



Codage Conjoint Source-Canal des Sources Vidéo

Georgia Feideropoulou

► To cite this version:

Georgia Feideropoulou. Codage Conjoint Source-Canal des Sources Vidéo. Traitement du signal et de l'image [eess.SP]. Télécom ParisTech, 2005. Français. NNT: . pastel-00001294

HAL Id: pastel-00001294

<https://pastel.hal.science/pastel-00001294>

Submitted on 22 Nov 2010

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

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



Thèse

présentée pour obtenir le grade de docteur

de l'Ecole Nationale Supérieure
des Télécommunications

Spécialité : Signal et Images

Codage Conjoint Source-Canal des Sources Vidéo

Georgia Feideropoulou

Soutenue le 4 Avril 2005 devant le jury composé de

Michel Barlaud
Emanuele Viterbo

Rapporteurs

Pierre Duhamel
Catherine Lamy-Bergot

Examineurs

Béatrice Pesquet-Popescu
Jean-Claude Belfiore

Directeurs de thèse

Abstract

Due to the wide application of novel wireless technologies, the transmission of video sequences over wireless channels is one of the topics of recent research. Efficient compression and low-error transmission are among the most important goals of a coding scheme. The separated optimisation of the source and channel coder due to its high degree of complexity remains, however, an impractical idea for real-time applications. The research for an alternative solution leads to the joint optimisation of the source and channel coder.

We propose a joint source-channel coding scheme, developed for video sequences, which consists of a vector quantization based on lattice constellations and a linear labelling minimizing simultaneously the source and the channel distortion. The linear labelling has already been proved to minimize the channel distortion on binary symmetric channels and the linear transforms based on lattice constellations of maximum diversity to minimize, at the same time, the distortion of Gaussian sources.

We study the dependencies between the wavelets coefficients of a $t + 2D$ video decomposition in order to efficiently exploit the linear transforms developed for Gaussian sources. As the source distribution of the subbands is not Gaussian we present the necessary modifications in order to obtain a robust coding scheme.

We propose a stochastic model to capture the dependencies between the wavelets coefficients and we use it to build an optimal mean square predictor for missing coefficients. We present two applications of this predictor on the transmission over packet networks : a quality enhancement technique for resolution scalable video bitstreams and an error concealment method.

In the case of transmission of video sequences over noisy channels, we scale appropriately the lattice-based quantizer to the video source dynamics and we use our stochastic model to reduce the side information sent on the channel.

We develop a robust joint source-channel coding scheme for transmission of video sequences over a Gaussian channel using uncoded and coded index assignment via Reed-Muller codes. We investigate the conditions requiring the use of a coded index assignment and we prove its superiority compared to an unstructured vector quantizer.

For a transmission over a flat-Rayleigh channel, we develop a coding scheme using our vector quantizer followed by a rotation matrix. We exploit the benefit of the increased diversity offered by the rotated constellations without adding any redundancy by error-correcting codes.

Résumé

L'objet de cette thèse est de proposer un codage conjoint source-canal de séquences vidéo pour la transmission sur des canaux sans fil.

D'après le théorème de séparation de Shannon, l'optimisation d'un système de transmission passe par l'optimisation séparée du codeur source et du codeur canal. Cependant, cette optimisation n'est valable que pour des blocs de données de taille infiniment longue, ce qui se traduit en une complexité élevée des deux codeurs, prohibitive pour des systèmes temps réel. La solution la plus récente dans ce domaine est d'optimiser des combinaisons de blocs de la chaîne de transmission afin de diminuer la complexité sans sacrifier les performances. Le but dans un tel schéma de codage conjoint source-canal est de minimiser la distorsion globale du système.

Le système de codage conjoint source-canal proposé dans cette thèse est fondée sur un quantificateur vectoriel structuré et une assignation linéaire d'étiquette qui minimisent simultanément la distorsion canal et la distorsion source. Le quantificateur vectoriel qui est construit à partir de constellations provenant de réseaux de points, lesquels satisfont la propriété de diversité maximale, minimise la distorsion source d'une source gaussienne. La distorsion canal est également minimisée par l'étiquetage linéaire.

Nous avons étudié les dépendances entre les coefficients d'ondelettes provenant d'une décomposition $t + 2D$, avec ou sans estimation de mouvement afin d'étendre le schéma du codage conjoint source-canal, développé pour les sources gaussiennes, dans le domaine vidéo où la distribution des coefficients est loin d'être gaussienne.

Nous proposons un modèle doublement stochastique afin de capturer ces dépendances et ensuite nous proposons également des estimateurs robustes des paramètres de ce modèle. Nous appliquons ce modèle à la protection des erreurs pour prédire les coefficients perdus et améliorer ainsi la qualité de vidéo.

Nous présentons les modifications nécessaires afin d'avoir un système efficace et robuste en présence de bruit. Dans le cas d'un canal gaussien, nous développons deux systèmes, un avec un étiquetage linéaire non codé et l'autre avec un étiquetage linéaire codé utilisant des codes de Reed-Muller. Nous comparons ces deux schémas de codage avec un schéma non-structuré dont l'étiquetage est adapté au canal et avec un coder vidéo scalable.

Dans le cas d'un canal de Rayleigh non-sélectif à évanouissements indépendants le schéma devient encore plus robuste lorsque nous utilisons une matrice de rotation avant la transmission sur le canal. Cette matrice de rotation augmente la diversité du système et conduit à des performances remarquables.

Remerciements

je tiens à remercier tout d'abord M. Barlaud et E. Viterbo d'avoir accepté d'être rapporteurs, ainsi que P. Duhamel et C. Lamy-Bergot pour avoir gentilement accepté d'être examinateurs de mon jury de thèse.

Je remercie vivement ma professeur Beatrice Pesquet-Popescu pour son soutien pendant ces trois années de thèse, pour ses conseils, et pour son caractère optimiste qui m'a donné la force de continuer même lorsque le chemin semblait sans issue. Merci Beatrice : tu étais toujours présente à mes côtes. Sans toi cette thèse n'aurait jamais vu le jour.

Je tiens à remercier tout autant mon professeur J.C. Belfiore : ses conseils scientifiques avisés et sa bonne humeur permanente sont pour beaucoup dans l'accomplissement de cette thèse.

Je remercie aussi le professeur J. Fowler pour ses conseils scientifiques qui m'a donné durant son passage à l'E.N.S.T.

Les personnes que l'on rencontre durant une thèse sont beaucoup plus importantes que le travail lui-même. Je pourrais écrire une page pour chacune mais au final mes remerciements seraient plus longs que ma thèse!!!

Et comme ma thèse ne se réduisait pas à Télécom, je tiens à remercier Olivier, Marceau, Georges, Paul, Joan, Filip, Gwladys, Katerina, Rena pour toutes les soirées détendues que j'ai pu partager avec eux.

Merci à la famille grecque de Télécom...à Manolis (vrai maître de philosophie), à Chris-le Grec (j'avoue...que les premiers jours étaient difficiles!!), à Vassilis (ça va, ça vient!!). Merci Yianni (à l'Olympe le Samedi???) pour tous les discussions sérieuses!... sur nos problèmes de vie;)).

Merci à Theo (liberté...à la cuisine roumaine;)) pour sa patience les jours difficiles;), à Greg (efkaristo Boly) optimisateur incontournable sur pogany!!!..., à Maria (Ange!!!) pour l'ambiance joyeuse même quand les simules ne marchaient pas..., à Lionel un vrai Breton-avec un déficit de mémoire;)...mais bon...personne n'est parfait...

Merci à Slim pour ses encouragements les jours où ça n'allait pas, à Cleo, à Mathieu (Bonjour Mathieu), à Thomas, à toute l'équipe d'audio dont je n'ai pas vraiment profité...ce que je regrette :(mais la fin était dure;).

L'équipe de la "folie" avec tous ses ti-punch, avec son uno, ses discussions, son merle resteront inoubliables... Merci Mohamed -le chef de A303- pour les pauses pleines de musique libanaise, Alireza pour tous les poèmes sur les reines!! Merci Francois pour l'ambiance de fwansswa que tu as apporté au labo avec tes auteurs et ton canard! Merci Gossass...pour les nuits blanches, pour les tours sur les Champs en dansant le M.I.A. Merci à Bruno qui était toujours partant pour un petit pot...qui se terminait en une longue soirée...

Merci Mariam, tu étais une colloque sans partager le même appart, merci pour toutes les "décisions" importantes et surtout rebelles que nous prenions autour d'un thé et d'un

kir ;) et qui ne dureraient qu'une soirée... ;)

Et vous deux, Sab et Seb je ne sais pas comment vous remercier...je vous remercie seulement pour avoir été tout le temps présents...pour m'avoir supportée pendant les hauts et les bas (et il y en avait beaucoup!!) et pour m'avoir montré que l'on peut se faire de vrais amis meme quand on ne parle pas la même langue... je n'oublierai jamais...nous trois devant du vin et des assiettes de fromage...au merle en demandant "noir desir" et "louise attaque"...et vos vrais efforts pour m'intégrer dans votre vie...Euxaristw poly...

Il reste quatre personnes qui étaient en Grèce...mais heureusement ils ne m'ont jamais laissée tomber même si j'étais loin...Merci à Soula (sous la mer) à Dimitris (takis), à Veia et à Marina (marinero) qui m'ont prouvé que la distance n'est pas un obstacle devant les vraies amitiés.

Un grand merci à mes parents qui ont supporté brièvement mon absence et qui ne se sont jamais opposés à mes désirs...

Table des matières

Résumé de la thèse.	1
1 Statistical models for wavelet coefficients	1
1.1 Introduction	1
1.2 Overview of statistical wavelet models for still images	2
1.2.1 Interscale Models	2
1.2.2 Intrасcale Models	3
1.2.3 Composite Models	3
1.3 Simoncelli's Joint Statistical Model	3
1.4 Introduction to the $t + 2D$ scheme of decompositon of a video sequence . .	5
1.5 Overview of statistical wavelet models for video	7
1.5.1 Marginal distribution of the spatio-temporal wavelets coefficients . .	7
1.5.2 Extension of models for still images to the video domain	9
1.6 Conclusion	11
2 Spatio-temporal modelling of the wavelet coefficients in a $t + 2D$ scheme	13
2.1 Introduction	13
2.2 Conditional histograms of wavelet coefficients in a $t + 2D$ scheme	13
2.3 Double Stochastic Model	14
2.4 Model Estimation	19
2.4.1 <i>Least squares (LS)</i>	20
2.4.2 <i>Maximum-likelihood (ML)</i>	21
2.4.3 <i>A more Efficient Criterion (EC)</i>	22
2.5 Illustration Examples	23
2.6 Conclusion	25
3 Application of the statistical model to error concealment and quality enhancement of video	27
3.1 Introduction	27
3.2 Prediction Method	28
3.3 Model-Based Quality Enhancement of Scalable Video	29
3.4 Error concealment in the Spatio-temporal wavelet domain	30
3.5 Error concealment of scalable bitstreams	35
3.6 Conclusion	37

4	Overview of Joint Source-Channel coding schemes	39
4.1	Introduction	39
4.2	Problem Statement	40
4.2.1	Vector Quantization	41
4.2.2	Index assignment (IA)	42
4.2.3	Channel coding	47
	Reed-Muller codes	47
	Rate-Punctured Convolutional Codes	47
4.2.4	Channel decoding	48
	Decoding Rate Punctured Convolutional codes with the Viterbi algorithm .	48
	Decoding Reed-Muller codes	48
4.3	Source-optimized channel coding	50
4.4	Channel-optimized source coding	55
4.5	Other combined optimizations	57
4.6	Conclusion	58
5	Joint Source-Channel coding based on linear labelling and rotated constel- lations	59
5.1	Introduction	59
5.2	Linear Labelling and Joint Source-Channel coding	60
5.2.1	Linear labelling to minimize channel distortion	60
5.2.2	Minimisation of source distortion-case of Gaussian sources	62
5.3	Maximum component diversity constellation	63
5.4	Construction of rotated $\mathbb{Z}^n$ -lattices of dimension $n = (p-1)/2$ and mixture constructions	64
5.4.1	Mixture constructions of rotated $\mathbb{Z}^n$ -lattices	65
5.5	Rotated $\mathbb{Z}^n$ lattices from cyclotomic fields where n is a power of 2	66
5.6	Sphere Decoder	67
5.7	Performances of rotated BPSK over a Rayleigh fading channel	69
5.8	Conclusion	70
6	Joint source-channel coding of a video on a Gaussian channel	71
6.1	Introduction	71
6.2	Bit allocation algorithm	72
6.2.1	A general method of bit allocation for spatio-temporal subbands . .	72
6.2.2	A robust bit allocation algorithm	74
6.3	Coding algorithm	75
6.3.1	Source codebook construction	76
6.3.2	Scaling the lattice constellation to the source dynamics	78
6.4	Application of joint source-channel coding scheme of a video over a Gaus- sian channel	81
6.4.1	A First Attempt	81
6.4.2	Application of the bit allocation algorithm	84
6.5	Calculation of the end-to-end distortion in the noisy system	88
6.6	Vector quantization by linear mapping of a block code	89

6.6.1	Index assignment using $RM(r, m)$	90
6.6.2	Overall results when $RM(r, m)$ is used for index assignment	94
6.7	Structured vs. unstructured codebook with index assignment	95
6.8	Comparison with a MC-EZBC protected with rate punctured convolutional codes	100
6.9	Conclusion	103
7	Joint source-channel coding of a video on a flat-Rayleigh fading channel	105
7.1	Introduction	105
7.2	Linear Labelling and Puntured Convolutional Codes on a flat-Rayleigh fading Channel	106
7.3	Joint source-channel coding using rotations prior to transmission over a flat-Rayleigh channel	108
7.4	Comparison with the MC-EZBC protected by rate punctured convolutional codes over a flat-Rayleigh channel	111
7.5	Conclusion	113
	Conclusions and Perspectives.	119
A	Basic definitions in Algebraic Number theory	123
A.1	General definitions	123
A.1.1	Properties of Ideals	126
A.2	Ideal Lattices	126
A.2.1	Construction of rotated $-\mathbb{Z}^n$ lattices with full diversity by Ideal lattices	127
A.2.2	Cyclotomic fields	128
B	Vectorial extension of the double stochastic model	129
C	Rotation Matrices	131
	Bibliography.	132

Table des figures

1	Histogrammes conditionnels des coefficients du deuxième niveau temporel, du premier niveau spatial et d'orientation verticale. A gauche : sans estimation de mouvement et à droite avec d'estimation/compensation de mouvement.	5
2	Description du système	8
3	Trames reconstruites de la séquence "hall-monitor" à 566 Kbs. Première ligne : canal gaussien avec SNR=4.33 dB. Deuxième ligne : canal gaussien avec SNR=6.75 dB. Première colonne : schéma de codage JSC. Deuxième colonne : schéma de codage JSC+RM. Troisième colonne : schéma de codage MCA.	16
1.1	Conditional histogram for a fine scale horizontal coefficient. Conditioned on the parent (same location and orientation, coarser scale).	4
1.2	Example of the log-domain conditional histogram of the fine scale coefficient conditioned on its parent (same location and orientation, coarser scale)	4
1.3	Example of the log-domain conditional histogram of the fine scale coefficient conditioned on a linear combination of neighboring coefficient magnitudes.	5
1.4	Spatio-temporal neighbors of a wavelet coefficient in a video sequence (the original Group of Frames is decomposed over four temporal levels). Dependencies are highlighted with its spatial (parent, cousins, aunts) and spatio-temporal neighbors (temporal parent and temporal aunts). Temp i stands for the i -th temporal decomposition level.	6
1.5	Approximation by a Gaussian mixture of the vertical subband in the first spatial and temporal resolution level.	10
1.6	Approximation by a Generalized Gaussian of the vertical subband in the first spatial and temporal resolution level.	11
2.1	With motion estimation : Conditional histograms of coefficients in the vertical subband at the first temporal and first spatial resolution level. 2.1(a) : Conditioned on the spatial parent. 2.1(b) : Conditioned on the spatial neighbors. 2.1(c) : Conditioned on the spatio-temporal parent. 2.1(d) : conditioned only on the spatio-temporal neighbors. 2.1(e) : Conditioned on the spatial and spatio-temporal neighbors.	15

2.2	Without motion estimation : Conditional histograms of coefficients in the vertical subband at the first temporal and first spatial resolution level. 2.2(a) : Conditioned on the spatial parent. 2.2(b) : Conditioned on the spatial neighbors. 2.2(c) : Conditioned on the spatio-temporal parent. 2.2(d) : conditioned only on the spatio-temporal neighbors. 2.2(e) : Conditioned on the spatial and spatio-temporal neighbors.	16
2.3	With motion estimation : Conditional histograms of coefficients in the vertical subband at the first spatial resolution level on both spatial and spatio-temporal neighbors. 2.3(a) : First temporal resolution level. 2.3(b) : Second temporal resolution level. 2.3(c) : Third temporal resolution level.	17
2.4	Without motion estimation : Conditional histograms of coefficients in the vertical subband at the first spatial resolution level on both spatial and spatio-temporal neighbors. 2.4(a) : First temporal resolution level. 2.4(b) : Second temporal resolution level. 2.4(c) : Third temporal resolution level.	18
2.5	Left : vertical detail subband at the highest spatial resolution of the first temporal decomposition level for "hall_monitor" sequence. Right : simulated subband, using the conditional law given in Eq. (2.6).	25
3.1	MSE of the spatial reconstruction of the detail frames at each temporal resolution level in a GOF.	30
3.2	PSNR improvement for a GOF of 16 frames of the "foreman" and "hall-monitor" CIF sequences, when we predict the <i>finest frequency</i> subbands at different temporal resolution levels.	31
3.3	Zoom in a temporal detail frame at the first temporal resolution level. Left : original frame. Center : reconstructed detail frame when we predict its <i>finest resolution</i> subbands. Right : reconstructed frame when we set them to zero.	32
3.4	MSE of the spatial reconstruction of the first temporal detail frame when losing the horizontal subband at different temporal resolution levels and for three spatial resolution levels. SP i stands for the i -th spatial resolution level and TEMP i for the i -th temporal decomposition level.	32
3.5	First temporal detail frame at the first temporal resolution level. Left : original frame. Center : reconstructed detail frame when we predict the lost horizontal subband of the <i>second spatial resolution level</i> . Right : reconstructed detail frame when we set it to zero.	33
3.6	PSNR improvement (prediction <i>vs.</i> setting to zero) of a reconstructed GOF of the original sequence "foreman" in CIF format, 30 fps, when we loose the horizontal subbands of the <i>second spatial resolution level</i> at each temporal resolution level.	34
3.7	Temporal detail frame at the first temporal resolution level. Left : original frame. Center : reconstructed frame obtained by predicting the (lost) <i>second spatial resolution level</i> . Right : reconstructed frame when we set to zero the details corresponding to the lost packet.	35

3.8	Improvement of the PSNR of the reconstructed GOF of the original sequence "foreman" when we loose the horizontal subbands of the second spatial resolution level at each temporal resolution level and we predict them and the finest spatial resolution subbands.	36
3.9	First temporal detail frame at the first temporal resolution level. Left : original subband. Center : reconstructed frame when we predict the second and then the first spatial resolution levels. Right : reconstructed frame when we set lost subbands to zero.	37
4.1	General block diagram of the transmission system.	40
5.1	System description	62
5.2	A simple transmission scheme involving rotated constellations.	69
5.3	BER of rotated multidimensional Z^n lattice constellations constructed as in [4] for different SNR over a Rayleigh fading channel.	69
5.4	BER of rotated multidimensional Z^n lattice constellations constructed as in [34] for different SNR over a Rayleigh fading channel.	70
6.1	System description.	75
6.2	Histograms of real subbands vs GGD with the same parameters. Up : vertical subband at second spatial/ second temporal resolution level of a detail frame. Down : vertical subband at first spatial resolution level of the approximation frame.	79
6.3	PSNR of the frames of a GOF on a Gaussian channel of $SNR = 4.33$ dB with four different types of coding.	83
6.4	Reconstructed Frames. First line : reconstructed frame in a noiseless environment. Second line, left : reconstructed frame when no protection by a code is used over a Gaussian channel with $SNR=4.33$ dB. Second line, right : the reconstructed frame when the coarsest spatial resolution level of the temporal approximation frame is protected (Gaussian channel with $SNR=4.33$ dB).	84
6.5	Reconstructed Frames of "foreman" at 1104 kbs with motion estimation and "hall-monitor" at 566 kbs without motion estimation. First line : reconstructed frames in a noiseless environment. Second line : reconstructed frames over a Gaussian channel with $SNR=6.75$ dB. Third line : reconstructed frames over a Gaussian channel with $SNR=8.0$ dB.	87
6.6	Reconstructed frames of "foreman" at 1104 kbs with motion estimation and "hall-monitor" at 566 kbs without motion estimation when RM is partially applied as index assignment. First line : reconstructed frames in a noiseless environment. Second line : reconstructed frames after transmission over a Gaussian channel with $SNR=4.33$ dB. Third line : reconstructed frames after transmission over a Gaussian channel with $SNR=6.75$ dB.	96

6.7	Reconstructed Frames of "foreman" at 1104 kbs with motion estimation and "hall-monitor" at 566 kbs without motion estimation using GLA and minimax index assignment. First line : reconstructed frames in a noiseless environment. Second line : reconstructed frames after transmission over a Gaussian channel with SNR=4.33 dB. Third line : reconstructed frames after transmission over a Gaussian channel with SNR=6.75 dB.	99
6.8	Unequal Error Protection of a MC-EZBC using RPC codes.	101
7.1	Increasing of the diversity with rotated constellation. Left : Diversity $L = 1$. Right : Diversity $L = 2$	106
7.2	Performance of the $R_p = 3/4$ punctured convolutional code over a Rayleigh fading channel without and with CSI embedded in the Viterbi decoder. . .	107
7.3	Reconstructed Frames. First line : reconstructed frame in a noiseless environment. Second line : reconstructed frame when no protection by a code is used over a Rayleigh fading channel with SNR=10.00 dB. Third line, Left : the reconstructed frame when only the lowest frequency spatial subband of the temporal approximation frame is protected and no CSI is supposed Right : Same type of protection as above but with CSI embedded in the Viterbi algorithm. (Rayleigh channel with SNR=10.00 dB).	109
7.4	System model with rotation prior fading channel.	110
7.5	Reconstructed Frames. First line : reconstructed frame in a noiseless environment. Second line : reconstructed frame after transmission over a Rayleigh fading channel with SNR=9.00 dB ; Left : when no rotation matrix is used. Right : with rotation matrix prior transmission. Third line : reconstructed frame after transmission over a Rayleigh fading channel with SNR=11.00 dB ; Left : when no rotation matrix is used. Right : with rotation matrix prior transmission.	114
7.6	Reconstructed Frames. First line : reconstructed frame in a noiseless environment. Second line : reconstructed frame after transmission over a Rayleigh fading channel with SNR=9.00 dB ; Left : when no rotation matrix is used. Right : with rotation matrix prior transmission. Third line : reconstructed frame after transmission over a Rayleigh fading channel with SNR=11.00 dB ; Left : when no rotation matrix is used. Right : with rotation matrix prior transmission.	115

Liste des tableaux

1	Distorsion de la source D_s et distorsion totale D sur un canal Gaussien de $SNR = 4.33$ dB, pour les cas d'étiquetage non codé et codé.	14
2	PSNR moyen (en dB) de la séquence reconstruite "hall-monitor" à différents débits et différents SNR sur un canal gaussien.	15
3	PSNR moyen de la séquence "hall-monitor" reconstruite, transmis sur un canal Rayleigh non-sélectif, avec et sans rotation imposée avant la transmission.	17
1.1	Estimation of the parameters of a vertical subband in different spatial and temporal resolution levels modeled by a Gaussian mixture.	9
1.2	Estimation of the parameters of a vertical subband in different spatial and temporal resolution levels modeled by a GGD.	9
2.1	Without motion estimation : Parameter estimation : the first column indicates the various spatio-temporal neighbors whose weights are estimated (see Fig.1.4). The second column indicates the value of the true parameters, the next four ones the mean square error of the estimation by the four proposed methods over 50 realizations.	24
2.2	With motion estimation : Parameter estimation : the first column indicates the various spatio-temporal neighbors whose weights are estimated (see Fig.1.4). The second column indicates the value of the true parameters, the next four ones the mean square error of the estimation by the four proposed methods over 50 realizations.	24
3.1	MSE of the reconstruction of the first frame of the original sequence, when prediction of a subband at different spatial resolution levels and at different temporal resolution levels is used, compared with setting to zero the lost coefficients.	33
3.2	MSE of a reconstructed frame of the original sequence when we lose the first or second spatial resolution level of a temporal detail frame at each temporal resolution level.	34
3.3	MSE of the reconstructed first frame of a GOF of original sequence "foreman" when we lose the second spatial resolution level of the first temporal detail frame at each temporal resolution level.	37
6.1	Quantization and modelling of different spatio-temporal subbands.	77

6.2	Quantization results for different spatio-temporal subbands, after classification of the subband coefficients.	80
6.3	Average PSNR of the "hall-monitor" and "foreman" CIF sequences, coded by the four schemes and for two different states of the Gaussian channel. .	83
6.4	Coding rates and corresponding matrices of quantization.	85
6.5	Average PSNR of the "hall-monitor" and "foreman" CIF sequences, without motion estimation.	86
6.6	Average PSNR of the "hall-monitor" and "foreman" CIF sequences, with motion estimation/compensation in the temporal transform.	86
6.7	Source distortion D_s and global distortion D (at $SNR = 4.33$ dB) at a bitrate 295 kbs for different subbands of the "hall-monitor" sequence without motion estimation, fr designates the frame to which the examined subband belongs to.	89
6.8	Source distortion D_s and global distortion D (at $SNR = 4.33$ dB) at a bitrate 566 kbs for different subbands of the hall-monitor sequence without motion estimation, fr shows the frame that the examined subband belongs to.	90
6.9	Source distortion D_s and global distortion D (at $SNR = 4.33$ dB) at a bitrate 1027 kbs for different subbands of the hall-monitor sequence without motion estimation, fr shows the frame that the examined subband belongs to.	91
6.10	Source distortion D_s and global distortion D (at $SNR = 4.33$ dB) at a bitrate 1550 kbs for different subbands of the hall-monitor sequence without motion estimation, fr shows the frame that the examined subband belongs to.	92
6.11	Source distortion D_s and global distortion D (at $SNR = 4.33$ dB) at a bitrate 295 kbs for different subbands of the "hall-monitor" sequence, when a $RM(2, 4)$ is chosen as the space of index assignment.	92
6.12	Source distortion D_s and global distortion D (at $SNR = 4.33$ dB) at a bitrate 566 kbs for different subbands of the "hall-monitor" sequence, when a $RM(2, 4)$ is chosen as the space of index assignment.	92
6.13	Source distortion D_s and global distortion D (at $SNR = 4.33$ dB) at a bitrate 1027 kbs for different subbands of the "hall-monitor" sequence, when a $RM(2, 4)$ is chosen as the space of index assignment.	93
6.14	Source distortion D_s and global distortion D (at $SNR = 4.33$ dB) at a bitrate 1550 kbs for different subbands of the hall-monitor sequence, when a $RM(2, 4)$ is chosen as the space of index assignment.	93
6.15	Average PSNR of the "hall-monitor" CIF sequence, without motion estimation with and without $RM(2, 4)$ chosen as the index assignment.	93
6.16	Average PSNR of the "hall-monitor" and "foreman" CIF sequences, without motion estimation and index assignment with $RM(2, 4)$	94
6.17	Average PSNR of the "hall-monitor" and "foreman" CIF sequences, with motion estimation/compensation and index assignment with $RM(2, 4)$	95

6.18	Average PSNR of the reconstructed sequence "hall-monitor" without motion estimation for different bitrates and different SNR. JSC denotes our first coding scheme with structured quantizer and uncoded assignment, $JSC + RM$ the second one with partial coded index assignment and finally, MCA describes the algorithm using unstructured quantizer and minimax index assignment.	98
6.19	Average PSNR of the reconstructed sequence "foreman" without motion estimation for different bitrates and different SNR. JSC denotes our first coding scheme with structured quantizer and uncoded assignment, $JSC + RM$ the second one with partial coded index assignment and finally, MCA describes the algorithm using unstructured quantizer and minimax index assignment.	100
6.20	Average PSNR of the reconstructed sequence "hall-monitor" with motion estimation for different bitrates and different SNR. JSC denotes our first coding scheme with structured quantizer and uncoded assignment, $JSC + RM$ the second one with partial coded index assignment and finally, MCA describes the algorithm using unstructured quantizer and minimax index assignment.	101
6.21	Average PSNR of the reconstructed sequence "foreman" with motion estimation for different bitrates and different SNR. JSC denotes our first coding scheme with structured quantizer and uncoded assignment, $JSC + RM$ the second one with partial coded index assignment and finally, MCA describes the algorithm using unstructured quantizer and minimax index assignment.	102
6.22	Average PSNR of the reconstructed sequence "hall-monitor" for different bitrates and different SNR. JSC+RM denotes our coding scheme with partial coded index assignment and MC-EZBC is the standard coder followed by UEP by RPC codes.	103
6.23	Average PSNR of the reconstructed sequence "foreman" for different bitrates and different SNR. JSC+RM denotes our coding scheme with partial coded index assignment and MC-EZBC is the standard coder followed by UEP by RPC codes.	104
7.1	Average PSNR of "foreman" CIF sequence, coded by the four schemes and for three different states of the Rayleigh channel, with or without CSI embedded in the Viterbi decoder.	108
7.2	Average PSNR of the reconstructed "hall-monitor" CIF sequence without motion estimation when transmitted over a flat-Rayleigh channel, with and without rotation matrix applied prior to transmission.	111
7.3	Average PSNR of the reconstructed "foreman" CIF sequence without motion estimation when transmitted over a flat-Rayleigh channel, with and without rotation matrix applied prior to transmission.	112
7.4	Average PSNR of the reconstructed "hall-monitor" sequence with motion estimation when transmitted over a flat-Rayleigh channel, with and without rotation matrix applied prior to transmission.	112

7.5	Average PSNR of the reconstructed "foreman" sequence with motion estimation when transmitted over a flat-Rayleigh channel, with and without rotation matrix applied prior to transmission.	113
7.6	Average PSNR of the reconstructed sequence "hall-monitor" with motion estimation for different bitrates and different SNR. JSC denotes our coding scheme with the rotation matrix and MC-EZBC is the standard coder followed by UEP by RPC codes.	116
7.7	Average PSNR of the reconstructed sequence "foreman" with motion estimation for different bitrates and different SNR. JSC denotes our coding scheme with the rotation matrix and MC-EZBC is the standard coder followed by UEP by RPC codes.	116

List of Notations

Statistical models for wavelet coefficients

c : the current coefficient or "child" when is referred to a coefficient at the coarse scale and similar orientation.

P : the "parent" of the current coefficient which means the coefficient at the coarse scale and similar orientation.

Q : the spatio-temporal neighborhood consisting of the magnitudes of coefficients at other orientations, scales and spatial locations of the current coefficient.

Spatio-temporal modelling of the wavelet coefficients in a $t + 2D$ scheme

$a_{n,m}$: the square of the magnitude of coefficient in the (n, m) position.

$l_{n,m}$: the predictor of the $a_{n,m}$.

$\sigma_{n,m}^2$: the variance of the coefficient in the (n, m) position.

Overview of Joint Source-Channel coding schemes

$\mathbf{x}$: d -dimensional source vector.

M : size of source codebook.

$\mathbf{y}$: d -dimensional reproduction vector.

$\mathbf{b}$: n -bit binary codeword.

$RM(r, m)$: Reed-Muller code of order r and m a positive integer $0 \leq m \leq r$

Joint Source-Channel coding based on linear labelling and rotated constellations

D : total distortion

D_s : source distortion

D_c : channel distortion

$\mathbf{G}_{d,n}$: $d \times n$ ($d \leq n$) quantization matrix.

$\mathbf{G}'_{n,r} : n \times r$ ($r \leq n$) quantization matrix.

$\mathbf{U}_n : n \times n$ generator matrix of full-rotated $\mathbb{Z}^n$ lattice constellations where n is a power of two.

$R^n : n \times n$ generator matrix of full-rotated $\mathbb{Z}^n$ lattice constellations with $n = (p - 1)/2$ where $p \geq 5$ is a prime.

Joint Source-Channel coding of a video on a Gaussian channel

β : parameter which scales the lattice constellation to the source dynamics. R_p : rate of the punctured convolutional code

Joint source-channel coding of a video on a flat-Rayleigh fading channel

$\mathbf{a}$: vector of random fadings coefficients with Rayleigh distribution.

List of abbreviations

AWGN : Additive White Gaussian Noise
BCJR : Bahl, Cooke, Jelinek and Ravin trellis
BER : Bit Error Rate
BSC : Binary Symmetric Channel
COVQ : Channel Optimized Vector Quantizer-Quantization
DCT : Discrete Cosine Transformation
EZW : Embedded Zerotree Wavelet coder
GGD : Generalized Gaussian Distribution
GLA : Generalized Lloyd Algorithm
GOF : Group Of Frames
HMT : Hidden Markov Tree
IA : Index Assignment
JSC : Joint Source-Channel
JSC : Joint Source-Channel with coded index assignment by Reed-Muller codes
LS : Least Squares criterion
MC : Motion-Compensation
MC-EZBC : Motion-Compensated Embedded Zerotree Block Coder
MCA : Minimax Cover Algorithm
MEE : Maximum Entropy Encoder
ML : Maximum Likelihood criterion
MLS : Modified Least Squares criterion
MSE : Mean Square Error
MSGM : Minimal Span Generator Matrix
PSNR : Peak Signal-to-Noise Ratio
RM : Reed-Muller linear code
RPC : Rate-Punctured Convolutional codes **SNR** : Signal-to-Noise Ratio
SPIHT : Set Partitioning In Hierarchical Trees
SQNR or **SQR** : Signal-to-Quantization-Noise Ratio
UEP : Unequal Error Protection
VQ : Vector Quantizer-Quantization

Résumé de la thèse

Introduction-Motivation

L'objet de cette thèse est de proposer un codage conjoint source-canal de séquences vidéo pour la transmission sur des canaux sans fil.

D'après le théorème de séparation de Shannon, l'optimisation d'un système de transmission passe par l'optimisation séparée du codeur source et du codeur canal. Cependant, cette optimisation n'est valable que pour des blocs de données de taille infiniment longue, ce qui se traduit en une complexité élevée des deux codeurs, prohibitive pour des systèmes temps réel. La solution la plus récente dans ce domaine est d'optimiser des combinaisons de blocs de la chaîne de transmission afin de diminuer la complexité sans sacrifier les performances. Le but dans un tel schéma de codage conjoint source-canal est de minimiser la distorsion globale du système.

Les travaux dans ce domaine se répartissent essentiellement en deux catégories : l'optimisation du codeur de canal en fonction de la source [61], [62], [44], [12], [88] et l'optimisation du codeur de source en fonction du canal [21], [19], [23], [22]. Dans le premier cas, le code source est généré pour un canal sans bruit. Le codeur de canal est ensuite optimisé en fonction de la source afin de minimiser la distorsion globale du système. Dans le deuxième cas, le code de la source est directement optimisé pour un canal bruité connu, afin de minimiser la distorsion globale du système. Une des caractéristiques cruciales pour la performance du système, surtout lorsque l'on utilise un quantificateur vectoriel, est l'étiquetage binaire, c.a.d. l'assignation d'un mot de code source à un mot de code canal. De façon générale, il est souhaitable que les mots de code proches en distance euclidienne correspondent à des étiquettes binaires proches en distance de Hamming ; ainsi, si un bit est erroné, la distorsion engendrée reste faible.

Le système de codage conjoint source-canal proposé dans cette thèse est fondée sur un quantificateur vectoriel structuré et une assignation linéaire d'étiquette qui minimisent simultanément la distorsion canal et la distorsion source. Le quantificateur vectoriel est construit à partir de constellations provenant de réseaux de points, lesquels satisfont la propriété de diversité maximale et minimisent la distorsion source. La distorsion canal est également minimisée par l'étiquetage linéaire [46].

Codage conjoint source-canal de sources gaussiennes

Soit $\mathbf{x}$ un vecteur de source d -dimensionnel et $C = \{\mathbf{y}_1, \mathbf{y}_1, \dots, \mathbf{y}_M\}$ un dictionnaire d -dimensionnel de taille M composé de vecteurs de code $\mathbf{y}_i$ (ou représentants). Le diction-

naire est généré par la partition de l'espace $\mathbb{R}^d$ en M cellules de quantification distinctes S_i , $1 \leq i \leq M$ et en associant chaque cellule à un seul vecteur de code. Une fonction $A : \mathbf{y}_i \rightarrow \mathbf{b}_j$ relie un représentant $\mathbf{y}_i$ à un mot de code binaire $\mathbf{b}_j$ (index). Ce mot de code de taille $n = \log_2 M$ est supposé transmis sur le canal.

Sous l'hypothèse qu'un codeur maxentropique est utilisé (les mots de code sont dans ce cas équiprobables), la distorsion totale D est donnée par :

$$D = D_s + D_c$$

où D_s est la distorsion de la source liée à l'erreur de quantification et D_c est la distorsion du canal liée à une mauvaise interprétation du mot reçu.

Dans [46], il est montré que la distorsion d'un canal binaire symétrique est minimisée lorsque l'étiquetage est linéaire. Partant de ce résultat les auteurs de [5] ont construit des transformations linéaires qui, de plus, minimisent asymptotiquement la distorsion d'une source Gaussienne. L'étiquetage linéaire est décrit par une matrice $\mathbf{G}_{d,n}$ qui doit transformer une suite de variables aléatoires identiquement distribuées en une suite de variables aléatoires de même distribution de probabilité que la source.

Soit $\mathbf{U}_n$ une matrice $n \times n$. Les lignes et les colonnes de $\mathbf{U}_n$ sont notées par L_{in} et C_{jn} respectivement, où $1 \leq i, j \leq n$. Si J est un sous-ensemble de $\{1, \dots, n\}$, $C_{jn}(J)$ est une troncation de la colonne j -th de $\mathbf{U}_n$ selon le sous-ensemble d'indices de J . En fonction de $\mathbf{U}_n$ il est possible d'étiqueter linéairement $BPSK_n = \{-1, +1\}^n$ sur un nouvel ensemble dénoté $\mathbf{U}_n \cdot BPSK_n$. A J fixé, lorsque n augmente, on obtient un dictionnaire $S_n(J)$ ayant comme représentants :

$$y_n = \sum_{j=1}^n b_j C_{jn}(J)$$

où $\mathbf{b} = (b_1, \dots, b_n)^t \in BPSK_n$.

Si l'on veut obtenir une famille de matrices $\mathbf{U}_n$ telle que $S_n(J)$ soit un dictionnaire asymptotiquement Gaussien et qui minimise la distorsion de la source D_s lorsque $n \rightarrow \infty$, il est montré dans [5] que $\mathbf{U}_n$ doit être orthogonale avec des coefficients qui tendent uniformément vers 0 lorsque $n \rightarrow \infty$. Dans ce cas, l'étiquetage,

$$\mathbf{b} \in BPSK_n \rightarrow (\mathbf{G}_{d,n} \mathbf{b} \in S_n(J)),$$

où $\mathbf{G}_{d,n} = \mathbf{U}_n(J)$, est linéaire et le dictionnaire obtenu est asymptotiquement Gaussien.

Ces deux conditions peuvent s'obtenir à partir de constellation de réseaux de points qui possèdent la propriété de diversité maximale. Alors, si $\mathbf{U}_n$ est une matrice $n \times n$ qui génère une constellation à diversité maximale, on construit $\mathbf{G}_{d,n}$ en prenant n'importe quel ensemble de d lignes de $\mathbf{U}_n$. Les mêmes propriétés s'obtiennent lorsque l'on choisit des colonnes de la matrice $\mathbf{U}_n$. On appelle $\mathbf{G}'_{n,r}$ $n \times r$ ($r \leq n$) la matrice construite de telle façon.

Constellations à diversité maximale

Avant d'introduire la contribution de cette thèse, nous présentons un résumé de propriétés des constellations de réseaux de points à diversité maximale. L'utilisation de systèmes à diversité maximale est plutôt liée aux canaux à évanouissements pour lesquels

le décodeur, s'il dispose de plusieurs répliques du signal émis, est capable de reproduire l'information. Ceci justifie la popularité des diversité en temps, diversité en fréquence ou la diversité d'antennes pour de canaux à évanouissements. Les derniers temps, la diversité de modulation, qui n'est qu'une modulation tournée, devient, également, un outil assez populaire. Elle correspond au nombre minimal de composantes différentes entre toute paire de points d'une constellation.

La théorie algébrique des nombres permet de construire des constellations multidimensionnelles extraites de réseaux de points à diversité maximale. Dans [4], [33], [34], [9], nous trouvons que les réseaux provenant du plongement canonique dans les corps de nombres algébriques réels garantissent à la constellation obtenue une diversité maximale égale à la dimension de la constellation. Si l'on applique le plongement canonique à un idéal particulier de cet anneau, on obtient une version tournée $\mathbb{Z}_{n,n}$ du réseau $\mathbb{Z}_n$. Plus précisément, dans [4], [34], le corps choisi est le sous-corps réel d'un corps cyclotomique.

Dans cette thèse, nous utilisons deux constructions différentes de $\mathbb{Z}_{n,n}$ provenant du plongement canonique d'un idéal d'anneau de sous-ensemble réel dans un corps cyclotomique. La première [4] construit des réseaux $\mathbb{Z}_n$ tournés où n est de la forme $n = (p-1)/2$ et $p \geq 5$ est un nombre premier. On obtient de cette façon des constellations de taille $n = 2, 3, 5, 6, 8, 9, 11, 14, 15, 18, 20, 21, \dots$. Si, ensuite, on mélange des constellations produites de cette façon, on obtient de nouvelles constellations de taille $n = 10, 16, 22, 24, \dots$. La deuxième construction [34] produit des réseaux $\mathbb{Z}_n$ tournés, où n est une puissance de 2. Cette méthode est plus simple que la première puisqu'elle ne nécessite pas de réduction de base. De plus, elle est plus appropriée pour le traitement vidéo, comme nous le verrons par la suite, où en raison de l'existence de dépendances entre les coefficients d'ondelettes, il est préférable que la dimension des vecteurs de source soit un multiple de quatre. Nous avons comparé ces deux méthodes de constructions de constellations tournées sur un canal Rayleigh, pour différentes dimensions, et avons constaté qu'elles donnaient, quasiment, les mêmes performances.

Afin d'obtenir ces performances, nous avons utilisé comme algorithme de décodage le décodage universel par sphères [92] qui utilise le critère du maximum de vraisemblance. Cet algorithme est applicable aussi bien sur le canal Gaussien additif (AWGN) que sur le canal de Rayleigh. Son idée principale est la suivante : il limite la recherche, parmi tous les points du réseau, du celui qui est le plus proche du vecteur reçu, aux points du réseau qui se trouvent dans une sphère de rayon $\sqrt{C}$ centrée sur le point reçu. Sur un canal de Rayleigh, où l'on suppose que les évanouissements sont indépendants grâce à un entrelaceur, la complexité devient un peu plus élevée, puisque le réseau change à chaque vecteur de coefficients d'évanouissements. De plus, sur un canal Rayleigh, le choix du rayon de décodage devient un facteur important pour la vitesse de l'algorithme. Il est ainsi souhaitable de pouvoir adapter le choix du rayon en fonction des valeurs des coefficients d'évanouissement. Cependant, la complexité de cet algorithme limite son utilisation à des dimensions de réseau inférieures ou égales à 32.

Codage conjoint source-canal de séquences vidéo

Le but de cette thèse est d'étendre le schéma du codage conjoint source-canal, développé pour les sources gaussiennes, dans le domaine vidéo où la distribution des sources est loin d'être gaussienne. Afin d'atteindre ce but nous avons étudié, dans un premier temps, les dépendances entre les coefficients d'ondelettes provenant d'une décomposition $t + 2D$, avec ou sans estimation de mouvement.

Modèle Statistique

Dans ce cadre là, nous avons étudié des modèles stochastiques, développés pour les images fixes puis étendus au domaine vidéo. Pour les images fixes, nous avons classé ces modèles dans trois catégories, en fonction du type de dépendances entre les coefficients d'ondelettes qu'ils peuvent prendre en compte. Dans la première catégorie, nous trouvons des modèles qui supposent une grande corrélation d'amplitudes de coefficients d'ondelettes à travers des différents niveaux de décomposition spatiale [78], [74], [73]. Le deuxième cas contient des modèles où la corrélation d'amplitudes des coefficients d'ondelettes s'exploite dans la même sous-bande [60], [53] et enfin nous trouvons des modèles qui considèrent ces deux types de dépendance, dans la même sous-bande et dans les sous-bandes de différents niveaux spatiaux [84], [85], [52]. Dans le domaine vidéo, nous avons remarqué que les modèles utilisés [45], [16], [55], [63] sont une extension des modèles développés pour des images fixes.

Ensuite, nous avons étendu le modèle, développé pour des images fixes [84], [85] et qui appartient à la troisième catégorie. Dans [84], [85] Simoncelli a constaté que, pour des images fixes, les coefficients d'ondelettes ne sont pas indépendants et qu'il existe une grande corrélation entre le carré de l'amplitude d'un coefficient avec celui de ses voisins dans la même sous-bande, avec celui du coefficient au niveau spatial supérieur de même orientation (parent); au niveau spatial supérieur, d'orientations adjacentes et de même positions ("aunts"); enfin, au même niveau spatial, d'orientations adjacentes et de même positions ("cousins"). Il a proposé un modèle statistique conjoint où la probabilité conditionnelle d'un coefficient est gaussienne avec une variance qui dépend du carré de l'amplitude des coefficients, présentés avant, qui constituent le "voisinage".

$$P(c|\vec{Q}) = \mathcal{N}(0; \sum_k w_k Q_k^2 + \alpha^2)$$

où $Q = (Q_k)_k$ est le voisinage. Les poids w_k et le paramètre α sont estimés avec le critère des moindres carrés.

Dans notre étude, afin d'examiner les dépendances entre les coefficients d'ondelettes d'une décomposition $t + 2D$, nous avons considéré un voisinage étendu. Après avoir observé les histogrammes d'un coefficient en fonction de plusieurs voisinages différents, nous avons constaté que le modèle de Simoncelli est aussi valable dans le domaine vidéo mais qu'en même temps les coefficients sont dépendants des coefficients qui se trouvent dans les sous-bandes du niveau temporel supérieur. Ainsi, nous avons ajouté dans le voisinage considéré le parent spatio-temporel, les (haut et gauche) voisins du parent spatio-temporel et les "aunts" spatio-temporelles. Dans la Fig. 1, nous présentons le type d'histogrammes (en

domaine log-log) d'un coefficient en fonction du nouveau voisinage considéré, c'est-à-dire les voisins spatiaux et spatio-temporels.

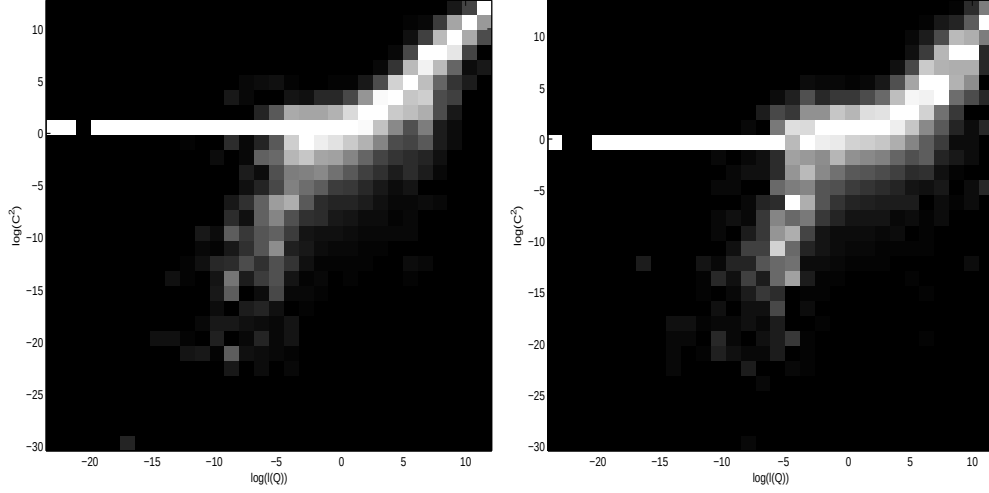


FIG. 1 – Histogrammes conditionnels des coefficients du deuxième niveau temporel, du premier niveau spatial et d'orientation verticale. À gauche : sans estimation de mouvement et à droite avec d'estimation/compensation de mouvement.

Nous avons remarqué que la partie droite de la distribution est concentrée autour d'une pente unitaire alors que la partie gauche est plutôt concentrée autour de l'axe ligne horizontale. De plus, nous avons observé que le nouveau voisinage permet une meilleure concentration autour de ces deux lignes que l'utilisation d'un voisinage totalement spatial.

Soit $c_{n,m}$ le coefficient considéré et $Q_k(n, m)$ ses voisins. La prédiction de $a_{n,m} = |c_{n,m}|^2$ peut être donnée par :

$$l_{n,m} = \sum_k w_k |Q_k(n, m)|^2, \quad (1)$$

où $\mathbf{w} = (w_k)_k$ est le vecteur de poids.

La Fig. 1 suggère l'utilisation du modèle suivant :

$$\log a_{n,m} = \log(l_{n,m} + \alpha) + z_{n,m}, \quad (2)$$

où $z_{n,m}$ est un bruit additif. Dans le cas où $l_{n,m}$ prend des valeurs importantes, la dépendance entre $\log a_{n,m}$ et $\log l_{n,m}$ est presque linéaire, ceci étant en accord avec la partie droite de la Fig. 1. Enfin, la constante α est utilisée afin de décrire la partie gauche du log-log histogramme conditionnel.

Par la suite, nous obtenons :

$$|c_{n,m}| = (l_{n,m} + \alpha)^{1/2} e^{z_{n,m}/2} \quad (3)$$

et si nous réintroduisons le signe, nous avons :

$$c_{n,m} = (l_{n,m} + \alpha)^{1/2} e^{z_{n,m}/2} s_{n,m}, \quad (4)$$

où $s_{n,m} \in \{-1, 1\}$. Nous supposons que le bruit est gaussien, c'est-à-dire : $\beta_{n,m} = e^{z_{n,m}/2} s_{n,m} \sim \mathcal{N}(0, 1)$. Ainsi, $g(c_{n,m} | l_{n,m}) \sim \mathcal{N}(0, l_{n,m} + \alpha)$, et nous aboutissons finalement à une distribution gaussienne conditionnelle :

$$g(c_{n,m} | \sigma_{n,m}^2) = \frac{1}{\sqrt{2\pi}\sigma_{n,m}} e^{-\frac{c_{n,m}^2}{2\sigma_{n,m}^2}} \quad (5)$$

où

$$\sigma_{n,m}^2 = \sum_k w_k |Q_k(n, m)|^2 + \alpha \quad (6)$$

et $\mathbf{Q}(n, m) = (Q_k(n, m))_k$ est le vecteur de voisins.

Ainsi, la probabilité conditionnelle des coefficients dans une sous-bande est gaussienne, avec une variance dépendante de l'ensemble des voisins spatio-temporels.

Ensuite, nous avons proposé des nouveaux estimateurs de paramètres de ce modèle (poids w_k et α) et avons montré qu'ils ont des performances statistiques meilleures que l'estimateur par le moindres carrés utilisé par Simoncelli.

Applications du modèle statistique

Nous avons utilisé ce modèle statistique pour prédire les coefficients des sous-bandes qui n'ont pas été reçus par le décodeur. En effet, le décodeur reçoit d'abord les niveaux spatio-temporels supérieurs et ensuite l'utilisation du modèle peut prédire les restes. La loi conditionnelle suivie par les coefficients est utilisée pour établir un estimateur d'erreur quadratique moyenne optimal de l'amplitude de chaque coefficient, en fonction de ses voisins spatio-temporels. L'expression mathématique du prédicteur est donné par :

$$|\hat{c}_{n,m}| = E\{|c_{n,m}| \mid \mathbf{Q}(n, m)\} = \int_{-\infty}^{\infty} |c_{n,m}| g(c_{n,m} \mid \sigma_{n,m}^2) dc_{n,m}.$$

Ainsi, nous obtenons :

$$|\hat{c}_{n,m}| = \sqrt{\frac{2}{\pi}} \sigma_{n,m},$$

où $\sigma_{n,m}$ est donné par Eq. 6.

Le choix de voisins spatio-temporels est fait d'une telle façon afin d'éviter la propagation d'erreurs. En supposant que le décodeur reçoit en premier le niveau spatial supérieur de chaque trame, le voisinage de chaque coefficient est constitué du parent spatial, des "aunts" spatiaux, du parent spatio-temporel et ses voisins, et à la fin des "aunts" spatio-temporelles. Cependant, cette méthode de prédiction ne tient pas compte du signe de coefficients. En fait, nous le séparons de l'amplitude de coefficients et nous supposons qu'il est bien protégé.

Ce prédicteur s'est révélé d'une importance capitale dans deux applications que nous avons considéré en compression vidéo $t + 2D$.

Dans la première application, nous avons supposé que le décodeur recevait trois niveaux spatiaux de chaque trame et qu'il devait prédire le niveau spatial de la résolution spatiale

supérieure. Nous avons comparé notre méthode à celle d'un décodeur simple qui aurait mis à zéro les coefficients de ce niveau. Nous avons alors observé une décroissance significative de l'erreur de reconstruction et en conséquence, une amélioration du PSNR des trames reconstruites, en utilisant notre méthode de prédiction.

La deuxième application est intéressante pour la transmission des flux vidéo scalables sur des réseaux IP où nous considérons le problème de perte de paquets d'information. Nous avons distingué trois cas d'intérêt lors de la construction de ces paquets et en conséquence à leur perte :

1. une sous-bande spatiale par paquet ;
2. toutes les sous-bandes de même niveau spatial et temporel et de même orientation, dans un paquet ;
3. toutes les sous-bandes de même niveau spatial de chaque trame, dans un paquet.

1. Nous avons considéré la perte d'une sous-bande (ou d'un paquet) dans différents niveaux spatiaux et temporels et sa prédiction en utilisant notre estimateur. L'erreur de reconstruction et la qualité des trames reconstruites ont montré l'avantage de notre méthode comparée à un décodeur simple qui aurait mis les coefficients du paquet perdu à zéro. Comme prévu, la perte d'un paquet au dernier niveau temporel influence l'erreur de reconstruction plus que la perte d'un paquet dans un autre niveau temporel.

2. Nous avons considéré la perte d'un paquet qui contient toutes les sous-bandes du même niveau spatial et temporel et de même orientation. Nous avons observé que l'utilisation de notre prédicteur dans le cas de la perte d'un paquet au premier niveau temporel conduit à une amélioration du PSNR (de plus de 2.5 dB). En effet, les erreurs de prédiction ne se propagent pas à travers de la synthèse temporelle.

3. Nous avons considéré la perte d'un paquet qui contient toutes les sous-bandes du même niveau spatial dans chaque trame. Comme prévu, à cause de la propagation de l'erreur de prédiction, la perte d'un paquet au dernier niveau spatial et dernier niveau temporel conduit à la détérioration de la qualité de reconstruction de la séquence vidéo.

Dans la troisième application nous avons proposé un scénario beaucoup plus difficile que les précédents. Nous avons considéré que la largeur de la bande passante ne permet pas l'émission des niveaux spatiaux de haute résolution, comme dans la première application, mais que, de plus, quelques paquets des niveaux spatiaux supérieurs sont perdus. Ainsi, dans cette application, nous devons prédire le paquet perdu et le niveau spatial de haute résolution. Les sous-bandes de ce niveau seront prédites par des voisins spatio-temporels qui ont été aussi prédits. Nous avons reconsidéré les trois techniques de construction de paquets, effectué dans la deuxième application. Les résultats de nos simulations ont montré que notre estimateur a la capacité de prédire avec succès non seulement le paquet perdu, mais aussi les niveaux spatiaux de haute résolution à partir de la prédiction du paquet perdu.

Construction du dictionnaire de source pour des séquences vidéo

Dans la première partie de cette thèse où nous avons étudié les dépendances statistiques entre les coefficients d'ondelettes, nous avons aussi constaté que la distribution marginale des sous-bandes spatio-temporelles n'est pas gaussienne. Elle est plus pointue à l'origine

et a des queues plus lourdes qu'une gaussienne généralisée. Nous avons observé qu'un mélange de gaussiennes peut en effet être un modèle mieux adapté à ces sous-bandes vidéo.

Le but est de construire un dictionnaire de source qui se base sur la structure du quantificateur développé pour des sources gaussiennes mais qui soit approprié pour des sources vidéo provenant d'une décomposition $t + 2D$. De plus, il doit être robuste en présence de bruit.

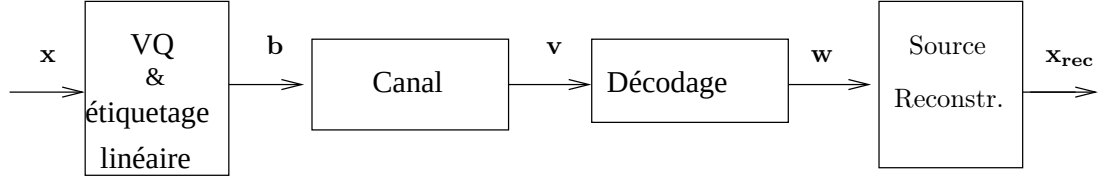


FIG. 2 – Description du système

Les éléments principaux de notre schéma du codage, illustrés dans la Fig. 2, sont les suivants :

- une décomposition temporelle est appliquée sur un groupe de trames, avec ou sans estimation/compensation de mouvement, puis une décomposition spatiale basée sur des filtres biorthogonaux est alors effectuée. $\mathbf{x}$ est le vecteur de coefficients d'ondelettes à l'entrée du quantificateur.
- une quantification vectorielle et un étiquetage linéaire, décrit par la matrice $\mathbf{G}_{d,n}$, produisent les mots du code du canal $\mathbf{b} \in BPSK_n$, qui vont être envoyés sur le canal.
- les canaux étudiés sont gaussien et de Rayleigh avec des évanouissements indépendants.

Dans un premier temps, nous avons développé un algorithme itératif, afin de construire le dictionnaire de source avec des représentants $\mathbf{y} = \mathbf{G}_{d,n}\mathbf{b}$, en minimisant l'expression suivante :

$$\min_{\mathbf{b}} E\|\mathbf{x} - \beta\mathbf{G}_{d,n}\mathbf{b}\|^2 \quad (7)$$

β est un paramètre important car il adapte la constellation à la distribution de la source. Nous avons observé qu'à cause de la nature des sous-bandes, l'utilisation d'un seul β par sous-bande n'est pas suffisante. Nous avons ainsi considéré comme une meilleure adaptation, l'utilisation de deux β par sous-bande, un pour les petits coefficients et un pour les grands coefficients. En utilisant deux β par sous-bande, le gain de codage a été prouvé plus élevé. De plus, nous avons remarqué que les deux valeurs de β diffèrent suffisamment entre elles, confirmant ainsi qu'un seul β n'était pas la solution la plus appropriée.

Le défi de cette décision était de trouver une façon efficace pour classifier les coefficients de la sous-bande en deux parties. Le décodeur doit être en position de calculer le seuil, afin de ne pas être obligé de transmettre un bit supplémentaire pour indiquer la classe à laquelle chaque vecteur appartient. De plus, ce seuil doit être robuste en présence du bruit. Ainsi, nous avons fait appel à notre modèle statistique afin de séparer la sous-bande. Nous

avons estimé la variance moyenne de chaque vecteur, basé sur un voisinage proprement choisi. Les voisins de chaque coefficient, utilisé par le modèle, ont été déjà codés. Nous classifions les vecteurs, en fonction de la valeur de la variance qui offre le gain du codage total de la sous-bande le plus élevé. La même méthode est suivie par le décodeur et le bit supplémentaire n'est alors plus nécessaire.

Cependant, cette méthode de classification peut être appliquée que aux sous-bandes de trames de détails. Pour la sous-bande de basse-fréquence de trames de détail ainsi que pour les sous-bandes de trames d'approximation, comme nous n'avons pas suffisamment d'information pour estimer la variance, nous sommes obligés de transmettre un bit supplémentaire.

Premier résultats sur des canaux Gaussien et Rayleigh

Dans un premier temps, nous avons mis en oeuvre notre schéma de codage sur un canal Gaussien et ensuite sur un canal Rayleigh avec des évanouissements indépendants. Dans ces deux applications, nous avons concaténé partiellement de codes convolutifs perforés.

Nous rappelons qu'un code convolutif perforé est un code à rendement élevé provenant d'un code à rendement faible après élimination périodique de certains symboles. Le grand avantage de ces codes réside dans la complexité de leur décodage (avec l'algorithme de Viterbi) qui reste presque égale à celle d'un code à rendement faible.

Nous avons considéré une décomposition temporelle à 4 niveaux en utilisant des filtres de Haar et une décomposition spatiale en ondelettes à 2 niveaux. Une estimation/compensation de mouvement est appliquée, basée sur l'algorithme de "block matching".

Afin de quantifier les sous-bandes de trames de détails et les sous-bandes de trames d'approximation, nous avons utilisé respectivement une matrice $\mathbf{G}'_{16,4}$ de taille 16×4 et une matrice $\mathbf{G}_{2,16}$ de taille 2×16 .

Dans le cas d'un canal Gaussien, nous avons utilisé un code convolutif perforé de rendement $R_p = 8/9$, provenant d'un code $R = 1/2$ et de mémoire $m = 8$ et dans le cas d'un canal Rayleigh, un code perforé $R_p = 3/4$ et de mémoire $m = 4$.

Nous avons effectué quatre types différents de codage en fonction du degré de protection désiré :

Type 1 : toutes les sous-bandes de trames d'approximation sont protégées par le code perforé.

Type 2 : seulement le niveau spatial supérieur de trames d'approximation est protégé par le code perforé.

Type 3 : seulement la sous-bande de basse-fréquence de trame d'approximation est protégé par le code perforé.

Type 4 : aucune protection n'est imposée.

Au niveau du décodeur, pour les sous-bandes protégées par le code perforé nous avons utilisé l'algorithme de Viterbi et pour les sous-bandes, où aucune protection n'est imposée, nous avons utilisé du décodage ferme. Sur le canal Rayleigh, nous avons distingué deux cas : avec ou sans connaissance du canal, incorporée au calcul de la métrique de Viterbi.

Nous avons remarqué que le schéma de codage sans protection s'est montré assez

robuste en présence de bruit et qu'avec une petite protection par le code perforé (Type 3) l'amélioration du PSNR est assez élevée sans une grande augmentation du débit. Plus précisément, nous obtenons une différence en PSNR de 4 dB pour la séquence "hall-monitor" et de 2.5 dB pour la séquence "foreman" sur un canal Gaussien de $SNR = 4.33$ dB. Sur le canal Rayleigh de $SNR = 10.0$ dB, nous obtenons une différence en PSNR de 5.27 dB pour la séquence "foreman" sans connaissance du canal et 6.14 dB quand on l'incorpore dans l'algorithme de Viterbi. Le débit lorsque on n'impose aucune protection est de 2.34 Mbps et pour le type de codage "Type 3" est de 2.35 Mbps.

Algorithme d'allocation de débit

Les résultats précédents nous ont convaincu que notre système de codage est assez robuste en présence de bruit. Cependant, sans l'utilisation d'un algorithme d'allocation de débit, notre schéma présente des performances assez modeste.

Dans une transmission de vidéo, on peut mettre en évidence différents niveaux d'importance d'information. Une décomposition $t + 2D$, à cause de la scalabilité imposée, tient en compte ces différences. Nous avons développé un algorithme d'allocation de débit qui prend en compte la contrainte de non-négativité des débits dans les différentes sous-bandes et qui alloue les bits en fonction de l'importance des sous-bandes.

Soit I le nombre total de sous-bandes et k une sous-bande quelconque $k \in \{1, 2, \dots, I\}$. Soit N le nombre total de coefficients dans un groupe de trames et n_k le nombre total de coefficients dans la sous-bande k . σ_k^2 est la variance de la sous-bande k et r_k les bits par coefficient dans la sous-bande k . Si nous supposons que la distorsion de surcharge de quantification est négligeable et sous l'hypothèse de haute résolution, un modèle approximatif d'erreur de la quantification et ainsi de la distorsion de la sous-bande peut être donné par :

$$D_k \approx \sigma_k^2 2^{-2r_k}$$

Comme la distorsion du canal est minimisée, à cause de l'étiquetage linéaire, nous considérons que la distorsion qui doit être minimisée en fonction d'un débit total $R = \sum_{k=1}^I \frac{r_k n_k}{N}$ est donnée par :

$$\min_R D = \sum_{k=1}^I D_k = \sum_{k=1}^I \sigma_k^2 2^{-2r_k}$$

Ce problème de minimisation sous contraintes peut être remplacé par un problème de minimisation non-contraint avec l'aide des multiplieurs de Lagrange.

$$J_\lambda = D + \lambda \left(\sum_{j=1}^I \frac{r_j n_j}{N} - R \right)$$

Cependant, l'application de la solution de la minimisation au-dessus, peut donner de mauvais résultats, particulièrement aux bas débits, car elle ne tient pas en compte de la contrainte $r_k \geq 0$.

L'algorithme que nous avons proposé est ainsi présenté ci-dessous :

1. Classifier les sous-bandes dans un ordre décroissant, en fonction de leur variance :
 $+\infty \geq \sigma_1^2 \geq \sigma_2^2 \geq \dots \geq \sigma_l^2 \geq 0$
2. Initialiser l'index d'itérations : $l = 1$
3. Calculer le

$$\lambda_l = -2 \ln 2 \cdot 2^{-\frac{2NR}{M_l}} \prod_{k=1}^l (\sigma_k^2)^{\frac{n_k}{M_l}}$$

où $M_l = \sum_{i=1}^l n_i$

4. Si λ_l satisfait l'inégalité suivante :

$$-2 \ln 2 \sigma_l^2 < \lambda_l \leq -2 \ln 2 \sigma_{l+1}^2$$

poser $l = l + 1$ et aller à l'étape 3, sinon arrêter l'algorithme.

L'algorithme n'alloue des bits qu'aux sous-bandes dont la variance est inférieure à la variance de la position $l + 1$ de la première étape de l'algorithme. Le résultat final de cet algorithme est donné par :

$$r_k = \begin{cases} \frac{N}{M_l} R + \frac{1}{2} \log_2 \left(\frac{\sigma_k^2}{\prod_{k=1}^l (\sigma_k^2)^{\frac{n_k}{M_l}}} \right), & k \in \{1, \dots, l\} \\ 0, & k \in \{l+1, \dots, I\} \end{cases}$$

Ce résultat est obtenu par application immédiate du théorème de dualité de Fenchel en optimisation convexe [54].

Application d'algorithme d'allocation du débit

L'algorithme d'allocation du débit, présenté ci-dessus, définit les tailles de $\mathbf{G}_{d,n}$ et $\mathbf{G}'_{n,r}$ qui minimisent la distorsion totale du système. Cependant, les choix de $\mathbf{G}_{d,n}$ et $\mathbf{G}'_{n,r}$ sont limités par le degré de complexité et par les dépendances parmi les coefficients d'ondelettes. Afin d'exploiter les relations des coefficients avec leur voisins spatiaux et spatio-temporels, les dimensions d de $\mathbf{G}_{d,n}$ et n de $\mathbf{G}'_{n,r}$ sont multiple de quatre. Cependant, afin de pouvoir traiter des hauts et des bas taux de codage, nous nous permettons d'utiliser d et n , respectivement, égaux à deux. De plus, afin d'avoir une complexité basse, les dimensions n de $\mathbf{G}_{d,n}$ et r de $\mathbf{G}'_{n,r}$ sont limités à 16.

Ainsi, les taux de codage possibles sont : 8bits/coeff (2×16), 4bits/coeff (4×16), 2bits/coeff (8×16), 1bits/coeff (16×16), 0.5bits/coeff (16×8), 0.25bits/coeff (16×4) et 0.125bits/coeff (16×2).

Nous avons testé les débits totaux suivants, en bits par pixel : 0.1bpp, 0.16bpp, 0.33bpp et 0.48bpp. Cependant, à cause des bits supplémentaires qui sont nécessaires pour indiquer la classe de vecteurs de trames d'approximation et du codage de vecteurs de mouvement quand on applique une estimation de mouvement, le débit total peut ne pas être le même pour deux groupes différents de trames ou pour deux séquences différentes.

Notons qu'aucun code correcteur d'erreurs n'a été concaténé à la sortie du quantificateur.

Nous avons présenté des résultats de simulations sur un canal Gaussien de $SNR = 4.33\text{dB}$, de $SNR = 6.75\text{dB}$ et de $SNR = 8.0\text{dB}$, pour deux séquences différentes, "hall-monitor" et "foreman" avec et sans estimation de mouvement.

Nous avons remarqué que l'utilisation de l'algorithme d'allocation de débit a apporté une flexibilité accrue à notre système de codage. De plus, nous avons observé que le nouveau schéma atteint des performances supérieures à celles obtenues sans l'utilisation de l'algorithme, pour des débits totaux inférieurs. Le système de codage s'est montré assez robuste en présence de bruit et lorsque $SNR = 8.0\text{ dB}$ il atteint les performances d'une transmission sans bruit.

Dans cette première tentative, en utilisant l'algorithme d'allocation de débit, les mots du code $\mathbf{b}$ envoyés sur le canal Gaussien (Fig. 2) appartiennent à une $BPSK_n = \{-1, +1\}^n$ non-codée. Le "défaut" de ce schéma de codage était qu'en présence d'environnements très bruités ($SNR=4.33\text{ dB}$) l'augmentation du débit n'a pas offert une amélioration significative du PSNR des trames reconstruites, ce qui n'est pas le cas sur un canal sans bruit.

Afin de comprendre et finalement résoudre ce problème, nous avons calculé la distorsion totale du système D avec et sans bruit. Ainsi, quand le canal est sans bruit, la distorsion totale est égale à la distorsion de la source D_s et en présence du bruit, elle est égale à $D = D_s + D_c$, où D_c est la distorsion du canal. Après avoir observé différents tableaux de ces deux valeurs, D en cas du bruit et D_s , pour des sous-bandes différentes, nous avons remarqué que la distorsion D était assez élevée pour quelques sous-bandes en la comparant avec la distorsion D_s de mêmes sous-bandes. La distorsion du canal D_c pour ces sous-bandes était le facteur dominant. De plus, nous avons observé que ce comportement était plus souvent apparu dans les sous-bandes d'énergie élevée, c'est-à-dire dans les sous-bandes de trames d'approximation. Cependant, aux débits élevés, nous avons observé le même comportement du système aussi pour la sous-bande de basse-fréquences de quelques trames de détails, spécialement aux niveaux temporels supérieurs.

Schéma du codage avec un étiquetage linéaire codé sur un canal gaussien

Nous avons résolu les problèmes cités précédemment grâce à un étiquetage linéaire codé imposé sur les sous-bandes dont on a indentifié la nécessité d'une protection accrue. Cette solution sans augmenter le débit total, rend le système puissant en présence du bruit. Afin d'utiliser un étiquetage codé nous avons fait appel aux codes de Reed-Muller.

Rappelons qu'un code de Reed-Muller (r, m) d'ordre r est constitué de l'ensemble des $n' = 2^m$ -uples binaires qui représentent toutes les fonctions binaires polynomiales de degré inférieur ou égal à r , $0 \leq r \leq m$. Sa dimension k est donnée par $k = \sum_{l=0}^r \binom{m}{l}$ et sa distance minimale de Hamming par 2^{m-r} .

Dans ce nouveau schéma de codage, les mots de code $\mathbf{b}$, envoyés sur le canal Gaussien, appartiennent aux 2^k 2^m -uples binaires possibles d'un code de Reed-Muller. Cet étiquetage codé n'est imposé que sur quelques sous-bandes. De plus, nous avons imposé comme contrainte que le taux de codage de chaque sous-bande reste le même que dans le cas non codé et, par conséquent, le débit total ne change pas.

Le choix des paramètres (n', k) du code de Reed-Muller est basé sur un compromis

entre les trois conditions suivantes :

1. la capacité de correction d'erreurs du code
2. n' doit être égal à n (resp. r) de $\mathbf{G}_{d,n}$ (resp. $\mathbf{G}'_{n,r}$) afin de ne pas augmenter le débit.
3. la dimension du code k doit être assez proche de n (resp. r) de $\mathbf{G}_{d,n}$ (resp. $\mathbf{G}'_{n,r}$) afin que le nombre 2^k de mots de code possibles soit assez proche du nombre 2^n (resp. 2^r) de mots de code possibles dans le cas non codé.

Les étapes du nouvel algorithme de codage sont les suivants :

- Calcul de la distortion D_s (environnement sans bruit) et de la distorsion totale D du canal bruité dans le cas d'un étiquetage non codé.
- Si la différence entre elles est élevée, ce qui veut dire que la distorsion D_c est dominante, on limite l'assignation d'étiquettes aux mots du code d'un Reed-Muller (r, m) (en tenant compte des trois conditions présentées ci-dessus).
- Sinon appliquer un étiquetage non codé.

Au décodeur, quand on reçoit des mots de code non codé, on applique un décodage ferme et quand on reçoit des mots de code codé, on applique l'algorithme de Viterbi à travers d'un treillis BCJR [57]. Un code linéaire, comme les codes convolutifs, peut se représenter sous la forme d'un treillis. Afin de diminuer la complexité de l'algorithme de Viterbi, la meilleure représentation sous forme de treillis est celle qui a le plus petit nombre de branches. Dans [57] nous trouvons que le treillis qui minimise le nombre de branches est celui qui a été introduit par Bahl, Cooke, Jelinek et Ravin [2] et qui est isomorphe au "treillis minimal" présenté, plus tard par Forney [24]. Dans une représentation par un treillis BCJR, chaque mot de code d'un code linéaire $C(n, k)$, correspond à un chemin de longueur n . McEliece [57] a introduit une famille de matrices, que nous avons aussi utilisée, qui facilite la construction de ce treillis. Cependant, cette méthode ne reste valable que pour des codes de courte longueur. Dans le cas d'un code long, nous proposons d'utiliser la méthode présentée dans [25], [26] qui, bien que sous-optimale, conserve une faible complexité et des performances assez proches d'un décodage en treillis.

Dans notre système de codage, les sous-bandes dont on a prouvé qu'elles nécessitaient plus de protection ont été quantifiées par $\mathbf{G}_{d,16}$ ou $\mathbf{G}'_{n,16}$. Ainsi, l'utilisation d'un code de Reed-Muller RM(2,4), où $n' = 16$ et $k = 11$, est un compromis suffisant pour les trois conditions présentées précédemment. Dans ce cas-là, l'utilisation d'un décodage avec le treillis BCJR reste encore réalisable et pas trop complexe.

Nous avons observé qu'en utilisant le RM(2,4), la distorsion totale D du système bruité est diminuée mais qu'au même temps la distorsion de la source a augmenté car l'espace de mots de code possibles a diminué de 2^{16} à 2^{11} . Le PSNR des trames reconstruites a suffisamment augmenté, spécialement dans des environnements très bruités (SNR=4.33) dB.

Dans le Tableau 1 nous donnons un court exemple des sous-bandes qui ont besoin plus de protection, pour la séquence "hall-monitor" à 566 kbs, sans estimation de mouvement. Les 4-ème et 5-ème colonnes représentent, respectivement, les valeurs de D_s et D sur un canal gaussien de $SNR = 4.33$ dB avec un étiquetage non codé. On remarque que leur différence est assez élevée. Les 6-ème et 7-ème colonnes représentent, respectivement, les valeurs de D_s et D sur le même canal ($SNR = 4.33$ dB) avec un étiquetage codé en utilisant RM(2,4). On remarque, comme prévu, que la distorsion totale D a beaucoup

diminué et que la distorsion de la source D_s a augmenté.

"hall-monitor" 566 kbs						
Trame d'approximation temporelle						
			étiquetage non codé		étiquetage codé	
niveau temporel	niveau spatial	orientation	D_s	D	D_s	D
4ème	2ème	LL	11.63	6888.86	350.50	378.18
4ème	2ème	Vertical	137.17	844.27	542.07	549.89
4ème	2ème	Horizontal	116.69	937.40	550.58	555.5

TAB. 1 – Distorsion de la source D_s et distorsion totale D sur un canal Gaussien de $SNR = 4.33$ dB, pour les cas d'étiquetage non codé et codé.

Comparaison avec un quantificateur vectoriel non-structuré

Nous avons comparé les deux systèmes de codage, avec et sans étiquetage codé, avec un quantificateur vectoriel optimal du point de vue de la distorsion source et non structuré (GLA [50]) où, pour l'assignation des étiquettes, nous avons appliqué un algorithme nommé "Minimal Cover Algorithm" (MCA) [70]. MCA est un bon algorithme sous-optimal qui utilise le critère de minimisation d'erreur dans un sens minimax et non pas un critère de minimisation de la distorsion moyenne. De cette façon cet algorithme produit des mots de code mieux adaptés au pire cas d'un canal sans, en même temps, sacrifier la performance moyenne.

Nous avons appliqué l'algorithme MCA sur les dictionnaires non-structurés, générés par le GLA sur nos séquences vidéo. Les débits totaux de ce schéma non-structuré sont égaux à ceux utilisés dans les schémas structurés.

Dans le Tableau 2, nous donnons un exemple des performances obtenues, pour la séquence d'"hall-monitor" avec estimation de mouvement à différents débits. Nous présentons le PSNR moyen de la séquence reconstruite, avec trois systèmes de codage différents. JSC représente notre premier système de codage, où nous avons utilisé notre quantificateur structuré avec un étiquetage linéaire non codé et JSC+RM notre deuxième système de codage où un étiquetage codé, avec des codes de Reed-Muller, est appliqué sur une partie des sous-bandes. Au final, MCA est le schéma de codage non-structuré, généré par le GLA en combinaison avec l'algorithme d'assignation d'étiquettes qui utilise le critère d'erreur minimax.

Nous remarquons que l'utilisation d'un étiquetage codé, pour les sous-bandes basse-fréquence, est bien justifiée par l'augmentation du PSNR par rapport au cas non codé. Le schéma de codage JSC+RM est beaucoup plus robuste en présence de bruit que le JSC, spécialement dans le cas où nous avons $SNR = 4.33$ dB.

Le schéma de codage non-structuré, comme prévu, offre une meilleure performance en absence de bruit. Le PSNR de la séquence reconstruite est beaucoup plus élevé que dans le cas du système structuré. Mais en augmentant le niveau de bruit, les performances du schéma non-structuré MCA se dégradent dramatiquement, même avec l'application d'un bon algorithme d'assignation d'étiquettes. En même temps, on peut observer que

"hall-monitor" 1550 kbs			
	sans bruit	SNR=4.33 dB	SNR=6.75 dB
JSC	34.66	26.90	32.84
JSC+RM		32.04	33.49
MCA	40.35	24.35	33.29
"hall-monitor" 1027 kbs			
	sans bruit	SNR=4.33 dB	SNR= 6.75 dB
JSC	32.81	26.57	31.54
JSC+RM		30.90	32.03
MCA	38.18	23.65	32.15
"hall-monitor" 566 kbs			
	sans bruit	SNR=4.33 dB	SNR=6.75 dB
JSC	30.53	26.01	29.77
JSC+RM		29.30	30.05
MCA	34.96	23.58	31.43
"hall-monitor" 295 kbs			
	sans bruit	SNR= 4.33 dB	SNR=6.75 dB
JSC	26.68	23.44	26.21
JSC+RM		24.43	26.21
MCA	32.45	23.53	30.04

TAB. 2 – PSNR moyen (en dB) de la séquence reconstruite "hall-monitor" à différents débits et différents SNR sur un canal gaussien.

les performances des deux schémas structurés (spécialement de JSC+RM) restent assez proche de leurs performances sans bruit même pour des canaux très bruités, ce qui n'est absolument pas le cas pour le schéma non-structuré. Sur des canaux très bruités ($SNR = 4.33$ dB), JSC+RM obtient les meilleures performances parmi les trois schémas de codage.

Dans la Fig. 3, nous présentons un exemple de trame reconstruite, dans le cas d'un canal gaussien du $SNR = 4.33$ dB et $SNR = 6.75$ dB qui a été codée avec les trois systèmes de codage présentés au-dessus, JSC, JSC+RM et MCA. La supériorité de la qualité de la trame codée par JSC+RM même à $SNR = 4.33$ dB est remarquable.

Schéma de codage avec matrice de rotation sur un canal Rayleigh

Le canal de Rayleigh est le modèle le plus proche d'un canal radio-mobile. Les erreurs à la réception se produisent lorsque l'atténuation est trop forte. Cependant, si le décodeur dispose de différentes copies du signal, émises sur des canaux à évanouissements indépendants, la probabilité que l'atténuation soit aussi forte sur toutes les copies diminue largement. La construction de telles répliques correspond à une technique de diversité.

Parmi les différentes techniques classiques de diversité, on trouve la diversité en temps, la diversité en fréquence ou la diversité d'antennes. Une technique, un peu moins classique



FIG. 3 – Trames reconstruites de la séquence "hall-monitor" à 566 Kbs. Première ligne : canal gaussien avec SNR=4.33 dB. Deuxième ligne : canal gaussien avec SNR=6.75 dB. Première colonne : schéma de codage JSC. Deuxième colonne : schéma de codage JSC+RM. Troisième colonne : schéma de codage MCA.

mais également assez efficace, est la diversité de modulation, qui, comme nous l'avons déjà présenté, correspond au nombre minimal de composantes dont deux symboles de la constellation diffèrent. La diversité de modulation peut être produite par des réseaux de points tournés.

Dans notre étude, nous avons utilisé un canal Rayleigh non sélectif à évanouissements indépendants, où l'indépendance est obtenue en ajoutant un entrelaceur et un desentrelaceur au niveau respectivement de l'émetteur et du récepteur.

Dans les premiers tests que nous avons réalisé sur un canal de Rayleigh, nous avons remarqué qu'une BPSK sans protection présente des performances assez faibles. Dans un deuxième temps, nous adoptons l'idée de la rotation de la constellation considérée, afin que sa diversité augmente sans ajouter de la redondance.

La matrice $\mathbf{G}_{d,n}$ (ou $\mathbf{G}'_{n,r}$) qui génère le dictionnaire source structuré, provient du choix des d lignes ou r colonnes d'une matrice $\mathbf{U}_n$ qui construit des constellations de diversité maximale. Ces constellations ne sont que de constellations $\mathbb{Z}^n$ tournées, de diversité $L = n$. Notre nouveau système de codage est basé sur l'idée suivante : avant la transmission sur le canal de Rayleigh, nous concaténons une matrice de rotation $\mathbf{U}_n$ (resp. $\mathbf{U}_r$) à $\mathbf{G}_{d,n}$ (resp. $\mathbf{G}'_{n,r}$). Au décodeur, nous utilisons le décodeur universel par sphères [92] qui utilise le critère du maximum de vraisemblance. Ce système profite d'une rotation à diversité maximale avant la transmission et d'un décodage à maximum de vraisemblance à la réception.

Dans le Tableau 3, nous donnons un exemple de performances obtenues, en PSNR moyen, de la séquence "hall-monitor", sans estimation de mouvement, à différents débits, sur un canal de Rayleigh non sélectif à évanouissements indépendants avec un $SNR = 9.0$ dB et un $SNR = 11.0$ dB. Nous présentons, aussi, le cas où aucune matrice de rotation n'est imposée avant la transmission. L'utilisation d'une matrice de rotation conduit à un gain de qualité de 5 à 10 dB.

"hall-monitor" 1550 kbs		
sans bruit=34.66 dB	Sans rotation	Avec rotation
SNR=9.0	22.82	33.38
SNR=11.0	24.51	34.47
"hall-monitor" 1027 kbs		
sans bruit=32.81 dB	Sans rotation	Avec rotation
SNR=9.0	22.74	32.01
SNR=11.0	24.37	32.74
"hall-monitor" 566 kbs		
sans bruit=30.53 dB	Sans rotation	Avec rotation
SNR=9.0	22.49	29.96
SNR=11.0	24.01	30.47
"hall-monitor" 295 kbs		
sans bruit=26.68 dB	Sans rotation	Avec rotation
SNR=9.0	20.43	26.36
SNR=11.0	21.80	26.65

TAB. 3 – PSNR moyen de la séquence "hall-monitor" reconstruite, transmis sur un canal Rayleigh non-sélectif, avec et sans rotation imposée avant la transmission.

Conclusions

Dans cette thèse, nous avons présenté un système de codage conjoint source-canal pour des séquences vidéo. L'évaluation du système est faite sur des canaux Gaussien et Rayleigh non-sélectif à évanouissements indépendantes.

Notre système de codage se base sur un quantificateur vectoriel structuré, qui provient de réseaux de points de diversité maximale, et sur un étiquetage linéaire. Il a déjà été prouvé que ce système minimise, en même temps, la distorsion canal et la distorsion source pour des sources Gaussiennes. Cette thèse a étendu ce travail aux sources provenant d'une décomposition $t + 2D$ de séquences vidéo, qui ne suivent pas une distribution gaussienne.

Dans un premier temps, nous avons présenté nos études dans le domaine de la décomposition $t + 2D$ des séquences vidéo et ceci dans le but d'exploiter les dépendances hiérarchiques des coefficients d'ondelettes. Nous avons proposé un modèle qui considère que la probabilité conditionnelle d'un coefficient est gaussienne, avec une variance qui est dépendante d'un ensemble des voisins spatio-temporels. Nous avons également proposé

différents estimateurs de paramètres de ce modèle et nous avons montré qu'ils ont des performances statistiques meilleures que l'estimateur des moindres carrés.

La loi conditionnelle suivie par les coefficients est utilisée pour établir un prédicteur des coefficients non reçus par le décodeur. Nous avons présenté différentes applications du prédicteur dans le cas de la vidéo.

Ensuite, nous avons effectué des modifications nécessaires sur le système de codage développé pour des sources gaussiennes, afin d'obtenir un système approprié pour des séquences vidéo et robuste en présence de bruit. Nous avons présenté un algorithme d'allocation de débit qui tient compte de la contrainte de non-négativité des débits et qui offre une flexibilité à notre système de codage. Nous avons considéré deux cas d'intérêt sur un canal gaussien. Dans le premier, nous avons utilisé un étiquetage linéaire non codé et dans le deuxième, un étiquetage linéaire partiellement codé en utilisant les codes de Reed-Muller. Nous avons prouvé que ce dernier système de codage a des bonnes performances même dans des environnements très bruités. Nous avons comparé les deux systèmes avec un quantificateur non-structuré où nous avons appliqué une assignation d'étiquettes adapté au canal. Le système de codage avec notre quantificateur structuré et l'étiquetage codé a prouvé sa robustesse dans des environnements très bruités par rapport au système avec le quantificateur non-structuré.

Au final, nous avons développé un système de codage de séquences vidéo sur un canal Rayleigh non-sélectif à évanouissements indépendants. Nous avons exploité la propriété de la diversité de modulation, fournie par des constellations tournées. Puis nous avons développé un système de codage qui utilise le même quantificateur que dans le cas du canal gaussien, concaténé à la matrice de rotation appropriée. Enfin, nous avons présenté les performances de ce système et nous avons constaté sa supériorité en le comparant à un système qui n'utilise pas une matrice de rotation.

Contributions de la thèse

Les contributions apportées par cette thèse dans le domaine du codage vidéo sont les suivantes :

- Modèle stochastique développés pour la décomposition $t + 2D$ de séquences vidéo
- Estimateur robuste des paramètres du modèle
- Prédicteur optimal des coefficients non reçus par le décodeur
- Codage conjoint source-canal avec un quantificateur vectoriel structuré et étiquetage linéaire noncodé adapté aux séquences vidéo
- Codage conjoint source-canal avec un quantificateur vectoriel structuré et étiquetage linéaire, partiellement codé adapté aux séquences vidéo
- Codage conjoint source-canal avec un quantificateur vectoriel structuré et une matrice de rotation pour des séquences vidéo sur un canal Rayleigh

Organisation de la thèse

Ce document est divisé en 7 chapitres et des annexes.

Le chapitre 1 présente un état de l'art sur les modèles stochastiques utilisés dans le domaine d'images fixes. Nous donnons, ensuite, une analyse plus étendue du modèle, proposé pour des images fixes, que nous avons choisi d'utiliser dans le domaine de la

vidéo. Les modèles stochastiques utilisés dans le domaine de la vidéo sont, en général, des extensions des modèles utilisés dans le domaine des images fixes.

Dans le chapitre 2, nous étudions les dépendances hiérarchiques des coefficients d'une décomposition $t + 2D$ d'une séquence vidéo, en présence ou pas d'une estimation de mouvement. Nous proposons un modèle doublement stochastique afin de capturer ces dépendances et ensuite nous proposons également des estimateurs robustes des paramètres de ce modèle. Nous présentons finalement des résultats qui confirment l'efficacité de nos estimateurs.

Dans le chapitre 3, nous proposons un prédicteur des coefficients perdus provenant de notre modèle stochastique. Nous présentons des applications de ce prédicteur afin d'améliorer la qualité de la vidéo reçue et de prédire des paquets qui seraient perdus dans les réseaux à commutation de paquets.

Dans le chapitre 4 nous présentons un bref état de l'art sur les techniques de codage conjoint source-canal. La classification de ces méthodes nous permet, à la fin du chapitre, de situer notre travail parmi les approches de quantification vectorielle adaptée au canal.

Dans le chapitre 5 nous présentons un schéma de codage conjoint source-canal développé pour des sources gaussiennes. Ce schéma utilise un quantificateur vectoriel provenant de réseaux de points à diversité maximale suivi d'un étiquetage linéaire. Pour cela, nous introduisons les réseaux de points à diversité maximale, ainsi que deux méthodes différentes permettant de les construire.

Dans le chapitre 6 nous étendons le schéma de codage conjoint source-canal développé pour des sources gaussiennes au domaine de la vidéo. Nous présentons les modifications nécessaires afin d'avoir un système efficace et robuste en présence de bruit sur un canal gaussien. Nous développons deux systèmes, un avec un étiquetage linéaire non codé et l'autre avec un étiquetage linéaire codé en utilisant les codes de Reed-Muller. A la fin de ce chapitre, nous comparons ces deux schémas de codage avec un schéma non-structuré dont l'étiquetage est adapté au canal et avec un coder vidéo scalable.

Le chapitre 7 présente les performances de notre système de codage sur un canal de Rayleigh non-sélectif à évanouissements indépendants. Le schéma devient encore plus robuste lorsque nous utilisons une matrice de rotation avant la transmission sur le canal. Cette matrice de rotation augmente la diversité du système et donne des performances remarquables.

Enfin, nous donnons nos conclusions sur les différentes idées présentées dans ce document ainsi que quelques perspectives.

Chapitre 1

Statistical models for wavelet coefficients

In this Chapter, we present an overview of statistical models for wavelet coefficients that have already been used for still images and we analyze the model that we choose to extend to a $t + 2D$ wavelet decomposition of a video sequence. At the end of this chapter, we make a brief summary of models that have already been used in the video domain.

1.1 Introduction

Wavelet-based techniques have already become a very powerful tool in video coding due to their complete spatio/temporal/ SNR/complexity scalability and their increased robustness in error-prone environments. However, in order to obtain successful coding algorithms, the understanding of the dependencies between the wavelets coefficients is necessary.

Hierarchical dependencies between the wavelet coefficients have been largely used for still images [51], as well for coding in methods like EZW [78] and SPIHT [74], as for denoising [77], [84]. They rely on a quad-tree model which has been thoroughly investigated, leading to a joint statistical characterization of the wavelet coefficients [85], [86]. This way of treatment comes as a natural consequence of the properties of the wavelets. Wavelet subband coefficients follow highly non-Gaussian marginal statistics with sharper peaks at zero and more extensive tails. In addition, wavelets act as edge detectors, which are represented by large magnitude coefficients. These coefficients tend to occur at neighboring spatial locations and also, at the same relative spatial locations of subbands at adjacent scales and orientation. Even if the wavelet coefficients are decorrelated, *they are not statistically independent*.

In still images, various wavelet models exist in order to capture the dependencies between the coefficients [51], which can be extended to the video domain. We present them in the following section.

1.2 Overview of statistical wavelet models for still images

Based on the classification of models in [51], we present the various models separated into three different groups.

1. *Interscale Models*, where the magnitudes of the wavelet coefficients are strongly correlated across scales ;
2. *Intrascale Models*, where strong dependencies in the form of spatial clusters exist between wavelet coefficients inside each subband ;
3. *Composite Dependency Models*, where both types of the above dependencies may be combined.

1.2.1 Interscale Models

This type of dependencies was the base of the Shapiro's embedded zerotree wavelet coder [78] (EZW) . In this work, it is shown that even if the wavelet transform produces nearly uncorrelated coefficients, there exist additional dependencies between the squares (or magnitudes) of parents and children. The coefficient at the coarse scale is called the *parent* and all coefficients corresponding to the same spatial location at the next finer scale of similar orientation are called *children*. In addition, an "optimal" estimate, such as the conditional expectation of the child's magnitude given the parent's magnitude would, with high probability, predict that the child's magnitude would be smaller than the magnitude of its parent.

Combining the self similarity across scales with a clever scheme for set partitioning in hierarchical trees (SPIHT), in [74], they have developed an even better coder. Briefly, we mention that they search for sets in spatial-orientation trees in a wavelet transform. Indeed, these trees define the spatial relationship on the hierarchical pyramid. Each node of the tree corresponds to a pixel and its direct descendants correspond to the pixels of the same spatial orientation in the next finer level of the pyramid. The pixels in the highest level of the pyramid are the tree roots. Then, they partition these trees into sets defined by the level of the highest significant bit in a bit-plane representation of their magnitudes and they code and transmit bits associated with the highest remaining bit planes first.

Almost the same property, that the relative magnitude of a wavelet coefficient is related to the magnitude of its parent, has led the authors in [73] to use a hidden Markov tree model (HMT) in order to capture this dependency. The HMT models the non-Gaussian marginal pdf as a two-component Gaussian mixture where each component is labeled by a hidden state signifying whether the coefficient is small or large. The Gaussian component corresponding to the small state has a relatively small variance, capturing the peakiness around zero, while the component corresponding to the large state has relatively large variance, capturing the heavy tails. The persistence of wavelet coefficient magnitudes across scales is modeled by linking these hidden states across scales in a Markov tree. A state transition matrix for each link quantifies statistically the degree of persistence of large/small coefficients.

1.2.2 Intrascale Models

A double stochastic process is the way of modelling the image wavelet coefficient in [60]. The wavelet coefficients are assumed to be conditionally independent zero-mean Gaussian random variables given their variances. These variances are modeled as identically distributed, highly correlated random variables. Assuming that the image pixels are corrupted by additive white Gaussian noise (AWGN), they perform an approximate maximum a posteriori estimation (MAP) of the variance of each coefficient using the observed noisy data *only* in a local neighborhood and a prior model for the variance. In order to obtain this prior model they follow two ways. As a first choice, they treated the variances as deterministic quantities and computed approximate ML estimates and as a second, they used an exponential distribution as a prior for the underlying variance field and computed approximate MAP estimates of the variances.

The above model was inspired by a simpler approach that was proposed in [53], the Estimation-Quantization coder (EQ), where, indeed, the wavelets coefficients are modeled as Generalized Gaussian (GG) distributions and the local variances are considered as unknown deterministic parameters. However, both of them consider as a neighborhood of the coefficient only the adjacent coefficients in the same subband.

1.2.3 Composite Models

The works in this domain try to capture both inter and intra scale dependencies. In [52], in each subband, except the fine scale, the coefficients are partitioned in two classes based on the magnitude of their parents : the set of coefficients that have significant (bigger than a threshold) parents and the set of coefficient that have insignificant parents. The two classes have quite different statistics. The histogram of the first one is proved to be quite spread and a Laplacian distribution provides a good fit. On the other hand the histogram of coefficients with insignificant parents is highly concentrated around zero and the intrascale model, proposed in the EQ coder [53] is used.

Simoncelli in [84], [85], [86], at first, introduces a model for the marginal statistics of the wavelet coefficients. The model supposes that each wavelet subband consists of independent identically distributed random variables drawn by a generalized Laplacian (or "stretched exponential") distribution. Next, as the coefficients of the wavelet tranform are not independent, he proposes a joint statistical model. He assumes a strong correlation between the squared magnitude of a wavelet coefficient and those of its parent and neighbors and develops a prediction scheme based on this assumption. We will presenet this model in more detail in the next section, as we choose it for extension to the video domain.

1.3 Simoncelli's Joint Statistical Model

Simoncelli in [84], [85], is based on the observation that large magnitude coefficients tend to occur at neighboring spatial locations and also at the same relative locations of subbands at adjacent scales and orientations. At the beginning, he considered the conditional histogram of two coefficients representing information at adjacent scales but the

same orientation and spatial location. Fig.1.1, shows this conditional histogram $H(c|P)$ of the "child" coefficient conditioned on a coarser-scale "parent" coefficient, where brightness corresponds to the probability, except that each column has been independently rescaled to fill the full range of display intensities.

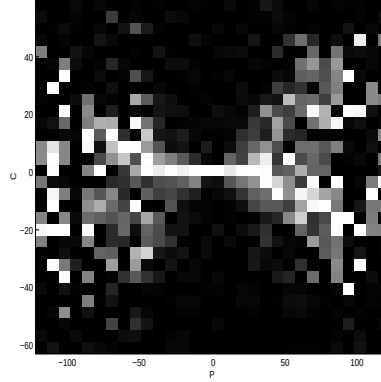


FIG. 1.1 – Conditional histogram for a fine scale horizontal coefficient. Conditioned on the parent (same location and orientation, coarser scale).

The histogram illustrates two significant aspects of the relationship between the two coefficients. First, they are second-order decorrelated, since the expected value of c is approximately zero for all values of P . Second, the variance of the conditional histogram of c clearly depends on P . Thus, although c and P are uncorrelated, they are still statistically dependent.

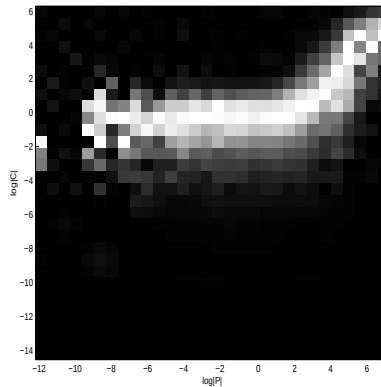


FIG. 1.2 – Example of the log-domain conditional histogram of the fine scale coefficient conditioned on its parent (same location and orientation, coarser scale)

Next, he observed the conditional histogram of these two coefficients in the log domain where the structure of their relationship becomes more apparent. Fig. 1.2 shows the conditional histogram $H(\log_2(c^2)|\log_2(P^2))$. The right side of the distribution is unimodal and concentrated along a unit-slope line which means that in this region the conditional expectation, $E(c^2|P^2)$, is approximately proportional to P^2 . In addition, vertical cross

sections have approximately the same shape for different values of P^2 and finally, the left side of the distribution is concentrated about an horizontal line, suggesting that c^2 is independent of P^2 in this region. He observed that the qualitative form of these statistical relationships also holds for pairs of coefficients at adjacent spatial locations ("siblings"), adjacent orientations ("cousins") and adjacent orientations at a coarser scale ("aunts").

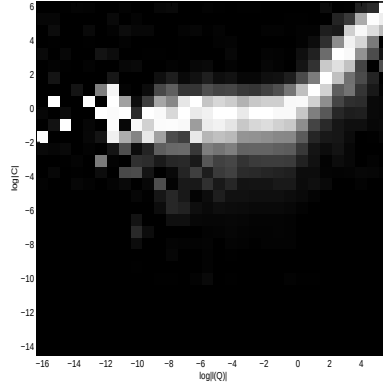


FIG. 1.3 – Example of the log-domain conditional histogram of the fine scale coefficient conditioned on a linear combination of neighboring coefficient magnitudes.

Finally, he observed the histogram of $\log_2(c^2)$ conditioned on a linear combination of the squares of eight adjacent coefficients in the same subband, two coefficients at other orientations and a coefficient at a coarser scale. According to Fig. 1.3, he suggests a simple Markov model, in which the density of a coefficient c , is conditionally Gaussian with variance a linear function of the squared coefficients in a local neighborhood :

$$P(c|\vec{Q}) = \mathcal{N}(0; \sum_k w_k Q_k^2 + \alpha^2)$$

where the neighborhood $Q = (Q_k)_k$ consists of the magnitudes of coefficients at other orientations and adjacent scales, as well as adjacent spatial locations, as mentioned above. The weights $(w_k)_k$ and the α used, are chosen to be least-square optimal.

1.4 Introduction to the $t + 2D$ scheme of decomposition of a video sequence

Before we present the statistical models used in the video domain, we have to introduce the notions of a $t + 2D$ decomposition of a video sequence. In a $t + 2D$ scheme, the input video signal is decomposed into spatio-temporal subbands by temporal analysis and a spatial wavelet transform, as it is illustrated in the Fig. 1.4.

In the case of a small motion, signal energy is concentrated in the temporal low subband, otherwise, the spatial frequencies are shared along the temporal axis and the temporal high subbands contain significant energy. The temporal analysis can be motion-compensated (MC) or without motion estimation. We are going to investigate both of

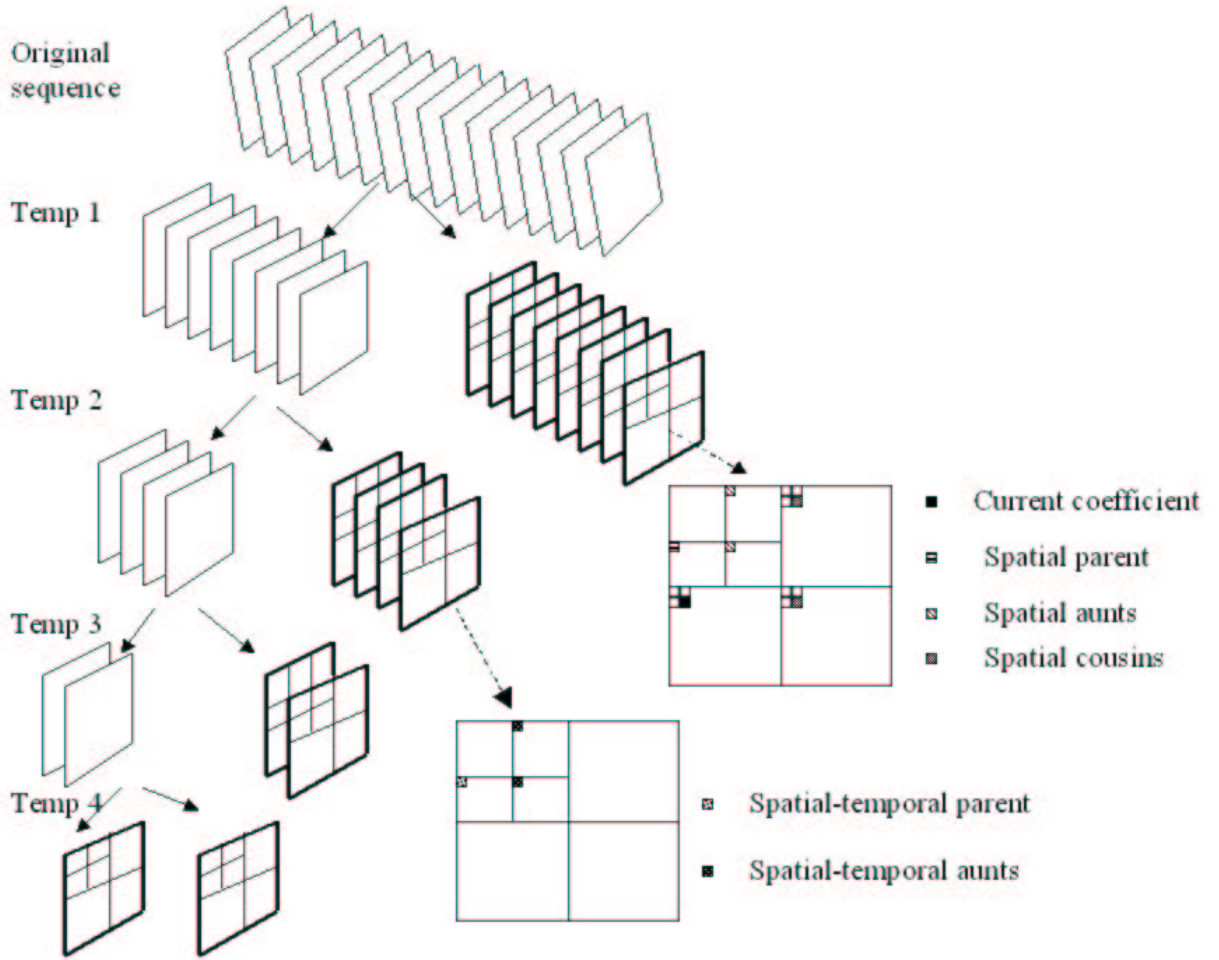


FIG. 1.4 – Spatio-temporal neighbors of a wavelet coefficient in a video sequence (the original Group of Frames is decomposed over four temporal levels). Dependencies are highlighted with its spatial (parent, cousins, aunts) and spatio-temporal neighbors (temporal parent and temporal aunts). Temp i stands for the i -th temporal decomposition level.

them in the following Chapters. MC temporal analysis means that the temporal filter is performed along the motion trajectory [43]. If a Haar filterbank is used, a pair of temporal low- and high- subbands is generated for each two successive input frames by filtering and decimation.

In the case of a motion-compensated coding, the block matching is used more frequently as in any other motion estimation technique. It is based on partitioning a frame into nonoverlapped, equally spaced, fixed size, small rectangular blocks and assuming that all the pixels in the block experience the same translation motion. Various issues are involved in the implementation of the block matching, such as selection of block sizes, matching criteria, search strategies, matching accuracy and its limitations. A detailed description of this estimation technique and its related issues can be found in [79]. We

only mention that in our $t + 2D$ decomposition with ME/MC, we used 8×8 as block size, Mean Square Error (MSE) as matching criterion, full search strategy and pixel or half-pixel matching accuracy.

A problem that arises from the motion compensated technique is the double-connected and unconnected pixels related to uncovered and covered areas. Two regions in the current frame may correspond, by MC, to the same (uncovered) region in the reference frame. This uncovered area appears as double-connected during temporal filtering. Choi's solution [15] for these pixels is to associate them with the first block encountered during MC. Concerning unconnected pixels, the original value is inserted into the low-frequency temporal subband for the previous image and a displaced frame difference value is taken for the current one. However, in a nonlinear lifting formulation of the temporal analysis [68], among the pixels connected to the same pixel in the reference subband, it is chosen the one that minimizes the energy of the high frequency temporal subband.

In Fig. 1.4 the spatial wavelet transform is also illustrated, consisting of two spatial resolution levels. In addition, we present the spatial and spatio-temporal neighbors of a coefficient belonging to the finest spatial resolution level of the finest temporal resolution level. In the group of spatial neighbors we find the coefficients at adjacent spatial locations (which are not illustrated in the Figure, "Up" and "Left"), the coefficients at adjacent orientations ("cousins") and adjacent scales ("parent" and "aunts"). The group of spatio-temporal neighbors consists of coefficients in adjacent scales but in the coarser temporal resolution level ("spatio-temporal parent" and "spatio-temporal aunts"). The role of these neighbors will be explained in the next Chapter.

1.5 Overview of statistical wavelet models for video

1.5.1 Marginal distribution of the spatio-temporal wavelets coefficients

In the video domain, as for still images, various approximations of the marginal distribution of the spatio-temporal coefficients were proposed. However, the models that have sufficiently been used, are based on the Generalized Gaussian distribution and Gaussian mixture distribution.

Lets consider a **Gaussian mixture distribution** :

$$w(n, m) \sim (1 - \epsilon) \mathcal{N}(0, \sigma_0^2) + \epsilon \mathcal{N}(0, \sigma_1^2)$$

with $\epsilon \in [0, 1]$.

Meaning that a hidden binary random variable q is supposed such that :

$$p(q = 1) = \epsilon, p(q = 0) = 1 - \epsilon \text{ and } p(w/q = 0) \sim \mathcal{N}(0, \sigma_0^2), p(w/q = 1) \sim \mathcal{N}(0, \sigma_1^2)$$

Several approaches could be envisaged for the parameter estimation :

1. an Estimation- Maximization (EM) algorithm ;
2. a cumulant-band method : the cumulants of the observed image are matched to their theoretical expressions. Usually, the fourth and the sixth order cumulants are

exploited in this method, using the expression :

$$cum_4 w = 3\epsilon(1 - \epsilon)(\sigma_1^2 - \sigma_0^2)^2$$

$$cum_6 w = 15\epsilon(1 - \epsilon)(1 - 2\epsilon)(\sigma_1^2 - \sigma_0^2)^3$$

together with the expression of the variance :

$$\sigma_w^2 = (1 - \epsilon)\sigma_0^2 + \epsilon\sigma_1^2$$

3. a moment - based method. This is based on the calculation of the p -th order absolute mean values of $w(m, n)$:

$$E\{|w(n, m)|^p\} = C_p[(1 - \epsilon)\sigma_0^p + \epsilon\sigma_1^p]$$

where C_p is a constant depending on p . Closed form expressions of the estimated parameters are obtained from first, second and third absolute mean values of the estimated driving noise. This method was preferred, since it is known to be more robust than the above ones when estimating parameters from real data. We have :

$$\mu_{p,m} = E\{|w(n, m)|^p\} = \int |w|^p p(w) dw = (1 - \epsilon) \int |w|^p p(w/q = 0) dw + \epsilon \int |w|^p p(w/q = 1) dw$$

$$\mu_{p,m} = (1 - \epsilon)\mu_p(\sigma_0^2) + \epsilon\mu_p(\sigma_1^2)$$

where $\mu_p(\sigma^2)$ denotes the p -th order moment of a Gaussian distribution. It is well known that this moment only depends on the variance σ^2 , $\mu_p(\sigma^2) = C_p\sigma^p$, and the constant C_p can easily be computed, i.e. for $p = 1$ and $p = 3$: ($C_2 = 1$) $C_1 = \sqrt{\frac{2}{\pi}}$, $C_3 = \frac{2\sqrt{2}}{\sqrt{\pi}}$.

if we denote by $A = \frac{\mu_{1,w}}{\sqrt{\frac{2}{\pi}}}$, $B = \mu_{2,w}$, $C = \frac{\mu_{3,w}}{\sqrt{\frac{2}{\pi}}}$, then the following relations can be

exploited to compute σ_0 and σ_1 :

$$\sigma_0 + \sigma_1 = \frac{C - AB}{B - A^2} \text{ and } \sigma_0\sigma_1 = \frac{AC - B^2}{B - A^2}$$

The mixing parameter ϵ is then found from the estimated variance, i.e. :

$$\sigma_w^2 = B = (1 - \epsilon)\sigma_0^2 + \epsilon\sigma_1^2$$

Now, let's consider a **Generalized Gaussian distribution** (GGD) of zero mean which is given explicitly by :

$$p(x) = a \exp(-|bx|^\gamma)$$

with

$$a = \frac{b\gamma}{2\Gamma(\frac{1}{\gamma})} \text{ and } b = \frac{1}{\sigma} \sqrt{\frac{\Gamma(\frac{3}{\gamma})}{\Gamma(\frac{1}{\gamma})}}$$

where σ^2 and γ are the variance and the shape parameter respectively. $\Gamma(\cdot)$ is the usual Gamma function. The shape parameter γ in the GGD determines the decay rate of the density function. Note that $\gamma = 2$ yields the Gaussian density function and $\gamma = 1$ the Laplacian one. The smaller values of shape parameter γ correspond to more peaked distributions.

In Table 1.1 and Table 1.2 we show the estimated parameters when we approximate a subband by Gaussian mixture and GGD respectively. By the approximation as GGD, we notice, clearly, that the subbands in a $t + 2D$ decomposition with or without motion estimation, have a shape parameter that is smaller enough from the one of the Gaussian or Laplacian distribution.

<i>With motion estimation</i>					<i>Without motion estimation</i>			
	<i>Temp1</i>		<i>Temp3</i>		<i>Temp1</i>		<i>Temp3</i>	
	SPI	SPII	SPI	SPII	SPI	SPII	SPI	SPII
ϵ	0.0344	0.0362	0.0752	0.0730	0.1085	0.0681	0.1029	0.1790
σ_0	1.6177	1.6403	3.4539	4.6980	1.9631	2.8145	3.7721	6.3080
σ_1	11.7733	13.379	25.7711	37.4821	12.5025	23.6552	27.4019	44.2373

TAB. 1.1 – Estimation of the parameters of a vertical subband in different spatial and temporal resolution levels modeled by a Gaussian mixture.

<i>With motion estimation</i>					<i>Without motion estimation</i>			
	<i>Temp1</i>		<i>Temp3</i>		<i>Temp1</i>		<i>Temp3</i>	
	SPI	SPII	SPI	SPII	SPI	SPII	SPI	SPII
γ	0.77	0.61	0.59	0.56	0.70	0.50	0.54	0.50
σ	0.98	0.83	1.24	1.25	0.82	0.55	1.25	1.71

TAB. 1.2 – Estimation of the parameters of a vertical subband in different spatial and temporal resolution levels modeled by a GGD.

In Fig. 1.5 and Fig. 1.6 we show the original histogram of a spatio-temporal subband and its approximation by a Gaussian mixture and a GGD respectively. We notice that the GGD cannot really approximate the extended tails of the distribution of the subband and that the mixture of Gaussian distribution cannot really approximate the sharp peak around the origin. This motivates us to look for more accurate models.

1.5.2 Extension of models for still images to the video domain

Generally, the statistical models used in the video domain represent extensions of the models used for still images. In this Section we present, briefly, some of them.

In [45], we find a 3-D extension of the set partitioning in hierarchical trees (SPIHT) algorithm. It consists of a partial ordering by magnitude of the 3-D wavelet transformed video with 3-D set partitioning algorithm, an ordering bit plane transmission of refinement

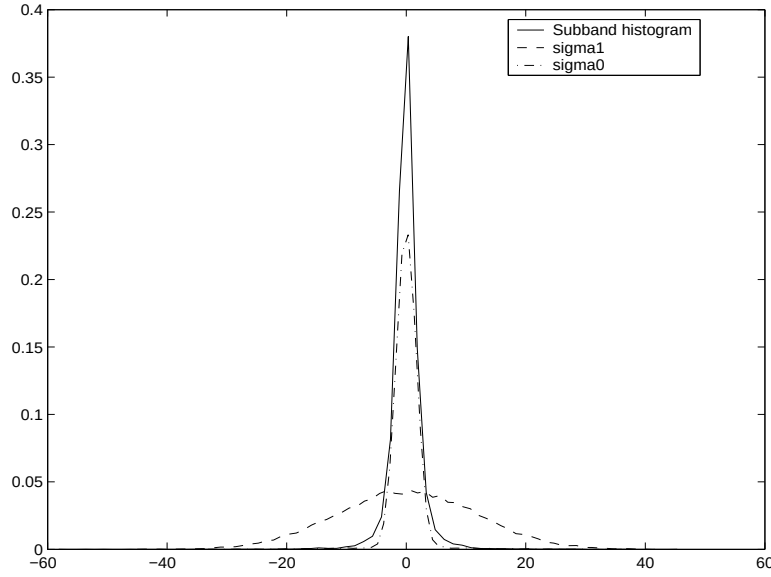


FIG. 1.5 – Approximation by a Gaussian mixture of the vertical subband in the first spatial and temporal resolution level.

bits and an exploitation of self similarity across spatio-temporal orientation trees. In this case, a node is a block with eight adjacent pixels, two extending to each of the three dimensions.

In [16], the probability distribution of two adjacent pixels in subband signals are modeled by bivariate generalized Gaussian distribution (GGD). They based on the observation that even if the subband decomposition generates statistically uncorrelated subbands, within each subband adjacent pixels are still correlated and dependent with each other. Using multivariate *pdf* models, the nonlinear higher order statistical relationships between the signal elements can be exploited but there exist the trade-off between the number of subbands to decompose and the dimension of the pdf model to use. They found that bivariate pdf models are suitable for subband decomposition used for scalar quantization. The orientation (horizontal/vertical) of the adjacent pixel pairs is determined from the subband they reside in. The default option is horizontal pixel. But for subbands with high horizontal, low vertical spatial frequencies they choose vertically adjacent pairs because they are more correlated.

In [55], a multi-modal Laplacian distribution is used to model the spectral distribution of a high frequency subband and a Gibbs random field is employed to model the spatial constraints. They accept that the coefficient distribution of a high frequency subband is approximately Laplacian, and they show that the composite distribution of multiple Laplacians appears consistent with the overall Laplacian distribution, especially in the tail parts. Then, they use a Gibbs random field in order to enforce the neighborhood constraints which means that a pixel is quantized according to not only its own value, but more importantly, the neighborhood spatial constraints. The parametrization of the Gibbs random field is tuned according to the resolution level and the preferential direction of each individual subband. The preferential direction is the direction along which the structures

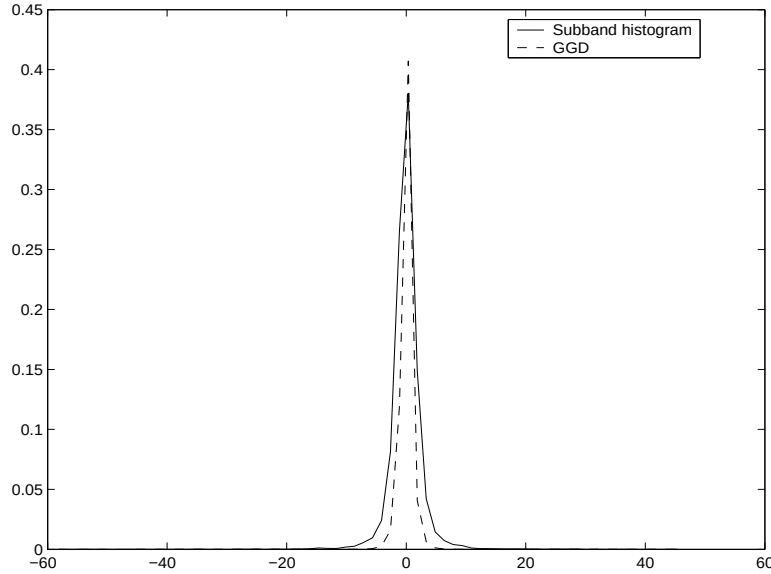


FIG. 1.6 – Approximation by a Generalized Gaussian of the vertical subband in the first spatial and temporal resolution level.

are aligned and is perpendicular to the filtering direction.

A recent different approach is considered in [63]. In order to extend the SPIHT [74] to scan vectors of wavelet coefficients and to use refinement vector quantization techniques, they propose models for describing the statistics of adjacent coefficients within the same subband. Coefficients in 2×1 windows in each subband are grouped together to form vectors of dimension 2, which are subsequently classified into nine classes according to octavely decreasing vector magnitudes thresholds. They observed that the higher magnitude wavelet vectors are heavily oriented along the $(1, 1)$ axis which is not the case for the lower magnitude vectors where the distribution is uniform for all orientations at the same magnitude. Along the same orientation the vector distribution becomes sparse as the magnitude increases. They deduced that the distribution of 2-D vectors of neighboring wavelets are such that the contours of equal probability density are concentric circles at the lower magnitudes, but become increasingly elongated along the $(1, 1)$ orientation, as the magnitude increases. As larger and larger contours are traversed away from the origin, the probability density decreases. In their first model, they assume that the constant probability surfaces are mathematical ellipsoids, with eccentricity increasing outward from the origin and in the second model the constant probability surfaces are elongated in shape but are not mathematical ellipsoids, because the eccentricity varies along the surface.

1.6 Conclusion

We presented various models used in still images and in the video domain, in order to capture the dependencies between the wavelet coefficients. We notice the non-Gaussian marginal *pdf* of the subband wavelet coefficient which is a common characteristic for still

images and video decomposition. For this reason and based on observations via histograms, we detailed the model that we choose to extend to the video domain. This model considers the density of a coefficient conditionally Gaussian with variance a linear function of squared coefficients in a local neighborhood. The neighborhood, in the model for still images, consists of coefficients in adjacent spatial locations, orientations and scales of the current one.

In the next chapter we show the adaptation/extension of this model to the video domain and the observations that lead us to choose it.

Chapitre 2

Spatio-temporal modelling of the wavelet coefficients in a $t + 2D$ scheme

In this Chapter, we explore the statistical dependencies between the wavelet coefficients in a $t + 2D$ scheme of decomposition of a video sequence, with and without motion estimation. We propose a doubly stochastic model to capture these dependencies and a new robust estimator of the parameters of our model.

2.1 Introduction

Among the statistical models that we have already presented in the previous chapter, we pay a special attention to the model of Simoncelli. The basic motivation is that in a $t + 2D$ scheme of decomposition, the temporal analysis leads to have more confidence to a composite model than to an intrascale or interscale model in order to capture all kind of dependencies. For this $t + 2D$ decomposition, shown in Fig. 1.4, an extended spatio-temporal neighborhood can be considered. In addition to the spatial neighbors, we take into account additional dependencies with the spatio-temporal parent, its neighbors and the spatio-temporal “aunts” (see Fig. 1.4).

2.2 Conditional histograms of wavelet coefficients in a $t + 2D$ scheme

In order to precise the model, let us consider a spatio-temporal subband and let $(c_{n,m})_{1 \leq n \leq N, 1 \leq m \leq M}$ be the NM coefficients in this subband. For a given coefficient $c_{n,m}$, we denote by $Q_k(n, m)$ its “neighbors” (k being the index over the considered set of neighbors). Similarly to Chapter 1, let the prediction of $a_{n,m} = |c_{n,m}|^2$ be :

$$l_{n,m} = \sum_k w_k |Q_k(n, m)|^2, \quad (2.1)$$

where $\mathbf{w} = (w_k)_k$ is the vector of weights which are chosen to be least-square optimal.

The high-order statistical dependence involved by this relation can be illustrated *via* conditional histograms of coefficient magnitudes. In order to get a full image of the dependencies existing in a $t + 2D$ scheme, different combinations of the neighborhood of the current coefficient are considered. In addition, the current coefficient is examined in various temporal resolution levels. Thus, for a given coefficient the histograms are supposed conditioned on :

1. only its spatial parent.
2. a neighborhood consisting only of spatial neighbors : adjacent spatial coefficients (Up and Left), adjacent coefficients in spatial orientations ("cousins"), its spatial parent and adjacent coefficients in spatial orientations and scales ("aunts").
3. only its spatio-temporal parent : adjacent coefficient in scale at the coarser temporal resolution level.
4. a neighborhood consting only of spatio-temporal neighbors : its spatio-temporal parent, spatio-temporal parent's spatial neighbors (Up and Left) and adjacent coefficients in spatial orientations and scales at the coarser temporal resolution level ("spatio-temporal aunts")
5. a neighborhood consisting of the spatial and the spatio-temporal coefficients mentioned above.

In Fig. 2.1 and Fig. 2.2 we present the conditional histograms in log-log scales, conditioned to a mean square linear prediction of squared spatio or/and spatio-temporal neighbors, for coefficients in spatio-temporal subbands, with and without motion estimation respectively.

These conditional histograms are strongly related to the conditional histograms of Simoncelli for still images. We can remark that the right side of the distribution is unimodal and concentrated along a unit-slope line and that the left side of the distribution is concentrated about an horizontal line, as for still images. We observe that conditioning on both spatial and spatio-temporal neighbors enforces a better appearence of the linear part and a better concentration of the coefficients along the horizontal line. Especially, in the case without motion estimation, we remark that the linearity of the right part of the histogram conditioning on both spatial and spatio-temporal neighbors is more clear than in the case with motion estimation.

In Fig. 2.3 and Fig. 2.4 we present the conditional histograms on both spatial and spatio-temporal neighbors for different temporal resolution levels, with and without motion estimation respectively.

Based on the observations presented above, in the next section, we present our doubly stochastic model used in order to capture sufficiently the dependencies between the wavelets in a motion-compensated $t + 2D$ scheme.

2.3 Double Stochastic Model

From now on, we call "neighbors" or "spatio-temporal neighbors" both types of spatial and spatio-temporal neighbors.

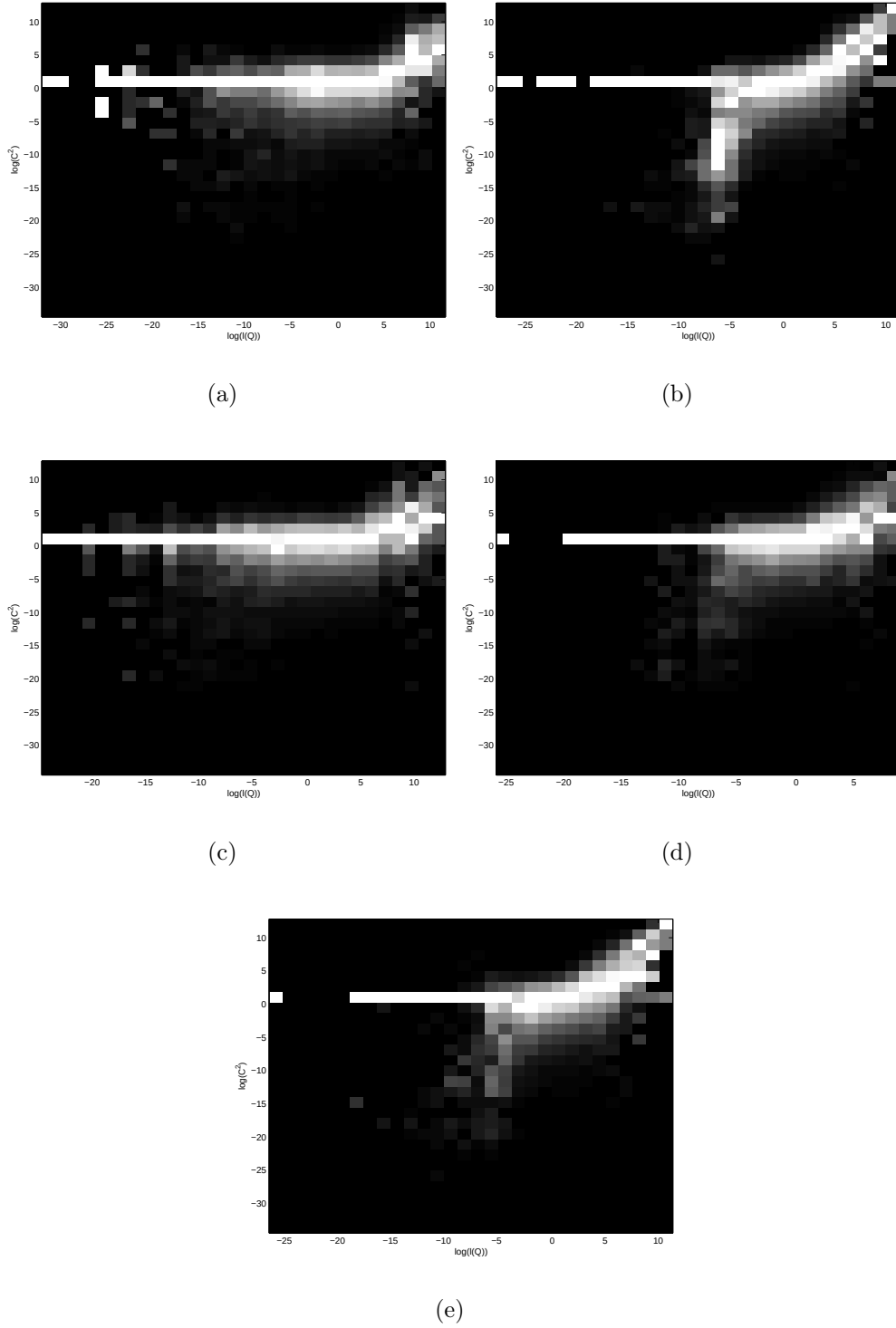


FIG. 2.1 – With motion estimation : Conditional histograms of coefficients in the vertical subband at the first temporal and first spatial resolution level. 2.1(a) : Conditioned on the spatial parent. 2.1(b) : Conditioned on the spatial neighbors. 2.1(c) : Conditioned on the spatio-temporal parent. 2.1(d) : conditioned only on the spatio-temporal neighbors. 2.1(e) : Conditioned on the spatial and spatio-temporal neighbors.

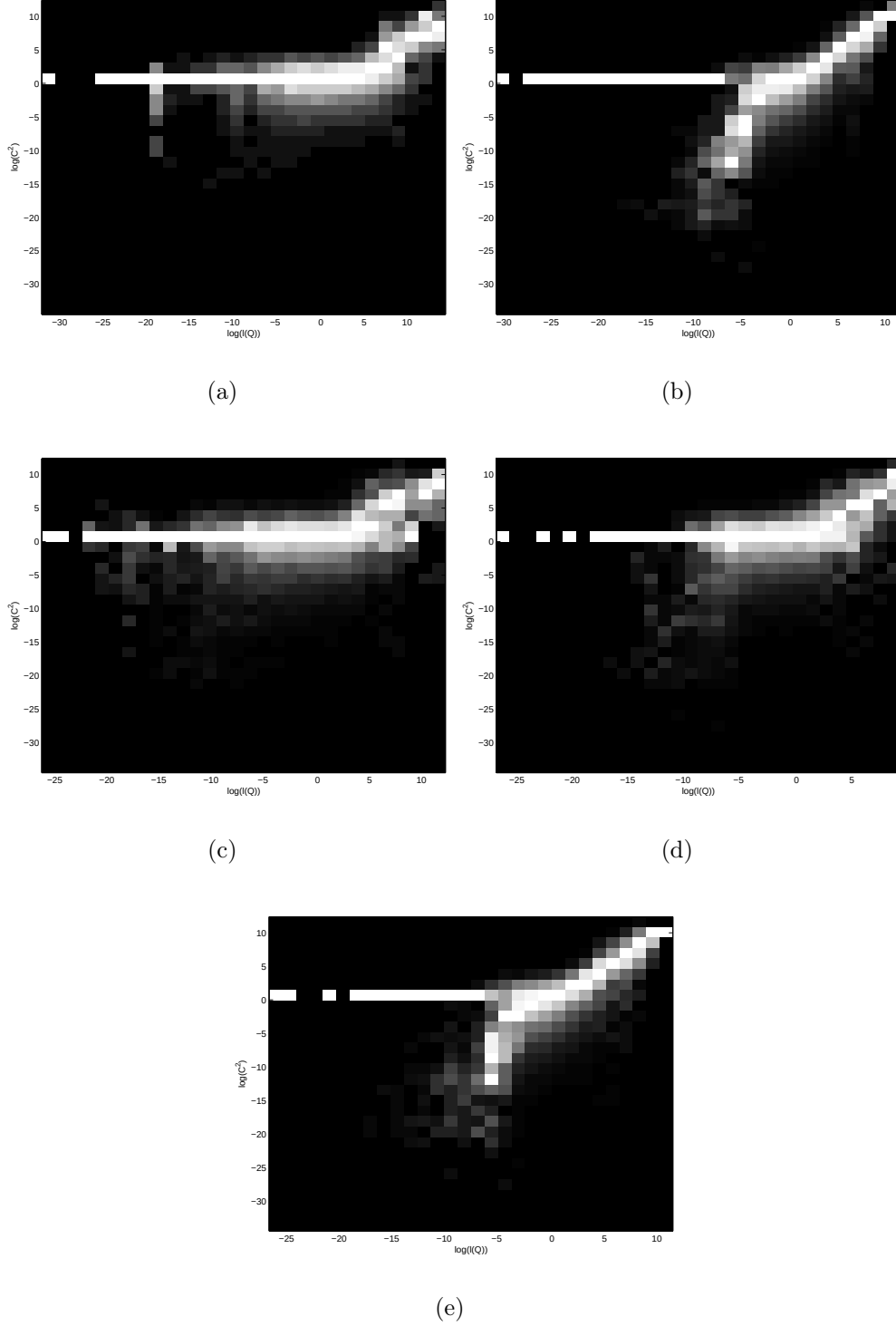


FIG. 2.2 – Without motion estimation : Conditional histograms of coefficients in the vertical subband at the first temporal and first spatial resolution level. 2.2(a) : Conditioned on the spatial parent. 2.2(b) : Conditioned on the spatial neighbors. 2.2(c) : Conditioned on the spatio-temporal parent. 2.2(d) : conditioned only on the spatio-temporal neighbors. 2.2(e) : Conditioned on the spatial and spatio-temporal neighbors.

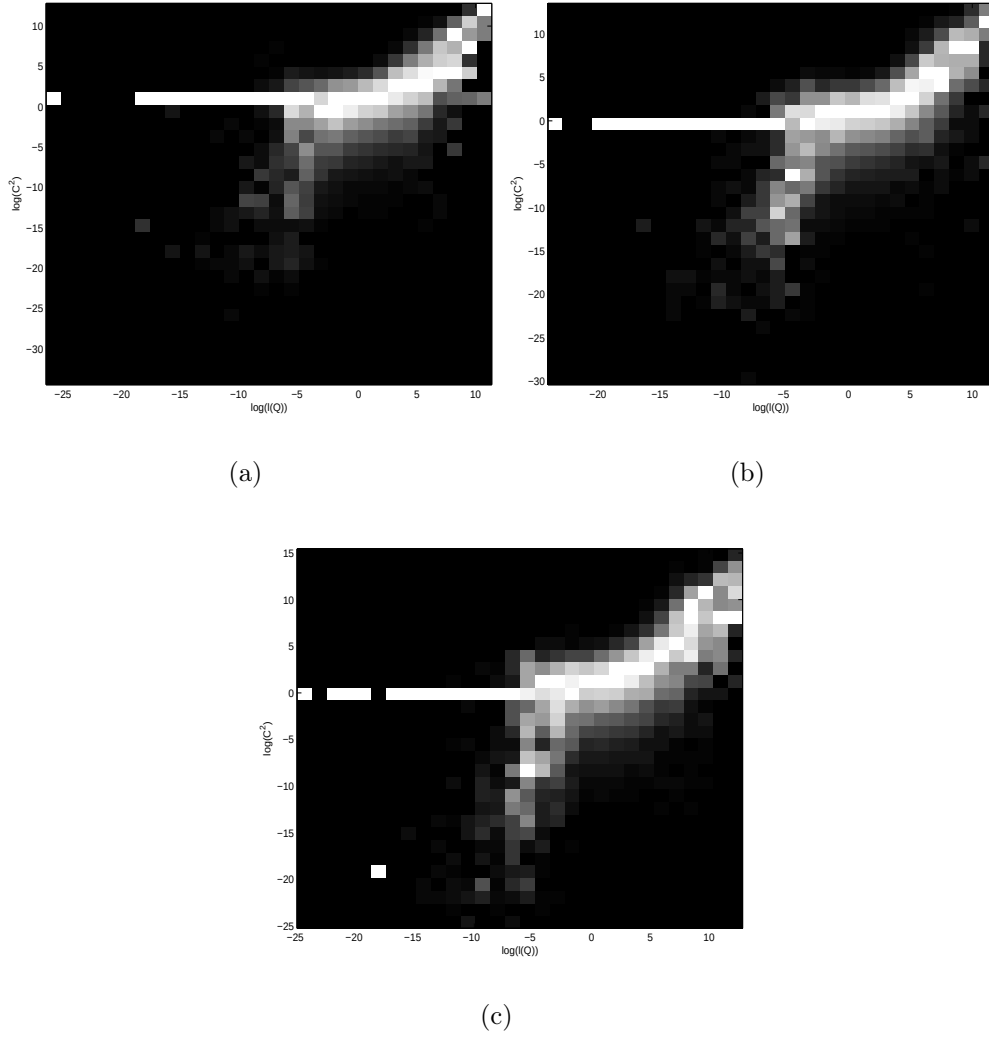


FIG. 2.3 – With motion estimation : Conditional histograms of coefficients in the vertical subband at the first spatial resolution level on both spatial and spatio-temporal neighbors. 2.3(a) : First temporal resolution level. 2.3(b) : Second temporal resolution level. 2.3(c) : Third temporal resolution level.

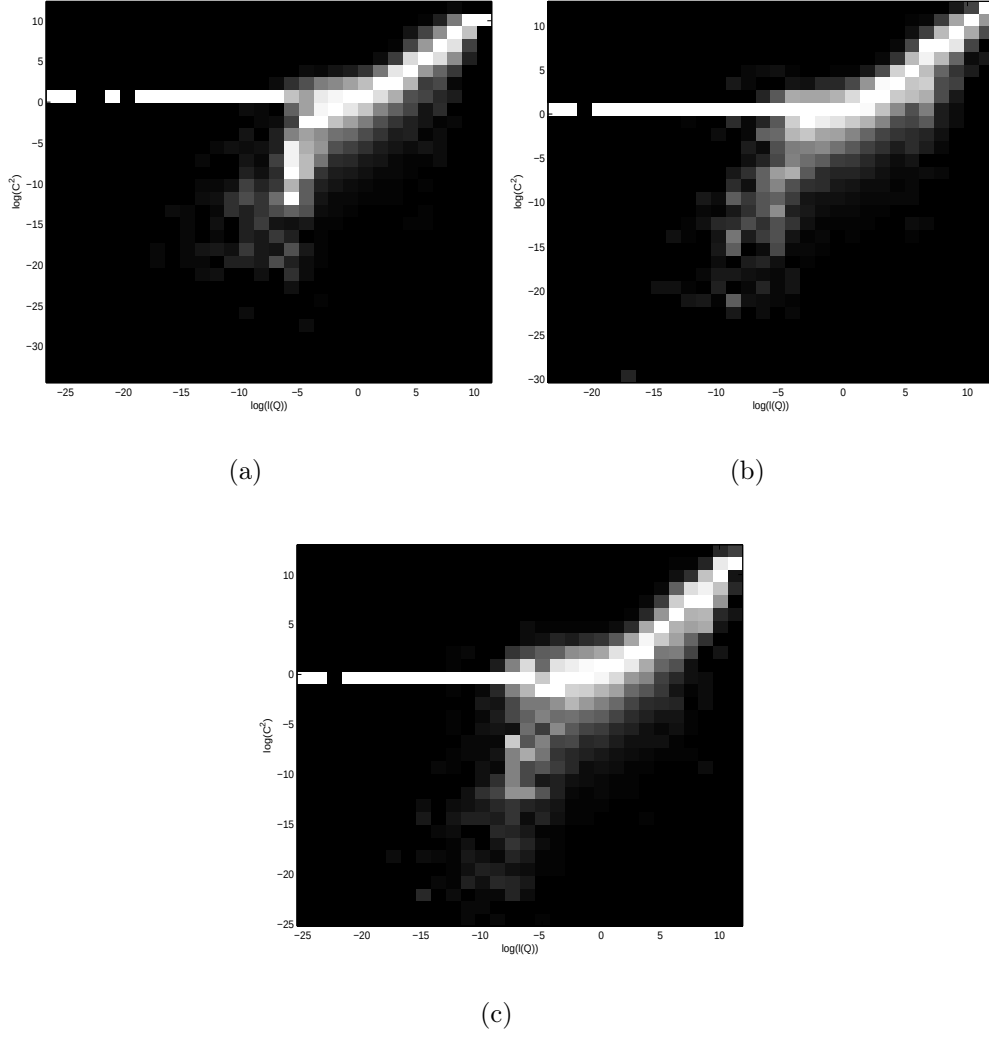


FIG. 2.4 – Without motion estimation : Conditional histograms of coefficients in the vertical subband at the first spatial resolution level on both spatial and spatio-temporal neighbors. 2.4(a) : First temporal resolution level. 2.4(b) : Second temporal resolution level. 2.4(c) : Third temporal resolution level.

Let $c_{n,m}$ be the current coefficient and $Q_k(n, m)$ its “neighbors” and let the prediction of

$a_{n,m} = |c_{n,m}|^2$ be :

$$l_{n,m} = \sum_k w_k |Q_k(n, m)|^2, \quad (2.2)$$

where $\mathbf{w} = (w_k)_k$ is the vector of weights.

Fig. 2.1 suggests us to consider the following model :

$$\log a_{n,m} = \log(l_{n,m} + \alpha) + z_{n,m}, \quad (2.3)$$

where $z_{n,m}$ is an additive noise. When $l_{n,m}$ takes large values, the dependence between $\log a_{n,m}$ and $\log l_{n,m}$ is approximately linear, which is in agreement with the right part of the plot in Fig. 2.1 and Fig. 2.2. In the meantime, the constant α is useful to describe the flat left part of the log-histogram.

This model amounts to :

$$|c_{n,m}| = (l_{n,m} + \alpha)^{1/2} e^{z_{n,m}/2} \quad (2.4)$$

and by reintroducing the sign, we have :

$$c_{n,m} = (l_{n,m} + \alpha)^{1/2} e^{z_{n,m}/2} s_{n,m}, \quad (2.5)$$

where $s_{n,m} \in \{-1, 1\}$. We suppose the noise to be normal, that is $\beta_{n,m} = e^{z_{n,m}/2} s_{n,m} \sim \mathcal{N}(0, 1)$. Then, $g(c_{n,m}|l_{n,m}) \sim \mathcal{N}(0, l_{n,m} + \alpha)$ which leads to a Gaussian conditional distribution for the spatio-temporal coefficients of the form :

$$g(c_{n,m}|\sigma_{n,m}^2) = \frac{1}{\sqrt{2\pi}\sigma_{n,m}} e^{-\frac{c_{n,m}^2}{2\sigma_{n,m}^2}} \quad (2.6)$$

where

$$\sigma_{n,m}^2 = \sum_k w_k |Q_k(n, m)|^2 + \alpha \quad (2.7)$$

and $\mathbf{Q}(n, m) = (Q_k(n, m))_k$ is the vector of neighbors.

Observing the Fig. 2.1 and Fig. 2.2, one can remark the increase of the variance of the model with the conditioning value which leads to this **double stochastic model**, in which we consider the conditional probability law of the coefficients in a given subband to be Gaussian, with variance depending on the set of the spatio-temporal neighbors.

2.4 Model Estimation

In order to estimate the parameters

$$\boldsymbol{\theta} = \begin{pmatrix} \mathbf{w} \\ \alpha \end{pmatrix}$$

of the model, we use the wavelet coefficients $(\mathbf{c}_{n,m})_{(1 \leq n \leq N, 1 \leq m \leq M)}$ (where N, M represent the image size) to build several criteria and compare their estimation performances. Ideally, a criterion $J_{N,M}(\boldsymbol{\theta})$ should satisfy some nice properties, such as :

1. A parameter estimator should be such that $\hat{\boldsymbol{\theta}}_{N,M} = \arg \min_{\boldsymbol{\theta}} J_{N,M}(\boldsymbol{\theta})$;
2. $J_{N,M}(\boldsymbol{\theta}) \longrightarrow J(\boldsymbol{\theta})$ when $N, M \rightarrow \infty$, the convergence being almost sure (or, at least, in probability);
3. $J(\boldsymbol{\theta}) \geq J(\boldsymbol{\theta}_0)$, with $\boldsymbol{\theta}_0$ being the vector of the true parameters;

These conditions define what is called a “contrast” in statistics [36]. However, they may be difficult to satisfy in practice, and one can therefore require slightly weaker constraints to be satisfied. In the sequel, we shall check whether the following alternative two constraints are satisfied by the proposed criteria :

$$\mathbb{E} \{ J_{N,M}(\boldsymbol{\theta}) \} \geq \mathbb{E} \{ J_{N,M}(\boldsymbol{\theta}_0) \} \quad (2.8)$$

$$J_{N,M}(\boldsymbol{\theta}_0) \longrightarrow J(\boldsymbol{\theta}_0) \quad \text{in probability, when } N, M \rightarrow \infty. \quad (2.9)$$

In the above equation, $E\{\cdot\}$ denotes the mathematical expectation. Let us now introduce the criteria and discuss their properties with respect to the above constraints.

2.4.1 Least squares (LS)

The criterion proposed in [85] is a least mean squares one, which can be written as :

$$J_{N,M}(\boldsymbol{\theta}) = \frac{1}{NM} \sum_{n=1}^N \sum_{m=1}^M (c_{n,m}^2 - \sigma_{n,m}^2(\boldsymbol{\theta}))^2. \quad (2.10)$$

$$J_{N,M}(\boldsymbol{\theta}) = \frac{1}{NM} \sum_{n=1}^N \sum_{m=1}^M (c_{n,m}^4 - 2 \cdot c_{n,m}^2 \cdot \sigma_{n,m}^2(\boldsymbol{\theta}) + \sigma_{n,m}^4(\boldsymbol{\theta})). \quad (2.11)$$

In order to examine the behavior of this criterion to the above constraints presented in Eq.(2.8) and Eq.(2.9), we follow the next steps :

$$\mathbb{E}\{J_{N,M}(\boldsymbol{\theta})\} = \frac{1}{NM} \sum_{n=1}^N \sum_{m=1}^M \mathbb{E}\{c_{n,m}^4 - 2 \cdot c_{n,m}^2 \cdot \sigma_{n,m}^2(\boldsymbol{\theta}) + \sigma_{n,m}^4(\boldsymbol{\theta})\}. \quad (2.12)$$

$$\mathbb{E}\{J_{N,M}(\boldsymbol{\theta})\} = \frac{1}{NM} \sum_{n=1}^N \sum_{m=1}^M \mathbb{E}_{\mathbf{Q}_{n,m}}\{\mathbb{E}\{c_{n,m}^4 - c_{n,m}^2 \cdot \sigma_{n,m}^2(\boldsymbol{\theta}) + \sigma_{n,m}^4(\boldsymbol{\theta})\} / \mathbf{Q}_{n,m}\}. \quad (2.13)$$

where $\mathbf{Q}_{n,m}$ is the vector of neighbors of the coefficient $c_{n,m}$. Thus,

$$\mathbb{E}\{J_{N,M}(\boldsymbol{\theta})\} = \frac{1}{NM} \sum_{n=1}^N \sum_{m=1}^M \mathbb{E}_{\mathbf{Q}_{n,m}}\{3 \cdot \sigma_{n,m}^4(\boldsymbol{\theta}_0) - \sigma_{n,m}^2(\boldsymbol{\theta}_0) \cdot \sigma_{n,m}^2(\boldsymbol{\theta}) + \sigma_{n,m}^4(\boldsymbol{\theta})\}. \quad (2.14)$$

$$\mathbb{E}\{J_{N,M}(\boldsymbol{\theta})\} = \frac{1}{NM} \sum_{n=1}^N \sum_{m=1}^M \mathbb{E}_{\mathbf{Q}_{n,m}}\{[\sigma_{n,m}^2(\boldsymbol{\theta}) - \sigma_{n,m}^2(\boldsymbol{\theta}_0)]^2 + 2 \cdot \sigma_{n,m}^4(\boldsymbol{\theta}_0)\} \quad (2.15)$$

$$\mathbb{E}\{J_{N,M}(\boldsymbol{\theta})\} \geq \frac{1}{NM} \sum_{n=1}^N \sum_{m=1}^M \mathbb{E}_{\mathbf{Q}_{n,m}} \{2 \cdot \sigma_{n,m}^4(\boldsymbol{\theta}_0)\} = \frac{1}{NM} \sum_{n=1}^N \sum_{m=1}^M \mathbb{E}\{J_{N,M}(\boldsymbol{\theta}_0)\} \quad (2.16)$$

By the Eq.(2.16) we see that the constraint of the Eq.(2.8) is satisfied for this criterion. On the contrary, the constraint described by the Eq. 2.9 holds only subject to some additional ergodicity conditions on $c_{n,m}^2 - \sigma_{n,m}^2(\boldsymbol{\theta}_0)$.

2.4.2 Maximum-likelihood (ML)

We propose the use of an approximate maximum likelihood estimator :

$$\hat{\boldsymbol{\theta}} = \arg \max_{\boldsymbol{\theta}} \prod_{n=1}^N \prod_{m=1}^M g(c_{n,m} | \sigma_{n,m}^2(\boldsymbol{\theta})).$$

This is equivalent to minimize the following criterion :

$$J_{N,M}(\boldsymbol{\theta}) = \frac{1}{NM} \sum_{n=1}^N \sum_{m=1}^M \left\{ \frac{c_{n,m}^2}{\sigma_{n,m}^2(\boldsymbol{\theta})} + \log \sigma_{n,m}^2(\boldsymbol{\theta}) \right\}.$$

As above we have :

$$\mathbb{E}\{J_{N,M}(\boldsymbol{\theta})\} = \frac{1}{NM} \sum_{n=1}^N \sum_{m=1}^M \mathbb{E}_{\mathbf{Q}_{n,m}} \left\{ \frac{\sigma_{n,m}^2(\boldsymbol{\theta}_0)}{\sigma_{n,m}^2(\boldsymbol{\theta})} + \log \sigma_{n,m}^2(\boldsymbol{\theta}) \right\}. \quad (2.17)$$

$$\mathbb{E}\{J_{N,M}(\boldsymbol{\theta})\} = \frac{1}{NM} \sum_{n=1}^N \sum_{m=1}^M \mathbb{E}_{\mathbf{Q}_{n,m}} \left\{ \frac{\sigma_{n,m}^2(\boldsymbol{\theta}_0)}{\sigma_{n,m}^2(\boldsymbol{\theta})} + \log \frac{\sigma_{n,m}^2(\boldsymbol{\theta})}{\sigma_{n,m}^2(\boldsymbol{\theta}_0)} + \log \sigma_{n,m}^2(\boldsymbol{\theta}_0) \right\}. \quad (2.18)$$

Setting as $u = \frac{\sigma_{n,m}^2(\boldsymbol{\theta})}{\sigma_{n,m}^2(\boldsymbol{\theta}_0)}$ we examine the behaviour of the following equation :

$$f(u) = \frac{1}{u} + \log u \quad (2.19)$$

which is minimized when $u = 1$ and thus $f(u) \geq 1 \rightarrow \frac{\sigma_{n,m}^2(\boldsymbol{\theta}_0)}{\sigma_{n,m}^2(\boldsymbol{\theta})} + \log \frac{\sigma_{n,m}^2(\boldsymbol{\theta})}{\sigma_{n,m}^2(\boldsymbol{\theta}_0)} \geq 1$

Finally we have :

$$\mathbb{E}\{J_{N,M}(\boldsymbol{\theta})\} \geq \frac{1}{NM} \sum_{n=1}^N \sum_{m=1}^M \{1 + \mathbb{E}_{\mathbf{Q}_{n,m}} \{\log \sigma_{n,m}^2(\boldsymbol{\theta}_0)\} = \mathbb{E}\{J_{N,M}(\boldsymbol{\theta}_0)\} \quad (2.20)$$

Again, it is easy to verify that this criterion satisfies the constraint Eq. 2.8 but the constraint Eq. 2.9 requires ergodicity conditions on $\log \sigma_{n,m}^2(\boldsymbol{\theta}_0)$.

2.4.3 A more Efficient Criterion (EC)

The two previous criteria don't guarantee an asymptotical consistency and thus in order to satisfy also the second constraint Eq. 2.9 we introduce a more efficient criterion (EC) defined by :

$$J_{N,M}(\boldsymbol{\theta}) = \frac{1}{NM} \sum_{n=1}^N \sum_{m=1}^M \left(\frac{\gamma |c_{n,m}|^\beta}{\sigma_{n,m}^\beta(\boldsymbol{\theta})} - 1 \right)^2 \quad (2.21)$$

where γ and β are two real positive parameters.

For a very large number of coefficients ($N, M \rightarrow \infty$), according to the law of large numbers the criterion $J_{N,M}(\boldsymbol{\theta}_0)$ converges in probability to the following expression :

$$J(\boldsymbol{\theta}_0) = \gamma^2 \mathbb{E} \left\{ \frac{|c_{n,m}|^{2\beta}}{\sigma_{n,m}^{2\beta}(\boldsymbol{\theta}_0)} \right\} - 2\gamma \mathbb{E} \left\{ \frac{|c_{n,m}|^\beta}{\sigma_{n,m}^\beta(\boldsymbol{\theta}_0)} \right\} + 1. \quad (2.22)$$

Besides, we have $\mathbb{E}\{|c_{n,m}|^\beta \mid \mathbf{Q}(n, m)\} = C_\beta^c \sigma_{n,m}^\beta(\boldsymbol{\theta}_0)$, where

$$C_\beta^c = 2 \int_0^\infty u^\beta g(u|1) du.$$

The expression (2.21) thus leads to :

$$\mathbb{E}\{J_{N,M}(\boldsymbol{\theta}) \mid (\mathbf{Q}(n, m))_{1 \leq n \leq N, 1 \leq m \leq M}\} = \frac{1}{NM} \sum_{n=1}^N \sum_{m=1}^M \left[\gamma^2 C_{2\beta}^c \frac{\sigma_{n,m}^{2\beta}(\boldsymbol{\theta}_0)}{\sigma_{n,m}^{2\beta}(\boldsymbol{\theta})} - 2\gamma C_\beta^c \frac{\sigma_{n,m}^\beta(\boldsymbol{\theta}_0)}{\sigma_{n,m}^\beta(\boldsymbol{\theta})} + 1 \right] \quad (2.23)$$

The parameter γ should be chosen so as to guarantee that $\mathbb{E}\{J_{N,M}(\boldsymbol{\theta}_0)\} \leq \mathbb{E}\{J_{N,M}(\boldsymbol{\theta})\}$ for all $\boldsymbol{\theta}$, with equality if and only if $\boldsymbol{\theta} = \boldsymbol{\theta}_0$. This condition is satisfied if

$$\boldsymbol{\theta} \mapsto \gamma^2 C_{2\beta}^c \frac{\sigma_{n,m}^{2\beta}(\boldsymbol{\theta}_0)}{\sigma_{n,m}^{2\beta}(\boldsymbol{\theta})} - 2\gamma C_\beta^c \frac{\sigma_{n,m}^\beta(\boldsymbol{\theta}_0)}{\sigma_{n,m}^\beta(\boldsymbol{\theta})} + 1$$

is minimum for $\boldsymbol{\theta} = \boldsymbol{\theta}_0$.

Setting $x = \frac{\sigma_{n,m}^\beta(\boldsymbol{\theta}_0)}{\sigma_{n,m}^\beta(\boldsymbol{\theta})}$ the previous relation becomes :

$$C_{2\beta}^c \gamma^2 x^2 - 2\gamma C_\beta^c x + 1 = C_{2\beta}^c \gamma^2 (x^2 - 2 \frac{C_\beta^c}{C_{2\beta}^c \gamma} x) + 1$$

After some simple calculations, we get :

$$C_{2\beta}^c \gamma^2 (x - \frac{C_\beta^c}{C_{2\beta}^c \gamma})^2 + 1 - \frac{(C_\beta^c)^2}{C_{2\beta}^c}$$

which must be minimized when $x = 1$, and thus $\gamma = \frac{C_\beta^c}{C_{2\beta}^c}$.

We can notice that due to the Gaussian assumption in the particular case $\beta = 2$, we get

$C_2^c = 1$, $C_4^c = 3$. In this case, the criterion in (2.23) is equivalent to a *Modified Least Squares (MLS)* criterion, leading to the minimization :

$$\sum_{n,m} \left(\frac{c_{n,m}^4}{3\sigma_{n,m}(\boldsymbol{\theta})^4} - 2 \frac{c_{n,m}^2}{\sigma_{n,m}(\boldsymbol{\theta})^2} \right).$$

One of the advantages of the third criterion (EC) over the former two (LS, ML) is that no additional ergodicity conditions are required for Eq. 2.9 to be satisfied. In the next section, we provide evidence through Monte Carlo simulations for the improved mean square estimation error achieved by the new criterion.

2.5 Illustration Examples

In order to illustrate the previous theoretical results, we consider two cases of decomposition : with and without motion estimation. In both of them, a temporal Haar decomposition of the video sequence, is applied on groups of 16 frames, with 4 temporal and 4 spatial resolution levels and a spatial multiresolution analysis based on the biorthogonal 9/7 filters. In the case of motion estimation/compensation a lifting-based motion-compensated Haar temporal decomposition is used [68], with a full search block matching algorithm and half-pel motion accuracy.

The spatio-temporal neighborhood consists of 12 coefficients of the current one : its Up and Left neighbors, its spatial parent, aunts and cousins, and its spatio-temporal parent together with its Up and Left neighbors, and spatio-temporal aunts. In order to check the validity of our model, the parameters estimated by least mean squares on a given subband have been used to generate a Gaussian random field having the same conditional probability density as our model.

Based on the synthetic data, the different estimators, presented in the previous section, have been compared and the parameter values estimated over 50 realizations are presented : in Table 2.1 in case without motion estimation and in Table 2.2 in case with motion estimation. A critical point in the estimation is that, in order to keep the variance of the model positive, we need to constrain the weights to be positive.

As we can notice from these tables, in both cases, the EC with $\beta = 2$ proves to be the most robust and with the best performance compared to the LS and ML criteria especially for the neighbors which are more significant. On the other hand, if we want a neighborhood that includes more of spatio-temporal neighbors than spatial, it is recommended to choose the EC with $\beta = 1$. We are going to see in the next Chapter, that the application of our stochastic model to a prediction method suggests the use of EC with $\beta = 1$.

Another example for the case with motion estimation is illustrated in Fig. 2.5. We show the real subband (which is, in this case, the vertical detail subband at the highest spatial resolution of the first temporal decomposition level for "hall-monitor" sequence) and a typical simulated one (with the parameters estimated by MLS criterion).

Neighbors' weights	True parameters	LS	ML	EC($\beta = 2$)	EC($\beta = 1$)
wUp	0.4916	0.1529	0.0326	0.0322	0.0408
wLeft	0.0991	0.0656	0.0414	0.0225	0.0380
wcous1	0.3431	0.3157	0.0708	0.0284	0.0516
wcous2	0.0041	0.0274	0.0053	0.0048	0.0074
wpar	0.0078	0.0093	0.0056	0.0093	0.0069
waunt1	0.0133	0.0124	0.0053	0.0055	0.0046
waunt2	0.0009	0.0037	0.0013	0.0085	0.0040
wpartm	0.0032	0.0046	0.0050	0.0091	0.0043
wLeftpartm	0.0024	0.0030	0.0021	0.0036	0.0021
wUppartm	0.0011	0.0082	0.0022	0.0041	0.0024
waunt1tm	0.0081	0.0138	0.0046	0.0080	0.0041
waunt2tm	0.0010	0.0023	0.0010	0.0050	0.0023
α	0.3508	0.8058	0.1241	0.0649	0.1161

TAB. 2.1 – Without motion estimation : Parameter estimation : the first column indicates the various spatio-temporal neighbors whose weights are estimated (see Fig.1.4). The second column indicates the value of the true parameters, the next four ones the mean square error of the estimation by the four proposed methods over 50 realizations.

Neighbors' weights	True parameters	LS	ML	EC($\beta = 2$)	EC($\beta = 1$)
wUp	0.1008	0.0670	0.0405	0.0315	0.0369
wLeft	0.2940	0.1158	0.0272	0.0266	0.0318
wcous1	0.0389	0.0331	0.0132	0.0249	0.0141
wcous2	0.2121	0.1487	0.0487	0.0326	0.0496
wpar	0.0073	0.0350	0.0070	0.0101	0.0071
waunt1	0.0358	0.0230	0.0070	0.0106	0.0080
waunt2	0.0685	0.0273	0.0048	0.0067	0.0045
wpartm	0.0341	0.0147	0.0042	0.0045	0.0051
wLeftpartm	0.0012	0.0079	0.0022	0.0032	0.0033
wUppartm	0.0054	0.0065	0.0048	0.0056	0.0049
waunt1tm	0.0013	0.0034	0.0024	0.0058	0.0021
waunt2tm	0.0002	0.0092	0.0011	0.0009	0.0008
α	0.3663	0.5121	0.0975	0.0465	0.0774

TAB. 2.2 – With motion estimation : Parameter estimation : the first column indicates the various spatio-temporal neighbors whose weights are estimated (see Fig.1.4). The second column indicates the value of the true parameters, the next four ones the mean square error of the estimation by the four proposed methods over 50 realizations.

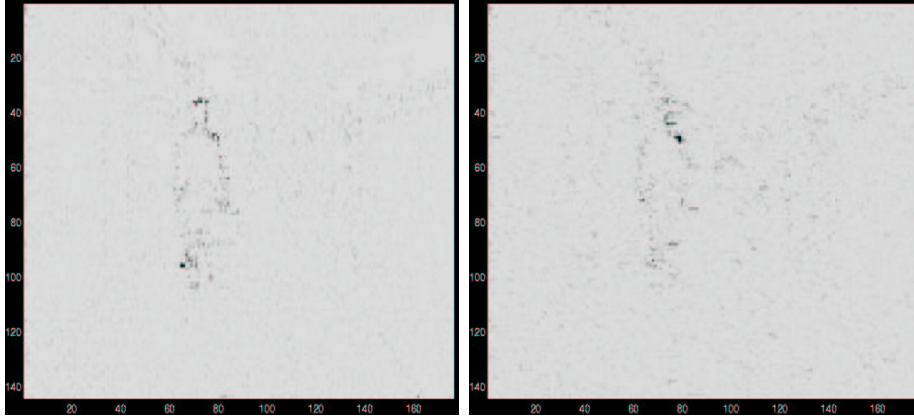


FIG. 2.5 – Left : vertical detail subband at the highest spatial resolution of the first temporal decomposition level for “hall_monitor” sequence. Right : simulated subband, using the conditional law given in Eq. (2.6).

2.6 Conclusion

We present a doubly stochastic model of the wavelet coefficients in a $t + 2D$ decomposition of a video sequence where the conditional probability law of the coefficients is Gaussian with variance that depends on the set of spatio-temporal neighbors. We observed the dependencies between the wavelet coefficients in two cases : with and without motion estimation. In the neighborhood of the model used for still images we added a set of new coefficients which exposes the dependencies across the temporal axis.

We proposed a new general criterion of the parameters of our model which is found more robust than the already proposed least-square criterion. We gave analytical comparative tables of its performance in both cases : with and without motion estimation.

After having established a well-defined model of dependencies between the wavelet coefficients and estimating efficiently its parameters, we will show in the next Chapter, the application of it to a prediction method. We will present the gain of using this prediction method to the quality enhancement of the scalable video when packets loss occurred.

Chapitre 3

Application of the statistical model to error concealment and quality enhancement of video

In this Chapter, we present a prediction method which takes advantage of the well-situated relations between the wavelet coefficients, captured by our stochastic model. We, also, present two applications of it to : the quality improvement of scalable video and the error concealment when packet losses occur during video transmission.

3.1 Introduction

In packet networks without QoS (quality of service), even considering a strong channel protection for the most important parts of the bitstream, some of the packets will be lost during the transmission due to the network congestion or bursts of errors. In this case, an error concealment method should be applied by the decoder in order to improve the quality of the reconstructed sequence.

There exist a plethora of error concealment methods of video, most of them applying directly on the reconstructed sequences (for a comparative review, see [81]). Approaches exploiting the redundancy along the temporal axis try to conceal the corrupted blocks in the current frame by selecting suitable substitute blocks from the previous frames. This approach can be reinforced by introducing data partitioning techniques [90] : data in the error prediction blocks are separated in motion vectors and DCT coefficients, which are unequally protected. This way, if the motion vector data are received without errors, the missing blocks are set to their corresponding motion compensated blocks. However, the loss of a packet usually results in the loss of both the motion vectors and the DCT coefficients. So, many concealment methods first estimate the motion vectors associated with a missing block using the motion vectors of adjacent blocks [32], [47].

Spatial error concealment methods restore the missing blocks only based on the information decoded in the current frame. To restore the missing data, several methods can be used : minimization of a measure of variations (e.g. gradient or Laplacian) between adjacent pixels [93], each pixel in the damaged block is interpolated from the corresponding

pixels in its four neighboring blocks such that the total squared border error is minimized [41], or the missing information is interpolated utilizing spatially correlated edge information from a large local neighborhood [89]. Statistical models like the Markov random fields (MRF) have also been proposed for error concealing in video [75], [82]. These methods estimate the missing pixels by exploiting spatial or spatio-temporal constraints between pixels in the original sequence. Note that such approaches can also be employed to estimate missing motion vectors [98].

Our error concealment method proposed is based on the statistical model presented in Chapter 2, applied on the transformed wavelet domain. As we have pointed out it is a spatio-temporal multiscale model, exhibiting the correlation between discontinuities at different resolution levels in the error prediction (temporal detail) frames.

3.2 Prediction Method

The stochastic model presented in the previous Chapter can be applied to the prediction of the subbands that are not received by the decoder. Indeed, a spatio-temporal multi-resolution analysis (MRA) as described in Chapter 2 naturally provides a hierarchical subband structure, allowing to transmit information by decreasing order of importance. The decoder receives, therefore, the coarser spatio-temporal resolution levels first and then with the help of the spatio-temporal neighbors, can predict the finest resolution ones. The conditional law of the coefficients exhibited in Eq. 2.6 is used to build an optimal mean square error estimator of the magnitude of each coefficient, given its spatio-temporal ancestors. This leads to the following predictor :

$$|\hat{c}_{n,m}| = E\{|c_{n,m}| \mid \mathbf{Q}(n,m)\} = \int_{-\infty}^{\infty} |c_{n,m}| g(c_{n,m} \mid \sigma_{n,m}^2) dc_{n,m}. \quad (3.1)$$

After some simple calculations, we get the optimal estimator expression :

$$|\hat{c}_{n,m}| = \sqrt{\frac{2}{\pi}} \sigma_{n,m}, \quad (3.2)$$

with $\sigma_{n,m}$ given in Eq. 2.7 and the model parameters estimated using the criterion in Eq. 2.22.

The choice of the spatio-temporal neighbors used by the predictor, in the context of a scalable bitstream, has been made in such a way to avoid error propagation. Supposing the coarser spatial level of each frame is received (for example, it can be better protected against channel errors), we restrict the choice of the coefficients $Q_k(n,m)$ in our model to the spatial parent, spatial aunts and the spatio-temporal parent, its neighbors and the spatio-temporal aunts of the current coefficient. As the bitstream is resolution scalable, all these spatio-temporal ancestors belong to the spatio-temporal subbands that have already been received by the decoder and can therefore be used in a causal prediction.

Note that our statistical model and therefore the proposed prediction do not concern the sign of the coefficients. As the sign of the coefficients remains an important information, data partitioning can be used to separate it from the magnitude of the coefficients, in order to better protect it in the video bitstream. Efficient algorithms for encoding the

sign of wavelet coefficients are already available (see, for example, [18]). In the sequel, we shall consider therefore that the sign has been correctly decoded.

3.3 Model-Based Quality Enhancement of Scalable Video

In the first application, we consider scalable video transmission over heterogeneous networks and we are interested in improving the spatial scalability properties. In this case, the adaptation of the bitstream to the available bandwidth can lead to discarding the finest spatial detail subbands during the transmission. However, if the decoder has display size and CPU capacity to decode in full resolution, the lack of the finest frequency details would result in a low quality, oversmoothed, reconstructed sequence. We propose to use the stochastic model developed in the previous Chapter to improve the rendering of the spatio-temporal details in the reconstructed sequence. Thus, the decoder will receive the coarser spatial resolution levels at each temporal level and predict with the help of the spatio-temporal neighbors the finest resolution ones. We propose to use for the prediction the optimal mean square error estimator of the magnitude of each coefficient, given its spatio-temporal ancestors presented in Section 3.2.

Note that this strategy can also be seen as a quality scalability, since bitrate reduction is achieved by not transmitting the finest frequency details.

In order to apply this method, as we can remark from the Table 2.2, it is more convenient to use the EC Criterion with $\beta = 1$. Its performance in the considered neighborhood is slightly better than for the same criterion with $\beta = 2$.

For simulations, we have considered the spatio-temporal neighborhood consisting of the 8 coefficients mentioned before. We send the three low-resolution spatial levels of each temporal detail frame and predict the highest resolution detail subbands using our model. We compare this procedure with the reconstruction of the full resolution using the finest spatial detail subbands set to zero, which would be the reconstruction strategy of a simpler decoder.

In Fig. 3.1 we present the Mean Square Error (MSE) of the spatial reconstruction of each temporal detail frame at different temporal resolution levels. One can remark the significant decrease in reconstruction error by using the proposed prediction strategy. Another observation is related to the MSE value in itself, which is highest at the last temporal resolution level. This is related to the higher energy of the low-resolution temporal detail subbands.

In Fig. 3.2 we present the PSNR improvement of the reconstructed sequence obtained by predicting the finest frequency subbands at all the temporal resolution levels with our model, instead of setting them to zero. As we can see, for two different sequences the PSNR improvement varies between 1.3 dB and 2.7 dB.

In Fig. 3.3, we present the reconstructed temporal detail frames of the first temporal resolution level of the “hall-monitor” sequence. The first one is the real reconstructed temporal detail frame, the second one is the reconstructed frame when we predict the finest subbands and the third one is the reconstructed frame when we set its finest subbands to zero. As we can see, the later image proved to be more blurred than the real one and the

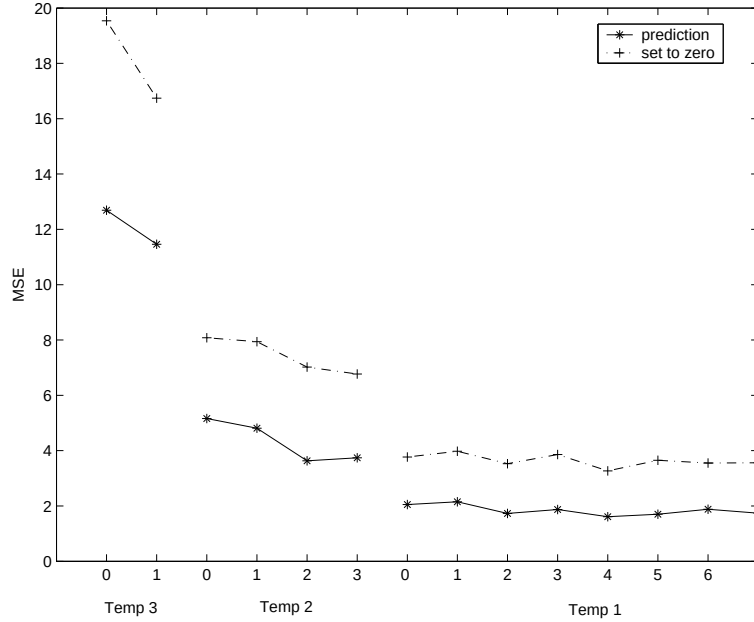


FIG. 3.1 – MSE of the spatial reconstruction of the detail frames at each temporal resolution level in a GOF.

frame reconstructed with the help of our model has sharper edges and outlines.

3.4 Error concealment in the Spatio-temporal wavelet domain

The application we consider in this section is the transmission of scalable video bits-stream over IP networks, prone to packet losses. The packetization strategy will highly influence the error concealment methods that we need to apply. Indeed, depending on the application and on the level of protection desired (and the overhead allowed for error protection), several strategies of packetization can be envisaged for the spatio-temporal coefficients :

1. one spatial subband per packet
2. all subbands with the same spatial resolution and orientation in one packet
3. all subbands at the same spatial level in each temporal detail frame in one packet.

We further analyse the influence of losing a packet at different spatio-temporal levels in each one of these settings and the ability of our prediction model to provide error concealment.

1. First, we analyse the concealment ability of our model when the packetization method consists of taking one subband per packet. In this case, if a spatio-temporal subband is lost, we predict it with the help of the neighbors of the coarser spatial and temporal resolution levels that we assume have been received by the decoder without losses. In Fig. 3.4 we present the mean square error (MSE) of the reconstruction of a

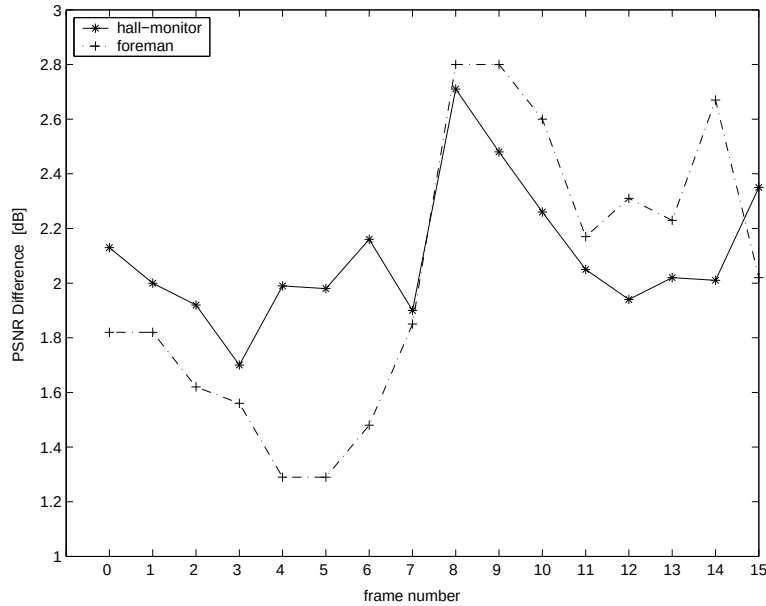


FIG. 3.2 – PSNR improvement for a GOF of 16 frames of the "foreman" and "hall-monitor" CIF sequences, when we predict the *finest frequency* subbands at different temporal resolution levels.

detail frame when we lose a subband at different spatial and temporal resolution levels. The MSE of the reconstructed frames using the prediction based on the statistical model is better than the one obtained by setting to zero the coefficients corresponding to the lost subband. Note also that, as expected, the loss of a subband at the last temporal resolution level influences more the MSE of the reconstructed frame than at any other temporal level.

An interesting point that comes out from these results is that the spectral behaviour of a temporal detail frame is different from that of still images. One can see from Fig. 3.4 that the energy of the subbands at different spatial resolution levels does not decay across the scales, as remarked for still images, but the medium and high frequency levels have more power than the lowest frequency one. This is due to the fact that the frames we are studying represent temporal prediction errors, therefore containing spatial patterns very similar to edges, whose energy is concentrated at rather high spatial frequencies.

Another useful point is to see how the prediction of a subband at different spatial resolution levels influences the reconstruction of a frame in the original sequence. Thus, in Table. 3.1 we present the MSE of the reconstruction of a frame in the original sequence when we lose a subband of a temporal detail frame at a given temporal resolution level (numbered 1, 2, 3) and at each spatial resolution level (denoted by I, II, III).

In this case, the reconstruction quality using our optimal predictor is proved to be superior to the reconstruction performed with the details corresponding to the lost subbands set to zero. We, also, notice that, as expected, the loss of a subband at the third temporal resolution level is more damaging for the reconstruction than at another temporal level.

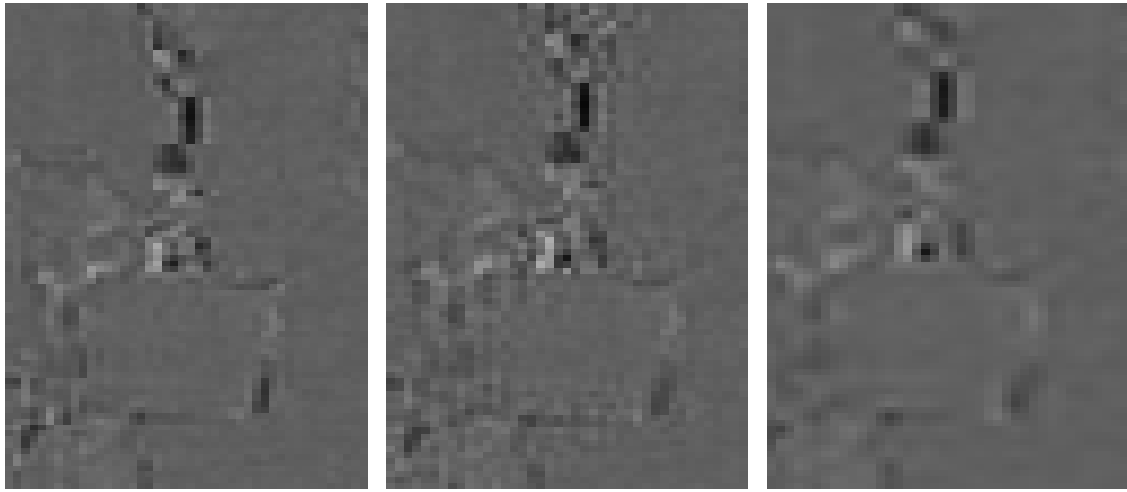


FIG. 3.3 – Zoom in a temporal detail frame at the first temporal resolution level. Left : original frame. Center : reconstructed detail frame when we predict its *finest resolution* subbands. Right : reconstructed frame when we set them to zero.

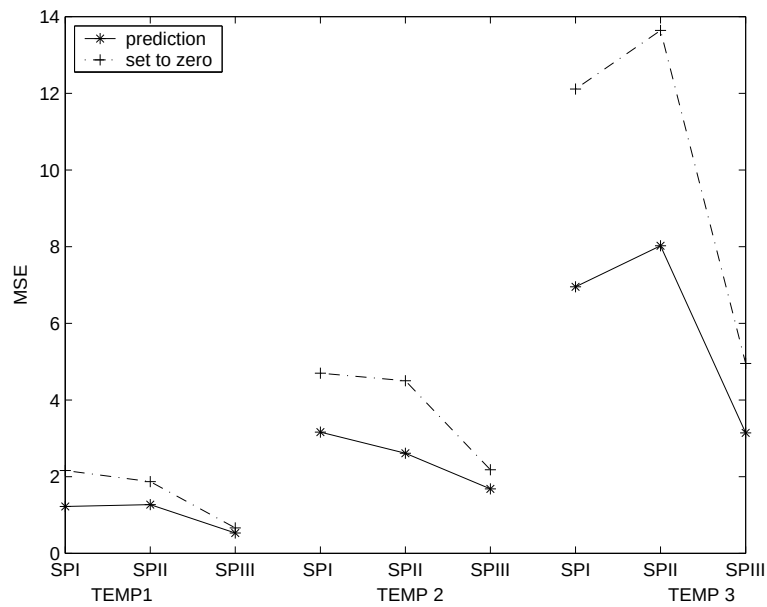


FIG. 3.4 – MSE of the spatial reconstruction of the first temporal detail frame when losing the horizontal subband at different temporal resolution levels and for three spatial resolution levels. SPi stands for the i -th spatial resolution level and $TEMPi$ for the i -th temporal decomposition level.

	<i>Temp1</i>			<i>Temp2</i>			<i>Temp3</i>		
	SPI	SPII	SPIII	SPI	SPII	SPIII	SPI	SPII	SPIII
set to zero	0.93	0.81	0.28	0.86	0.89	0.42	1.14	1.32	0.49
prediction	0.53	0.56	0.22	0.63	0.51	0.33	0.68	0.77	0.31

TAB. 3.1 – MSE of the reconstruction of the first frame of the original sequence, when prediction of a subband at different spatial resolution levels and at different temporal resolution levels is used, compared with setting to zero the lost coefficients.

In Fig. 3.5 we show a detail of a reconstructed frame at the first temporal resolution level, assuming that a subband at the second spatial resolution level was lost.

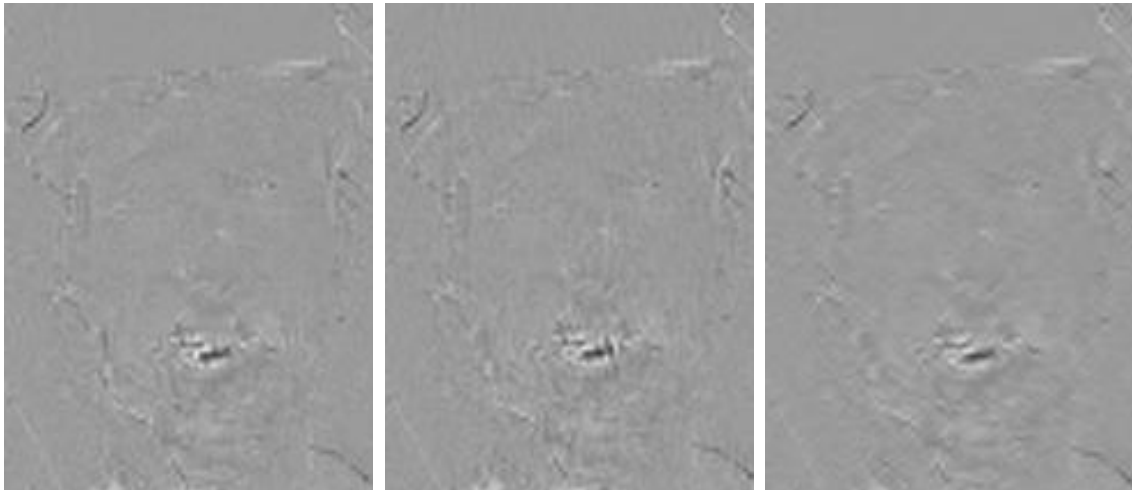


FIG. 3.5 – First temporal detail frame at the first temporal resolution level. Left : original frame. Center : reconstructed detail frame when we predict the lost horizontal subband of the *second spatial resolution level*. Right : reconstructed detail frame when we set it to zero.

2. Next, we consider the packetization technique in which all the subbands of the same spatial resolution and orientation level at the same temporal resolution level belong to a packet. In Fig. 3.6 we present the PSNR improvement of the reconstructed sequence assuming that we lose a packet at each temporal level. We notice here that our model leads to a higher improvement of the PSNR (up to 2.5 dB) when the lost packet is at the first temporal resolution level, where the prediction errors do not propagate through the temporal synthesis procedure.

3. The third method of packetization considered consists of taking the subbands of the same spatial resolution level in each temporal detail frame in one packet. In Table. 3.2 we present the MSE of a reconstructed frame of the original sequence in case we lose a spatial resolution level (first or second) of a temporal detail frame at different temporal resolution

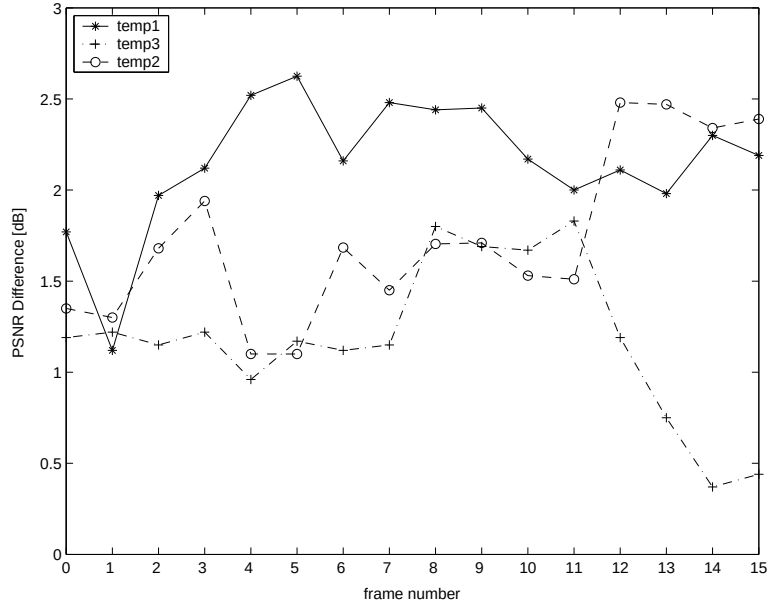


FIG. 3.6 – PSNR improvement (prediction *vs.* setting to zero) of a reconstructed GOF of the original sequence “foreman” in CIF format, 30 fps, when we loose the horizontal subbands of the *second spatial resolution level* at each temporal resolution level.

levels. We remark that, as we move to coarser temporal resolution levels, the loss of the coarser spatial resolution level becomes more significant. This could be expected, as the coefficients of a coarser temporal and spatial resolution level are bigger than at a finer one and so even a small error at the prediction becomes important in the reconstruction of the original frames.

Fig. 3.7 compares the reconstruction of a temporal detail frame by the proposed method with the one that consists of setting to zero the coefficients corresponding to the lost packet. We remark the oversmoothing resulting from the latter method and the good visual rendering of the high-frequency details obtained with the proposed method.

	<i>Temp1</i>		<i>Temp2</i>		<i>Temp3</i>	
	SPI	SPII	SPI	SPII	SPI	SPII
set to zero	1.59	1.12	1.52	1.27	2.52	2.58
prediction	0.85	0.83	1.02	1.15	1.74	1.97

TAB. 3.2 – MSE of a reconstructed frame of the original sequence when we lose the first or second spatial resolution level of a temporal detail frame at each temporal resolution level.

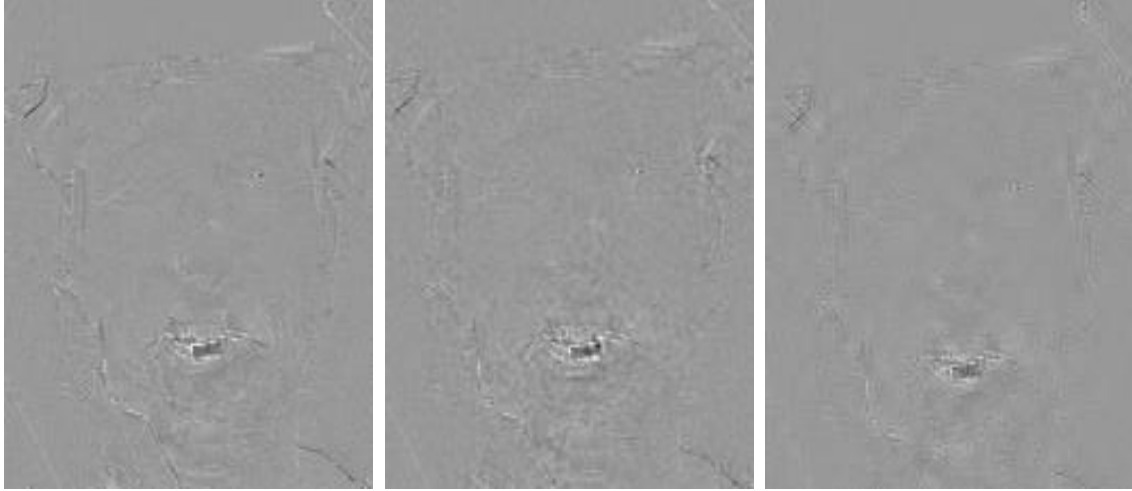


FIG. 3.7 – Temporal detail frame at the first temporal resolution level. Left : original frame. Center : reconstructed frame obtained by predicting the (lost) *second spatial resolution level*. Right : reconstructed frame when we set to zero the details corresponding to the lost packet.

3.5 Error concealment of scalable bitstreams

In the previous simulation results, we have assumed that except the lost packet, all the other subbands have been correctly received by the decoder. Here, we consider an even worse scenario : bandwidth reduction during the transmission requires to cut from the bitstream the finest detail subbands and, in addition, some packets are lost from the remaining bitstream. The main difference with the previous situation is that we need to predict not only the lost packet, but also the finest spatial resolution level. Some of the subbands in this level will be predicted based on spatio-temporal neighbors that also result from a prediction. As this procedure inherently introduces a higher error, we show by simulation results that the reconstruction of the full resolution video sequence has better quality than what we can obtain by a “naïve” decoder (which, as in the previous section, would set to zero all the unknown coefficients, so actually the first and the second spatial level).

We next examine the error concealment ability of our model in the same three packetization strategies as in Section 3.4 :

1. For the first packetization strategy (one subband per packet), the mean square error (MSE) of a frame in the original sequence when we lose a subband at the second spatial resolution level at different temporal resolution levels is computed. The difference in MSE when using our prediction method compared with the “naive” decoder is about 1 at the first temporal resolution level and about 1.5 for the second and the third temporal resolution level. This variation can be explained by the fact that the loss of a subband at the second and third spatial resolution level influences more the reconstruction of the original frame as this loss affects the spatial neighbors used in the reconstruction of the finest spatial subbands at the same temporal resolution level and also the spatio-temporal

neighbors of the finest spatial subbands at the next finer temporal resolution level.

2. In the second case (a packet includes all the subbands of the same orientation and spatial resolution level, at the same temporal resolution level), Fig. 3.8 illustrates the PSNR improvement of the original sequence in case we predict the finest resolution subbands after having predicted a lost packet at a coarser resolution level, compared with the case where all these lost subbands are set to zero.

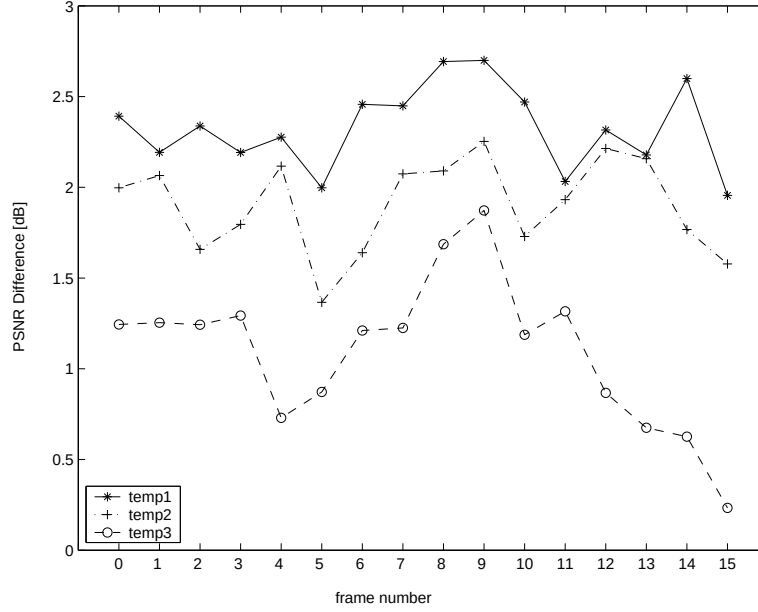


FIG. 3.8 – Improvement of the PSNR of the reconstructed GOF of the original sequence “foreman” when we loose the horizontal subbands of the second spatial resolution level at each temporal resolution level and we predict them and the finest spatial resolution subbands.

The higher improvement of the PSNR at the finest temporal resolution level is due to the fact that in this case the loss of the packet influences only the reconstruction of the temporal detail frame at this temporal resolution level. On the contrary, a loss at any other temporal level influences also the prediction of the subbands at finer temporal resolution.

3. At the end, we examine the third packetization technique (a packet includes all the subbands at a given spatial resolution level for each temporal detail frame). Table.3.3 illustrates the MSE when in the reconstruction of the original sequence we predict the lost second spatial resolution level as well as the finest ones compared to the case where both of these spatial resolution levels are considered to be lost and set to zero. We remark that even in the case where we lose the whole second spatial resolution level, our model is able to succesfully predict it from the received subbands and based on this to predict also the finer spatial resolution level.

Fig. 3.9 shows the reconstructed images when we lose the second spatial resolution

	Temp 1	Temp 2	Temp 3
set to zero	3.50	4.34	6.63
prediction	3.20	3.61	5.95

TAB. 3.3 – MSE of the reconstructed first frame of a GOF of original sequence “foreman” when we lose the second spatial resolution level of the first temporal detail frame at each temporal resolution level.

level of a temporal detail frame at the first temporal resolution level. Compared with Fig. 3.7, the frame obtained using the prediction method keeps almost the same amount of details, while the image obtained by setting to zero all the lost subbands suffered an even worse degradation.

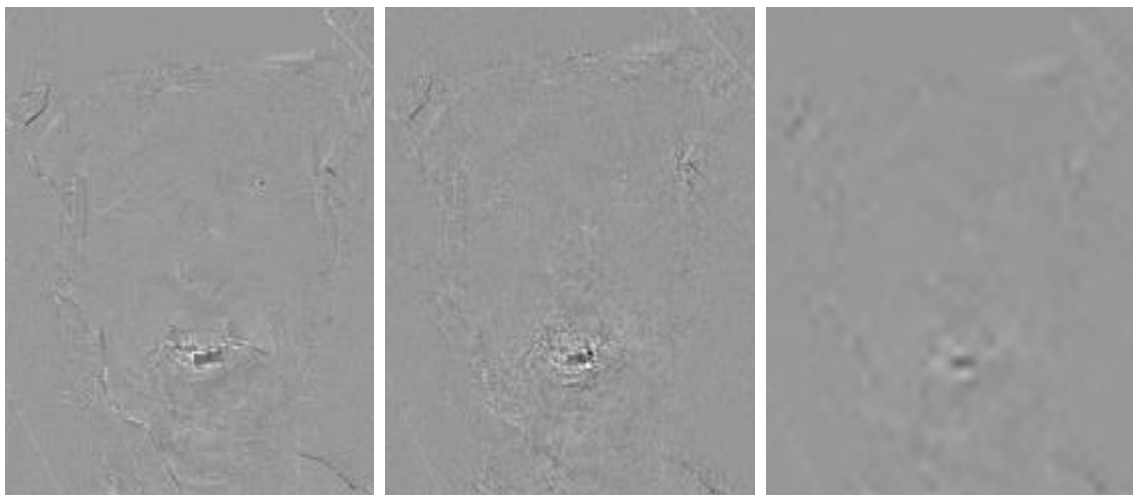


FIG. 3.9 – First temporal detail frame at the first temporal resolution level. Left : original subband. Center : reconstructed frame when we predict the second and then the first spatial resolution levels. Right : reconstructed frame when we set lost subbands to zero.

3.6 Conclusion

We have deduced an optimal MSE predictor for the lost coefficients and used these theoretical results in two applications to scalable video transmission over packet networks. In the first application, we have shown significant quality improvement achieved by this technique in spatio-temporal resolution enhancement. In the second one, we have proved the error concealment properties conferred by our stochastic model on a scalable video bitstream, under different packet loss conditions and with different packetization strategies.

Chapitre 4

Overview of Joint Source-Channel coding schemes

In this Chapter, we present an overview of the state-of-the art techniques in joint source-channel coding. We provide a classification of these methods, which allows us to connect our work to the existing approaches in channel-constrained vector quantization.

4.1 Introduction

In the previous chapters we investigated the behaviour of our source coder without taking account the channel constraints. The transmission was, indeed, supposed noiseless. However, in a typical transmission scheme, the output of the source encoder has to be protected against the errors caused by the channel noise. Generally, this protection is done by adding redundancy to the output of the source encoder through channel coding (or error correcting coding).

Shannon's separability theorem states that if the minimum achievable source coding rate of a given source is below the capacity of a channel, then the source can be reliably transmitted through the channel. In addition, it states that the source and the channel encoder can be separated in such a way that the source coding rate reduction takes place in the source encoder and the protection against channel errors in the channel encoder. Thus, the source coding and channel coding can be treated separately without any loss of performance for the overall system. However this separation theorem is justifiable only considering that the blocks of source and channel symbols are appropriately long and for arbitrarily complex encoders and decoders. The concatenation of a source coder, followed by a channel coder, which are separately optimized is called a *tandem scheme*.

In practical situations, there are limitations on the system complexity and on the length of source and channel coders, which make this separation questionable. Within the past few decades a solution to this problem came from an approach which consists of combined source-channel coding. Its objective is to include both source and channel coding modules into the same processing block in order to reduce the complexity of the overall system while, in the meantime, under non-ideal conditions, the performance increases. Typically, combined source-channel coding has been focused on either optimizing channel

coding with respect to the source, denoted by *source-optimized channel coding*, or on optimizing source coding with respect to the channel, which is called *channel-optimized source coding*.

In this Chapter we present an overview of the different approaches studied in the domain of combined source-channel coding. This overview is not exhaustive and we only investigate methods which are close to the one we propose. A large diversity of approaches exists for this problem and a much larger space would be necessary to review them all [95], [13], [56], [37], [27], [28].

4.2 Problem Statement

Before we present some of the previous works in the domain of joint source-channel coding, we would like to give a very general presentation of the transmission block diagram (see Fig. 4.1). Next we will try to present briefly each component of this scheme and its influence on the end-to-end distortion. It is important to notice that the basic idea of the joint source-channel coding is to include both source and channel coding in the same processing block to reduce the complexity of the overall system compared to the tandem scheme. Indeed, most of the works done in this domain are based on combining some of the blocks in the Fig. 4.1 into a single block and/or omitting some of them.

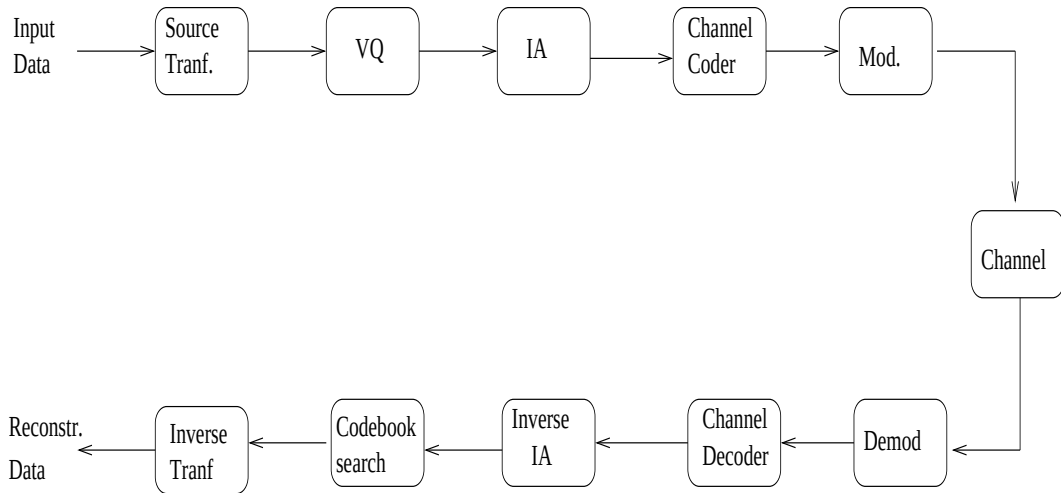


FIG. 4.1 – General block diagram of the transmission system.

There exists a vast literature on image transforms and their use in image or video coding. But in this chapter, we notice only that the most commonly used spatial transforms are the Discrete Cosine Transformation (DCT) or the wavelet transformation where for video, this is applied either on the temporal prediction error in a hybrid DPCM scheme or on the temporal detail frames of a (MC) temporal subband decomposition. As a general requirement, we also note that the source transform has to be unitary (or almost unitary).

4.2.1 Vector Quantization

Let $\mathbf{x} = [x_1, x_2, \dots, x_d]$ be a d -dimensional source vector. The design of a codebook $C = \{\mathbf{y}_1, \mathbf{y}_2, \dots, \mathbf{y}_M\}$ of size M consists in partitioning the d -dimensional space of the random vector $\mathbf{x}$ into M non overlapping regions or cells $\{S_i, 1 \leq i \leq M\}$ and associating with each cell S_i a d -dimensional vector $\mathbf{y}_i$ which is called the reconstruction vector. A function $A : \mathbf{y}_i \rightarrow \mathbf{b}_j$ associates a binary codeword (or index) $\mathbf{b}_j$ of size $n = \log_2 M$ to each reconstruction vector. The binary codeword is the one supposed to be transmitted over the channel. If we consider a noiseless channel, the total distortion is only due to the quantization.

The most commonly used distortion measure is the Mean Square Error (MSE) :

$$d(\mathbf{x}, \mathbf{y}_i) = \|\mathbf{x} - \mathbf{y}_i\|^2$$

Two approaches are quite common for the design of a vector quantizer. The first one is based on the generalized Lloyd algorithm (GLA) or Linde-Buzo-Gray (LBG) [50] algorithm which is also a particular case of clustering algorithms. The second one is based on the use of some subset of a lattice to force highly structured codebook.

Generalized Lloyd Algorithm

The GLA is a descent algorithm, when run on either a probabilistic model or on a training set, and it can always be used to improve a codebook in terms of reducing distortion. It iteratively improves a codebook by alternately optimizing the encoder for the decoder, using a minimum distortion or nearest neighbor mapping and the decoder for the encoder, replacing the old codebook by generalized "centroids". For squared error, centroids are the Euclidean mean of the input vectors mapped into a given index. Note that the GLA is almost always based on a training set of typical data rather than on mathematical models of the data (which can be however exploited in the algorithm, through analytical formulas for some simple pdfs).

VQ with constrained structures

Based on the classification in [17] we present some different constrained structures of the VQ.

1. Lattice VQ : A lattice Λ_n in $\mathbb{R}^n$ is composed of all integral combinations of a set of linearly independent vectors $\mathbf{a}_i$ that span the space such that

$$\Lambda = \{y | y = u_1 \mathbf{a}_1 + u_2 \mathbf{a}_2 + \dots + u_n \mathbf{a}_n\}$$

where u_i are integers.

Lattice VQ is widely used in image and video coding ([1], [3], [97], [71], [72], [80], [69]). Two important issues in lattice VQ concern the *region of support*, which is the subset of the lattice that will actually be coded and the *scaling of the lattice*, which has the effect of increasing or decreasing the size of the basic quantization cell. Using lattice points as codewords avoids the task of designing the codebook. Thus, lattice vector quantizers offer the possibility of a substantial reduction in computational

and storage complexity over unstructured full-search VQ designed by the GLA algorithm. Best known lattices for low dimensions ($n \leq 8$) are the root lattices A_n ($n \geq 8$), D_n ($n \geq 2$), E_n ($n = 6, 7, 8$), the Barnes-Wall lattice Λ_{16} and the Leech lattice Λ_{24} . From another point of view in [30], [31], [72] the superiority of the cubic Z lattice over the E_8 and Leech lattices in the case of generalized Gaussian sources is established. The interest of this latter remark is that the generalized Gaussian distribution is a very common way of modelling the source data in the area of image and video.

2. Tree-Structured VQ : The codeword is selected by a sequence of binary minimum distortion decisions comparing the input vector to stored vector reproductions at each available node. Code trees can be balanced if all indexes have the same length or unbalanced in the opposite case. Compared with a full-search unstructured VQ, the search complexity of a balanced tree is linear in the bit rate instead of exponential but at the cost of a roughly doubled memory size.
3. Multistage VQ : The encoding task is divided into several stages. The first stage performs a relatively crude encoding of the input vector using a small codebook. Then, a second-stage quantizer operates on the error vector between the original vector and the quantized first stage output. The quantized error vector provides a refinement to the first approximation. In the multistage VQ the encoding complexity increases linearly with the dimension-rate product.
4. Predictive VQ : The encoder makes a prediction of the incoming vector based on previously encoded vectors. The difference between the actual input vector and its prediction is called the residual vector. This residual is vector quantized. The prediction is often a simple linear predictor that takes a weighted average of nearby previously encoded coefficients.
5. Trellis-Coded VQ : It can be thought of as having a supercodebook that is partitioned into a collection of subcodebooks. At any time, the encoder may have only a few of the subcodebooks available. A popular choice is to have the superbook be a lattice and the subcodebooks be sublattices. The encoder has a collection of possible states and allowed transitions between those states. Each of the transitions between states corresponds to one of the subcodebooks. The encoder in a given state can view a single source vector, looking at all the available subcodebooks and find the nearest neighbor. The encoder communicates this word to the decoder with a two-part binary vector. The first part tells the encoder's next state by sending an index of which allowable transition is being taken. The second part indexes which of the possible codewords in that particular codebook is chosen.

4.2.2 Index assignment (IA)

We have shown that, when an d -dimensional vector is fed to the quantizer, a n -bit binary codeword is produced. This codeword is said to be the index of the vector used for signal reconstruction at the receiver. The decoder receives a codeword j and produces an d -dimensional reconstruction vector $\mathbf{y}_j$. In the index assignment task we assign indices, i.e. codewords, to the codevectors so as to reduce the effect of channel errors. Index assignment

of the codevectors does not affect the average distortion, in the absence of channel noise, while in its presence, the assignment plays an important role in determining the overall VQ performance.

One approach to the problem is to design a quantizer without regard to the channel but then to code the resulting indices in a way that ensures that small Hamming distance of the channel codewords corresponds to small distortion between the resulting reproduction codewords (with no explicit channel coder). Index assignment algorithms are necessarily suboptimal since the search problem is NP-hard.

In [94] an algorithm, named Binary Switching Algorithm, which performs pseudo-Gray coding on a given VQ codebook and gives as output a locally optimal assignment of binary indexes, is presented. This index assignment provides an improvement in the average distortion due to channel errors, over an arbitrary index assignment. The main idea involves iteratively switching the positions of two codevectors to reduce the total distortion after each switch. Each codevector is assigned a cost function which measures the total contribution to bit error distortion when the particular codevector is selected by the encoder, assuming a certain permutation. The codevectors are sorted in decreasing order of their cost values and the one with the largest cost, i.e. $\mathbf{y}_0$, is selected as candidate to be switched first. This vector is temporarily switched with each of the other codevectors to determine the potential decrease in terms of the total distortion, following each switch. The codevector which yields the greatest decrease when switched with $\mathbf{y}_0$ is then switched permanently with it. The algorithm is repeated with the next highest cost and so on till no decrease in total distortion is performed when every codevector is switched with every other codevector. As the input of this algorithm is an initial index assignment, it is proposed a preprocessing technique where the highest probability codevectors are distributed throughout the codebook and close vectors in terms of distortion function are placed as their nearest neighbors.

In [20] a simulated annealing type algorithm for optimal assignment of binary codewords to codevectors is presented. Simulated annealing is a Monte Carlo algorithm which mimics the objective function associated with the combinatorial optimization problem as the energy of the physical system and by slowly reducing an appropriately defined effective temperature, seeks the minimum energy state. In the case of the index assignment problem, the objective function is the channel distortion, i.e. $D_c(\mathbf{b})$, for a specific permutation of the codewords $\mathbf{b}$. The simulated annealing algorithm works as follows : **1.** define an initial high temperature $T = T_0$ and choose randomly an initial state $\mathbf{b}$; **2.** Choose a permutation of the state $\mathbf{b}$, i.e. $\mathbf{b}'$ and compute $\Delta D_c = D_c(\mathbf{b}') - D_c(\mathbf{b})$. If $\Delta D_c < 0$ then replace $\mathbf{b}$ by $\mathbf{b}'$, if $\Delta D_c \geq 0$ then replace $\mathbf{b}$ by $\mathbf{b}'$ with probability $\exp\{-\Delta D_c/T\}$; **3.** Decrease the temperature T and go back to **2**, unless T is below some prescribed freezing temperature or if it appears that a stable state is reached. This process allows the algorithm to climb out of local minima when the temperature is high in the hope that as the system is cooled the state falls in the global minimum or somewhere close to it, out of which the algorithm should not be able to climb. However, the convergence to the global minimum depends on the choice of the initial value of the temperature T_0 and a sufficiently slow cooling schedule which is not suitable for practical applications.

Another approach of careful assignment of binary codes to VQ codewords in order to minimize the distortion due to channel errors is found in [70]. They propose a polynomial

time algorithm for computing a good suboptimal minimax index assignment. They formulate the minimax nonredundant channel coding task as a graph problem and they adopt the simplifying assumption that no more than one bit is in error for any received n binary index. The proposed algorithm provide good minimax distortion without sacrificing the average performance. The iterative technique is initialised to the index assignment produced by GLA algorithm with splitting for initialisation. An extended presentation of this algorithm will be presented in Chapter 6 (Section 6.7) as we will use it for comparison purposes.

In [46] it is shown that for binary symmetric channels, the channel distortion is minimized if the codebook can be expressed as a linear transform of an hypercube. The index assignment problem is then regarded as a problem of linearizing the codebook. The Hadamard transform and a fidelity criterion, highly correlated to the channel, are used in order to find the optimal index assignment given an arbitrary vector quantizer. This work will be presented, in detail, in the next Chapter.

In [39] is presented a vector quantizer by a linear mapping of a block code assuming a memoryless binary-symmetric channel. The reconstruction vectors of the codebook are given as a linear mapping of a binary block code in systematic form obtained from the index bits transmitted on the channel. The vector quantizer is defined by the following constraint on the reproduction vectors :

$$\mathbf{y}_j = \mathbf{T}\mathbf{b}_j$$

where $\mathbf{T}$ is a real-valued matrix of dimension $d \times (n + 1)$. The vector $\mathbf{b}_j$ is a codevector in an (n, k) binary linear block code in systematic form augmented with a +1 in the first position. Thus, the quantizer consists of a mapping from an $n + 1$ -dimensional hypercube to a d -dimensional source space. Only the corners of the hypercube specified by the block code are occupied by a block codevector. For a given block code, the mapping matrix $\mathbf{T}$ completely specifies the codebook. Only the information bits $k + 1$ are transmitted on the channel and the parity bits are generated at the receiver from the received information bits. They investigate two cases :

1. A maximum-entropy quantizer (MEE) where the indices are assumed to be equiprobable. It is shown that the minimum possible distortion is obtained if we can find an index assignment that has all the mapping energy at information bits. The mapping energy of a component in the block codevector is defined as the squared length of the corresponding column in $\mathbf{T}$. A sample iterative algorithm is used for the mapping matrix. The mapping is updated every time a sample $\mathbf{x}$ from the training database is quantized. The update of the mapping matrix is given by : $\mathbf{T}^{(m+1)} = \mathbf{T}^m - \frac{\partial d(\mathbf{x}, \mathbf{T}^m \mathbf{b}(\mathbf{x}))}{\partial \mathbf{T}} f(m)$ where m is an iteration counter and $f(m)$ is an annealing function specifying the step size.
2. When the entropy of the quantizer is not maximum, parity bits are exploited to obtain a better performance. A block-iterative optimization algorithm for the mapping matrix is proposed. This algorithm counterpart the GLA for the squared-error distortion method with only difference that the centroid condition is replaced with the optimality condition $E[\mathbf{b}(\mathbf{X})\mathbf{b}^T(\mathbf{X})] \cdot \mathbf{T}^T = E[\mathbf{b}(\mathbf{X})\mathbf{X}^T]$ which constitutes d systems of equations of size $n + 1$.

In [58] the subject of affine index assignment is revisited due to the low implementation complexities of such codes. Natural Binary Code, Folded Binary Code, Two's Complement Code and Gray Code are all "affine" functions in a vector space over binary field. Specifying an affine index assignment requires only $O(n^2)$ bits for a 2^n - point quantizer, as opposed to $O(n2^n)$ bits for an unconstrained index assignment. Linear index assignment is a special case of affine index assignment. The authors study the performances of general affine index assignments when explicit block channel is used on a binary-symmetric channel or when no explicit channel coding is used on binary asymmetric channels.

In [87] the index assignment problem for arbitrary VQ- no entropy constraint- and arbitrary (possibly nonbinary) discrete memoryless channels is investigated. It is shown that the reason why the Hadamard transform over a binary-symmetric channel is a powerful tool for IA (see [46], [39], [58]), is due to the fact that the Hadamard transform is the eigentransform of the transition matrix of a binary symmetric channel. For the general discrete memoryless channels the diagonalization of the transition matrix (or of a matrix related to the transition matrix) is still a key tool in order to derive general expressions having the known Hadamard transform-based formulas as special cases. Let us consider a d -dimensional M -level VQ codebook with $\{\mathbf{y}(j)\}_{j=0}^{M-1}$ codevectors. The encoder is the mapping from $\mathbb{R}^d$ to $I_M = \{0, 1, \dots, M-1\}$ according to $\mathbf{X} \in S_i \rightarrow I = i$, where $\mathbf{X} \in \mathbb{R}^d$ is the source vector and the sets $\{S_i\}_{i=0}^{M-1}$ form a partition of the Euclidean space $\mathbb{R}^d$. Let $\mathbf{Q}$ define the channel transition matrix with $Q(j|i) = Pr(J = j|i = i)$ the probability that the decoder receives the index j when the index i has been sent. The channel distortion can be written as :

$$D_c = E\|\mathbf{y}(I) - \mathbf{y}(J)\|^2 = \sum_{i=0}^{M-1} P(i) \sum_{j=0}^{M-1} Q(j|i) \|\mathbf{y}(i) - \mathbf{y}(j)\|^2 = Tr\{\mathbf{Y}\mathbf{G}\mathbf{Y}^T\}$$

where $P(i) = Pr(I = i) = Pr(\mathbf{X} \in S_i)$, $\mathbf{Y}$ is a matrix containing the codevectors as columns and where

$$\mathbf{G} = \sum_{i=0}^{M-1} \sum_{j=0}^{M-1} P(i)Q(j|i)(\mathbf{e}_i - \mathbf{e}_j)(\mathbf{e}_i - \mathbf{e}_j)^T$$

where $\mathbf{e}_i$ is a size M column vector with a one in position i and zero otherwise. Let $\{\mathbf{g}_i\}_{i=0}^{M-1}$ and $\{\gamma_i\}_{i=0}^{M-1}$ be the eigenvectors and the eigenvalues respectively of the matrix $\mathbf{G}$. Then

$$\mathbf{G} = \sum_{i=0}^{M-1} \gamma_i \mathbf{g}_i \mathbf{g}_i^T \text{ and thus :}$$

$$D_c = Tr\{\mathbf{Y}\mathbf{G}\mathbf{Y}^T\} = \sum_{i=0}^{M-1} \gamma_i \|\mathbf{Y}\mathbf{g}_i\|^2 \quad (4.1)$$

This expression for the channel distortion is valid over an arbitrary discrete memoryless channel and for arbitrary VQ.

Always in [87], it is shown that in the case of a binary symmetric channel the transition matrix $\mathbf{Q}$ has the eigenvalues $\lambda_i = (1 - 2q)^{w(i)}$ and the eigenvectors $\mathbf{q}_i = M^{-1/2} \mathbf{h}(i)$, where $w(i)$ is the Hamming weight of the natural binary representation of the integer i and $\mathbf{h}(i)$

is the i th column of an $M \times M$ Hadamard matrix $\mathbf{H}$. For a maxentropic quantizer, $\mathbf{G} = 2M^{-1}(\mathbf{I}_M - \mathbf{Q})$ ($\mathbf{I}_M$ the $M \times M$ identity matrix) with eigenvalues $\gamma_i = 2M^{-1}[1 - (1 - 2q)^{w(i)}]$ and eigenvectors $\mathbf{g}_i = \mathbf{q}_i = M^{-1/2}\mathbf{h}(i)$. Employing (4.1) for full entropy encoding over a binary symmetric channel we find the same expression of the channel distortion depending on the Hadamard transform of the codebook, as presented in [46], [39].

The expression (4.1) suggests that D_c is low if the rows of $\mathbf{Y}$ are orthogonal to the eigenvectors $\mathbf{g}_i$ that correspond to large eigenvalues. This kind of VQ is referred in [87] as *channel-constrained VQ*, since the choice of the constraint on the VQ is dictated by the channel and the purpose is the channel robustness. If we assume that the eigenvalues and the corresponding eigenvectors are reordered so that $0 \leq \gamma_0 \leq \gamma_1 \cdots \gamma_{M-1}$ then constraining the codebook $\mathbf{Y}$ to have rows that lie in the subspace spanned by $\{\mathbf{g}_i\}_{i=0}^{K-1}$, for some $K < M$, should make the codebook robust against channel errors. However, in (4.1), $\mathbf{G}$ depends on the index assignment given to the codevectors since $\mathbf{G}$ is defined in terms of the probabilities $\{P(i)\}$. In order to find an expression which does not depend on the index assignment of the codevectors the author of [87] employs a permutation matrix $\mathbf{F}$ to represent the employed index assignment. Let $P = \text{diag}(P(0), P(1), \dots, P(M-1))$ and define a matrix $\mathbf{D}$ as $D_{ij} = \|\mathbf{y}(i) - \mathbf{y}(j)\|^2$ using the initial ordering of the codevectors. Then the channel distortion corresponding to the initial IA is : $D_c = \text{Tr}\{\mathbf{PQD}\}$ and by using a particular IA represented by $\mathbf{F}$ is equivalent to rearranging the elements of both $\mathbf{P}$ and $\mathbf{D}$ according to $\mathbf{P} \rightarrow \mathbf{FPF}^T$ and $\mathbf{D} \rightarrow \mathbf{FDF}^T$. The D_c can be now expressed by : $D_c = \text{Tr}\{\mathbf{PF}^T\mathbf{Q}^T\mathbf{FD}\}$ and the IA problem is now to find a permutation matrix that minimizes this last expression of D_c . Let $\mathbf{K}$ be a symmetric matrix $\mathbf{K} = \{\mathbf{DP} + (\mathbf{DP})^T\}$ and β_{min} its lowest eigenvalue. Then the modified version $\tilde{\mathbf{K}} = \mathbf{K} + \beta_{min}\mathbf{I}_M$ must be positive semidefinite and $\tilde{\mathbf{K}}$ can be factored into squared roots according to $\tilde{\mathbf{K}} = \mathbf{S}^T\mathbf{S}$. Then for channels with symmetric transition matrices and letting the eigendecomposition

of $\mathbf{Q}$ be : $\mathbf{Q} = \sum_{i=0}^{M-1} \lambda_i \mathbf{q}_i \mathbf{q}_i^T$, the optimal permutation matrix $\mathbf{F}^*$ is given by :

$$\mathbf{F}^* = \arg \min_{\mathbf{F} \in A} \sum_{i=0}^{M-1} \lambda_i \|\mathbf{SF}^T \mathbf{q}_i\|^2$$

where A is the set of all possible permutation matrices. Thus, the IA problem is equivalent to the problem of reordering the columns of the matrix $\mathbf{S}$ in a such a way that the rows of $\mathbf{S}$ are as "orthogonal as possible" to eigenvectors of $\mathbf{Q}$ corresponding to large eigenvalues λ_i . The geometrical interpretation of the best IA is the one that "concentrates" the rows of $\mathbf{S}$ as much as possible to the subspace spanned by the eigenvalues of $\mathbf{Q}$ that correspond to the low eigenvalues.

Another approach is based on simultaneous optimization of quantizer and IA : "Channel optimized vector quantization" (COVQ) introduced by Farvardin [22], [21] which indeed uses a distortion measure involving both the quantization error and the error due to channel perturbation. COVQ assumes precise knowledge of the channel characteristics. This knowledge is used to modify the VQ codebook and the encoding rule so as to achieve optimal performance under the prescribed channel conditions. However, COVQ incurs some performance degradation under clean channel conditions. An extended presentation of this technique will be presented in Section 4.4 of this chapter.

4.2.3 Channel coding

Channel coding consists of various methods that add some protection to the message obtained from the source coding process. This is done by adding some redundancy to the message which will be used later in the channel decoder to detect and to correct the errors due to the channel noise. There are two main groups of channel coders, the block coders and the convolutional coders. However, we restrict our presentation here to a branch of the block coders which is the Reed-Muller (RM) linear codes and to a branch of the convolutional codes which is the Rate Punctured Convolutional (RCP) codes, due to their flexibility and their increasing usage in the domain of channel coding.

Reed-Muller codes

A Reed-Muller $RM(r, m)$ of order r is a linear code of blocklength $n = 2^m$, dimension

$$k = \sum_{l=0}^r \binom{m}{l}$$

and minimum distance $d_{min} = 2^{m-r}$, where m is a positive integer and $0 \leq r \leq m$. The error correction capability of an $RM(r, m)$ is $\max(0, 2^{m-r-1} - 1)$ errors. An $RM(r, m)$ is the set of all binary vectors of length n associated with the Boolean polynomials of degree at most r . A Boolean polynomial is a linear combination of Boolean monomials with coefficients in the Galois field of order 2, $GF(2)$. A Boolean monomial p in the variables $x_1, \dots, x_m$ is an expression of the form, $p = x_1^{r_1} x_2^{r_2} \dots x_m^{r_m}$ where $r_i \in \{0, 1, 2, \dots\}$ and $1 \leq i \leq m$.

It is interesting to note that RM codes can be defined recursively and based on a construction which is also referred as $(u, u+v)$ construction. If C_i is a $[n, k_i, d_i]$ for $i = 1, 2$ linear code then the code defined by :

$$C = \{(u, u+v) : u \in C_1, v \in C_2\}$$

is a $[2n, k_1+k_2, \min\{2d_1, d_2\}]$ linear code. Then, the generator matrix of $C : G = G_{RM(r,m)}$ is given by $C_1 : G_1 = G_{RM(r,m-1)}$ and $C_2 : G_2 = G_{RM(r-1,m-1)}$ as follows :

$$G_{RM(r,m)} = \begin{bmatrix} G_{RM(r,m-1)} & G_{RM(r,m-1)} \\ \mathbf{0} & G_{RM(r-1,m-1)} \end{bmatrix}$$

Thus, $RM(0, m)$ is just a repetition of either zeros or ones of length 2^m and at the other extreme, the m -order $RM(m, m)$ consists of all binary strings of length 2^m .

Rate-Punctured Convolutional Codes

A rate punctured convolutional (RPC) code [38] is a high-rate code obtained by periodic elimination of specific code symbols from the output of a low-rate encoder. The resulting high-rate code depends on both the low-rate code, called the original or *the mother code*, and on the number and the specific positions of the punctured symbols. The pattern of punctured symbols is called perforation pattern and is conveniently described in a matrix form called the *perforation matrix*.

Consider a high-rate code $R = k/n$ from a given mother code of low-rate $R = 1/N$. From every $N * k$ code symbols corresponding to the encoding of k information bits by the original encoder, a number of $S = (Nk - n)$ symbols are deleted according to a chosen perforation pattern. The perforation pattern can be expressed as matrix having N rows and k columns with only binary elements corresponding to the deleting ("0") or keeping ("1") of the corresponding code symbol of the original encoder.

The main advantage of the RPC codes is that the encoder and decoder can be kept the same and by changing only the puncturing matrix we can obtain different rate codes. Thus, the complexity is significantly reduced.

4.2.4 Channel decoding

There are various methods to decode a linear code as RM but, in this Section, we present only the methods that we shall use in the sequel. In the case of RPC codes, the Viterbi decoder is the most famous and we provide below only a brief presentation of it.

Decoding Rate Punctured Convolutional codes with the Viterbi algorithm

A general representation of a RPC or a convolutional code is the trellis form which, indeed, is a linear time sequencing of events. In the case of RCP which comes from a mother code of type $R = 1/N$ with constraint length L , there are 2^{L-1} states in the vertical axis and we connect each state to the next state by the allowable codewords for that state depending on the input bit. In general, Viterbi examines the metrics of all paths entering each state and keeps only the path with the smallest metric, called surviving path, for each state. Thus, there are 2^{L-1} surviving paths at each stage and 2^{L-1} metrics, one for each surviving path. For a soft-decision decoding the metric is the Euclidean distance instead of the Hamming distance used for hard-decision.

When considering the general case of a convolutional case of $R = K/N$ the number of computations performed at each stage increases exponentially with the K and L . Also, the decoding delay for a long information sequence is usually too long and a simple modification is, at any given time t , to retain only the most recent b decoded information bits in each surviving sequence. b is chosen sufficiently large ($b \geq 5L$) in order to have a negligible degradation in the performance of the algorithm.

The RPC codes have the beautiful property to be decoded as a $R = 1/N$ convolutional code. The decoding is performed on the trellis of the mother code where the only modification consists of discarding the metric increments corresponding to the punctured code symbols. Given the perforation pattern of the code, this can be readily performed by inserting dummy data into the positions corresponding to the deleted code symbols. In the decoding process this dummy data is discarded by assigning them usually zero. By this way, the high rate punctured codes involve none of the complexity of the straightforward decoding of high rate convolutional codes.

Decoding Reed-Muller codes

As the convolutional code, a linear block code can be represented in a trellis form. We based our brief presentation on the work of McEliece [57] who also makes a nice

description of the previous works in this field. In order to keep the complexity of the Viterbi algorithm low, the "best" trellis representation for a given linear block code is the one with the fewest edges. He shows that the trellis introduced by Bahl, Cooke, Jelinek and Ravin [2] in 1974 (BCJR) uniquely minimizes the edge count and is isomorphic to the Forney "minimal trellis" [24]. In the BCJR trellis every codeword in a code $C(n, k, d)$ corresponds to a path of length n . The BCJR trellis or the "minimal-trellis" are based on an important set of subcodes of the code C , called *past* and *future* subcodes. A past subcode consists of all codewords whose nonzero components are in the "past" and a future subcode of all codewords whose nonzero components are in the "future" relative to the current time.

McEliece introduces the class of "trellis-oriented" or "minimal-span" generator matrices for block codes which facilitates the construction of the BCJR trellis. If $\mathbf{x} = (x_1, x_2, \dots, x_n)$ is a nonzero binary n -vector, its *left index*, denoted $L(\mathbf{x})$, is the smallest index i such that $x_i \neq 0$ and its *right index*, denoted $R(\mathbf{x})$, is the largest index i such that $x_i \neq 0$. The *span* of $\mathbf{x}$, $Span(\mathbf{x})$, is the discrete "interval" $\{L(\mathbf{x}), L(\mathbf{x}) + 1, \dots, R(\mathbf{x})\}$ and spanlength of $\mathbf{x}$, is the number of elements in $Span(\mathbf{x})$. The spanlength of a matrix is the sum of the spanlength of the rows. Among all generator matrices of a binary linear code C , those for which the spanlength is as small as possible are called minimal span generator matrices (MSGM). A set of binary vectors $\{\mathbf{x}_1, \dots, \mathbf{x}_k\}$ has the left-right property (LR property) if $L(\mathbf{x}_i) \neq L(\mathbf{x}_j)$ and $R(\mathbf{x}_i) \neq R(\mathbf{x}_j)$, whenever $i \neq j$. A matrix G is an MSGM if and only if it has the LR property. In order to produce an MSGM, McEliece presents a "greedy algorithm" which uses a sequence of elementary row operations. A MSGM is a "trellis-oriented" matrix in the sense that the past and future subcodes needed to construct the BCJR trellis can be read directly.

Unfortunately, the construction of the trellis of even a second-order RM becomes quickly very complex. In [25], [26] we find an interesting approach of soft-decision decoding of linear block codes which even if it is a suboptimal method, it is proved simple and efficient. This algorithm consists of the following steps :

1. reordering of the symbols of the received sequence in decreasing order, based on a measure of reliability which for BPSK transmission is $|r_i|$, where $\mathbf{r} = (r_1, r_2, \dots, r_N)$ is the received sequence
2. hard-decision decoding of the reordered received sequence which provides a codeword expected to have as few information bits in error as possible
3. reprocessing and improving the hard-decision decoding progressively in stages until the practically optimal error performance or a desired error performance is achieved

The reprocessing stages follow directly from the statistics of the noise after ordering. A sufficient condition to recognize the Maximum Likelihood Decoded codeword in order to stop the algorithm, is developed. In addition, they can determine the number of reprocessing stages needed to decode up to a particular Bit Error Rate (BER). By this way, they guarantee the error performance associated with the Maximum-Likelihood Decoding for a given BER.

In [26], it is introduced an algorithm for coset decoding of the $(u, u + v)$ construction presented above. The code $C[2n, k_1 + k_2, \min\{2d_1, d_2\}]$ can be viewed as the union of 2^{k_2} cosets of the $(2n, k_1, 2d_1)$ repetition code C' with generator matrix $[G_1 G_1]$. Then each

coset can be decoded independently, after removing the contribution of its representative from the received sequence. The final decision is obtained by comparing the 2^{k_2} decoding costs of each coset decoding. Indeed, the decoding of the repetition code C' can be realized by decoding the code C_1 with the algorithm consisting of the above three steps applying a different type of reliability measure.

4.3 Source-optimized channel coding

In source-optimized channel coding, the source code is designed for a noiseless channel. A channel code is then designed for the source code to minimize end-to-end distortion over a given channel, which is typically a binary symmetric channel (BSC), an additive white Gaussian noise (AWGN) channel with a given modulation, or a time-varying channel. When no explicit channel coder is used, the way that a VQ, designed for a noiseless channel, becomes robust in the presence of noise is by applying a good index assignment as presented in Section 4.2.2.

When an explicit coder is used, a great attention is paid to optimal partitioning of a given amount of resources among the source and channel coder. By "optimal" we mean a partitioning that results in minimum expected value of distortion D which, in general, is chosen to be the mean square error. Traditionally, the bit allocation algorithms consider the objective function D as the sum of individual user functions, i.e. d_k , such that the sum of individual user resources, i.e. n_k , does not exceed the resource limit R_{target} :

$$\min_{n_k} D = \sum_{k=1}^K d_k(n_k) \quad \text{subject to} \quad \sum_{k=1}^K n_k \leq R_{target}$$

where K is the number of users.

This constraint problem can be treated as an unconstrained minimization problem with the help of Lagrange multipliers and can be stated as :

$$\min_{n_k} \sum_{k=1}^K [d_k(n_k) + \lambda n_k]$$

The function $Q(\lambda) = \sum_{k=1}^K \lambda n_k$ can be considered as a penalty function as it adds a penalty line of slope λ to each of the K distortion functions. In [83] an algorithm treating the problem of how the multiplier or the slope of the penalty lines should be adjusted, is presented. It is an integer programming algorithm which terminates with an optimal solution or an error-bounded approximate solution. This algorithm is based on the notion of singular points, a special set of multiplier values such that the optimal set solutions is non-unique. Geometrically, such multipliers create a slope on the source distortion curves such that a pair of adjacent convex-hull data points on a particular distortion function are simultaneously minimum. The neighboring singular points always share one solution and there can be no additional solutions between the neighboring singular points. Thus, the idea of this algorithm is that instead of sweeping the multiplier value λ from zero to infinity continuously in search for an operational rate that is close to the target rate,

it is sufficient to look only at the singular points, since they alone lead to all possible Lagrangian solutions anyway.

In [61] and [62] Modestino and *al.* provide an early approach of combined source-channel coding using in each case different scalar quantizer. In [61] a 2-D differential pulse code modulation (DPCM) for image coding followed by convolutional channel coding and in [62] a discrete cosine transform-coded images with convolutional codes on AWGN are presented. When the transmission is supposed error-free and no channel code is used, the reconstructed image quality increases as the transform block size increases. For any fixed block size, the performance is monotonically increasing in the average number of bits transmitted per pixel. On the other hand, when the channel is very noisy there is a complete reversal behavior compared to that of the noiseless channel. For low SNR the smaller block sizes are preferred in order to achieve the best reconstructed image quality. The explication is obvious : more bits transmitted per block results in a larger number of quantizer levels decoded incorrectly, lowering the system performance. Based on these remarks, they conclude that the protection of the output of the quantizer against the channel noise is necessary. They indicate that an intelligent choice of the size of quantization levels and the degree of protection of the bits passing through the channel, increases radically the system performance in noisy channels. However, they note that their approach is suboptimal in the sense that each distortion (of the quantizer and the channel) is separately optimized.

In [44] a tree-structured vector quantizer (TSVQ) designed for a noiseless channel is followed by a rate-compatible punctured convolutional (RCPC) code, with the bit allocation and channel code optimized via an exhaustive search over all possible channel code choices. They expressed the end-to-end expected distortion as the sum of the source coder distortion and the channel distortion under the assumptions that the channel errors are independent from the source and that the source encoder is designed to meet the centroid condition. The source distortion is precomputed for the chosen class of source encoders and in order to compute the channel distortion they model the combination of the convolutional channel encoder, channel and channel decoder as a modified Gilbert noise channel. In the case of a wireless channel they model the channel by a finite state channel which is a varying binary symmetric channel (BSC) with crossover probabilities determined by a multi-state Markov process.

In [88] for each subband, the source coding rate as well as the level of protection quantized by the channel coding rate, are jointly chosen to minimize the total end-to-end mean squared distortion suffered by the source. The optimal choice of source and channel coding rates depends on the state of the physical channel. These results are extended to transmission over fading channels using a finite state model where every state corresponds to an AWGN channel with a suitable noise variance. A spatio-temporal subband decomposition followed by vector quantization of the subband coefficients forms the source coding approach. The VQ indexes of each coded subband are interleaved and protected using RCPC codes. Interleaving aids in the analytical computation of the channel-induced distortion by making the equivalent channel memoryless. The optimal allocation of source and channel coding rates is stated as a constrained optimization problem.

The joint source-channel approach considered is to choose the set of source coding rates $\mathbf{R}_s = \{R_{s,1}, \dots, R_{s,K}\}$ and the set of channel coding rates $\mathbf{R}_c = \{R_{c,1}, \dots, R_{c,K}\}$ in

order to minimize the overall distortion in the source subject to an overall rate constraint. M is the number of the subbands. The overall distortion is a function of the source and channel coding rates and can be expressed as : $D(\mathbf{R}_s, \mathbf{R}_c) = \sum_{i=1}^K f_i D_i(R_{s,i}, R_{c,i})$ where f_i is the fraction of coefficients in subband i and $D_i(R_{s,i}, R_{c,i})$ is the total distortion per sample of subband i depending on the source and channel coding rates for that subband.

Thus, a general form of the minimization problem can be stated as :

$$\min_{\{(R_{s,i}, R_{c,i}) \in R_{s,i} \times R_{c,i}, i=1, \dots, K\}} \sum_{i=1}^K f_i D_i(R_{s,i}, R_{c,i})$$

subject to

$$\sum_{i=1}^K f_i \frac{R_{s,i}}{R_{c,i}} \leq R_{target}$$

It is not necessary to compute the functions $D_i(R_{s,i}, R_{c,i})$ directly under the conditions that the source coder satisfies the centroid condition and that the channel errors are independent of the source codewords. Then, the overall distortion is decomposed into the sum of source-coding and channel induced distortions :

$$D_i(R_{s,i}, R_{c,i}) = D_{s,i}(R_{s,i}) + D_{c,i}(R_{s,i}, R_{c,i})$$

$D_{s,i}(R_{s,i})$ is the operational source rate-distortion curve for subband i and $D_{c,i}(R_{s,i}, R_{c,i})$ in an operational rate-distortion curve reflecting channel distortion in subband i as a function of the channel coding rate and the source codebook. The $D_{s,i}(R_{s,i})$ is computed during optimization of the source quantizers using training data. If we assume that the channel can be modeled by a binary symmetric channel of crossover probability ϵ then the channel induced distortion can be expressed as $D_{c,i}(R_{s,i}, R_{c,i}) = \sum_u \sum_v p(u)p(v/u)d(u, v)$, $u, v \in 1, 2, \dots, C$ where C is the cardinality of the codebook and $d(u, v)$ is the per sample distortion between the source samples corresponding to code vectors with indexes u and v . The probability $p(u)$ is computed from the source statistics when the source codebook is generated. The transition probability $p(v/u)$ is a function of the channel state, the channel coding rate $R_{c,i}$ and the convolutional code itself. By this way a composite rate-distortion curve for the i th subband is derived taking into consideration both source and channel distortions. Let us denote this curve by $D_i(R_i)$ where R_i is the total coding rate per sample of the subband i . In practice, $D_i(R_i)$ is computed as the lower hull of all the values of $D_i(R_{s,i}, R_{c,i})$ obtained by sweeping over $(R_{s,i}, R_{c,i})$ which means that :

$$D_i(R_i) = \min_{\{(R_{s,i}, R_{c,i}) \in R_{s,i} \times R_{c,i} : R_{s,i}/R_{c,i} \leq R_i\}} D_i(R_{s,i}, R_{c,i})$$

In order to optimally allocate the source and channel coding rates for the subbands, they consider the ensemble of composite source-channel rate-distortion curves of all subbands and choose an operating point $\{R_1, R_2, \dots, R_K\}$ that results in minimum $\sum_{i=1}^K f_i D_i(R_i)$. Corresponding to this constrained minimization problem, an unconstrained minimization problem can be stated as :

$$\min_{\mathbf{R} \in \prod_{i=1}^K R_i} \{D_i(R_i) + \lambda R_i\}$$

where r_i is the support of $D(R_i) : \{(R_{s,i}, R_{c,i}) \in R_{s,i} \times R_{c,i} : R_{s,i}/R_{c,i} \leq R_i\}$. Let $\mathbf{R}^*(\lambda)$ be the vector of total coding rates that solve the above unconstrained problem. This solution is the solution to the original constrained problem if $\sum_{i=1}^K f_i R^*(\lambda) = R_{target}$.

The optimal solution depends on the SNR of the physical channel and is estimated by the receiver and relayed to the transmitter. The transmitter has to choose the appropriate allocation based on this information and encode the source.

In [11] and few years later in [12] a 3D subband coding using multirate quantization and bit-sensitivity approach over noisy channels is considered. An analytical expression for the end-to-end transmission distortion is found for the case of a scalable subband coding scheme with rate-compatible punctured convolutional (RCPC) codes. The source coder consists of spatio temporal subbands obtained by a 3D subband coding which are successively refined via layered quantization techniques and finally coded by a conditional arithmetic coding. In the case of channel errors unequal error protection (UEP) of the source bits is applied using RCPC codes. The problem of optimum partitioning of source and channel coding bits is based on two assumptions. First, all bits within the same quantization layer must receive the same level of protection and second, higher quantization layers never receive more protection than the lower ones. An unconstrained optimization is performed based on Lagrange multipliers. The optimization problem, in these works, is based on two constraints.

$$\min_{n_k, m_k} D = \sum_{k=1}^K d_k(n_k, m_k)$$

subject to

$$\sum_{k=1}^K n_k \leq R_s^{target} \quad \text{and} \quad \sum_{k=1}^K m_k \leq B - R_s^{target} = R_c^{target}$$

where $d_k(n_k, m_k)$ are the subband distortion functions, m_k is the distribution of channel bits used to protect n_k source bits in subband k and B is the total bit budget. R_s^{target} and R_c^{target} is the target source and channel rate respectively. Thus, the corresponding Lagrangian problem is :

$$\min_{n_k, m_k} \sum_{k=1}^K [d_k(n_k, m_k) + \lambda n_k + \mu m_k]$$

If there exist multipliers λ and μ such that the source and channel budgets are satisfied with equality, then the optimal solution to the Lagrangian problem is also the optimal solution to the original problem. In [12] an extended analysis of the estimation of the Lagrange multipliers is presented. The algorithm is based on the one presented in [83] with an extension of the notion of singular points to two dimensions. It is shown that when the error probability of the channel increases, the total number of quantization layers selected decreases. In addition, in good channel conditions the high frequency layers are mostly dropped due to the low error sensitivities of high frequency components. On the other hand, in poor channel conditions the number of layers of low frequency subbands is reduced. Finally, it is shown that the above optimized codec with UEP offered by the RCPC outperforms the case where equal error protection is used.

In [10], a method for optimal rate allocation for standard video encoders which is based upon the assumption of dependence between the video frames is presented. To limit the complexity of this otherwise difficult problem models of the operational distortion-rate characteristics as well as models of the channel code BER as a function of the bandwidth extension are proposed. A *H.263* coded video segment is used with one *I* frame and one or more *P* frames. The compressed video is channel coded using rate-compatible punctured systematic recursive convolutional (RCPSRC) codes. The work is focused on a frame-level rate allocation and a single quantization parameter per frame is selected. Thus, $\mathbf{q} = \{q_1, q_2, \dots, q_K\}$ is the quantization vector for an *K*-frame GOP where $q_i \in \{1, 2, \dots, 31\}$ is the quantization parameter for the *i*th frame. Similarly, the channel encoder assigns a selected channel coding rate r_i to each frame. Then $\mathbf{r} = \{r_1, r_2, \dots, r_K\}$ denotes the channel code rate selection for the *K*-frame sequence. The formulation of the general problem of the minimization can be stated as :

Given a set Q of admissible quantizers and a set R of admissible channel coding rates, find $\mathbf{q}^*$ ($q_i^* \in Q$) and $\mathbf{r}^*$ ($r_i^* \in R$) for $i = 1, \dots, K$, such that :

$$(\mathbf{q}^*, \mathbf{r}^*) = \underset{\mathbf{q}, \mathbf{r}}{\operatorname{argmin}} D_{S+C}(\mathbf{q}, \mathbf{r})$$

subject to

$$\frac{1}{K} \sum_{i=1}^K \frac{R_s^{(i)}(q_i)}{r_i} \leq \frac{1}{N} R_{S+C}$$

where $D_{S+C}(\mathbf{q}, \mathbf{r})$ is the average overall *K*-frame sequence distortion, R_{S+C} is the overall rate constraint and $R_s^{(i)}(q_i)$ is the source rate function for the frame *i*.

In order to perform an optimal rate allocation through the above minimization, a model for the RCPSRC codes is given and it is shown that the model closely fits simulated performance for a range of channel coding rates. In addition, they introduce a two-parameter polynomial model to compute the distortion for the *k*th distribution to the *i*th frame where $1 \leq k \leq i$. In general, the results showed that more rate should be allocated to earlier frames in a sequence than to later frames. For a fixed source coding rate, it was shown that there is a significant advantage to unequal error protection allowing variable channel coding rates between frames in a sequence.

In [96] the authors cascade the 3D-SPIHT [45] video coder with rate-compatible punctured convolutional codes in combination with an automatic repeat request (ARQ) strategy for transmission over a BSC. The 3D-SPIHT bitstream is first partitioned into blocks of equal length and then each block is passed through a cyclic redundancy code (CRC) parity checker to generate *c* parity bits. Next, *m* bits are padded at the end of each block to flush the memory of the RCPC coder. The final block is encoded by the RCPC code before the transmission through the channel. The Viterbi decoder searches the "best path" which satisfies both the lowest path metric and the check sum equations. When the check bits indicate an error in the block the decoder usually fixes it by finding the path with the next lowest metric. If the decoder fails to decode the received block within a certain depth of trellis depth, a technique of ARQ is used which send a negative acknowledgment to the transmitter, requesting re-transmission of the same block. As the main default of an ARQ technique is the transmission delay incompatible for example with wireless and

real time communications and the bandwidth of the feedback channel, in this work the ARQ is limited to one stage.

In the same work [96], for a total transmission rate of 2.53Mbps the authors, using equal error protection (EEP), they add a RCPC of rate $r = 2/7$ for the transmission over a BSC of crossover probability $p = 10^{-1}$, a RCPC of rate $r = 2/3$ over a BSC with $p = 10^{-2}$ and a RCPC of rate $r = 8/9$ over a BSC with $p = 10^{-3}$. They notice that without ARQ they often encountered incomplete decoding and the average PSNR obtained depends on where the incomplete decoding occurs in the 3D-SPIHT bitstream. In the case of the "Football" sequence over a BSC of $p = 10^{-2}$ the average PSNR is 7.6 dB, higher using ARQ and for the same sequence over a BSC of $p = 10^{-3}$, it is 4.6 dB higher. However, we remark from the experiments reported in the same work that the average PSNR over a BSC of $p = 10^{-3}$ using ARQ is only 0.7 dB higher than the average PSNR over a BSC of $p = 10^{-2}$ using ARQ too.

In [14] the authors divide the 3D wavelet coefficients into several groups, according to their spatial and temporal relationships and then encode each group independently using the 3D-SPIHT algorithm. They use the same procedure of channel coding as in [96], without using the ARQ method. By coding the wavelet coefficients into multiple and independent bitstreams, any single bit error affects only one of these streams, while the others are received unaffected. Thus, in this work, when the decoder fails to decode the received packet within a certain depth of the trellis depth, it stops decoding that stream. The decoding procedure continues until either the final packet has arrived or a decoding failure has occurred in all sub-streams. When they compare their scheme, where no ARQ is used, with the 3D SPIHT/RCPC+ARQ [96] the average PSNR is just 1 – 2 dB lower.

In [99] an Unequal Error Protection (UEP) of a 3D-SPIHT bitstream is proposed. The channel codes used on unequally long segments of the bitstream are the ones also used in [96]. Smaller rates of RCPC code are used for the more sensitive and significant sub-blocks to provide a better error protection, while larger rates of RCPC code are used for insensitive and insignificant sub-block. The authors notice that when the noise exceeds a certain level for a fixed rate of the RCPC code, the performance of Equal Error Protection scheme will deteriorate immediately while the UEP scheme provide a better trade-off between the coder's performance and a very low SNR.

4.4 Channel-optimized source coding

In channel-optimized source coding, the codebook is designed by minimizing the distortion criterion including channel errors. Therefore, the codebook is designed for a specific channel.

An early attempt in this area was done by Farvardin and Vaishampayan [22] where instead of considering the quantizer and the channel encoder separately, they concentrated on designing an *encoder* whose input is the source output and whose output is the channel input. First for a fixed decoder they develop necessary conditions for the optimality of the encoder function and then fixing the encoder function they develop necessary conditions for the optimality of the decoder. The pair of the conditions can be used to establish a set of conditions for the entire system. For a fixed decoder, each input is classified into the

cell with the least expectation of distortion and for a fixed encoder, the optimal decoder is described by the conditional expectation of the input given the channel output. But this set of conditions need not satisfy the sufficient conditions for the system's optimality. The final solution is only a locally optimum solution. An interesting point of their algorithm is that when the channel is very noisy in certain cases the total number of quantization regions should be smaller than the total number of available codewords. However a whole new codebook might have to be designed in order to optimally accommodate each new channel condition.

A generalization of this algorithm for vector quantizers comes from the same authors in [21]. A vector quantizer has been optimized for a given set of crossover probabilities of the codeword indices. These codeword indices are generally mapped to binary strings and the crossover probabilities are then functions of the channel's bit error probability. Let us consider a d -dimensional, M -level VQ and a discrete memoryless channel with input and output alphabets $I = \{1, 2, \dots, M\}$. Consider an encoder mapping $\gamma : \mathbb{R}^d \rightarrow I$ which is described in terms of a partition $L = \{S_1, S_2, \dots, S_M\}$ of $\mathbb{R}^d$ according to $\gamma(\mathbf{x}) = i$, if $\mathbf{x} \in S_i$, $i \in I$, where $\mathbf{x}$ is a d -dimensional source output vector. The channel index assignment is a one-to-one mapping $b : I \rightarrow I$, which assigns to the encoder output i an index $i' = b(i) \in I$, which is then delivered to the channel. The channel is described by the probability $P(j|i')$ which denotes the probability that the index j is received given that i' is transmitted. Finally, the decoder mapping $g : I \rightarrow \mathbb{R}^d$ is described in terms of a finite reproduction codebook $C = \{c_1, c_2, \dots, c_M\}$ according to $g(j) = c_j$, $j \in I$. Then the average distortion per sample is given by :

$$D(L, C, b) = \frac{1}{d} \sum_{i=1}^M \int_{S_i} p(\mathbf{x}) \left\{ \sum_{j=1}^M P(j|b(i)) d(\mathbf{x}, c_j) \right\}$$

where $p(\mathbf{x})$ is the d -fold probability density function of the source and $d(\mathbf{x}, \mathbf{y})$ is a non-negative distortion measure, which in the case of a squared error criterion is equal to $d(\mathbf{x}, \mathbf{y}) = \|\mathbf{x} - \mathbf{y}\|^2$. For a fixed b and a fixed C the minimization of the average distortion leads to the optimum partition $L^* = \{S_1^*, S_2^*, \dots, S_M^*\}$ such that :

$$S_i^* = \left\{ \mathbf{x} : \sum_{j=1}^M P(j|b(i)) \|\mathbf{x} - \mathbf{c}_j\|^2 \leq \sum_{j=1}^M P(j|b(l)) \|\mathbf{x} - \mathbf{c}_j\|^2 \quad \forall l \right\}, \quad i \in I \quad (4.2)$$

For a fixed b and a fixed L the optimum codebook $C^* = \{c_1^*, c_2^*, \dots, c_M^*\}$ must satisfy :

$$c_j^* = \frac{\sum_{i=1}^M P(j|b(i)) \int_{S_i} \mathbf{x} p(\mathbf{x})}{\sum_{i=1}^M P(j|b(i)) \int_{S_i} p(\mathbf{x})} \quad (4.3)$$

The successive application of the equations (4.2) and (4.3) results in a sequence of encoder-decoder pairs for which the corresponding average distortions form a nonincreasing sequence of nonnegative numbers which has to converge. The result of this process is locally optimum as in the scalar case presented above. This kind of system is usually called in the literature a channel-optimized vector quantizer (COVQ). For the case of VQ, as for the scalar case, when the channel is noisy, the number of nonempty encoding regions associated with the encoder-decoder pair may turn to be smaller than the cardinality of the

codebook M . If we denote by y_i the nonempty S_i , where $y_i = \sum_{j=1}^M P(j|b(i))c_j$, in [21] it is shown that the hyperplane H_{il} which separates the regions S_i^* and S_l^* is not necessarily the perpendicular bisector of the cord connecting c_i and c_j , as for a VQ developed for a noiseless channel case, but it is perpendicular to (but not necessarily the bisector of) the cord connecting y_i and y_l .

In [19] it is noticed that while index assignment would be unnecessary if the COVQ achieves the globally optimal source code design, attention to the index assignment problem in practical COVQ design does yield some performance gains by finding a better local optimum.

A joint source channel codec for images using trellis coded quantization and convolutional codes is presented in [6]. Their approach is based on the joint trellis quantization/modulation (TCQ/TCM) of [23] in which the same Ungerboeck trellis is used for the TCQ and TCM. The TCM is done by applying the same Ungerboeck trellis used for TCQ source coding and for soft-decision convolutional channel coding. Methods to ensure the best possible mapping between Euclidean distance distortion in the source codec and Hamming distance in the channel codec are also presented.

4.5 Other combined optimizations

Another interesting approach is the joint design of channel-optimized and Multicarrier Modulation (MCM) schemes. MCM provides the flexibility to change the power, the modulation and the channel encoding of each individual subchannel, so that different degrees of error protection may be provided for different bits according to their protection. In [42] the problem of combined quantizer and modulation design when the quantizer outputs are to be transmitted via noisy channels, is studied. Their iterative algorithm has three steps. First, for a fixed power allocation and a fixed decoder they optimize the encoder (the quantizer and the mapping). Second, for a fixed encoder and fixed power allocation they optimize the decoder and third, they allocate the power of the MCM system. This approach has one step more than those in the COVQ presented above and determines the optimal combined source-channel code for fixed overall power and number of bits. However this algorithm depends on the initial conditions and can converge to a locally optimum solution. Their results show an improvement on a AWGN channel over the Lloyd scalar quantizer, LGB vector quantizer and the COVQ transmitted over single-carrier or multicarrier channels, especially under very noisy conditions.

Another joint design of a VQ and a multidimensional constellation, under an average power constraint is presented in [29]. However, the authors notice when the resulting constellations in each modulation interval are constrained to be binary, the problem to solve becomes more practical and simple. They treat the problem of joint optimization of the transmission energy allocation (TEA) and index assignment. Their iterative optimization algorithm starts with an initial index assignment, then evaluates the sensitivities of the various bits and applies the energy allocation algorithm and reassigns indexes to the codevectors via binary switching which has been modified to include the effect of unequal bit error rates. Application of this algorithm to a source-optimized VQ has shown that the solution is susceptible to poor local minima as it heavily depends on the initial

index assignment. Their solution to this latter problem is to design a channel-optimized VQ for high level of channel noise and then gradually reduce the level of channel noise assumed for the design. This technique is called *noisy channel relaxation* (NCR). Initializing the iterations at a very high level of channel noise makes it easier for the system to find a good initial IA which is then tracked and reoptimized as the noise level is reduced. The final iterations are performed for a noiseless channel, thereby yielding a source optimized VQ. When the information about the state of the channel is available at the receiver they show that the overall performance can be further improved by developing a switched-encoder-based system capable of adapting itself to time-varying characteristics of the channel.

In [35] three techniques of combined source-channel coding are presented. The first is a channel-optimized source code : COVQ optimized for a given equal-error-protection convolutional code with soft decision decoding. The second is a source-optimized channel code : RCPC channel coding optimized for a VQ design based on a noiseless channel. Finally the third technique is an iterative algorithm combining the design strategies of the previous two. It is shown that all these joint design algorithms surpass the performance of their corresponding predecessors by optimizing the bit allocation between the source and channel codes for the given signal-to-noise ratio (SNR). This optimal allocation depends on the source statistics, the channel quality and the modulation. The distortions of all the three techniques are approximately the same because the optimal bit allocation for each code design results in a channel code which removes most of the channel errors. In addition, the three techniques have the same sensitivity to channel mismatch and the third one is less sensitive than the others to a suboptimal bit allocation. This robustness is achieved at the expense of higher off-line design complexity.

4.6 Conclusion

In this Chapter, we presented the different components of a transmission system and the roles that each one has in a joint source-channel coding scheme. The main goal of a joint source-channel scheme is to efficiently combine some of these components in order to reduce the complexity and approach the performance of a tandem scheme.

We briefly introduced the existing approaches and classified them into three main categories ; a first one developing a source-optimized channel coding, a second one developing a channel-optimized source coding and finally a third category, combining one of them with the optimization of other components of the transmission chain. The presentation was not intended to be exhaustive. Its main goal was to give a general panorama of joint source-channel coding field and to classify the large number of proposed methods, in the literature, into three main approaches.

In the next Chapter, we will present our main contribution in this domain, which is a joint source-channel coding scheme developed for iid Gaussian sources. This scheme is closer to the idea of a channel-constrained VQ with the difference that the codebook is a structured one, resulting from lattice constellations which in the same time minimize the source distortion, too.

Chapitre 5

Joint Source-Channel coding based on linear labelling and rotated constellations

In this Chapter, we present a Joint Source-Channel coding scheme developed for Gaussian sources which minimizes simultaneously the channel distortion due to linear labelling and the source distortion due to linear transforms resulting from rotated $\mathbb{Z}^n$ constellations. We introduce the "maximum component diversity" constellations and study their performances when used in conjunction with the sphere decoder to transmit over a Rayleigh channel.

5.1 Introduction

As we have already mentioned, the basic idea in Joint Source-Channel (JSC) coding is that the source and channel processing are jointly optimized. A traditional approach to the JSC coding is to follow the source optimization point of view. First the source codebook is matched to the source statistics in order to minimize the distortion due to the source. Then, the labelling is optimized in order to minimize the distortion due to the channel. These methods use algorithms based on stochastic optimization which limit the size of the source codebook due to complexity constraints and the channel distortion remains important.

A different and more interesting approach [5] follows the channel point of view. First the channel distortion is minimized and then the algorithm is tuned to minimize also the source distortion. This technique is based on a set of linear transforms which minimize the channel distortion along with the distortion of the Gaussian sources. It was proved [46] that the channel distortion on a binary symmetric channel is minimized by a linear labelling. Thus, the set of the linear transforms in [5] minimizes the channel distortion and at the same time it builds a source codebook which mimics the source distribution in order to minimize the source distortion, too.

In the sequel we present the way that both optimizations are conducted.

5.2 Linear Labelling and Joint Source-Channel coding

5.2.1 Linear labelling to minimize channel distortion

Let a d -dimensional vector $\mathbf{x}$ be the input of a vector quantizer, producing an n -bit binary codeword i , which is the index of the vector used for signal reconstruction at the receiver. The set of $M = 2^n$ codewords defines the codebook. After transmission through the channel, the decoder receives a codeword j and produces a d -dimensional vector $\mathbf{y}_j$. The encoder is said to be a *maximum-entropy encoder* (MEE) when all codewords are used with equal probability [46]. This encoder facilitates the distortion analysis but represents a difficult task for index assignment, as we cannot focus on finding a good index assignment for just a subset of frequently used codewords.

The total distortion D can be expressed as :

$$D = D_s + D_c + D_{mixed}$$

where D_s is the source distortion due to the quantizing effects only, and D_c is the distortion caused by misinterpretation of the received codeword and hence, D_c is dependent on the index assignment. The mixed term is zero for optimum codebooks designed for noiseless channels since then the VQ points are the centroids of the reconstruction cells [20]. This term is also zero when the size of codebook approaches infinity even when the VQ points are not positioned in the centroids.

When the codevectors are not the centroids of their respective encoding regions, which means that the structure of the quantizer dictates the placement of the codevectors and the encoding regions are chosen to satisfy the Nearest Neighbor Condition, then the Minkowski inequality can be used to bound the distortion [59] as :

$$D \leq (\sqrt{D_s} + \sqrt{D_c})^2$$

However, in general the D_{mixed} is supposed equal to zero as the error in the assumption is small [39].

The distribution of the VQ points in signal space depends on the source distribution and in the case of a COVQ also on how likely indices are to be confused with other indices when transmitted on a noisy channel. The probability of confusing codewords is again dependent on the indices given to the VQ points and is best expressed by the Hamming distance between indices. A way of combining the logical description (i.e. indices) and the physical description (i.e. position in the signal space) is to move the logical hypercube, spanned by the codebook's binary codewords, to the signal space with preserved topology. The M VQ points can then be described by a transform, or a mapping, of the M vertices of the hypercube. In [46] it is shown that in the case of a binary symmetric channel the Hadamard transform is suitable for this task.

The source codebook $\mathbf{Y} = [\mathbf{y}_0 \ \mathbf{y}_1 \ \cdots \ \mathbf{y}_{M-1}]$ can be viewed as a function of $(b_1 \dots b_n) \in \{+1, -1\}^n$ representing the index assignment. In [46] it is shown that each coordinate of the $\mathbf{y}_i$ can be expressed as a discrete Volterra series and so that the full codebook $\mathbf{Y}$ can be represented as :

$$\mathbf{Y} = \mathbf{T}\mathbf{H} \tag{5.1}$$

where the $M \times M$ matrix $\mathbf{H}$ is a normalized Sylvester-style Hadamard matrix and $\mathbf{T} = [\mathbf{t}_0 \ \mathbf{t}_1 \ \cdots \ \mathbf{t}_{M-1}]$ is the Hadamard transform of $\mathbf{Y}$.

The $\mathbf{t}_0$ is called the codebook offset, the $\mathbf{t}_1, \mathbf{t}_2, \mathbf{t}_4, \dots, \mathbf{t}_{M/2}$ form the linear part of the Volterra series and the remaining vectors the nonlinear part. A codebook is linear when its nonlinear part is zero.

A Hadamard matrix is a symmetrical matrix with orthogonal rows and columns and so $\mathbf{H}$ is its own inverse. The inverse Hadamard transform is thus :

$$\mathbf{T} = \frac{1}{M} \mathbf{Y} \cdot \mathbf{H}$$

which shows that there exists a matrix $\mathbf{T}$ for an arbitrary codebook $\mathbf{Y}$.

The channel distortion can be expressed by means of the norms of the vectors $\mathbf{t}$. For an M-point MEE, this channel distortion due to noise on a BSC with bit-error probability q is :

$$D_c = 2\sigma_{VQ}^2 - 2 \sum_{i=1}^{M-1} (1 - 2q)^{w_h(i)} \cdot \|\mathbf{t}_i\|^2$$

where $w_h(i)$ is the Hamming weight of the base 2 representation of i and σ_{VQ}^2 is the variance of the source codebook, which can be expressed in terms of the transform vectors

$$\text{as : } \sigma_{VQ}^2 = \sum_{i=1}^{M-1} \|\mathbf{t}_i\|^2.$$

In [46], the authors employ a measure on how dominant the linear part of the general nonlinear transform in Eq. (5.1) is. They define a linearity index λ as :

$$\lambda = \frac{1}{\sigma^2} \sum_{l=0}^{n-1} \|\mathbf{t}_{2^l}\|^2$$

The range of λ is $0 \leq \lambda \leq 1$ where $\lambda = 1$ denotes a purely linear transform. The linear index is independent of the channel bit-error probability q .

In the same work [46], it is shown that on binary symmetric channels, the channel distortion D_c is minimized if there exists an index assignment giving $\lambda = 1$. In that case the channel distortion is given by :

$$D_c = 4q\sigma_{VQ}^2$$

where σ_{VQ}^2 is the variance of the source codebook and q the error probability.

Thus, the *labelling should be linear in order to minimize the channel distortion*.

In joint source-channel decoding, the receiver aims at minimizing the channel distortion D_c rather than the error probability, which means that the detector has to associate to the received codeword a point in the source space which minimizes D_c .

Following the data flow presented in Fig. 5.1 and assuming the case of linear labelling, we denote by $\mathbf{P}$ the matrix representing the inverse of the linear labelling l . Then, at the transmitter side we have :

$$\mathbf{y}_j = \mathbf{P}\mathbf{b}_j$$

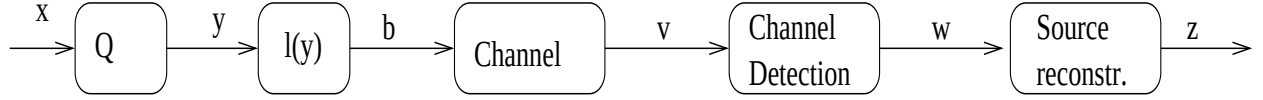


FIG. 5.1 – System description

and at the receiver side :

$$\mathbf{z}_j = \mathbf{P}\mathbf{w}_j$$

Assuming a maxentropic quantizer, the channel distortion as a function of the channel codebook is given by :

$$D_c = \frac{1}{M} \sum_{j=0}^{M-1} E(\|\mathbf{z}_j - \mathbf{y}_j\|^2) = \frac{1}{M} \sum_{j=0}^{M-1} E(\|\mathbf{w}_j - \mathbf{b}_j\|_{\mathbf{S}}^2)$$

where $\|\cdot\|_{\mathbf{S}}^2$ is the Euclidean norm defined by the symmetric matrix $\mathbf{S} = \mathbf{P}^t\mathbf{P}$. The expression derived above depends on the channel itself and on the detection criterion.

Case of the Gaussian channel with binary inputs

In the case of hard decision detection, the Gaussian channel with binary inputs is transformed into a binary symmetric channel (BSC) with probability $Q(1/\sigma_b)$, where σ_b^2 is the variance of the Gaussian noise. Thus, the channel distortion is :

$$D_c = 4 \cdot Q(1/\sigma_b) \cdot \sigma_{VQ}^2$$

where σ_{VQ}^2 is the variance of the quantized source which under these assumptions is almost equal to the source variance and the function $Q(x)$ can be related to the complement of the error function $erfc$ by the following expression : $Q(x) = \frac{1}{2}erfc(\frac{x}{\sqrt{2}})$.

5.2.2 Minimisation of source distortion-case of Gaussian sources

The next step is the construction of a linear labelling that minimizes the source distortion D_s . In [5] the investigation of this step is developed for memoryless zero-mean Gaussian sources with variance one. The linear labelling represented by the matrix $\mathbf{G}_{d,n}$ of size $(d \times n)$ must transform identically distributed random variables into random variables (the source codebook) which have to mimic the source distribution.

Let $\mathbf{U}_n$ be an $n \times n$ matrix. The rows and the columns of $\mathbf{U}_n$ are denoted by L_i and C_j respectively, where $1 \leq i, j \leq n$. If J is some subset of $\{1, \dots, n\}$, then $C_j(J)$ is the truncation of the j -th column of $\mathbf{U}_n$ according to the subset of indices J . By means of $\mathbf{U}_n$, one can linearly map $BPSK_n = \{-1, +1\}^n$ onto a new set $\mathbf{U}_n \cdot BPSK_n$. Allowing j to increase while J remains fixed, we get a codebook $S_n(J)$ with codevectors :

$$\mathbf{y} = \sum_{j=1}^n C_j(J)b_j,$$

where $\mathbf{b} = (b_1, \dots, b_n)^t \in BPSK_n$.

If the subset J contains for example d elements, then the components of a d -dimensional codevector $\mathbf{y}$ in terms of $\mathbf{G}_{d,n}$ can be expressed as :

$$y_i = \sum_{j=1}^d \sum_{k=1}^n G_{i,j} b_k,$$

where $\mathbf{b} = (b_1, \dots, b_n)^t \in BPSK_n$.

In [5] it is shown that in order to insure the Gaussianity of $\mathbf{y}$ it is necessary that all the components $G_{i,j} b_j$ of the summation to be nonzero, for any vector $\mathbf{b}$. Let b_j be independent and uniformly distributed. Then the d -dimensional variables $\xi_j = C_j(J) b_j$ are independent and zero mean with covariance matrices $\mathbf{\Gamma}_j$ which can be expressed as :

$$\mathbf{\Gamma}_j = \sigma^2 C_j(J) C_j^T(J)$$

where σ^2 is the variance of b_j . Then,

$$\sum_{j=1}^n \mathbf{\Gamma}_j = \sum_{j=1}^n \sigma^2 C_j(J) C_j^T(J) = \text{Cov}(\mathbf{y}) = \sigma^2 \mathbf{U}(J) \mathbf{U}^T(J)$$

Suppose that $\mathbf{\Gamma}$ is a $d \times d$ matrix, symmetric and positive definite, and

1. $\forall \epsilon > 0 \quad \sum_{j=1}^n Pr(\|\xi_j\| > \epsilon) \rightarrow 0$
2. $\sum_{j=1}^n \mathbf{\Gamma}_j \rightarrow \mathbf{\Gamma}$

Then, $\mathbf{y} \sim \mathcal{N}(0, \mathbf{\Gamma})$ in law when $n \rightarrow \infty$.

In [5] it is shown that if $\mathbf{U}_n$ is orthogonal, with coefficients going uniformly to 0 as $n \rightarrow \infty$ then the two above conditions are satisfied.

Then the mapping

$$\mathbf{b} \in BPSK_n \rightarrow (\mathbf{G}_{d,n} \mathbf{b} \in S_n(J)),$$

where $\mathbf{G}_{d,n} = \mathbf{U}_n(J)$ is linear, and allows building an asymptotically Gaussian source dictionary which minimizes the D_s as $n \rightarrow \infty$.

This property can be obtained with "maximum component diversity" constellations. Thus, if the $\mathbf{U}_n$ is a $n \times n$ matrix of a "maximum component diversity" constellation, then the $\mathbf{G}_{d,n}$ matrix is constructed by taking any set of d rows of $\mathbf{U}_n$. Similar properties are achieved when selecting columns of the matrix $\mathbf{U}_n$ and we shall denote by $\mathbf{G}'_{n,r}$ the $n \times r$ ($r \leq n$) matrices constructed this way.

In the sequel, we present the basic properties of the maximum component diversity constellations and the different constructions that have been proposed.

5.3 Maximum component diversity constellation

Maximum component diversity constellations or full diversity rotations have been studied in the context of fading channels [34]. This is motivated by the fact that on fading channels, when the channel is in deep fade, errors occur systematically, and the

information is lost for the receiver. The common way to avoid it is to provide the receiver with several replicas of the information. These replicas are supposed to be subjected to independent fadings and the combination of them can restore the information. Besides the common diversity techniques such as time and frequency diversity, we can also speak of *modulation diversity* [8] techniques when special multidimensional signal constellations having lattice structure are used. They provide the receiver with an order of diversity dependent on the number of dimensions of the signal constellation. The diversity order of a multidimensional signal set is the minimum number of distinct components between any two constellation points.

Given a $\mathbb{Z}^n$ -lattice constellation, the desired modulation diversity is obtained by applying a suitable rotation using algebraic number-theoretical tools presented in the Appendix A. An interesting feature of the rotation operation is that the rotated signal set has exactly the same performance as the nonrotated one when used over a pure additive white Gaussian noise channel (AWGN). Therefore, the interest of this construction will become obvious in the case of transmissions over Rayleigh channels, as we shall present in Chapter 7.

In a fading channel, we desire the maximum component diversity, which means the maximum number of distinct components between two points in a signal constellation. In [34], [9] it is shown that lattices constructed by the canonical embeddings of a totally real algebraic number field result in the maximum diversity $L = n$, equal to the dimension of the lattice constellation. The maximum component diversity constellations used in our JSC coding scheme, in order to minimize the source distortion, are based on the theory of the ideal lattices resulting from the maximal real subfield of a cyclotomic field. In Appendix A we present some algebraic number-theoretical tools allowing to construct ideal lattices.

5.4 Construction of rotated $\mathbb{Z}^n$ -lattices of dimension $n = (p - 1)/2$ and mixture constructions

In [4] we find a construction of rotated $\mathbb{Z}^n$ lattices on the ring of integers of the maximal real subfield of a cyclotomic field. This construction gives a rotated $\mathbb{Z}^n$ -lattice of dimension $n = (p - 1)/2$, where $p \geq 5$ is a prime. Thus, the rotated lattices resulting from this method can have dimension $n = 2, 3, 5, 6, 8, 9, 11, 14, 15, 18, 20, 21, \dots$

Recalling the definitions and notations introduced in the Appendix A.2, then if $a = (1 - \zeta)(1 - \zeta^{-1})$, we have that

$$\frac{1}{p} \text{Tr}(axy)$$

is isomorphic to the unit form of degree n by choosing a simple new base $\omega'_1, \dots, \omega'_n$, where $\omega'_n = \omega_n$ and $\omega'_j = \omega_j + \omega'_{j+1}$.

Considering the n field embeddings defined by :

$$\sigma_k(\omega_j) = \zeta^{kj} + \zeta^{-kj} = 2 \cos\left(\frac{2\pi kj}{p}\right)$$

then the lattice generated by the ring of integers has an $n \times n$ generator matrix M whose elements are given by : $M_{k,j} = 2 \cos(\frac{2\pi k j}{p})$

The element a is represented by the diagonal matrix : $A = \text{diag}(\sqrt{\sigma_k(a)})$.
The basis transformation matrix from $\{\omega_j\}$ to $\{\omega'_j\}$ is given by :

$$T = \begin{pmatrix} 1 & 1 & \cdots & 1 & 1 \\ 0 & 1 & 1 & \cdots & 1 \\ \vdots & & & & \vdots \\ 0 & \cdots & 0 & 1 & 1 \\ 0 & 0 & \cdots & 0 & 1 \end{pmatrix}$$

Finally, the rotated $\mathbb{Z}^n$ -lattice generator matrix is given by :

$$R^n = \frac{1}{\sqrt{p}} T M A$$

5.4.1 Mixture constructions of rotated $\mathbb{Z}^n$ -lattices

In order to build rotated $\mathbb{Z}^n$ -lattices in higher dimensions, in [4] is proposed a way of combine different generator matrices $R^{(j)}$.

Especially, let K be the compositum of N Galois extensions K_j of degree n_j (i.e. the smallest field containing all K_j) with coprime determinant, $(d_{K_i}, d_{K_j}) = 1, \forall i \neq j$. Assume that there exists an a_j such that the trace form over K_j , $\text{Tr}(a_j xy)$ is isomorphic to the unit form of degree n_j for $j = 1, \dots, N$. Then the form over K

$$\text{Tr}(a_1 xy) \otimes \cdots \otimes \text{Tr}(a_N xy)$$

is isomorphic to the unit form of degree $n = \prod_{j=1}^N n_j$.

Thus, the lattice generator matrix can be obtained as the tensor product of the generator matrices $R^{(j)}$ corresponding to the forms $\text{Tr}(a_j xy)$ for $j = 1, \dots, N$:

$$R = R^{(1)} \otimes \cdots \otimes R^{(N)}$$

Let $K = K_1 K_2$ be the composition of two Galois extensions of degree n_1 and n_2 , with coprime discriminant. The discriminant of K is $d_K = d_{K_1}^{m_1} d_{K_2}^{m_2}$, where $m_j = n/n_j, j = 1, 2$.

Then the minimum product distance is :

$$d_{p,min} = \frac{1}{\sqrt{d_{K_1}^{m_1} d_{K_2}^{m_2}}}$$

By this way the authors of [4] construct lattices in dimensions such as $n = 10, 12, 16, 22, 24, \dots$.

The dimension $n = 4$ is obtained by combining lattices coming from $K = \mathbb{Q}(\sqrt{p=5})$ and $K_1 = \mathbb{Q}(\sqrt{2})$. Note also that if $K = \mathbb{Q}(\sqrt{d})$ with a square-free positive d and $\{1, \omega\}$ the integral basis of O_K , then the corresponding generator matrix is :

$$M = \begin{pmatrix} 1 & 1 \\ \omega & \sigma(\omega) \end{pmatrix}$$

If there exists a totally positive element a such that $N(a) = 1/\det(M)^2$, we obtain the square lattice generator matrix :

$$R = M \text{diag}(\sqrt{a}, \sqrt{\sigma(a)})$$

In [4] it is shown that the existence of such an element a is guaranteed for a family of quadratic fields given by $K = \mathbb{Q}(\sqrt{p})$ for all primes p such that $p \equiv 1 \pmod{4}$. A square lattice can also be obtained from other quadratic fields such as $\mathbb{Q}(\sqrt{2})$.

5.5 Rotated $\mathbb{Z}^n$ lattices from cyclotomic fields where n is a power of 2

We present in this section a different construction of rotated $\mathbb{Z}^n$ lattices [34], which are also based on the use of ideals of rings of maximal real cyclotomic fields. They yield a diversity of order n .

Let n to be the dimension of the requested lattice. Recalling the properties of cyclotomic fields presented in the Appendix A.2.2, we are interested here in the real maximal cyclotomic field $K = \mathbb{Q}(2 \cos \frac{2\pi}{m})$ of degree $s = \frac{\phi(m)}{2} = 2n$, where $m \geq 5$ and $\phi(\cdot)$ is the Euler function. We accept only the fields for which m can be expressed as a power of 2, i.e. $m = 2^r$.

Such fields will have an integral basis :

$$\omega_0 = 1, \quad \omega_l = 2 \cos \frac{2\pi l}{m}, \quad 1 \leq l \leq s-1$$

Let $\{h_1 = 1, h_2, \dots, h_s\}$ be the integers satisfying $1 \leq h < m/2$ and which are prime with m , $\gcd(h, m) = 1$. The transformation

$$\sigma_j : \begin{cases} \omega_0 \rightarrow \omega_0 \\ \omega_l \rightarrow 2 \cos \frac{2\pi l}{m} h_j \end{cases} \quad 1 \leq l \leq s-1$$

is an automorphism of the cyclotomic real field K and the s automorphisms $\{\sigma_1 = 1, \sigma_2, \dots, \sigma_s\}$ form the Galois group Γ of K . The group Γ is isomorphic to $\mathbb{Z}_m^*/\langle -1 \rangle$ where $\mathbb{Z}_m^*$ is the multiplicative group of the invertible elements in $\mathbb{Z}_m$. When $m = 2^r$, the group $\mathbb{Z}_m^*$ is isomorphic to $\langle 5 \rangle \oplus \langle -1 \rangle$. In that case Γ is cyclic and it is isomorphic to $\langle 5 \rangle$.

The construction of the rotated $\mathbb{Z}^n$ lattice passes through the following steps :

- First Step : take the generator matrix $\Omega = \sigma_i(\omega_j)_{1 \leq i, j+1 \leq s}$ of dimension $2n$

$$\Omega = \begin{pmatrix} \sigma_1(\omega_0) & \sigma_1(\omega_1) & \cdots & \sigma_1(\omega_i) & \cdots & \sigma_1(\omega_{s-1}) \\ \sigma_2(\omega_0) & \sigma_2(\omega_1) & \cdots & \sigma_2(\omega_i) & \cdots & \sigma_2(\omega_{s-1}) \\ \vdots & \vdots & \cdots & \vdots & \cdots & \vdots \\ \sigma_s(\omega_0) & \sigma_s(\omega_1) & \cdots & \sigma_s(\omega_i) & \cdots & \sigma_s(\omega_{s-1}) \end{pmatrix}.$$

- Second Step : take the odd rows and odd columns (first row and column starting with the index 0) of this matrix and construct the final generator matrix U_n .

– Third Step : normalize $\mathbf{U}_n$ with $1/\sqrt{2^{r-2}}$

Indeed, taking the odd rows and columns, the constructed family is orthogonal and all its elements have the same norm.

Let us see now an example, in order to better understand the above method.

Let us build a rotated $\mathbb{Z}^2$ lattice. Thus, with $n = 2$ we search for the maximal real of a cyclotomic field of degree $s = 2n = 4 = \frac{\phi(m)}{2}$, $\phi(m) = 8$. The integer m whose Euler function gives 8 and is a power of 2 is $m = 16 = 2^4$. Finally we have the $K = \mathbb{Q}(2 \cos(\frac{2\pi}{16}))$ the requested real cyclotomic field, with integral basis :

$$\omega_0 = 1, \quad \omega_1 = 2 \cos \frac{2\pi}{16} = \sqrt{2 + \sqrt{2}}, \quad \omega_2 = 2 \cos \frac{2\pi}{16} 2 = \sqrt{2}, \quad \omega_3 = 2 \cos \frac{2\pi}{16} 3 = \sqrt{2 - \sqrt{2}}$$

The integers $\{h_1 = 1, h_2, \dots, h_4\}$ that satisfy $1 \leq h < 8$ and $\gcd(h, 16) = 1$ are : $h_1 = 1, h_2 = 3, h_3 = 5$ and $h_4 = 7$.

Thus, the automorphisms $\{\sigma_1, \sigma_2, \sigma_3, \sigma_4\}$ can be expressed as follows :

$$\begin{aligned} \sigma_1(\omega_0) &= \sigma_2(\omega_0) = \dots = \sigma_4(\omega_0) = 1 \\ \sigma_1(\omega_1) &= 2 \cos \frac{2\pi}{16} h_1, \sigma_1(\omega_2) = 2 \cos \frac{2 \cdot 2\pi}{16} h_1, \sigma_1(\omega_3) = 2 \cos \frac{2 \cdot 3\pi}{16} h_1 \\ \sigma_2(\omega_1) &= 2 \cos \frac{2\pi}{16} h_2, \sigma_2(\omega_2) = 2 \cos \frac{2 \cdot 2\pi}{16} h_2, \sigma_2(\omega_3) = 2 \cos \frac{2 \cdot 3\pi}{16} h_2 \\ \sigma_3(\omega_1) &= 2 \cos \frac{2\pi}{16} h_3, \sigma_3(\omega_2) = 2 \cos \frac{2 \cdot 2\pi}{16} h_3, \sigma_3(\omega_3) = 2 \cos \frac{2 \cdot 3\pi}{16} h_3 \\ \sigma_4(\omega_1) &= 2 \cos \frac{2\pi}{16} h_4, \sigma_4(\omega_2) = 2 \cos \frac{2 \cdot 2\pi}{16} h_4, \sigma_4(\omega_3) = 2 \cos \frac{2 \cdot 3\pi}{16} h_4 \end{aligned}$$

By its automorphisms we create the 4×4 matrix Ω as described in the First Step.

$$\Omega = \begin{pmatrix} 1 & \sqrt{2 + \sqrt{2}} & \sqrt{2} & \sqrt{2 - \sqrt{2}} \\ 1 & \sqrt{2 - \sqrt{2}} & -\sqrt{2} & -\sqrt{2 + \sqrt{2}} \\ 1 & -\sqrt{2 - \sqrt{2}} & -\sqrt{2} & \sqrt{2 + \sqrt{2}} \\ 1 & -\sqrt{2 + \sqrt{2}} & \sqrt{2} & -\sqrt{2 - \sqrt{2}} \end{pmatrix}.$$

By choosing the odd rows and columns and after normalization we finally have :

$$\mathbf{U}_2 = \frac{1}{2} \begin{pmatrix} \sqrt{2 - \sqrt{2}} & -\sqrt{2 + \sqrt{2}} \\ -\sqrt{2 + \sqrt{2}} & -\sqrt{2 - \sqrt{2}} \end{pmatrix}.$$

We can easily see that $\mathbf{U}_2$ is orthogonal and so **no basis reduction is necessary** in order to obtain a rotation, which makes this method simpler to use than the one presented in [4].

In this way we can directly build lattices of dimension n which is power of 2, i.e. $n = 2, 4, 8, 16, 32$, etc... The rotation matrices for the cases $n = 2, 4, 8, 16$ are provided in Appendix C.

5.6 Sphere Decoder

In order to evaluate the performances of the rotated constellations over a Rayleigh channel, we need an efficient decoding algorithm. Therefore, in this Section we briefly

present the universal lattice code sphere decoder [92] which, indeed, is a maximum-likelihood (ML) decoding algorithm. We will not enter in the details of this algorithm, we will only present its basic features. The decoder is based on a bounded distance search among the lattice points falling inside a sphere centered at the received point. This algorithm was developed for fading channels with perfectly known channel state information (CSI), but it is also applicable to the AWGN channel.

Let us assume independent fading channels with perfect CSI known to the receiver. Then the ML decoding requires the minimization of the metric :

$$m(\mathbf{x}|\mathbf{r}, \alpha) = \sum_{i=1}^n |r_i - \alpha_i x_i|^2,$$

where $\mathbf{r} = \boldsymbol{\alpha} * \mathbf{x} + \mathbf{n}$ is the received vector and $\mathbf{n} = (n_1, n_2, \dots, n_n)$ is the noise vector which has Gaussian distributed independent random variables components with zero mean and variance N_0 . If $\mathbf{x}$ is one of the transmitted lattice points, then it can be expressed as $\mathbf{x} = \mathbf{u}M$, where M is the generator matrix of the lattice and $\mathbf{u} = (u_1, u_2, \dots, u_n)$ is the integer component vector to which the information bits are mapped. The random independent fading coefficients $\boldsymbol{\alpha} = (\alpha_1, \alpha_2, \dots, \alpha_n)$ have unit second moment and $*$ denotes the component-wise product. The fading coefficients α_i are supposed independent random variables, which is achievable in practice by the use of a component interleaver.

The sphere-decoder algorithm searches through the points of the lattice Λ which are found inside a sphere of given radius $\sqrt{C}$ centered at the received point. By this way, only the lattice points within the distance $\sqrt{C}$ from the received point are considered in the metric minimization. The problem for a AWGN channel where $\alpha_i = 1, i = 1, \dots, n$ becomes therefore :

$$\min_{\mathbf{x} \in \Lambda} \|\mathbf{r} - \mathbf{x}\| = \min_{\mathbf{w} \in \mathbf{r} - \Lambda} \|\mathbf{w}\|$$

which means searching the shortest vector $\mathbf{w}$ in the translated lattice $\mathbf{r} - \Lambda$ in the n -dimensional Euclidean space R^n .

In the case of a fading channel, the situation becomes a little more complicated but can be easily adjusted to the previous minimization. Let us consider the lattice Λ_c with generator matrix $M_c = M \text{diag}(\alpha_1, \alpha_2, \dots, \alpha_n)$. In this new lattice each component has been changed by a factor α_i and thus a point can be written as $\mathbf{x}^{(c)} = (x_1^{(c)}, \dots, x_n^{(c)}) = (\alpha_1 x_1, \dots, \alpha_n x_n)$. The metric to minimize becomes :

$$m(\mathbf{x}|\mathbf{r}, \alpha) = \sum_{i=1}^n |r_i - \alpha_i x_i^{(c)}|^2.$$

Applying the decoding algorithm to the lattice Λ_c for the received point $\mathbf{r}$, we obtain the decoded point $\hat{\mathbf{x}}^{(c)} \in \Lambda_c$, which has the same integer components $(\hat{u}_1, \dots, \hat{u}_n)$ as $\hat{\mathbf{x}} \in \Lambda$. We can easily see that for each received point there is a different lattice Λ_c .

An advantage of this algorithm is that the vectors with norm greater than the given radius are never tested. Thus, the choice of the radius is crucial, especially in the fading channels where due to the deep fades many points fall inside the search sphere and the decoding can be very slow. In this case C is adapted according to the values of the fading coefficients α_i .

5.7 Performances of rotated BPSK over a Rayleigh fading channel

In the previous sections we presented different constructions of full diversity rotations of dimension n as that the minimization of the source distortion in our joint source-channel coding scheme is based on these linear transforms, as explained in Section 5.2.

However, before introducing our scheme, we present the performances of these rotated constellations over a Rayleigh fading channel. The original constellation is a BPSK. We compare the performances of its transmission over a Rayleigh channel with and without rotation. Let us suppose the transmission system presented in the Fig.5.2.

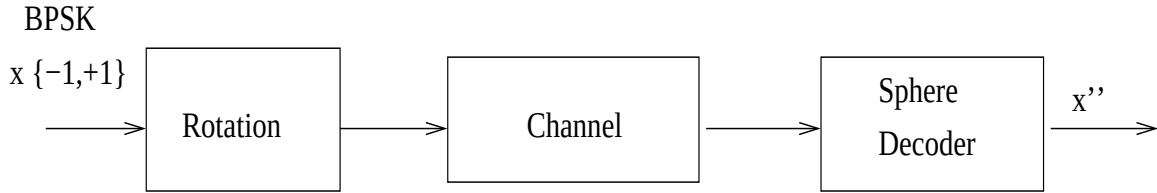


FIG. 5.2 – A simple transmission scheme involving rotated constellations.

In Fig. 5.3 and 5.4 we show the performances obtained for several dimensions, when we used matrices constructed as in [4] and [34], respectively. The advantage of using a rotated constellation before a Rayleigh channel is obvious through the following figures. Also, we can notice the role of the diversity in the performance improvement. Apart from these general remarks, we can see that the performance of the two lattice constellation constructions that have been tested ([4] and [34]) are very close.

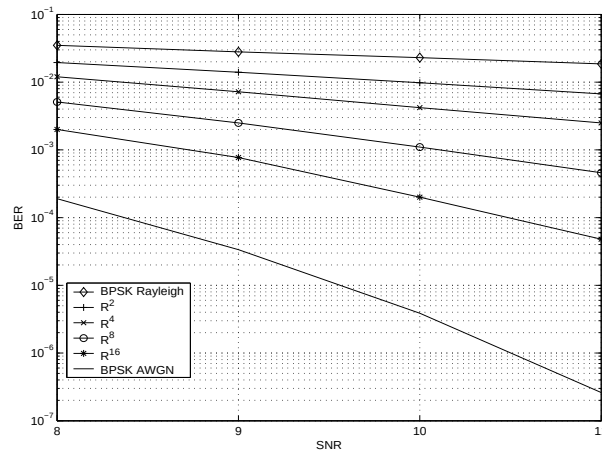


FIG. 5.3 – BER of rotated multidimensional Z^n lattice constellations constructed as in [4] for different SNR over a Rayleigh fading channel.

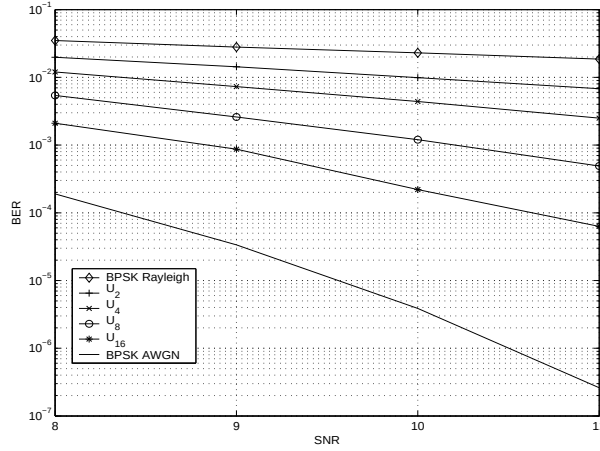


FIG. 5.4 – BER of rotated multidimensional $\mathbb{Z}^n$ lattice constellations constructed as in [34] for different SNR over a Rayleigh fading channel.

5.8 Conclusion

In this chapter, we presented a different approach for the Joint Source-Channel coding, which follows the channel point of view. The channel distortion is first minimized due to the linear labelling, and linear transforms minimize simultaneously the source distortion of Gaussian sources. We have shown that these linear transforms are, indeed, rotated multidimensional $\mathbb{Z}^n$ lattice constellations with full diversity and exploited their construction based on the theory of ideal lattices. We compared their performance over a Rayleigh fading channel with the help of the Maximum-Likelihood sphere decoder.

We deduced that, even if the constructions found in [4] have slightly better theoretical performance than those found in [34], the simplicity of the last method makes it more attractive for practical constructions. In addition, in the video domain, as the wavelet coefficients exhibit strong relations with their spatial and spatio-temporal neighbors, it is preferred that the dimensions of d in $\mathbf{G}_{d,n}$ and n in $\mathbf{G}'_{n,r}$ to be multiple of four. However, both of them come from the same square matrix $n \times n$ which represents a rotated $\mathbb{Z}^n$ -lattice. The constructions in [34] directly produce the necessary dimensions compared to the constructions in [4].

In the next Chapter, we present the extension of the JSC coding scheme developed for Gaussian sources to the video domain. The quantization matrices $\mathbf{G}_{d,n}$ and $\mathbf{G}'_{n,r}$ come from the square $n \times n$ matrix $\mathbf{U}_n$ constructed in 5.5.

Chapitre 6

Joint source-channel coding of a video on a Gaussian channel

In this Chapter, we present the construction of a joint source-channel coding scheme for video sequences based on a structured vector quantizer with linear labelling. To this end, we first introduce an efficient bit allocation algorithm taking into account the positivity of the desired bitrate solution. Next, the construction of a robust source codebook, specifically designed for video sources is developed. Two different cases are further studied : in the first one, the index of the codewords is transmitted uncoded, while in the second one Reed-Muller codes are used to transmit the index. We evaluate the performances of both schemes on a Gaussian channel and compare them with an unstructured vector quantizer also using an index assignment procedure.

6.1 Introduction

The approach for joint source-channel coding that we presented in the previous chapter has been initially developed for Gaussian sources. The subbands of the $t+2D$ decomposed video sequences to which we want to apply this coding procedure are far away from being Gaussian, and this irrespective of the fact that the temporal filtering is performed with or without motion estimation/compensation. Thus, the extension of this JSC coding scheme to the video domain requires significant modifications.

A first step in building our video coding scheme is the introduction of an efficient bit allocation algorithm, taking into account the positivity of the desired bitrate solution. Next, the construction of a robust source codebook, specifically designed for video sources, is developed, exploiting the previously introduced statistical model. We analyse two procedures for transmitting the codewords, the first one with uncoded index assignment and the second one with coded index assignment using Reed-Muller codes for increasing the transmission robustness.

We evaluate the robustness of our coding scheme in the presence of noise under different states of the Gaussian channel and different bitrates. Finally we compare our structured quantizer with linear index assignment with an unstructured vector quantizer also using an index assignment procedure.

6.2 Bit allocation algorithm

The study of bit allocation addresses the general problem of finding the optimal distribution of resources (i.e., bits) among a set of competing users (e.g., quantizers) that minimizes the objective function (e.g., distortion), subject to some global resource constraints.

In this section, we develop an algorithm for optimal partitioning of source and channel coding bits for the scalable video decomposition described in the previous chapters. By "optimal", we mean a partitioning which results in minimum expected value of the end-to-end distortion, which we choose to be the mean squared error (MSE).

As we have already argued, the channel distortion D_c is minimized by the linear labelling. Thus, the minimization of the global distortion for a video sequence $D = D_s + D_c$, subject to an overall rate constraint, will be focused on the minimization of the source distortion, D_s , resulting from quantization errors.

6.2.1 A general method of bit allocation for spatio-temporal subbands

We introduce a general method of rate allocation in the framework of a video transmission system. Note that excellent surveys of bit allocation approaches, both model-based and general ones, are presented in [65], [83], [91]. We consider a wavelet based spatio-temporal decomposition of the video sequence and Q levels of quantization.

Let $f_{n,m}$ represent a frame of the video sequence of size $N \times M$. If we assume that the overload distortion is negligible and that the high resolution assumption is valid, then an approximated model of the quantization error is given by :

$$E\{[f_{n,m} - Q(f_{n,m})]^2\} \approx K 2^{-2r} \sigma_{f_{n,m}}^2,$$

where $Q()$ is the quantization function, K is a constant (determined by the pdf of the normalized random variable $f_{n,m}/\sigma_{f_{n,m}}$ and the type of quantizer which is used) and $r = \log_2 Q$ is the number of bits encoding the Q quantization levels, in case of a uniform scalar quantization.

We consider j_m levels of spatial decomposition of the temporal subband frames. If $\tilde{f}_{n,m}$ represents the reconstructed frame and $d_j^s, \hat{d}_j^s$ the detail subbands with orientation s ($s \in \{\mathbf{Horizontal}, \mathbf{Vertical}, \mathbf{Diagonal}\}$) at the spatial level j before and after quantization respectively, and $c_{j_m}, \hat{c}_{j_m}$ the approximation subband at the coarsest spatial level j_m , before and after quantization respectively, then the quantization distortion of the frame $f_{n,m}$ is given by :

$$\sum_{n=0}^{N-1} \sum_{m=0}^{M-1} \{(f_{n,m} - \tilde{f}_{n,m})^2\} = \sum_{j=0}^{j_m} \sum_{s \in H,V,D} \sum_{k=0}^{\frac{N}{2^j}-1} \sum_{l=0}^{\frac{M}{2^j}-1} [d_j^s(k,l) - \hat{d}_j^s(k,l)]^2 + \sum_{k=0}^{\frac{N}{2^{j_m}}-1} \sum_{l=0}^{\frac{M}{2^{j_m}}-1} [c_{j_m}(k,l) - \hat{c}_{j_m}(k,l)]^2$$

In addition, by the law of large numbers, we can estimate the statistical averages $E\{\cdot\}$ by a spatial mean :

$$E\{(f_{n,m} - \tilde{f}_{n,m})^2\} \approx \frac{1}{NM} \sum_{n=0}^{N-1} \sum_{m=0}^{M-1} \{(f_{n,m} - \tilde{f}_{n,m})^2\}$$

Thus,

$$E\{(f_{n,m} - \tilde{f}_{n,m})^2\} \approx \frac{1}{NM} \left[\sum_{j=0}^{j_m} \sum_{s \in H,V,D} \frac{N}{2^j} \frac{M}{2^j} E\{(d_j^s(k,l) - \hat{d}_j^s(k,l))^2\} + \frac{N}{2^{j_m}} \frac{M}{2^{j_m}} E\{(c_{j_m}(k,l) - \hat{c}_{j_m}(k,l))^2\} \right]$$

$$E\{(f_{n,m} - \tilde{f}_{n,m})^2\} \approx \sum_{j=0}^{j_m} 2^{-2j} \sum_{s \in H,V,D} E\{(d_j^s(k,l) - \hat{d}_j^s(k,l))^2\} + 2^{-2j_m} E\{(c_{j_m}(k,l) - \hat{c}_{j_m}(k,l))^2\}$$

If the wavelet transform is orthogonal, which is (approximately) true in most cases of the filters generally used, we can say that :

$$E\{(f_{n,m} - \tilde{f}_{n,m})^2\} \approx \sum_{j=0}^{j_m} 2^{-2j} \sum_{s \in H,V,D} K_{j,s} 2^{-2r_{j,s}} \sigma_{d_j^s}^2 + 2^{-2j_m} K_{j_m,0} 2^{-2r_{j_m,0}} \sigma_{c_{j_m}}^2$$

where $r_{j,s}$ (resp., $r_{j_m,0}$) represents the number of bits per coefficient in the subband s (resp. the approx. subband) at the spatial level of decomposition j (resp., j_m). We suppose that $K_{j,s} = K_{j_m,0} = K$.

The bit allocation algorithm consists of minimizing the above distortion subject to an overall rate constraint $\bar{r}$. That means we minimize the following expression :

$$\min \sum_{j=0}^{j_m} 2^{-2j} \sum_{s \in H,V,D} K_{j,s} 2^{-2r_{j,s}} \sigma_{d_j^s}^2 + 2^{-2j_m} K_{j_m,0} 2^{-2r_{j_m,0}} \sigma_{c_{j_m}}^2$$

subject to the rate constraint :

$$\bar{r} = \sum_{j=0}^{j_m} 2^{-2j} \sum_{s \in H,V,D} 2^{-2r_{j,s}} + 2^{-2j_m} 2^{-2r_{j_m,0}}$$

Corresponding to the above constrained minimization problem, an unconstrained minimization problem using Lagrange multipliers can be stated as :

$$\arg \min \left\{ \sum_{j=0}^{j_m} 2^{-2j} \sum_{s \in H,V,D} K_{j,s} 2^{-2r_{j,s}} \sigma_{d_j^s}^2 + 2^{-2j_m} K_{j_m,0} 2^{-2r_{j_m,0}} \sigma_{c_{j_m}}^2 + \lambda \left(\sum_{j=0}^{j_m} 2^{-2j} \sum_{s \in H,V,D} 2^{-2r_{j,s}} + 2^{-2j_m} 2^{-2r_{j_m,0}} - \bar{r} \right) \right\}$$

This minimization leads to the following solution :

$$r_{j,s} = \bar{r} + \frac{1}{2} \log_2 \left[\frac{\sigma_{j,s}^2}{\prod_{j,s} (\sigma_{j,s}^2)^{2^{-2j}} (\sigma_{j_m,0}^2)^{2^{-2j_m}}} \right] \quad (6.1)$$

Based on the above solution we can compute the distortion of each subband :

$$D_{j,s} = E\{(d_j^s - \hat{d}_j^s)^2\} = K 2^{-2r_{j,s}} \sigma_{d_j^s}^2 = K 2^{-2\bar{r}} \prod_{j,s} (\sigma_{j,s}^2)^{2^{-2j}} (\sigma_{j_m,0}^2)^{2^{-2jm}}$$

We can notice that the above distortion is independent of the level of the spatial decomposition and the position of the subband. Thus, if I is the total number of subbands then for each subband i , where $i \in \{1, \dots, I\}$ the equation (6.1) becomes :

$$r_i = \bar{r} + \frac{1}{2} \log_2 \left[\frac{\sigma_i^2}{\prod_j (\sigma_j^2)^{\frac{n_j}{n}}} \right] \quad (6.2)$$

where now we denote with n the total number of coefficients of the frame and with n_j the total number of coefficients of the subband j .

We also have to notice that the above algorithm can be applied not only to a frame but also to a GOF or to the whole video sequence.

6.2.2 A robust bit allocation algorithm

The algorithm presented in the previous section can lead to very bad results, especially at low bitrates. Indeed, it neglects the practical requirement that the bitrate $r_i \geq 0$ and leads to "optimal" solutions that admit negative values of r_i .

An early attempt towards finding the solution of this problem has been presented in [76], where some basic techniques of functional analysis were used to take into account the nonnegativity constraint. It supposed that the quantizer function is strictly convex, which means that the quantizer distortion decreases smoothly as the number of bits increases and the rate of decrease of distortion drops monotonically to zero as the number of bits increases without bound. Another approach to the problem, found in [83], takes another step forward. This algorithm is not based on prior assumptions about the nature of the given set of quantizers and the solution takes into account a nonnegativity integer constraint. The algorithm is based on Lagrangian method for integer allocation and it is widely used.

Our bit allocation algorithm takes, also, into account the nonnegativity constraint and distributes the given bitrate according to the decreasing order of the variance of the subbands and by this way the importance of each subband is taken into consideration.

The proposed iterative algorithm follows the next steps :

1. Order the subbands by decreasing variance :
 $+\infty \geq \sigma_1^2 \geq \sigma_2^2 \geq \dots \geq \sigma_I^2 \geq 0$
2. Initialize the index of iterations : $l = 1$
3. At each iteration calculate

$$\lambda_l = -(2 \ln 2) 2^{-\frac{2NR}{M_l}} \prod_{k=1}^l (\sigma_k^2)^{\frac{n_k}{M_l}}$$

where $M_l = \sum_{i=1}^l n_i$

4. If λ_l satisfies the following inequality :

$$-2 \ln 2 \sigma_l^2 < \lambda_l \leq -2 \ln 2 \sigma_{l+1}^2$$

then set $l = l + 1$ and go to step 3, else exit.

The final allocation result is given by :

$$r_k = \begin{cases} \frac{N}{M_l} \bar{r} + \frac{1}{2} \log_2 \left(\frac{\sigma_k^2}{\prod_{k=1}^l (\sigma_k^2)^{\frac{n_k}{M_l}}} \right), & k \in \{1, \dots, l\} \\ 0, & k \in \{l+1, \dots, I\} \end{cases}$$

This result is obtained by direct application of the theorem of duality of Fenchel [54].

6.3 Coding algorithm

In this section we present the JSC coding scheme developed for transmitting scalable video over a noisy channel. The main components of our coding scheme are the following :

1. a motion compensated temporal-subband decomposition applied on groups of frames (GOF), followed by a spatial wavelet transform of the temporal subbands.
2. a bit-allocation procedure, as described in Section 6.2.
3. a vector quantization of spatio-temporal coefficients with a linear mapping between the source codebook and the coded symbols sent on the channel as described in Chapter 5. Fig. 6.1 summarizes this step together with the corresponding decoding procedure.

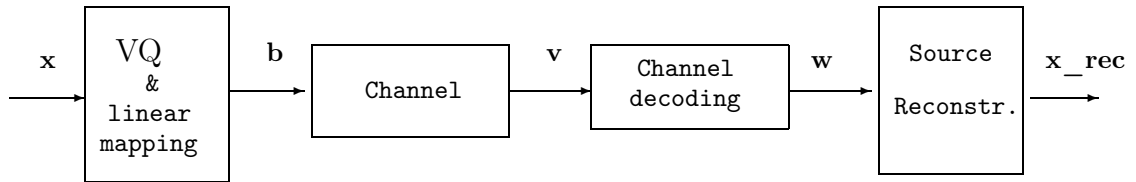


FIG. 6.1 – System description.

In Fig. 6.1, $\mathbf{x}$ designates the vectors formed from wavelet coefficients representing the input to our JSC coding scheme and $\mathbf{b} \in \{+1, -1\}^n$ the codewords, representing the index assignment, that will be transmitted through the channel.

The vector quantization and the linear labelling are based on the theory presented in the Sections 5.2 and 5.5. That means that the quantization matrices are based on rotated

$\mathbf{Z}^n$ lattices generated by the ideals of rings of maximal real cyclotomic fields, where n is a power of two. In addition, as presented in Section 5.2, using the appropriate $\mathbf{G}_{d,n}$ matrix we can build a source codebook which mimics the source distribution and minimizes thus the source distortion. The channel distortion is minimized as the labelling is chosen to be linear.

However, the above theory is designed for iid Gaussian sources, which means that it requires significant modifications in order to be efficiently applied to a video source. As already discussed in Chapter 1, the subband coefficients of a $t+2D$ wavelet decomposition are not Gaussian and the histograms of the detail subband frames are observed to have sharper peaks and heavier tails than the Generalized Gaussian distributions. In the rest of this section, we present the necessary modifications in order to take into account these features.

6.3.1 Source codebook construction

First, we develop a general algorithm in order to find the source codebook for each subband. Consider the system illustrated in Fig. 6.1. We find the source codebook for each subband by minimizing the following expression :

$$\min_{\mathbf{b}} E \|\mathbf{x} - \beta \mathbf{G}_{d,n} \mathbf{b}\|^2 \quad (6.3)$$

where $\mathbf{b} = (b_1, \dots, b_n)^t \in BPSK_n$ and $\mathbf{G}_{d,n}$ is the matrix obtained as explained in Section 5.2. β is an important parameter which scales the lattice constellation to the source dynamics. In order to find the parameter β and the source codebook with vectors $\mathbf{y} = \mathbf{G}_{d,n} \mathbf{b}$, the following iterative algorithm is applied :

1. Initialisation :
Set $\beta = \beta^0$, $D_0 = \infty$ and let T be a threshold, $T > 0$.
At iteration i , $i \geq 1$:
2. Find $\mathbf{b}^{(i)}$ minimizing the expression in Eq. (6.3).
3. Compute the average distortion

$$D_i = \frac{1}{p} \sum_{k=1}^p \|\mathbf{x}_k - \beta^{(i-1)} \mathbf{G}_{d,n} \mathbf{b}_k^{(i)}\|^2$$

where p is the number of source vectors.

4. If $\frac{D_{i-1} - D_i}{D_{i-1}} < T$ exit, else update β :

$$\beta^{(i)} = \frac{\sum_{k=1}^p \mathbf{y}_k^{(i)T} \mathbf{x}_k}{\sum_{k=1}^p \|\mathbf{y}_k^{(i)}\|^2}$$

and go to step 2.

A special attention should be paid to the initialisation of β . Assuming that the initial codeword vector $\mathbf{y}$ is with zero mean, we find β^0 by solving the following optimisation in the matrix space :

$$\min \|\beta^2 \mathbf{R}_y - \mathbf{R}_x\|_F$$

where $\mathbf{R}_y, \mathbf{R}_x$ are the autocorrelation matrices of the codeword and the source vectors respectively and $\|\cdot\|_F$ is the Frobenius norm of a matrix. Then :

$$\beta^0 = \sqrt{\frac{\text{Tr}(\mathbf{R}_x \mathbf{R}_y^T)}{\text{Tr}(\mathbf{R}_y \mathbf{R}_y^T)}}.$$

We have to note that the performance of a quantizer is often specified in terms of signal-to-quantization-noise ratio, sometimes denoted by SQNR or SQR. It is defined by normalizing the signal power by the quantization error power, D , and taking a scaled logarithm :

$$SQR = 10 \log_{10} \frac{E(X^2)}{D}$$

measured in decibels (dB). It is usually assumed that the input process has zero mean and hence that the signal power is the same as the variance σ_X^2 of the signal samples.

In Table 6.1 we show some results of the application of the above algorithm to several subbands of the CIF "hall-monitor" test sequence, 30 fps. No motion estimation was used for the temporal filtering step. The coding rate is chosen to be $r = 0.5$ bits/dim, using a matrix $G'_{16,8}$. In the last column, we also added the estimated parameter of GGD, γ , if we assume that the subbands are modelled as GGD. We remark that the SQR are extremely low.

"hall-monitor" $r = 0.5$ bits/dim					
Temporal detail frames					
temp. level.	spatial level	spatial orientation	SQR (dB)	β	γ
4	2	Approximation	0.13	12.31	0.3
4	2	Vertical	0.31	3.75	0.54
4	1	Vertical	0.57	1.48	0.7
2	2	Approximation	0.18	5.16	0.4
2	2	Vertical	0.35	2.8	0.63
2	1	Vertical	0.51	1.33	0.73
Temporal approximation frames					
temp. level.	spatial level	spatial orientation	SQR (dB)	β	γ
4	2	Approximation	0.41	175.68	0.91
4	2	Vertical	0.28	41.23	0.5
4	1	Vertical	0.16	11.89	0.31

TAB. 6.1 – Quantization and modelling of different spatio-temporal subbands.

In Fig. 6.2 we present the histograms of the real subbands versus the plots of GGD with the same parameters. We notice that the modelling by GGD does not perfectly fit to

actual distributions. As we have already remarked in Section 1.5.1, the histograms of the detail subband frames have sharper peaks and heavier tails than Generalized Gaussian distributions.

Based on the above remark, we can explain the low SQR presented in Table 6.1. Indeed, the histograms show that the subband coefficients having respectively very small and very large magnitude gather together at the two extremes of the distribution and the attempt to capture both behaviours in a single marginal law does not lead to the most efficient model. Thus, it appeared to us more appropriate to capture separately the small and large coefficients in order to benefit from a better coding gain. This was the idea that leads us to estimate two different β parameters in the same subband, corresponding to a mixture of Gaussian implicit modelling. However, our method explicitly makes use only of the conditional law exhibit in Section 2.3 when setting the threshold for the above classification.

6.3.2 Scaling the lattice constellation to the source dynamics

The question that we shall tackle now is how to find an efficient way of classifying the subband coefficients or, in other words, how to find the most appropriate threshold. This threshold should be available at the decoder but in a way that does not increase the bitrate. Sending one value per subband appears therefore as too budget consuming. We shall try to compute it from the encoded coefficients, and thus not having to send it as side-information. From another point of view, this threshold should be robust in the presence of noise.

In Chapter 2 we proposed an efficient model, where the conditional probability law of the coefficients in a given subband is considered Gaussian with variance depending on the set of spatio-temporal neighbors :

$$g(x_{n,m}|\sigma_{n,m}^2) = \frac{1}{\sqrt{2\pi}\sigma_{n,m}} e^{-\frac{x_{n,m}^2}{2\sigma_{n,m}^2}},$$

where

$$\sigma_{n,m}^2 = \sum_k w_k |\mathbf{p}_k(n, m)|^2 + \alpha,$$

$x_{n,m}$ is the current coefficient in a given spatial subband and $\mathbf{p}_k$ is the vector of neighbors. The parameter estimation of this model has been described in Chapter 2.

The classification method of the coefficients in the temporal detail frames will be based on the above stochastic model of the spatio-temporal dependencies between the wavelet coefficients. The spatio-temporal neighborhood consists of coefficients that have already been received by the decoder when the current coefficient is decoded, which ensures that the decoder can do the same computations as the encoder.

Based on Fig. 1.4 presented in Chapter 1, we choose as spatio-temporal neighborhood for the current coefficient the set of coefficients consisting of its spatial parent, its spatial aunts, its spatio-temporal parent, its spatio-temporal parent's spatial neighbors (Up and Left) and its spatio-temporal aunts. By this way, every coefficient in a group of four (2×2

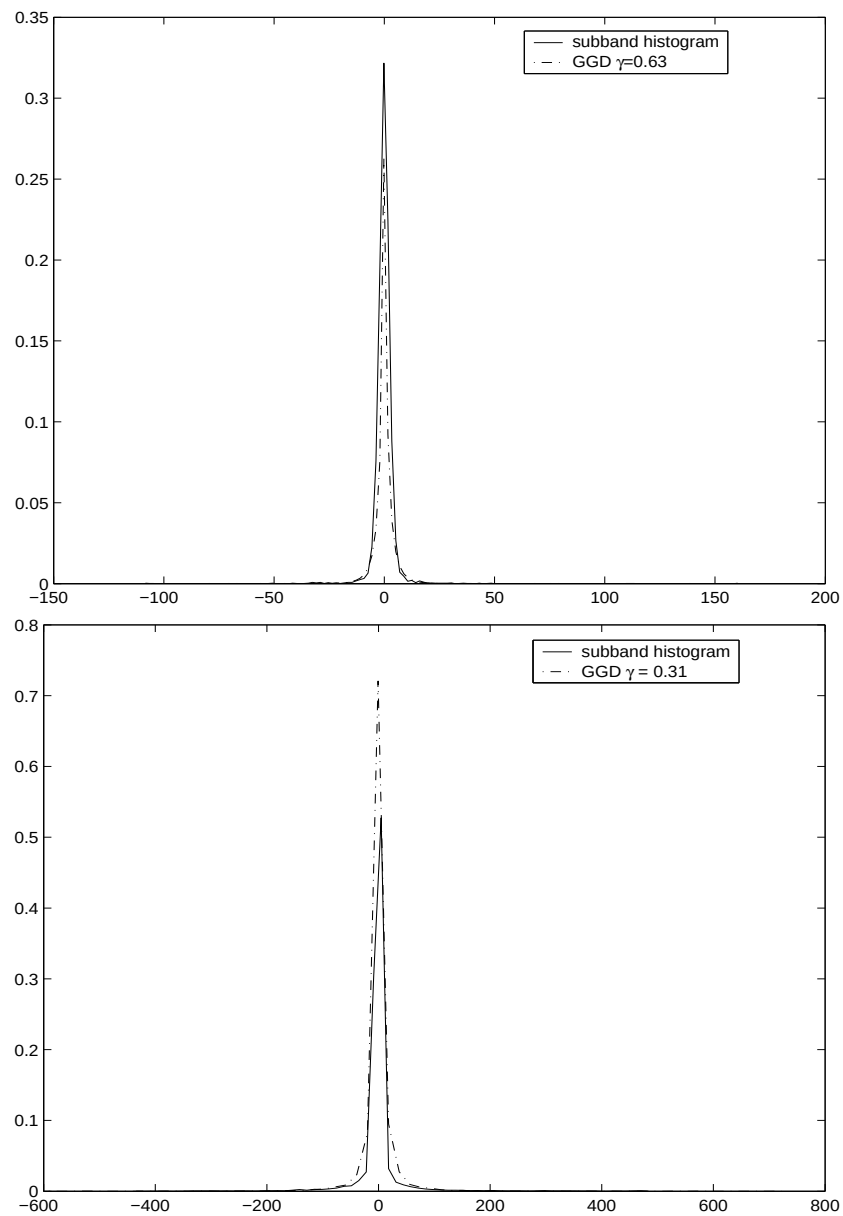


FIG. 6.2 – Histograms of real subbands vs GGD with the same parameters. Up : vertical subband at second spatial/ second temporal resolution level of a detail frame. Down : vertical subband at first spatial resolution level of the approximation frame.

block) has the same conditional variance with the other three. For the coefficients of the subbands in the coarsest spatial resolution level (except the approximation one), the absence of the above neighborhood forces us to take as neighbors the coefficient in the same position in the approximation subband and its neighbors (Up and Left).

The procedure for threshold computation and subband coding is the following : we estimate the average value $\sigma_{n,m}$ of the subband vectors at the encoder based on the already coded spatio-temporal neighbors. Next, we classify them taking as threshold the value maximizing the *total* coding gain of the entire subband. The same procedure is applied at the decoder and, in this way, no side information is required to indicate the class to which the current vector belongs to.

However, for the approximation frames, as well as for the spatial approximation subband of the temporal detail frames, the classification procedure is different, since the spatio-temporal neighborhood does not have enough coefficients to estimate $\sigma_{n,m}$. In this case, the subband vectors are classified according to the norm of the uncoded vectors and thus an additional bit is necessary to be transmitted in order to give to the decoder the appropriate class information.

In Table 6.2 we present the obtained results, again on the "hall-monitor" sequence, by dividing the subband coefficients in two classes. β_1 and β_2 represent the scaling parameters of the lattice constellation corresponding to the two classes of the subband. The procedure to compute them remains the same as for the case of one β per subband.

We can notice the significant increase of coding gain compared to the results in Table 6.1. Another interesting comment concerns the values of β_1 and β_2 , which are quite different. Together with the coding gain, this indicates that the classification of the coefficients into two classes was indeed more appropriate than trying to accommodate a single scaling parameter per subband.

"hall-monitor" $r = 0.5$ bits/dim					
Temporal detail frames					
temp level.	spatial level	spatial orientation	SQR (dB)	β_1	β_2
4	2	Approximation	1.08	4.19	121.58
4	2	Vertical	0.91	2.27	19.97
4	1	Vertical	0.87	1.17	4.86
2	2	Approximation	1.18	1.91	44.76
2	2	Vertical	0.76	1.98	12.79
2	1	Vertical	0.68	1.09	3.56
Temporal approximation frames					
temp level.	spatial level	spatial orientation	SQR (dB)	β_1	β_2
4	2	Approximation	1.56	81.01	409.21
4	2	Vertical	0.89	25.62	179.24
4	1	Vertical	0.63	6.99	84.15

TAB. 6.2 – Quantization results for different spatio-temporal subbands, after classification of the subband coefficients.

6.4 Application of joint source-channel coding scheme of a video over a Gaussian channel

In this Section, we present the simulation results when the video sequences are transmitted over a Gaussian channel. The mostly used performance criterion in the video domain is the Peak Signal-to-Noise Ratio (PSNR), which for a frame (with pixels represented on 255 levels of gray) is calculated as :

$$PSNR = 10 \log_{10} \frac{255^2}{MSE}$$

where MSE is the mean square reconstruction error.

In our simulations, we considered a temporal Haar decomposition applied on GOFs of 16 frames, with 4 temporal and 2 spatial resolution levels. The spatial multiresolution analysis is based on the biorthogonal 9/7 filters. Note that more performant temporal (and even spatial) decompositions have been recently proposed in the literature [66], [67]. Our approach can be applied with no additional difficulties to such representations, since the main characteristics of the spatio-temporal coefficients still remain valid. In what follows, we make the distinction between two cases :

1. the temporal decomposition involves the motion estimation/compensation. The motion estimation/compensation in the Haar temporal decomposition uses a full-search block matching algorithm with full pixel accuracy.
2. the temporal decomposition is performed without motion estimation.

In the first case, the temporal transform better compacts the energy of the video sequence, taking advantage of a better temporal prediction. However, the resulting motion vector fields (MVF) need to be encoded and transmitted. For an efficient 5/3 temporal transform, the number of MVF is even double compared with the motion-compensated Haar transform. Therefore, at low bitrates or for sequences presenting slow motion, the trade-off is in favour of not having to encode the MVF. Moreover, encoding the motion vectors raises, in our case, the tricky problem of their robust transmission. While one can argue they represent a small percentage of the bitstream, and therefore can be very well protected, this remains a challenging problem for efficient coding and transmission over noisy channels. Therefore, in the sequel, we provided simulations as well with motion estimation/compensation, to show an upper bound of performance for our scheme, as well without motion estimation/compensation, which is the structure currently fully compatible with our approach.

6.4.1 A First Attempt

In a very early attempt, we have checked the performance of our coding scheme without involving the bit allocation algorithm proposed in Section 6.2. While this approach sacrifices the compression efficiency, this feasibility test proved that our JSC was quite promising especially concerning its robustness capabilities.

In order to quantize the temporal detail frames, the wavelet coefficients are grouped in 16-dimensional vectors and we use a 16×4 matrix $\mathbf{G}'_{16,4}$ to map them to the channel

codewords. As better protection for the frames in the temporal approximation subband is necessary, we use a matrix $\mathbf{G}_{2,16}$ of size 2×16 . Both matrices are extracted from the same $\mathbf{U}_{16}$ full diversity matrix.

To construct $\mathbf{G}'_{16,4}$, we take the first 4 columns of $\mathbf{U}_{16}$ and to construct $\mathbf{G}_{2,16}$, the first two rows are retained. We have remarked that better performance is achieved in practice when consecutive rows or columns of the $\mathbf{U}_{16}$ are chosen.

In this early attempt, even though the chosen linear labelling minimizes the channel errors, we enhanced the robustness over noisy environments of our scheme by implementing an additional protection of the approximation frames using a channel code. We applied to this part of the bitstream a Rate Punctured Convolutional Code [7].

Briefly, recall that a punctured code is a high-rate code $R_p = k_p/n_p$ formed by periodic elimination of specific code symbols of the output of a low-rate code $R_o = 1/n_o$. An important advantage of using this type of code is that decoding with the Viterbi algorithm is hardly more complex than for the original code.

Based on [7], we used a punctured convolutional code $R_p = 8/9$ which depends on a convolutional code $R_o = 1/2$ (561, 753) of memory $m = 8$. The perforation matrix [7] is given by

$$P = \begin{bmatrix} 1 & 1 & 1 & 0 & 0 & 0 & 0 & 0 \\ 1 & 0 & 0 & 1 & 1 & 1 & 1 & 1 \end{bmatrix}$$

The use of Rate Punctured Convolutional Codes is a classical way of treatment for protection against the channel errors, as we have already presented in Chapter 4. However, in our application, we only use them for a partial protection of subbands of the temporal approximation frames.

Four coding scenarii have been examined, according to the degree of protection which is desired :

Type 1 : all the spatial subbands of the temporal approximation frames are protected by the punctured code.

Type 2 : only the subbands at the coarsest spatial resolution level of the approximation frames are protected.

Type 3 : only the lowest frequency spatial subband of the approximation frames is protected.

Type 4 : there is no additional protection. The robustness is achieved by the use of only our JSC coding method.

A hard decision criterion is applied at the decoder for the temporal detail frames and for the subbands of the approximation frames that are not protected by the punctured code. For the part of the bitstream where the punctured code was used, a soft decision decoding by the Viterbi algorithm is applied.

For our tests, we considered CIF (352×288) test sequences at 30 fps. Thus, when considering also the redundancy introduced by the punctured code, the total transmission rate for the *Type 1* coding is 2.53 Mbs, for the *Type 2* is 2.39 Mbs, for the *Type 3* is 2.35 Mbs and finally for the *Type 4* is 2.34 Mbs.

In Fig. 6.3, we compare the PSNR (averaged over 30 noisy realizations) of a GOF in the video test sequence "hall-monitor" obtained with the four types of coding mentioned above over a noisy Gaussian channel with SNR = 4.33 dB. We also compare them with the PSNR obtained on a noiseless channel. We remark that the proposed vector-quantization scheme

offers, in itself, good robustness to noise and then, even a small additional protection of the approximation frames, like in the *Type 3* scenario, leads to significant increase in performance without affecting proportionally the bitrate.

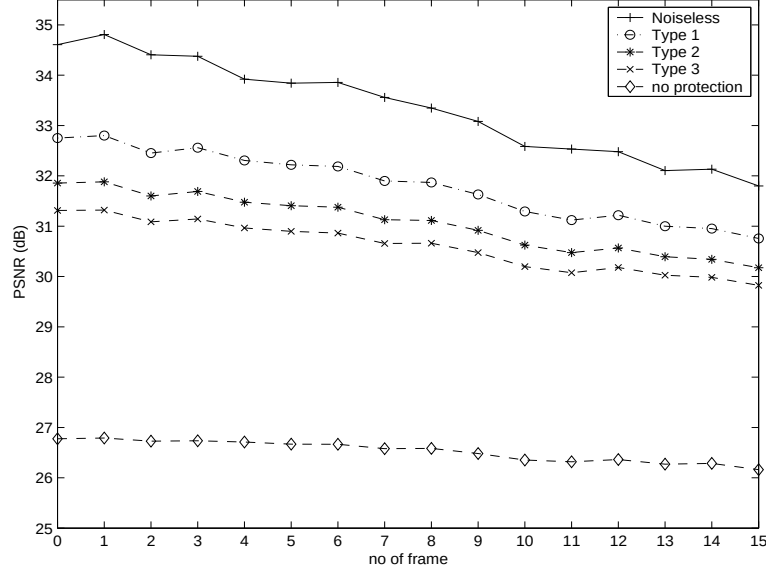


FIG. 6.3 – PSNR of the frames of a GOF on a Gaussian channel of $SNR = 4.33$ dB with four different types of coding.

In Table 6.3 we present the average PSNR of the "hall-monitor" and "foreman" video sequences coded with the four different coding scenarii and under two different noise states ($SNR=4.33$ dB and $SNR=6.75$ dB) of a Gaussian channel. We can remark that the performance of our joint source-channel coding scheme on a Gaussian channel with $SNR=6.75$ dB is high, and with a small protection of the coefficients in the approximation frames, we attain the limits of a noiseless scheme for the first two scenarii and loose only about 1.5 dB compared with the noiseless case in the Type 4 coding. Obviously, considering higher channel SNRs would be worthless, since we almost attained the noiseless performance with the first scenario.

"hall-monitor" (noiseless average PSNR=33.34 dB)				
	Type 1	Type 2	Type 3	Type 4
SNR=4.33	31.81	31.06	30.60	26.52
SNR=6.75	33.33	33.17	33.02	31.94
"foreman" (noiseless average PSNR=31.51 dB)				
	Type 1	Type 2	Type 3	Type 4
SNR=4.33	28.64	28.51	28.41	25.90
SNR=6.75	31.50	31.48	31.44	30.40

TAB. 6.3 – Average PSNR of the "hall-monitor" and "foreman" CIF sequences, coded by the four schemes and for two different states of the Gaussian channel.

In Fig. 6.4, we show the reconstructed frames first in a noiseless environment and then after transmission over a Gaussian channel with $\text{SNR}=4.33$ dB. For this latter situation, the second and third images show the reconstructed frames using *Type 4* and *Type 2* coding schemes respectively. Note that the very small protection introduced in the second scheme is enough to approach the noiseless quality (which is quite bad, as no bit allocation was used to optimise the coding algorithm). Without protection at all, in these very severe channel conditions, our scheme still works, and we remark a graceful degradation. Note that in these noise conditions, most systems cannot display an acceptable reconstructed image.



FIG. 6.4 – Reconstructed Frames. First line : reconstructed frame in a noiseless environment. Second line, left : reconstructed frame when no protection by a code is used over a Gaussian channel with $\text{SNR}=4.33$ dB. Second line, right : the reconstructed frame when the coarsest spatial resolution level of the temporal approximation frame is protected (Gaussian channel with $\text{SNR}=4.33$ dB).

6.4.2 Application of the bit allocation algorithm

The previous results showed that our JSC coding scheme is attractive and robust enough in the presence of noise. However, without the application of a bit allocation

algorithm its performance remains too low. In addition, in what follows, **no additional error correcting codes are used**.

The bit allocation algorithm presented in the Section 6.2 indicates the size of the $\mathbf{G}_{d,n}$ or $\mathbf{G}'_{n,r}$ which minimizes the end-to-end distortion. The choices of the $\mathbf{G}_{d,n}$ or $\mathbf{G}'_{n,r}$ are, however, limited by the complexity and the dependences of the spatio-temporal coefficients. As already explained, the spatio-temporal coefficients exhibit strong relations with their spatial or spatio-temporal neighbors. Thus, in order to capture these relations, the dimensions d in $\mathbf{G}_{d,n}$ or n in $\mathbf{G}'_{n,r}$ are preferred to be multiple of four. However, in order to attain high or low coding rates we permit us taking d or r equal to two. In addition, in order to keep the complexity low, we have limited the dimensions n in $\mathbf{G}_{d,n}$ and r in $\mathbf{G}'_{n,r}$ to 16.

Following the above "conditions", in Table 6.4, we present the possible coding rates and corresponding matrices of quantization. Note again that all these matrices are extracted from the same $\mathbf{U}_{16}$ full diversity matrix and that in the case of a coding rate equal to 1 bits/coeff. the quantization matrix is the $\mathbf{U}_{16}$ itself. .

cod. rate (bits/coeff)	Matrix of quantiz.
8	$\mathbf{G}_{2,16}$
4	$\mathbf{G}_{4,16}$
2	$\mathbf{G}_{8,16}$
1	$\mathbf{G}_{16,16} = \mathbf{U}_{16}$
0.5	$\mathbf{G}'_{16,8}$
0.25	$\mathbf{G}'_{16,4}$
0.125	$\mathbf{G}'_{16,2}$

TAB. 6.4 – Coding rates and corresponding matrices of quantization.

The global bitrates per pixel tested are : $0.1bpp$, $0.16bpp$, $0.33bpp$ and $0.48bpp$. However, due to the additional bits sent in order to indicate the class where the vectors of the approximation frames belong to and due to the coding of the motion vectors when motion estimation is applied, the global bitrate cannot be the same for different sequences.

Experiments without motion estimation

In Table 6.5 we present the average PSNR of two test sequences ("hall-monitor" and "foreman") on a Gaussian channel without motion estimation under different noise states and under different bitrates. The decoder uses in all cases a hard decision criterion.

From Table 6.5 we can remark that the JSC coding scheme becomes much more flexible when the bit allocation algorithm is applied. The overall coding scheme remains very efficient in the presence of noise. In addition, when $SNR = 8.0$ dB one can notice that the performance of the scheme almost attains the noiseless quality.

Experiments involving motion estimation

In Table 6.6 we present the average PSNR of two test sequences ("hall-monitor" and "foreman") with motion estimation, on a Gaussian channel under different noise states

"hall-monitor"				
bitrate (kbs)	295	566	1027	1550
noiseless	26.68	30.53	32.81	34.66
SNR=4.33 dB	23.44	26.01	26.57	26.90
SNR=6.75 dB	26.21	29.75	31.54	32.84
SNR=8.0 dB	26.58	30.38	32.54	34.25
"foreman"				
bitrate (kbs)	307	576	983	1410
noiseless	25.07	27.56	28.81	30.14
SNR=4.33 dB	22.98	24.71	25.04	25.44
SNR=6.75 dB	24.86	27.13	28.22	29.36
SNR=8.0 dB	25.06	27.49	28.66	29.98

TAB. 6.5 – Average PSNR of the "hall-monitor" and "foreman" CIF sequences, without motion estimation.

and under different bitrates. The decoder uses in all cases a hard decision criterion.

"hall-monitor"				
bitrate (kbs)	317	565	1113	1660
noiseless	27.88	31.78	34.38	35.59
SNR=4.33 dB	23.91	26.41	26.96	27.23
SNR=6.75 dB	27.22	30.81	32.68	33.47
SNR=8.0 dB	27.76	31.58	34.05	35.09
"foreman"				
bitrate (kbs)	376	524	1104	1530
noiseless	30.07	30.78	33.30	34.51
SNR=4.33 dB	25.53	25.66	26.27	26.49
SNR=6.75 dB	29.27	29.86	31.72	32.63
SNR=8.0 dB	29.91	30.58	32.98	34.05

TAB. 6.6 – Average PSNR of the "hall-monitor" and "foreman" CIF sequences, with motion estimation/compensation in the temporal transform.

It is interesting to compare the Table 6.6 to Table 6.3 presented in Sect. 6.4.1, where the bit allocation algorithm was not applied. We can notice that a better distribution of the available rate among the subbands gives a really higher performance even at bitrates lower than those used in 6.4.1.

Fig. 6.5 illustrates the reconstructed frames of "hall-monitor" at 566 kbs without motion estimation and of "foreman" at 1104 kbs with motion estimation, first in a noiseless environment and then after transmission over a Gaussian channel with SNR=6.75 dB and SNR=8.0 dB. One can remark the graceful degradation with the noise level, due to the efficient allocation, and also with the bitrate, due to the scalability of the scheme.



FIG. 6.5 – Reconstructed Frames of "foreman" at 1104 kbs with motion estimation and "hall-monitor" at 566 kbs without motion estimation. First line : reconstructed frames in a noiseless environment. Second line : reconstructed frames over a Gaussian channel with SNR=6.75 dB. Third line : reconstructed frames over a Gaussian channel with SNR=8.0 dB.

6.5 Calculation of the end-to-end distortion in the noisy system

In the previous results, where the bit allocation algorithm is applied but no error correcting codes are added, we can easily remark that, especially in the case of $SNR = 4.33$ dB, increasing the bitrate does not offer a significant improvement to the PSNR of the reconstructed video sequence, which is not the case when the channel is noiseless. This can be explained if we look carefully at the end-to-end distortion in both cases. When no noise enters the system, the measured end-to-end distortion is the *source distortion* D_s . In the presence of noise, the measured end-to-end distortion, called *global distortion* D , includes, also, the contribution of the *channel distortion*.

In the following tables we present the actual values of these two quantities for different bitrates of the "hall-monitor" sequence without motion estimation. The Tables 6.7-6.10 show the D_s and D for different subbands (four temporal and two spatial decomposition levels) including the subbands where these two quantities have a significant difference.

The distortion of a subband $N \times M$ in both cases is calculated as :

$$D, D_s = \frac{\sum_{k=0}^{NM} (x_k - \hat{x}_k)^2}{NM}$$

where x_k is the wavelet coefficient before quantization and $\hat{x}_k$ is the quantized wavelet coefficient in the noiseless case. In the presence of noise, $\hat{x}_k$ is the reconstructed wavelet coefficient after quantization and channel decoding.

Remark Table 6.7 :

In the detail frames the contribution of the channel distortion is not important. On the contrary, if we pay attention to the approximation subband of the temporal approximation frame we can clearly see that this part of the bitstream is not sufficiently protected against the noise.

Remark Table 6.8 :

In the temporal detail frames, the contribution of the channel distortion is still not important. In the temporal approximation frame, the Approximation and the Vertical and Horizontal subbands of the second spatial decomposition level are not enough well protected.

Remark Table 6.9 :

In this case, we notice that except of the whole second spatial decomposition level of the temporal approximation frame and the vertical subband of the first spatial level of the temporal approximation frame, the approximation at the fourth temporal decomposition level in the temporal detail frame should be better protected.

Remark Table 6.10 :

Here, two subbands in the temporal detail frames and most of the subbands in the temporal approximation frame should be better protected.

We can also make almost the same remarks in the other cases, i.e. the "foreman" sequence without motion estimation and for both of them with motion estimation. In a high SNR environment the channel distortion, as expected, is found important only for a smaller number of subbands.

"hall-monitor" 295 kbs					
Temporal detail frames					
temporal level	spatial level	orientation	cod. rate (bits/coeff)	D_s	D
4	2	Approx.	1	416.69	561.54
3	2	Approx. (fr=1)	1	440.90	498.80
2	2	Approx. (fr=1)	0.125	221.71	221.73
1	2	Approx. (fr=4)	0.125	35.45	35.46
Temporal approximation frame					
temporal level	spatial level	orientation	cod. rate (bits/coeff)	D_s	D
4	2	Approx.	4	1708.39	13802.08
4	2	Vertical	2	2441.00	3982.33
4	2	Horizontal	2	1443.80	2173.68
4	2	Diagonal	2	426.6	619.95
4	1	Vertical	1	1272.71	1314.70
4	1	Horizontal	1	792.82	827.12

TAB. 6.7 – Source distortion D_s and global distortion D (at $SNR = 4.33$ dB) at a bitrate 295 kbs for different subbands of the "hall-monitor" sequence without motion estimation, fr designates the frame to which the examined subband belongs to.

6.6 Vector quantization by linear mapping of a block code

The Tables in the previous section prove that parts of the bitstream, especially in the temporal approximation frames, have a small source distortion compared with the end-to-end distortion, and therefore they should be protected against the errors induced by the channel. The conventional way, that we have already tested in Section 6.4.1, is to add an error-control channel coding, where redundancy is introduced in a controlled manner prior to transmission. By this way the transmission rate is however increased.

Another approach, which does not increase the transmission rate, and we shall develop in the sequel, is to find a better linear index assignment, by inserting the error-correcting code directly in this step. Previously, we supposed that the quantizer codevectors are mapped linearly to uncoded codewords, in this section we assume that the codeword belongs to a (n', k) linear block code and more specifically to a Reed-Muller $RM(r, m)$ block code, where $n' = 2^m$.

In general it is desirable to have short block code, producing codevectors that are separated as much as possible not only in terms of Hamming distance, but in terms of Euclidean distance too. The RM codes satisfy this constraint and as it is known, they have been largely used for the efficient construction of lattices. In addition, they provide the property of symmetry and they proved a powerful component in the error-control theory.

In our coding scheme, the choice of the specific RM code is based on the trade-off

"hall-monitor" 566 kbs					
Temporal detail frames					
temporal level	spatial level	orientation	cod. rate (bits/coeff)	D_s	D
4	2	Approx.	2	416.69	561.54
3	2	Approx. (fr=1)	1	143.93	172.91
2	2	Approx. (fr=1)	0.5	130.69	137.51
1	2	Approx. (fr=4)	0.125	35.45	35.77
Temporal approximation frame					
temporal level	spatial level	orientation	cod. rate (bits/coeff)	D_s	D
4	2	Approx.	8	11.63	6888.86
4	2	Vertical	4	137.17	844.27
4	2	Horizontal	4	116.69	937.40
4	2	Diagonal	2	426.60	704.78
4	1	Vertical	2	362.67	543.26
4	1	Horizontal	2	235.81	386.85
4	1	Diagonal	0.125	105.47	106.08

TAB. 6.8 – Source distortion D_s and global distortion D (at $SNR = 4.33$ dB) at a bitrate 566 kbs for different subbands of the hall-monitor sequence without motion estimation, fr shows the frame that the examined subband belongs to.

between the following three conditions :

- error correction capability of the code
- the blocklength of the code n' should be equal to n or r of $\mathbf{G}_{d,n}$ or $\mathbf{G}'_{n,r}$ respectively, in order not to increase the bitrate
- the dimension of the code k should be quite close to n or r of $\mathbf{G}_{d,n}$ or $\mathbf{G}'_{n,r}$ respectively, so that the number of 2^k possible codewords is close to the 2^n or 2^r possible codewords of the uncoded case, in order not to increase the source distortion.

6.6.1 Index assignment using $RM(r, m)$

In order to improve the index assignment and to benefit from a soft decoding method we impose a $RM(r, m)$ code as the mapping space of the codevectors. Regarding carefully the coding rate of the subbands where it is preferable to use a coded index assignment in the previous tables, we can remark that a $RM(2, 4)$ is the most flexible compromise between the three conditions exposed above.

The $RM(2, 4)$ is a linear code of blocklength $n' = 16$, dimension $k = 11$ and minimum distance $d_H = 4$. In this case, as the code is still short, we can use the Maximum-Likelihood decoding provided by the minimal trellis. Obviously, the number of the possible codewords is reduced to 2^{11} instead of 2^{16} possibilities for the uncoded case. The codevectors of the subbands which, according to the tables presented above, are not very well protected (i.e., $D_s \ll D$), are mapped to codewords belonging to this code. We stress again that, in this

"hall-monitor" 1027 kbs					
Temporal detail frames					
temporal level	spatial level	orientation	cod. rate (bits/coeff)	D_s	D
4	2	Approx.	4	44.22	344.88
3	2	Approx. (fr=1)	2	143.93	221.52
2	2	Approx. (fr=1)	2	38.45	58.84
1	2	Approx. (fr=4)	0.25	32.52	32.82
Temporal approximation frame					
temporal level	spatial level	orientation	cod. rate (bits/coeff)	D_s	D
4	2	Approx.	8	11.63	6519.56
4	2	Vertical	4	137.17	987.47
4	2	Horizontal	4	116.69	987.23
4	2	Diagonal	4	34.73	328.78
4	1	Vertical	4	22.90	163.96
4	1	Horizontal	2	235.81	382.55
4	1	Diagonal	1	71.42	74.02

TAB. 6.9 – Source distortion D_s and global distortion D (at $SNR = 4.33$ dB) at a bitrate 1027 kbs for different subbands of the hall-monitor sequence without motion estimation, fr shows the frame that the examined subband belongs to.

case, we do not add a linear block code, but we impose an index assignment where the codewords belong to a specific code. The double benefit of this approach is the protection that a linear code offers against errors, combined with the non increase of the global bitrate.

By imposing this specific code, it is obvious that keeping the *same source coding rate* and comparing to the case where the index assignment was uncoded,

- the source distortion increases as, now, we have 2^{11} possible codewords instead of 2^{16}
- the end-to-end distortion of the noisy system decreases
- the PSNR of the reconstructed video sequence increases at a given BER
- the total bitrate does not change.

We present now, the D_s and D when a $RM(2, 4)$ is applied to the subbands which, based on the previous tables, are proved to be not enough protected against channel errors.

An interesting remark concerning these last tables is related to the column of the D . Comparing it to the D of the uncoded case we see clearly that this kind of index assignment results in a lower value of the end-to-end distortion of the noisy system.

In Table 6.15 we present the average PSNR of the reconstructed video sequence when RM code is used to improve the index assignment.

From this example, we can clearly see that this approach is very robust in the presence of noise and significantly improves the PSNR of the reconstructed sequence.

"hall-monitor" 1550 kbs					
Temporal detail frames					
temporal level	spatial level	orientation	cod. rate (bits/coeff)	D_s	D
4	2	Approx.	2	44.22	331.2
3	2	Approx. (fr=1)	1	14.12	113.2
2	2	Approx. (fr=1)	0.5	38.45	57.62
1	2	Approx. (fr=4)	0.125	26.61	27.27
Temporal approximation frame					
temporal level	spatial level	orientation	cod. rate (bits/coeff)	D_s	D
4	2	Approx.	8	11.63	7492.25
4	2	Vertical	8	137.17	885.6
4	2	Horizontal	8	116.69	968.68
4	2	Diagonal	4	34.73	274.34
4	1	Vertical	4	22.90	171.77
4	1	Horizontal	4	19.41	157.67
4	1	Diagonal	2	25.97	38.76

TAB. 6.10 – Source distortion D_s and global distortion D (at $SNR = 4.33$ dB) at a bitrate 1550 kbs for different subbands of the hall-monitor sequence without motion estimation, fr shows the frame that the examined subband belongs to.

"hall-monitor" 295 kbs				
Temporal approximation frame				
temporal level	spatial level	orientation	D_s	D
4	2	Approx.	8469.00	8513.14

TAB. 6.11 – Source distortion D_s and global distortion D (at $SNR = 4.33$ dB) at a bitrate 295 kbs for different subbands of the “hall-monitor” sequence, when a $RM(2, 4)$ is chosen as the space of index assignment.

"hall-monitor" 566 kbs				
Temporal approximation frame				
temporal level	spatial level	orientation	D_s	D
4	2	Approx.	350.50	378.18
4	2	Vertical	542.07	549.89
4	2	Horizontal	550.58	555.5

TAB. 6.12 – Source distortion D_s and global distortion D (at $SNR = 4.33$ dB) at a bitrate 566 kbs for different subbands of the “hall-monitor” sequence, when a $RM(2, 4)$ is chosen as the space of index assignment.

"hall-monitor" 1027 kbs				
Temporal detail frames				
temporal level	spatial level	orientation	D_s	D
4	2	Approx.	196.97	226.94
Temporal approximation frame				
temporal level	spatial level	orientation	D_s	D
4	2	Approx.	350.50	382.53
4	2	Vertical	542.07	564.41
4	2	Horizontal	550.58	551.27
4	2	Diagonal	170.47	171.73
4	1	Vertical	109.69	129.44

TAB. 6.13 – Source distortion D_s and global distortion D (at $SNR = 4.33$ dB) at a bitrate 1027 kbs for different subbands of the “hall-monitor” sequence, when a $RM(2, 4)$ is chosen as the space of index assignment.

"hall-monitor" 1550 kbs				
Temporal detail frames				
temporal level	spatial level	orientation	D_s	D
4	2	Approx.	196.97	200.62
3	2	Approx. (fr=1)	64.61	65.62
Temporal approximation frame				
temporal level	spatial level	orientation	D_s	D
4	2	Approx.	350.50	403.33
4	2	Vertical	542.07	556.91
4	2	Horizontal	550.58	558.93
4	2	Diagonal	170.47	170.96
4	1	Vertical	109.69	109.92
4	1	Horizontal	88.00	88.19

TAB. 6.14 – Source distortion D_s and global distortion D (at $SNR = 4.33$ dB) at a bitrate 1550 kbs for different subbands of the hall-monitor sequence, when a $RM(2, 4)$ is chosen as the space of index assignment.

"hall-monitor"				
bitrate(kbs)	295	566	1027	1550
noiseless	26.68	30.53	32.81	34.66
SNR=4.33 dB	23.44	26.01	26.57	26.90
SNR=4.33 dB with RM	24.43	29.3	30.9	32.04

TAB. 6.15 – Average PSNR of the "hall-monitor" CIF sequence, without motion estimation with and without $RM(2, 4)$ chosen as the index assignment.

6.6.2 Overall results when $RM(r, m)$ is used for index assignment

All the previous remarks show that, in order to have a powerful JSC coding scheme with no additional error-correcting codes, the choice of an index assignment in the space of the $RM(r, m)$ is the appropriate method. Encoding then involves the following steps :

- In each subband we calculate the D_s (corresponding to a noiseless environment) and the end-to-end distortion D for the given noisy channel, as would be obtained with an uncoded index assignment.
- If the difference between them is high, which means that the D_e is significant, we restrict the index assignment to codewords of the $RM(r, m)$ (taking into consideration the three conditions presented previously); otherwise, we apply the mapping to uncoded codewords.
- At the decoder side : when we receive coded codewords, we apply soft decision decoding via the Viterbi algorithm with the BCJR trellis; otherwise, when we receive uncoded codewords, we apply a hard decision decoding.

In the remaining of this section, we present the performances obtained when the last procedure is applied to two different test sequences, for different bitrates and for different channel conditions. As in the previous Section, we distinguish two cases in our simulations : with and without motion estimation.

Experiments without motion estimation

Table 6.16 presents the average PSNR of the two test sequences obtained by applying the above procedure. If we compare it with the Table 6.5, we can clearly see the major increase of performance, especially in the case of low SNR. Note that the bitrate does not change, but the end-to-end distortion of the noisy system decreases with the use of $RM(2, 4)$. As we noticed previously, this index assignment is applied only on some selected subbands, and as the SNR becomes higher, the number of subbands benefiting from this protection decreases. Our JSC scheme, even with uncoded codewords, provides a good robustness to noise, as it appears especially from the results obtained for both of the sequences at $SNR = 6.75$ dB and lower bitrate (in these cases, we do not use any code).

"hall-monitor"				
bitrate (kbs)	295	566	1027	1550
noiseless	26.68	30.53	32.81	34.66
SNR=4.33 dB	24.43	29.30	30.90	32.04
SNR=6.75 dB	26.21	30.05	32.03	33.49
"foreman"				
bitrate (kbs)	307	576	983	1410
noiseless	25.07	27.56	28.81	30.14
SNR=4.33 dB	23.34	26.83	27.75	28.61
SNR=6.75 dB	24.86	27.27	28.40	29.57

TAB. 6.16 – Average PSNR of the "hall-monitor" and "foreman" CIF sequences, without motion estimation and index assignment with $RM(2, 4)$.

Experiments involving motion estimation

Table 6.17 presents the average PSNR for the two test sequences when motion estimation and RM as partial index assignment are used. As previously, if we compare these results with those in Table 6.6, we remark that the choice of the coded index assignment RM(2,4) is justified by the considerable increase in PSNR of the reconstructed sequence, as well as the increased robustness of the new coding scheme in the presence of noise, as compared to the uncoded case.

"hall-monitor"				
bitrate(kbs)	317	565	1113	1660
noiseless	27.88	31.78	34.38	35.59
SNR=4.33 dB	25.02	30.29	32.08	33.45
SNR=6.75 dB	27.22	31.15	33.36	34.18
"foreman"				
bitrate (kbs)	376	524	1104	1530
noiseless	30.07	30.78	33.30	34.51
SNR=4.33 dB	29.05	29.51	31.33	32.03
SNR=6.75 dB	29.40	30.01	32.00	32.85

TAB. 6.17 – Average PSNR of the "hall-monitor" and "foreman" CIF sequences, with motion estimation/compensation and index assignment with RM(2,4).

Fig. 6.6 illustrates the reconstructed frames of "hall-monitor" at 566 kbs without motion estimation and of "foreman" at 1104 kbs with motion estimation, first in a noiseless environment and then after transmission over a Gaussian channel with SNR=4.33 dB and SNR=6.75 dB when $RM(2,4)$ is applied as index assignment on part of the subbands.

A general remark for the results presented above is that the partial coded index assignment is proved to be necessary for the subbands with high energy. Most of these subbands are part of the temporal approximation frames, coarsest spatial resolution level. However, as the bitrate increases, we remark that the approximation subband of some temporal detail frames, especially at the coarser temporal resolution levels, also needs to be protected. If we compare the Fig. 6.6, where we presented visually the result of this coding scheme, with the Fig. 6.5, where the index assignment is uncoded, we remark that protecting the subbands with high energy has as consequence that there are no more "spots" (corresponding to decoding errors) in the reconstructed frames, even at low SNR.

6.7 Structured vs. unstructured codebook with index assignment

In [70], as we have already mentioned in Chapter 4, a good suboptimal minimax index assignment achieved by a polynomial-time algorithm is presented. They consider a finite source codebook $C = \{c_1, \dots, c_N\}$ consisting of N elements. The codeword c_i occurs with probability $p(i)$. Let us denote by $d(c_i, c_j)$ the distance between the codewords c_i and c_j



FIG. 6.6 – Reconstructed frames of "foreman" at 1104 kbs with motion estimation and "hall-monitor" at 566 kbs without motion estimation when RM is partially applied as index assignment. First line : reconstructed frames in a noiseless environment. Second line : reconstructed frames after transmission over a Gaussian channel with $\text{SNR}=4.33$ dB. Third line : reconstructed frames after transmission over a Gaussian channel with $\text{SNR}=6.75$ dB.

in the sense of the squared Euclidean norm. The authors adopt a fixed length binary code of n bits per codeword, so $N = 2^n$ is the size of the codebook.

Any assignment of a binary code from $\{0, 1\}^n$ to each codeword in C is simply obtained by a permutation π of the integers $\{1, 2, \dots, N\}$ and association of binary index i to codeword $c_{\pi(i)}$. Let Π denote the set of $N!$ permutations.

The goal of the nonredundant coding task is to find a permutation π in order to minimize the message distortion introduced by channel noise. In [70] the authors adopt a minimax error as measure of distortion and not an average distortion performance criterion. However, in each case the optimal permutation, intuitively, result in a large Hamming distance between the binary codewords for c_i and c_j when $d(c_i, c_j)$ is large. Adopting a minimax error criterion designs a channel code against the worst case performance. However, their algorithm, named Minimax Cover Algorithm (MCA), does not sacrifice the average performance. Indeed, they formulate the minimax nonredundant channel coding task as a graph problem, and they develop an heuristic algorithm for obtaining good suboptimal solutions. The initial codebook is generated using the GLA and the square Euclidean distance is used in training the codebook.

The authors in [70] note that the minimax error criterion is more tractable for image coding than an average error criterion, as the visual perception of quality is not proportional to the distortion amplitude and larger errors contribute much more to perceived degradation. Thus, we choose this implementation to compare with our coding scheme which, on the contrary, is based on a structured codebook with linear index assignment. We applied the minimax algorithm to the codebook produced by GLA from our video sequences. The total bitrates (in kbs) for this unstructured scheme are the same as the ones produced by our structured schemes. The size of the unstructured codebook for each subband is as prescribed by the bit allocation algorithm presented previously. However, due to the random generation of the codebook via GLA, we have implemented four different generations- four different codebooks- and we have chosen, as final, the codebook that gives an average PSNR closer to the mean average PSNR of all four, in case of low SNR.

In Tables 6.18, 6.19, we present the results of the Minimax Cover Algorithm (MCA), compared with our two coding schemes applied to two video sequences without motion estimation, the first using uncoded assignment (JSC) and the second one using partially coded assignment with $RM(2, 4)$ (JSC+RM).

In Tables 6.20, 6.21 we performed the same comparison, with the only difference that motion estimation/compensation is applied to the two test video sequences.

As expected, due to the unstructured nature of the codebook generated by GLA, in the noiseless case, the PSNR of the MCA scheme is quite high as compared to that of the structured-codebook schemes. On the other hand, as the channel noise increases, the PSNR for the unstructured MCA drops dramatically despite the index assignment provided by the minimax algorithm. Meanwhile, one can clearly observe that both our coding schemes, especially the one using coded assignment, remain quite close to their noiseless performance, even for very noisy channels (i.e. low SNR). Indeed, except for "foreman" when no motion estimation is applied at very low bitrate, JSC+RM using partially coded index assignment outperforms both the uncoded JSC approach as well as the unstructured MCA technique when the channel SNR is very low.

"hall-monitor" 1550 kbs			
channel state	noiseless	SNR=4.33 dB	SNR=6.75 dB
JSC	34.66	26.90	32.84
JSC+RM		32.04	33.49
MCA	40.35	24.35	33.29
"hall-monitor" 1027 kbs			
channel state	noiseless	SNR=4.33 dB	SNR= 6.75 dB
JSC	32.81	26.57	31.54
JSC+RM		30.90	32.03
MCA	38.18	23.65	32.15
"hall-monitor" 566 kbs			
channel state	noiseless	SNR=4.33 dB	SNR=6.75 dB
JSC	30.53	26.01	29.77
JSC+RM		29.30	30.05
MCA	34.96	23.58	31.43
"hall-monitor" 295 kbs			
channel state	noiseless	SNR= 4.33 dB	SNR=6.75 dB
JSC	26.68	23.44	26.21
JSC+RM		24.43	26.21
MCA	32.45	23.53	30.04

TAB. 6.18 – Average PSNR of the reconstructed sequence "hall-monitor" without motion estimation for different bitrates and different SNR. JSC denotes our first coding scheme with structured quantizer and uncoded assignment, *JSC+RM* the second one with partial coded index assignment and finally, MCA describes the algorithm using unstructured quantizer and minimax index assignment.

Fig. 6.7 illustrates the reconstructed frames of "hall-monitor" at 566 kbs without motion estimation and of "foreman" at 1104 kbs with motion estimation, first in a noiseless environment and then after transmission over a Gaussian channel with SNR=4.33 dB and SNR=6.75 dB when the unstructured scheme with MCA is used. One can easily remark the degradation of the quality of the reconstructed frames as compared to Fig. 6.5, where *JSC + RM* is used.



FIG. 6.7 – Reconstructed Frames of "foreman" at 1104 kbs with motion estimation and "hall-monitor" at 566 kbs without motion estimation using GLA and minimax index assignment. First line : reconstructed frames in a noiseless environment. Second line : reconstructed frames after transmission over a Gaussian channel with $\text{SNR}=4.33$ dB. Third line : reconstructed frames after transmission over a Gaussian channel with $\text{SNR}=6.75$ dB.

"foreman" 1410 kbs			
channel state	noiseless	SNR=4.33 dB	SNR=6.75 dB
JSC	30.14	25.44	29.36
JSC+RM		28.61	29.57
MCA	36.10	25.02	32.60
"foreman" 983 kbs			
channel state	noiseless	SNR=4.33 dB	SNR= 6.75 dB
JSC	28.81	25.04	28.22
JSC+RM		27.75	28.40
MCA	34.14	25.02	31.65
"foreman" 576 kbs			
channel state	noiseless	SNR=4.33 dB	SNR=6.75 dB
JSC	27.56	24.71	27.13
JSC+RM		26.83	27.27
MCA	31.96	24.79	30.35
"foreman" 307 kbs			
channel state	noiseless	SNR= 4.33 dB	SNR=6.75 dB
JSC	25.07	22.98	24.86
JSC+RM		23.34	24.86
MCA	29.77	23.64	28.49

TAB. 6.19 – Average PSNR of the reconstructed sequence "foreman" without motion estimation for different bitrates and different SNR. JSC denotes our first coding scheme with structured quantizer and uncoded assignment, *JSC+RM* the second one with partial coded index assignment and finally, MCA describes the algorithm using unstructured quantizer and minimax index assignment.

6.8 Comparison with a MC-EZBC protected with rate punctured convolutional codes

We consider a packet-based motion-compensated embedded zerotree block coder (MC-EZBC). The input video sequence is compressed with an MC-EZBC based encoder and the resulting file is partitioned into packets of unequal length. Each packet contains the information corresponding to a spatial resolution level of each temporal subband frame. We consider a five-level spatial decomposition with the resulting bitstream that represents the low-frequency temporal detail subband kept alone in one packet. The resulting packets are encoded by RPC codes before the transmission through the channel. If the decoder fails to decode the received packet it drops this packet and the decoding procedure continues with the next packet received.

An unequal error protection (UEP) is applied to the packets using RPC codes. The RPC codes chosen are the $R_p = 2/3, 3/4, 7/8$ with memory $m=6$ and mother code $R = 1/2$ [38]. In the temporal detail frames the distribution of the coding protection is illustrated in Fig. 6.8, while in the temporal approximation frame all the packets are

"hall-monitor" 1660 kbs			
channel state	noiseless	SNR=4.33 dB	SNR=6.75 dB
JSC	35.59	27.23	33.47
JSC+RM		33.45	34.18
MCA	40.31	24.15	33.13
"hall-monitor" 1113 kbs			
channel state	noiseless	SNR=4.33 dB	SNR= 6.75 dB
JSC	34.38	26.96	32.68
JSC+RM		32.08	33.36
MCA	38.03	23.90	32.54
"hall-monitor" 565 kbs			
channel state	noiseless	SNR=4.33 dB	SNR=6.75 dB
JSC	31.78	26.41	30.81
JSC+RM		30.29	31.15
MCA	35.36	23.78	31.61
"hall-monitor" 317 kbs			
channel state	noiseless	SNR= 4.33 dB	SNR=6.75 dB
JSC	27.88	23.91	27.22
JSC+RM		25.02	27.22
MCA	32.81	23.54	30.31

TAB. 6.20 – Average PSNR of the reconstructed sequence "hall-monitor" with motion estimation for different bitrates and different SNR. JSC denotes our first coding scheme with structured quantizer and uncoded assignment, *JSC+RM* the second one with partial coded index assignment and finally, MCA describes the algorithm using unstructured quantizer and minimax index assignment.

protected by the RPC code of rate $R_p = 2/3$.

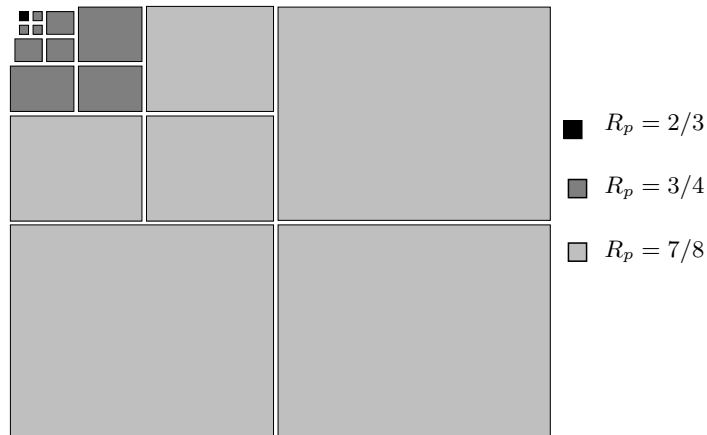


FIG. 6.8 – Unequal Error Protection of a MC-EZBC using RPC codes.

"foreman" 1530 kbs			
channel state	noiseless	SNR=4.33 dB	SNR=6.75 dB
JSC	34.51	26.49	32.63
JSC+RM		32.03	32.85
MCA	38.72	25.62	34.06
"foreman" 1104 kbs			
channel state	noiseless	SNR=4.33 dB	SNR= 6.75 dB
JSC	33.30	26.27	31.72
JSC+RM		31.33	32.00
MCA	37.44	25.59	33.16
"foreman" 524 kbs			
channel state	noiseless	SNR=4.33 dB	SNR=6.75 dB
JSC	30.78	25.66	29.86
JSC+RM		29.51	30.01
MCA	34.80	25.22	32.08
"foreman" 376 kbs			
channel state	noiseless	SNR= 4.33 dB	SNR=6.75 dB
JSC	30.07	25.53	29.27
JSC+RM		29.05	29.40
MCA	33.64	24.99	31.38

TAB. 6.21 – Average PSNR of the reconstructed sequence "foreman" with motion estimation for different bitrates and different SNR. JSC denotes our first coding scheme with structured quantizer and uncoded assignment, $JSC + RM$ the second one with partial coded index assignment and finally, MCA describes the algorithm using unstructured quantizer and minimax index assignment.

The total bitrates (in kbs) of the "foreman" and "hall-monitor" CIF test sequences with the MC-EZBC+RPC coding scheme are the same as the ones produced by our structured scheme.

In Tables 6.22, 6.23 we present the results of the MC-EZBC+RPC coding scheme compared with our JSC+RM coding scheme applied to two video sequences with motion estimation. The noiseless performance of the MC-EZBC scheme is quite high compared to our coding scheme. However, when the channel noise increases (SNR=4.33 dB), even if we do not remark the same level of degradation as for the unstructured quantizer presented in the previous Section, we note a higher performance of the JSC+RM coding scheme (about 0.77 – 4.22 dB) except for the lowest bitrate of the "hall-monitor" sequence. In addition, the MC-EZBC+RPC coding scheme appears to have a difference in PSNR of about 5.75 – 13.27 dB from its noiseless performance while the JSC+RM coding scheme has a degradation only of 1.02 – 2.86 dB when the channel noise is high ($SNR = 4.33$ dB).

"hall-monitor" 1660 kbs			
channel state	noiseless	SNR=4.33 dB	SNR=6.75 dB
JSC+RM	35.59	33.45	34.18
MC-EZBC	42.16	32.46	40.79
"hall-monitor" 1113 kbs			
channel state	noiseless	SNR=4.33 dB	SNR= 6.75 dB
JSC+RM	34.38	32.08	33.36
MC-EZBC	41.13	27.86	40.26
"hall-monitor" 565 kbs			
channel state	noiseless	SNR=4.33 dB	SNR=6.75 dB
JSC+RM	31.78	30.29	31.15
MC-EZBC	39.30	27.85	38.90
"hall-monitor" 317 kbs			
channel state	noiseless	SNR= 4.33 dB	SNR=6.75 dB
JSC+RM	27.88	25.02	27.22
MC-EZBC	36.94	26.36	36.33

TAB. 6.22 – Average PSNR of the reconstructed sequence "hall-monitor" for different bitrates and different SNR. JSC+RM denotes our coding scheme with partial coded index assignment and MC-EZBC is the standard coder followed by UEP by RPC codes.

6.9 Conclusion

In this Chapter, we presented the necessary modifications to the joint source-channel coding scheme introduced in Chapter 5, in order to be extended to video sequences and we explained step-by-step the reasons that led us to them.

In an early attempt, we evaluated the performance of our coding scheme with the partial addition of an RPC code, which convinced us that the coding scheme is quite robust in the presence of noise. We then developed a bit allocation algorithm which takes into account the nonnegativity constraint and provides an increased flexibility to our coding scheme. We presented the performance of our coding scheme using this bit allocation algorithm and uncoded index assignment on a Gaussian channel for different SNR and different bitrates. Several remarks, in this case, led us to a second approach which consists in partially applying a coded index assignment using Reed-Muller codes. This later approach led us to very good results, especially at low SNR.

Finally, we compared our coding schemes to an unstructured quantizer with minimax index assignment and a standard video coder with unequal error protection by RPC codes. We remarked that our coding schemes, especially the one with coded index assignment, proved to be more robust especially over very noisy channels. The proposed structured coders obtain between 1 and 5 dB greater average PSNR over very noisy channels compared to the unstructured quantizer and between 0.77 and 4.22 dB compared to the standard video coder. This feature makes our coding schemes attractive for wireless applications.

"foreman" 1530 kbs			
channel state	noiseless	SNR=4.33 dB	SNR=6.75 dB
JSC+RM	34.51	32.03	32.85
MC-EZBC	39.41	30.15	38.73
"foreman" 1104 kbs			
channel state	noiseless	SNR=4.33 dB	SNR= 6.75 dB
JSC+RM	33.30	31.33	32.00
MC-EZBC	38.00	29.08	37.35
"foreman" 524 kbs			
channel state	noiseless	SNR=4.33 dB	SNR=6.75 dB
JSC+RM	30.78	29.51	30.01
MC-EZBC	34.49	28.74	34.10
"foreman" 376 kbs			
channel state	noiseless	SNR= 4.33 dB	SNR=6.75 dB
JSC+RM	30.07	29.05	29.40
MC-EZBC	32.59	26.27	32.11

TAB. 6.23 – Average PSNR of the reconstructed sequence "foreman" for different bitrates and different SNR. JSC+RM denotes our coding scheme with partial coded index assignment and MC-EZBC is the standard coder followed by UEP by RPC codes.

Chapitre 7

Joint source-channel coding of a video on a flat-Rayleigh fading channel

In this Chapter, we first evaluate the performance of our structured quantizer with linear index assignment on a flat-Rayleigh channel. We introduce a new scheme, where the previous joint source-channel coding scheme is concatenated with a rotation matrix prior to transmission over a Rayleigh channel. We illustrate by simulation results the increase in performance of this new scheme and compare it with standard video coder over the Rayleigh channel.

7.1 Introduction

A good model for several terrestrial mobile radio channels is a flat-Rayleigh fading channel. The multipath propagation environment and the time variations of the channel due to the relative movement of the transmitter and the receiver are the main features. If all the frequency components in the transmitted signal are affected by the same random attenuation and phase shift then the channel is called *frequency-flat*. This happens when the duration of the modulated symbol is much greater than the delay spread caused by multipath propagation. The delay spread is just the difference between the largest and the smallest among the delays of the various paths. If, in addition, the channel varies very slowly with respect to the symbol duration, then the fading remains approximately constant during the transmission of one symbol.

On a fading channel, errors occur in reception when the channel attenuation is large. However, if we can supply to the receiver several replicas of the same information signal transmitted over independently fading channels, the probability that all the signal components will fade simultaneously is reduced considerably. For this reason, several diversity techniques have been developed, such as frequency diversity, time diversity techniques or diversity techniques based on multiple receiving antennas.

Another recent method, as we have already presented, is the use of multidimensional modulation schemes with an inherent diversity order, given by the Hamming distance distribution of the constellation points. The common way to increase the modulation diversity is to apply a certain rotation to a classical signal constellation in such a way

that any two points achieve the maximum number of distinct components.

Fig 7.1 illustrates the increase of diversity when a 4-PSK is rotated [9]. The second constellation, the rotated one, offers a better protection against the noise, since no two points collapse together.

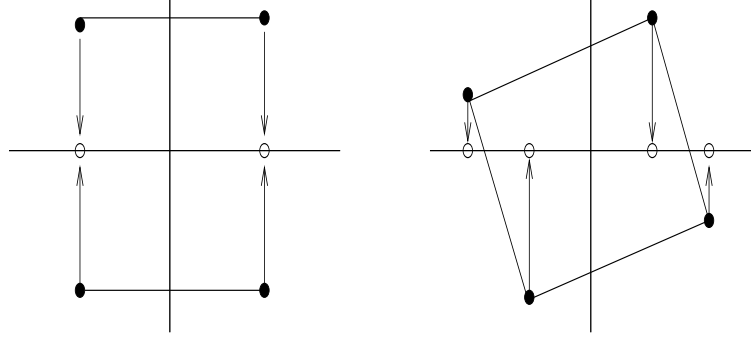


FIG. 7.1 – Increasing of the diversity with rotated constellation. Left : Diversity $L = 1$. Right : Diversity $L = 2$.

In this Chapter, we present a first attempt using, partially, rate-punctured convolutional codes in order to protect our coding scheme on the fading channel and next, we propose another method using rotation matrices prior to the transmission, which is proved to be very efficient.

7.2 Linear Labelling and Puntured Convolutional Codes on a flat-Rayleigh fading Channel

In this early attempt, as for the Gaussian channel presented in Section 6.4.1, we do not take into consideration the bit allocation algorithm and we apply a partial additional protection by a rate punctured convolutional code. The inherent property of our quantization matrix, resulting from "maximum component diversity" lattice constellations, the punctured convolutional code and a perfect interleaving provide a coding scheme with high diversity.

We recall that the four coding scenarii we examined, as in the case of Gaussian channel, are the following :

Type 1 : all the spatial subbands of the temporal approximation frames are protected by the punctured code.

Type 2 : only the subbands at the coarsest spatial resolution level of the approximation frames are protected.

Type 3 : only the lowest frequency spatial subband of the approximation frames is protected.

Type 4 : there is no additional protection. The robustness is achieved by the use of only our JSC coding method.

As in the case of the Gaussian channel in Section 6.4.1, when no algorithm of bit allocation is applied, we use a 16×4 matrix $\mathbf{G}'_{16,4}$ in order to quantize the wavelet coefficients

of the subbands of the temporal detail frames and a 2×16 matrix $\mathbf{G}_{2,16}$ to quantize the coefficients of the subbands of the temporal approximation frames.

In this case we used a code of rate $R_p = 3/4$, coming from a mother code of rate $R = 1/2$ and memory $m = 4$. In Fig. 7.2 we present its performance over a flat-Rayleigh fading channel, when the channel state information (CSI) is known and when it is not.

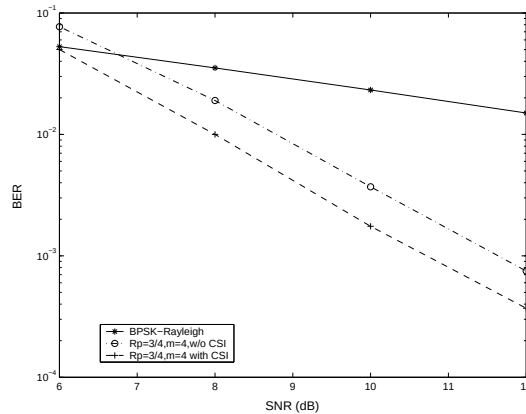


FIG. 7.2 – Performance of the $R_p = 3/4$ punctured convolutional code over a Rayleigh fading channel without and with CSI embedded in the Viterbi decoder.

In the sequel, we consider a memoryless channel. This assumption can be approached in practice by means of perfect interleaving performed at the coordinate level. By this way, the combination of the interleaving and the convolutional coding provides an even higher channel diversity.

A hard decision criterion is applied at the decoder for the temporal detail frames and for the subbands of the approximation frames that are not protected by the punctured code. For the part of the bitstream where the punctured code was used, a soft decision decoding by the Viterbi algorithm is applied. We make the distinction between two cases in the use of the Viterbi decoder. In the first we suppose that we don't have information about the channel state ("without CSI") and in the second the CSI is embedded in the path metric of the Viterbi algorithm ("with CSI").

In a noiseless environment and for the CIF test sequence "foreman" at 30 fps we get an average PSNR equal to 31.51dB.

In Table 7.1 we present the PSNR averaged over the noise realizations of a GOF in our video test sequence obtained with the four types of coding mentioned above and for the cases where the CSI is embedded or not in the Viterbi decoder. We notice that the proposed quantization scheme in itself offers good robustness to noise and even a small additional protection of the approximation frames, e.g in *Type 3*, significantly increases the performance. In addition we remark the significant increase in coding performance, especially for lower SNR, when we supposed that the CSI can be estimated and embedded in the path metric of the Viterbi decoder.

Fig. 7.3 illustrates the reconstructed frames, first in a noiseless environment and then after transmission over a Rayleigh fading channel with SNR=10.00 dB. For this latter case we present the reconstructed frames using the *Type 4* coding scheme, meaning no

Without CSI				
	Type 1	Type 2	Type 3	Type 4
SNR=8.0	23.85	23.79	23.71	21.37
SNR=10.0	28.73	28.51	28.23	22.96
SNR=12.0	30.74	30.47	30.13	24.53
with CSI				
	Type 1	Type 2	Type 3	Type 4
SNR=8.0	25.83	25.68	25.50	21.37
SNR=10.0	29.81	29.48	29.10	22.96
SNR=12.0	31.06	30.76	30.40	24.53

TAB. 7.1 – Average PSNR of "foreman" CIF sequence, coded by the four schemes and for three different states of the Rayleigh channel, with or without CSI embedded in the Viterbi decoder.

additional protection (second line, first image) and the *Type 2* coding scheme, without and with CSI. **Note that all these frames have been reconstructed after being affected by the same noise realization.**

7.3 Joint source-channel coding using rotations prior to transmission over a flat-Rayleigh channel

The results of the previous section belong to an early attempt without even applying a bit allocation algorithm. However, they clearly show that a BPSK modulation is significantly damaged during its transmission through the fading channel when no protection is applied. In this Section, we develop the idea of rotating the given constellation in order to increase its diversity order without adding any redundancy. An interleaver/deinterleaver is required to assure that the components of the received symbol are affected by independent fading.

The matrix $\mathbf{G}_{d,n}$ (or $\mathbf{G}'_{n,r}$) which generates our structured codebook, as already presented, is the result of choosing d rows, or (r columns) of an $n \times n$ U_n generator matrix of "maximum component diversity" lattice constellation or, other way speaking, of a full rotated $\mathbb{Z}^n$ constellation. Thus, the idea is the following : prior to the transmission through the channel we apply the matrix of rotation $\mathbf{U}_n$ or $\mathbf{U}_r$ corresponding to the $\mathbf{G}_{d,n}$ or $\mathbf{G}'_{n,r}$ respectively, and the decoding is based on the sphere decoder. In this way, we obtain the benefit of the increased diversity due to the rotation and the benefit of a maximum likelihood decoding method at the decoder.

In Fig. 7.4, $\mathbf{x}$ is the vector of wavelet coefficients that will be quantized and linearly labeled to a vector $\mathbf{b} \in \{-1, +1\}^n$ through the matrix $\mathbf{G}_{d,n}$ (or $\mathbf{G}'_{n,r}$), $\mathbf{u}$ is the vector of real values that is obtained by applying the rotation matrix $\mathbf{U}_n$ or $\mathbf{U}_r$.

The set of all points $\{\mathbf{u} = \mathbf{b}\mathbf{U}_n, \mathbf{b} \in \{-1, +1\}^n\}$ (or $\{\mathbf{u} = \mathbf{b}\mathbf{U}_r, \mathbf{b} \in \{-1, +1\}^r\}$) belongs to the n -dimensional (or r -dimensional) cubic lattice $\mathbb{Z}_{n,L=n}$ (or $\mathbb{Z}_{r,L=r}$) with generator matrix $\mathbf{U}_n$ or $\mathbf{U}_r$ and diversity L . The $\mathbb{Z}_{n,L=n}$ (or $\mathbb{Z}_{r,L=r}$) are rotated $\mathbb{Z}^n$ (or $\mathbb{Z}^r$) lattices



FIG. 7.3 – Reconstructed Frames. First line : reconstructed frame in a noiseless environment. Second line : reconstructed frame when no protection by a code is used over a Rayleigh fading channel with SNR=10.00 dB. Third line, Left : the reconstructed frame when only the lowest frequency spatial subband of the temporal approximation frame is protected and no CSI is supposed Right : Same type of protection as above but with CSI embedded in the Viterbi algorithm. (Rayleigh channel with SNR=10.00 dB).

that guarantee the maximum degree of diversity.

The channel is a flat-fading Rayleigh channel and we assume that the CSI is known at the receiver. The independent fading is satisfied by the assumption of using the interleaver. Thus, the received vector after deinterleaving is given by :

$$\mathbf{r} = \mathbf{a} * \mathbf{u} + \mathbf{n}'$$

where $\mathbf{a}$ is the vector of the random fading coefficients with Rayleigh distribution and $E[a_i^2] = 1$ and $\mathbf{n}'$ is the noise vector with Gaussian distributed independent random variables with zero mean and variance N_0 .

The sphere decoder with perfect CSI, as presented in a previous chapter, minimizes the following metric :

$$m(\mathbf{u}|\mathbf{r}, \mathbf{a}) = \sum_{i=1}^n |r_i - a_i u_i|^2$$

Thus, as the Sphere decoder extract the integer components of the $\mathbf{u}$ lattice point, we take the corresponding $\hat{\mathbf{b}}$ vector and, after a simple multiplication with $\mathbf{G}_{d,n}$ (or $\mathbf{G}'_{n,r}$), the $\hat{\mathbf{x}}$ vector of the wavelet coefficients.

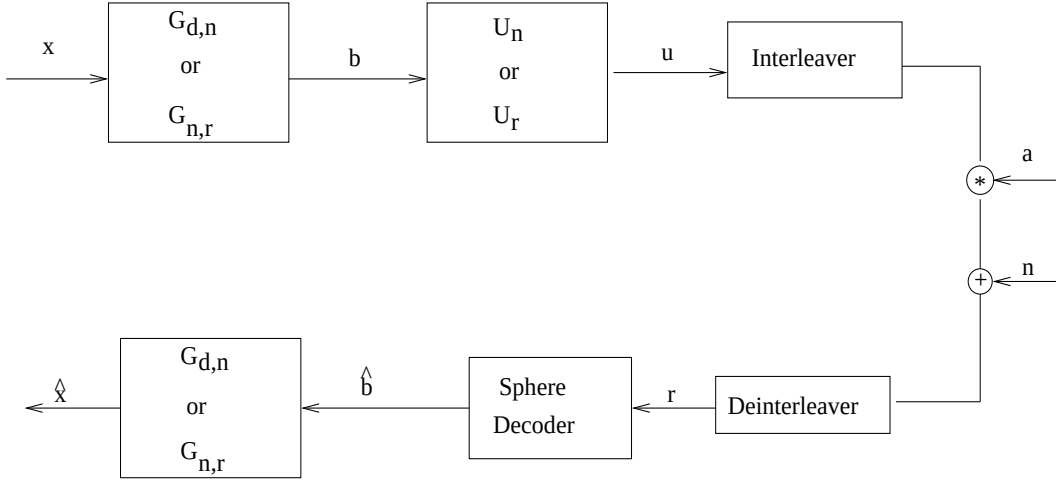


FIG. 7.4 – System model with rotation prior fading channel.

It is obvious that we can use a rotation matrix whose dimension is higher than the dimension of the vector at the output of the quantizer, in order to achieve a higher order of diversity, but this, unfortunately, will increase the complexity of the decoder. Moreover, keeping the dimension less than or equal to 16, as we did in the case of the Gaussian channel, allows us to use the sphere decoder algorithm (remind that its complexity limited us to use dimensions less than or equal to 32). However, we could use some improvements of the sphere decoder algorithm or another lattice decoder, i.e. in [48], [49] a mean-square universal lattice decoder suitable for large dimensions up to 1024 is presented. Finally, at the place of the matrix U_n we can use other matrices of rotation [4], which yield maximum component diversity.

In Tables 7.2, 7.3 we present some simulations results, obtained for the same bitrates and sequences as in the case of the Gaussian channel when no motion estimation is applied. We present, also, the case where no rotation is used in order to make clear the major benefit of using a rotation matrix prior to the transmission on the fading channel.

In Tables 7.4, 7.5 we performed the same comparison, with the only difference that motion estimation/compensation is applied to the two test video sequences.

It is obvious that, when using a rotation matrix, no redundancy by an error-correcting code is anymore necessary.

"hall-monitor" 1550 kbs		
noiseless=34.66 dB	Without rotation	With rotation
SNR=9.0	22.82	33.38
SNR=11.0	24.51	34.47
"hall-monitor" 1027 kbs		
noiseless=32.81 dB	Without rotation	With rotation
SNR=9.0	22.74	32.01
SNR=11.0	24.37	32.74
"hall-monitor" 566 kbs		
noiseless=30.53 dB	Without rotation	With rotation
SNR=9.0	22.49	29.96
SNR=11.0	24.01	30.47
"hall-monitor" 295 kbs		
noiseless=26.68 dB	Without rotation	With rotation
SNR=9.0	20.43	26.36
SNR=11.0	21.80	26.65

TAB. 7.2 – Average PSNR of the reconstructed "hall-monitor" CIF sequence without motion estimation when transmitted over a flat-Rayleigh channel, with and without rotation matrix applied prior to transmission.

In Fig. 7.5 and 7.6 the reconstructed frames of the two sequences "hall-monitor" and "foreman" after transmission over a flat-Rayleigh fading channel are illustrated. We present the reconstructed frames for two cases : with rotation matrix prior transmission and without one. The quality of the reconstructed frames, when a rotation matrix is used, is very high compared to the opposite case, even for very low SNR.

7.4 Comparison with the MC-EZBC protected by rate punctured convolutional codes over a flat-Rayleigh channel

In this Section we use the MC-EZBC+RPC coding scheme, presented in the Section 6.8 of the previous Chapter, for comparison reasons over a flat-Rayleigh fading channel. We performed UEP as illustrated in the Fig. 6.8 for the subbands of the temporal detail

"foreman" 1410 kbs		
noiseless=30.14 dB	Without rotation	With rotation
SNR=9.0	21.93	29.56
SNR=11.0	23.45	30.10
"foreman" 983 kbs		
noiseless=28.81 dB	Without rotation	With rotation
SNR=9.0	21.86	28.41
SNR=11.0	23.27	28.78
"foreman" 576 kbs		
noiseless=27.56 dB	Without rotation	With rotation
SNR=9.0	21.60	27.28
SNR=11.0	22.89	27.54
"foreman" 307 kbs		
noiseless=25.07 dB	Without rotation	With rotation
SNR=9.0	20.57	24.94
SNR=11.0	21.70	25.03

TAB. 7.3 – Average PSNR of the reconstructed "foreman" CIF sequence without motion estimation when transmitted over a flat-Rayleigh channel, with and without rotation matrix applied prior to transmission.

"hall-monitor" 1660 kbs		
noiseless=35.59 dB	Without rotation	With rotation
SNR=9.0	23.12	33.93
SNR=11.0	24.86	35.46
"hall-monitor" 1113 kbs		
noiseless=34.38 dB	Without rotation	With rotation
SNR=9.0	22.84	33.13
SNR=11.0	24.56	34.23
"hall-monitor" 565 kbs		
noiseless=31.78 dB	Without rotation	With rotation
SNR=9.0	22.71	31.07
SNR=11.0	24.27	31.72
"hall-monitor" 317 kbs		
noiseless=27.88 dB	Without rotation	With rotation
SNR=9.0	20.66	27.48
SNR=11.0	22.06	27.85

TAB. 7.4 – Average PSNR of the reconstructed "hall-monitor" sequence with motion estimation when transmitted over a flat-Rayleigh channel, with and without rotation matrix applied prior to transmission.

"foreman" 1530 kbs		
noiseless=34.51 dB	Without rotation	With rotation
SNR=9.0	22.35	33.05
SNR=11.0	24.06	34.41
"foreman" 1104 kbs		
noiseless=33.30 dB	Without rotation	With rotation
SNR=9.0	22.29	32.22
SNR=11.0	23.97	33.23
"foreman" 524 kbs		
noiseless=30.78 dB	Without rotation	With rotation
SNR=9.0	22.10	30.16
SNR=11.0	23.67	30.73
"foreman" 376 kbs		
noiseless=30.07 dB	Without rotation	With rotation
SNR=9.0	22.04	29.53
SNR=11.0	23.54	30.03

TAB. 7.5 – Average PSNR of the reconstructed "foreman" sequence with motion estimation when transmitted over a flat-Rayleigh channel, with and without rotation matrix applied prior to transmission.

frames while all the packets containing the bitstream corresponding to the subbands from the temporal approximation frames are protected equally using the RPC of rate $R_p = 2/3$ and $m = 6$ [38].

A soft decision decoding by the Viterbi algorithm is performed with the CSI embedded in the path metrics.

In Table 7.6, 7.7 we present the results of the MC-EZBC+RPC coding scheme compared with our JSC coding scheme with rotation matrix applied to two video sequences with motion estimation over a flat-Rayleigh fading channel. The average PSNR of the JSC coding scheme with a rotation matrix is between 2 and 5 dB higher than the one gained with the MC-EZBC+RPC coding scheme for a low SNR ($SNR = 9.0$ dB). However, even if for an $SNR = 11.0$ dB the MC-EZBC+RPC coding scheme outperforms the JSC coding scheme, especially at low bitrates, the difference from its noiseless performance (between 0.85 and 7.74 dB) is much higher than the one of the JSC coding scheme (between 0.03 and 0.15 dB).

7.5 Conclusion

In this Chapter, we evaluated the performance of our coding scheme on a flat-Rayleigh channel, which is a good model for wireless channels. First, we examined the conventional way, which consists of adding redundancy by an error-correcting code, a RPC code in our case. The results, in this case, proved quite promising.

Next, taking into account the role of diversity in a fading environment, we adopted a



FIG. 7.5 – Reconstructed Frames. First line : reconstructed frame in a noiseless environment. Second line : reconstructed frame after transmission over a Rayleigh fading channel with $\text{SNR}=9.00$ dB ; Left : when no rotation matrix is used. Right : with rotation matrix prior transmission. Third line : reconstructed frame after transmission over a Rayleigh fading channel with $\text{SNR}=11.00$ dB ; Left : when no rotation matrix is used. Right : with rotation matrix prior transmission.



FIG. 7.6 – Reconstructed Frames. First line : reconstructed frame in a noiseless environment. Second line : reconstructed frame after transmission over a Rayleigh fading channel with $\text{SNR}=9.00$ dB ; Left : when no rotation matrix is used. Right : with rotation matrix prior transmission. Third line : reconstructed frame after transmission over a Rayleigh fading channel with $\text{SNR}=11.00$ dB ; Left : when no rotation matrix is used. Right : with rotation matrix prior transmission.

"hall-monitor" 1660 kbs		
channel state	SNR=9.0 dB	SNR=11.0 dB
JSC	33.93	35.46
MC-EZBC	30.36	34.42
"hall-monitor" 1113 kbs		
channel state	SNR=9.0 dB	SNR=11.0 dB
JSC	33.13	34.23
MC-EZBC	28.22	34.36
"hall-monitor" 565 kbs		
channel state	SNR=9.0 dB	SNR=11.0 dB
JSC	31.07	31.72
MC-EZBC	27.57	34.33
"hall-monitor" 317 kbs		
channel state	SNR=9.0 dB	SNR=11.0 dB
JSC	27.48	27.85
MC-EZBC	25.21	33.89

TAB. 7.6 – Average PSNR of the reconstructed sequence "hall-monitor" with motion estimation for different bitrates and different SNR. JSC denotes our coding scheme with the rotation matrix and MC-EZBC is the standard coder followed by UEP by RPC codes.

"foreman" 1530 kbs		
channel state	SNR=9.0 dB	SNR=11.0 dB
JSC	33.05	34.41
MC-EZBC	29.79	33.18
"foreman" 1104 kbs		
channel state	SNR=9.0 dB	SNR=11.0 dB
JSC	32.22	33.23
MC-EZBC	28.14	32.90
"foreman" 524 kbs		
channel state	SNR=9.0 dB	SNR=11.0 dB
JSC	30.16	30.73
MC-EZBC	27.64	32.13
"foreman" 376 kbs		
channel state	SNR=9.0 dB	SNR=11.0 dB
JSC	29.53	30.03
MC-EZBC	27.56	31.74

TAB. 7.7 – Average PSNR of the reconstructed sequence "foreman" with motion estimation for different bitrates and different SNR. JSC denotes our coding scheme with the rotation matrix and MC-EZBC is the standard coder followed by UEP by RPC codes.

different approach, which consists of adding a rotation matrix prior to the transmission. The rotation of a BPSK constellation increases its diversity order and, contrary to the addition of redundancy by error-correcting codes, does not increase the bitrate. This approach proved to have excellent results. The performance of the entire coding scheme attains the noiseless performance even at very low SNR.

Finally, we compared the performance of our coding scheme with a state-of-the-art scalable video coder using UEP by RPC codes over a flat-Rayleigh channel. Our coding scheme proved to be more robust than the standard video coder. Moreover, it preserves its performances close to its noiseless quality even for low SNR Rayleigh channels.

Conclusions and Perspectives

In this thesis, a joint source-channel coding scheme based on a structured vector quantizer and linear index assignment has been proposed and its applications to video coding of $t + 2D$ decomposed video sequences have been investigated.

This approach for joint source-channel coding scheme has been developed for Gaussian sources and proved to minimize simultaneously the channel and the source distortion. The channel distortion is minimized due to the linear labelling and the source distortion due to linear transforms based on "maximum component diversity" lattice constellations. These linear transforms produce structured codebooks which mimic the Gaussian distribution of the source. The purpose of this thesis was to extend the previous work to the video domain, where the source distribution is not Gaussian, and to develop a robust coding scheme in the presence of noise.

Thus, an extended investigation of the spatial and spatiotemporal dependencies of the wavelet coefficients was first presented. A stochastic model which considers the conditional probability law of the coefficients in a given spatiotemporal subband to be Gaussian, with variance depending on a set of spatiotemporal neighbors, was proposed. We provided new estimators for this model with improved statistical performances. Then, we used this model to build an optimal mean square predictor for missing coefficients, which was further exploited in two applications of transmission over packet networks : a quality enhancement technique for resolution scalable video bitstreams and an error concealment method. Both of them are applied directly to the missing subband of the spatiotemporal decomposition in contrast with the usual error concealment methods which are rather applied to the reconstructed sequences. Simulation results showed significant quality improvement achieved by this technique for a scalable video bitstream, with different packetization strategies.

An interesting perspective at this point could be the study of sign prediction methods of the wavelet coefficients in $t + 2D$ decompositions of video sequences. In addition, we can also consider a vectorial extension of our proposed double stochastic model. However, as we present briefly in the Appendix B, the number of the parameters characterizing this model and that we have to estimate, increases rapidly with the number of considered neighbors. This could be a drawback compared with the scalar approach.

As the marginal probability distribution of the coefficients in the spatio-temporal subbands is not Gaussian, we performed significant modifications to the joint source-channel coding scheme developed for Gaussian sources in order to apply it into a practical context. We partitioned each subband into two classes of coefficients, allowing to achieve a better scaling of the lattice constellation in each class and we used our stochastic model to predict the partition threshold both at the coder and the decoder side. For most of

the subbands, this technique avoids the transmission of side-information and thus the bitrate increase. An iterative algorithm was developed to produce the two codebooks and respectively the two scaling parameters per subband.

In early attempts, on both Gaussian and Rayleigh channels, we checked the performance of our coding scheme by adding a partial protection on the bitstream by Rate Punctured Convolutional codes. The coding scheme had been proved to behave quite efficiently by itself and the partial addition of the RCP codes gave a very robust scheme in the presence of noise. However, in these attempts, the distribution of the bits among the subbands and the choice of the RPC had been made almost "intuitively" and based on the remarks on the energy of the subbands, made in the first part of this thesis. Thus, an obvious extension could be the development of a bit allocation algorithm, which allocates "optimally" a given bitrate among the possible source coding rates offered by the quantization matrices and different possible RPC codes, in order to minimize the end-to-end distortion of the system. However, this kind of treatment belongs to the almost classic methods in the joint source-channel coding domain.

As the linear index assignment had been proved to minimize the channel distortion, we developed an iterative bit allocation algorithm which results in an optimal codebook allocation, subject to a global bit rate and a nonnegativity constraint. The application of this algorithm to the choice of the quantization matrices made our coding scheme more flexible and efficient in the presence of noise.

We proposed two coding schemes for the Gaussian channel : the first one with uncoded index assignment and the second one with partial coded index assignment. In our implementation, this was performed using Reed-Muller codes. No error-correcting codes were added explicitly. In the coded case, we restricted the mapping space of the codevectors to codewords belonging to a Reed-Muller code. The performance of the latter coding scheme on a Gaussian channel was proved to be high even in very noisy environments (4,33 dB of SNR). In addition, we compared our coding schemes with an unstructured quantizer provided by GLA with careful index assignement by a minimax algorithm. The experiments proved that our coding schemes, even in a very noisy environment, remain quite close to their noiseless performance, which is not the case for an unstructured quantizer. We had remarked almost the same behavior compared to a standard coder protected by RPC codes.

In the transition from the coding scheme with uncoded index assignment to the coding scheme with partial coded index assignment, our goal was the decrease of the end-to-end distortion of the noisy system without changing the source coding rate used in the uncoded case. However, we could improve our noiseless performance using a larger source codebook, employing a non-exhaustive search among the possible codewords sent on the channel in order to keep the complexity low. A larger codebook would permit us to use a bigger set of Reed-Muller codes and the development of a suitable bit allocation algorithm could optimally distribute the given bitrate among the source and the channel coder. In addition, we could imagine the use of non-binary modulation schemes and more specific lattice constellations whose construction is based on the Reed-Muller codes, i.e. construction B.

Finally, in the last chapter of the thesis, we proposed a very robust coding scheme for flat-Rayleigh channels. Taking into consideration the important role of diversity, we

added a rotation matrix prior to the transmission through the channel. The performance of this coding scheme was excellent, even in a very noisy environment, compared to the coding scheme without the rotation matrix. The proposed scheme obtain between 5 and 10 dB increase in average PSNR over a Rayleigh channel of $SNR = 9.0$ dB compared to the scheme where no rotation matrix is used. We also compared our coding scheme with a standard video coder protected unequally by RPC codes. For low SNR our coding scheme outperforms the standard coder. However, we could further improve the diversity gain by applying a single rotation of higher dimension to a group of codewords. Moreover, in order to go to even lower SNRs over Rayleigh channels, we could use the Reed-Muller codes before the rotation matrix as in the Gaussian case.

Annexe A

Basic definitions in Algebraic Number theory

In this Appendix, based on the work in [64], we give some definitions in Algebraic number theory in order to explain the construction of full-diversity rotated cubic lattices using the theory of ideal lattices. In what follows :

Let $\mathbb{Z}$ be the set of rational integers, $\mathbb{Q}$ the set of rational numbers $\mathbb{Q} = \{\frac{a}{b}, a, b \in \mathbb{Z}\}$ and $\mathbb{C}$ the set of complex numbers.

A.1 General definitions

Field extension

Let L and K be two fields. If $L \subseteq K$, then we say that K is a field extension of L and it is denoted by K/L . Indeed, K has a natural structure as a vector space over L , where the addition is addition in K and the scalar multiplication of $\alpha \in L$ with $u \in K$ is just $\alpha u \in K$.

The dimension of K as a vector space over L is called degree of K over L and if it is finite we say that K is a finite extension of L .

A finite extension of $\mathbb{Q}$ is called a number field.

Let $K/\mathbb{Q}$ be a field extension and $\alpha \in K$. If there exists a non-zero irreducible monic polynomial $p \in \mathbb{Q}[X]$ such that $p(\alpha) = 0$, we say that α is algebraic over $\mathbb{Q}$ and such a polynomial is called the minimal polynomial of α over $\mathbb{Q}$.

If all the elements of K are algebraic then we say that K is an algebraic extension of $\mathbb{Q}$. If K is a number field, then $K = \mathbb{Q}(\theta)$ for some algebraic number $\theta \in K$ is a $\mathbb{Q}$ -vector space generated by the powers of θ . If K has degree n then $\{1, \theta, \theta^2, \dots, \theta^{n-1}\}$ is a basis of K and the degree of the minimal polynomial of θ is n .

Ring of integers of K : O_K

$\alpha \in K$ is an algebraic integer if it is a root of a monic polynomial with coefficients in $\mathbb{Z}$. The set of algebraic integers of K is a ring called the ring of integers of K denoted O_K . The ring of integers O_K of K forms a $\mathbb{Z}$ -module of rank n which, indeed, is a linear vector space of dimension n over $\mathbb{Z}$.

It possesses a $\mathbb{Z}$ -basis $\{\omega_1, \dots, \omega_n\}$ so that we can write any element of O_K as $x = \sum_{i=1}^n a_i \omega_i$ with $a_i \in \mathbb{Z}$, $i = 1, \dots, n$ and n is the degree of K .

A $\mathbb{Z}$ -basis of O_K is called an *integral basis of K or O_K* .

$\mathbb{Q}$ -homomorphism

Let $K/\mathbb{Q}$ and $L/\mathbb{Q}$ be two field extensions of $\mathbb{Q}$. We call $\sigma : K \rightarrow L$ a $\mathbb{Q}$ -homomorphism if σ is a ring homomorphism that satisfies $\sigma(\alpha) = \alpha$ for all $\alpha \in \mathbb{Q}$. If A is a ring, a ring homomorphism is a map $\psi : A \rightarrow A$ that satisfies for all $a, b \in A$:

- $\psi(a + b) = \psi(a) + \psi(b)$
- $\psi(ab) = \psi(a)\psi(b)$
- $\psi(1) = 1$

A $\mathbb{Q}$ -homomorphism $\sigma : K \rightarrow \mathbf{C}$ is called an *embedding* of K into $\mathbf{C}$.

Let $K = \mathbb{Q}(\theta)$ be a number field of degree n over $\mathbb{Q}$. There are exactly n embeddings of K into $\mathbf{C}$: $\sigma_i : K \rightarrow \mathbb{C}$, $i = 1, \dots, n$, defined by $\sigma_i(\theta) = \theta_i$, where θ_i are the distinct roots in $\mathbb{C}$ of the minimum polynomial of θ over $\mathbb{Q}$.

Trace and norm of $x \in K$: $Tr(x)$ and $N(x)$

Let $x \in K$. The elements $\sigma_1(x), \sigma_2(x), \dots, \sigma_n(x)$ are called the conjugates of x .

The trace of x over $\mathbb{Q}$ is defined as :

$$Tr(x) = \sum_{i=1}^n \sigma_i(x)$$

and the norm of x :

$$N(x) = \prod_{i=1}^n \sigma_i(x)$$

An interesting remark is that for any $x \in K$, we have $N(x)$ and $Tr(x) \in \mathbb{Q}$. If $x \in O_K$, we have $N(x)$ and $Tr(x) \in \mathbb{Z}$.

Galois extension of $\mathbb{Q}$: $Gal(K/\mathbb{Q})$

Let $K = \mathbb{Q}(\theta)$ be an extension of $\mathbb{Q}$ of degree n .

If the minimal polynomial of θ over $\mathbb{Q}$ has all its roots in K and splits into linear factors over K , then we say that K is a Galois extension of $\mathbb{Q}$. The set of field automorphisms is called the Galois group of K over $\mathbb{Q}$:

$$Gal(K/\mathbb{Q}) = \{\sigma : K \rightarrow K | \sigma(x) = x, \forall x \in \mathbb{Q}\}$$

Let $\{\sigma_1, \sigma_2, \dots, \sigma_n\}$ be the n embeddings of K into $\mathbb{C}$. Let r_1 be the number of embeddings with image in $\mathbb{R}$, the field of real numbers, and $2r_2$ the number of embeddings with image in $\mathbb{C}$ so that

$$r_1 + 2r_2 = n$$

The pair (r_1, r_2) is called *signature* of K . If $r_2 = 0$ we have a totally real algebraic number field. If $r_1 = 0$ we have a totally complex algebraic number field.

We call canonical embedding $\sigma : K \rightarrow \mathbb{R}^{r_1} \times \mathbb{C}^{r_2}$ (we can identify $\mathbb{R}^{r_1} \times \mathbb{C}^{r_2}$ with $\mathbb{R}^n$) the isomorphism defined by :

$$\sigma(x) = (\sigma_1(x), \dots, \sigma_{r_1}(x), \sigma_{r_1+1}(x), \dots, \sigma_{r_1+r_2}(x)) \in \mathbb{R}^n$$

The definition of canonical embedding establishes a one-to-one correspondence between the elements of an algebraic number field of degree n and the vectors of the n -dimensional Euclidean space.

Discriminant of K : d_K

Let $\{\omega_1, \dots, \omega_n\}$ be an integral basis of O_K . The discriminant of K is defined as :

$$d_K = \det[\sigma_j(\omega_i)]^2$$

Note that the discriminant is independent of the choice of a basis and the discriminant of a number field belongs to $\mathbb{Z}$.

The interest of the discriminant is that it defines a first invariant of a number field, a property of the field that does not depend on the way it is presented.

Algebraic lattices

Let $\{\omega_1, \dots, \omega_n\}$ be an integral basis of K . The n vectors $\mathbf{v}_i = \sigma(\omega_i) \in \mathbb{R}^n$, $i = 1, \dots, n$ are linearly independent, so they define a full rank algebraic lattice $\Lambda = \sigma(O_K)$.

The lattice $\Lambda = \sigma(O_K)$ can be expressed by means of its generator matrix M :

$$\Lambda = \{\mathbf{x} = \boldsymbol{\lambda}M \in \mathbb{R}^n | \boldsymbol{\lambda} \in \mathbb{Z}^n\}$$

The lattice generator matrix M is given by :

$$M = \begin{pmatrix} \sigma_1(\omega_1) & \cdots & \sigma_{r_1}(\omega_1) & \Re\sigma_{r_1+1}(\omega_1) & \Im\sigma_{r_1+1}(\omega_1) & \cdots & \Re\sigma_{r_1+r_2}(\omega_1) & \Im\sigma_{r_1+r_2}(\omega_1) \\ \sigma_1(\omega_2) & \cdots & \sigma_{r_1}(\omega_2) & \Re\sigma_{r_1+1}(\omega_2) & \Im\sigma_{r_1+1}(\omega_2) & \cdots & \Re\sigma_{r_1+r_2}(\omega_2) & \Im\sigma_{r_1+r_2}(\omega_2) \\ \vdots & \cdots & \vdots & \vdots & \vdots & \cdots & \vdots & \vdots \\ \sigma_1(\omega_n) & \cdots & \sigma_{r_1}(\omega_n) & \Re\sigma_{r_1+1}(\omega_n) & \Im\sigma_{r_1+1}(\omega_n) & \cdots & \Re\sigma_{r_1+r_2}(\omega_n) & \Im\sigma_{r_1+r_2}(\omega_n) \end{pmatrix}$$

$\Re, \Im$ denote respectively the real and imaginary part of a complex number.

The volume of the fundamental parallelotope of Λ is given by :

$$\text{vol}(\Lambda) = \det(\Lambda) = |\det(M)| = 2^{-r_2} \sqrt{|d_K|}$$

The algebraic lattices exhibit a diversity :

$$L = r_1 + r_2$$

The basic idea to build an algebraic lattice has been the existence of a $\mathbb{Z}$ -basis in K . Since it is known that O_K has such a basis (O_K is a free $\mathbb{Z}$ -module of rank n) we can embed it into $\mathbb{R}^n$ so as to obtain an algebraic lattice. However there exists other subsets of O_K that also have this structure of free $\mathbb{Z}$ -module of rank n . These are the ideals of O_K .

A.1.1 Properties of Ideals

Ideal of a commutative ring O_K : I

We recall that an ideal I of a commutative ring O_K is an additive subgroup of O_K which is stable under multiplication by O_K , $aI \subseteq I$ for all $a \in O_K$

Prime ideal I of O_K :

An ideal I of O_K is called prime (or maximal) if the quotient ring O_K/I is a field.

Principal ideal I of O_K :

An ideal is principal if it is of the form :

$$I = (x) = xO_K, \quad x \in I$$

That means that these ideals have the special property of being generated by only one element.

Norm of an ideal of O_K :

Let $I = (x)O_K$ be a principal ideal of O_K . Its norm is defined by :

$$N(I) = |N(x)|$$

Every ideal $I \neq \{0\}$ of O_K has a $\mathbb{Z}$ -basis $\{v_1, \dots, v_n\}$ where n is the degree of K .

Thus, an algebraic lattice Λ' built from an ideal $I \subset O_K$ gives a sublattice of the algebraic lattice built from O_K .

A.2 Ideal Lattices

Based on the notions presented above, we are now ready to see the notion of the ideal lattice and its minimum product distance. We always remain in the case of totally real number fields.

Let K be a totally real number field of degree n .

Ideal Lattice : (I, q_a)

An ideal lattice is a lattice (I, q_a) , where I is an O_K -ideal and

$$q_a = I \times I \rightarrow \mathbb{Z}, \quad q_a(x, y) = \text{Tr}(axy) \quad \forall x, y \in I$$

where $a \in K$ is totally positive.

If $\{\omega_1, \dots, \omega_n\}$ is a $\mathbb{Z}$ -basis of I then the generator matrix M of the lattice $\Lambda = \{\mathbf{x} = \lambda M | \lambda \in \mathbb{Z}^n\}$ is given by :

$$M = \begin{pmatrix} \sqrt{a_1}\sigma_1(\omega_1) & \sqrt{a_2}\sigma_2(\omega_1) & \cdots & \sqrt{a_n}\sigma_n(\omega_1) \\ \vdots & \vdots & \cdots & \vdots \\ \sqrt{a_1}\sigma_1(\omega_n) & \sqrt{a_2}\sigma_2(\omega_n) & \cdots & \sqrt{a_n}\sigma_n(\omega_n) \end{pmatrix}$$

where $a_j = \sigma_j(a), \forall j$. In addition, $G = MM^t = \{Tr(a\omega_i\omega_j)\}_{i,j=\{1,\dots,n\}}$

Since the Gram matrix G is a trace form, this shows that the generator matrix as given above indeed defines an ideal lattice.

In the case of ideal lattices the determinant of the lattice is related both to the discriminant d_K and to the norm of the ideal I . Thus :

$$\det(\Lambda) = N(\alpha)N(I)^2|d_K|.$$

Minimum product distance

Let Λ be an n -dimensional lattice with full diversity $L = n$.

The minimum product distance of $\mathbf{x} = (x_1, \dots, x_n)$ from the origin is defined as :

$$d_{p,min} = \min_{\mathbf{x} \in \Lambda} d_p(\mathbf{x})$$

where

$$d_p(\mathbf{x}) = \prod_{i=1}^n |x_i|.$$

If I is a principal ideal of O_K then

$$\min_{x \in I \setminus \{0\}} N(x) = N(I)$$

The minimum product distance of an ideal lattice of determinant $D = \det(\Lambda)$ defined over I is :

$$d_{p,min}(\Lambda) = \sqrt{\frac{D}{d_K}}$$

In order to compare different lattices the determinant D is normalized to be 1.

A.2.1 Construction of rotated $-\mathbb{Z}^n$ lattices with full diversity by Ideal lattices

If we focus on the construction of a particular rotated lattice $\mathbb{Z}^n$, in terms of ideal lattices, this means that we are looking for a number field K of degree n and an ideal $I \subseteq O_K$ such that $\Lambda = (I, q_a)$ is equivalent to $\mathbb{Z}^n$, $n \geq 2$. The $G = MM^t$ should represent the identity matrix which means that the trace form is isomorphic to the unit form and that the corresponding generator matrix M becomes an *orthogonal matrix* after basis reduction. We can also see that a lattice $\Lambda' = (I, q_a)$ over $I \subseteq O_K$ is a sublattice of $\Lambda = (O_K, q_a)$. The idea is that in a given lattice Λ there may be a sublattice which is $\mathbb{Z}^n$. A useful but not sufficient criterion to find the $\mathbb{Z}^n$ lattice is the lattice determinant. The determinant of $\mathbb{Z}^n$ is 1 and a scaled version of $\mathbb{Z}^n$ is of the form $(\sqrt{c}\mathbb{Z}^n)$ for some integer c , so that its determinant is $\det(G) = c^n$. A necessary but not sufficient condition is :

$$N(I)^2 N(a) d_K = c^n$$

and if we assume that $I = O_K$ this simplifies to :

$$N(a) d_K = c^n.$$

A.2.2 Cyclotomic fields

As all the constructions of the rotated $\mathbb{Z}^n$ lattice that we are going to present in the sequel are based on the maximal real subfield of a cyclotomic field, we give here some basic definitions of it and the ideal of its ring of integers.

A cyclotomic field is a number field $K = \mathbb{Q}(\zeta_p)$ generated by an p -th root of unity, $\zeta = \zeta_p = e^{2i\pi/p}$.

But the field $\mathbb{Q}(\zeta)$ is not real, so we examine the field $K = R \cap \mathbb{Q}(\zeta)$ which is totally real [33].

Thus, the $K = \mathbb{Q}(\zeta + \zeta^{-1})$ is the maximal real subfield of $\mathbb{Q}(\zeta)$ of degree $(p-1)/2$ when p is a prime.

The lattices are built via the ring of integers of K , $O_K = \mathbb{Z}[\zeta + \zeta^{-1}]$, whose integral basis is $\{\omega_j = \zeta^j + \zeta^{-j}\}_{j=1}^n$. There are n embeddings of K in $\mathbb{C}$ given by $\sigma_k(\omega_j) = \zeta^{kj} + \zeta^{-kj}$. Based on the necessary but not sufficient condition for building the $\mathbb{Z}^n$, we deduce that $a = (1 - \zeta)(1 - \zeta^{-1})$.

Annexe B

Vectorial extension of the double stochastic model

We extend in this Appendix the scalar model of conditional dependencies between coefficients to a vectorial model. We exploit the fact that groups of 2×2 coefficients in a subband have the same ancestors. Therefore, our vectors will consist of such groups of coefficients $C_{\mathbf{p}}$ where $\mathbf{p} = \{(2n, 2m), (2n+1, 2m), (2n, 2m+1), (2n+1, 2m+1)\}$.

Let their autocovariance matrix be denoted by $\mathbf{\Gamma}(\mathbf{p})$. Following the same idea as in the scalar case, we propose a conditional Gaussian model for the vectors :

$$\mathcal{L}(C_{\mathbf{p}}|\mathcal{V}(\mathbf{p})) = \mathcal{N}(0, \mathbf{\Gamma}(\mathbf{p})) \quad (\text{B.1})$$

where

$$\mathbf{\Gamma}(\mathbf{p}) = \sum_{\mathbf{k} \in \mathcal{V}(\mathbf{p})} Q_{\mathbf{p}}^2(\mathbf{k}) \mathbf{W}_{\mathbf{k}} + \boldsymbol{\alpha}$$

with $\boldsymbol{\alpha} = \text{diag}(\alpha_1, \alpha_2, \alpha_3, \alpha_4)$, $\mathcal{V}(\mathbf{p})$ is the neighborhood of the vector $C_{\mathbf{p}}$, $Q_{\mathbf{p}}(\mathbf{k})$ $\{Q_{\mathbf{p}}(\mathbf{k}), \mathbf{k} \in \mathcal{V}(\mathbf{p})\}$ are values of the spatio-temporal neighbors of $\mathbf{p}$ and $\mathbf{W}_{\mathbf{k}}$ are the matrices of weights that have to be estimated.

We suppose that $\alpha_1 = \alpha_2 = \alpha_3 = \alpha_4 = \alpha$, as there is no reason to grant a privilege to some of the coefficients in the vector. As the matrix $\mathbf{\Gamma}(\mathbf{p})$ represents an autocovariance matrix, for $\mathbf{W}_{\mathbf{k}}$ positive definite and $\alpha \geq 0$, the constraint of positive definiteness on $\mathbf{\Gamma}(\mathbf{p})$ is satisfied.

For the sake of simplicity, we consider a parametric form of the matrices $\mathbf{W}_{\mathbf{k}}$, corresponding to 2D AR models. They can be expressed therefore as :

$$\mathbf{W}_{\mathbf{k}} = \sigma_{\mathbf{k}}^2 \begin{pmatrix} 1 & \rho_{\mathbf{k}}^1 & \rho_{\mathbf{k}}^2 & \rho_{\mathbf{k}}^1 \rho_{\mathbf{k}}^2 \\ \rho_{\mathbf{k}}^1 & 1 & \rho_{\mathbf{k}}^1 \rho_{\mathbf{k}}^2 & \rho_{\mathbf{k}}^1 \\ \rho_{\mathbf{k}}^2 & \rho_{\mathbf{k}}^1 \rho_{\mathbf{k}}^2 & 1 & \rho_{\mathbf{k}}^2 \\ \rho_{\mathbf{k}}^1 \rho_{\mathbf{k}}^2 & \rho_{\mathbf{k}}^1 & \rho_{\mathbf{k}}^2 & 1 \end{pmatrix}$$

where $\rho_{\mathbf{k}}^1$, $\rho_{\mathbf{k}}^2$ are respectively the horizontal and vertical correlation coefficients for the $\mathbf{k}$ th neighbor.

However, if we maintain the same number of neighbors as in the Chapter 2 (12 neighbors), then the number of the parameters that we should estimate becomes considerably high. Indeed, for each neighbor we should estimate the $\sigma_{\mathbf{k}}^2$, $\rho_{\mathbf{k}}^1$ and $\rho_{\mathbf{k}}^2$.

Thus, we restrict the neighborhood to include only the spatio-temporal parent and the autocovariance matrix can, now, be expressed as :

$$\mathbf{\Gamma}(\mathbf{p}) = Q_{\mathbf{p}}(0)^2 \mathbf{W}_0 + \alpha \mathbf{I}, \quad \alpha \geq 0$$

where $Q_{\mathbf{p}}(0)$ is the value of the spatio-temporal parent of the current vector of coefficients with spatial index $\mathbf{p}$.

Thus, the Least Square criterion which has to be minimized becomes :

$$\sum_{\mathbf{p}} \|\mathbf{C}_{\mathbf{p}} \mathbf{C}_{\mathbf{p}}^T - \mathbf{\Gamma}(\mathbf{p})\|_F^2 = \sum_{\mathbf{p}} \text{Tr}[(\mathbf{C}_{\mathbf{p}} \mathbf{C}_{\mathbf{p}}^T - \mathbf{\Gamma}(\mathbf{p}))^2]$$

Annexe C

Rotation Matrices

We present, here, the numerical values of the rotation matrix $\mathbf{U}_n$, for some typical values of n used in simulations : $n = 2, 4, 8, 16$.

$$\mathbf{U}_2 = \begin{bmatrix} 0.38 & -0.92 \\ -0.92 & -0.38 \end{bmatrix}$$

$$\mathbf{U}_4 = \begin{bmatrix} 0.59 & -0.14 & -0.69 & -0.39 \\ 0.14 & -0.39 & 0.59 & -0.70 \\ -0.40 & 0.69 & -0.14 & -0.59 \\ -0.70 & -0.59 & -0.40 & -0.14 \end{bmatrix}$$

$$\mathbf{U}_8 = \begin{bmatrix} 0.48 & 0.32 & 0.05 & -0.24 & -0.44 & -0.50 & -0.39 & -0.15 \\ 0.39 & -0.24 & -0.48 & 0.05 & 0.50 & 0.15 & -0.44 & -0.32 \\ 0.24 & -0.50 & 0.32 & 0.15 & -0.48 & 0.39 & 0.05 & -0.44 \\ 0.05 & -0.15 & 0.24 & -0.32 & 0.39 & -0.44 & 0.48 & -0.50 \\ -0.15 & 0.39 & -0.50 & 0.44 & -0.23 & -0.05 & 0.32 & -0.48 \\ -0.32 & 0.44 & 0.15 & -0.50 & 0.05 & 0.48 & -0.24 & -0.39 \\ -0.44 & -0.05 & 0.39 & 0.48 & 0.16 & -0.32 & -0.50 & -0.24 \\ -0.50 & -0.48 & -0.44 & -0.39 & -0.32 & -0.24 & -0.15 & -0.05 \end{bmatrix}$$

$$\mathbf{U}_{16} = \begin{bmatrix}
 0.35 & 0.32 & 0.26 & 0.18 & 0.09 & -0.02 & -0.12 & -0.21 & -0.28 & -0.33 & -0.35 & -0.34 & -0.30 & -0.24 & -0.15 & -0.05 \\
 0.33 & 0.18 & -0.05 & -0.26 & -0.35 & -0.28 & -0.09 & 0.15 & 0.32 & 0.34 & 0.21 & -0.02 & -0.24 & -0.35 & -0.30 & -0.12 \\
 0.30 & -0.01 & -0.32 & -0.28 & 0.05 & 0.33 & 0.26 & -0.07 & -0.34 & -0.24 & 0.12 & 0.35 & 0.21 & -0.15 & -0.35 & -0.18 \\
 0.26 & -0.21 & -0.30 & 0.15 & 0.33 & -0.09 & -0.35 & 0.01 & 0.35 & 0.05 & -0.34 & -0.12 & 0.32 & 0.18 & -0.28 & -0.24 \\
 0.21 & -0.33 & -0.02 & 0.34 & -0.18 & -0.24 & 0.32 & 0.05 & -0.35 & 0.15 & 0.26 & -0.30 & -0.09 & 0.35 & -0.12 & -0.28 \\
 0.15 & -0.34 & 0.28 & -0.02 & -0.26 & 0.35 & -0.18 & -0.12 & 0.33 & -0.30 & 0.05 & 0.24 & -0.35 & 0.21 & 0.09 & -0.32 \\
 0.09 & -0.24 & 0.33 & -0.35 & 0.28 & -0.15 & -0.02 & 0.18 & -0.30 & 0.35 & -0.32 & 0.21 & -0.05 & -0.12 & 0.26 & -0.34 \\
 0.01 & -0.05 & 0.08 & -0.12 & 0.15 & -0.18 & 0.21 & -0.24 & 0.26 & -0.28 & 0.30 & -0.32 & 0.33 & -0.34 & 0.35 & -0.35 \\
 -0.05 & 0.15 & -0.24 & 0.30 & -0.34 & 0.35 & -0.33 & 0.28 & -0.21 & 0.12 & -0.02 & -0.09 & 0.18 & -0.26 & 0.32 & -0.35 \\
 -0.11 & 0.30 & -0.35 & 0.24 & -0.01 & -0.21 & 0.34 & -0.32 & 0.15 & 0.09 & -0.28 & 0.35 & -0.26 & 0.05 & 0.18 & -0.33 \\
 -0.18 & 0.35 & -0.15 & -0.21 & 0.35 & -0.12 & -0.24 & 0.34 & -0.09 & -0.26 & 0.33 & -0.05 & -0.28 & 0.32 & -0.01 & -0.30 \\
 -0.24 & 0.28 & 0.18 & -0.32 & -0.12 & 0.34 & 0.05 & -0.35 & 0.02 & 0.35 & -0.09 & -0.33 & 0.15 & 0.30 & -0.21 & -0.26 \\
 -0.28 & 0.12 & 0.35 & 0.09 & -0.30 & -0.26 & 0.15 & 0.35 & 0.05 & -0.32 & -0.24 & 0.18 & 0.34 & 0.02 & -0.33 & -0.21 \\
 -0.32 & -0.09 & 0.21 & 0.35 & 0.26 & -0.05 & -0.30 & -0.33 & -0.12 & 0.18 & 0.35 & 0.26 & -0.02 & -0.28 & -0.34 & -0.15 \\
 -0.34 & -0.26 & -0.12 & 0.05 & 0.21 & 0.32 & 0.35 & 0.30 & 0.18 & 0.02 & -0.15 & -0.28 & -0.35 & -0.33 & -0.24 & -0.09 \\
 -0.35 & -0.35 & -0.34 & -0.33 & -0.32 & -0.30 & -0.28 & -0.26 & -0.24 & -0.21 & -0.18 & -0.15 & -0.12 & -0.08 & -0.05 & -0.02
 \end{bmatrix}$$

Bibliographie

- [1] M. Antonini, M. Barlaud and P. Mathieu, "Image coding using lattice vector quantization of wavelet coefficients", IEEE Int. Conf. on Acoustics, Speech and Signal Processing, ICASSP'91, vol. 4, pp. 2273-2276, Apr. 1991.
- [2] L. R. Bahl, J. Cocke, F. Jelinek and J. Raviv, "Optimal decoding of linear codes for minimizing symbol error rate", IEEE Trans. on Information Theory, vol. 20, pp. 284-287, 1974.
- [3] M. Barlaud, P. Sole, T. Gaidon, M. Antonini and P. Mathieu, "Pyramidal lattice Vector Quantization for multiscale image coding", IEEE Trans. on Image Processing, vol. 3, no. 4, pp. 367-381, Jul. 1994.
- [4] E. Bayer-Fluckiger, F. Oggier and E. Viterbo, "New Algebraic Constructions of Rotated Z^n -Lattice Constellations for the Rayleigh Fading Channel", IEEE Trans. on Information Theory, vol. 50, no. 4, pp. 702-714, Apr. 2004.
- [5] J. C. Belfiore, X. Giraud and J. Rodriguez, "Optimal linear labelling for the minimization of both source and channel distortion", in Proc. IEEE Int. Symp. Information Theory, pp. 204, Sorrento, Italy, Jun. 2000.
- [6] B. Belzer, J. D. Villasenor and B. Girod, "Joint source channel coding of images with trellis coded quantization and convolutional codes", IEEE Proc. Int. Conf. Image Processing, Washington, DC, Oct. 1998.
- [7] Y. Bian, A. Popplewell and J. J. O'Reilly, "New very high rate punctured convolutional codes", IEEE Electronics Letters, vol 30, no 14, pp. 1119-1120, July 1994.
- [8] J. Boutros and E. Viterbo, "Signal Space Diversity : A Power-and-Bandwidth-Efficient Diversity Technique for the Rayleigh Fading Channel", IEEE Trans. on Information Theory, vol. 44, no. 4, pp. 1453-1467, July 1998.
- [9] J. Boutros, E. Viterbo, C. Rastello and J. C. Belfiore, "Good lattice constellations for both Rayleigh fading and Gaussian channels", IEEE Trans. on Information Theory, vol. 42, pp. 502-518, Mar. 1996.
- [10] M. Bystrom and T. Stockhammer, "Dependent Source and Channel Rate Allocation for Video Transmission", IEEE Trans. on Wireless Commun., vol. 3, no. 1, Jan. 2004.
- [11] G. Cheung and A. Zakhor, "Joint source/channel coding of scalable video over noisy channels", IEEE Int. Conference Proc. on Image Processing, vol. 3, pp. 767-770, Lausanne, Switzerland, Sept. 1996.
- [12] G. Cheung and A. Zakhor, "Bit allocation for Joint Source/Channel Coding of Scalable Video", IEEE Trans. on Image Processing, vol. 9, no. 3 pp. 340-356, Mar. 2000.

- [13] Jui-Chiu Chiang, "Codage source-canal conjoint pour la transmission robuste de video sur des canaux radio-mobiles", PhD dissertation, Université de Paris-Sud, Dec. 2004.
- [14] S. Cho and W. A. Pearlman, "A Full-Featured, Error-Resilient, Scalable Wavelet Video Codec Based on the Set Partitioning in Hierarchical Trees (SPIHT) Algorithm", IEEE Trans. on Circuits and Systems for Video Techn., vol. 12, no. 3, pp. 157-171, Mar. 2002.
- [15] S. -J Choi and J. W. Woods, "Motion-Compensated 3D Subband Coding of Video", IEEE Trans. on Image Processing, vol. 8, no. 2, pp. 155-167, Feb. 1999.
- [16] M. Z. Coban and R. M. Mersereau, "Adaptive subband video coding using bivariate generalized Gaussian distribution model", IEEE Int. Conference. Proc. on Acoustics, Speech and Signal Processing., ICASSP'96, vol. 4, pp. 1990-1993, 1996.
- [17] P. C. Cosman, R. M. Gray and M. Vetterli, "Vector Quantization of Image Subbands : A Survey", IEEE Trans. on Image Processing, vol. 5, no. 2, pp. 202-225, Febr. 1996.
- [18] A. Deever and S. Hemami, "Efficient Sign Coding and Estimation of Zero-Quantized Coefficients in Embedded Wavelet Image Codecs", IEEE Trans. on Image Processing, vol. 12, no. 4, pp. 420-430, Apr. 2003.
- [19] M. Effros, "Robustness to channel variation in source coding for transmission across noisy channels", Proc. IEEE ICASSP, pp. 2961-2964, Apr. 1997.
- [20] N. Farvardin, "A Study of Vector Quantization for Noisy Channels", IEEE Trans. on Information Theory, vol. 36, no. 4, Jul. 1990.
- [21] N. Farvardin and Vaishampayan, "On the performance and complexity of channel-optimized vector quantizers", IEEE Trans. Information Theory, vol. 37, pp. 155-160, Jan. 1991.
- [22] N. Farvardin and V. Vaishampayan, "Optimal quantizer design for Noisy channels : an approach to combined Source- Channel Coding", IEEE Trans. on Information Theory, vol. IT33, pp. 827-838, Nov. 1987.
- [23] T. R. Fischer and M. W. Marcellin, "Joint trellis coded quantization/modulation", IEEE Trans. Commun, vol. 39, no. 2, pp. 172-176, Feb. 1991.
- [24] G. D. Forney, Jr., "Codet codes-Part II : Binary lattices and related codes", IEEE Trans. on Information Theory, vol. 34, pp. 1152-1187, Sept. 1988.
- [25] M. P. C. Fossorier and S. Lin, "Soft-decision decoding of linear block codes based on ordered statistics", IEEE Trans. on Information Theory, vol. 41, pp. 1379-1396, Sept. 1995.
- [26] M. P. C. Fossorier and S. Lin, "Computationally Efficient Soft-Decision Decoding of Linear Block Codes Based on Ordered Statistics", IEEE Trans. on Information Theory, vol. 42, no. 3, pp. 738-750, May 1996.
- [27] A. Gabay, O. Rioul and P. Duhamel, "Joint source-channel coding using structured oversampled filters banks applied to image transmission", IEEE Int. Conf. on Acoustics, Speech and Signal Processing, ICASSP'2001, vol. 4, pp. 2581-2584, May 2001.

- [28] A. Gabay, P. Duhamel and O. Rioul, "Real BCH codes as joint source channel codes for satellite images coding", IEEE Int. Conf. on Global Telecommunications, GLOBECOM'00, vol.2, pp. 820-824, Dec. 2000.
- [29] S. Gadkari and K. Rose, "Vector Quantization with Transmission Energy Allocation for Time-Varying Channels", IEEE Trans. on Commun., vol. 47, no. 1, pp. 149-157, Jan. 1999.
- [30] Z. Gao, F. Chen, B. Belzer and J. Villasenor, "A comparison of the Z , E_8 and Leech lattices for image subband quantization", IEEE Data Compression Conf. pp. 312-321, Mar. 1995.
- [31] Z. Gao, B. Belzer and J. Villasenor, "A comparison of the Z , E_8 and Leech lattices for quantization of low-shape-parameter generalized Gaussian sources", IEEE Signal Proc. Letters, vol. 2, no. 10, pp. 197-199, Oct. 1995.
- [32] M. Ghanbari and V. Seferidis, "Cell loss concealment in ATM video codecs", IEEE Trans. Circuits Systems Video Technol., vol. 3, pp. 238-247, Jun. 1993.
- [33] X. Giraud and J. C. Belfiore, "Constellations Matched to the Rayleigh Fading Channel", IEEE Trans. on Information Theory, vol. 42, no. 1, pp. 106-115, Jan. 1996.
- [34] X. Giraud, E. Boutillon and J. C. Belfiore, "Algebraic Tools to Build Modulation Schemes for Fading Channels", IEEE Trans. on Information Theory, vol. IT 43, pp. 938-952, May 1997.
- [35] A. J. Goldsmith and M. Effros, "Joint Design of Fixed-Rate Source Codes and Multi-resolution Channel Codes", IEEE Trans. on Commun., vol. 46, no. 10, pp. 1301-1312, Oct. 1998.
- [36] C. Gouriéroux and A. Monfort, "Statistics and Econometric Models", Cambridge University Press, 1995.
- [37] J. Rodriguez Guisantes, "Codage Conjoint Source-canal", PhD dissertation, E.N.S.T. Paris, Dec. 1997.
- [38] D. Haccoun and G. Begin, "High rate punctured convolutional codes for Viterbi and sequential decoding", IEEE Trans. on Commun., vol. 37, pp. 1113-1125, 1989.
- [39] R. Hagen and P. Hedelin, "Robust Vector Quantization by a Linear Mapping of a Block Code", IEEE Trans. on Information Theory, vol. 45, no. 1, pp. 200-218, Jan. 1999.
- [40] J. Hagenauer, "Rate-Compatible Punctured Convolutional Codes (RCPC codes) and their Applications", IEEE Trans. Commun. vol. 36 pp. 389-399, Apr. 1988.
- [41] S.S Hemami and T.H.Y. Meng, "Transform coded image reconstruction exploiting interblock correlations", IEEE Trans. on Image Processing, vol. 4, pp. 1023-1027, Jul. 1995.
- [42] K. -P. Ho and J. M. Kahn, "Joint Design of a Channel-Optimized Quantizer and Multicarrier Modulation", IEEE Trans. on Commun., vol. 46, no.10, Oct. 1998.
- [43] S. T. Hsiang and J. W. Woods, "Invertible 3D analysis/synthesis system for video coding with half-pixel accurate motion compensation", Proc. VCIP'99, SPIE Vol. 3653, pp. 537-546, 1999.

- [44] H. Jafarkhani, P. Ligdas and N. Farvardin, "Adaptive rate allocation in a joint source/channel coding framework for wireless channels", in Proc. IEEE VTC'96, pp. 492-496, Apr. 1996.
- [45] B. -J Kim, Z. Xiong and W. A. Pearlman, "Low Bit-Rate Scalable Video Coding with 3-D Set Partitioning in Hierarchical Trees (3-D SPIHT)", IEEE Trans. on Circuits and Systems for Video Technology, vol. 10, no. 8, pp. 1374-1387, Dec. 2000.
- [46] P. Knagenhjelm and E. Agrell, "The Hadamard Transform- A tool for index Assignment", IEEE Trans. on Information Theory, vol. IT 42, pp. 1139-1151, Jul. 1996.
- [47] W. M. Lam, A. R. Reibman, and B. Liu, "Recovery of lost or erroneously received motion vectors", Proc. Int. Conf. Acoust. Speech, Signal Processing, vol. V, pp. 417-420, 1993.
- [48] C. Lamy and J. Boutros, "On Random Rotations Diversity and Minimum MSE Decoding of Lattices", IEEE Trans. on Information Theory, vol.46, no. 4, pp. 1584-1589, Jul.2000.
- [49] C. Lamy, "Communications à grande efficacité spectrale sur le canal à évanouissements", PhD dissertation, E.N.S.T. Paris, Avr. 2000.
- [50] Y. Linde, A.Buzo and R. M. Gray, "An algorithm for vector quantizer design", IEEE trans. Commun., vol. COM-28, pp. 84-95, Jan. 1980.
- [51] J. Liu and P. Moulin, "Information-Theoretic Analysis of Interscale and Intrascale Dependencies Between Image Wavelet Coefficients", IEEE Trans. on Image Processing, vol 10, no 11, pp 1647-1658, Nov. 2001.
- [52] J. Liu and P. Moulin, "Image denoising based on scale-space mixture modelling for wavelet coefficients", in Proc. ICIP'99, Kobe, Japan, pp. I. 386-390, Oct.1999 .
- [53] S. LoPresto, K. Ramchandran and M. T. Orchard, "Image coding based on mixture modelling of wavelet coefficients and a fast estimation-quantization framework", in Data Compression Conf. '97, Snowbird, UT, pp. 221-230, 1997.
- [54] D. G. Luenberger, "Optimization by vector space methods", John Wiley and sons, Inc. 1969.
- [55] J. Luo, C. W. Chen, K. J. Parker and T. S. Huang, "Adaptive quantization with spatial constraints insubband video compression using wavelets", IEEE Int. Conference Proc. on Image Processing, ICIP'95, vol. 1, pp. 594-597, 1995.
- [56] S. Marinkovic and C. Guillemot, "Joint source ans channel coding/decoding by concatenating an oversampled filter bank code and redundant entropy code", IEEE Global Telecommunications Conf., GLOBECOM'04, vol. 4, Nov. 2004.
- [57] R. J. McEliece, "On the BCJR Trellis for linear Block Codes", IEEE TRans. on Information Theory, vol. 42, no. 4, pp. 1072-1091, Jul. 1996.
- [58] A. Mehes and K. Zeger, "Binary Lattice Vector Quantization with Linear Block Codes and Affine Index Assignments" IEEE Trans. on Information Theory, vol. 44, no. 1, pp. 79-94, Jan. 1998.
- [59] A. Mehes and K. Zeger, "Performance of Quantizers on Noisy Channels Using Structured Families of Codes", IEEE Trans. on Information Theory, vol.46, no. 7, Nov. 2000.

- [60] M. K. Mihcak, K. Ramchandran and P. Moulin, "Low-Complexity Image Denoising Based on Statistical Modeling of Wavelet Coefficients", *IEEE Signal Processing Lett*, vol 6, no 12, Dec. 1999.
- [61] J. W. Modestino, D. G. Daut, "Combined source-channel coding of images", *IEEE Trans. Commun.*, vol. COM-27, pp. 1644-1659, Nov.1979.
- [62] J. W. Modestino, D. G. Daut and A. L. Vickers, "Combined source-channel coding of images using the block cosine transform", *IEEE Trans. Commun.*, vol. COM-27, pp. 1261-1274, Sept.1981.
- [63] D. Mukherjee and S. K. Mitra, "Vector SPIHT for Embedded Wavelet Video and Image Coding", *IEEE Trans. on Circuits and Systems for Video Technology*, vol. 13, no. 3, pp. 231-246, Mar. 2003.
- [64] F. Oggier and E. Viterbo, "Algebraic number theory and its applications to code design for Rayleigh fading channels", Technical Report, May 2004.
- [65] C. Parisot, "Allocations basées modèles et transformée en ondelettes au fil de l'eau pour le codage d'images et de vidéos", PhD dissertation, Université de Nice-Sophia Antipolis, Jan. 2003.
- [66] G. Pau, C. Tillier and B. Pesquet-Popescu, "Motion compensation and scalability in lifting-based video coding", *Signal Processing : Image Communication*, special issue on Wavelet Video Coding Elsevier/EURASIP, vol. 19, pp. 577-600, Aug. 2004.
- [67] M. Pereira, M. Antonini, M. Barlaud, "Channel adapted scan-based multiple description video coding", *IEEE Int. Conf. on Multimedia and Expo, ICME '02*, vol. 2, pp. 609-612, Aug. 2002.
- [68] B. Pesquet-Popescu and V. Bottreau, "Three-dimensional lifting schemes for motion compensated video compression", *Proc. ICASSP 2001*, Salt Lake City, May 2001.
- [69] M. S. Postol, "Some new lattice quantization algorithms for video compression coding", *IEEE Trans. on Circuits and Systems for Video Technology*, vol. 12, no. 1, pp. 53-60, Jan. 2002.
- [70] L. C. Potter, Da-Ming Chiang, "Minimax Nonredundant Channel Coding", *IEEE Trans. on Commun.* vol. 43, no.4, pp. 804-811, Apr. 1995.
- [71] P. Rault and F. Kossentini, "Robust subband image coding for wireless transmission" *IEEE Int. Conf. on Image Proc., ICIP'99*, vol. 33, pp. 374-378, Oct. 1999.
- [72] P. Raffy, M. Antonini and M. Barlaud, "Distortion-rate Models for entropy-coded lattice vector quantization", *IEEE Trans. on Image Processing*, vol. 9, no. 12, pp. 2006-2017, Dec. 2000.
- [73] J.K. Romberg, H. Choi, and R.G. Baraniuk, "Bayesian Tree-Structured Image Modeling Using Wavelet-Domain Hidden Markov Models", *IEEE Trans. on Image Processing*, vol. 10, no. 7, pp. 1056-1068, Jul. 2001.
- [74] A. Said and W. A. Pearlman, "A new, fast and efficient image codec based on set partitioning in hierarchical trees", *IEEE Trans. on Circuits and Systems for Video Technology*, vol. 6, no. 3, pp. 243-250, Jun. 1996.
- [75] P. Salama, N. Shroff, E. J. Coyle, and E. J. Delp, "Error concealment in encoded video streams", *Int. Conference Proc. on Image Processing*, vol. 1, pp. 23-26, 1995.

- [76] A. Segall, "Bit Allocation and Encoding for Vector Sources", IEEE Trans. on Information Theory, vol. IT 22, no. 2, pp.162-169, Mar. 1976.
- [77] L. Sendur, and Ivan W. Selesnick, "Bivariate Shrinkage Functions for Wavelet-Based Denoising Exploiting Interscale Dependency", IEEE Trans. on Signal Processing, Vol. 50, No. 11, pp. 2744-2756, Nov.2002.
- [78] J. Shapiro, "Embedded Image Coding Using Zerotrees of Wavelet Coefficients", IEEE Trans. on Signal Processing, vol. 41, no. 12, pp. 3445-3462, Dec. 1993.
- [79] Yun Q. Shi and Huifang Sun, "Image and Video Compression for Multimedia Engineering : fundamentals, algorithms and standards", CRC Press.
- [80] D.G Sampson, E. A. B. da Silva and M. Chanbari, "Wavelet lattice quantization for low bit rate video coding", Int. Conf. on Communications, ICC'95, vol. 3, pp. 1423-1427, Jun. 1995.
- [81] S. Shirani, F. Kossentini, R. Ward, "Error concealment methods, A comparative study", IEEE Canadian Conf. on Electrical and Computer Engineering, pp. 835-840, May 1999.
- [82] S. Shirani, F. Kossentini, R. Ward, "A Concealment Method for Video Communications in an Error-Prone Environment", IEEE Journal on Selected Areas in Commun., vol. 18, no. 6, pp. 1122-1128, Jun. 2000.
- [83] Y. Shoham and A. Gersho, "Efficient bit allocation for an arbitrary set of quantizers", IEEE Trans. Acoust. Speech Signal. Proc., vol. 36, no. 9, pp : 1445-1453, Sept. 1988.
- [84] E. P. Simoncelli, "Bayesian denoising of visual images in the wavelet domain", Lecture Notes in Statistics, Chap. 18, Springer-Verlag, New York., vol. 141, pp. 291-308, Mar. 1999.
- [85] E. P. Simoncelli, "Modeling the joint statistics of images in the wavelet domain", Proc. SPIE 44th Annual Meeting, Denver Colorado, vol. 3813, pp. 188-195, Jul. 1999.
- [86] E.P. Simoncelli and R. Buccigrossi, "Image compression via joint statistical characterization in the wavelet domain", IEEE Trans. on Image Processing, vol. 8, pp. 1688-1701, Dec. 1999.
- [87] M. Skoglund, "On Channel-Constrained Vector Quantization and Index Assignment for Discrete Memoryless Channels", IEEE Trans. on Information Theory, vol. 45, no. 7, pp. 2615-2622, Nov. 1999.
- [88] M. Srinivasan and R. Chellappa, "Adaptive source-channel subband video coding for wireless channels", IEEE J. Select. Areas Commun., vol. 16. Dec. 1998.
- [89] H. Sun and W. Kwok, "Concealment of damaged block transform coded images using projection onto convex sets", IEEE Trans. on Image Processing, vol. 4, pp. 470-477, Apr. 1995.
- [90] R. Talluri, "Error resilient video coding in the MPEG-4 standard", IEEE Commun. Magazine, vol. 26, pp. 112-119, Jun. 1999.
- [91] D. S. Taubman and M. W. Marcellin, "JPEG2000 : Image Compression Fundamentals, Standards and Practice", Kluwer Int. Series in Engineering and Computer Science.

-
- [92] E. Viterbo and J. Boutros, "A universal lattice code decoder for fading channels", IEEE Trans. on Information Theory, vol. 45, no. 5, pp. 1639-1642, Jul. 1999.
 - [93] Y. Wang, Q. F. Zhu, and L. Shaw, "Maximally smooth image recovery in transform coding", IEEE Trans. Commun., vol. 41, pp. 1544-1551, Oct. 1993.
 - [94] K. Zeger and A. Gersho, "Pseudo-Gray Coding", IEEE Trans. on Commun. vol. 38, no. 12, pp. 2147-2158, Dec. 1990.
 - [95] S. B. Zahir Azami, P. Duhamel, O. Rioul, "Combined Source-Channel Coding : Panorama of Methods", CNES Workshop on Data Compression, Toulouse, Nov. 1996.
 - [96] Z. Xiong, B. -J. Kim and W. A. Pearlman, "Progressive video coding for noisy channels", in IEEE Int. Conf. Image Proc., ICIP'98, vol. 1, pp. 334-337, Oct. 1998.
 - [97] Z. M. Ysof and I. Fischer, "An entropy-coded lattice vector quantization for transform and subband image coding", IEEE Trans. on Image Processing, vol. 5, pp. 289-298, Febr. 1996.
 - [98] Y. Zhang and Kai-Kuang Ma, "Error Concealment for Video Transmission With Dual Multiscale Markov Random Field Modeling", IEEE Trans. on Image Processing, vol. 12, no. 2, pp. 236-242, Feb. 2003.
 - [99] Z. Zhang, G. Liu and Y. Yang, "Progressive source-channel coding video for unknown noisy channels", IEEE Int. Conf. on Acoustics, Speech and Signal Proc., ICASSP'02, vol.3, pp.2493-2496, May 2002.

List of publications

G. Feideropoulou, J. Fowler, B. Pesquet-Popescu and J. C. Belfiore, "Joint Source-Channel coding of scalable video with partially coded index assignment using Reed-Muller codes", IEEE Int. Conf. on Image Processing, ICIP'05, Geneoa-Italy, Sept. 2005.

G. Feideropoulou, B. Pesquet-Popescu and J. C. Belfiore, "Bit Allocation Algorithm for Joint Source-Channel Coding of t+2D Video Sequences", IEEE ICASSP'05, Philadelphia, Mars 2005.

G. Feideropoulou, B. Pesquet-Popescu and J. C. Belfiore, "Joint Source- Channel coding of scalable video", IEEE GLOBECOM'04, Dallas, Nov. 2004.

G. Feideropoulou, B. Pesquet-Popescu and J. C. Belfiore, "Joint Source- Channel coding of scalable video on a Rayleigh fading channel", IEEE Int. Workshop on Multimedia and Signal Processing (MMSP), Siena, Italy, Sept. 2004.

G. Feideropoulou, B. Pesquet-Popescu, "Model-Based Quality Enhancement of Scalable Video", SPIE VCIP, San Jose, CA, Jan. 2004.

G. Feideropoulou, B. Pesquet-Popescu, "A Model-Based Error Concealment Method for Scalable Video", IEEE Int. Symposium on Signal Processing and Information Technology (ISSPIT), Darmstadt, Germany, Dec. 2003.

G. Feideropoulou and B. Pesquet-Popescu and J-C. Belfiore et G. Rodriguez, "Non-linear modelling of wavelet coefficients for a video sequence", Proc. of IEEE Workshop on Nonlinear Signal Processing, Grado, Italy, Jun. 2003.

G. Feideropoulou, B. Pesquet-Popescu and J. C. Belfiore, " Stochastic Modelling of the Spatio-Temporal Wavelet Coefficients. Application to Quality Enhancement and Error Concealment ", EURASIP Journal of Signal Processing and Applications (JASP), Jan. 2004.