur géographique) et des besoins technologiques du réseau.

3.4 Interfaces et signalisations en GSM

Comme le montre la figure 1.2, il y a plusieurs interfaces et signalisations impliquées dans le système GSM. Ici, nous en discuterons brièvement leur principe.

3.4.1 Interfaces

Interface air U_m : elle constitue la pièce centrale et la plus importante des interfaces dans la plupart des systèmes mobiles. L'importance de cette interface résulte du fait que c'est la seule interface à qui l'abonné mobile est exposé, et la qualité de cette interface est cruciale pour le succès du réseau mobile. La qualité de cette interface dépend principalement de l'utilisation efficace du spectre de fréquence qui lui est assigné. Elle s'appuie sur la signalisation LAPD_m (Modified Link Access Protocol for D-Channel).

Interface Abis : existe entre le BTS et le BSC et s'appuie sur le protocole LAPD, qui est utilisée pour assurer le trafic des données et les trames de signalisation.

Interface A : existe entre BSC et MSC, et s'appuie sur le protocole sémaphore N°7 du CCITT. Elle est utilisée pour assurer le trafic des données et les trames de signalisation. Le sous système radio et le sous système réseau, communiquent par l'intermédiaire de l'interface A.

3.4.2 Signalisations

LAPD_m : (Modified Link Access Protocol for D-Channel) est une version modifiée et optimisée du LAPD signalant pour l'interface air du GSM. La structure d'armature composée de 23 bytes est présente en trois formats : A, B et Abis. Les deux formats A et B sont employés pour la liaison montante et la liaison descendante, alors que l'Abis est employé seulement pour la liaison descendante.

SS7 : (Signalling System No 7) constitue la base de tout le trafic de signalisation sur toutes les interfaces du NSS. Il est employé entre le BSC et le MSC. Il est capable de contrôler l'information de signalisation dans les réseaux complexes. Il est employé pour la configuration d'appel, le management, et les dispositifs tels que le roaming, l'authentification, le renvoi d'appel,...etc.

X.25 : Cette interface relie le BSC au centre d'exploitation et de maintenance (OMC). Elle possède la structure en 7 couches du modèle OSI [1].

4. 3G (UMTS)

Les réseaux mobiles de troisième génération sont conçus pour les communications multimédia, augmentant de ce fait la qualité d'image et de vidéo, et augmentant les débits dans les réseaux publics et privés. Dans les forums de standardisation, la technologie WCDMA a émergé comme l'interface air la largement adoptée en 3G. Les spécifications ont été créées par le 3GPP et le nom WCDMA est employé couramment pour les deux opérations FDD et TDD. Les réseaux 3G se composent de deux parties : le réseau d'accès radio (RAN) et le réseau de noyau (NC). A son tour le RAN se compose des parties radio et transmission (voir figure 1.3).

4.1 Radio Access Network (RAN)

Les éléments principaux dans la présente partie du réseau sont la station de base (BS) et le contrôleur de réseau de radio (RNC). Les fonctions principales incluent la gestion des ressources par radio et la gestion des télécommunications.

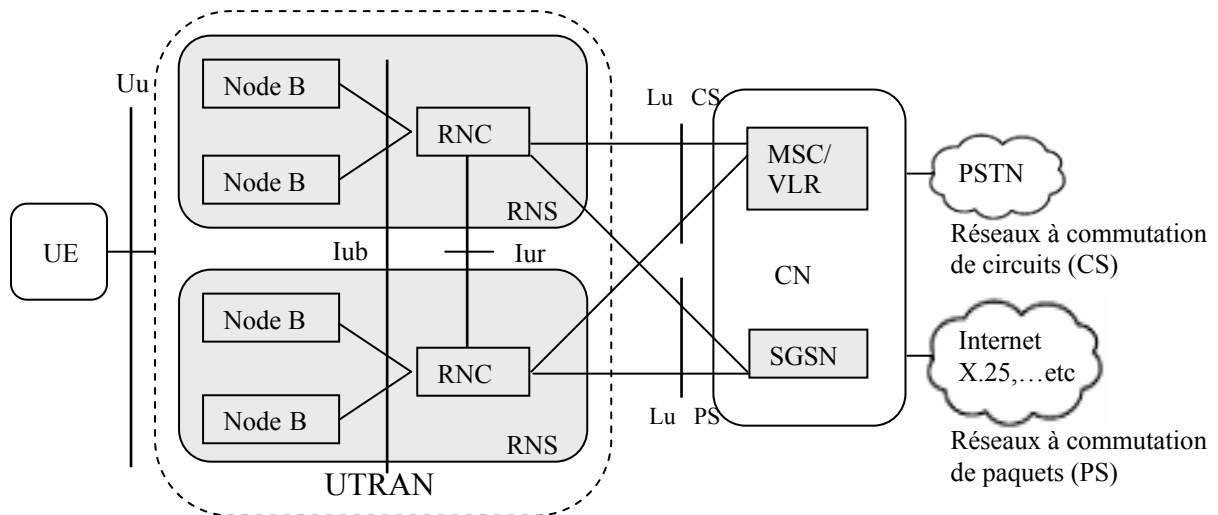


Fig.1.3 — Architecture 3G (UMTS)

4.1.1 Base Station (BS)

La station de base dans les réseaux 3G est également connue sous la dénomination noeud B. La station de base est une entité importante comme une interface entre le réseau et l'interface air. Comme dans les réseaux de deuxième génération, la transmission et la réception des signaux de la station de base sont faites par les antennes omnidirectionnelles ou directrices. Les fonctions principales des BS incluent le codage de canal, l'interfoliage (Interleaving), l'adaptation de taux, l'étalement,...etc., avec le traitement de l'interface air.

4.1.2 Radio Network Controller (RNC)

Ceci agit en tant qu'interface entre la station de base et le réseau de noyau. Le RNC est responsable de la commande des ressources par radio. En outre, à la différence dans le GSM, le RNC en même temps que la station de base peut manipuler toutes les fonctions de ressource par radio sans la participation du réseau de noyau. Les fonctions principales du RNC impliquent la commande de charge et de congestion des cellules, la commande d'admission et l'attribution de code, l'acheminement des données entre les interfaces I_{ub} et I_{ur} ...etc.

4.2 Core Network (CN)

Le réseau noyau dans les réseaux 3G se compose de deux domaines : domaine de commutation de circuit (CS) et domaine de commutation de paquets (PS). La partie CS traite le trafic en temps réel et la partie PS traite sa partie complémentaire. Ces deux domaines sont reliés à d'autres réseaux (par exemple CS au PSTN, et PS au réseau d'IP public).

La conception du protocole de l'UE et de l'UTRAN est basée sur la nouvelle technologie WCDMA, mais les définitions de CN ont été adoptées des caractéristiques du GSM. Les éléments importants de la CN sont WMSC, VLR, HLR, MGW, du côté de CS c'est SGSN (noeud de soutien du serveur GPRS) et GGSN (noeud de soutien de passage GPRS) du côté de PS.

4.2.1 WCDMA Mobile Switching Centre (WMSC) and VLR (Visitor Location Register)

Le commutateur et la base de données sont responsables des activités de control d'appel. WMSC est employé pour les transactions de CS, et la fonction du VLR contient l'information sur l'abonné visitant la région, qui inclut l'endroit du mobile dans une région.

4.2.2 Gateway Mobile Switching Centre (GMSC)

C'est l'interface entre le réseau mobile et les réseaux CS externes. Ceci établit les communications entrantes et sortantes dans le réseau. Elle trouve également le correct WMSC/VLR pour la connexion du chemin d'appel.

4.2.3 Home Location Register (HLR)

C'est la base de données dans laquelle se trouvent toutes les informations inhérentes à l'utilisateur mobile et le type de services auxquels il est souscrit. Une nouvelle entrée de base de données est faite quand un nouvel utilisateur est ajouté au système. Elle stocke également la localisation UE dans le système.

4.2.4 Serving GPRS Support Node (SGSN)

Le SGSN maintient une interface entre le RAN et le domaine PS du réseau. Il est principalement responsable des issues de gestion de mobilité comme l'enregistrement et la mise à jour de l'UE, des activités de pagination connexes, et des problèmes de sécurité pour le réseau PS.

4.2.5 Gateway GPRS Support Node (GGSN)

Il agit en tant qu'interface entre le réseau 3G et les réseaux PS extérieurs. Ses fonctions sont semblables au GMSC dans le domaine CS du CN, mais applicables pour le domaine PS.

4.3 Network Management System in 3G Networks

Pendant que la technologie des réseaux évoluait de la 2G à la 3G, les systèmes de gestion du réseau ont suivi cette évolution. Les NMS dans les systèmes 3G seront capables

de contrôler la commutation de paquets de données par comparaison avec les données de voix et avec les circuits à commutation dans les systèmes 2G. Les systèmes de gestion dans la 3G doivent être plus efficaces (c'est-à-dire plus de travail possible des NMS plutôt que de visiter les sites). Ces systèmes pourront optimiser la qualité du système généralement d'une manière plus efficace. On s'attend à ce qu'également les systèmes de gestion manipulent la multi technologie (c'est-à-dire 2G à 3G) et les environnements à plusieurs fournisseurs.

4.4 Interfaces et signalisation en 3G

Comme le montre la figure 1.3, il y a quelques interfaces et signalisations impliquées dans les systèmes 3G. Ici, nous en discuterons brièvement leurs principes.

4.4.1 Interfaces

Interface air U_u /WCDMA : U_u ou l'interface air du WCDMA est l'interface la plus importante dans les réseaux 3G. Elle travaille avec les principes du WCDMA où tous les utilisateurs sont assignés un code, qui varie avec la transaction. Suivant les indications du chapitre 2, chaque utilisateur a un code de propagation séparé. A la différence avec le GSM, ici chaque utilisateur emploie la même bande de fréquence.

Interface I_{ub} : C'est l'interface qui relie la station de base (BS) au RNC. Ceci est normalisé comme une interface ouverte, à la différence du GSM où l'interface entre le BTS et la BSC n'est pas ouverte. Ceci se compose des liens de signalisation et des points de terminaisons communs du trafic. Chacun de ces derniers points de terminaisons trafic est commandé par des liens de signalisation consacrés, et, est capable de commander plus d'une cellule. La signalisation d'interface d' I_{ub} , connue sous le nom de la partie d'application du noeud B (NBAP), a deux constituants importants. Ce sont NBAP communs (C-NBAP) et NBAP consacrés (D-NBAP), qui sont employés pour des liens de signalisation communs et dédiés respectivement.

Interface I_{ur} : C'est une interface unique dans les réseaux de WCDMA. Elle est entre deux RNC. Il n'y avait aucune interface dans le réseau GSM comme telle (c'est-à-dire entre deux BSC). Cette interface a été conçue pour soutenir la fonctionnalité inter RNC du soft-handover. Cette interface soutient également la mobilité entre les RNCs, le trafic et la gestion des ressources des canaux communs et dédiés (par exemple transfert des mesures entres cellules). Le protocole de signalisation utilisé pour cette interface est RNSAP (Radio Network System Application Part).

4.4.2 Signalisations

La signalisation dans les réseaux 3G est de trois plans : le plan de transport, le plan de contrôle et le plan d'utilisateur. Un réseau 3G est basé sur trois couches qui contiennent les flux de données (Modèle OSI). Ces trois couches sont : transport, contrôle et utilisateur suivant les indications de la figure 1.4.

Le plan de transport est le moyen de fournir la connexion entre un UE et le réseau (c'est-à-dire une interface air). Il contient trois couches : physique, de données et réseau. La couche physique est la couche de WCDMA (TDD/FDD), alors que la couche de données est responsable de la configuration, du maintien, protection, corrections d'erreurs ...etc. La couche de données contrôle également la couche physique. La couche réseau contient fondamentalement les fonctions qui sont exigées pour le plan de contrôle et de transport.

Le plan de contrôle contient la signalisation liée aux services manipulés par le réseau. L'interface I_{ub} est maintenue par le NBAP ; dans l'interface I_u c'est RANAP ; tandis que dans I_{ur} c'est RNSAP. Pour signaler entre l'application et la destination sur la couche physique, la signalisation du plan d'utilisateur est utilisée ; par exemple sur l'interface U_u c'est le DPDCH [1].

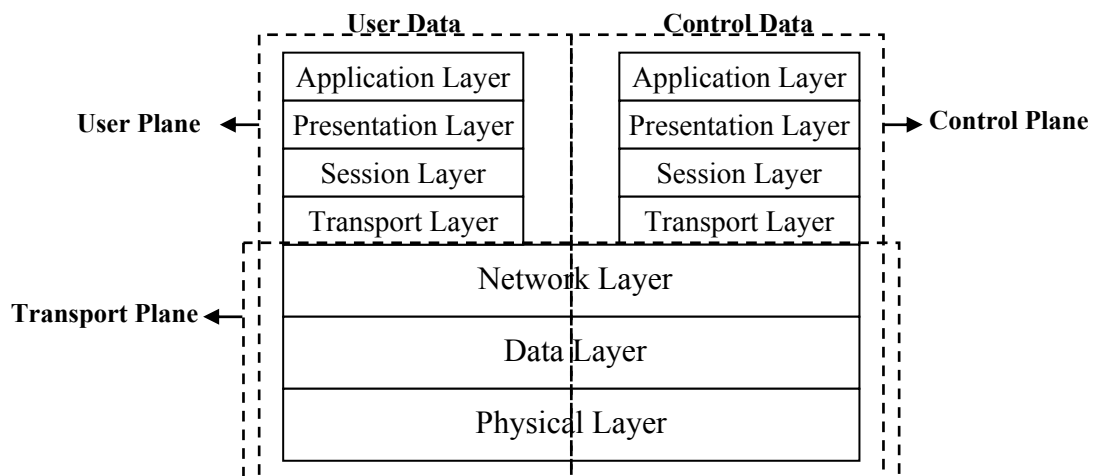


Fig.1.4 — Le modèle OSI en 3G

5. Modulations

Dans beaucoup de situations, il est nécessaire de transmettre des signaux numériques, en général sous la forme d'une séquence binaire. Les signaux numériques présentent en effet plusieurs propriétés intéressantes pour les télécommunications : souplesse de traitement, signal à états discrets donc moins sensibles aux bruits et simple à régénérer, utilisation de codes correcteurs d'erreur, cryptage de l'information ...etc. En

revanche, nous verrons ultérieurement qu'à quantité d'informations transmise identique, un signal numérique nécessite une bande de fréquence nettement plus importante.

Pour des raisons identiques au cas de signaux analogiques, ces signaux numériques modulent une porteuse sinusoïdale afin de présenter, soit des caractéristiques compatibles avec le canal de transmission utilisé (exemple des modems), soit pour transmettre plusieurs signaux simultanément. Toutefois, c'est l'explosion de la téléphonie mobile et de la télévision numérique qui suscite une étude de ce type de modulations.

5.1 Binary Phase Shift Keyed (BPSK)

La BPSK utilise une phase pour représenter un 1 et une autre pour représenter un 0 (binaire). Elle est utilisée pour transmettre jusqu'à 1 Mb/s.

5.2 Gaussian Minimum Phase-shift Keying (GMSK)

La GMSK est la modulation utilisée pour les signaux en GSM. C'est une méthode de modulation dérivée du déphasage minimum (MSK). Ainsi, elle est basée sur la modulation de fréquence. La modulation GMSK fonctionne avec deux gammes de fréquences et peut basculer facilement entre les deux. L'avantage principal de la GMSK est qu'elle ne contient pas de partie de la modulation d'amplitude, et la largeur de bande exigée de la fréquence de transmission est de 200 kHz, qui est une largeur de bande acceptable par les normes du GSM. C'est la modulation utilisée dans les réseaux GSM et GPRS.

5.3 Octagonal Phase-shift Keying (8-PSK)

La raison derrière le perfectionnement de manipulation de données dans les réseaux 2G tel que le GPRS est l'introduction de l'octogonale PSK (8-PSK). Dans cet arrangement, le signal modulé peut porter 3 bits par symbole modulé à travers le signal radio par rapport à 1 bit dans la modulation GMSK. Mais cette augmentation du flux de données est au coût d'une diminution de la sensibilité du signal radio. Par conséquent les débits les plus élevés sont fournis dans une couverture partielle. C'est la modulation utilisée dans les réseaux GPRS et EDGE.

5.4 Quarter Phase-Shift Keying (QPSK)

Dans la modulation PSK, la phase de la forme d'onde transmise est changée au lieu de sa fréquence, aussi, le nombre de changements de phase est de deux. Un pas vers l'avant est la condition que le nombre de changements de phase est plus de deux (c'est-à-dire quatre), qui est le cas avec la QPSK. Ceci permet à la porteuse de porter quatre bits au

lieu de deux, effectivement doublant la capacité de la porteuse. Pour cette raison, la QPSK est la modulation choisie dans des réseaux 3G ; CDMA [1].

5.5 Modulation par code (CDMA et étalement de spectre)

De même que deux porteuses en quadratures peuvent transmettre deux signaux distincts, il est possible de moduler le signal binaire utile de débit D , par un signal binaire pseudo-aléatoire et de débit D_a ($D_a \gg D$). En prenant des signaux pseudo aléatoires dits orthogonaux entre eux, il est possible de transmettre, sans brouillage, plusieurs signaux utiles simultanément.

Le spectre du signal modulé est alors très large ce qui permet de rendre le signal robuste vis-à-vis des brouillages apparaissant sur une bande étroite de fréquences ou dus aux échos. Ce type de modulation a été retenu pour les téléphonies cellulaires de troisième génération (UMTS).

5.6 Modulation par fréquences orthogonales (OFDM)

De même que deux porteuses en quadratures dans le domaine temporel, il est possible de générer des porteuses orthogonales dans le domaine spectral. Cette solution a été retenue pour la diffusion TV numérique terrestre. Le signal modulé est alors robuste vis-à-vis des brouillages dus aux échos. Cette technique est encore utilisée récemment dans les réseaux de 3.5G et 4G [3].

6. Techniques de duplexage

Un autre paramètre extrêmement important est le duplexage, c'est-à-dire la technique qui va permettre de faire la différence entre deux sens de liaison. On parlera en général, dans un système cellulaire, de liaison montante quand il s'agit de l'émission du mobile vers la station de base et de liaison descendante pour l'émission de la station de base vers le mobile. Nous commencerons par détailler ces techniques et les raisons des choix effectués et ensuite, nous détaillerons les techniques d'accès multiple.

Il y a deux techniques de duplexage :

6.1 FDD : Frequency Division Duplexing, duplexage en fréquence. Dans ce mode, les voies montante et descendante sont sur des fréquences bien distinctes. C'est le choix du GSM et de l'UMTS (Fig. 1.5).

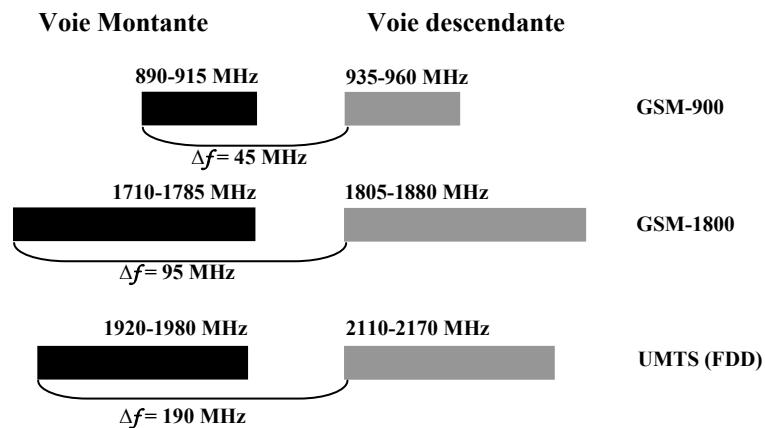


Fig.1.5 — Bandes GSM et UMTS (FDD)

6.2 TDD : Time Division Duplexing, duplexage temporel. Dans ce mode, les voies montante et descendante sont sur la même fréquence (Fig. 1.6), mais utilisent le canal alternativement (en général, d'abord la voie descendante et ensuite la voie montante). C'est le choix dans le cas des réseaux locaux sans fils. C'est aussi le cas d'une deuxième phase de l'UMTS pour les environnements urbains organisés en micro-cellules. Les deux avantages principaux du TDD sont [4] :

- Une plus grande simplicité de la partie RF, puisqu'on ne travaille que sur une fréquence à la fois (contre deux pour le choix FDD, une sur la voie montante, une sur la voie descendante).

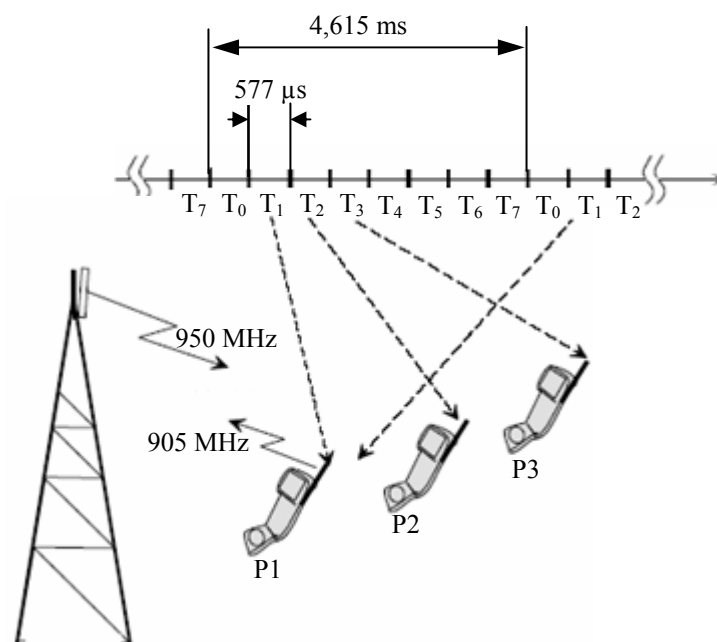


Fig.1.6 — Multiplexage temporel -GSM- (TDD)

- Le canal est réciproque (c'est-à-dire que la station de base voit le même canal de transmission que le terminal). De ce fait, l'adaptation de puissance est simple. (Adapter sa puissance d'émission pour avoir la bonne puissance à la réception et d'avoir donc le bon rapport signal/bruit au récepteur).

7. Techniques d'accès

Les principales techniques d'accès multiple utilisées en communications sans fil tirent le plus souvent leur origine dans les communications filaires, elles ont du être adaptées aux communications sans fils, dont les deux différences principales sont :

- Une bande passante limitée,
- Une communication non fiable (taux d'erreurs élevé, perte de lien,... etc.)

Le choix d'une technique d'accès (FDMA, TDMA, CDMA, ou une combinaison de deux d'entre elles) peut avoir un impact important sur les performances, la QoS (Quality of Service) et la capacité du système. Ce choix est tellement prépondérant (en tous cas dans l'esprit des concepteurs), qu'on dénomme souvent les systèmes en fonction de l'accès multiple [4].

7.1 L'accès multiple par répartition de fréquence (FDMA)

Dans un système FDMA pur, tous les utilisateurs peuvent transmettre leurs signaux simultanément, et sont distingués par leur fréquence d'émission (Fig. 1.7). Le FDMA est basé sur la plus ancienne technique de multiplexage connue : le multiplexage en fréquence, utilisé pour transmettre les signaux TV sur le câble, ou sur les canaux Hertziens classiques et satellite.

Le problème principal que le FDMA pose au concepteur des transceivers (émetteurs et récepteurs) est celui du canal adjacent. En effet, il faut absolument éviter de radier de la puissance hors de sa bande, sous peine de générer une interférence importante sur le canal occupant la fréquence voisine.

Ce problème est particulièrement important sur la voie montante. En effet, le signal reçu à la station de base par un mobile éloigné est nettement plus faible que le signal reçu à la station de base par un mobile proche. De ce fait, si le mobile proche génère une interférence importante, le signal du mobile éloigné risque d'être complètement noyé dans l'interférence générée par le mobile proche.

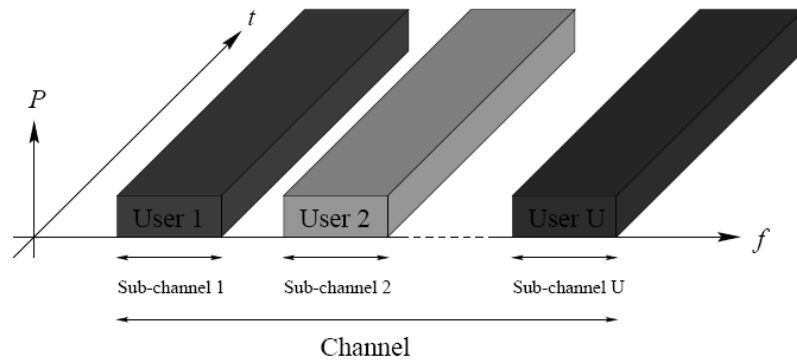


Fig.1.7 — L'accès multiple par répartition de fréquence (FDMA)

7.2 Accès multiple par répartition temporelle (TDMA)

Dans un système TDMA, les utilisateurs utilisent la même fréquence et prennent possession du canal chacun à leur tour.

Le TDMA est basé sur le multiplexage temporel utilisé par exemple en téléphonie, pour la concentration (numérique) des connexions entre centraux téléphoniques. L'avantage principal du TDMA est qu'il est facile pour un utilisateur de prendre possession de plusieurs tranches du multiplex temporel, et il est donc facile d'avoir des utilisateurs utilisant des débits de données différents. Le standard principal utilisant cette technique est le GSM, qui utilise une association de TDMA et de FDMA, sur une technique de duplexage TDD. Le DECT (Digital European Cordless Telephone) utilise également du TDMA/FDMA, mais avec une technique de duplexage TDD (puisque DECT est destiné aux micro-cellules d'environ 300 mètres de diamètre). De la même manière qu'en FDMA, le signal reçu sur la voie montante (c'est-à-dire à la station de base) peut être de puissance très différente pour un utilisateur éloigné et un utilisateur proche. Il convient donc également d'effectuer un contrôle de puissance.

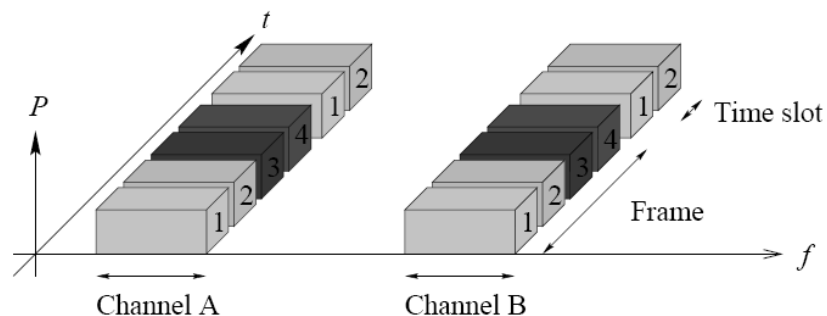


Fig.1.8 — L'accès multiple par répartition de temps (TDMA)

7.3 Accès multiple par répartition de codes (CDMA)

Le CDMA est basé sur la technique du spectre étalé, où chaque utilisateur se voit allouer un code PN (code aléatoire) différent, et est identifié par ce code. Outre que VITERBI, l'un des ardents défenseurs du CDMA pour la téléphonie mobile, et promoteur du premier standard (américain) IS-95, a démontré que le CDMA permettait une capacité accrue, l'une des raisons du succès du CDMA est sa souplesse. En effet, deux arguments en faveur du CDMA sont :

- L'absence du planning de fréquence. Si deux signaux ont des codes différents, il est possible de les séparer l'un de l'autre. Il suffit donc de faire en sorte que les utilisateurs des cellules adjacentes aient des codes différents pour les distinguer, et on a donc pas besoin de fréquences différentes pour chaque cellule.
- la facilité de mélanger des canaux voix et données. En effet, pour mélanger des canaux voix et données, il faut pouvoir mélanger des signaux de débit de données différents.

En spectre étalé, il suffit d'utiliser des signaux avec des facteurs d'étalement différents, ce qui est fait, par exemple, par les codes OVSF (Orthogonal Variable Spreading Factors) en UMTS.

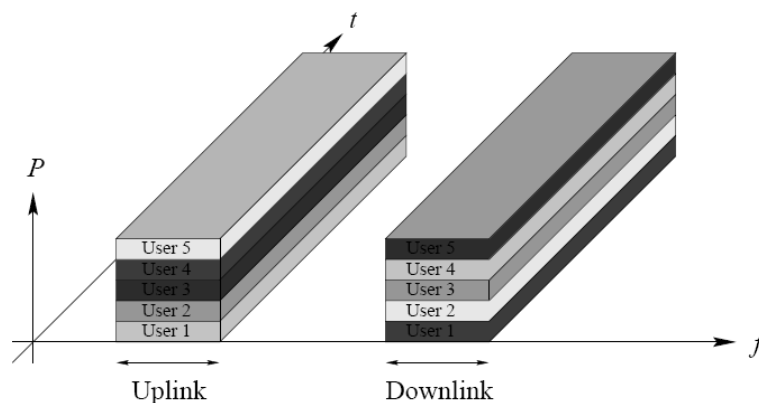


Fig.1.7 — L'accès multiple par répartition de code (CDMA)

7.4 Les performances comparées de TDMA/FDMA/CDMA

En première approximation (c'est-à-dire sans tenir compte des intervalles et/ou bandes de garde), les systèmes TDMA et FDMA ont les mêmes performances. En effet, les implications au niveau de l'interférence, et donc du planning de fréquences en cellulaire, est

le même. On peut alors se baser simplement sur l'efficacité spectrale dans une seule cellule, et que l'on répartisse les données en fréquence ou en temps ne change rien.

En CDMA, les choses sont un peu plus subtiles, puisqu'on peut utiliser les mêmes fréquences dans toutes les cellules. Par contre, on ne peut pas obtenir de codes orthogonaux pour les utilisateurs de cellules différentes (d'autant plus que ça demanderait de synchroniser les cellules entre elles, ce qui est très compliqué). Il faut donc faire une analyse complète (et statistique) des interférences, en fonction de la qualité des codes utilisés. En fonction de tout cela, certains prétendent que le CDMA permet une plus grande capacité.

Sans rentrer dans les détails, cette affirmation n'est correcte que si on utilise des algorithmes très puissants dans les récepteurs, ce qui n'est pas le cas pour le moment [4].

8. Conclusion

Nous avons vu dans ce chapitre les bases de communication dans les réseaux mobiles, leurs évolutions et leurs architectures, les techniques de modulations et de multiplexages appliquées, et plus principalement les techniques d'accès des utilisateurs.

Chaque réseau mobile à sa propre technique d'accès ; TDMA, FDMA, CDMA ou leurs fusions. De plus en plus, les réseaux de la dernière génération utilisent la technique CDMA qui représente plus de souplesse, et une facilité d'allocation pour plusieurs utilisateurs.

Le GSM repose sur les deux techniques TDMA et FDMA, mais l'UMTS repose sur la technique WCDMA l'une des variante du CDMA.

La technique CDMA et ses variantes seront développées dans le chapitre suivant.

DS-CDMA : Codage & détection

CHAPITRE 2

DANS CE CHAPITRE

- ☒ Introduction
- ☒ Etalement de spectre
- ☒ CDMA
- ☒ Codage
- ☒ Détection
- ☒ Conclusion

1. Introduction

Au début des années 20, Goldsmith utilisa la FM pour étaler le spectre d'une modulation AM afin de réduire l'effet du bruit et des trajets multiples sur une communication. Dès les années 1940, Wiener et Shannon ont développé des théories de traitement du signal dans le but d'utiliser la technique d'étalement de spectre à des fins de cryptage dévolu aux transmissions militaires. Dans plusieurs pays, l'armée met au point différents types d'étalement de spectre pour la guerre électronique. Les techniques restent très confidentielles. Les systèmes de l'époque, à tube, sont trop encombrants, chers et complexes pour être utilisés en télécommunication civile. L'utilisation de cette technique dans la radiolocalisation (GPS) lui permet d'entrer enfin dans le domaine public. Fin des années 90, l'étalement de spectre apparaît comme une solution de multiplexage d'avenir qui permet d'envisager les radiocommunications de masse [3].

2. Etalement du spectre

L'étalement de spectre (figure 2.1) est l'un des avantages mis en avant pour l'utilisation du CDMA dans le domaine des communications radiofréquences. En effet, la puissance d'un signal, après codage, est étalée sur toute la largeur de la bande de fréquence disponible. De ce fait deux caractéristiques importantes apparaissent :

- La puissance du signal étant étalée sur la bande spectrale disponible, le signal CDMA peut être confondu avec le bruit du canal et sera donc difficilement détectable par un utilisateur non concerné.
- Le signal CDMA étalé est plus résistant aux évanouissements sélectifs en fréquence.

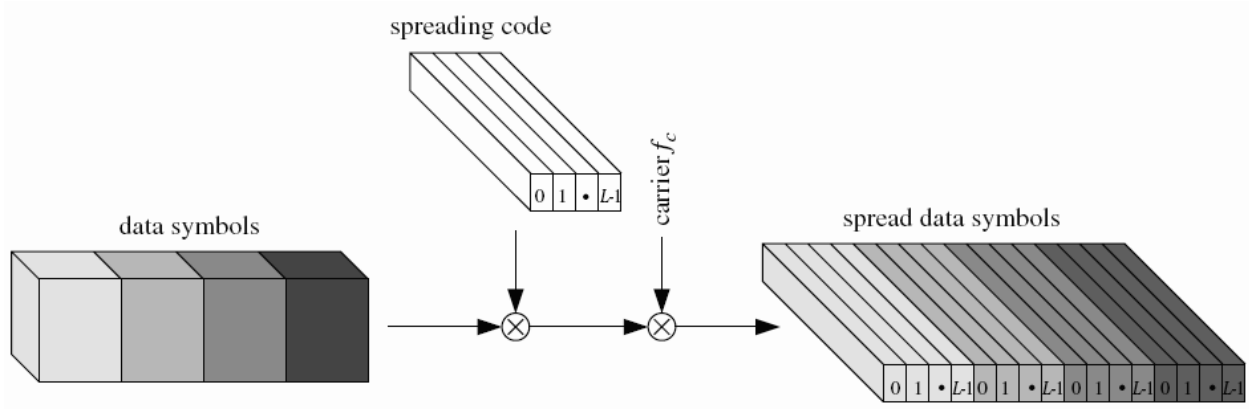


Fig.2.1 — Technique d'étalement de spectre

Cette technique consiste à multiplier chaque bit d'information par une séquence pseudo aléatoire binaire, notée PN, de rythme très supérieur à celui des données à transmettre, voir figure 2.2. En supposant que le rythme de la séquence d'étalement est N fois plus grand que le débit de données, il en résulte, en notant T_b la durée d'un bit d'information et T_c la durée d'un bit de la séquence d'étalement, que : $T_b = N T_c$ où N est appelé gain de traitement du système à étalement de spectre.

La figure 2.2 présente cette technique d'étalement de spectre par séquence directe sur l'émission du message binaire $d = (0; 1; 0)$ avec la séquence d'étalement PN = (0; 1; 0; 0; 1; 0; 1; 1): Pour simplifier les notations, nous considérons que les bits du message à transmettre d sont à valeurs dans $\{1, -1\}$ au lieu de $\{0, 1\}$. Les données à transmettre ($d_1; d_2; d_3$) = (1; -1; 1), représentées par le premier graphique, ont une durée d'émission de T_b . La séquence d'étalement PN est, quant à elle, représentée par le second graphique. Elle est répétée trois fois car, contrairement aux données à transmettre, les bits de la séquence PN ont une durée d'émission T_c plus petite que T_b . Dans notre exemple, nous avons $T_b = 8 T_c$. Le facteur d'étalement est donc $N = 8$. Ainsi, 8 bits de la séquence d'étalement PN sont nécessaires pour transmettre un bit de donnée d_i . Les bits transmis sur le canal sont donnés par le troisième graphique qui représente la multiplication des données par la séquence d'étalement PN ; ces bits ont une durée d'émission identique à celle des bits de la séquence PN, notée T_c . Ainsi, à chaque bit d'information peut correspondre, suivant la valeur du facteur d'étalement, soit une période de la séquence PN, soit une partie tronquée de cette période, soit plusieurs périodes.

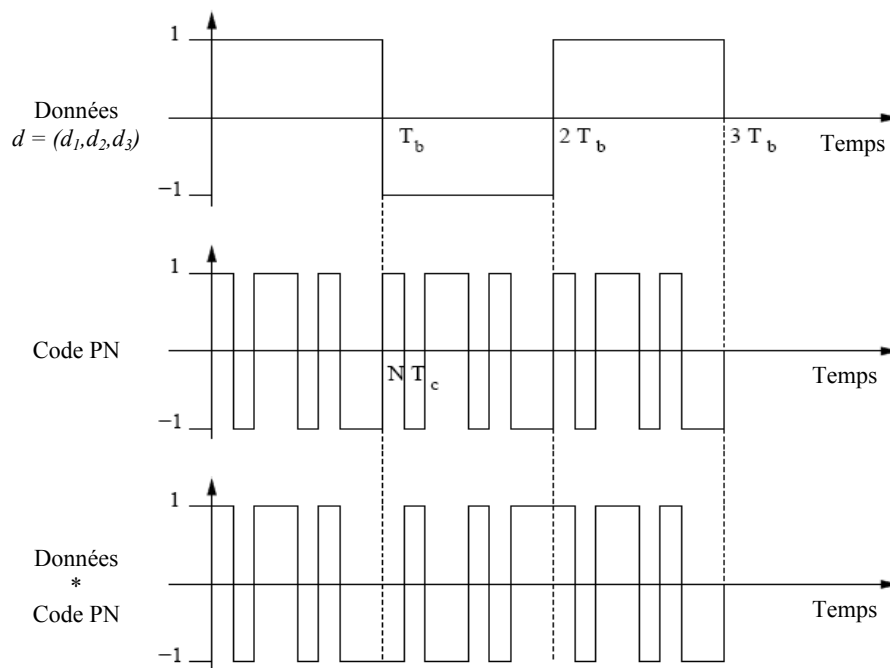


Fig.2.2 — Etalement de spectre par séquence directe (DSSS)

Dans le domaine fréquentiel, la notion d'étalement de spectre est représentée par la figure 2.3. La puissance du signal étalé, émis par étalement de spectre, est la même que celle du signal non étalé, émis sans étalement de spectre. Toutefois, elle n'est pas répartie sur la même largeur de bande de fréquence. On appelle alors un signal non étalé, un signal bande étroite, et un signal étalé, un signal large bande.

En réception, le signal est reconstruit en multipliant, dans un synchronisme parfait, le signal reçu par une séquence PN locale identique à celle utilisée à l'émission. La synchronisation de ces deux signaux est plus ou moins délicate suivant la technique utilisée pour le calcul de la corrélation [5].

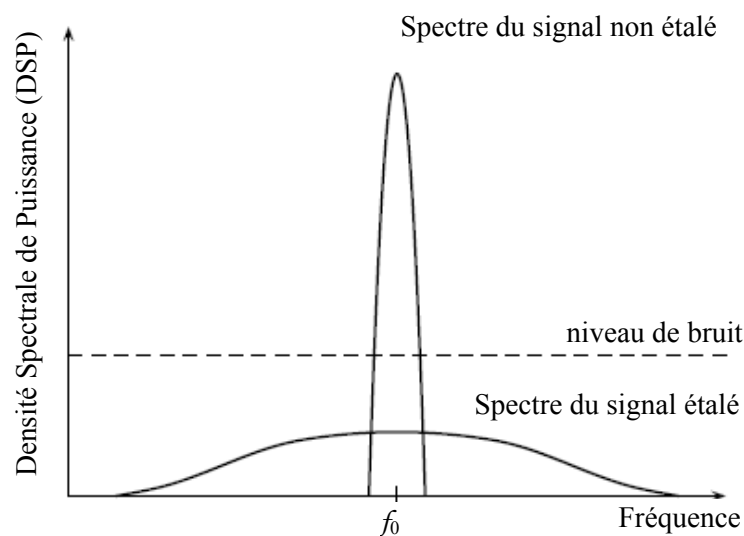


Fig.2.3 — Comparaison signal à bande étroite / signal à bande étalée

2.1 Pourquoi étaler le spectre ?

Considérons le théorème de Shannon et Hartley concernant la capacité d'un canal de communication :

$$C = B \log_2 \left(1 + \frac{S}{N} \right) \quad (2.1)$$

Dans cette équation, C représente la capacité maximale d'un canal en bits par seconde (bit/s), c'est le taux de transfert maximum pour un taux d'erreur binaire (Bit Error Rate, BER) nul, à condition qu'un procédé de codage adéquat puisse être trouvé. B étant la bande passante du canal en Hertz et $\frac{S}{N}$ le rapport de puissance signal/bruit.

On peut donc augmenter la capacité maximale en agissant sur la largeur de bande de façon linéaire et/ou en agissant sur le rapport/signal bruit de façon logarithmique.

A capacité maximale donnée (capacité maximale souhaitée), on peut réduire la bande et/ou diminuer le rapport signal/bruit en admettant un taux d'erreur non nul. Les erreurs peuvent être soit tout simplement ignorées soit corrigées par l'utilisation de protocoles de transmission de niveau supérieur. Au niveau de la formule, en fonction du type de bruit et du procédé de codage/décodage, on peut intégrer le BER sous la forme de l'addition d'une constante au rapport signal sur bruit.

Dans le cas du CDMA, le bruit est constitué principalement par les autres utilisateurs dont on cherchera à augmenter le nombre. Il en résulte qu'en règle générale un système CDMA opère sur des rapports signal/bruit faibles, voire très faibles. Par changement de base des logarithmes (base 2 vers base e), l'équation (2.1) devient :

$$\frac{C}{B} = \frac{1}{\ln(2)} \ln \left(1 + \frac{S}{N} \right) = 1.443 \ln \left(1 + \frac{S}{N} \right) \quad (2.2)$$

Si la puissance du signal est inférieure à la puissance du bruit, on peut simplifier et linéariser l'expression (2.1), en appliquant le développement en série de MacLaurin de $\ln(1+x)$:

$$\frac{C}{B} = 1.443 \left[\frac{S}{N} - \frac{1}{2} \left(\frac{S}{N} \right)^2 + \frac{1}{3} \left(\frac{S}{N} \right)^3 - \dots \right] \quad (2.3)$$

Puisque l'étalement du spectre permet un rapport $\frac{S}{N}$ très faible et que la puissance du signal utile pouvant être inférieure au niveau du bruit. Pour un $\frac{S}{N} \ll 1$, l'équation (2.1) devient alors :

$$\frac{C}{B} = 1.443 \left(\frac{S}{N} \right) \quad (2.4)$$

et par approximation on obtient :

$$\frac{C}{B} \approx \frac{S}{N} \quad (2.5)$$

La dépendance capacité/bande passante et signal/bruit est approximativement linéaire.

La bande étalée permet donc la transmission de signaux perturbés par d'autres signaux considérés alors comme du bruit, c'est-à-dire la transmission de signaux sur le même support. Le nombre de canaux utilisés à un instant donné pourra varier de façon souple puisque l'augmentation du nombre d'utilisateurs se traduira simplement par une augmentation, pour tous, du taux d'erreur. Ceci permet en téléphonie de maintenir une

qualité de service sensiblement égale pour tous, (plutôt qu'une dépréciation totale pour un utilisateur) ajustable et relativement facile [6].

2.2 Avantages et inconvénients

Comme nous l'avons vu, la technique d'étalement de spectre consiste à moduler le signal contenant l'information puis à l'étaler de manière à ce que le spectre du signal émis occupe une bande de fréquence très supérieure à celle nécessaire à la transmission de l'information. L'étalement de spectre, par rapport aux modulations à bande étroite, présente de nombreux avantages [5 et 7], voir :

- Bonne résistance aux perturbations bande étroite,
- Faible brouillage des émissions classiques à bande étroite,
- Insensibilité aux effets des trajets multiples,
- Faible probabilité d'interception,
- Multiplexage et adressage sélectif,
- Sécurité des communications.

Quelques inconvénients sont liés à cette technique :

- Encombrement spectral important qui rend souvent l'attribution de fréquences difficile. En effet, le signal a toujours la même puissance mais celle-ci est répartie différemment,
- Complexité accrue des systèmes qui rend leur coût plus élevé par rapport à celui des systèmes bande étroite,
- Nécessité d'avoir de bonnes méthodes de synchronisation permettant à la réception, de reconstruire le signal émis.

3. CDMA

Le CDMA permet aux différents utilisateurs de transmettre leurs données sur n'importe quelle fréquence et sans nécessiter de synchronisation entre eux. En effet, contrairement aux techniques TDMA et FDMA, la capacité de multiplexage du CDMA n'est pas limitée par des paramètres physiques (intervalles de temps disponibles, fréquences ou longueurs d'ondes utilisables...etc.) mais par la capacité à générer un maximum de séquences de codes, celles-ci étant choisies de manières à minimiser les Interférences d'Accès Multiple.

Les séquences de codes utilisées dans les systèmes CDMA sont composées d'une série d'impulsions nommées « chips » afin d'être distinguées des « bits » qui composent une séquence de données.

La figure 2.4 présente le principe du CDMA : les signaux provenant de M utilisateurs sont transmis autour d'une porteuse de fréquence f_0 . Nous avons à l'émission « l'empilement » des M spectres des signaux émis autour de la fréquence f_0 . En réception, seul le spectre du signal utile de l'utilisateur A est retenu, puisque les M-1 autres signaux sont à nouveau étalés par la séquence pseudo-aléatoire de l'utilisateur A.

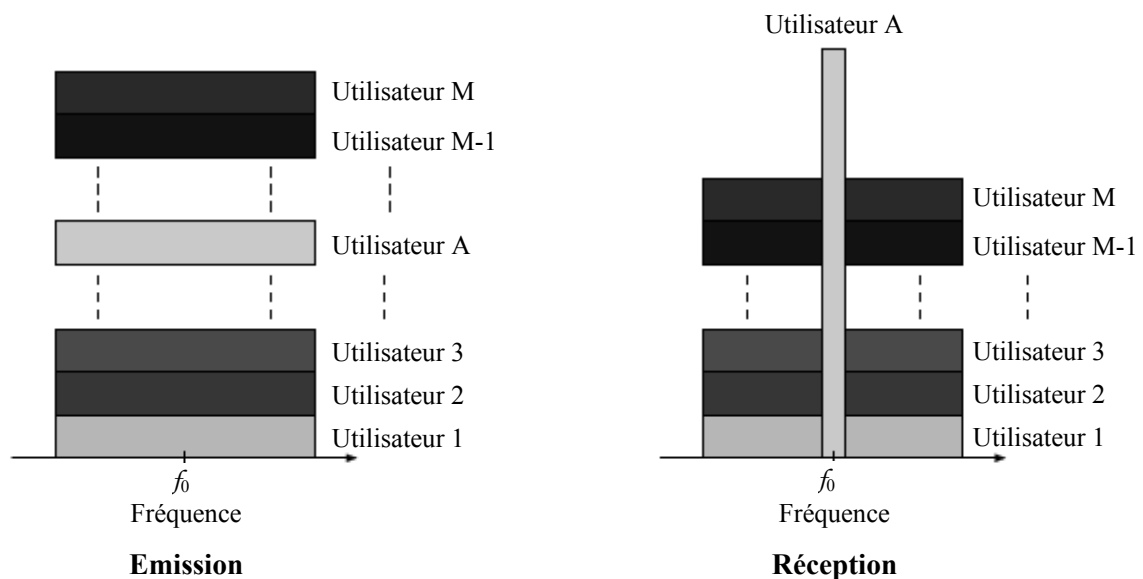


Fig.2.4 — Principe de la technique CDMA

Il existe deux types de CDMA : le CDMA synchrone (S-CDMA) et le CDMA asynchrone (A-CDMA). Dans le cas d'une communication station de base/utilisateur, le CDMA synchrone correspond à une communication de la station de base vers les utilisateurs, tandis que le CDMA asynchrone correspond à une communication de l'utilisateur vers la station de base.

Dans le cas du CDMA synchrone, l'orthogonalité des codes d'étalement est exploitée au maximum, tandis que dans le second cas, CDMA asynchrone, les décalages temporels existant entre les utilisateurs ne permettent pas l'utilisation de l'orthogonalité des codes d'étalement.

La technique CDMA englobe plusieurs types comme le montre la figure 2.5, pures et hybrides (la combinaison de plusieurs techniques en même temps citées dans le chapitre précédent, figure 2.6) :

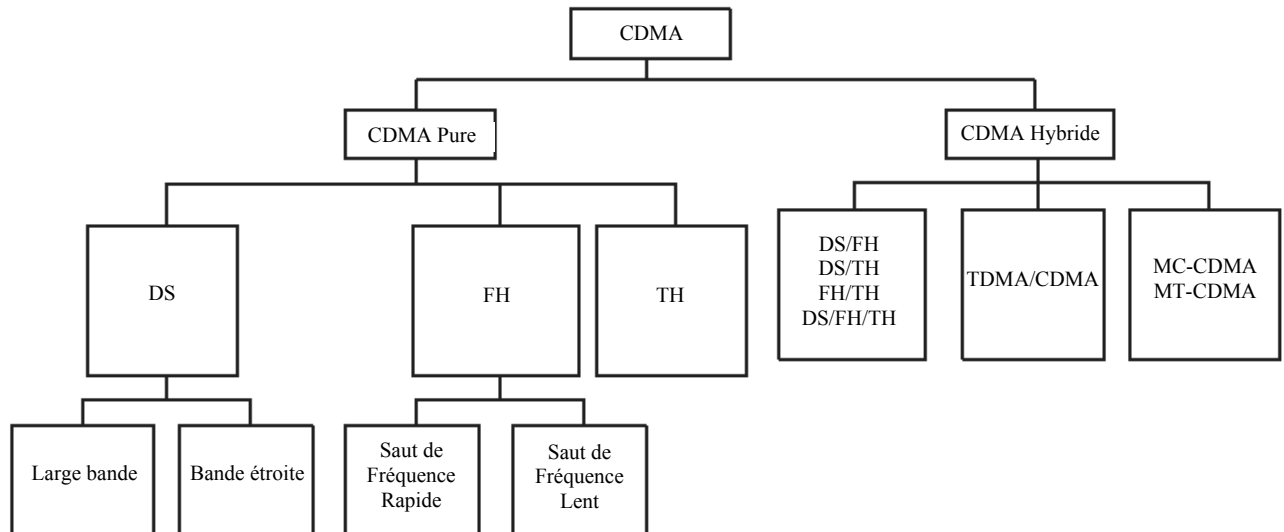


Fig.2.5 — Différents types de CDMA

3.1 FH-CDMA (Frequency Hopping CDMA)

Dans ce système, on fait de l'évasion de fréquence : la clé de chaque code utilisateur est réalisée pour une suite de fréquences qui feront alternativement office de porteuse. Ce système ressemble à un multiplexage fréquentiel dans lequel l'attribution des fréquences varierait rapidement (par rapport au débit d'informations à transmettre).

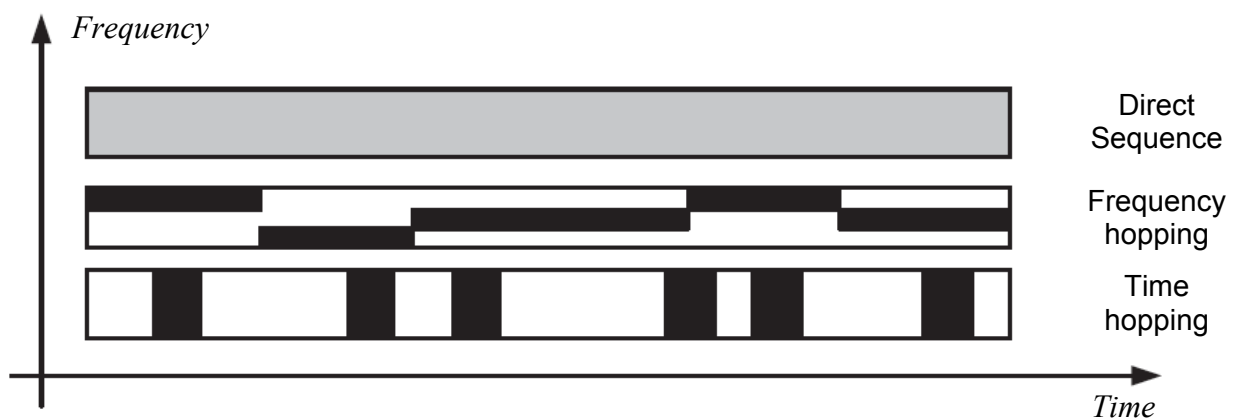


Fig.2.6 — Comparaison des techniques DS, FH et TH-CDMA

3.2 TH-CDMA (Time Hopping CDMA)

C'est un genre de multiplexage temporel ; le signal d'information est transmis avec interruption ou rafale (figure 2.6), le temps de transmission est indiqué par le code d'étalement. La technique a été développée à la fin des années 1940 comme première méthode de CDMA, et a été utilisée à des fins militaires.

3.3 DS-CDMA (Direct Sequence CDMA)

La technique de DS-CDMA est de loin la technique la plus populaire dans toutes les applications de communication de CDMA. Le tableau 2.1 montre que la DS-CDMA est la technique dominante dans presque tous les systèmes de 3G [8, 9, 10 et 11]. Il y a plusieurs raisons à sa popularité. D'abord, c'est la forme la plus simple de la technique CDMA et peut être mise en application relativement à un coût bas comparée à d'autres techniques, telles que FH-CDMA et TH-CDMA. Toutes les normes actuellement disponibles en communication mobile de la 3G sont basées sur la technique DS-CDMA à presque aucune exception. En second lieu, un système DS-CDMA peut également être conçu pour fonctionner compatiblement avec beaucoup de réseaux de transmission existants qui fonctionnent basiquement sur d'autres technologies d'accès multiple, telles que l'accès multiple temporel (TDMA), pour réaliser la soi-disant évolution sans collision. La proposition d'un système 3G ou d'un WCDMA a été basée principalement sur cette clause comme effort pour réaliser l'évolution sans heurt du système 2G travaillant sur la technologie TDMA (GSM) [12].

Standard	Bande de Fréquence (MHz)	Débit (bps)	Technique d'accès	Facteur d'étalement
IS-95	824-894 896-894	1.2288M	DS-CDMA	256
BLEUTOOTH	2400-2483.5	1M	FH-CDMA	79
UMTS	1900-2025 2110-2200	3.84M	DS-CDMA	4,8,...,256
CDMA2000	824-894 869-894	1.22883M 3.6864M	DS-CDMA	4,8,...,128 4,8,...,256
WLAN	2400-2484	11M	DS-CDMA	13
ZIGBEE	868-868.6 902-928 2400-2483.5	20K 40K 250K	DS-CDMA	1 10 16

Tableau.2.1 — Caractéristiques de quelques standards de télécommunication

Comme le montre la figure 2.2, un signal binaire modulé en phase BPSK (Binary Phase Shift Keying) est codé par une séquence pseudo-aléatoire ou Pseudo Noise (PN). Le résultat de ce codage est représenté par le signal produit des deux. Ce dernier est superposé aux autres signaux provenant des autres émetteurs et ayant subi un traitement similaire et est transporté par le canal de transmission (figure 2.7).

Le codage des données s'effectue donc de manière « directe », sans faire intervenir d'autres paramètres comme la fréquence ou la longueur d'onde.

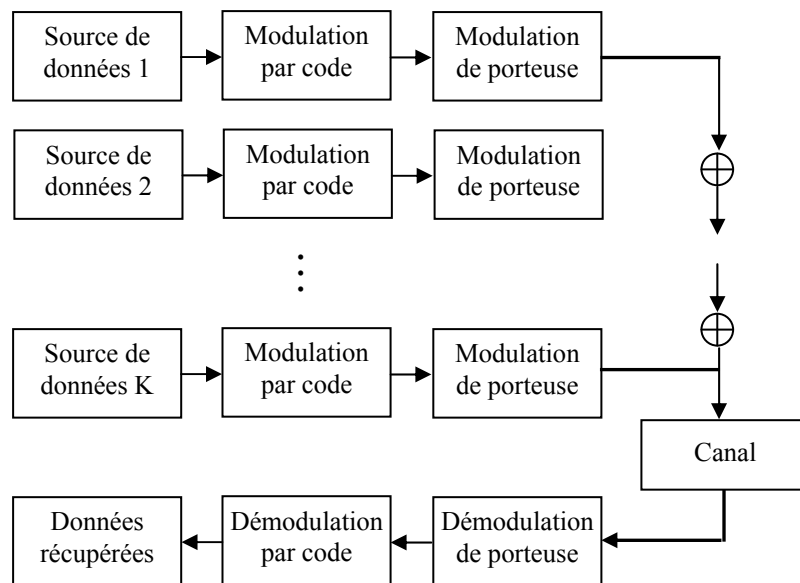


Fig.2.7 — Technique DS-CDMA (Uplink)

3.3.1 Avantages [5]

- Densité spectrale faible => non détectable,
- Densité spectrale constante et continue,
- Lutte contre les brouilleurs : Gain d'étalement,
- Plusieurs utilisateurs non coordonnés,
- Canal dispersif => Récepteur RAKE,
- Capacité Soft.

3.3.2 Inconvénients [5]

- Synchronisation difficile,
- Non synchronisation entre utilisateurs différents => Brouillages,
- Complexité du récepteur,

- Gain d'étalement limité,
- Difficulté de supprimer des bandes de fréquence,
- Interférences entre utilisateurs de la même cellule.

4. Codage

Plusieurs émissions peuvent cohabiter dans la même bande de fréquence dans la mesure où les codes d'étalement relatifs à chacun des signaux sont orthogonaux, c'est-à-dire dans la mesure où ils présentent une inter corrélation voisine de zéro (code de Hadamard, code de Walsh, code de Gold) [5]. La séquence d'étalement affectée à chaque signal constitue sa clé de codage. Ce signal ne peut être exploité que si le récepteur possède la même clé de codage.

Dans la plupart des systèmes de réseaux mobiles fondés sur le DS-CDMA, les codes utilisés pour effectuer le processus d'étalement sont de deux types : codes orthogonaux et codes pseudo aléatoires. Les deux genres de codes sont employés ensemble dans la liaison montante et descendante (Figure 2.8). Le même code est toujours employé pour l'étalement et le désétalement d'un signal. C'est possible parce que le processus d'étalement est réellement une opération XOR entre le flux de données et le code de l'étalement.

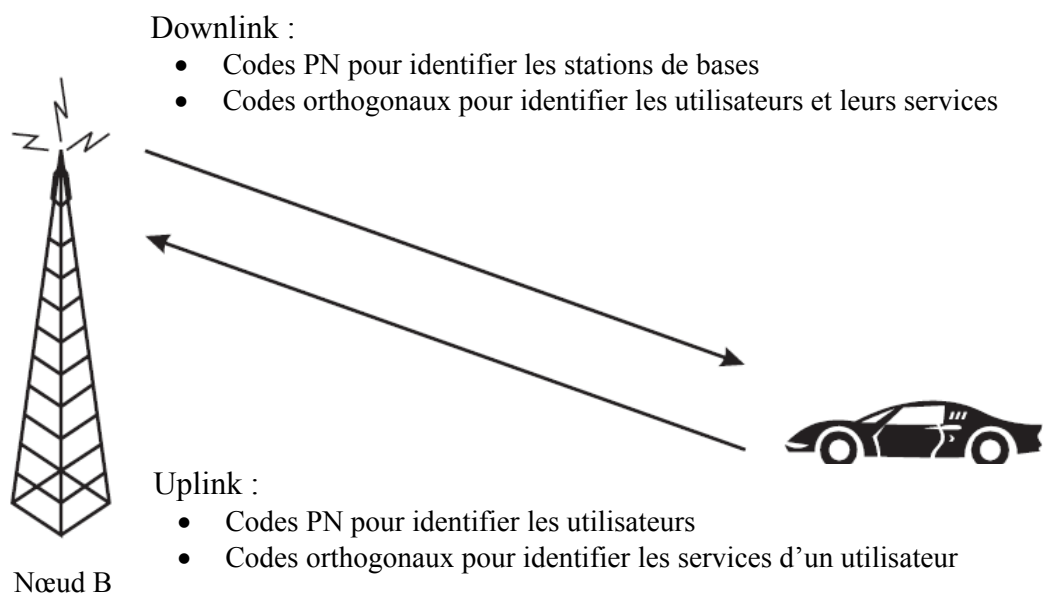


Fig.2.8 — Types de codes utilisés dans l'interface air

Le choix des codes d'étalement est dicté par leurs propriétés de corrélation, et plus précisément par leurs propriétés d'auto corrélation et d'inter corrélation [13]. Du point de vue

statistique, l'auto corrélation est une mesure de la correspondance entre un code et une version décalée de celui-ci. Par ailleurs, l'inter corrélation représente le degré de correspondance entre deux codes différents.

Le procédé de l'étalement dans l'UTRAN se compose de deux opérations séparées : canalisation et brouillage. La canalisation (spreading) emploie des codes orthogonaux et le brouillage (scrambling) emploie des codes PN (figure 2.9). La canalisation se produit avant le brouillage dans l'émetteur dans la liaison montante et l'inverse dans la liaison descendante [14].

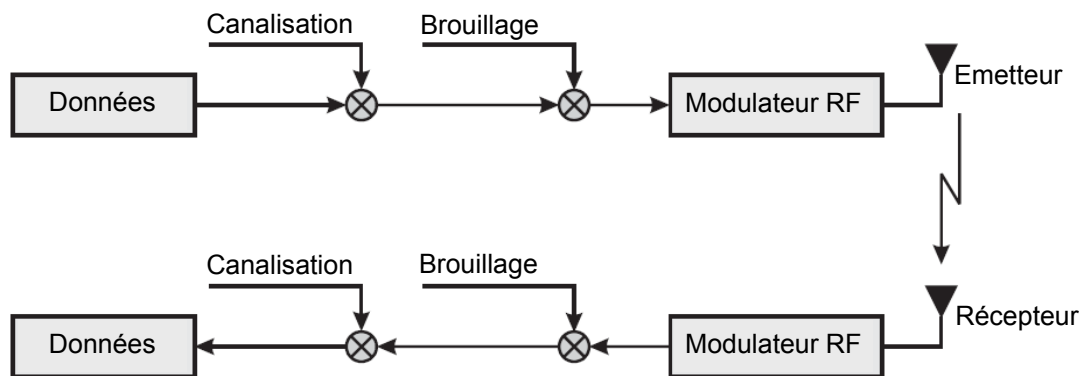


Fig.2.9 — Canalisation et brouillage

4.1 Codes Orthogonaux

La canalisation transforme chaque symbole de données en morceaux multiples. Ce rapport (nombre de chips/symbole) s'appelle le facteur de propagation **SF** (figure 2.10). Ainsi, c'est ce procédé qui augmente réellement la largeur de bande de signal.

La canalisation des codes (utilisation des codes orthogonaux) signifie que dans un environnement idéal les différents signaux n'interfèrent pas les uns sur les autres. Cependant, l'orthogonalité exige que les codes soient synchronisés. Par conséquent, elle peut être employée dans la liaison descendante pour séparer différents utilisateurs à moins d'une cellule, mais dans la liaison montante pour séparer seulement les différents services d'un utilisateur. Elle ne peut pas être employée pour séparer les différents utilisateurs de la liaison montante dans une station de base, car tous les mobiles ne sont pas synchronisés [14].

Comme son nom l'indique, ces codes, connus comme les codes de Walsh-Hadamard [15] sont mutuellement orthogonaux ; par conséquent leurs inter-corrélations sont théoriquement zéro. Mais, s'ils sont asynchrones, leurs inter-corrélations sont fortement dépendantes des paires des codes utilisées (entre zéro et une grande corrélation). En

WCDMA, ils sont connus comme codes d'étalement à facteur orthogonal variable (OVSF) et sont utilisés pour la canalisation dans les deux canaux montant (uplink) et descendant (downlink) [16].

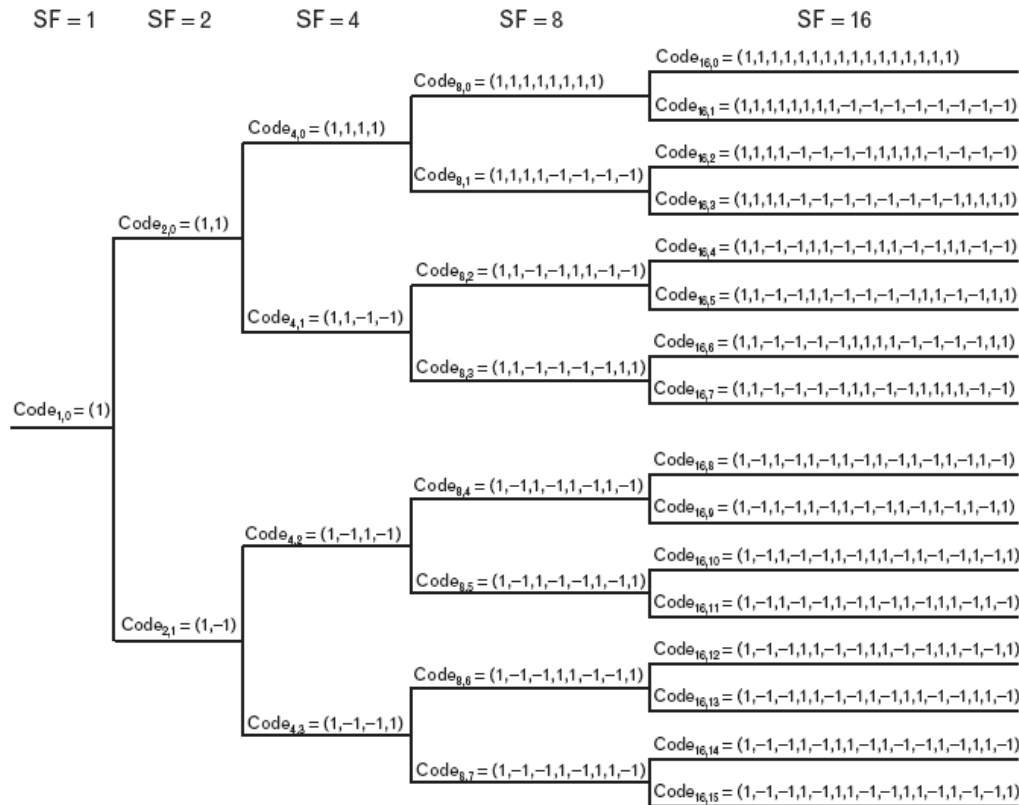


Fig.2.10 — Codes d'étalement à facteur orthogonal variable (OVSF)

4.2 Codes pseudo aléatoires (PN codes)

Un code pseudo aléatoire est une séquence de nombres binaires, qui apparaît être aléatoire, mais en réalité parfaitement déterministe [16]. La raison derrière l'utilisation des codes PN est : si les séquences de codes sont déterministes ; alors tout le monde peut avoir accès au canal. Par contre, si les séquences de codes sont vraiment aléatoires personne ne peut avoir accès au canal. Par conséquent, utiliser une séquence pseudo aléatoire rend le signal comme un bruit aléatoire pour tout le monde sauf pour l'émetteur et le récepteur voulu [17]. Les codes PN les plus utilisés sont :

- Codes à longueur maximale (M-codes),
- Codes Gold,
- Codes Kasami.

5. Détection

Le type d'émetteur-récepteur, utilisant l'étalement de spectre, est principalement utilisé dans les systèmes CDMA, c'est-à-dire dans un environnement multi-utilisateurs. Il se distingue aussi d'autres systèmes, comme le MBOK (M-ary Bi-Orthogonal Keying) [18 et 19] et le CCK (Complementary Code Keying) [20, 21 et 22], qui sont contraints de fonctionner simplement en mono utilisateur pour avoir des performances optimales.

Dans le CDMA, le MAI et le ISI sont deux facteurs qui limitent la capacité et l'exécution des systèmes DS-CDMA ; cependant l'exécution faible du détecteur conventionnel en atténuant ces interférences, et les besoins d'augmenter le flux de données dans de tels systèmes, ont mené à des recherches élaborées pour présenter des détecteurs fiables, susceptibles d'atténuer l'effet des interférences de MAI et ISI.

La classification de ces détecteurs est présentée dans la figure 2.11.

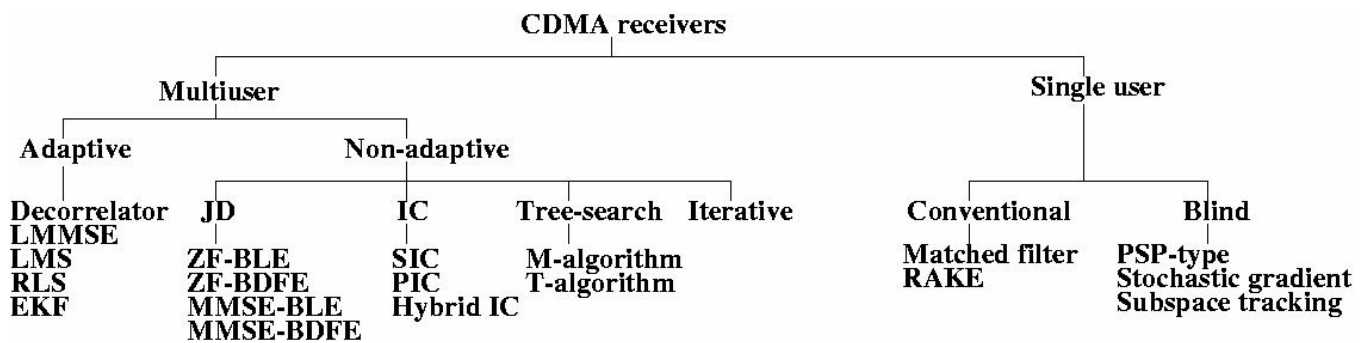


Fig.2.11 — Classification des détecteurs DS-CDMA

5.1 Détection mono-utilisateur

Dans la détection mono-utilisateur chaque signal d'un utilisateur désiré est détecté séparément sans prendre en compte les autres signaux d'utilisateurs. En d'autres termes aucune connaissance ou évaluation précise antérieure des paramètres des autres utilisateurs n'est exigée. Par conséquent, la technique à mono utilisateur est plus appropriée et pratique pour le cas de la liaison descendante (Downlink) due aux considérations de complexité et d'intimité. Plusieurs détecteurs ont été proposés dans la littérature pour le cas d'un mono utilisateur, tel que le récepteur RAKE, le MMSE et le MBER.

5.2 Détection multi-utilisateurs

Un récepteur multi-utilisateurs fait une évaluation basée sur l'observation de tous les signaux reçus par les utilisateurs, essayant d'évaluer l'interférence qu'ils créent au-dessus du signal de l'utilisateur désiré.

Un nombre significatif de paramètres de tous les utilisateurs actifs tels que la synchronisation du temps et de phase, séquence de propagation, puissance et état du canal devraient être connus ou estimés a priori.

L'utilisation de ce type de détection offre des avantages potentiels tels que l'amélioration significative de rendement, une utilisation du spectre plus efficace et une résistance d'effet du handover (near far effect). Ces gains d'exécution ont un effet positif sur l'énergie (augmentation de la vie de la batterie)

La détection multi-utilisateurs peut être classifiée dans deux catégories principales (figure 2.11) : Adaptative ou non adaptative [23].

6. Conclusion

CDMA est une technique d'accès multiple basée sur le procédé d'étalement de spectre, étalant ainsi le signal du message sur une bande relativement large en utilisant un code unique qui réduit les interférences et différencie entre les utilisateurs. Le CDMA n'a pas besoin de la répartition en temps ou en fréquence pour l'accès multiple. A noter que la variante DS-CDMA est la plus utilisée dans la plupart des systèmes 3G, elle améliore la capacité du système de communication.

Dans nos prochains chapitres on se basera sur cette technique et on essaiera d'améliorer et d'optimiser ses ressources.

Problèmes & égalisation du canal

CHAPITRE

3

DANS CE CHAPITRE

- ☒ Introduction
- ☒ Egalisation
- ☒ Simulations & résultats
- ☒ Conclusion

1. Introduction

En parcourant un trajet entre l'émetteur et le récepteur, le signal transmis est sujet à de nombreux phénomènes dont la plupart ont souvent un effet dégradant sur la qualité du signal. Cette dégradation se traduit en pratique par des erreurs dans les messages reçus qui entraînent des pertes d'informations pour l'utilisateur ou le système. Les dégradations du signal dues à la propagation en environnement mobile peuvent être classées en différentes catégories [6]:

- Les pertes de propagation dues à la distance parcourue par l'onde radio, ou l'affaiblissement de parcours (pathloss).
- Les atténuations de puissance du signal dues aux effets de masques (shadowing) provoqués par les obstacles rencontrés par le signal sur le trajet parcouru entre l'émetteur et le récepteur.
- Les atténuations de puissance du signal dues aux effets induits par le phénomène des trajets multiples.
- Les brouillages dus aux interférences créées par d'autres émissions. Ce type de pertes est très important dans les systèmes à réutilisation de fréquences.
- Les brouillages dus au bruit ambiant provenant d'émissions d'autres systèmes.

Les sources principales de dégradation dans les systèmes CDMA sont :

1.1 Interférence d'accès multiples (MAI)

L'interférence d'accès multiples (MAI) se produit quand un certain nombre d'utilisateurs partage un canal commun simultanément. Dans ce contexte, les signaux d'autres utilisateurs apparaissent comme une interférence pour un utilisateur donné.

1.2 Etats par trajets multiples de canal

Les états par trajet multiples de canal apparaissent dans les environnements où il y a réflexion, réfraction et dispersion des ondes radio par des bâtiments et d'autres obstacles synthétiques. Le signal transmis atteint le plus souvent le récepteur par l'intermédiaire de plus d'un chemin (figure 3.1).

1.3 Interférence d'inter symboles (ISI)

L'interférence d'inter symboles (ISI) se produit puisque, dans les systèmes de largeur de bande efficaces, l'effet de chaque symbole ou chip transmis au dessus d'un canal

dispersif se prolonge au delà de l'intervalle de temps employé pour représenter ce symbole. La déformation provoquée par le chevauchement résultant des symboles reçus s'appelle ISI (figure 3.2).

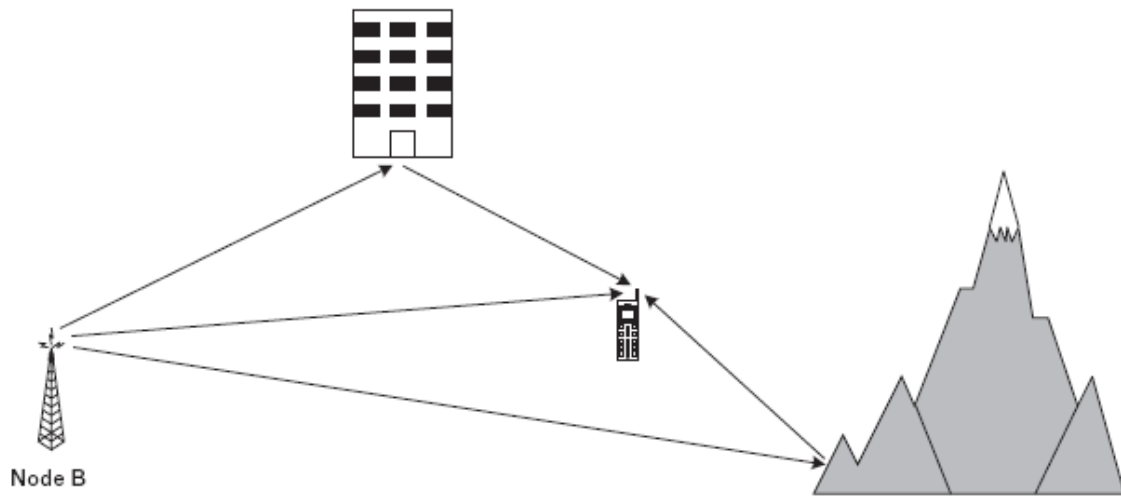


Fig.3.1 — Multi-trajets

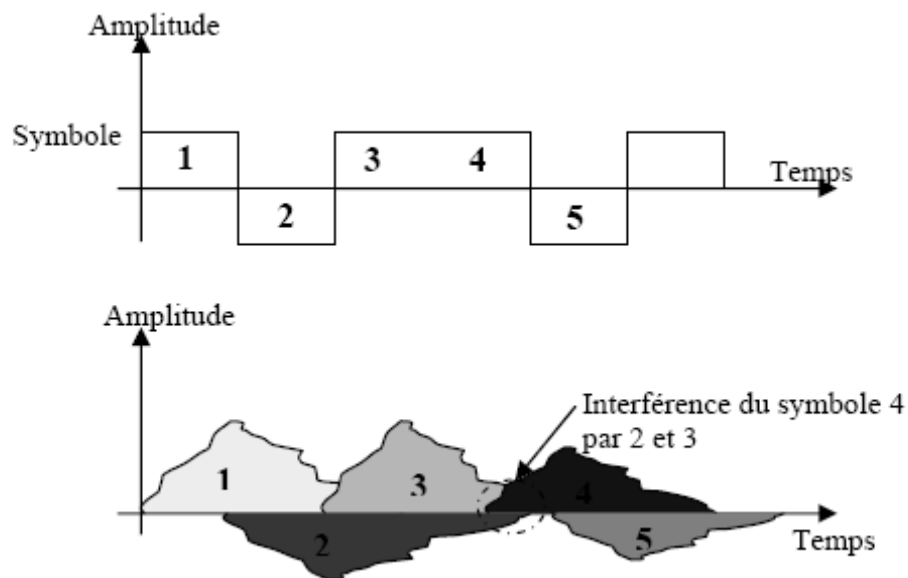


Fig.3.2 — Interférences inter-symboles

Ces sources perturbent l'orthogonalité parmi les codes des différents utilisateurs qui sont employés pour multiplexer les divers utilisateurs dans le système. Les récepteurs classiques (matched filter) ne sont pas efficaces pour combattre ces diverses sources d'interférence. Pour cela la suppression d'interférence a été l'objet de beaucoup d'efforts de recherche ; plusieurs techniques sont proposées, dont l'égalisation de canal [24].

2. Egalisation

Si le canal de transmission avait une atténuation constante et un déphasage linéaire sur la bande du signal, il ne modifierait pas la forme des impulsions émises et le récepteur recevrait tout simplement une version bruitée du signal émis. En pratique, ces deux conditions ne sont que très rarement vérifiées et la réponse du canal a besoin d'être égalisée pour éliminer la distorsion du signal reçu.

L'égalisation est une procédure qui est utilisée par les récepteurs des systèmes de communications numériques afin de réduire l'effet d'interférence entre symboles (ISI), due à la propagation du signal modulé à travers le canal. Pour empêcher une dégradation sévère des performances, il est nécessaire de compenser la distorsion des canaux. Cette compensation est effectuée par un filtre adaptatif.

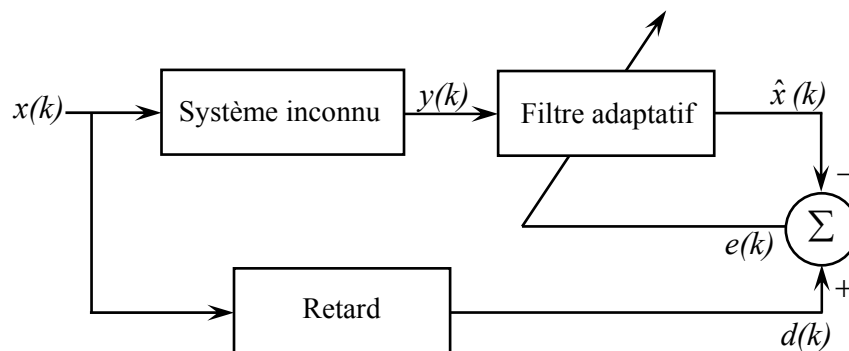


Fig.3.3 — Egalisation du canal

Cascader le filtre adaptatif avec le canal (système inconnu) fait converger le filtre adaptatif à une solution qui est l'inverse du canal. Si la fonction de transfert du canal est $H(z)$ et la fonction de transfert du filtre adaptatif est $G(z)$ (figure 3.3), l'erreur à mesurer entre le signal désiré et le signal du système cascadié atteint son minimum quand le produit de $H(z)$ et $G(z)$ est 1, $\mathbf{G(z) * H(z) = 1}$. Pour que cette relation soit vraie, $G(z)$ doit évaluer $1 / H(z)$, l'inverse de la fonction de transfert du canal.

Par ailleurs, la réponse du canal est en général inconnue et, de plus, susceptible de varier au cours du temps. Son égalisation nécessite alors un égaliseur adaptatif capable de s'adapter au canal et de poursuivre ses variations temporelles [24 et 25].

Il existe deux types d'égalisations :

- Egalisation avec apprentissage (Training)
- Egalisation sans apprentissage : Aveugle (Blind)

2.1 Egalisation avec apprentissage

Dans le cas où les propriétés du canal sont connues à l'avance, il est possible d'initialiser le filtre avant que la transmission commence. Si le canal n'est pas maîtrisé, les paramètres de l'égaliseur doivent être estimés. Les méthodes standards utilisent ce qu'on appelle des séquences d'apprentissage : l'émetteur émet une séquence de symboles connus par le récepteur. Ce procédé est appelé l'apprentissage (figure 3.4). Après l'estimation des paramètres, l'égaliseur est construit et l'égalisation du signal peut commencer [24].

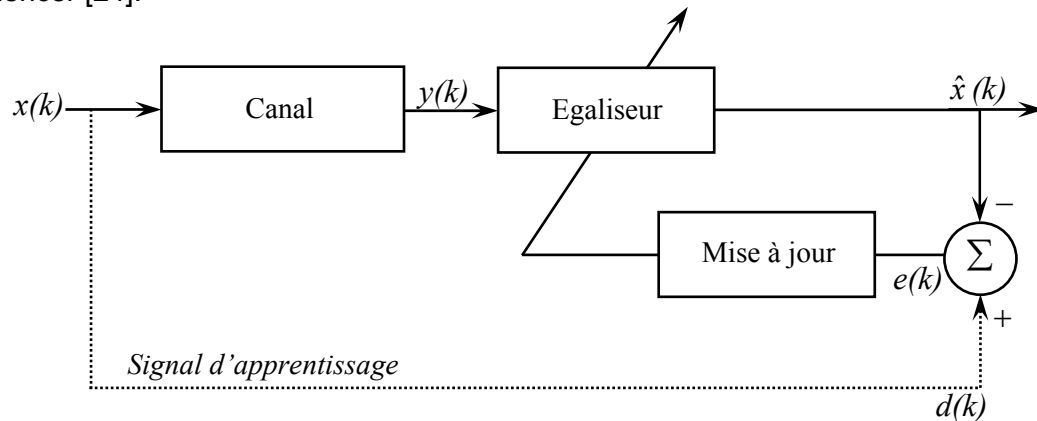


Fig.3.4 — Egalisation avec apprentissage

2.1.1 Egalisation adaptative

L'égalisation adaptative comporte le changement des paramètres du filtre (coefficients) en fonction du temps pour s'adapter aux changements du signal. Pendant les dernières trois décennies, les processeurs de traitement de signal numérique ont fait de grandes avancées dans la rapidité, et la miniaturisation. Par conséquent, les algorithmes de filtrage adaptatif en temps réel deviennent rapidement pratiques et essentiels pour le futur des communications.

L'égalisation adaptative opère en trois modes :

- Annulation d'écho.
- Egalisation de Canal.
- Suppression de bruit large bande / bande étroite.

Un filtre adaptatif est constitué de deux parties (figure 3.5) :

- Un filtre numérique à coefficients ajustables de nature FIR ou IIR (Wiener, Kalman ...etc.),

- Un algorithme de modification de ces coefficients ajustables basé sur un critère d'optimisation [26].

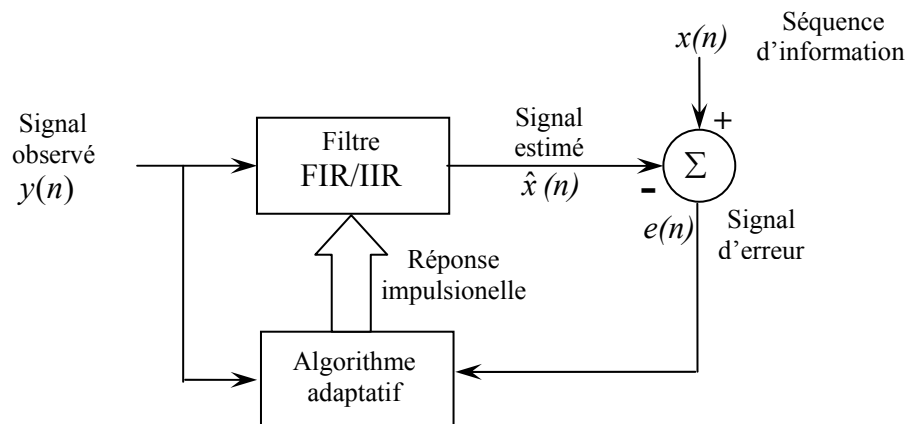


Fig.3.5 — Schéma d'un système de filtrage adaptatif

2.1.2 Filtre de Wiener

Les filtres de Wiener développés à partir de concepts temporels et non fréquentiels sont conçus pour minimiser l'erreur quadratique moyenne entre leur sortie et une sortie désirée. Ils sont dits optimaux au sens du critère de l'erreur quadratique moyenne et nous verrons que dans ce cas, les coefficients des filtres sont liés à la fonction d'auto corrélation du signal d'entrée et à l'inter corrélation entre les signaux d'entrée et de sortie désirée.

Problème d'estimation linéaire : La figure 3.6 illustre un problème courant d'estimation linéaire. $x(n)$ correspond au signal qui nous intéresse mais n'est pas directement accessible. Seul $y(n)$ l'est, et $y(n)$ est obtenu après passage de $x(n)$ dans un système linéaire suivi de l'addition d'un bruit.

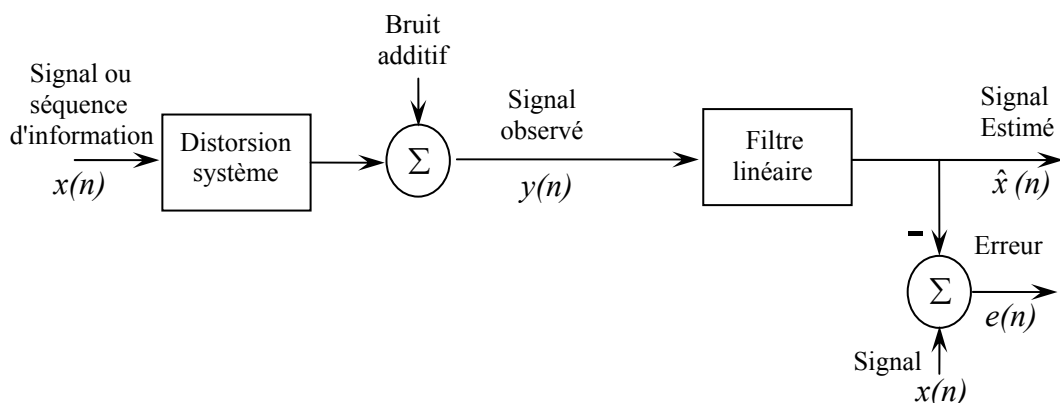


Fig.3.6 — Problème d'estimation linéaire

Le problème qui se pose est comment retrouver $x(n)$ à partir de $y(n)$. Une solution consiste à filtrer $y(n)$ de telle sorte que la sortie $\hat{x}(n)$ soit la plus proche possible de $x(n)$. On peut mesurer la qualité de l'estimation par $e(n)$ défini par :

$$e(n) = x(n) - \hat{x}(n) \quad (3.1)$$

Evidemment, plus $e(n)$ sera faible, plus l'estimation sera bonne. On cherche donc un filtre qui minimisera l'erreur. Il est pratique de chercher à minimiser $e^2(n)$ car c'est une fonction quadratique facilement dérivable. Par ailleurs, étant donné que les signaux intéressants sont aléatoires, la fonction coût qui sera à minimiser est l'erreur quadratique moyenne MSE (Mean Square Error) définie par :

$$\xi(n) = E(e^2(n)) \quad (3.2)$$

Le filtre optimal de Wiener correspond au filtre qui minimisera la MSE [27 et 28].

Filtre de Wiener de type FIR : On se limitera ici au calcul des filtres FIR (Finite Impulse Reponse) [27]. Appelons h , le filtre que nous recherchons et N la longueur de sa réponse impulsionnelle donnée avec une notation matricielle par :

$$h = [h_0 \quad h_1 \quad \dots \quad h_{N-1}]^T$$

Le signal estimé $\hat{x}(n)$ peut alors s'écrire :

$$\hat{x}(n) = \sum_{i=0}^{N-1} h_i y(n-i)$$

ou encore en introduisant la notation matricielle pour $y(n)$:

$$\hat{x}(n) = h^T y(n) \Leftrightarrow \hat{x}(n) = y^T(n) h \quad (3.3)$$

avec :

$$y(n) = [y(n) \quad y(n-1) \quad \dots \quad y(n-(N-1))]^T$$

En faisant l'hypothèse que les signaux $x(n)$ et $y(n)$ sont stationnaires, et si on introduit l'équation 3.3 dans l'équation 3.2, on arrive à la fonction coût suivante :

$$\begin{aligned} \xi &= E[(x(n) - h^T y(n))^2] \\ \Leftrightarrow \xi &= E[x^2(n) - 2h^T y(n)x(n) + h^T y(n)y^T(n)h] \\ \xi &= E[x^2(n)] - 2h^T \Phi_{yx} + h^T \Phi_{yy} h \end{aligned} \quad (3.4)$$

où Φ_{yy} est une matrice d'auto corrélation de taille NxN définie par :

$$\Phi_{yy} = E[y(n)y^T(n)] \quad (3.5)$$

et où Φ_{yx} est un vecteur d'inter corrélation de taille N défini par :

$$\Phi_{yx} = E[y(n)x(n)] \quad (3.6)$$

L'équation 3.4 montre que pour un filtre FIR, la fonction coût MSE dépend de la réponse impulsionnelle h . Pour en obtenir le minimum, il suffit de chercher les conditions d'annulation de la dérivée de la fonction coût par rapport aux variables qui sont les N points de la réponse impulsionnelle du filtre.

La dérivée de la fonction coût par rapport au $j^{\text{ème}}$ point de la réponse impulsionnelle est donnée par :

$$\frac{\partial \xi}{\partial h_j} = E \left[\frac{\partial}{\partial h_j} \{e^2(n)\} \right] = E \left[2e(n) \frac{\partial e(n)}{\partial h_j} \right]$$

En substituant dans cette équation $e(n)$ par les équations 3.1 et 3.3, on obtient l'expression suivante :

$$\frac{\partial \xi}{\partial h_j} = E \left[2e(n) \frac{\partial}{\partial h_j} \{x(n) - h^T y(n)\} \right]$$

En utilisant le fait que la sortie du filtre $h^T y(n)$ peut s'écrire comme une somme de N produits dont un seul contient le terme h_j , on arrive à l'expression suivante :

$$\begin{aligned} \frac{\partial \xi}{\partial h_j} &= E \left[2e(n) \frac{\partial}{\partial h_j} \{h_j y(n-j)\} \right] \\ \Leftrightarrow \frac{\partial \xi}{\partial h_j} &= E[-2e(n)y(n-j)] \end{aligned}$$

On cherche les conditions d'annulation de cette équation pour tous les $j = \{0, \dots, N-1\}$. Ceci nous donne un ensemble de N équations qui peut être écrit de façon matricielle en introduisant le vecteur gradient ∇ :

$$\nabla = \begin{bmatrix} \partial \xi / \partial h_0 \\ \partial \xi / \partial h_1 \\ \vdots \\ \partial \xi / \partial h_j \\ \vdots \\ \partial \xi / \partial h_{N-1} \end{bmatrix} = -2E \begin{bmatrix} y(n)e(n) \\ y(n-1)e(n) \\ \vdots \\ y(n-j)e(n) \\ \vdots \\ y(n-N+1)e(n) \end{bmatrix} = -2E \left\{ \begin{bmatrix} y(n) \\ y(n-1) \\ \vdots \\ y(n-j) \\ \vdots \\ y(n-N+1) \end{bmatrix} \cdot e(n) \right\} = -2E[y(n)e(n)]$$

En utilisant les équations 3.1 et 3.3 pour remplacer $e(n)$ on obtient :

$$\nabla = -2E[y(n)(x(n) - y^T(n)h)] = -2E[y(n)x(n)] + 2E[y(n)y^T(n)h]$$

qui devient en introduisant la matrice d'auto corrélation et le vecteur d'inter corrélation :

$$\nabla = -2\Phi_{yx} + 2\Phi_{yy}h \quad (3.7)$$

La réponse impulsionnelle optimale h_{opt} est celle qui annule cette équation, d'où :

$$\Phi_{yy}h_{opt} = \Phi_{yx} \quad (3.8)$$

Le filtre ainsi défini est appelé filtre FIR de Wiener. Il permet d'obtenir une erreur quadratique minimale entre $x(n)$ et son estimé $\hat{x}(n)$. Le filtre est lié à l'équation (3.9).

A partir de l'équation (3.4) que nous rappelons ici :

$$\xi = E[x^2(n)] - 2h^T \Phi_{yx} + h^T \Phi_{yy}h$$

on obtient

$$\xi_{min} = E[x^2(n)] - h_{opt}^T \Phi_{yx} \quad (3.9)$$

2.1.3 Types d'algorithmes adaptatifs

Les 31 algorithmes représentés dans la figure 3.7 incluent les gradients et les variantes de l'algorithme LMS ; les méthodes des moindres carrés récursives, la racine carrée conventionnelle, Trellis, et les formes des filtres rapides et transversales ; le domaine de fréquence, le domaine transformé, et les approches sous bande ; et les algorithmes de projection rapides [29].

La conception de la fonctionnalité adaptative d'un filtre commence par des spécifications avec lesquelles les algorithmes sont inclus dans la mise à jour des coefficients. Une taxonomie des algorithmes de filtrage adaptatif a été développée suivant la figure 3.7. Ce graphe connexe dépeint tous les algorithmes choisis pour l'implémentation. Les noeuds du côté droit contiennent les algorithmes du gradient stochastique et les

méthodes relatives, tandis que les noeuds du côté gauche sont des méthodes déterministes (projection et moindres carrés).

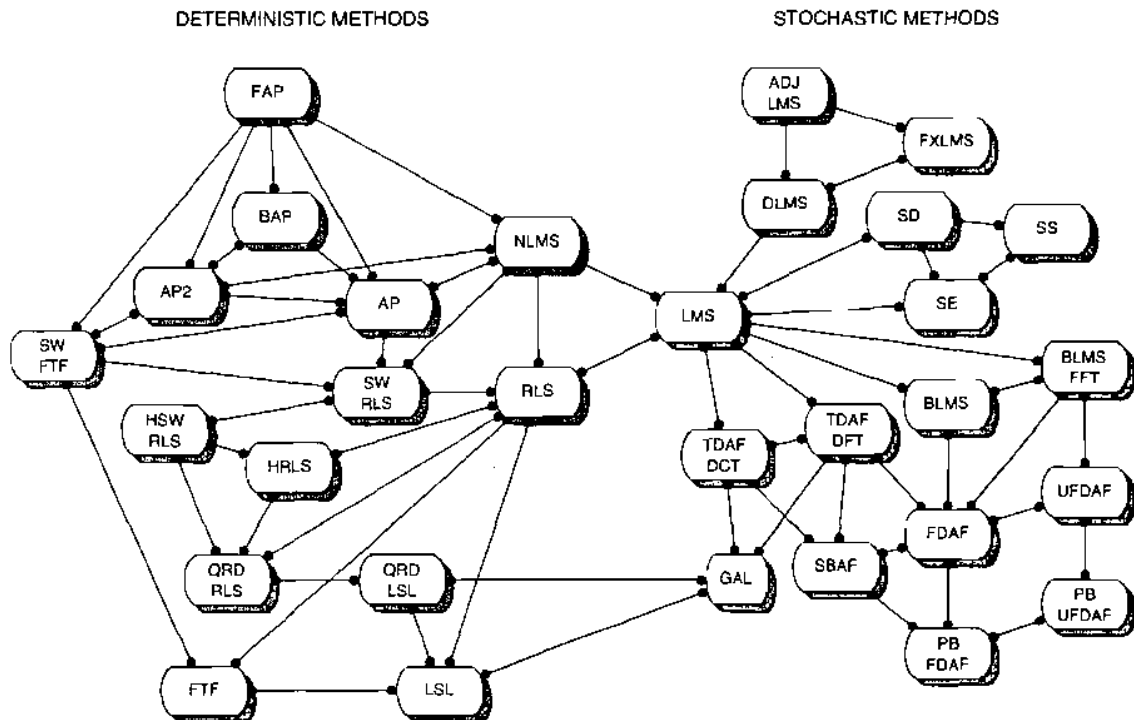


Fig.3.7 — Catégorisation des filtres adaptatifs

Les liens entre les méthodes indiquent des similitudes entre les paires d'algorithmes soit en termes de structure d'exécution ou en fonction de coût. Cette taxonomie a un certain nombre de points intéressants :

Les algorithmes les plus populaires et qui sont employés couramment sont : LMS, RLS, et NLMS. Ils sont situés au centre et ont de nombreux liens aux méthodes relatives. La plupart des algorithmes particuliers tels que le FTF et le FAP incluent le LMS et deviennent hybrides.

Il y a deux classes distinctes des algorithmes adaptatifs (méthodes stochastiques et déterministes) avec relativement peu de liens entre eux. Dans ces algorithmes les classes sont des groupes d'algorithmes qui sont semblables dans l'implémentation ou dans les principes à la base de leurs dérivations.

On peut utiliser le regroupement donné par cette taxonomie pour différencier entre les différents algorithmes présentés, le tableau 3.1 fait référence de regroupement et de classification de ces algorithmes [29].

Acronyme	Algorithme	Acronyme	Algorithme
<i>Simple Updates</i>		<i>Delayed Updates</i>	
LMS	Least-mean-square	DLMS	Delayed LMS
SE	Sign-error	FXLMS	Filtered-X LMS
SD	Sign-data	ADJLMS	Adjoint LMS
SS	Sign-sign	<i>Projections Methods</i>	
<i>Block and Frequency Domain Methods</i>		NLMS	Normalized LMS
BLMS	Block LMS	AP	Affine projection $O(N^3)$
BLMSFFT	Block LMS (FFT)	AP2	Affine projection $O(N^2)$
FDAF	Frequency-domain adaptive filter	FAP	Fast affine projection
UFDAF	Unconstrained FDAF	BAP	Block affine projection
PBFDAF	Partitioned-block FDAF	<i>Conventional Least-Squares Methods</i>	
PBUFDAF	Partitioned-block unconstrained FDAF	RLS	Recursive least-squares
<i>Transform Domain and Sub-band Methods</i>		SWRLS	Sliding-window RLS
TDAFDFT	Transform-domain adaptive filter (DFT)	<i>Square-Root Least-Squares Methods</i>	
TDAFDCT	Transform-domain adaptive filter (DCT)	QRDRLS	QR-decomposition RLS
SBAF	Sub-band adaptive filter	HRLS	Householder RLS
<i>Lattice Structures</i>		HSWRLS	Householder sliding-window RLS
GAL	Gradient adaptive lattice	<i>Fast transversal filter</i>	
LSL	Least squares lattice	FTF	Fast transversal filter
QRDLISL	QR-decomposition least-squares lattice	SWFTF	Sliding-window fast transversal filter

Tableau.3.1 — Algorithmes adaptatifs à travers les filtres adaptatifs.

2.1.4 Choix de l'algorithme

Quand les fonctions d'auto et d'inter corrélation ne sont pas connues (cas le plus courant), on va alors approcher le filtre optimal de Wiener en utilisant une boucle de retour et un algorithme de minimisation : c'est ce qu'on appelle le filtrage adaptatif (figure 3.5). Dans ce cas, on remplacera la connaissance des fonctions de corrélation par une phase d'apprentissage permettant de modifier itérativement la réponse impulsionnelle du filtre.

L'algorithme doit se préoccuper des facteurs suivants [30] :

- La rapidité de convergence qui sera le nombre d'itérations nécessaires pour converger «assez près» de la solution optimale de Wiener dans le cas stationnaire.
- La mesure de cette «proximité» entre cette solution optimale et la solution obtenue.
- La capacité de poursuite (tracking) des variations (non stationnarités) du processus. On examinera quels sont les algorithmes vraiment adaptatifs.
- La robustesse vis-à-vis du bruit.
- La complexité.
- La structure (modularité, parallélisme, ...).

- Les propriétés numériques (stabilité et précision) les meilleures possibles. Il doit être stable avec une précision satisfaisante.

Nous ne nous intéresserons dans le cadre de notre travail qu'aux trois premiers critères de choix, ce qui nous mène aux trois algorithmes les plus utilisés dans le domaine du traitement de signal et l'égalisation adaptative qui sont LMS, NLMS et RLS [29].

2.1.5 Algorithme LMS

L'algorithme LMS (Least Mean Squares) est un choix populaire dans beaucoup d'applications exigeant le filtrage adaptatif. Deux raisons principales de sa popularité : simplicité et complexité informatique réduite. En outre, il y a plusieurs variantes de l'algorithme qui peuvent être employées spécifiquement afin de résoudre différents types de problèmes qui sont inhérents à certaines applications.

La version de base du LMS est un cas spécial du filtre adaptatif du gradient descendant (steepest descent) bien connu. Le but de cette technique est de réduire au minimum une fonction de coût quadratique en mettant à jour itérativement des poids de sorte qu'ils convergent à la solution optimale.

De la méthode de gradient descendant, le vecteur de poids d'égalisation est donné par l'équation suivante :

$$h(n+1) = h(n) + 1/2\mu \left[-\nabla(E(e^2(n))) \right] \quad (3.10)$$

où :

μ est un paramètre crucial affectant la stabilité et le taux de convergence de l'algorithme LMS. Il représente le pas de descente de l'algorithme.

$e^2(n)$ est l'erreur quadratique moyenne entre la sortie $\hat{x}(n)$ et le signal de référence $x(n)$; elle est donnée par la formule suivante :

$$e^2(n) = [x^*(n) - h^T y(n)]^2 \quad (3.11)$$

où

$$e(n) = x(n) - \hat{x}(n)$$

et

$$\hat{x}(n) = h^T y(n) \Leftrightarrow \hat{x}(n) = y^T(n)h$$

Le vecteur gradient est donné par :

$$\nabla = -2\Phi_{yx} + 2\Phi_{yy}h$$

Dans la méthode du gradient descendant, le plus gros problème est le calcul impliqué dans la recherche des valeurs Φ_{yy} et Φ_{yx} des matrices en temps réel. Pour y remédier, l'algorithme LMS utilise les valeurs instantanées des matrices de covariance Φ_{yy} et Φ_{yx} au lieu de leurs valeurs réelles c'est-à-dire :

$$\Phi_{yy} = E[y(n)y^T(n)] \quad (3.12)$$

$$\Phi_{yx} = E[y(n)x^*(n)] \quad (3.13)$$

Par conséquent, la mise à jour du vecteur de poids d'égalisation peut être donnée par l'équation suivante :

$$h(n+1) = h(n) + \mu \cdot y(n)[x^*(n) - y^T(n)h(n)] \quad (3.14)$$

$$h(n+1) = h(n) + \mu \cdot y(n)e(n)^* \quad (3.15)$$

L'algorithme LMS est engagé à démarrer avec une valeur arbitraire $h(0)$ pour le vecteur de poids à $n = 0$. Les rectifications successives du vecteur de poids finalement conduisent à la valeur minimale de l'erreur quadratique moyenne.

Par conséquent, l'algorithme LMS peut se résumer par ces équations :

Initialiser avec :

$$\mathbf{h}(0) = \mathbf{0}$$

Pour toute séquence $n = 1, 2, \dots$ Faire :

$$\hat{x}(n) = \mathbf{h}^T(n-1) \mathbf{y}(n)$$

$$e(n) = x(n) - \hat{x}(n)$$

$$\mathbf{h}(n) = \mathbf{h}(n-1) + 2 \mu \mathbf{y}(n) e(n)$$

Convergence et stabilité de l'algorithme LMS : L'algorithme LMS engagé avec certaines valeurs arbitraires pour le poids est perçu comme vecteur de convergence : Si μ est choisie pour être très faible alors l'algorithme converge très lentement. Une grande valeur de μ peut conduire à une accélération de convergence, mais peut-être moins stable, autour de la valeur minimale. Habituellement μ est choisie dans la marge [27 et 28] :

$$0 < \mu < \frac{2}{\lambda_{\max}} \quad (3.16)$$

où $\lambda_{\max}$ représente la valeur propre maximale de la matrice d'auto corrélation Φ_{yy} .

La convergence de l'algorithme est inversement proportionnelle à la propagation des valeurs propres de la matrice d'auto corrélation Φ_{yy} . Pour des valeurs propres de Φ_{yy} qui sont très répandues, la convergence peut être lente.

2.1.6 Algorithme NLMS

Le principal inconvénient de l'algorithme LMS est qu'il est sensible à l'échelle de son entrée $y(n)$. Cela rend très difficile de choisir un taux d'apprentissage μ qui garantit la stabilité de l'algorithme. L'algorithme Normalised Least Mean Squares filtre (NLMS) est une variante de l'algorithme LMS qui permet de résoudre ce problème sans employer les évaluations de la fonction d'autocorrection de signal d'entrée. L'algorithme NLMS peut être résumé comme suit :

L'équation de mise à jour est :

$$h(n+1) = h(n) + \mu y(n) e^*(n)$$

Initialisation :

$$h(0) = 0$$

Calcul :

$$y(n) = [y(n) \quad y(n-1) \quad \dots \quad y(n-(N-1))]^T$$

$$e(n) = x^*(n) - \hat{x}(n)$$

$$\hat{x}(n) = h^T y(n)$$

$$h(n+1) = h(n) + \frac{\mu^* y(n) e^*(n)}{a + y(n)^H y(n)}$$

où $y(n)^H$ désigne la matrice Hermitienne transposée de la matrice $y(n)$

$$\mu = \frac{\mu^*}{a + y(n)^H y(n)}, \quad 0 < \mu^* < 2, \quad a \geq 0$$

Le NLMS converge beaucoup plus rapidement que le LMS. Dans quelques applications, la normalisation est universelle.

L'effet de grandes fluctuations au niveau de la puissance du signal d'entrée est compensé au niveau de l'adaptation, ainsi l'effet de la grande longueur de vecteur d'entrée est compensé en réduisant la taille des étapes de l'algorithme [27 et 28].

2.1.7 Algorithme RLS

Sachant que les propriétés statistiques nous sont inconnues, on ne va pas chercher à minimiser $E[e^2(n)]$ mais une somme finie d'erreur donnée au carré par :

$$\zeta = \sum_{k=0}^n (x(k) - \hat{x}(k))^2 \quad (3.17)$$

Quand cette fonction de coût est minimisée en utilisant une réponse impulsionnelle $h(n)$ associée à $\hat{x}(n)$, on obtient l'estimée des moindres carrés.

La réponse impulsionnelle est donc fonction des échantillons disponibles et non pas d'une moyenne statistique générale. Par analogie avec Wiener, elle est donnée par la relation :

$$R_{yy}(n)h(n) = r_{yx}(n) \quad (3.18)$$

où

$$R_{yy}(n) = \sum_{k=0}^n y(k)y^T(k) \quad (3.19)$$

et

$$R_{yx}(n) = \sum_{k=0}^n y(k)x(k) \quad (3.20)$$

La réponse impulsionnelle du filtre est donc à modifier à chaque nouvel échantillon. Pour limiter la taille du calcul, on passe par une équation récursive :

$$h(n) = h(n-1) + k(n)e(n) \quad (3.21)$$

où le vecteur du gain :

$$k(n) = \frac{R_{yy}^{-1}(n-1)y(n)}{(I + y^T(n)R_{yy}^{-1}(n-1)y(n))} \quad (3.22)$$

$$e(n) = x(n) - h^T(n-1)y(n) \quad (3.23)$$

et

$$R_{yy}^{-1}(n) = R_{yy}^{-1}(n-1) - k(n)y^T(n)R_{yy}^{-1}(n-1) \quad (3.24)$$

$$R_{yy}^{-1}(n) = R_{yy}^{-1}(n-1) - \frac{R_{yy}^{-1}(n-1)y(n)y^T(n)R_{yy}^{-1}(n-1)}{(I + y^T(n)R_{yy}^{-1}(n-1)y(n))} \quad (3.25)$$

Ces équations sont connues sous le nom de l'algorithme RLS [27 et 28]. Le détail de cet algorithme est donné ci-dessous.

Initialiser avec :

$$\mathbf{R}_{yy}(0) = \frac{1}{\delta} I_N, \quad \delta \text{ est un nombre positif et petit}$$

$$\mathbf{h}(0) = \mathbf{0}$$

Pour toute séquence $n = 1, 2, \dots$ Faire :

$$\hat{x}(n) = \mathbf{h}^T(n-1) \mathbf{y}(n)$$

$$e(n) = x(n) - \hat{x}(n)$$

$$\mathbf{R}_{yy}^{-1}(n) = \frac{1}{\alpha} \left(\mathbf{R}_{yy}^{-1}(n-1) - \frac{\mathbf{R}_{yy}^{-1}(n-1) \mathbf{y}(n) \mathbf{y}^T(n) \mathbf{R}_{yy}^{-1}(n-1)}{\left(\alpha + \mathbf{y}^T(n) \mathbf{R}_{yy}^{-1}(n-1) \mathbf{y}(n) \right)} \right)$$

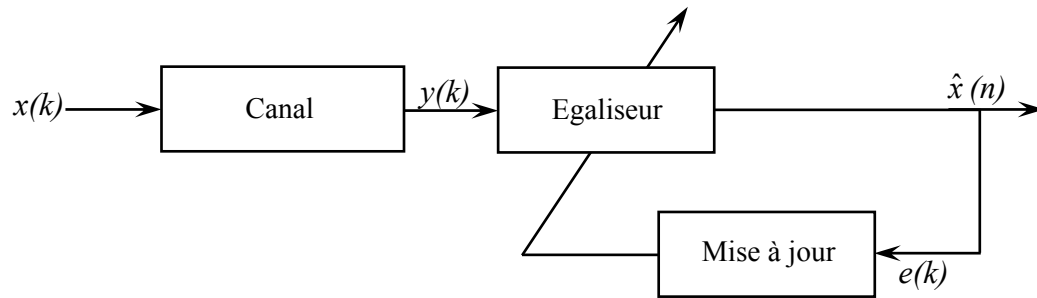
$$\mathbf{h}(n) = \mathbf{h}(n-1) + \mathbf{R}_{yy}^{-1}(n) \mathbf{y}(n) e(n)$$

2.2 Egalisation aveugle

Dans l'égalisation par filtrage adaptatif standard (exemple LMS) l'ajustement des coefficients de l'égaliseur est basé sur l'erreur qui est la différence entre la sortie du filtre et le signal désiré. Ce dernier est connu par l'émetteur et le récepteur (séquence d'apprentissage). L'inclusion de ces pilotes avec l'information transmise réduit le rendement du système.

Dans la plupart des applications à haut débit et à grande vitesse, le canal de propagation n'est pas connu à l'avance, et probablement est un modèle variant dans le temps, ce qui exige des solutions adaptatives dans la conception au niveau du récepteur. La séquence d'apprentissage est typiquement une partie fixe de la séquence de la source qui est connue a priori à l'émetteur et au récepteur, mais elle-même ne contient pas d'information.

Le GSM par exemple consacre jusqu'à 22% de son temps de transmission pour les pilotes d'apprentissage et l'UMTS jusqu'à 40% du temps pour ces séquences. Donc on est obligé de trouver une solution pour faire l'égalisation sans regard à ces pilotes (auto adaptatif) ; donc on parle alors de l'égalisation autodidacte ou aveugle « Blind Equalization ».

**Fig.3.8 — Égalisation aveugle**

La technique dite égalisation aveugle (figure 3.8) est basée sur quelques statistiques de la source. Le terme « Blind » signifie que le récepteur n'a aucune connaissance de la séquence transmise ou de la réponse impulsionnelle du canal. Seulement quelques propriétés statistiques ou structurales des signaux transmis et reçus sont exploitées pour adapter l'égaliseur. Par conséquent, les séquences d'apprentissage peuvent être exclues. Les méthodes aveugles tiennent compte des variations rapides du cheminement du canal et utilisent seulement des symboles de l'information.

Des algorithmes ont été proposés comme les statistiques du second degré (SOS), énergie minimum de rendement (MOE), et l'algorithme du module constant (CMA) qui est largement utilisé dans les techniques d'égalisation aveugle [24 et 31].

2.2.1 L'algorithme CMA (constant modulo algorithm)

L'algorithme du module constant (CMA) est développé indépendamment par Godard et Treichler Agee en 1980. L'implémentation du gradient stochastique ou le CMA est largement utilisée dans la pratique. L'algorithme de Godard est le plus couramment utilisé en égalisation autodidacte. Il ne nécessite ni la connaissance des données émises, ni celle des données décidées.

Le critère CM essaye de rapprocher la puissance du module de sortie d'égaliseur à une constante. Cette constante est choisie essentiellement en projetant tous les points sur un cercle. Cet algorithme est du type gradient descendant qui réduit au minimum la fonction de coût (appelée fonction de dispersion d'ordre P) suivante :

$$\xi_{Godard}^P(R) = E \left\{ \left(|y(n)|^P - \gamma_P \right)^2 \right\} \quad (3.26)$$

où γ_P est une constante réelle définie par :

$$\gamma_P = \frac{E\{|d_n|^{2P}\}}{E\{|d_n|^P\}^2} \quad (3.27)$$

La fonction de dispersion d'ordre 2 ($P=2$) est la plus utilisée en pratique, donc la fonction de coût devient :

$$\xi_{\text{Godard}}(R) = E\left\{\left(|y(n)|^2 - \gamma_2\right)^2\right\}$$

où $y(n)$ est la sortie de l'égaliseur et $u(n)$ est le signal transmis. γ_2 est souvent appelé la constante de dispersion.

Godard montre que la valeur de la constante de dispersion minimise la fonction de coût. La relation entre les critères CM et MSE est très étroite. *Treichler et Agee* démontrent que la minimisation des performances de la fonction de coût du module constant est équivalente à minimiser l'erreur des moindres carrés.

γ_2 est la constante de module (référence) qui dépend de la modulation utilisée pour la transmission de l'information par exemple BPSK, QPSK. Pour BPSK $\gamma_2 = 1$, et pour QPSK $\gamma_2 = \sqrt{2}$. Les symboles transmis sont : $1-j$, $1+j$, $-1+j$ et $-1-j$ (le module est le même $\sqrt{1^2 + 1^2} = \sqrt{2}$) [24 et 31].

2.2.2 Equation de mise à jour de CMA

Le processus d'adaptation est basé sur la même technique utilisée dans le LMS. L'ajustement des coefficients est basé sur l'algorithme du gradient descendant qui est défini par :

$$h(n+1) = h(n) - \frac{1}{2} \mu_n \nabla \xi_{\text{Godard}} \quad (3.28)$$

ξ_{Godard} : Fonction de coût qui doit être minimisée.

Calcul de la sortie du filtre : $y(n) = h^T(n)u(n)$

Calcul du signal d'erreur : $e(n) = (\gamma^2 - |y(n)|^2)y(n)$

Mise à jour du filtre : $h(n+1) = h(n) + \mu u(n) e(n)$

3. Simulations et résultats

Dans ce qui suit les simulations sont présentées pour illustrer les performances des algorithmes LMS, NLMS, RLS et CMA en utilisant le logiciel de simulation Matlab.

3.1 Principe :

On considère un canal complexe modélisé par un filtre RIF qui possède les caractéristiques suivantes $C(z) = 1 + (-0.25 + 0.34j)z^{-1} + (-0.15j)z^{-2}$ où z^{-n} représentent les différents retards, ce qui va mettre en évidence les problèmes du canal : les multi trajets, les problèmes d'accès multiples (MAI), et d'interférence des symboles (ISI).

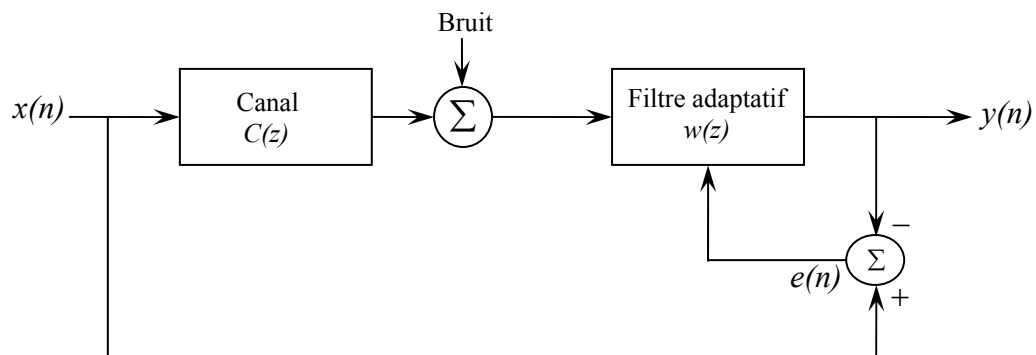


Fig.3.9 — Principe des simulations

On génère un signal aléatoire QPSK $x(n)$ d'une longueur de 3000 séquences (les symboles transmis sont $1-j$, $1+j$, $-1+j$ et $-1-j$), ce signal va être déformé avec les différents effets du canal, on y ajoute du bruit à la sortie du canal pour bien modéliser les effets réels, le bruit est calculé grâce au SNR (Signal to Noise Ratio - le rapport signal bruit).

Avec un filtre adaptatif de 10 coefficients $w(z)$ on essayera de reproduire le signal original avec la moindre erreur possible.

Les différents algorithmes d'apprentissage ; LMS, NLMS, RLS puis CMA, sont testés. Le critère d'apprentissage est l'erreur quadratique moyenne MSE entre le signal original ($x(n)$, désiré) et le signal de sortie du filtre ($y(n)$, estimé). On cherche à minimiser cette erreur en tenant compte des changements des coefficients du filtre $w(z)$ avec le temps, le pas d'apprentissage (μ) et les différentes constantes qui prennent part à l'apprentissage de chaque algorithme au fur et à mesure des itérations (nombre de séquences).

La figure 3.10 représente la constellation du signal à transmettre $x(n)$, et ce même signal après son passage au canal et l'ajout d'un bruit avec un SNR = 30, 20 et 15 dB.

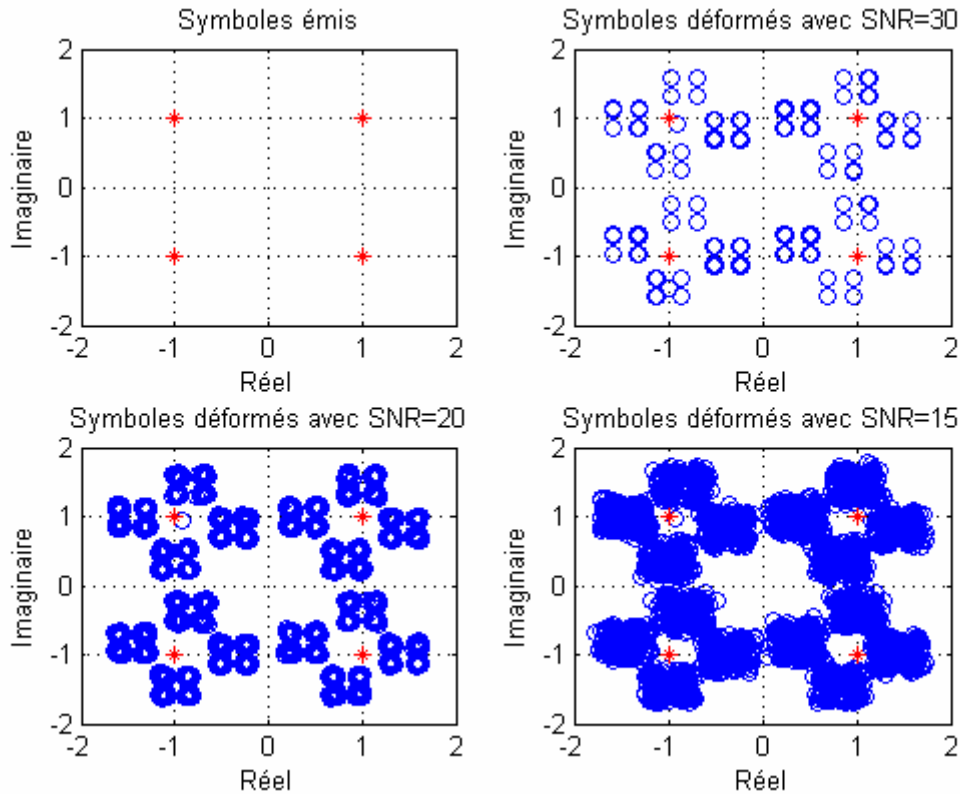


Fig.3.10 — Symboles à égaliser

3.2 Algorithme LMS

Pour un pas d'apprentissage approprié $\mu=0.03$, on a une bonne convergence de l'erreur MSE (figure 3.11) et des coefficients de l'égaliseur (figure 3.12). La figure 3.13 représente la poursuite des séquences estimées par l'égaliseur $y(n)$ et celles désirées $x(n)$. L'algorithme LMS présente une bonne poursuite des données avec une erreur très minime.

La rapidité de convergence est assurée avec le bon pas choisi μ , et c'est au bout de 100 itérations qu'on voit l'algorithme LMS se stabiliser.

La constellation présentée en figure 3.14 illustre le bon fonctionnement de l'algorithme LMS en ce qui concerne l'erreur d'adaptation, la convergence des séquences déformées par le canal et aussi son efficacité en matière d'adaptation, on voit bien que toutes les séquences tendent vers les séquences originales à part celles qui sont utilisées pour l'initialisation du filtre (0,0) et celles qui sont calculées graduellement pour arriver aux valeurs optimales.

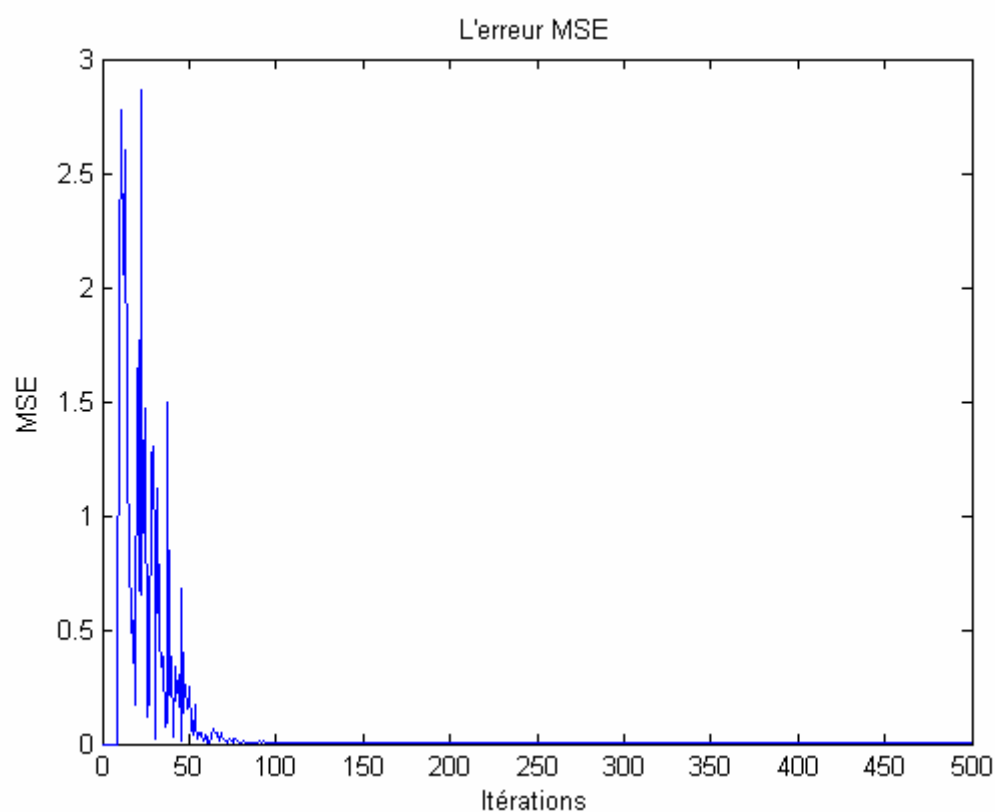


Fig.3.11 — L'erreur MSE du LMS sur 500 itérations

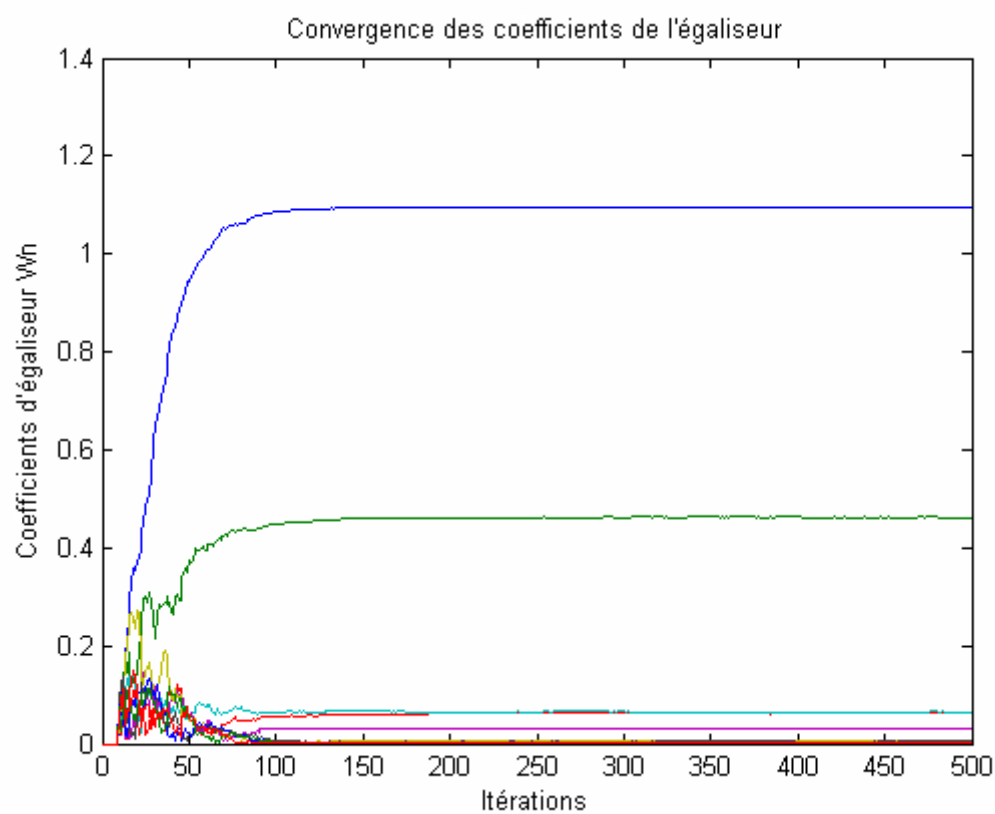


Fig.3.12 — Convergence des coefficients de l'égaliseur LMS sur 500 itérations

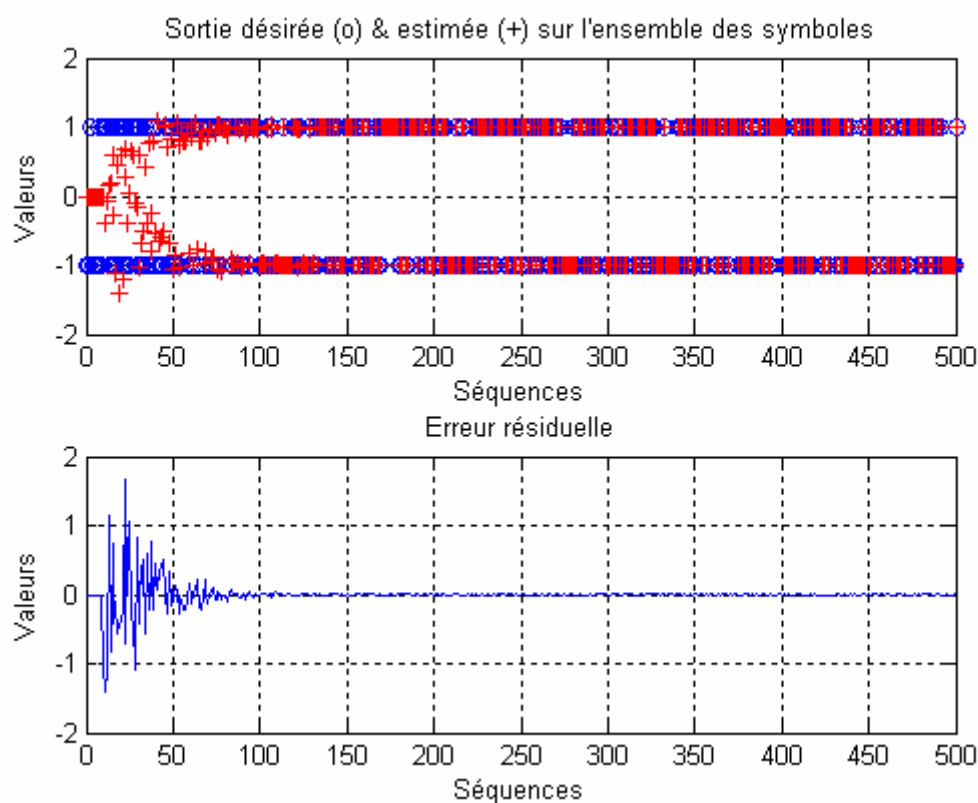


Fig.3.13 — Capacité de poursuite et erreur résiduelle du LMS sur 500 itérations

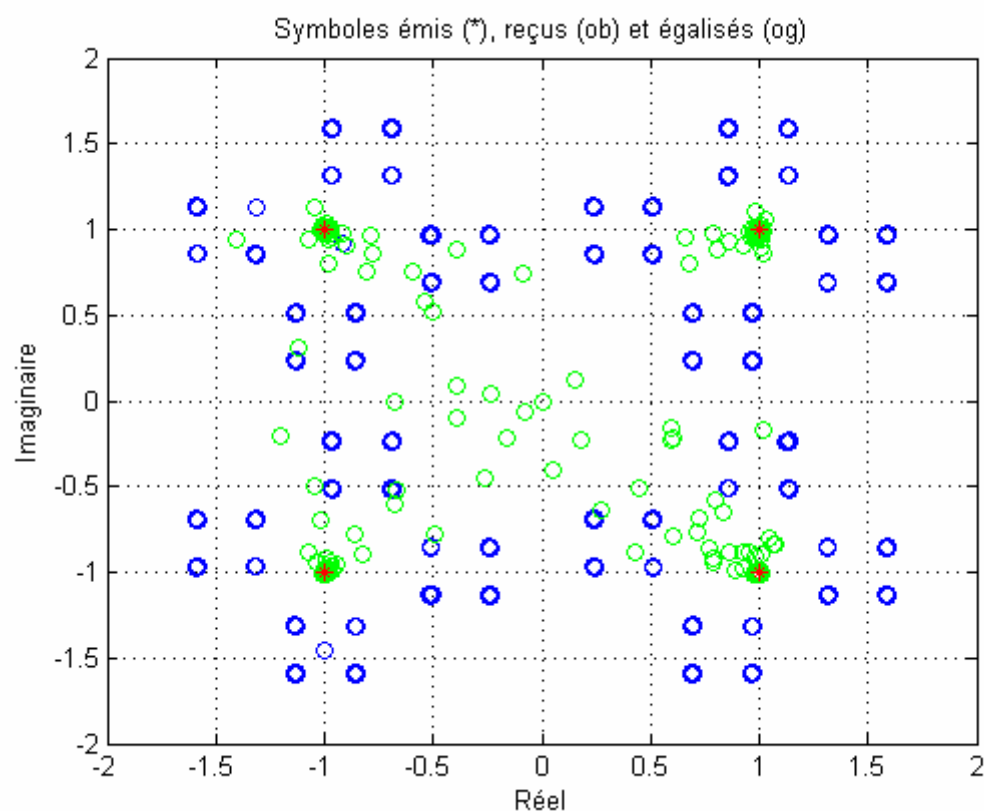


Fig.3.14 — Constellation des signaux émis, reçus (bleu) & égalisés (vert) avec le LMS

La condition pour que l'égaliseur LMS résolve la déformation du canal, est de prendre la fonction de transfert $w(z)=1/C(z)$ telle que la convolution $w(z)*C(z)=1$. La figure 3.15 illustre la convolution de l'égaliseur $w(z)$ avec le canal $C(z)$ qui est égale à 1.

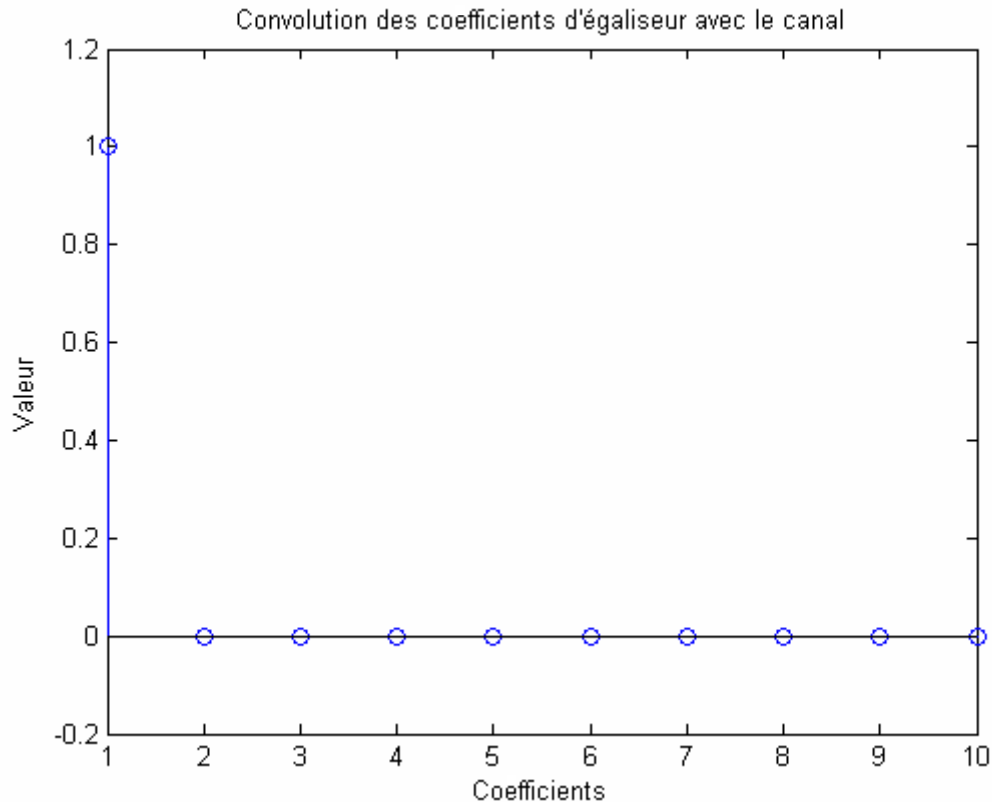


Fig.3.15 — Convolution du canal $C(z)$ avec l'égaliseur LMS

Pour un pas $\mu=0.09$ (figure 3.16) et un pas $\mu=0.0003$ (figure 3.17) l'erreur MSE, la valeur des coefficients du filtre $w(z)$ et la constellation du signal égalisé ont du diverger. La convolution du canal avec l'égaliseur n'étant pas à 1 ($w(z)*C(z)\neq 1$) ce qui montre l'importance du choix du pas d'apprentissage μ .

Pour le pas approprié $\mu=0.03$, et plusieurs rapports signal/bruit SNR=30, 20 et 15 dB, la convergence des coefficients est assurée (figure 3.18). L'erreur ne diverge pas, elle tend vers zéro à partir de 100 itérations. Tant que le bruit augmente l'erreur augmente aussi, le bruit a une grande influence sur la qualité de réception, mais en revanche le filtre LMS montre une bonne robustesse vis-à-vis des bruits externes.

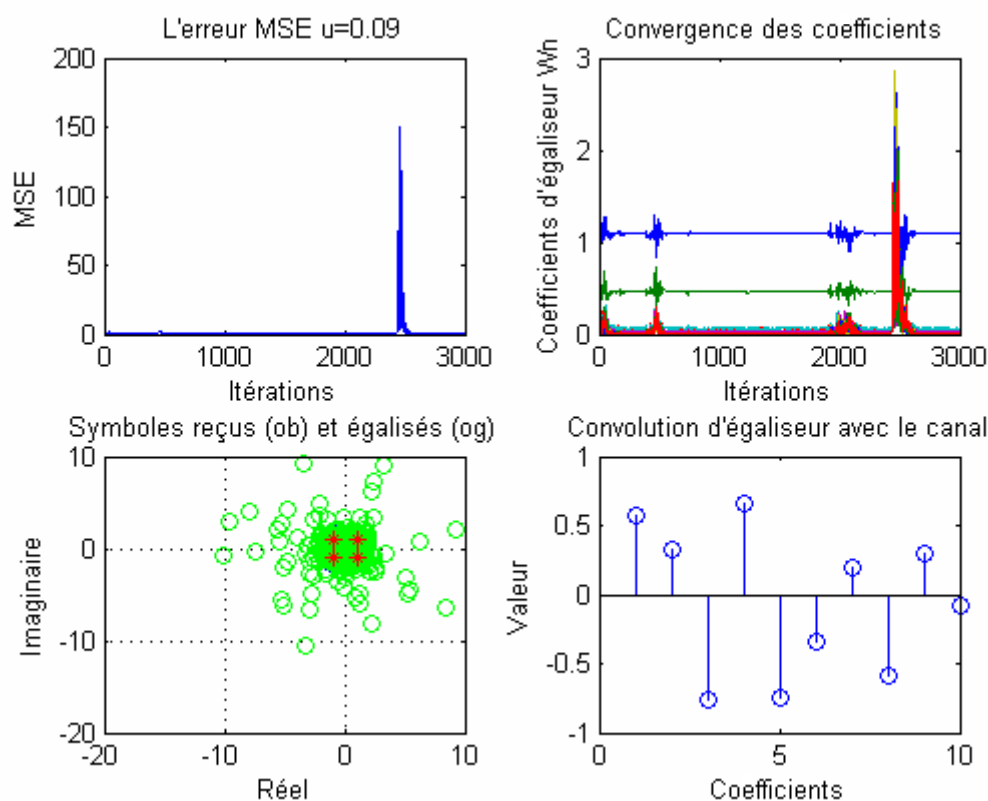


Fig.3.16 — Performances du LMS avec un pas $\mu = 0.09$

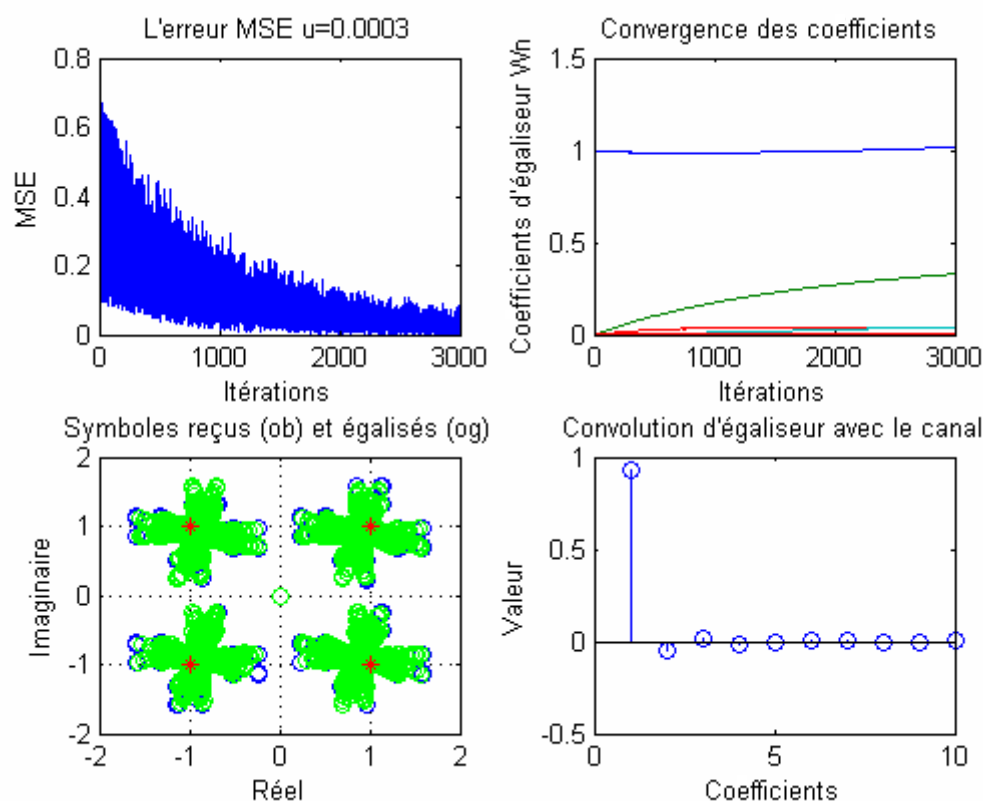


Fig.3.17 — Performances du LMS avec un pas $\mu = 0.0003$

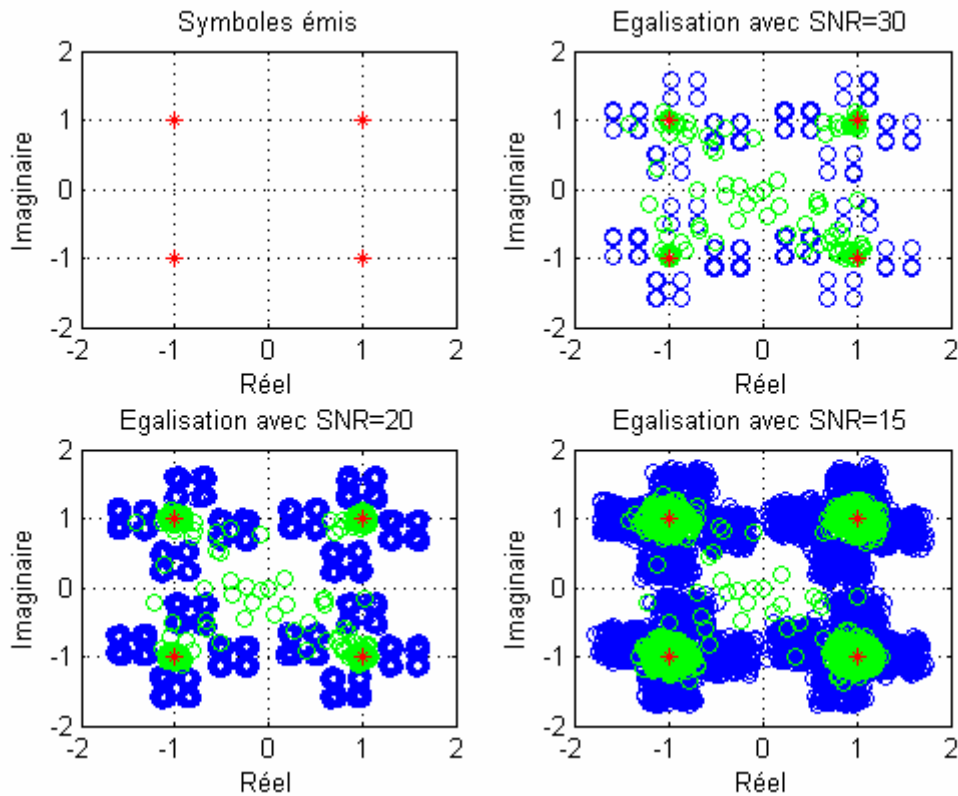


Fig.3.18 — Performances & robustesse du LMS avec des SNR=30, 20 & 15 dB

3.3 Algorithme NLMS

Pour la stabilité de la convergence de l'algorithme LMS, le facteur de convergence μ est soumis à une condition dépendante de la puissance du signal primaire à traiter. Pour s'affranchir de l'influence de $x(n)$ on utilise l'algorithme NLMS (LMS normalisé), où on normalise le signal reçu, ce qui est équivalent à changer μ . la condition de stabilité suivante est alors vérifiée : $0 < \mu < 2$.

Les résultats sur la convergence de l'erreur et les coefficients de l'égaliseur NLMS sont donnés par les figures 3.19 et 3.20.

La figure 3.21 représente la poursuite des séquences estimées par l'égaliseur $y(n)$ et celles désirées $x(n)$. L'algorithme NLMS présente une bonne poursuite des données avec une erreur minime.

La rapidité de convergence n'est pas assurée, et c'est au bout de 200 itérations qu'on voit l'algorithme NLMS se stabiliser, il est un peu lent par rapport au LMS.

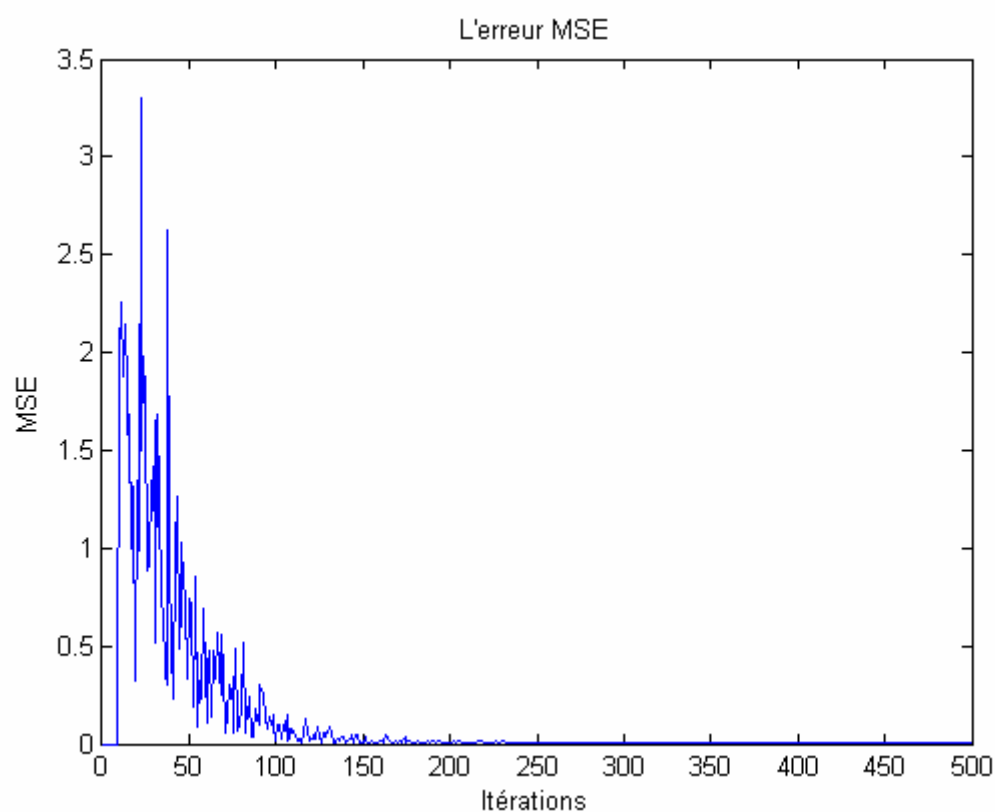


Fig.3.19 — L'erreur MSE du NLMS sur 500 itérations

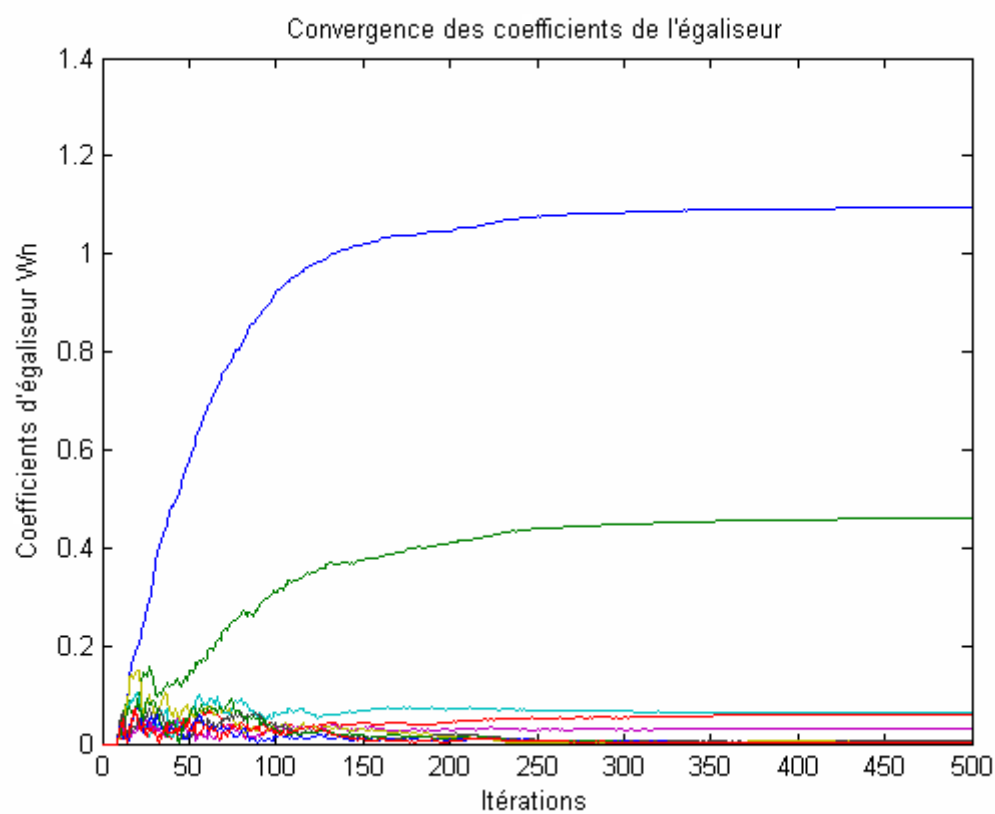


Fig.3.20 — Convergence des coefficients de l'égaliseur NLMS sur 500 itérations

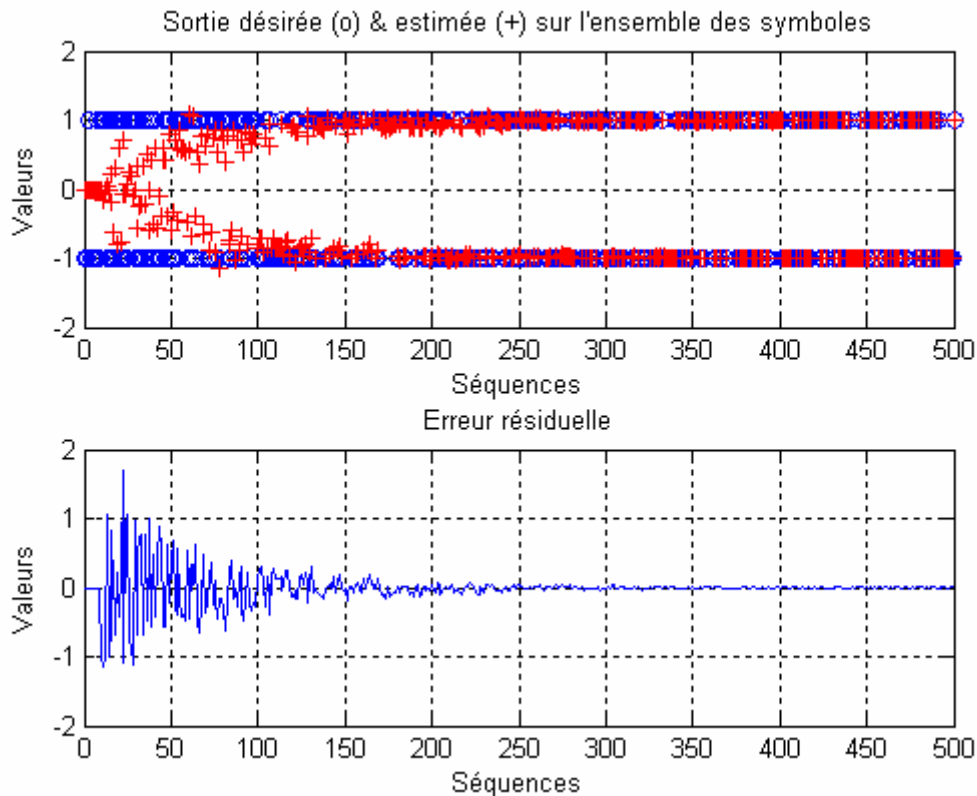


Fig.3.21 — Capacité de poursuite & l'erreur résiduelle du NLMS sur 500 itérations

La constellation présentée en figure 3.22 illustre le bon fonctionnement de l'algorithme NLMS en ce qui concerne l'erreur d'adaptation, la convergence des séquences déformées par le canal et aussi son efficacité en matière d'adaptation. On voit bien que toutes les séquences tendent vers les séquences originales à part celles qui sont utilisées pour l'initialisation du filtre (0,0) et celles qui sont calculées graduellement pour arriver aux valeurs optimales (leur nombre est élevé par rapport au LMS).

La condition pour que l'égaliseur NLMS résolve la déformation du canal, est de prendre la fonction de transfert $w(z)=1/C(z)$ telle que la convolution $w(z)*C(z)=1$. La figure 3.23 illustre la convolution de l'égaliseur $w(z)$ avec le canal $C(z)$ qui est égale à 1.

Pour plusieurs rapports signal/bruit SNR=30, 20 et 15 dB, la convergence des coefficients est assurée (figure 3.24), l'erreur ne diverge pas, elle tend vers zéro à partir de 200 itérations. Tant que le bruit augmente l'erreur augmente aussi, le bruit a une grande influence sur la qualité de réception, mais en revanche le filtre NLMS montre une bonne robustesse vis-à-vis des bruits externes.

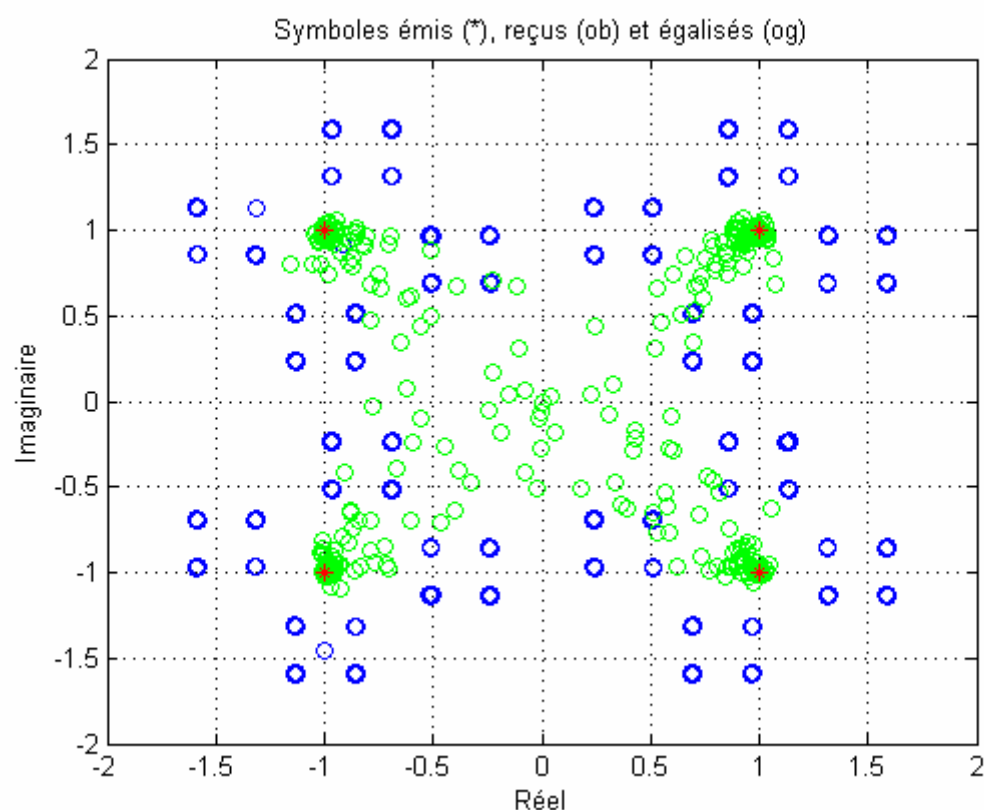


Fig.3.22 — Constellation des signaux émis, reçus (bleu) & égalisés (vert) avec le NLMS

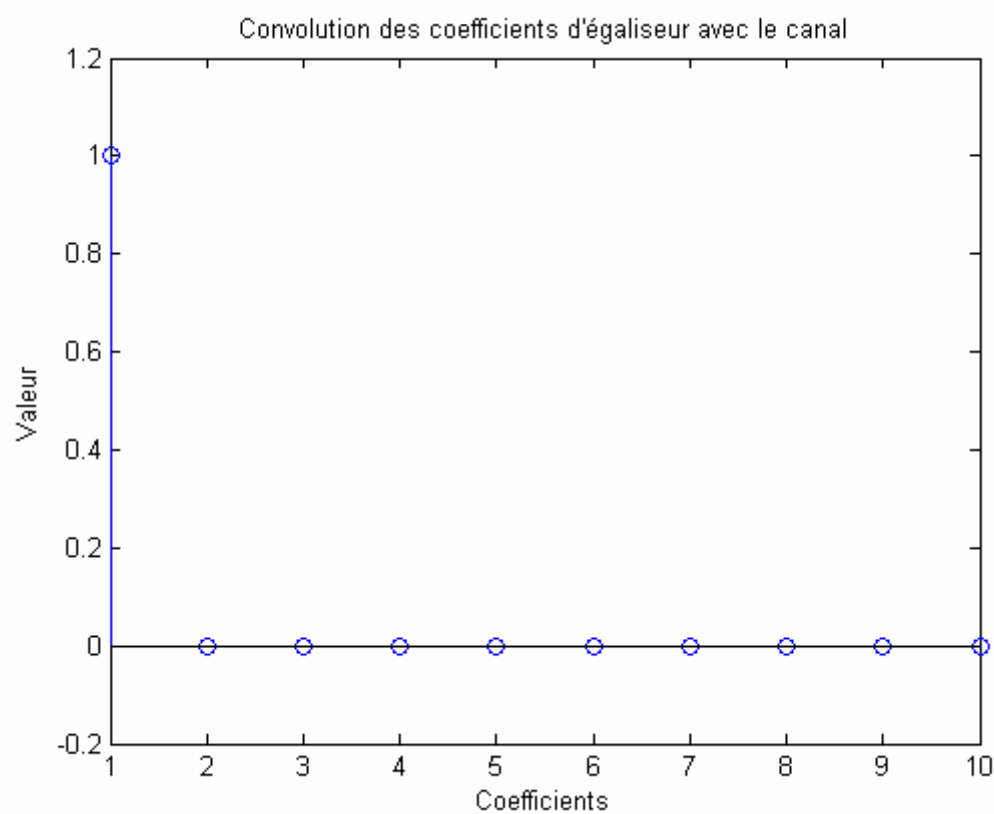


Fig.3.23 — Convolution du canal $C(z)$ avec l'égaliseur NLMS

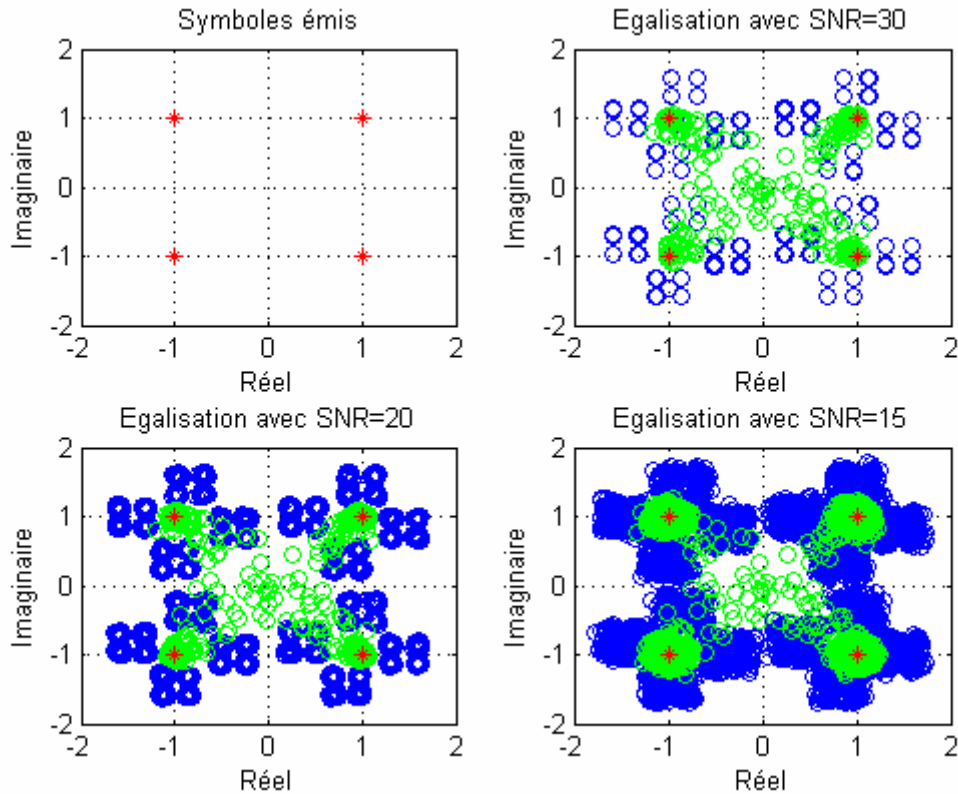


Fig.3.24 — Performances & robustesse du NLMS avec des SNR = 30, 20 & 15 dB

3.4 Algorithme RLS

Dans l'algorithme LMS, la correction appliquée à l'estimation précédente de ce produit est constituée de trois facteurs : la taille d'étape μ , le signal d'erreur $e(n)$ et le vecteur d'entrée $x(n)$. D'autre part, dans l'algorithme RLS cette correction se compose du produit de deux facteurs : l'erreur de l'estimation $e(n)$ et le vecteur de gain $k(n)$, (le vecteur de gain lui-même se compose de $R_{yy}^{-1}(n)$), l'inverse de la matrice de corrélation déterministe multiplié par le vecteur d'entrée $y(n)$. La principale différence entre les algorithmes LMS et RLS est donc la présence de $R_{yy}^{-1}(n)$.

L'application de cet algorithme sur le canal de transmission proposé est présentée par les figures 3.25 et 3.26 ; l'erreur MSE et les coefficients de l'égaliseur RLS sont explicités. La convergence de ces coefficients semble plus rapide. En effet, les valeurs des coefficients sont adaptées et stabilisées après environ 20 séquences.

La pertinente remarque pour ce type d'algorithme est la vitesse de convergence qui est très rapide et très stable comme le montre la poursuite dans la figure 3.27.

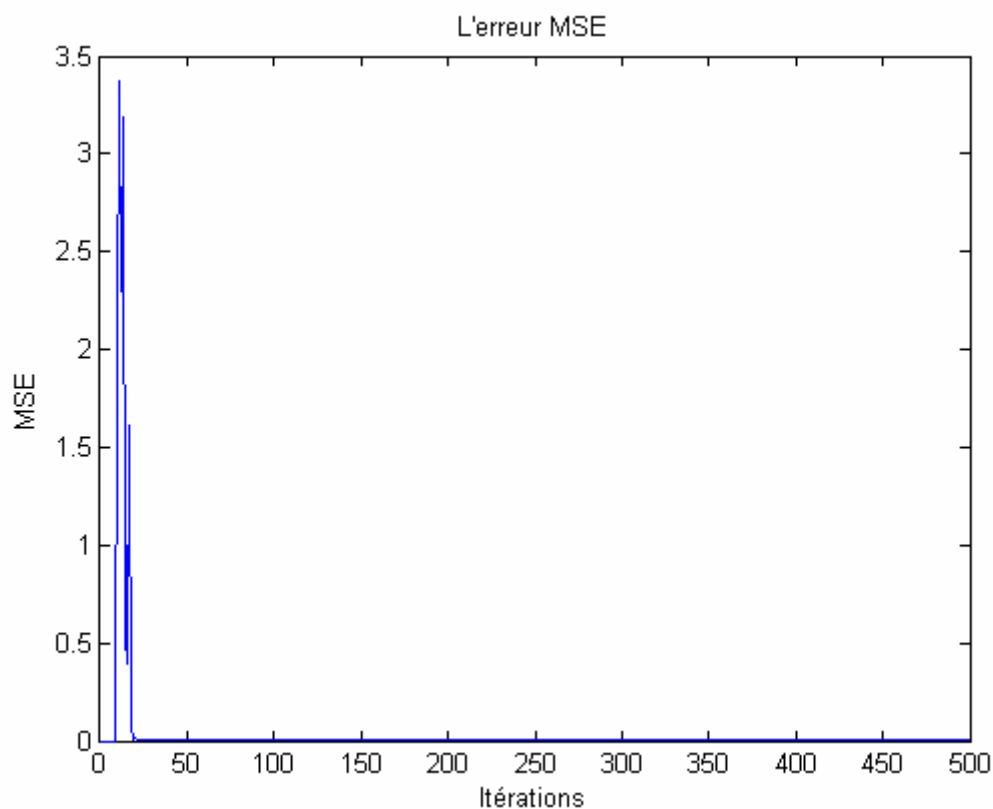


Fig.3.25 — L'erreur MSE du RLS sur 500 itérations

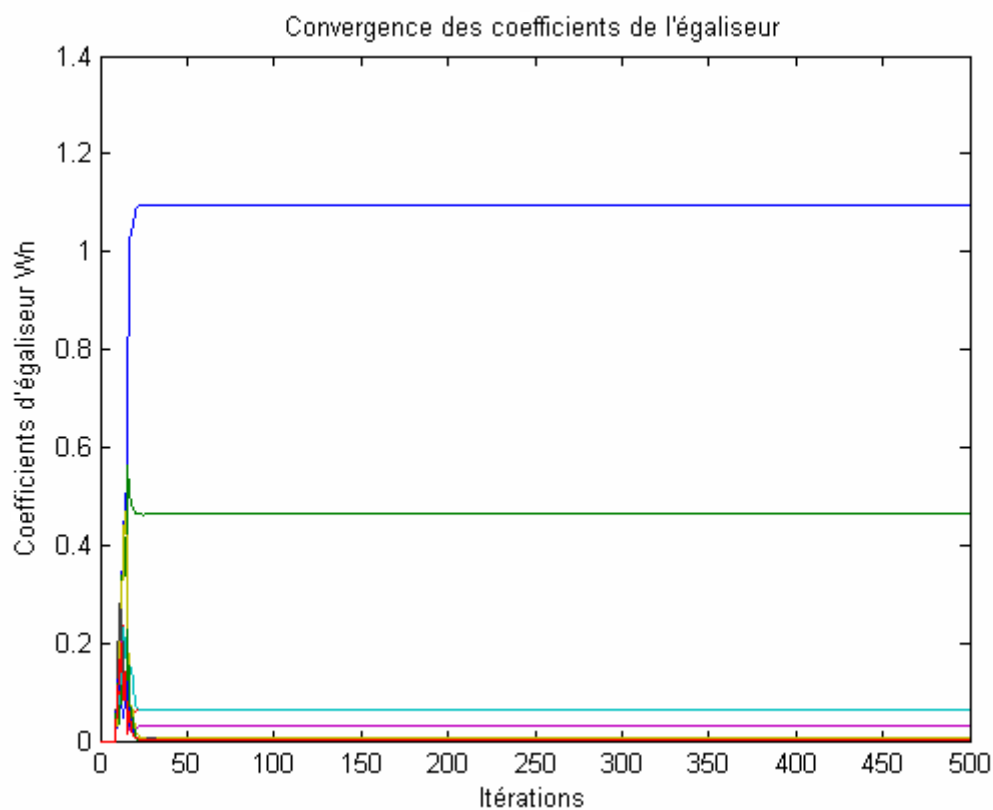


Fig.3.26 — Convergence des coefficients de l'égaliseur RLS sur 500 itérations

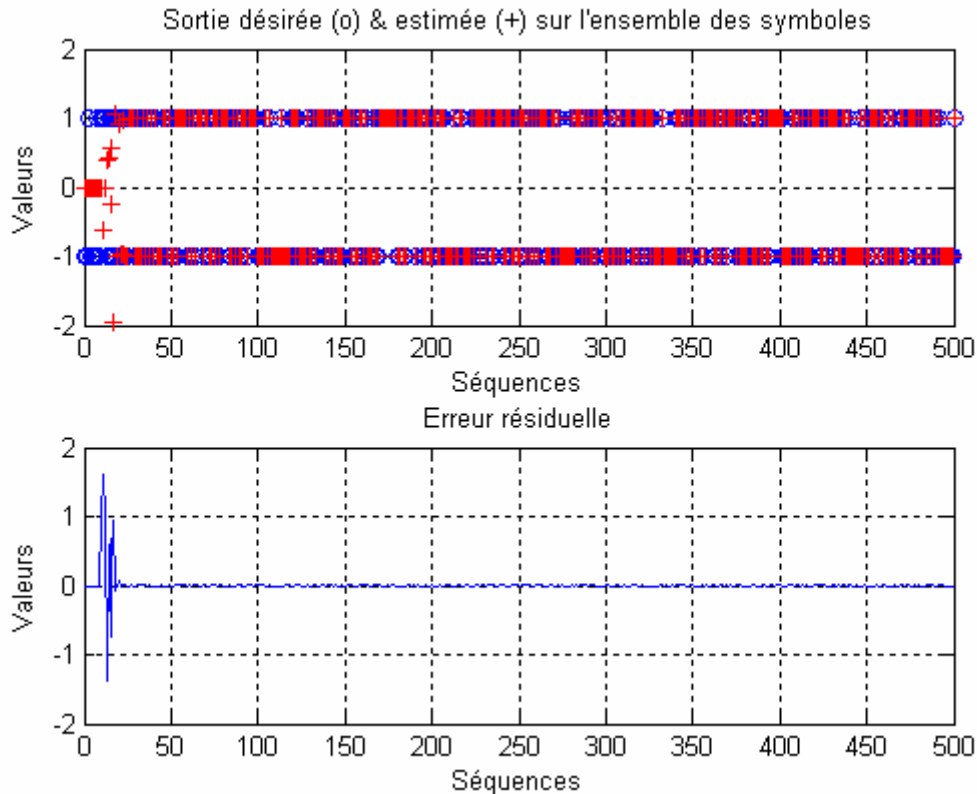


Fig.3.27 — Capacité de poursuite & l'erreur résiduelle du RLS sur 500 itérations

La constellation présentée en figure 3.28 illustre le bon fonctionnement de l'algorithme RLS en ce qui concerne l'erreur d'adaptation, la convergence des séquences déformées par le canal et aussi sa grande efficacité en matière d'adaptation par rapport au deux précédents. On voit bien que toutes les séquences tendent vers les séquences originelles à part 10 séquences qui sont utilisées pour l'initialisation du filtre (0,0) et celles qui sont calculées graduellement pour arriver jusqu'aux valeurs optimales (leur nombre n'est pas élevé par rapport aux LMS et NLMS ; il est de l'ordre de 10 séquences).

La condition pour que l'égaliseur RLS résolve la déformation du canal, est de prendre la fonction de transfert $w(z)=1/C(z)$ telle que la convolution $w(z)*C(z)=1$. La figure 3.29 illustre la convolution de l'égaliseur $w(z)$ avec le canal $C(z)$, elle est égale à 1.

Pour plusieurs rapports signal/bruit SNR=30, 20 et 15 dB, la convergence des coefficients est assurée (figure 3.30), l'erreur ne diverge pas, elle tend vers zéro à partir de 20 itérations ce qui est vraiment rapide. Tant que le bruit augmente l'erreur augmente aussi ; comme les deux précédents algorithmes, le bruit a une grande influence sur la qualité de réception, mais en revanche le filtre RLS montre une grande robustesse vis-à-vis des bruits externes.

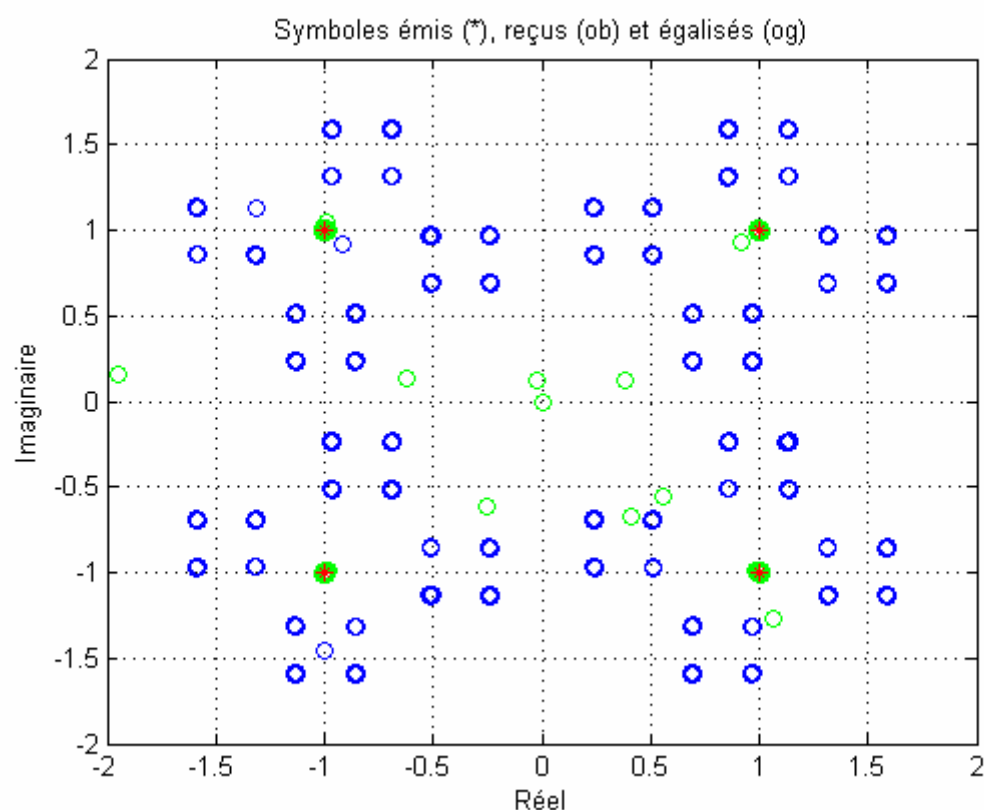


Fig.3.28 — Constellation des signaux émis, reçus (bleu) & égalisés (vert) avec le RLS

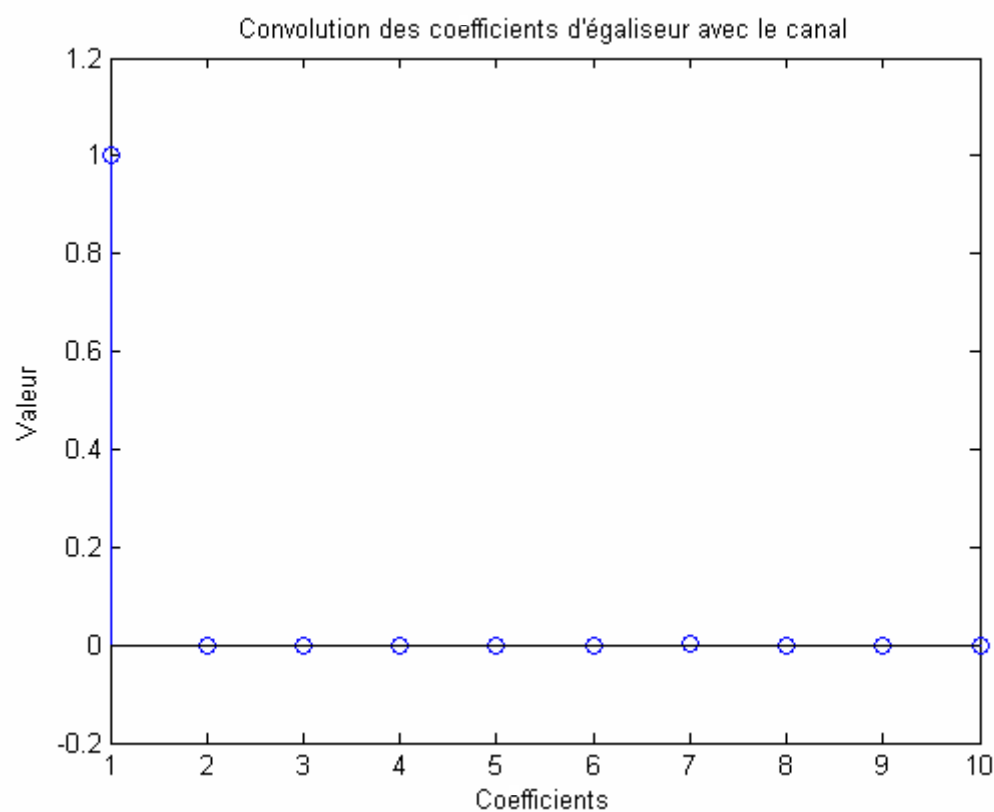


Fig.3.29 — Convolution du canal $C(z)$ avec l'égaliseur RLS

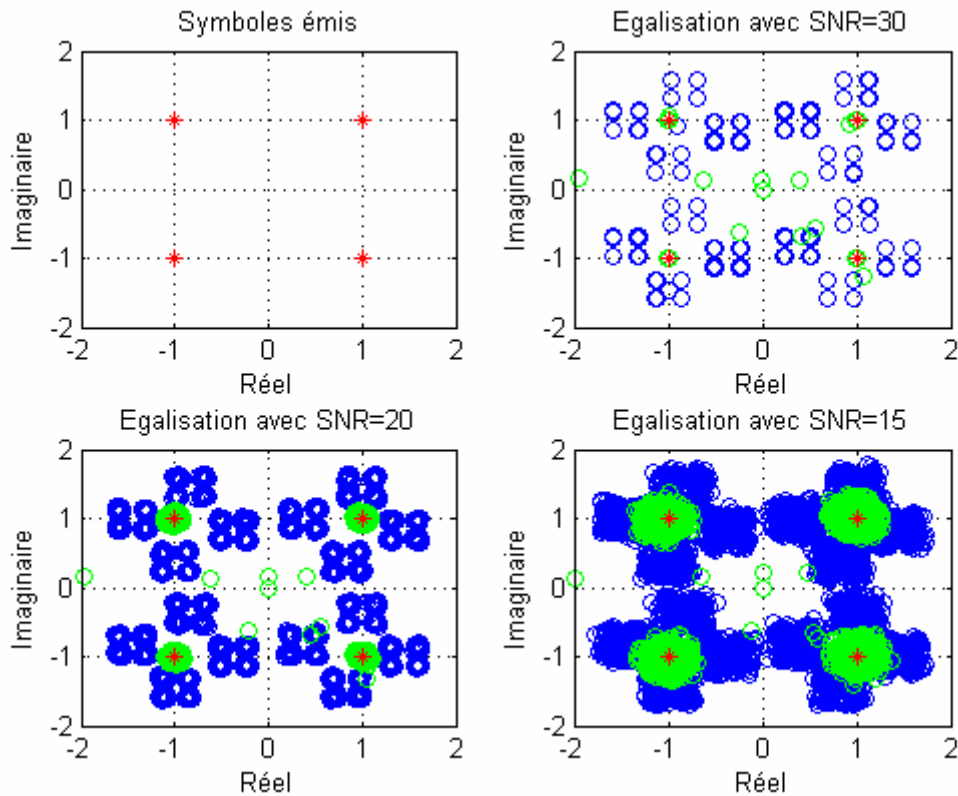


Fig.3.30 — Performances & robustesse du RLS avec des SNR = 30, 20 & 15 dB

3.5 CMA

Dans cette partie on analyse les performances d'un égaliseur aveugle (sans apprentissage) l'algorithme CMA est l'algorithme le plus utilisé dans ces applications. Le signal QPSK transmis dans le canal défini précédemment sera égalisé avec le critère de module constant CM, où seules les caractéristiques statistiques du signal sont connues au récepteur. On force le signal de sortie du canal à prendre une valeur fixe qui est dans notre cas $\sqrt{2}$ (le module du signal QPSK $\sqrt{1^2 + 1^2} = \sqrt{2}$).

On initialise l'égaliseur avec $w(1) = 1$. Les coefficients sont ajustés par l'équation de mise à jour. Les figures 3.31 et 3.32 montrent la convergence de la fonction de coût (MSE) et les coefficients w_n de l'égaliseur pour un pas approprié $\mu=0.009$.

La figure 3.33 représente la poursuite des séquences estimées par l'égaliseur $y(n)$ et celles désirées $x(n)$. L'algorithme CMA présente une poursuite acceptable des données avec une grande erreur au début et qui tend après à zéro.

La rapidité de convergence est assurée avec le bon pas choisi μ , et c'est au bout de 250 itérations qu'on voit l'algorithme CMA se stabiliser (lent par rapport aux LMS, NLMS et RLS).

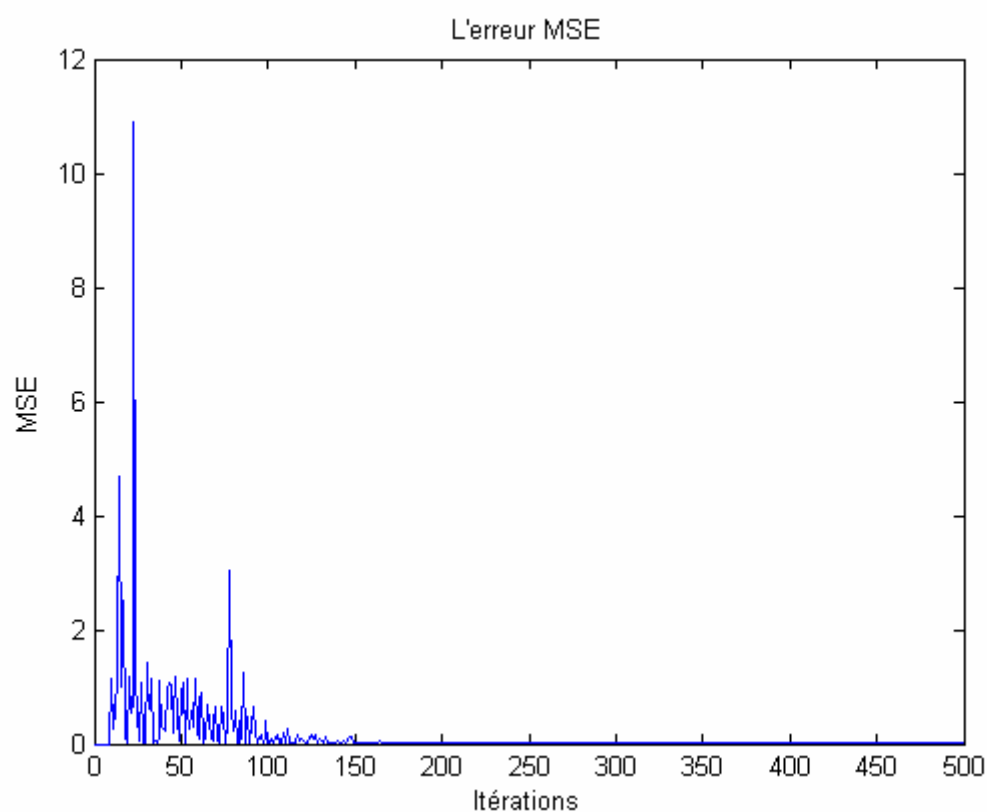


Fig.3.31 — L'erreur MSE du CMA sur 500 itérations

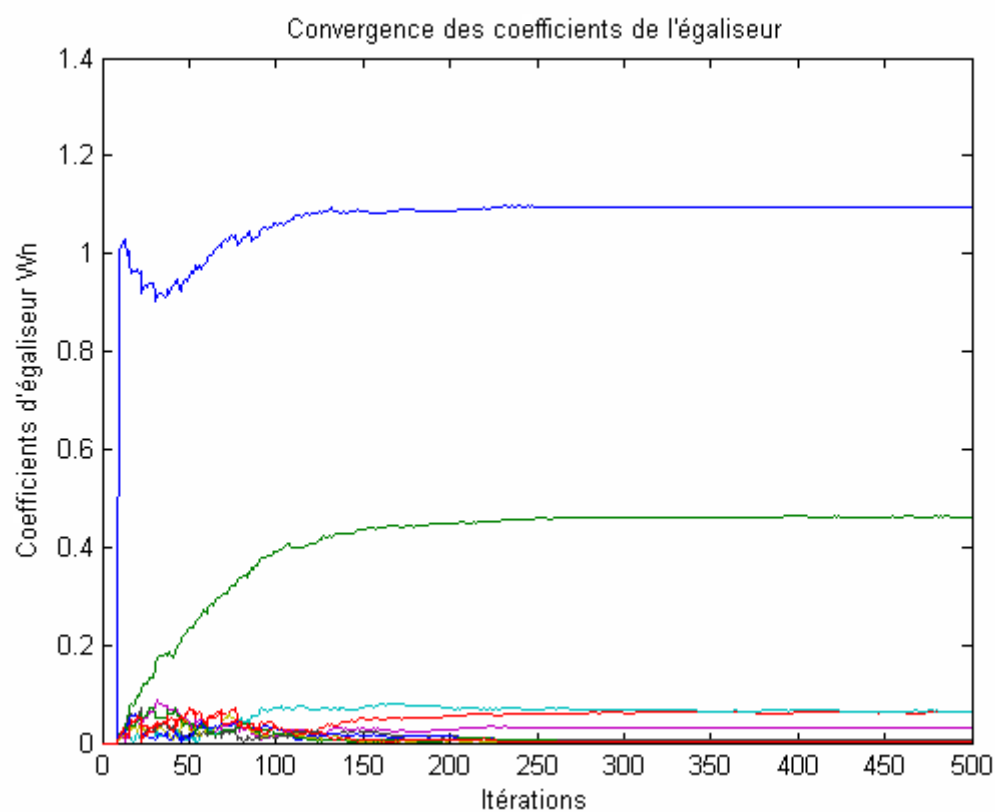


Fig.3.32 — Convergence des coefficients de l'égaliseur CMA sur 500 itérations

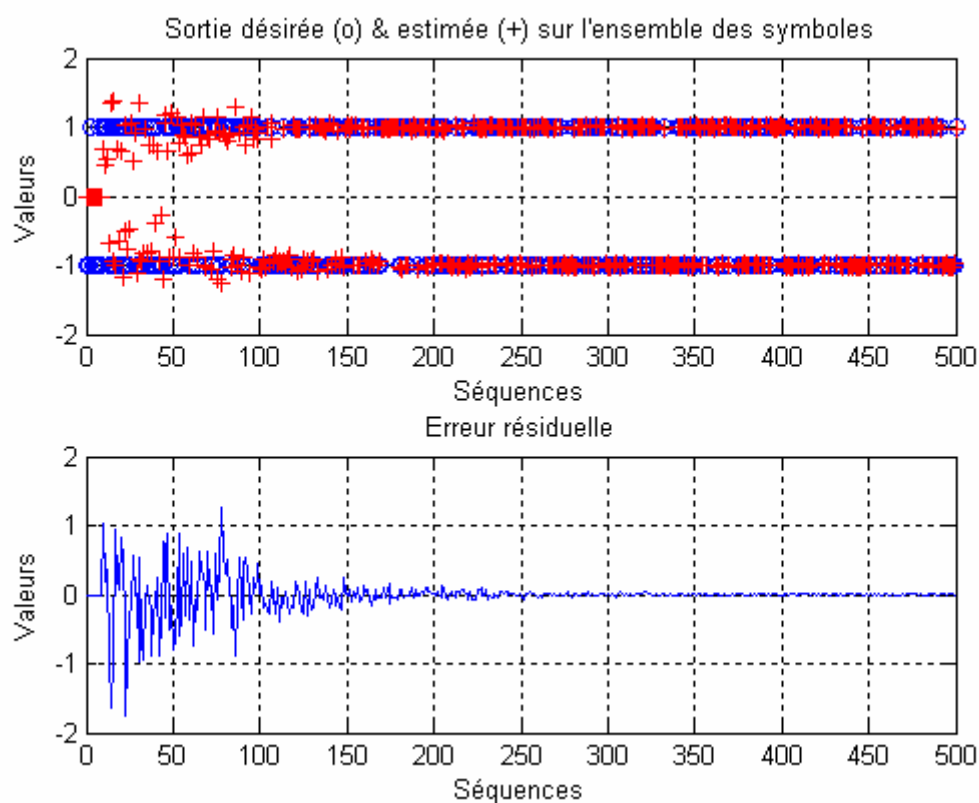


Fig.3.33 — Capacité de poursuite & l'erreur résiduelle du CMA sur 500 itérations

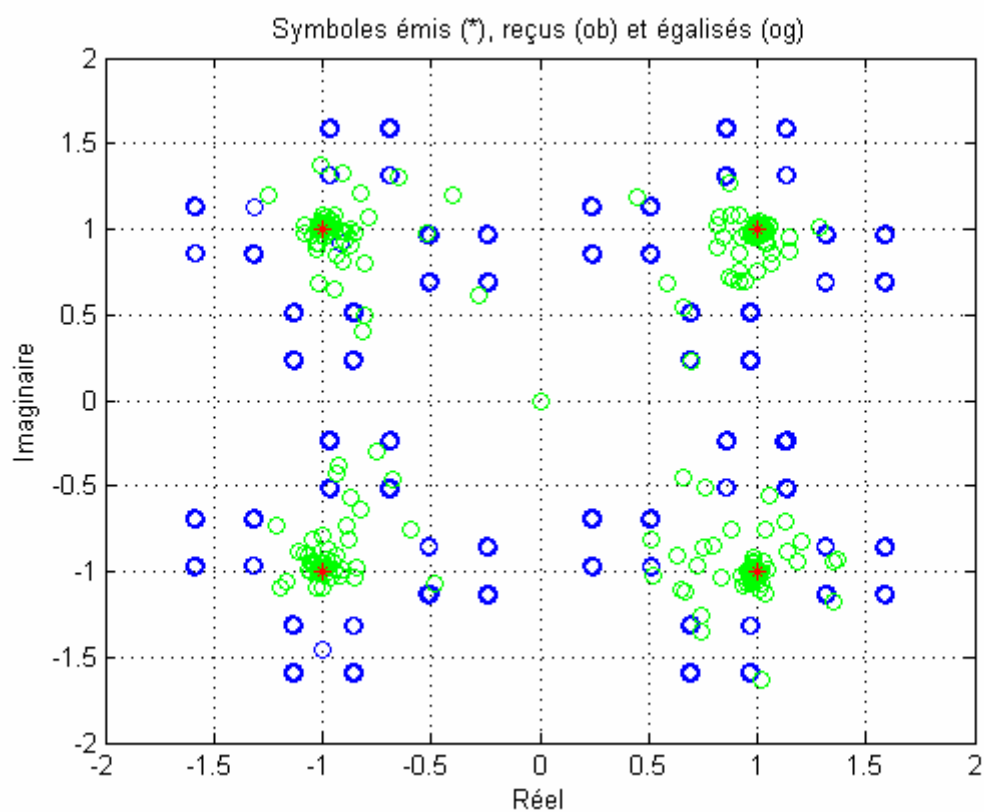


Fig.3.34 — Constellation des signaux émis, reçus (bleu) & égalisés (vert) avec le CMA

La constellation présentée en figure 3.34 illustre le bon fonctionnement de l'algorithme CMA (l'égalisation sans apprentissage) en ce qui concerne l'erreur d'adaptation, la convergence des séquences déformées par le canal et aussi son efficacité en matière d'adaptation aveugle, on voit bien que toutes les séquences tendent plus au moins vers les séquences originelles à part celles qui sont utilisées pour l'initialisation du filtre (0,0) et celles qui sont calculées graduellement pour arriver aux valeurs optimales (un grand nombre par rapport aux algorithmes précédents).

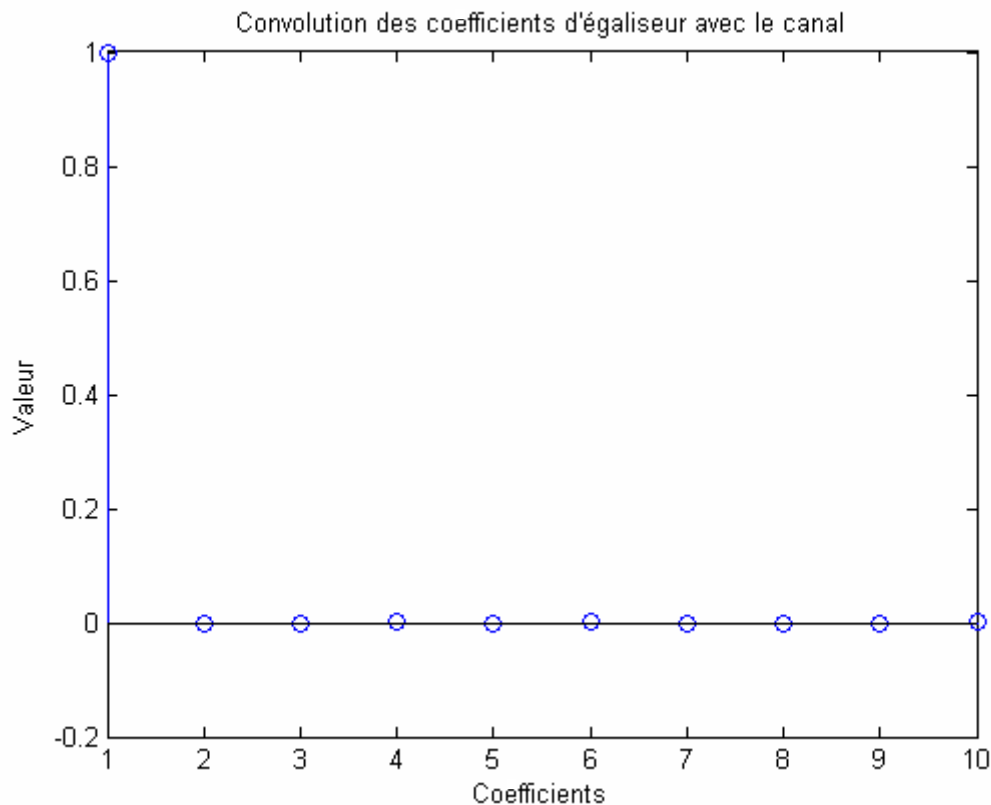


Fig.3.35 — Convolution du canal $C(z)$ avec l'égaliseur CMA

La condition pour que l'égaliseur CMA annule la déformation du canal, est de prendre la fonction de transfert $w(z)=1/C(z)$ telle que la convolution $w(z)*C(z)=1$. La figure 3.35 illustre la convolution de l'égaliseur $w(z)$ avec le canal $C(z)$ qui est égale à 1.

Comme la convergence de l'algorithme LMS dépend de la valeur du pas μ , (pas $\mu=0.023$ (figure 3.36) et un pas $\mu=0.0001$ (figure 3.37)) l'erreur MSE, la valeur des coefficients du filtre $w(z)$ et la constellation du signal égalisé ont du diverger. La convolution du canal avec l'égaliseur n'étant pas à 1 ($w(z)*C(z)\neq 1$) ce qui montre l'importance du choix du pas d'apprentissage μ .

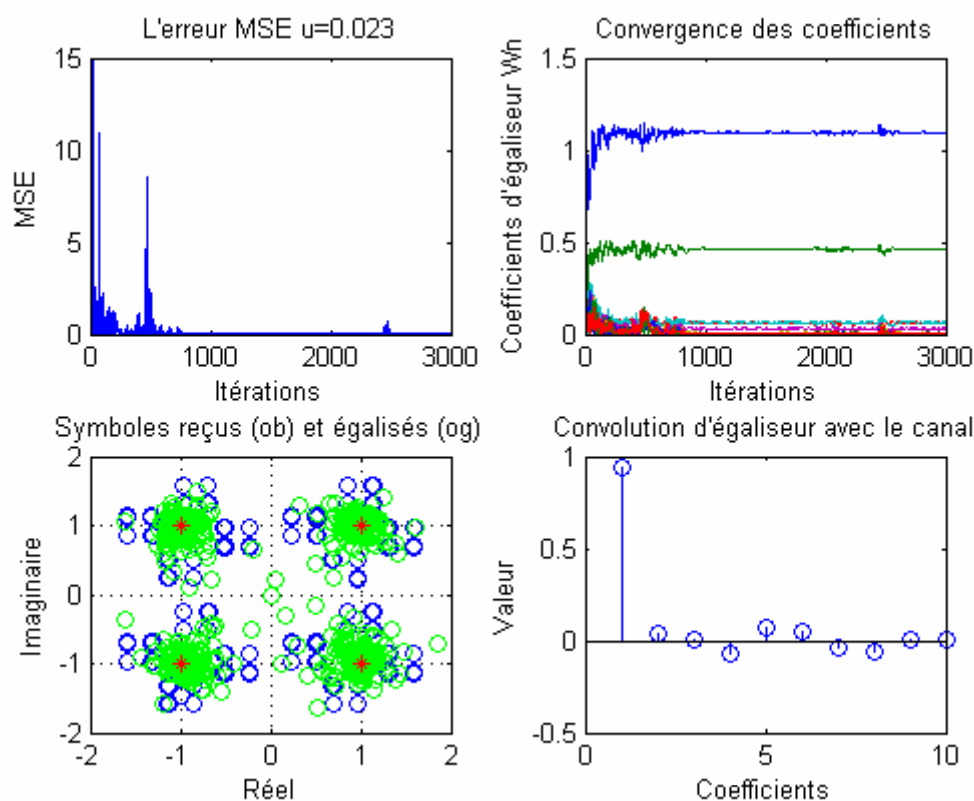


Fig.3.36 — Performances du CMA avec un pas $\mu = 0.023$

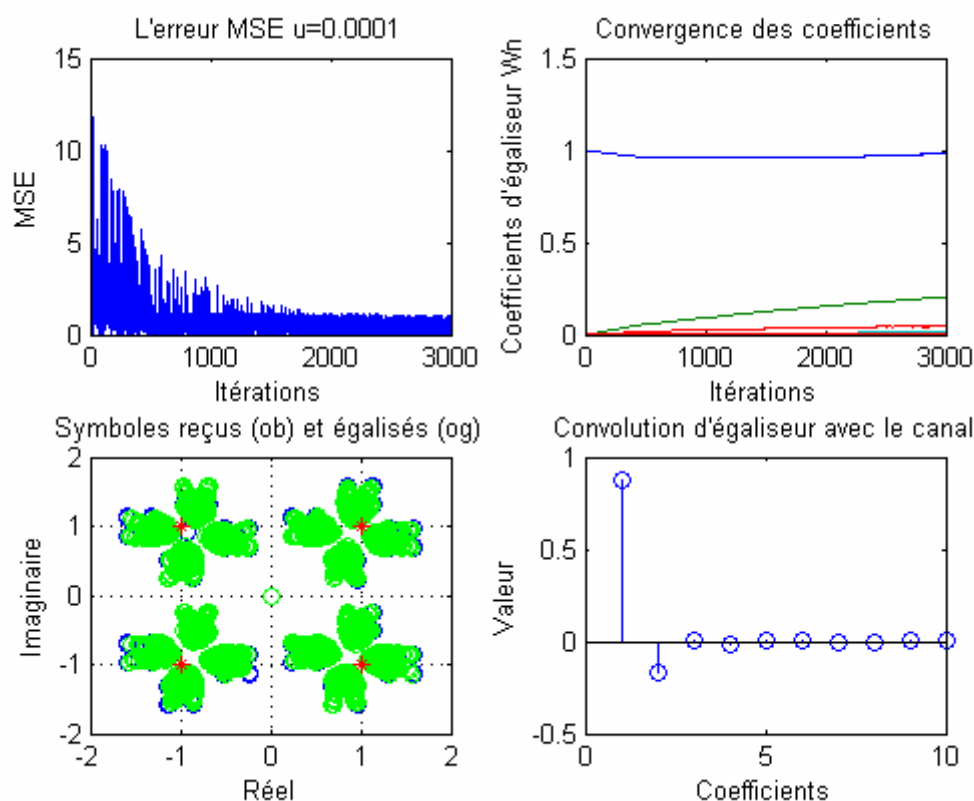


Fig.3.37 — Performance du CMA avec un pas $\mu = 0.0001$

Pour le pas approprié $\mu=0.009$, et plusieurs rapports signal/bruit SNR=30, 20 et 15 dB, la convergence des coefficients est assurée (figure 3.38), l'erreur ne diverge pas, elle tend vers zéro à partir de 250 itérations. Tant que le bruit augmente l'erreur augmente aussi, le bruit a une grande influence sur la qualité de réception, mais en revanche le filtre CMA se montre d'une bonne robustesse aux bruits externes.

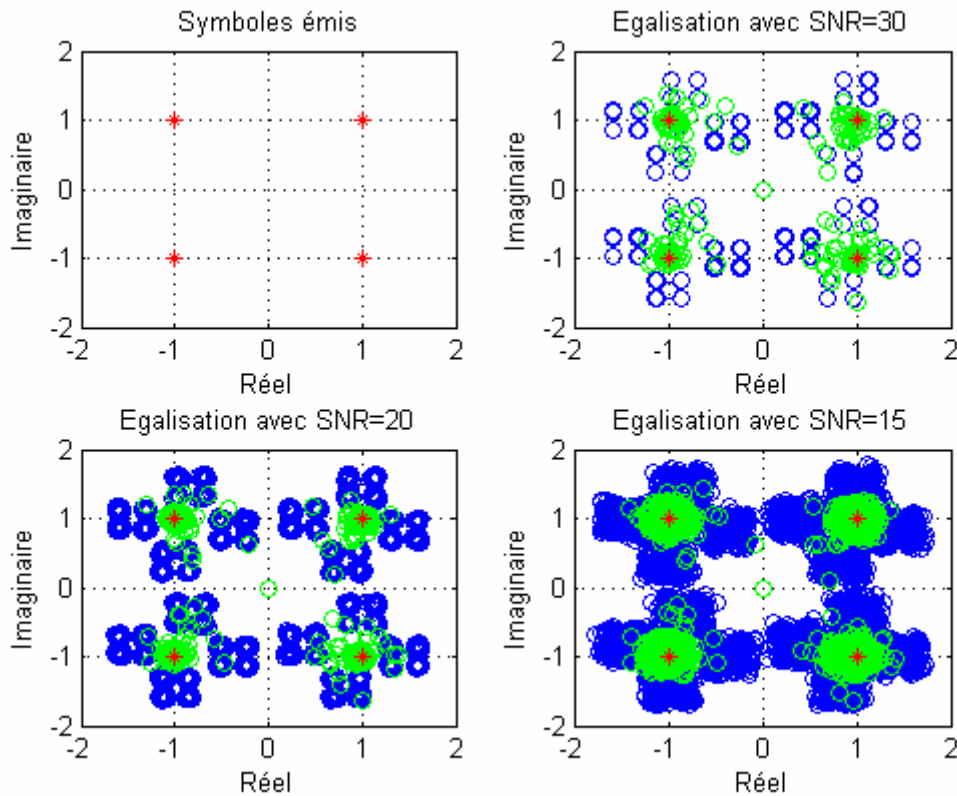


Fig.3.38 — Performances & robustesse du CMA avec des SNR = 30, 20 & 15 dB

4. Conclusion

Lors d'une transmission de données, le canal de transmission peut avoir plusieurs effets sur le signal transmis de l'émetteur au récepteur. Le canal est souvent symbolisé comme une source de bruit additif mais il peut aussi atténuer fortement certaines fréquences porteuses, on parle alors de fading sélectif. Le canal peut aussi avoir pour effet de "mélanger" les symboles transmis, on parle alors d'interférences entre symboles.

L'égalisation adaptative consiste à palier le mieux possible les déformations apportées par le canal de transmission au signal utile grâce à des filtres qui changent les coefficients de ces filtres en temps réel.

Dans ce chapitre, nous avons décrit les structures de certains égaliseurs en mettant en relief leurs algorithmes d'adaptation ainsi que leurs propriétés et performances. Plus précisément, les algorithmes LMS, NLMS, RLS et CMA sont explicités et dont les performances sont récapitulées ci-après :

LMS : Converge, rapide, stable, simple, robuste, peut tomber dans les minimums locaux.

NLMS : Converge, rapide, un pas normalisé pour une meilleure stabilité, robuste, une variante du LMS

RLS : Converge, très rapide, stable, robuste, une très bonne poursuite, calcul complexe et long.

CMA : Pour l'égalisation aveugle, basée sur le module constant des séquences transmises, convergence lente, stable et robuste, basé sur le LMS.

C'est la nature de l'application qu'on veut établir et les méthodes de calculs disponibles qui peuvent déterminer quel type d'algorithme choisir. A noter que les algorithmes LMS et NLMS nécessitent moins de calcul à chaque étape mais convergent plus lentement que l'algorithme RLS.

Le CMA converge lentement par rapport aux autres algorithmes, mais ne nécessite pas de séquences d'apprentissage.

Optimisation neuronale

CHAPITRE

4

DANS CE CHAPITRE

- ☒ Introduction
- ☒ L'optimisation
- ☒ Les réseaux de neurones artificiels (RNA)
- ☒ MultiLayer Perceptron (MLP)
- ☒ Radial Basis Function (RBF)
- ☒ Comparaison entre les égaliseurs MLP & RBF
- ☒ Conclusion

1. Introduction

Les techniques d'égalisations peuvent être classées en deux ensembles (figure 4.1) : égaliseurs linéaires et égaliseurs non linéaires [32]. Les techniques linéaires vues dans le chapitre précédent ont des performances limitées face à des canaux fortement dispersifs ou non linéaires ce qui a encouragé le développement de nouvelles techniques non linéaires.

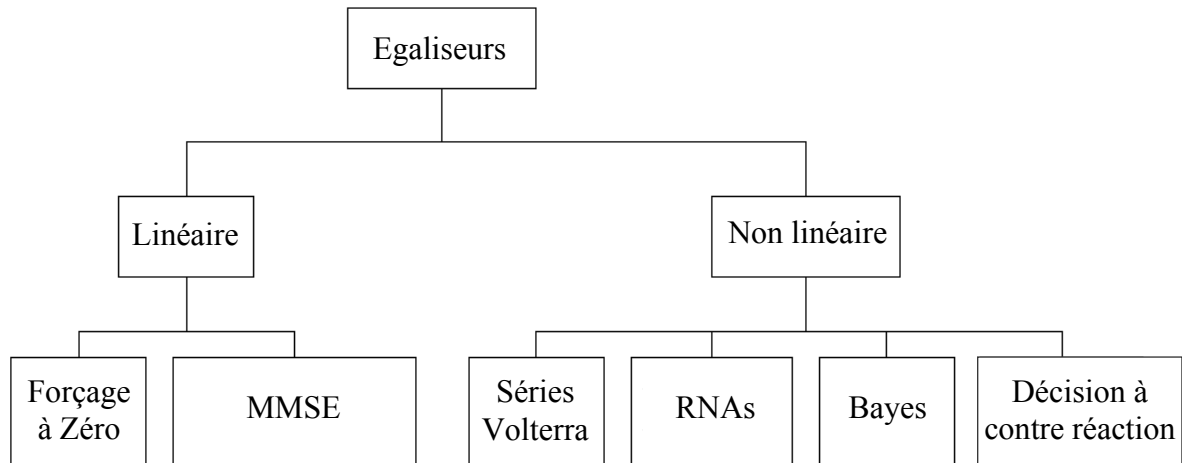


Fig.4.1 — Un groupement possible des techniques d'égalisations

1.1 Egalisation linéaire

L'égalisation linéaire est une technique commune pour annuler l'effet ISI et MAI qui consiste à adapter les coefficients d'un filtre transversal ou FIR jusqu'à ce qu'une fonction coût soit minimisée. Pour que le problème soit facile à formuler mathématiquement, la fonction coût doit être différentiable avec le respect des coefficients du filtre. Le choix le plus utilisé pour cette fonction coût est l'erreur quadratique moyenne (MSE).

Les algorithmes populaires employés pour cette égalisation sont le LMS et le RLS vus précédemment.

1.2 Egalisation non linéaire

Dans d'autres situations le canal peut avoir des distorsions non linéaires, si la distorsion est sévère, les égaliseurs linéaires auront de pauvres performances, dans ce cas les égaliseurs non linéaires doivent être employés mais, en général, la solution mathématique pour ce modèle est compliquée [33].

Une alternative a été trouvée dans les réseaux de neurones artificiels. Les réseaux de neurones comme le multilayer perceptrons (MLP) et le réseau radial basis function (RBF) ont été appliqués à l'égalisation du canal. Ceci a été fait en tournant le problème d'égalisation à un problème de classification : les données émises sont considérées comme des symboles qui appartiennent à un ensemble fini, le réseau comme un classifieur qui détermine quelle est la nature du symbole émit.

2. L'optimisation

2.1 Qu'est ce que l'optimisation ?

L'optimisation est un processus qui rend un concept plus performant en appliquant des variations sur l'état initial de ce concept et utiliser les informations disponibles pour améliorer son rendement.

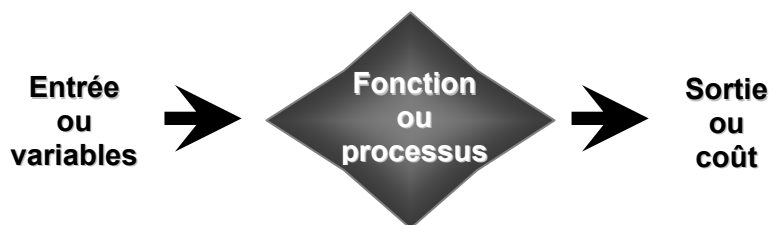


Fig.4.2 — Diagramme d'une fonction ou processus à optimiser

L'ordinateur est l'outil parfait pour l'optimisation tant que le processus qu'on cherche à optimiser ou les variables qui l'influencent peuvent être entrées en format numérique. L'alimentation de l'ordinateur avec des données appropriées permet à ce dernier de fournir une solution. Est-ce que c'est la seule solution ? La plupart du temps non. Est-ce la meilleure solution ? C'est une question difficile. L'optimisation est l'outil mathématique que nous comptons dessus pour obtenir ces réponses.

2.2 Catégories d'optimisation

La figure 4.3 divise les algorithmes d'optimisations en six catégories :

- 1- Trial & error optimisation réfère au processus qui ajuste les variables affectant la sortie sans trop savoir d'informations sur le processus qui produit les sorties. La découverte de la pénicilline comme un antibiotique est le résultat de cette approche d'optimisation. D'autre part une formulation mathématique décrit la fonction objective dans l'optimisation de fonction.

- 2- S'il y'a une seule variable l'optimisation est unidimensionnelle, un problème avec plusieurs variables requière une optimisation multidimensionnelle. Plusieurs approches d'optimisations multidimensionnelles généralisent une série d'optimisations unidimensionnelles.
- 3- L'optimisation "dynamique" veut dire que la sortie est une fonction de temps, pendant que "statique" veut dire que la sortie est indépendante du temps. Quel est le meilleur chemin pour aller au travail ? d'un point de vue distance ; le problème est statique, trouver le chemin le plus rapide ; le problème devient dynamique (de nos jours le chemin le plus court n'est pas nécessairement le plus rapide).
- 4- L'optimisation peut aussi être distinguée par des variables discrètes ou continues.
- 5- Les variables ont souvent des limites ou des contraintes, l'optimisation contrainte incorpore les égalités et les inégalités des variables dans la fonction de coût. L'optimisation sans contrainte permet aux variables de prendre n'importe quelle valeur.
- 6- Quelques algorithmes essaient de minimiser une fonction coût à partir d'initiales valeurs affectées aux variables, Ces algorithmes peuvent facilement tomber sur les minimums locaux, mais tendent à être rapides, d'autre part, les méthodes aléatoires emploient quelques calculs probabilistes pour trouver l'ensemble des variables ; elles tendent à être lentes mais elles ont un grand succès pour trouver les minimums globaux [34].

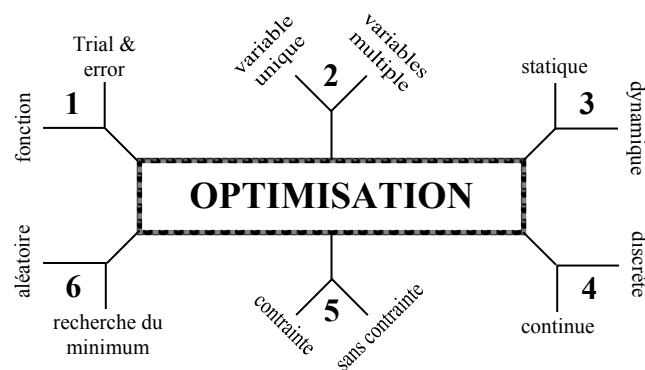


Fig.4.3 — Catégories d'optimisation

L'optimisation est guidée par une fonction coût associée à chaque variable dans l'ensemble des entrées à travers plusieurs techniques, méthodes et algorithmes de recherche. Ce genre de techniques est montré dans l'arbre suivant (figure 4.4). Cette figure montre les techniques naturelles entre autres procédures de recherche bien connues [35].

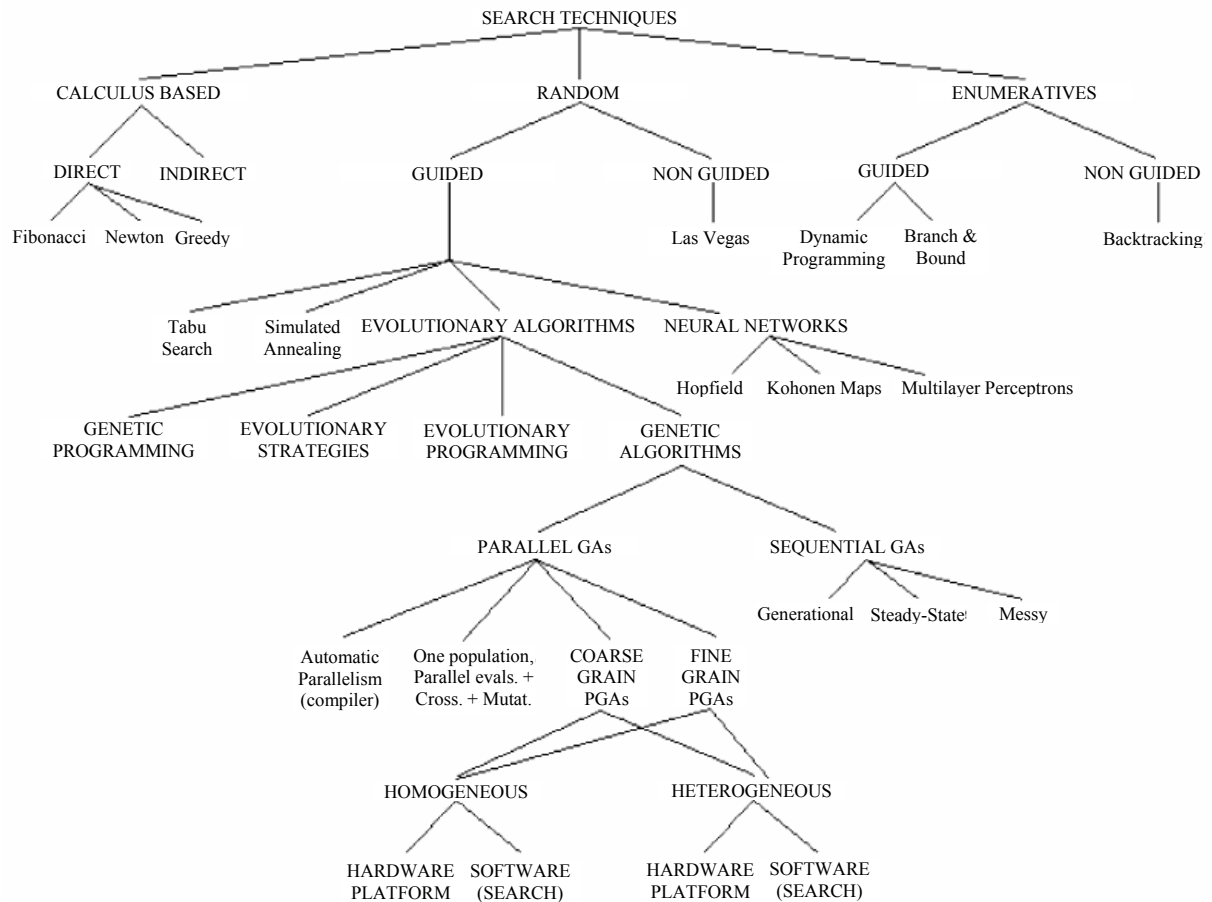


Fig.4.4 — Techniques de recherche et d'optimisation

De nos jours, les réseaux de neurones artificiels (ANN) et les algorithmes génétiques (GA) sont les modèles les plus avancés et vastement utilisés des méthodes d'évolution dans les systèmes d'intelligence artificielle [35].

Dans ce chapitre nous aborderons deux modèles de réseaux de neurones couramment utilisés (MLP et RBF) et nous essaierons de les adapter à notre problématique.

3. Les réseaux de neurones artificiels (ANN)

Inspirés des neurones biologiques, les ANN (artificiel neural networks) traitent l'information qu'ils reçoivent d'une manière analogue aux neurones du cerveau. Ils sont constitués de plusieurs éléments processeurs dits neurones artificiels (figure 4.5) reliés les uns aux autres par un réseau complexe.

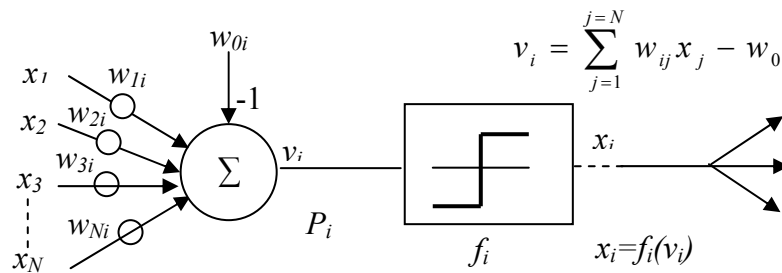


Fig.4.5 — Modèle de neurone artificiel

La sortie de chaque neurone peut être reliée en entrée à plusieurs autres neurones. Chaque neurone effectue une somme pondérée des signaux d'entrée modulés par une fonction dite activation (une fonction non linéaire) et génère une sortie qui sera appliquée aux autres neurones via des connexions pondérées. Avec cette simple structure et en choisissant un nombre approprié de neurones, ces réseaux sont capables d'approximer n'importe quelle fonction continue avec une certaine précision.

3.1 Architecture et fonctionnement

Le facteur déterminant le type d'un réseau de neurones est la nature des connexions entre ses cellules. Selon ce paramètre, les réseaux de neurones peuvent être classés en deux principales catégories suivant la structure des connexions : les réseaux non récurrents (statiques) et les réseaux récurrents (dynamiques). Dans un réseau de la première famille comme le montre la figure 4.6, les réseaux non récurrents ont une structure hiérarchique qui consiste en plusieurs couches ; une couche d'entrée, une ou plusieurs couches cachées et une couche de sortie.

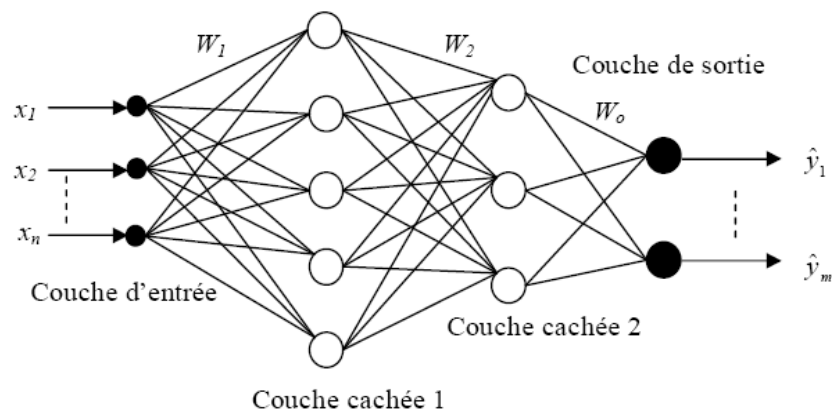


Fig.4.6 — Structure d'un réseau de neurones non récurrent (statique)

La sortie pour chaque neurone dépend uniquement des sorties des neurones précédents dont les signaux circulent de la couche d'entrée à la couche de sortie dans une seule direction. Pour les réseaux dynamiques, plusieurs neurones sont interconnectés pour organiser le réseau où la circulation de l'information est bidirectionnelle, tel que l'état global du réseau dépend de ses états. Dans la littérature de ce type de réseaux, on cite le modèle de Hopfield [36]. En conclusion, nous constatons que les critères motivants le choix du type de réseau sont la simplicité de mise en oeuvre et l'efficacité des algorithmes d'adaptation appelés à répondre aux performances désirées, quelque soient sa taille et sa complexité.

3.2 Fonctions de transfert

Le tableau 4.1 représente les différentes fonctions de transfert qui peuvent être utilisées comme fonction d'activation d'un neurone.

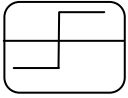
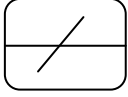
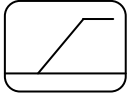
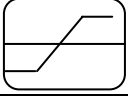
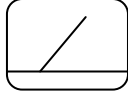
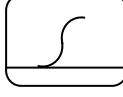
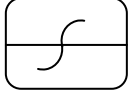
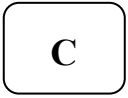
<i>Fonction</i>	<i>Relation d'entrée / sortie</i>	<i>Symbole</i>
<i>Seuil</i>	$a = 0 \text{ si } n < 0$ $a = 1 \text{ si } n \geq 0$	
<i>Seuil symétrique</i>	$a = -1 \text{ si } n < 0$ $a = 1 \text{ si } n \geq 0$	
<i>Linéaire</i>	$a = n$	
<i>Linéaire saturée</i>	$a = 0 \text{ si } n < 0$ $a = n \text{ si } 0 \leq n \leq 1$ $a = 1 \text{ si } n > 1$	
<i>Linéaire saturée Symétrique</i>	$a = -1 \text{ si } n < -1$ $a = n \text{ si } -1 \leq n \leq 1$ $a = 1 \text{ si } n > 1$	
<i>Linéaire positive</i>	$a = 0 \text{ si } n < 0$ $a = n \text{ si } n \geq 0$	
<i>Sigmoïde</i>	$a = \frac{1}{1 + e^n}$	
<i>Tangente hyperbolique</i>	$a = \frac{e^{-n} - e^n}{e^{-n} + e^n}$	
<i>Compétitive</i>	$a = 1 \text{ si } n \text{ maximum}$ $a = 0 \text{ autrement}$	

Tableau.4.1 — Fonctions de transferts d'un neurone

Les trois fonctions de transfert les plus utilisées sont les fonctions seuil, linéaire et sigmoïde [37].

3.3 Entraînement des réseaux de neurones

L'application des réseaux de neurones comprend typiquement deux phases : une phase d'apprentissage et une phase d'opération. L'apprentissage est un processus à travers lequel les paramètres (les poids de connexions) du réseau sont ajustés pour refléter l'information contenue dans la structure du réseau. Utilisant les valeurs d'entrée/sorties désirées, le réseau ajuste ses poids en se basant sur l'erreur de la sortie calculée. Une fois le réseau est entraîné, il représente une base de connaissance statique qui peut être appelée durant la phase d'opération.

L'algorithme de la rétropropagation est la méthode la plus utilisée pour l'entraînement des réseaux de neurones [38]. Il consiste à calculer les dérivées partielles d'une entité (norme d'erreur) par rapport aux paramètres du réseau (poids de connexions) en minimisant une fonction coût. Bien que cet algorithme soit le plus utilisé, il possède deux problèmes : une convergence lente et aussi l'estimation des paramètres risque de tomber dans un minimum local du critère d'optimisation. Pour contourner ces problèmes, plusieurs techniques d'apprentissage ont été adoptées dans la littérature telles que la méthode du moindre carré récurive, le filtre de Kalman, les algorithmes génétiques, ...etc.

4. MultiLayer Perceptron (MLP)

Le perceptron multicouches est un des réseaux de neurones les plus utilisés pour des problèmes d'approximation, de classification et de prédiction. Il est habituellement constitué de deux ou trois couches de neurones totalement connectés.

Le réseau perceptron multicouches est dans la famille générale des réseaux à propagation vers l'avant (Feedforward propagation), c'est-à-dire qu'en mode normal d'utilisation, l'information se propage dans un sens unique, des entrées vers les sorties sans aucune rétroaction. Son apprentissage est de type supervisé, par correction des erreurs. Dans ce cas uniquement, le signal d'erreur est rétropropagé vers les entrées pour mettre à jour les poids des neurones [37].

Le perceptron simple est la base du réseau perceptron multicouches. Un simple perceptron est un perceptron à une seule couche de neurones dont les fonctions d'activation sont de type seuils (Tableau 4.1). Les différentes règles d'apprentissage pour la correction des erreurs utilisées pour le MLP sont la règle LMS (Least Mean Square), l'algorithme de

rétropropagation (backpropagation), la méthode de Newton et la méthode du gradient conjugué. Dans notre travail nous nous basons sur l'algorithme de rétropropagation.

4.1 Algorithme d'apprentissage : La rétropropagation du gradient

L'algorithme BP (*back propagation*), qui est en fait une version généralisée du LMS mais basé sur l'algorithme de la descente du gradient, est appliqué à des structures non linéaires. Son apprentissage est le même que l'algorithme de la descente du gradient.

4.1.1 Etapes d'apprentissage

L'apprentissage consiste à ajuster les coefficients synaptiques pour que les sorties du réseau soient les plus proches possibles des sorties de l'ensemble d'entraînement. Donc il faut spécifier une règle d'apprentissage afin de l'utiliser pour l'adaptation de ces paramètres. On utilise pour cela la méthode de la rétropropagation qui se divise en deux étapes :

Une étape de propagation : qui consiste à présenter une configuration d'entrée au réseau, puis à propager cette entrée de proche en proche de la couche d'entrée à la couche de sortie en passant par les couches cachées.

Une étape de rétropropagation : qui consiste, après le processus de propagation, à minimiser l'erreur commise sur l'ensemble des exemples présentés, erreur considérée comme une fonction des poids synaptiques. Cette erreur représente la somme des différences au carré entre les réponses calculées et celles désirées pour tous les exemples contenus dans l'ensemble d'apprentissage. Simplement dit, la procédure de la rétropropagation repose sur l'idée de rétropropager vers les couches internes (cachées) l'erreur commise en sortie, d'où la méthode tire son nom.

Il est indispensable, avant d'entamer le principe de la rétropropagation, de définir une topologie du réseau, et les relations qui relient les entrées et les sorties d'une part, les entrées et les poids d'autre part [39].

4.1.2 Principe de la rétropropagation

La rétropropagation est basée sur l'adaptation des coefficients synaptiques (coefficients de pondérations) afin de minimiser une fonction de coût (performances) donnée par :

$$J(w) = \sum_{p=1}^T J_p(w) \quad (4.1)$$

$$J_p(w) = \frac{1}{2} \sum_{i=1}^n |y^d(t) - y(t)|^2 \quad (4.2)$$

où

$y^d(t)$: La sortie désirée du réseau.

$y(t)$: La sortie du réseau.

T : Le nombre d'exemples ou la longueur de l'ensemble d'entraînement.

La méthode du gradient consiste à modifier d'une quantité proportionnelle les coefficients synaptiques aux taux de changement de l'écart en fonction du changement de ce même coefficient.

On commence l'entraînement par un choix aléatoire des valeurs des poids, on présente l'entrée et la sortie désirée correspondante. A la sortie du réseau, l'erreur commise et le gradient de l'erreur par rapport à tous les poids sont calculés, ensuite les poids sont ajustés. Cette procédure est répétée jusqu'à ce que les sorties du réseau soient suffisamment proches (avec précision appropriée) des sorties désirées présentées. A ce niveau, l'apprentissage est achevé en donnant un réseau capable d'accomplir une tâche prévue.

4.2 Equations du réseau

Les états des différents neurones d'un réseau multicouches à L couches (couches cachées et couches de sorties) ayant n entrées et m sorties, sont donnés par les équations suivantes:

$$u_i^k(t) = f^k(p_i^k(t)) \quad (4.3)$$

$$p_i^k(t) = \sum_{j=0}^{N-1} w_{ij}^k u_j^k(t) \quad (4.4)$$

où

$$i = 1, 2, \dots, N_k$$

$$k = 1, 2, \dots, L$$

N_k : Nombre de neurones dans la $k^{\text{ème}}$ couche.

L : Nombre total de couches.

$u_i^k(t)$: La sortie du neurone i de la $k^{\text{ème}}$ couche.

$p_i^k(t)$: Potentiel somatique du neurone i de la $k^{\text{ème}}$ couche.

w_{ij}^k : Coefficient synaptique (poids) de la $j^{\text{ème}}$ entrée du neurone i de la couche k .

$$u_0^k(t) = 1 ; k = 1, 2, \dots, L \quad (4.5)$$

$$u_0^k(t) = x_i(t) ; i = 1, 2, \dots, n \quad (4.6)$$

$$u_0^k(t) = y_i(t) ; i = 1, 2, \dots, m \quad (4.7)$$

où

$x_i(t)$: Les entrées du réseau.

$y_i(t)$: Les sorties du réseau.

$f(x)$: La fonction d'activation.

4.2.1 Adaptation des poids

L'adaptation des poids (ajustement, mise à jour des coefficients synaptiques), se fait en se basant essentiellement sur la formule itérative suivante.

$$w_{ij}^k(n+1) = w_{ij}^k - \Delta w_{ij}^k \quad (4.8)$$

$$\Delta w_{ij}^k(n) = \mu \frac{\partial j(w)}{\partial w_{ij}^k(n)} \quad (4.9)$$

où :

n : Le numéro d'itération.

μ : Le pas d'apprentissage.

Il est à noter que le pas d'apprentissage μ influe sur la vitesse de convergence $\frac{\partial j(w)}{\partial w_{ij}^k(n)}$; sa valeur est généralement choisie en respectant un certain compromis entre la vitesse de convergence et la précision des résultats.

La dérivée partielle de la fonction coût par rapport aux poids w_{ij}^k , représente la vitesse de variation de l'erreur en fonction de la vitesse de variation des poids. Sur tout l'ensemble d'entraînement on a :

$$\Delta w_{ij}^k(n) = \mu \sum_{p=1}^T \frac{\partial j_p(w)}{\partial w_{ij}^k(n)} \quad (4.10)$$

Pour permettre d'implémenter l'algorithme, une expression pour la dérivée partielle de $j(w)$ par rapport à chaque poids du réseau pour un choix arbitraire d'une couche k est donnée par la relation (4.11).

$$\frac{\partial j_p(w)}{\partial w_{ij}^k} = \frac{\partial j_p(w) \partial u_j^k(t)}{\partial u_j^k(t) \partial w_{ij}^k} \quad (4.11)$$

$\frac{\partial j_p(w)}{\partial u(t)}$: représente la sensibilité de $j_p(w)$ par rapport à $y(t)$.

Pour la couche de sortie, la sensibilité est donnée par :

$$\frac{\partial j_p(w)}{\partial u(t)} = y_i^L(t) - y_i^d(t) \quad (4.12)$$

Cette expression est appelée erreur de sortie.

Pour les couches cachées :

$$\frac{\partial j_p(w)}{\partial u_i^k(t)} = \sum_{j=1}^{N_{k+1}} \frac{\partial j_p(w)}{\partial u_j^{k+1}(t)} \frac{\partial u_j^{k+1}(t)}{\partial u_i^k(t)} \quad (4.13)$$

Cette expression est appelée erreur de la couche cachée ou erreur équivalente.

L'expression $\frac{\partial u_j^i(t)}{\partial w_{ij}^k}$ est donnée comme suit :

$$\frac{\partial u_j^i(t)}{\partial w_{ij}^k} = f^k(p_j^k(t) \cdot u_i^{k-1}(t)) \quad (4.14)$$

De même, l'expression (4.11) s'écrit sous la forme :

$$\frac{\partial j_p(w)}{\partial w_{ij}^k} = \frac{\partial j_p(w)}{\partial u_i^k(t)} f^k(p_j^k(t) \cdot u_i^{k-1}(t)) \quad (4.15)$$

Les poids du réseau doivent être ajustés après la présentation de tous les exemples, pour minimiser l'erreur totale sur l'ensemble d'entraînement. Cependant, tous les poids peuvent être ajustés après la représentation de chaque exemple, ceci vient du fait que les corrections sont aussi faibles et la minimisation de $j(w)$ est une bonne approximation de la minimisation de $j_p(w)$. La variation des poids $\Delta w_{ij}^k(n)$ peut alors s'écrire ainsi :

$$\Delta w_{ij}^k(n) = \mu \frac{\partial j_p(w)}{\partial w_{ij}^k(n)} \quad (4.16)$$

4.2.2 Problèmes du BP

Bien que l'algorithme de rétropropagation soit l'algorithme le plus utilisé pour l'apprentissage supervisé des réseaux statiques multicouches, son implantation se heurte à plusieurs difficultés techniques. Car on ne trouve aucune méthode permettant d'éclaircir certaines ambiguïtés telles que:

- Trouver une architecture appropriée (nombre de couches, nombre de neurones).
- Choisir une taille et une qualité adéquate des exemples d'entraînement.
- Choisir des valeurs initiales satisfaisantes pour les poids et les paramètres d'apprentissage permettant d'accélérer la vitesse de convergence.
- Problème de la convergence vers un minimum local, (qui empêche la convergence et cause l'oscillation de l'erreur).

Pour éviter le problème d'oscillations, des chercheurs ont proposé une modification de la loi d'apprentissage donnée par l'expression (4.8), à laquelle ils proposent d'ajouter un terme appelé moment [39], donc la loi d'adaptation devient :

$$w_{ij}^k(n+1) = w_{ij}^k(n) - \mu \frac{\partial j_p(w)}{\partial w_{ij}^k(n)} + \alpha [w_{ij}^k(n) - w_{ij}^k(n-1)] \quad (4.17)$$

Avec : $0 \leq \alpha < 1$

4.3 Simulation & résultats

4.3.1 Valeurs complexes

Dans notre modèle cité au chapitre 3, les valeurs du signal et même le canal sont d'une nature complexe. L'algorithme CLMS (LMS complexe) a été introduit dans le chapitre précédent comme un filtre adaptatif linéaire complexe, sa convergence avec le MSE était adressée avec ses propriétés.

Une nouvelle problématique ; les réseaux de neurones ordinaires travaillent sur des valeurs réelles et non complexes, pour remédier à ce problème on doit concevoir un égaliseur complexe non linéaire.

Deux classes de filtres adaptatifs non linéaires complexes ont été présentées [40], fully-complex (entièrement complexe) et split-complex (complexe-divisé), les deux avec une approche duelle univariable.

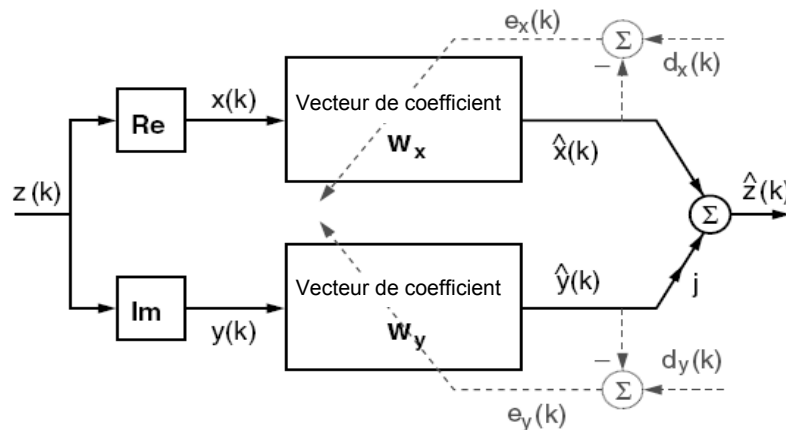


Fig.4.7 — Structure d'un filtre adaptatif split-complex (réel & duel univariable)

Quand on choisit une fonction non linéaire à la sortie du filtre adaptatif complexe (fonction tangente hyperbolique pour notre égaliseur MLP), on a besoin de choisir entre différentiable (fully-complex) et limite (boundedness – split-complex) ; ce choix dépend de

l'application : les filtres split-complex sont généralement plus utilisés dans les réseaux de neurones pour la classification, par contre, les filtres fully-complex constituent plus un choix naturel dans les filtres adaptatifs non linéaires.

Notre choix s'est porté sur la classe split-complex ; son principe est illustré dans la figure 4.7, où on doit traiter les parties réelles et imaginaires du signal indépendamment. Le problème devient alors réel, puis il suffira de rétablir la forme complexe à la sortie du filtre.

4.3.2 Principe

Le principe est le même décrit dans le chapitre précédent ; le même signal complexe généré avec 3000 séquences, déformé par le canal complexe régit par l'équation :

$C(z) = 1 + (-0.25 + 0.34j)z^{-1} + (-0.15j)z^{-2}$. Après, on y ajoute un bruit complexe à son tour, généré avec des SNR = 30, 20 et 15 dB (figure 3.10) pour illustrer l'effet du bruit sur cet égaliseur non linéaire, et le comparer aux méthodes traditionnelles vues précédemment.

Ici le traitement du problème d'égalisation est différent ; l'égaliseur est basé sur le MLP et l'entrée complexe du signal déformé est répartie en deux parties (Principe du split-complex) ; réelle et imaginaire, chaque partie est traitée à part et à la sortie on rend au signal sa forme complexe. Le critère d'apprentissage est l'erreur MSE. La figure 4.8 illustre le principe de la simulation.

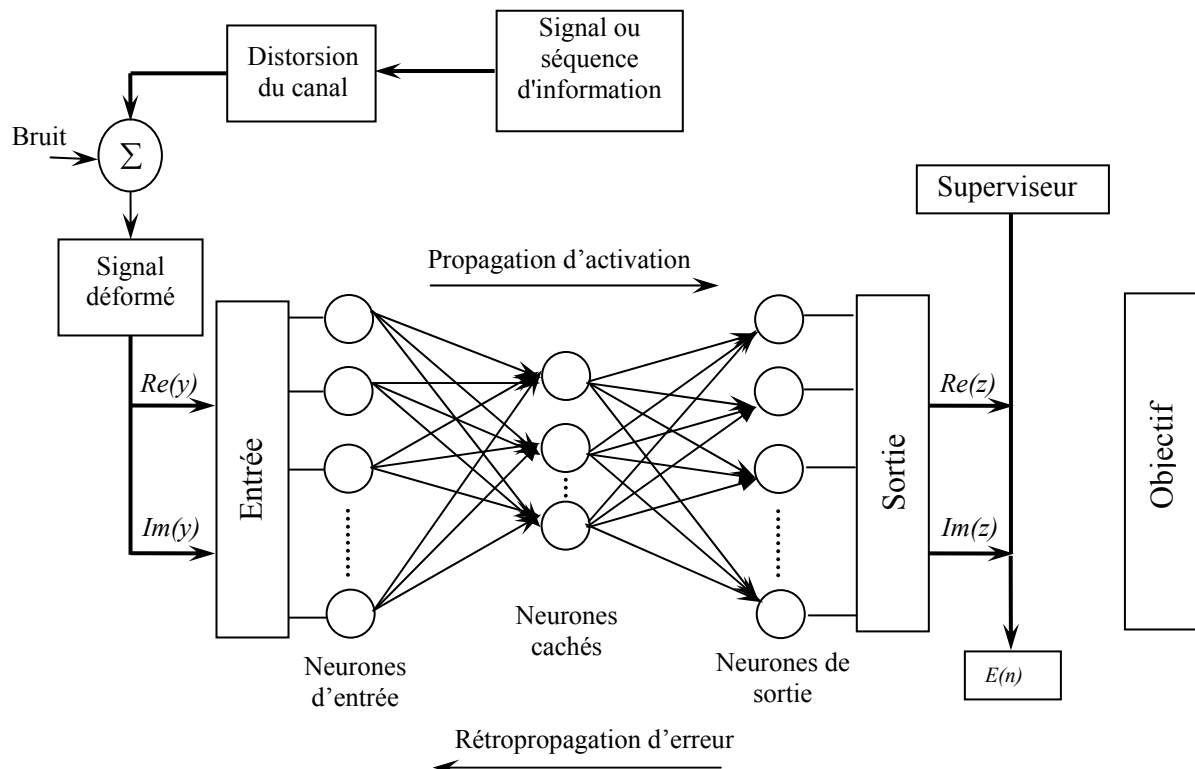


Fig.4.8 — Structure d'un égaliseur neuronal MLP

4.3.3 Etapes de simulation

Des étapes principales sont à suivre pour mettre en œuvre un réseau neuronal et la détermination de ses paramètres ce qui est notre cas avec le MLP qui est caractérisé par le nombre de couches cachées et leur nombre de neurones, le nombre d'itérations, le pas u et les poids w . Ces étapes sont :

- Choisir l'apprentissage. L'apprentissage sera par rétropropagation. La conception et la mise en œuvre de cet algorithme se résument en neuf étapes :

Etape 1 : Initialiser les poids w_{ij}^k et les seuils internes des neurones à de petites valeurs aléatoires.

Etape 2 : Présenter le vecteur d'entrée et de sortie désirée correspondante.

Etape 3 : Calculer la sortie du réseau en utilisant les expressions (4.3) et (4.4).

Etape 4 : Calculer l'erreur commise à la sortie en utilisant l'expression (4.12).

Etape 5 : Réinjecter l'erreur de sortie dans le réseau et calculer l'erreur dans les couches cachées en utilisant l'expression (4.13).

Etape 6 : Calculer le gradient de l'erreur par rapport aux poids utilisant l'expression (4.12).

Etape 7 : Ajuster les paramètres (poids) selon les lois itératives (4.8), (4.9) et (4.10).

Etape 8 : Si la condition sur l'erreur ou sur le nombre d'itération est atteinte, aller à l'étape 9, sinon recommencer avec un autre vecteur d'entrée et aller à l'étape 3.

Etape 9 : Fin.

- Fixer le nombre de couches cachées ; une dans notre cas.
- Choisir la fonction d'activation : dans notre cas «Tangente hyperbolique».
- Déterminer le nombre de neurones par couche cachée.

4.3.4 Résultats et discussions

Après les étapes précédentes le modèle développé pour le MLP est désormais prêt pour sa tâche qui est l'égalisation adaptative des séquences portées au chapitre précédent. Les séquences utilisées durant l'apprentissage ont un SNR=30 dB. Après les étapes d'apprentissage, de validation et de test de l'égaliseur neuronal MLP, on y traite les séquences qui ont un SNR=20 et 15 dB (figure 3.10).

Les séquences sont réparties entre apprentissage et validation ($\frac{2}{3}$ de l'ensemble des séquences) et test ($\frac{1}{3}$ de l'ensemble), chaque étape est évaluée avec le critère MSE.

Apprentissage : Lors de l'apprentissage et en changeant le nombre de neurones graduellement par pas de deux jusqu'à 20 neurones et en calculant à chaque fois l'erreur MSE, on pourra choisir le nombre de neurones qui aura l'erreur minimale (figure 4.9).

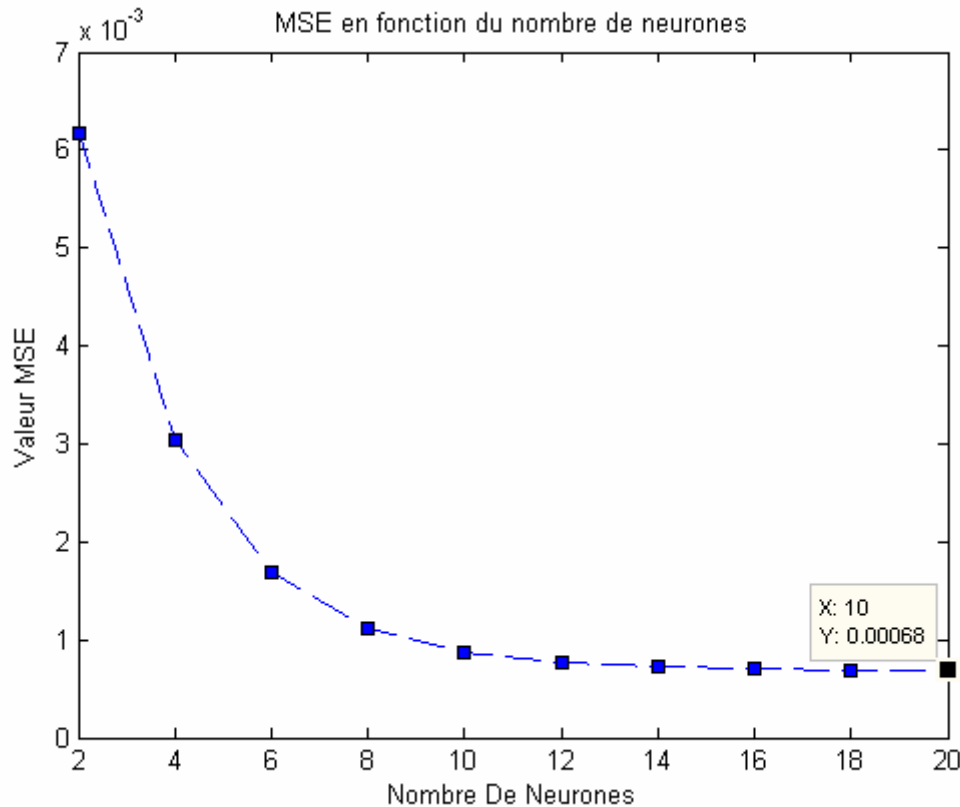


Fig.4.9 — MSE en fonction du nombre de neurones dans la couche cachée du MLP

Après avoir fixé le nombre optimal de neurones à notre réseau (20 neurones), on doit fixer le pas μ qui est la valeur du taux d'apprentissage. La figure 4.10 montre l'évolution de l'erreur MSE en fonction du pas μ , pour obtenir le meilleur pas on doit chercher l'erreur minimale ($\mu = 1$).

Validation : Après avoir fixé les valeurs des paramètres du réseau MLP (nombre de neurones dans la couche cachée et le pas μ) vient l'étape de la validation où les poids w_n des neurones d'entrées et des sorties sont calculés pour chaque séquence de l'ensemble de validation depuis une initialisation, aléatoire pour arriver après, à leurs meilleures valeurs qui donnent le résultat du MSE le plus inférieur sur tout l'ensemble.

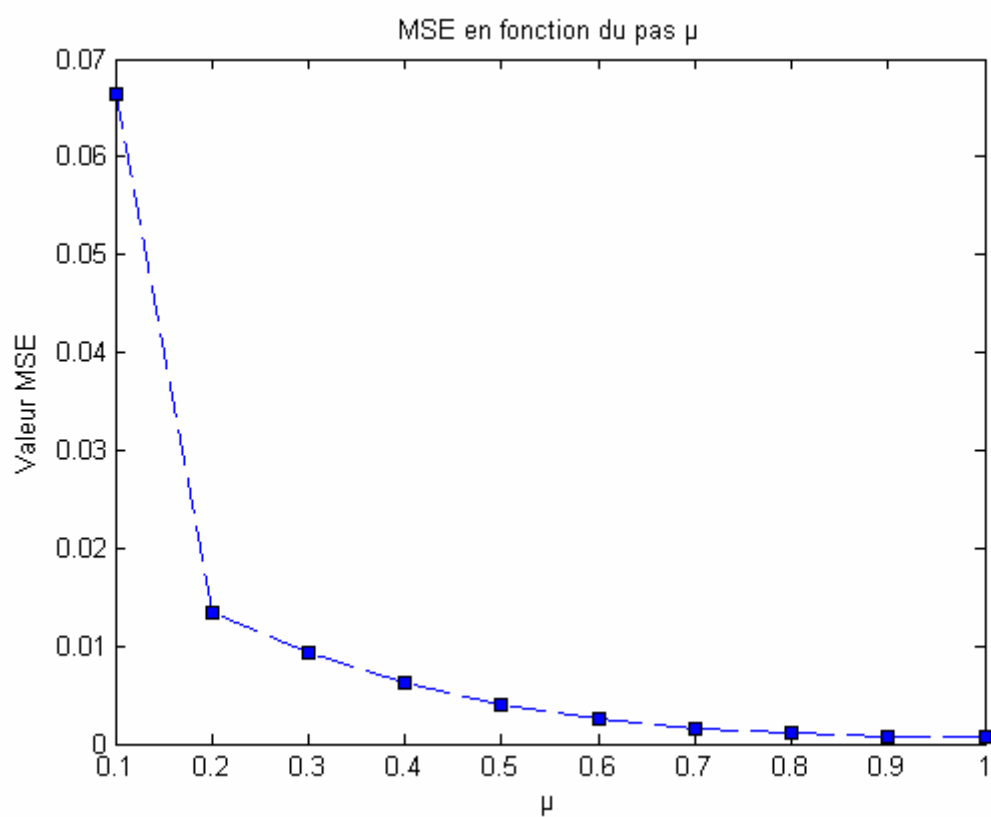


Fig.4.10 — MSE en fonction du pas d'apprentissage μ du MLP

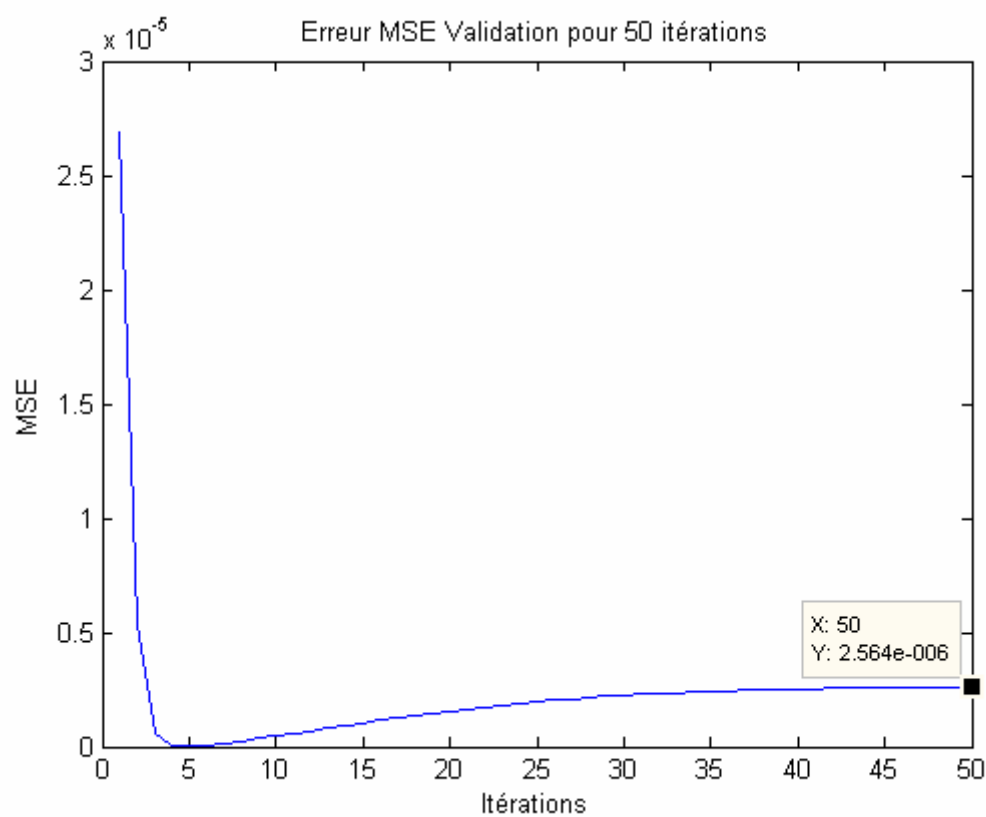


Fig.4.11 — MSE de validation du MLP en fonction de 50 itérations

Le nombre d'itérations qui augmente l'efficacité du réseau d'un côté, et qui le fait ralentir d'un autre côté est fixé aussi grâce au MSE, aucune règle ne définit comment fixer ce nombre, c'est juste les besoins ; rapidité ou exactitude de résultats. Les figures 4.11 et 4.12 montrent l'exemple à 50 et 100 itérations où on remarque que $MSE_{100 \text{ itér}} \approx \frac{1}{2} \cdot MSE_{50 \text{ itér}}$.

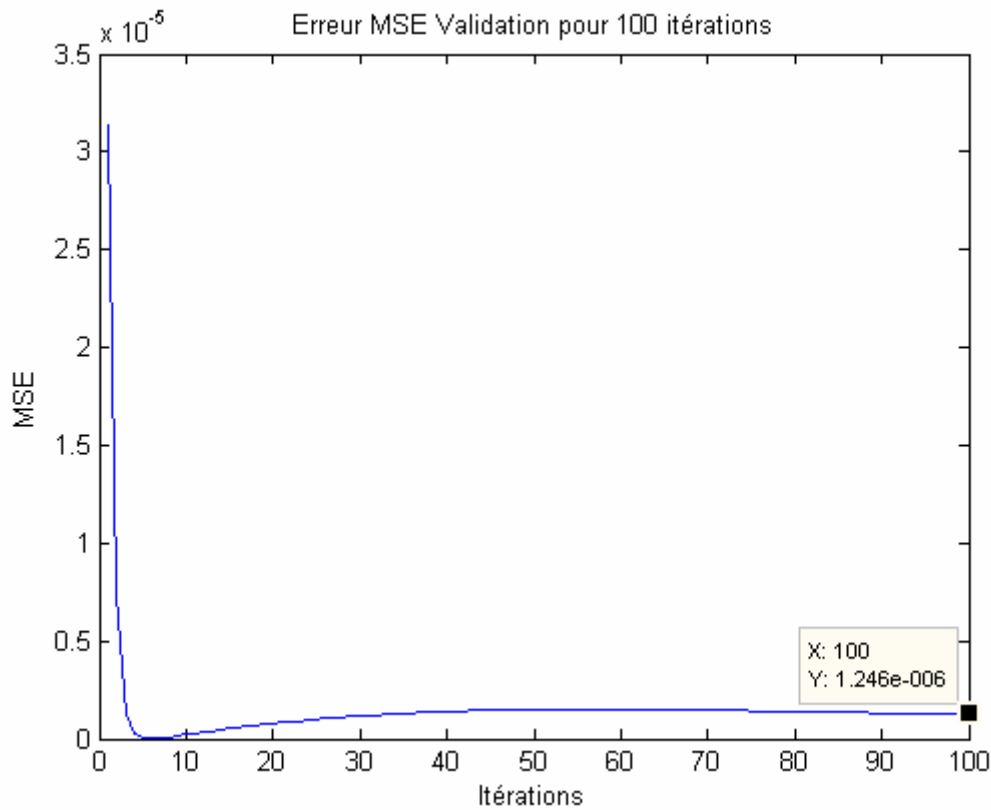


Fig.4.12 — MSE de validation du MLP en fonction de 100 itérations

Test : On peut finalement voir notre égaliseur à l'œuvre sur l'ensemble des tests. La figure 4.13 représente la poursuite des séquences estimées par l'égaliseur $y(n)$ et celles désirées $z(n)$. L'égaliseur MLP présente une très bonne poursuite des séquences avec une erreur très minime par rapport aux égaliseurs traditionnels LMS, NLMS, RLS et CMA.

La rapidité de convergence est assurée avec le bon pas μ et le nombre des itérations choisis ; l'égaliseur est stable, mais très lent.

La constellation présentée en figure 4.14 illustre le très bon fonctionnement de notre égaliseur en ce qui concerne l'erreur d'adaptation, la convergence des séquences déformées par le canal et aussi son efficacité en matière d'adaptation. On voit bien que toutes les séquences tendent vers les séquences originales. A noter que toutes les séquences sont égalisées correctement depuis le début pour arriver aux valeurs optimales sans nécessité d'initialisation ou temps d'adaptation.

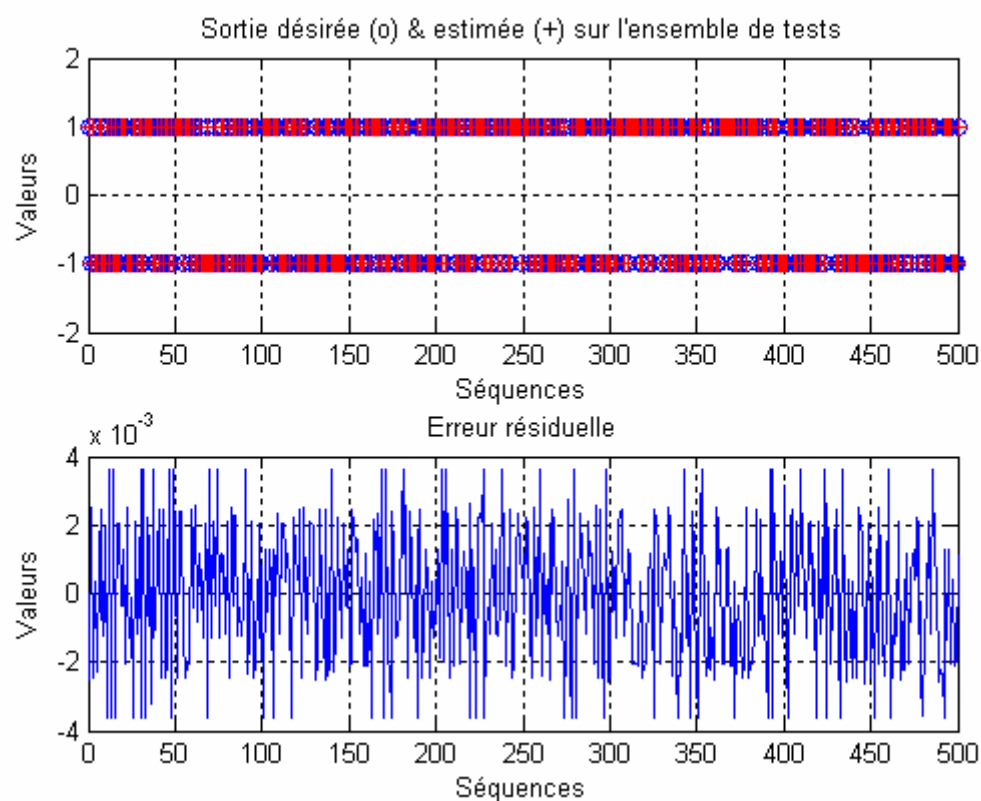


Fig.4.13 — Poursuite & erreur du MLP sur l'ensemble de tests sur 500 itérations

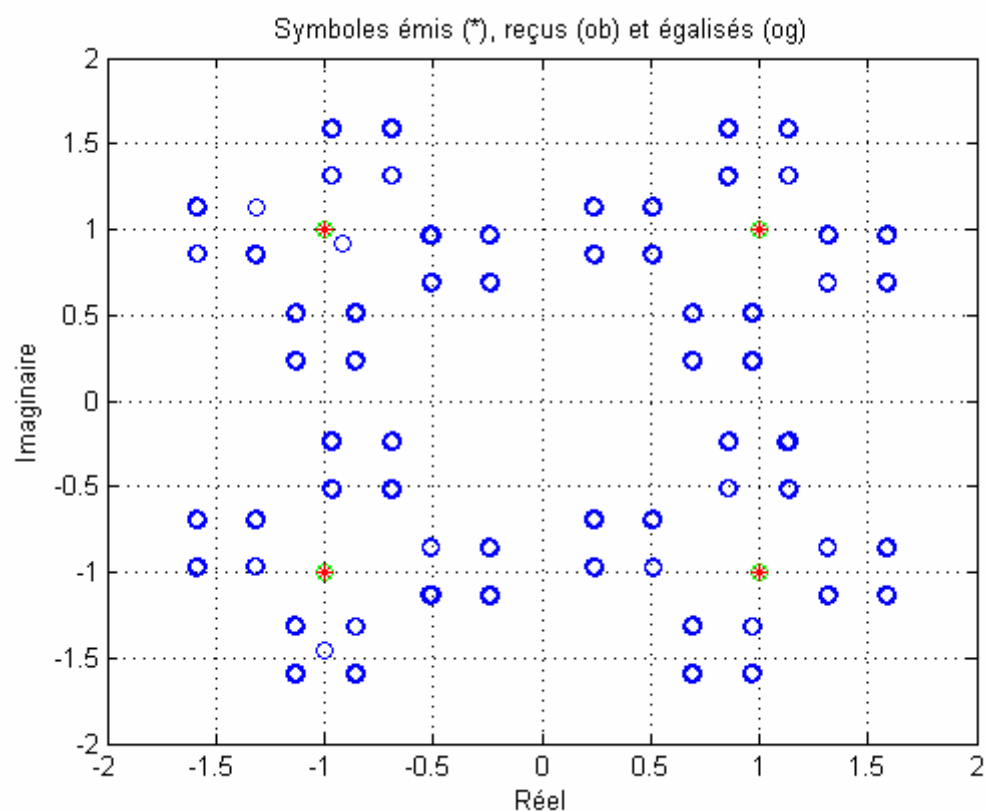


Fig.4.14 — Constellation des signaux émis, reçus (bleu) & égalisés (vert) avec MLP

Pour plusieurs rapports signal/bruit $\text{SNR}=30, 20$ et 15 dB (figure 4.15), l'erreur tend vers zéro. Tant que le bruit augmente l'erreur augmente aussi, mais l'égaliseur arrive à rassembler les séquences autour des séquences originales surtout en $\text{SNR}=15$ par rapport aux autres algorithmes. Le bruit a une grande influence sur la qualité de réception, mais en revanche l'égaliseur neuronal MLP développé montre une très grande résistance et une très bonne robustesse aux bruits externes.

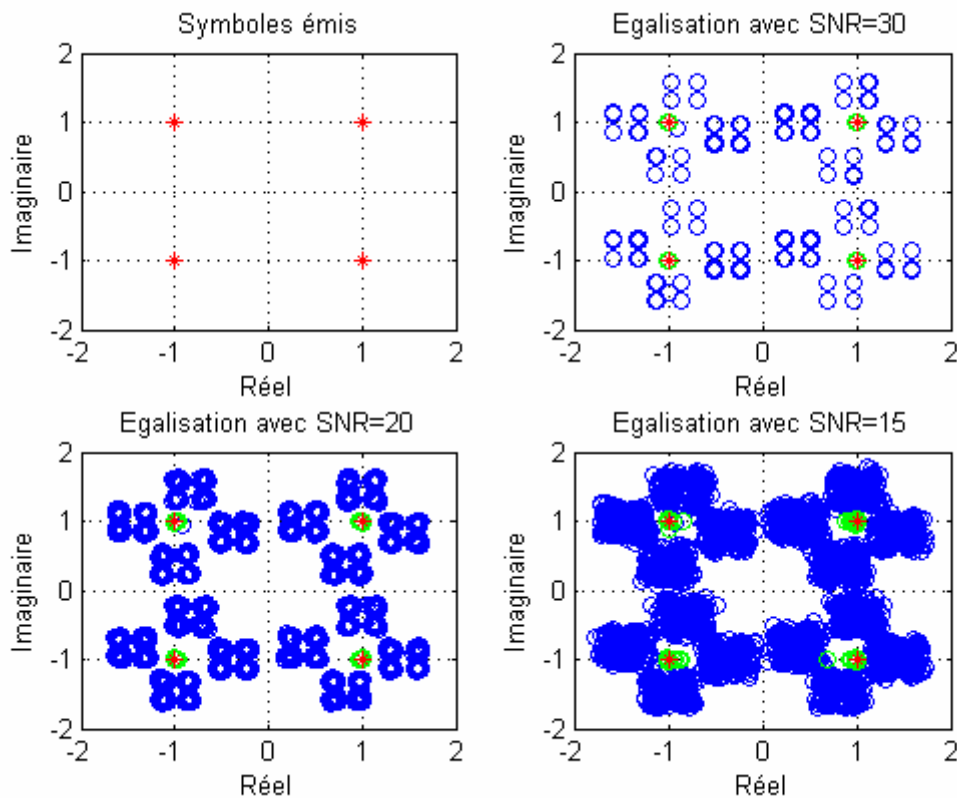


Fig.4.15 — Performances & robustesse du MLP avec des SNR=30, 20 & 15 dB

4.4 Conclusion

On a pu concevoir un égaliseur neuronal basé sur les réseaux de neurones MLP, ses performances sont meilleures que celles des égaliseurs linéaires complexes classiques supervisés en apprentissage (LMS, NLMS et RLS complexe). Sa fonction de transfert non linéaire traite mieux les non linéarités du canal qui causent les interférences inter symboles et les déformations du signal utile ; elle est aussi mieux située pour palier l'influence du bruit.

L'égaliseur MLP est robuste et performant, mais consomme des ressources importantes pour son apprentissage ; les séquences d'apprentissage ainsi un gain de temps très important pour fixer ses paramètres internes.

5. Radial Basis Function (RBF)

L'identification des paramètres du perceptron, appelés poids de connexions, est souvent réalisée par l'utilisation de l'algorithme de rétro propagation basé sur la méthode de la descente du gradient, dont l'objectif est la minimisation de l'erreur d'apprentissage. Cependant, la surface de l'erreur est souvent complexe et présente des caractéristiques peu satisfaisantes pour réaliser une descente du gradient, ce qui crée des inconvénients, tels que : la lenteur de la convergence, la sensibilité aux minima locaux et la difficulté à régler les paramètres d'apprentissage. Un autre type de réseau de neurones très utilisé, est le réseau de neurones à fonction radiale de base (RBF). Ce réseau a l'avantage d'être beaucoup plus simple que le perceptron tout en gardant la fameuse propriété d'approximation universelle de fonctions [41]. Les réseaux de neurones RBF sont très exploités dans l'optimisation et la synthèse des systèmes d'égalisation adaptative.

5.1 Types d'apprentissage

En apprentissage, il existe essentiellement deux types d'apprentissage : l'apprentissage supervisé et l'apprentissage non supervisé. Dans la majorité des réseaux neuronaux étudiés, l'apprentissage sera dit supervisé, car on impose des entrées fixes et on cherche à récupérer une sortie connue. On effectue alors la modification des poids pour retrouver cette sortie imposée. Malgré tout, il existe des réseaux à apprentissage non supervisé.

5.1.1 Apprentissage non supervisé

L'apprentissage non supervisé est le seul qui peut expliquer l'apprentissage des systèmes biologiques. Ce processus d'entraînement fait correspondre à une classe donnée de vecteurs d'entrée (qui ont une propriété commune) une sortie particulière, mais en premier temps, on ne peut pas connaître pour une classe de vecteurs d'entrée, la sortie correspondante. Nous avons plusieurs règles d'apprentissage non supervisées :

- Règle de Hebb
- Règle de Kohonen,
- Règle de Instar,
- Règle de Outstar.

5.1.2 Apprentissage supervisé

Dans ce type d'apprentissage, on cherche à imposer au réseau un fonctionnement donné, en forçant à partir des entrées présentées, les sorties du réseau à prendre des

valeurs données en modifiant les poids synaptiques. Le réseau se comporte alors comme un filtre dont les paramètres de transfert sont ajustés à partir des couples (entrée/sortie) présentés. L'adaptation des paramètres du réseau s'effectue à partir d'un algorithme d'optimisation, l'initialisation des poids synaptiques étant le plus souvent aléatoire [41].

Dans notre étude, nous avons utilisé :

- Apprentissage supervisé avec la règle du perceptron (utilisé précédemment),
- Apprentissage supervisé et non supervisé pour les réseaux à fonction radiale de base (RBF).

5.2 Apprentissage des réseaux à fonction radiale de base (RBF)

Les réseaux RBF peuvent être utilisés dans de vastes applications, puisque, principalement, ils peuvent approximer n'importe quelle fonction et leur entraînement est rapide comparé aux réseaux MLP. Cet apprentissage rapide vient en réalité du fait que les RBFs ont seulement deux couches (figure 4.16) de paramètres (centroïdes + largeurs et poids) et chaque couche peut être déterminée séquentiellement.

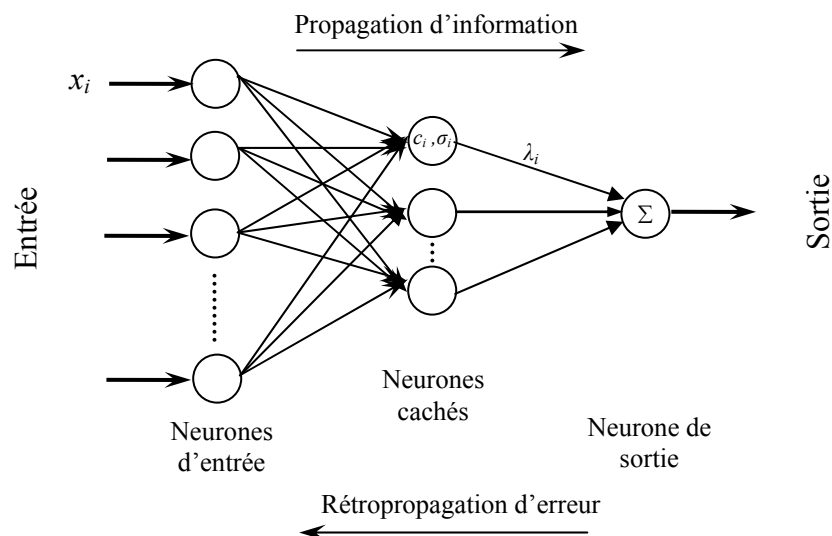


Fig.4.16 — Architecture des réseaux RBF

Le MLP est entraîné par des techniques supervisées : les poids sont calculés en minimisant une fonction coût non linéaire. Contrairement au MLP, l'entraînement du réseau RBF peut être divisé en une partie non supervisée et une autre partie linéaire supervisée. Les techniques de mise à jour des paramètres non supervisées sont relativement rapides et à propagation vers l'avant. Encore, la partie supervisée d'apprentissage consiste à résoudre

un problème linéaire, ce qui est également rapide, avec le bénéfice additionnel d'éviter le problème de minimums locaux rencontré souvent en utilisant le MLP. La procédure d'entraînement des réseaux RBF peut être décomposée en trois stages [42] :

1. Localiser les centres (C_i) des fonctions radiales (gaussiennes),
2. Déterminer leurs largeurs (σ_i),
3. Calculer les poids du réseau (λ_j) entre la couche de fonction radiale et la couche de sortie.

5.3 Equations du réseau

Le réseau RBF comporte deux couches de neurones. Les cellules de sortie effectuent une combinaison linéaire de fonctions non linéaires fournies par les neurones de la couche cachée. Les Méthodes d'approximations qui régissent le réseau sont :

$$F(x) = \sum_{j=1}^M \lambda_j \varphi_j(\|x - c_j\|) + \sum_{i=1}^d a_i x_i + b \quad (4.17)$$

où : M est le nombre des centres (C_j).

d est la dimension des variables d'entrée.

$\|(\cdot)\|$ dénote la distance euclidienne.

φ_j représente la fonction non linéaire du réseau qui est calculée par :

$$\varphi_j(\|x - c_j\|) = \exp\left(-\frac{\|x - c_j\|^2}{2\sigma_j^2}\right) \quad (4.18)$$

Les paramètres du réseau sont : C_j , σ_j et λ_j . Les centres C_j sont déterminés et adaptés respectivement en utilisant l'apprentissage non supervisé (competitive learning) avec :

$$d_2(x^i, C^j) = \sqrt{\sum_{k=1}^D (x_k^i - c_k^j)^2} = \|x^i - C^j\| = (x^i - C^j)^T (x^i - C^j) \quad (4.19)$$

$$C^k(t+1) = C^k(t) + \alpha(t)(x^i - C^k) \quad (4.20)$$

où $\alpha(t)$ est un facteur décroissant d'adaptation de temps, $0 < \alpha(t) < 1$.

σ_j est adapté grâce à la formule :

$$\sigma = \frac{d_{\max}}{\sqrt{2M}} \quad (4.21)$$

où M est le nombre de centroïdes.

$d_{\max}$ est la distance maximale entre une paire quelconque de centroïdes.

Une fois les C_j et $\sigma = \sigma_j$ sont estimés le problème devient linéaire, il faut donc utiliser la pseudo-inverse pour déterminer les λ_j [42].

5.4 Simulation & résultats

5.4.1 Principe

Le principe est le même décrit précédemment ; le même signal complexe généré avec 3000 séquences, déformé par le canal complexe est régi par l'équation :

$$C(z) = 1 + (-0.25 + 0.34j)z^{-1} + (-0.15j)z^{-2}$$

Après, on y ajoute un bruit complexe à son tour, généré avec des SNR = 30, 20 et 15 dB (figure 3.10) pour illustrer l'effet du bruit sur cet égaliseur non linéaire, et le comparer aux méthodes vues précédemment.

Ici le traitement du problème d'égalisation est le même que l'égaliseur MLP ; l'égaliseur est basé sur la RBF et l'entrée complexe du signal déformé est répartie en deux parties (Principe du split-complex) ; réelle et imaginaire, chaque partie est traitée à part et à la sortie on rend au signal sa forme complexe. Le critère d'apprentissage est l'erreur MSE. La figure 4.17 illustre le principe de la simulation.

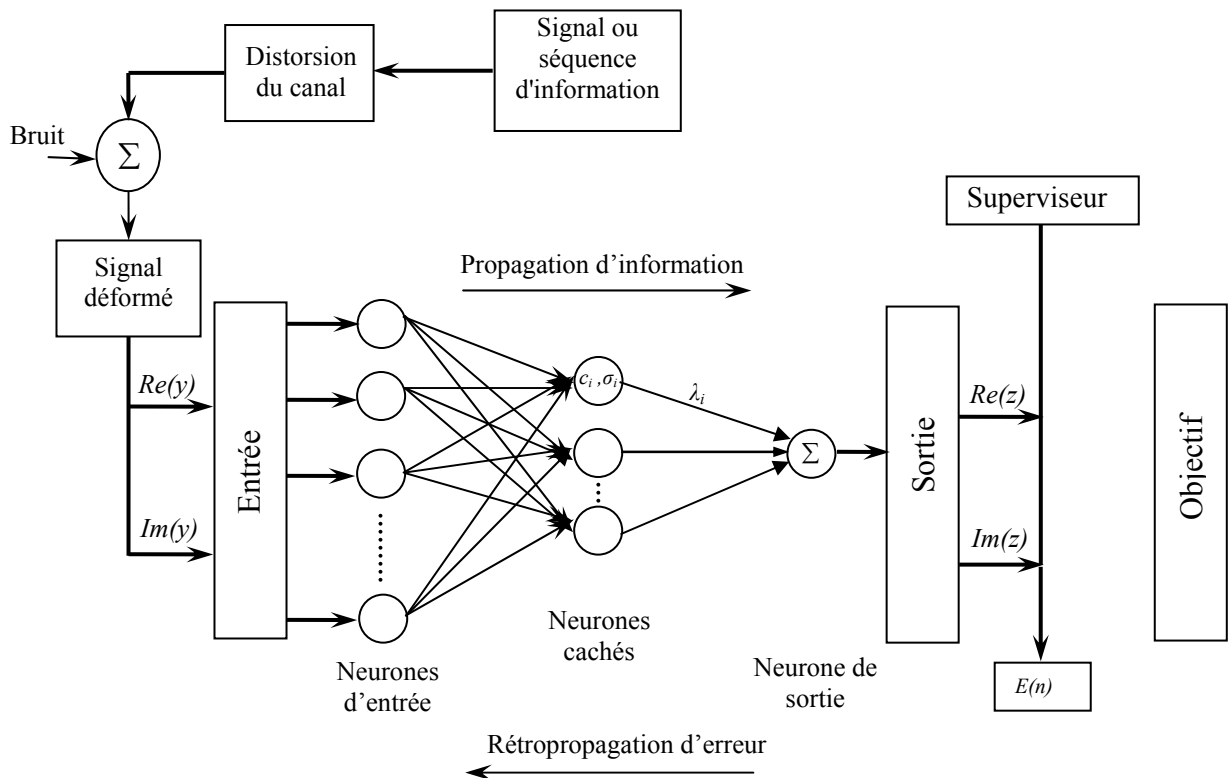


Fig.4.17 — Structure d'un égaliseur neuronal RBF

5.4.2 Etapes de mise en œuvre du réseau RBF

Le résumé des étapes et stratégies d'apprentissage sont comme suit :

- On fixe le nombre de neurones dans la couche cachée.
- On utilise une stratégie d'apprentissage traditionnelle qui est le calcul séparé basé sur les équations du réseau comme suit :
 - Les Centres des gaussiennes C_j
 - La largeur des gaussiennes σ_j
 } Apprentissage non supervisé.
- Les poids de connexions λ_j
- } Apprentissage supervisé.

Après on valide et on teste notre réseau pour les ensembles de validation et de tests.

5.4.3 Résultats et discussions

Après les étapes précédentes, le modèle développé pour le RBF est désormais prêt pour sa tâche qui est l'égalisation adaptative des séquences décrites au chapitre précédent. Les séquences utilisées durant l'apprentissage ont un SNR=30 dB, après les étapes d'apprentissage de validation et de test de l'égaliseur neuronal RBF, on y traite les séquences qui ont un SNR=20 et 15 dB (figure 3.10).

Les séquences sont réparties entre apprentissage et validation ($\frac{2}{3}$ de l'ensemble des séquences, $\frac{2}{3}$ et $\frac{1}{3}$ de cet ensemble respectivement) et test ($\frac{1}{3}$ de l'ensemble des séquences), chaque étape est évaluée avec le critère MSE.

Apprentissage : Lors de l'apprentissage, on utilise les séquences d'entrée pour l'apprentissage non supervisé (pas besoin des séquences de sortie correspondantes) des centres (C_j) en se servant des deux équations 4.19 et 4.20 (calcul de distances, et adaptation des centroïdes).

Après avoir fixé les centroïdes, on doit fixer la valeur σ qui est la largeur des gaussiennes. La figure 4.18 montre l'évolution de l'erreur MSE en fonction de la largeur σ . Pour obtenir la largeur optimale on doit chercher l'erreur minimale ($\sigma = 0.2$).

Après, en changeant le nombre de neurones graduellement par pas de deux jusqu'à 20 neurones et en calculant à chaque fois l'erreur MSE (figure 4.19), on pourra choisir le nombre de neurones qui aura l'erreur minimale. Le nombre de neurones optimal à notre réseau est 20 neurones.

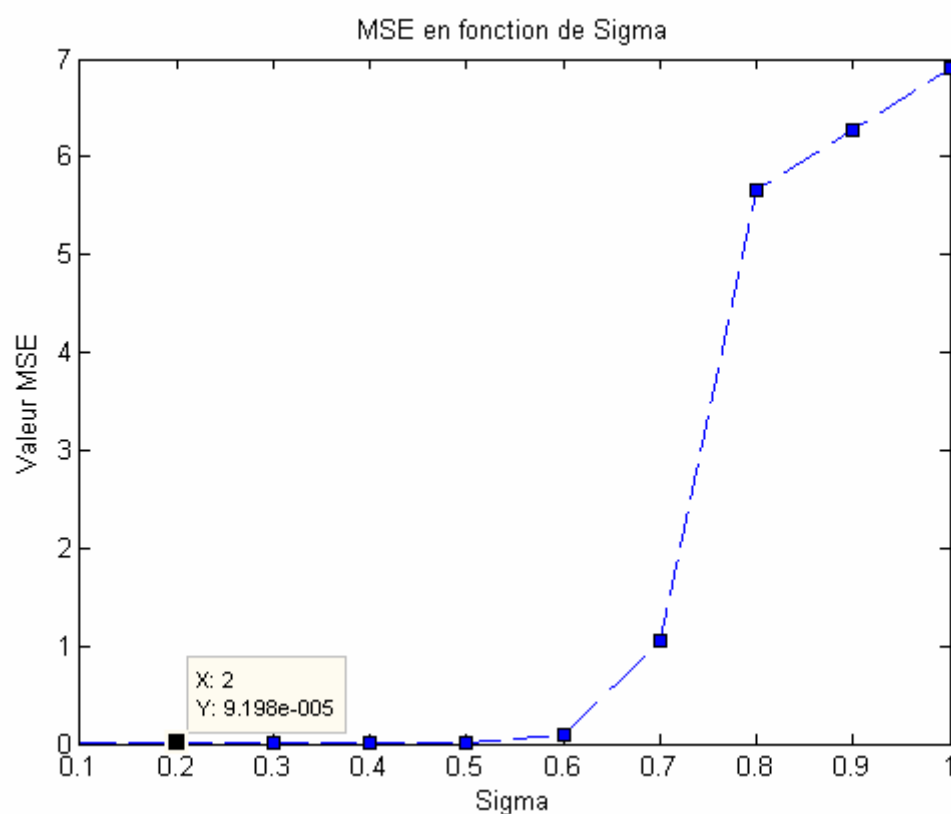


Fig.4.18 — MSE en fonction de la largeur des gaussiennes σ du RBF

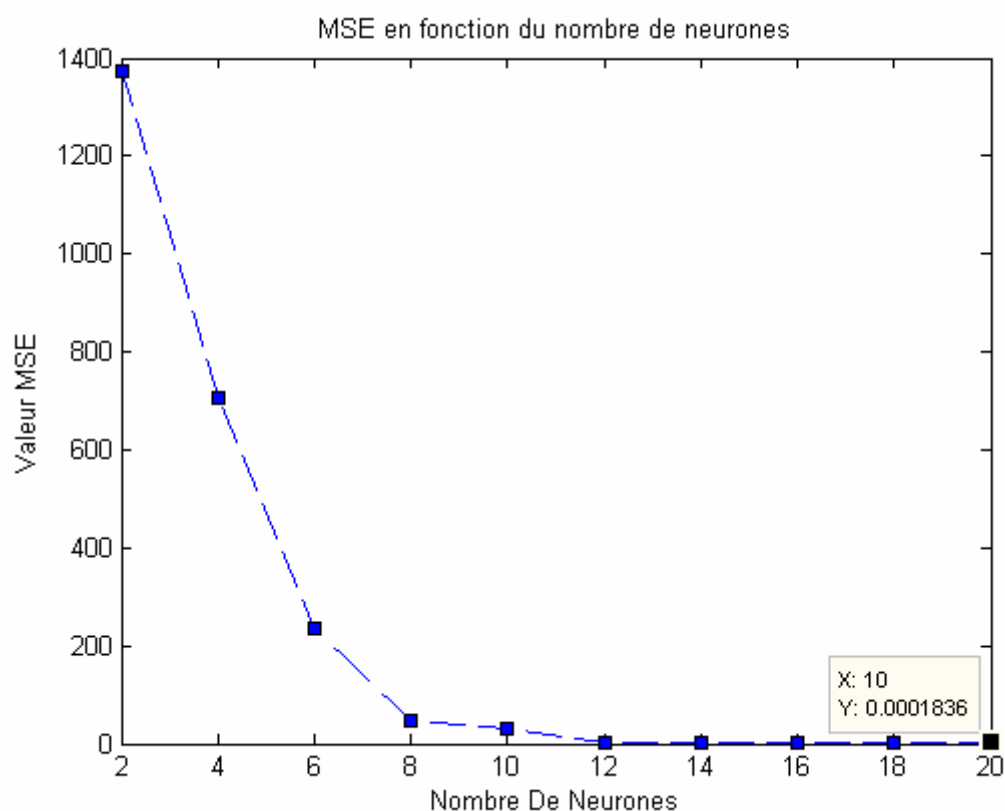


Fig.4.19 — MSE en fonction du nombre de neurones dans la couche cachée du RBF

Validation : Après avoir fixé les valeurs des paramètres du réseau RBF (Centroides, nombre de neurones dans la couche cachée et la valeur σ) vient l'étape de la validation où les poids λ_n des neurones de sortie sont calculés pour chaque séquence de l'ensemble de validation depuis une initialisation aléatoire pour arriver après à leurs meilleures valeurs qui donnent le résultat du MSE le plus inférieur sur tout l'ensemble.

Test : On peut finalement voir notre égaliseur à l'œuvre sur l'ensemble des tests, la figure 4.20 représente la poursuite des séquences estimées par l'égaliseur $y(n)$ et celles désirées $z(n)$. L'égaliseur RBF présente une très bonne poursuite des séquences avec une erreur très minime par rapport aux égaliseurs déjà vus ; traditionnels (LMS, NLMS, RLS et CMA) et neuronal (MLP). L'égaliseur est stable, mais très lent.

La constellation présentée en figure 4.21 illustre le très bon fonctionnement de notre égaliseur en ce qui concerne l'erreur d'adaptation, la convergence des séquences déformées par le canal et aussi son efficacité en matière d'adaptation. On voit bien que toutes les séquences tendent vers les séquences originales. A noter que toutes les séquences sont égalisées correctement depuis le début pour arriver aux valeurs optimales sans nécessité d'initialisation ou temps d'adaptation.

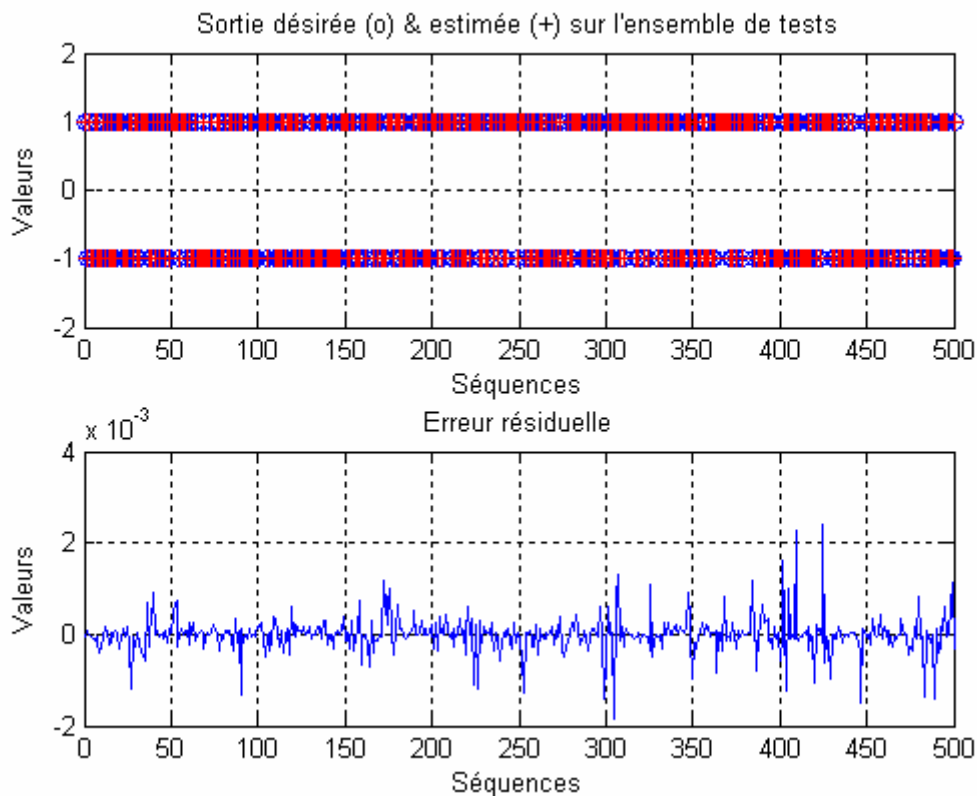


Fig.4.20 — Poursuite & erreur du RBF sur l'ensemble de tests sur 500 itérations

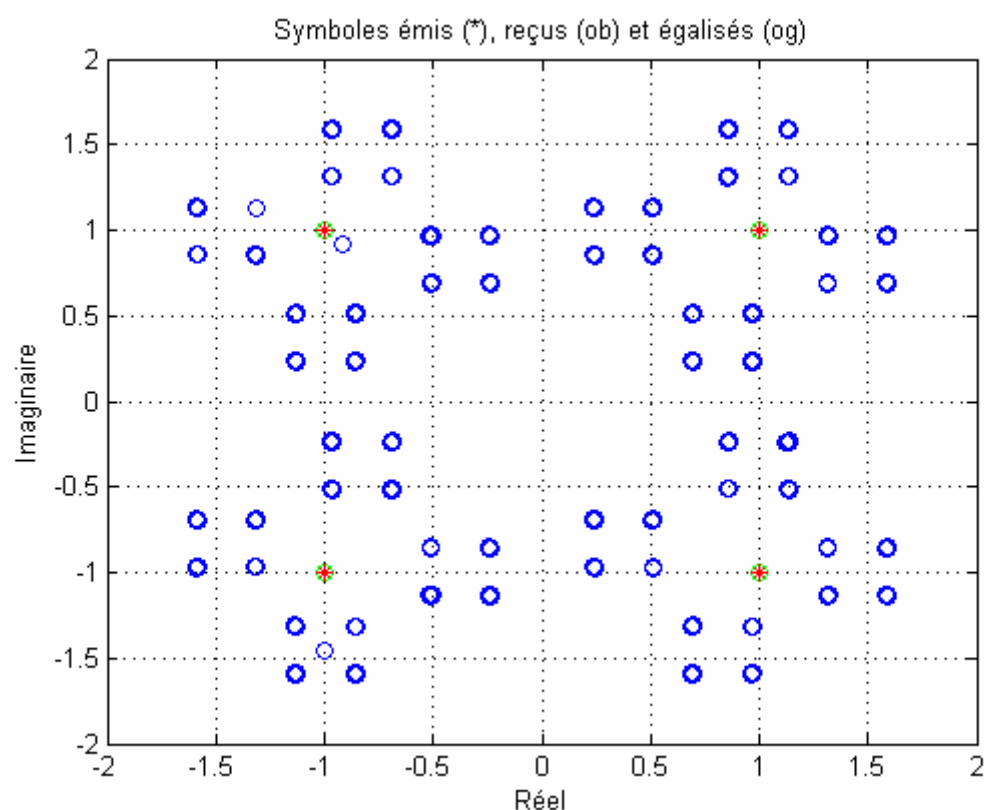


Fig.4.21 — Constellation des signaux émis, reçus (bleu) & égalisés (vert) avec le RBF

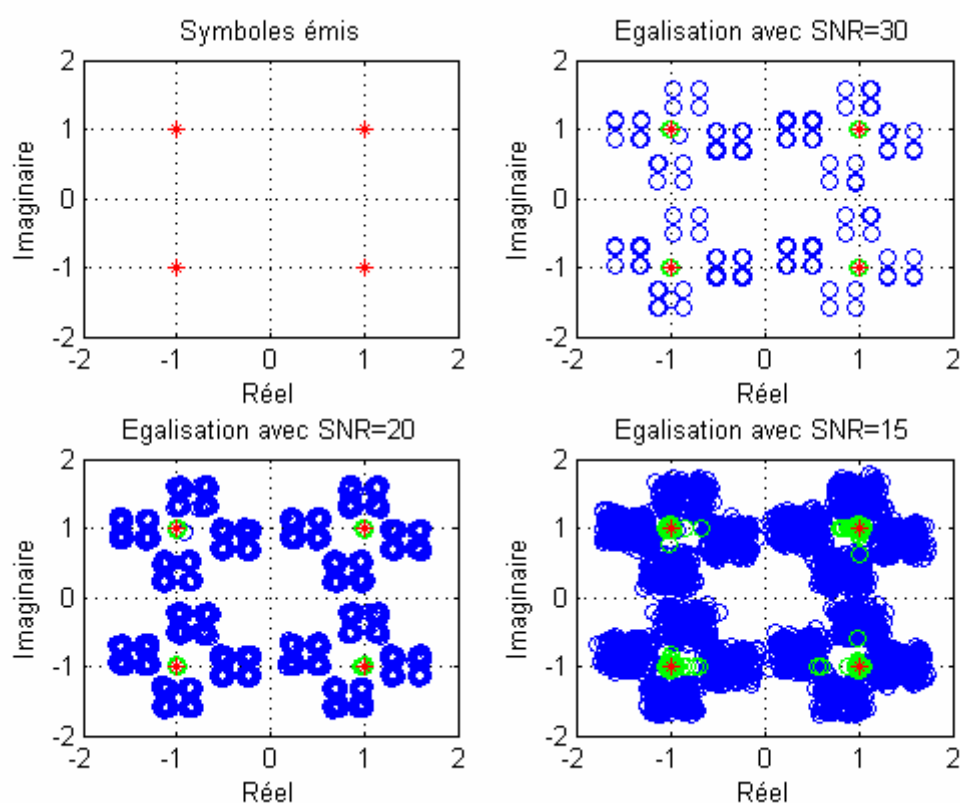


Fig.4.22 — Performances & robustesse du RBF avec des SNR=30, 20 & 15 dB

Pour plusieurs rapports signal/bruit SNR=30, 20 et 15 dB (figure 4.15), l'erreur tend vers zéro. Tant que le bruit augmente l'erreur augmente aussi, mais l'égaliseur arrive à rassembler les séquences autour des séquences originales surtout en SNR=15 par rapport aux autres algorithmes. Le bruit a une grande influence sur la qualité de réception, mais en revanche l'égaliseur neuronal RBF développé montre une très grande résistance et une très bonne robustesse aux bruits externes.

5.5 Conclusion

On a pu concevoir un égaliseur neuronal basé sur les réseaux de neurones RBF, ses performances sont meilleures que celles des égaliseurs linéaires complexes classiques supervisés en apprentissage (LMS, NLMS et RLS complexe). Sa fonction de transfert non linéaire traite mieux les non linéarités du canal qui cause les interférences inter symboles et les déformations du signal utile ; elle est aussi mieux située pour palier l'influence du bruit.

L'égaliseur RBF est robuste et performant, mais consomme des ressources importantes pour son apprentissage ; les séquences d'apprentissage offrent un gain de temps très important pour fixer ses paramètres internes par rapport aux méthodes classiques.

6 Comparaison entre les égaliseurs RBF et MLP

La figure 4.23 montre l'évolution de l'erreur résiduelle des séquences bruitées à 30, 20 et 15 dB respectueusement sur 500 séquences. À un rapport SNR=30 et 20 dB, l'égaliseur RBF montre de grandes performances par rapport à l'égaliseur MLP ; l'erreur qua-

SNR (dB)	MSE Validation (666 séquences)		MSE Test (999 séquences)		MSE Totale (2997 séquences)	
	<i>RBF</i>	<i>MLP</i>	<i>RBF</i>	<i>MLP</i>	<i>RBF</i>	<i>MLP</i>
30	$1.8420 \cdot 10^{-5}$	$2.9693 \cdot 10^{-4}$	$3.5031 \cdot 10^{-5}$	$3.2811 \cdot 10^{-4}$	$5.3686 \cdot 10^{-5}$	$4.4188 \cdot 10^{-4}$
20	0.0034	0.0048	0.0034	0.0056	0.0056	0.0105
15	0.4326	0.1534	0.1530	0.0248	0.4819	0.1777

Tableau.4.2 — MSE en Fonction du SNR pour les égaliseurs RBF et MLP sur l'ensemble des validation, test et le total

dratique minimale durant les étapes de validation et de tests de l'égaliseur RBF est bien inférieure à celle du MLP, ainsi que l'erreur totale (tableau 4.2). Par contre, pour un rapport SNR=15 dB ; l'égaliseur MLP est plus robuste et performant que l'égaliseur RBF durant toutes les étapes d'apprentissage.

Quand le canal est vraiment dispersif l'égaliseur MLP montre une grande robustesse à ce bruit additif, il est le plus convenable à ce genre de canal.

Quant à l'égaliseur RBF, il possède de grandes performances concernant les canaux légèrement et moyennement bruités.

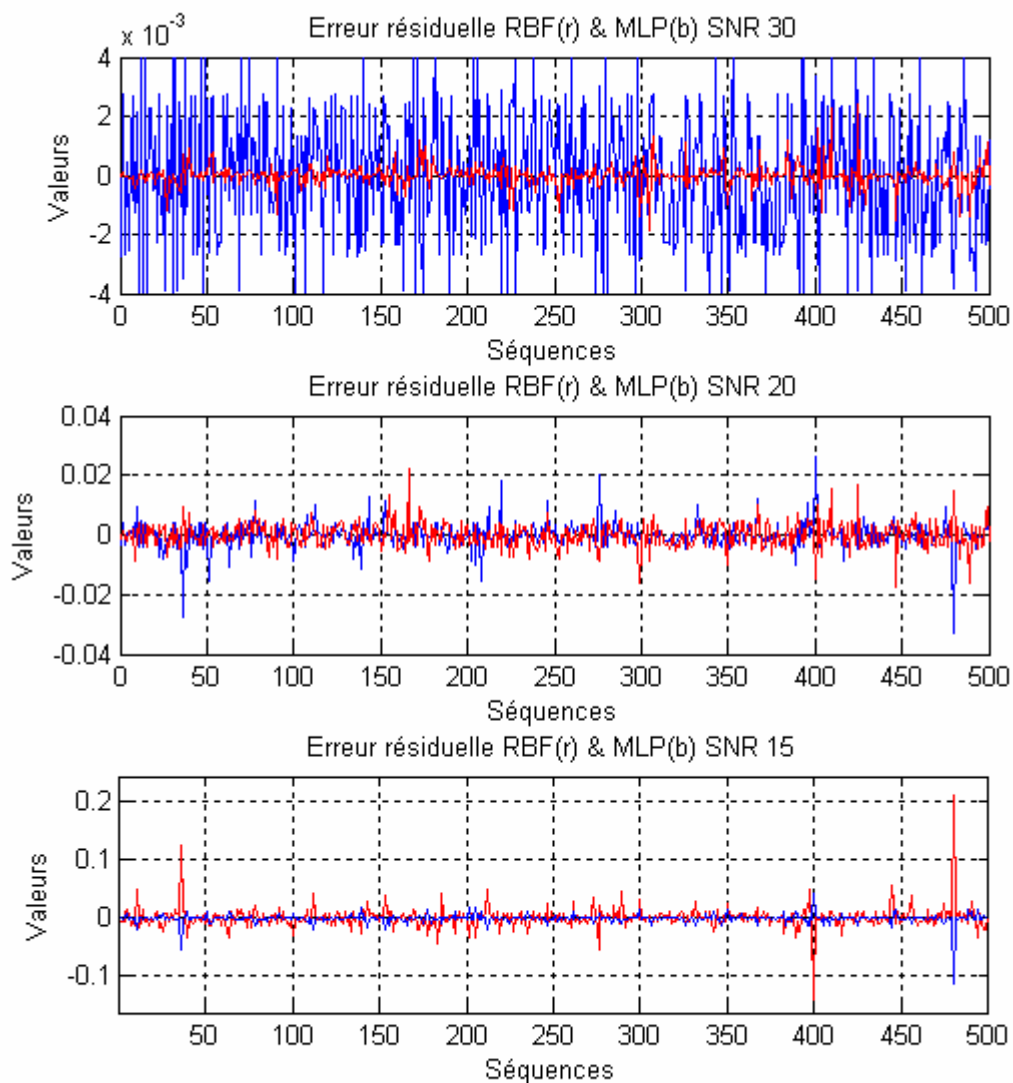


Fig.4.23 — Performances & robustesse des égaliseurs RBF (rouge) & MLP (bleu) avec des SNR=30, 20 et 15 dB sur 500 itérations

7 Conclusion

Plusieurs égaliseurs neuronaux ont été développés. La plupart d'eux sont une combinaison entre un filtre linéaire transversal conventionnel (LTF) et un réseau de neurone artificiel (RNA) qui peut être un perceptron multicouches (MLP), un réseau de fonctions radiales de base (RBF) ou autre réseau neuronal [43].

On a pu concevoir deux égaliseurs neuronaux non linéaires MLP et RBF qui ont donné de grandes performances par rapport aux égaliseurs classiques, en ce qui concerne la robustesse vis-à-vis du bruit, l'erreur du modèle, ou la reconstitution des séquences originales (élimination des facteurs ISI et MAI).

L'un des égaliseurs est plus robuste en canaux très dispersifs (MLP), l'autre est plus rapide et performant en matière d'égalisation en canaux légèrement et moyennement bruités.

Conclusion générale

Lors d'une transmission de données, le canal de transmission peut avoir plusieurs effets sur le signal transmis de l'émetteur au récepteur. Le canal est souvent symbolisé comme une source de bruit additif mais il peut aussi atténuer fortement certaines fréquences porteuses, on parle alors de fading sélectif. Le canal peut aussi avoir pour effet de "mélanger" les symboles transmis, on parle alors d'interférences entre symboles.

L'égalisation adaptative consiste à palier le mieux possible les déformations apportées par le canal de transmission au signal utile grâce à des filtres qui changent de coefficients en temps réel.

Ces filtres ont pour fonction principale d'inverser la réponse du canal de transmission de manière à ce que le couple « canal/égaliseur » puisse être considéré comme un canal idéal.

Dans cette dissertation nous avons traité deux techniques d'égalisation ; égalisation linéaire et non linéaire. Les propriétés et la convergence des filtres adaptatifs linéaires et non linéaires complexes ont été étudiés.

Les algorithmes les plus utilisés dans l'égalisation adaptative linéaire ; LMS et NLMS montrent une bonne convergence, mais sont moins complexes et moins rapides par rapport à l'algorithme RLS en ce qui concerne l'égalisation adaptative avec apprentissage. L'algorithme CMA utilisé dans l'égalisation adaptative aveugle, montre aussi une bonne stabilité mais s'avère trop lent, son avantage est qu'il n'a pas besoin de séquences d'apprentissages.

Nous nous sommes proposés de concevoir deux égaliseurs non linéaires basés sur deux modèles de réseaux de neurones MLP et RBF dans le but d'optimiser les performances du traitement, de poursuivre, et de minimiser l'erreur vis-à-vis des effets du canal et du bruit ascendant. Les deux égaliseurs neuronaux non linéaires MLP et RBF développés ont donné de grandes performances par rapport aux égaliseurs classiques, en ce qui concerne la robustesse vis-à-vis du bruit, l'erreur du modèle, ou la reconstitution des séquences originelles (élimination des facteurs ISI et MAI), mais une lenteur de traitement de données est remarquée durant le traitement. L'un des égaliseurs est plus robuste en canaux très dispersifs (MLP), l'autre est plus rapide et performant en matière d'égalisation en canaux légèrement et moyennement bruités.

Récemment, des travaux dans ce domaine expérimentent l'utilisation des algorithmes aveugles (non supervisés) pour entraîner les réseaux de neurones. En effet, l'égaliseur neuronal est encore un sujet d'étude assez vif. Comme continuité de ce travail, il serait intéressant d'améliorer et développer de nouvelles stratégies pour accélérer sa convergence et aussi pour avoir une régulation de la fonction de coût, ce qui doit réduire la complexité de calcul due au nombre de neurones.

Pour favoriser la convergence d'un égaliseur on utilise usuellement une séquence d'apprentissage c'est-à-dire, un préambule constitué de données connues du récepteur. Lorsque cela n'est pas possible, on est contraint d'utiliser des égaliseurs autodidactes (auto adaptatifs) ; donc on parle alors d'égalisation aveugle « blind equalization », qui base son traitement sur la connaissance des propriétés statistiques des signaux émis.

Cette approche fera l'objet de la suite de nos travaux, où seront également considérés les deux problèmes essentiels ; la vitesse de convergence, qui doit être la plus rapide possible, et l'obtention d'une faible erreur résiduelle en poursuite.

Références

- [1] Ajay R.Mishra, "Fundamentals Of Cellular Network Planning And Optimisation 2G/2.5G/3G... Evolution to 4G", Edition John Wiley & Sons, 2004.
- [2] Mohammed Yekhllef, "Etude des méthodes d'accès dans les réseaux mobiles", Thèse de Magister, Université de Batna, 2008.
- [3] A. Dupret, A. Fischer, "Cours de télécommunication", IUT de Villetaneuse, 2003.
- [4] Luc Deneire, "Réseaux locaux et personnels sans fil", GTR-Sophia-Antipolis, Avril 2003.
- [5] Laurent Dubreuil, "Amélioration de l'étalement de spectre par l'utilisation de codes correcteurs d'erreurs", Thèse de Doctorat, Université de LIMOGES, Octobre 2005.
- [6] Amel Aissaoui, "Synchronisation Adaptative du code PN Dans les système de communication DS/SS", Thèse de Doctorat, Université de Constantine, Juin 2008.
- [7] Mosa Ali Abu-Rgheff, "Introduction to CDMA Wireless Communication", Edition Elsevier, 2007.
- [8] 3GPP-201, "TS.25.201 UMTS; Physical layer-general description" 3GPP Technical Specification, Tech. Rep, Version 4.0.0, March 2001.
- [9] IEEE 11.B, "IEEE Standard 802.11b", IEEE Standardization, Tech. Rep., 1999.
- [10] IEEE 15.1, "IEEE Standard 802.15.1, Specification of the Bluetooth system", Version 1.2, IEEE Standardization, Tech. Rep, November 2003.
- [11] IEEE 15.4, "IEEE Standard 802.15.4", IEEE Standardization, Tech. Rep., May 2003.
- [12] Hsiao-Hwa Chen, "Next generation CDMA technologies", Edition John Wiley & Sons, 2007.
- [13] Sanchez J., Thioune M., "UMTS, Services, Architectures et WCDMA", Edition Lavoisier, 2001.
- [14] Juha Korhonen, "Introduction to 3G Mobile Communication", 2nd Edition, Edition Artech House, 2003.
- [15] E. H. Dinan and B. Jabbari, "Spreading codes for direct sequence CDMA and Wideband CDMA cellular networks", IEEE Communication Magazine, pp.48-54, September 1998.

- [16] Abdelouahab Bentercia, "Linear Interference Cancellation Structures in DS-CDMA Systems", Thèse de Doctorat, Université de Batna, 2008.
- [17] J. G. Proakis, "Digital Communication", 3rd Edition, Edition McGraw-Hill, 1995.
- [18] B. Pearson, "A Condensed Review of Spread Spectrum Techniques for ISM Band Systems", Application Note 9820, Intersil Corporation, May 2000.
- [19] Intersil Corporation, "Direct Sequence Spread Spectrum Baseband Processor", July 1999.
- [20] B. Pearson, "Complementary Code Keying Made Simple", Application Note 9850, Intersil Corporation, May 2000.
- [21] T. D. Chiueh, S. M. Li, "Trellis-Coded Complementary Code Keying for High-Rate Wireless LAN Systems", IEEE Communications Letters, Vol. 5, No. 5, pp.191-193, May 2001.
- [22] B. Ribov, G. Spasov, "Complementary Code Keying with PIC Based Microcontrollers For The Wireless Radio Communications", Proceedings of International Conference on Computer Systems and Technologies, Sofia (Bulgaria), June 2003.
- [23] L. Hanzo, L-L. Yang, E-L. Kuan and K. Yen, "Single and Multi-Carrier DS-CDMA : Multi-User Detection, Space-Time Spreading, Synchronisation, Networks and Standards", Edition Wiley-IEEE Press, 2003.
- [24] Y. Lekbir, Z. Khezzar, "Egalisation multi utilisateurs aveugle pour le Downlink DS-CDMA", Mémoire d'ingéniorat, Université de Batna, 2006.
- [25] www.mathworks.com
- [26] Paulo S.R. Diniz, "Adaptive Filtering : Algorithms and Practical Implementation", 3rd Edition, Edition Springer Science, 2008.
- [27] Hugues Benoit-Cattin, "Telecommunications Applications with the TMS320C50 DSP", INSA Lyon.
- [28] B. Mulgrew, P. Grant, J. Thompson, "Digital Signal Processing, Concepts and Applications", MacMillan Press, 1999.
- [29] S. C. Douglas, R. Losada, "Adaptive filters in Matlab : from novice to expert", Digital Signal Processing Workshop and the 2nd Signal Processing Education Workshop, IEEE Papers, 2002.
- [30] J.-F. Bercher & P. Jardin, "Introduction au filtrage adaptatif", ESIEE Paris, 2003.

- [31] Christophe Laot, "Egalisation autodidacte et turbo-égalisation : Application aux canaux sélectifs en fréquence", Thèse de Doctorat, Université de Rennes 1, 1997.
- [32] Rudolf Tanner, "Nonlinear receivers for DS-CDMA", Doctorate thesis, University of Edinburgh, September 1998.
- [33] Anna Isabel Esparcia Alcàzar, "Genetic programming for adaptive digital signal processing", Doctorate thesis, Glasgow University, 1999.
- [34] Randy .L. Haupt, Sue E. Haupt, "Practical Genetic Algorithms", 2nd Edition, Edition John Wiley & Sons, 2004.
- [35] S.N.Sivanandam, S.N.Deepa, "Introduction to Genetic Algorithms", Edition Springer, 2008.
- [36] Amar Mezache, "Optimisation de la détection décentralisée CFAR dans un Clutter Weibull utilisant les algorithmes génétiques et les réseaux de neurones flous", Thèse de Doctorat, Université de Constantine, 2009.
- [37] Marc Parizeau, "Réseaux de Neurones GIF-21140 et GIF-64326", Université de Laval, 2006.
- [38] F. Toshio and S. Takanori, "Theory and applications of neural networks for industrial control systems", IEEE Transactions on industrial electronics, Vol. 39, No. 6, December 1992.
- [39] Khireddine Melal, "Analyse des méthodes d'égalisation des techniques CDMA", Thèse de Magister, Université de Batna, Novembre 2008.
- [40] Danilo P. Mandic, Vanessa Su Lee Goh, "Complex Valued Nonlinear Adaptive Filters: Noncircularity, Widely Linear and Neural Models", Edition John Wiley & Sons, 2009.
- [41] Warda Barkat, "Etude et Synthèse des Caractéristiques de Réseaux d'antennes Imprimées Supraconductrices dans la Bande Millimétrique", Thèse de Doctorat, Université de Constantine, Septembre 2009.
- [42] Nabil Benoudjit, "Variable selection and neural networks for high-dimensional data analysis : application in infrared spectroscopy and chemometrics", Thèse de Doctorat, Université catholique de Louvain, Septembre 2003.
- [43] Corina Botoca, Georgeta Budura, "Symbol Decision Equalizer using a Radial Basis Functions Neural Network", Proceedings of the 7th WSEAS International Conference on Neural Networks, Cavtat, Croatia, pp79-84, June 12-14, 2006.