



Quantification vectorielle algébrique et arborescente.

Vincent Ricordel, Claude Labit

► To cite this version:

Vincent Ricordel, Claude Labit. Quantification vectorielle algébrique et arborescente.. Compression et représentation des signaux audiovisuels (CORESA), Feb 1996, Grenoble, France. hal-00451780

HAL Id: hal-00451780

<https://hal.science/hal-00451780>

Submitted on 31 Jan 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.

Quantification vectorielle algébrique et arborescente

Vincent Ricordel et Claude Labit

IRISA / INRIA Rennes
Campus de Beaulieu
35042 Rennes Cedex, France
fax: (+33) 99.84.71.71,
e-mail: ricordel@irisa.fr, labit@irisa.fr

1 Introduction

Notre étude vise à concevoir un nouveau schéma de Quantification Vectorielle (QV) devant prendre place au sein d'une chaîne de codage hybride pour la compression d'images. Une telle source vectorielle (c.a.d des vecteurs d'erreurs de prédiction transformées ou non) a sa distribution statistique généralement modélisée par une fonction de la famille des Gaussiennes généralisées. Cependant ce signal à coder n'est jamais stationnaire. C'est pourquoi nous retenons une technique d'apprentissage pour la construction du dictionnaire, ainsi une mise à jour des représentants peut être opérée au cours du temps à partir de Séquences d'Apprentissage (SA) représentatives de la statistique courante de la source [7][10]. Mais la construction du dictionnaire, et donc les opérations d'encodage-décodage, doivent être très rapides. Le coût calculatoire des techniques d'apprentissage classiques [8][6] (algorithme de Lloyd généralisé, LBG) rend ces méthodes inappropriées. La QV algébrique [5][1] est intéressante que si la source est stationnaire et telle que sa statistique autorise une troncature aisée du réseau.

2 Approche proposée et résumé de l'étude antérieure

L'innovation de notre approche repose alors sur la coopération bénéfique de 2 techniques de codage :

- la quantification simple et rapide avec les réseaux réguliers de points [3],

- la construction d'un dictionnaire arborescent non-équilibré suivant un compromis débit vs. distorsion [2][13].

Dans [11] [12], considérant un réseau pour lequel un algorithme de quantification rapide est connu [3] (c.a.d $Z^n, D_n, E_8, \Lambda_{16}$) :

- ce réseau est tronqué tel qu'il puisse être emboîté, en le contractant, dans son voronoï ;
- la hiérarchie de réseaux tronqués est obtenue en ajustant leur échelle telle que, un réseau de résolution supérieure s'emboîte dans le réseau de résolution juste inférieure ;
- un schéma simple de quantification multi-étages en découle ;
- un dictionnaire arborescent non-équilibré est construit par apprentissage (l'arbre est découpé ou élagué suivant le critère débit vs. distorsion de BFOS) [2].

Exactement l'approche par découpage de l'arbre est retenue car elle demeure adaptée à la construction du dictionnaire lorsque la dimension vectorielle et donc le nombre d'aires de l'arbre sont élevés : par un processus itératif, à l'issue de chaque boucle et suivant le critère débit vs. distorsion, une seule feuille de l'arbre est découpée (un seul emboîtement est réalisé). Cependant il faut noter que cette approche locale est sous-optimale comparée à celle globale d'élagage de l'arbre où, après une première étape de construction d'un arbre complet, à chaque boucle d'un processus itératif et suivant le critère débit vs. distorsion une branche est supprimée.

De plus l'énergie de troncature du réseau à emboîter est choisit minimale afin de réduire le nombre de nouveaux points représentants injectés à chaque nouvel emboîtement. Ainsi la découpe de l'espace se fait progressivement : les zones spatiales découpées sont localisées avec plus de précision, et les bits alloués avec justesse.

De premiers résultats ont illustrés comment ce Quantificateur Algébrique et Arborescent (QVAA) est adapté au codage de sources d'images différentielles ou hybrides car, pour un débit fixé, ce quantificateur découpe grossièrement les régions spatiales où se concentrent les vecteurs peu énergétiques et peut découper plus finement les régions moins denses où se situent les vecteurs source riches en information.

Notre étude vise à présent à déterminer le réseau algébrique le plus efficace dans ce contexte d'emboîtement. Ensuite nous présentons un test rapide de trie des vecteurs à coder, test devant être mis en place lorsque la source n'appartient pas à la SA utilisée pour construire le dictionnaire (ce test détecte les vecteurs dont la norme dépasse celle maximale envisagée par le dictionnaire). Enfin des résultats de codage d'une source réelle différentielle sont analysés (le codeur fonctionnant en boucle fermée).

3 Détermination du réseau optimal dans ce contexte d'emboîtement

Conway et Sloane [3] ont développés pour des réseaux réguliers de points (treillis) des algorithmes de quantification rapides. Pour notre application nous faisons notre choix parmi les

plus rapides qui sont $Z^k (k \geq 1)$, $D_k (k \geq 2)$, E_8 et Λ_{16} (k indiquant la dimension vectorielle). Nous précisons que D_4 , E_8 et Λ_{16} sont, pour leur dimension vectorielle respective et sous l'hypothèse haute-résolution, les treillis meilleurs quantificateurs (leur moment d'ordre 2 est minimal) [4].

Des courbes expérimentales débit vs. distorsion (entropie du dictionnaire vs. distorsion moyenne) sont représentées figure 1 a. Elles sont obtenues par le codage d'une source synthétique i.i.d normale ($\sigma^2 = 1$) et par découpage de l'arbre. Nous comparons donc Z^4 à D_4 , et Z^8 à D_8 . Dans les deux cas le réseau Z^k apparaît plus performant (de plus bas débits sont atteints, et la distorsion est inférieure). Les courbes de la figure 1 b obtenues avec Z^k ($k = 1, 2, 4, 8$ et 16), montrent que de plus bas débits sont évidemment atteints lorsque la dimension vectorielle croît.

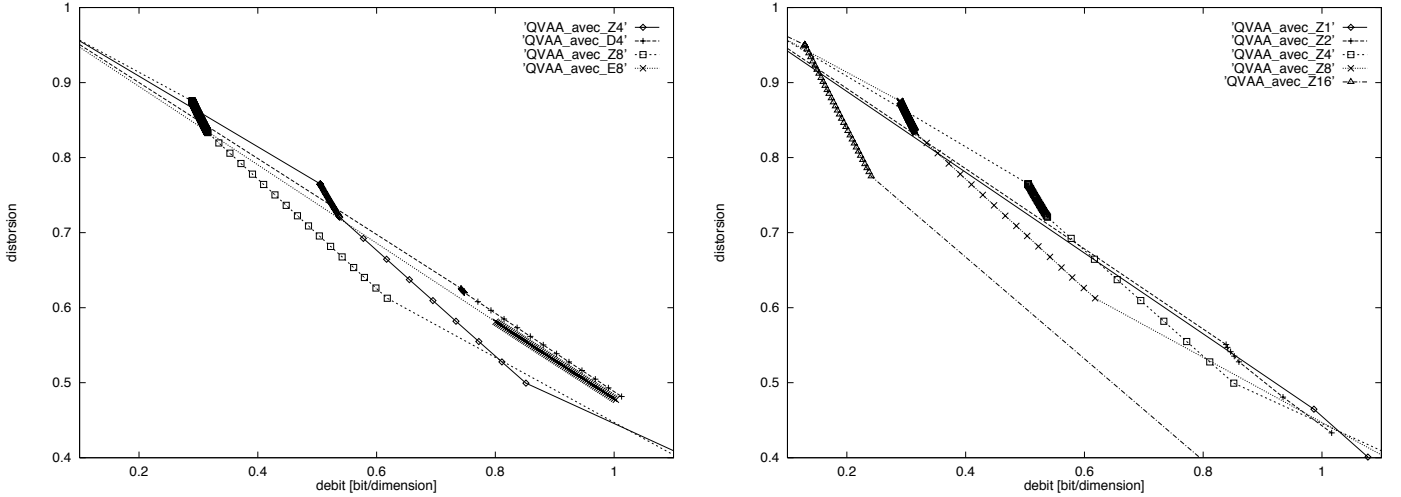


Figure 1: (a,b) Courbes débit vs. distorsion (entropie du dictionnaire vs. distorsion) obtenues par découpage du dictionnaire arborescent et mettant en jeu différents réseaux. La source iid obéit à une loi normale ($\sigma^2 = 1$).

Nous rappelons que l'emboîtement est l'opération qui consiste à inclure, dans un voronoï récepteur d'un treillis à une résolution donnée, un même treillis tronqué de résolution supérieure. L'emboîtement est optimal si le volume du voronoï récepteur n'est rempli que de voronoï entiers du treillis emboîté. Ce résultat est dû aux propriétés géométriques intrinsèques du treillis utilisé. La méthode de troncature décrite dans [11] [12] ne peut donc garantir qu'un emboîtement presque optimal où le voronoï récepteur est rempli d'un maximum de voronoï entiers du treillis emboîté.

Z^k se distingue car seul ce réseau fournit un emboîtement optimal. Afin d'appréhender l'intérêt de cette propriété, examinons la figure 2 qui présente l'exemple d'un emboîtement (optimal) avec Z^2 et celui (presque optimal) avec le réseau hexagonal, la source est distribuée uniformément dans les zones grisées considérées de même aire. Dans le cas du réseau cubique, les points injectés participent identiquement à la baisse de distorsion moyenne et à la hausse de débit entropique entraînées par ce découpage. Dans le cas hexagonal, 6 des 13 points injectés ont une probabilité d'occurrence moindre. Ceci se traduit par rapport au cas précédent, par une hausse supérieure du débit entropique et une baisse moindre de la

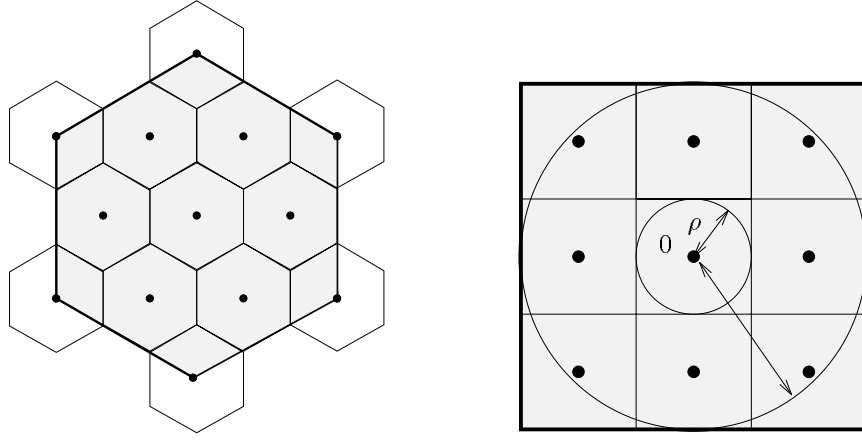


Figure 2: Exemples de 2 emboîtements : un optimal avec le réseau cubique, l'autre presque optimal avec le réseau hexagonal.

distorsion moyenne.

Notons que le fait de retenir le treillis cubique implique de ne tirer profit de l'"avantage de remplissage d'espace" de la QV sur la quantification scalaire [9]. Cependant il est montré que le gain espéré est peu significatif [9]. Enfin il faut remarquer que Z^k est le treillis le moins dense, le nombre de points injectés par emboîtement qui est aussi le nombre d'aires de l'arbre sera réduit, ce qui ajoute à la précision du découpage de l'espace vectoriel.

4 Quantification de sources n'appartenant pas à la séquence d'apprentissage

Le dictionnaire est construit à partir d'une SA représentative de la statistique de la source à coder. Il est néanmoins nécessaire de s'assurer, lors du codage de vecteurs non issus de cette séquence, que ceux-ci ont une norme inférieure à celle maximale envisagée pour le dictionnaire. Nous décrivons donc la mise au point d'un test rapide qui, placé à l'entrée du QVAA détecte les vecteurs marginaux qui seront codés à part.

Nous rappelons la formule du facteur introduit dans [11] [12] qui, appliqué aux coordonnées du vecteur source, permet de le projeter au sein du premier réseau tronqué :

$$F = \frac{3 \times \rho}{\sqrt{\mathcal{E}_{max}}}$$

avec

$$\mathcal{E}_{max} = \max_{\mathbf{x}} \left\{ \mathcal{E}(\mathbf{x}) = L_2(\mathbf{x}) = \sum_{i=1}^k x_i^2 \quad / \quad \mathbf{x} \in \text{SA} \right\}$$

où ρ est le rayon d'empilement du réseau ($\rho = 1/2$ pour Z^k) et, $\mathcal{E}_{max}$ l'énergie maximale d'un vecteur de la SA.

Le sous-ensemble des points du réseaux Z^k conservé est alors constitué de ceux dont le voronoï est entièrement ou partiellement à l'intérieur de la sphère de rayon $(3 \times \rho)$ car :

$$\sum_{i=1}^k (F \times x_i)^2 = F^2 \times \mathcal{E}(\mathbf{x}) = \frac{(3 \times \rho)^2}{\mathcal{E}_{max}} \times \mathcal{E}(\mathbf{x}) \leq (3 \times \rho)^2 \quad / \quad \mathbf{x} \in \text{SA}$$

Ce sous-espace tronqué est un cube, et la norme L_∞ des vecteurs $\mathbf{u}$ qui y sont inscrits est telle que (la figure 2 l'illustre) :

$$L_\infty(\mathbf{u}) = \max_{i=1,\dots,k} |u_i| \leq (3 \times \rho)$$

Alors la norme L_∞ , avant projection, d'un vecteur source $\mathbf{x}$ pouvant être coder par ce dictionnaire est telle que :

$$L_\infty(\mathbf{x}) = \max_{i=1,\dots,k} |x_i| \leq \sqrt{\mathcal{E}_{max}}$$

Le test de la norme L_∞ du vecteur est évidemment de complexité calculatoire inférieure à un test en norme L_2 .

5 Résultats expérimentaux

Des résultats de codage d'une source différentielle sont présentés sans aucune recherche d'un cas de prédiction optimale. Le schéma du codeur MICD est montré figure 3. L'image est balayée de haut en bas et de la gauche vers la droite et, en fonction de la dimension vectorielle choisie, 5 formes de blocs sont possibles. Le prédicteur fixe est grossier : la prédiction d'un bloc est construite par duplication de pixels de la ligne précédente (la première ligne est considérée transmise telle quelle).

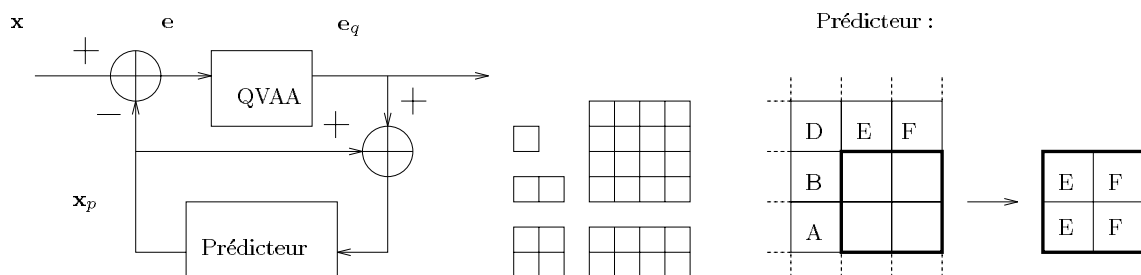


Figure 3: Schéma du codeur MICD manipulant des vecteurs, formes des blocs et principe du prédicteur.

Les SA sont calculées en boucle ouverte. Leurs tailles sont adaptées à la dimension vectorielle afin d'obtenir un rapport du nombre de vecteurs de la séquence sur celui des vecteurs représentants suffisant (typiquement supérieur à 100).

Observons premièrement les courbes débit vs. PSNR (entropie du dictionnaire vs. Peak Signal-to-Noise Ratio) de la figure 4 obtenues lors de la construction des dictionnaires. La prédiction étant trop rustre lorsque la dimension vectorielle croît, il faut s'attacher à observer les gains en PSNR apportés par la QVAA. Notons qu'un saut de la courbe indique le découpage de la région dense au centre de l'espace vectoriel, la série de découpages qui suit se produit à la périphérie de cette zone centrale.

Ces courbes confirment que de bas débits ne sont atteints qu'en mettant en jeu de hautes dimensions vectorielles ; cependant les tailles des SA nécessaires à la construction des dictionnaires deviennent vite considérables. Nous donnons à présent des résultats numériques

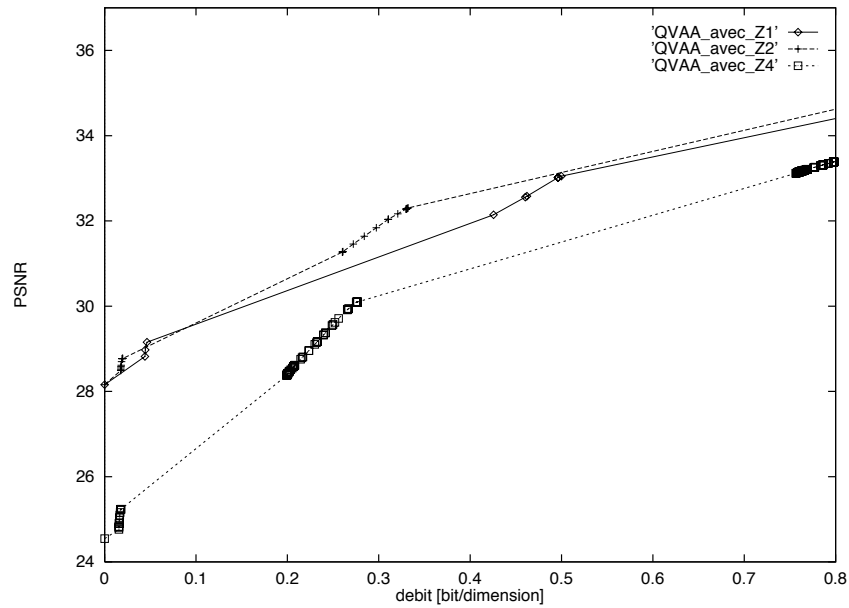


Figure 4: (a,b) Courbes débit vs. PSNR (entropie du dictionnaire vs. PSNR) obtenues lors de la construction des dictionnaires pour Z^k avec $k = 1, 2$ et 4 . Les SA sont respectivement constituées de 1, 1 et 4 images.

obtenus lors du codage des images LENA (512x512) et INTERVIEW (674x536). La encore, notons que la prédiction devient trop grossière lorsque la dimension vectorielle croît ce qui pénalise complètement ces résultats obtenus en boucle fermée, c'est le gain de quantification qu'il faut observer relativement au débit alloué. Nous rappelons les définitions du PSNR, du gain de prédiction (GP) et de celui de quantification (GQ) (les e_i sont les erreurs de prédiction, les e_{iq} celles quantifiées et N la taille de l'image) :

$$\text{PSNR [db]} = 10 \times \log_{10} \frac{255^2}{\frac{1}{N} \times \sum e_i^2} + 10 \times \log_{10} \frac{\sum e_i^2}{\sum (e_{iq} - e_i)^2} = \text{GP [db]} + \text{GQ [db]}$$

Les calculs sont effectués sur une machine SPARCStation 20, 75Mhz (compilateur acc).

image LENA			
dimension vectorielle	1	2	4
temps CPU contruction du dictionnaire	13,07 s	13,3 s	46,28 s
temps CPU encodage image	21,36	12,01	10,09
PSNR	28,225	26,451	25,85
GP	25,218	24,235	22,839
GQ	3,007	2,216	3,019
entropie des $\mathbf{e}_q$ [bpp]	0,828	0,627	0,358

image INTERVIEW			
dimension vectorielle	1	2	4
temps CPU contruction du dictionnaire	13,07 s	13,3 s	46,28 s
temps CPU encodage image	29,46	16,33	17,75
PSNR	28,037	26,126	25,553
GP	23,960	23,115	21,614
GQ	4,077	3,011	3,929
entropie des $\mathbf{e}_q$ [bpp]	0,941	0,709	0,429

Les images de la figure 5 illustrent comment le critère débit vs. distorsion agit : pour un débit fixé, les vecteurs peu énergétiques et de forte probabilité (typiquement ceux associés aux zones homogènes de l'image différentielle) sont grossièrement quantifiés, alors les vecteurs moins probables et plus énergétiques (typiquement ceux associés aux zones inhomogènes) peuvent être quantifiés plus finement.

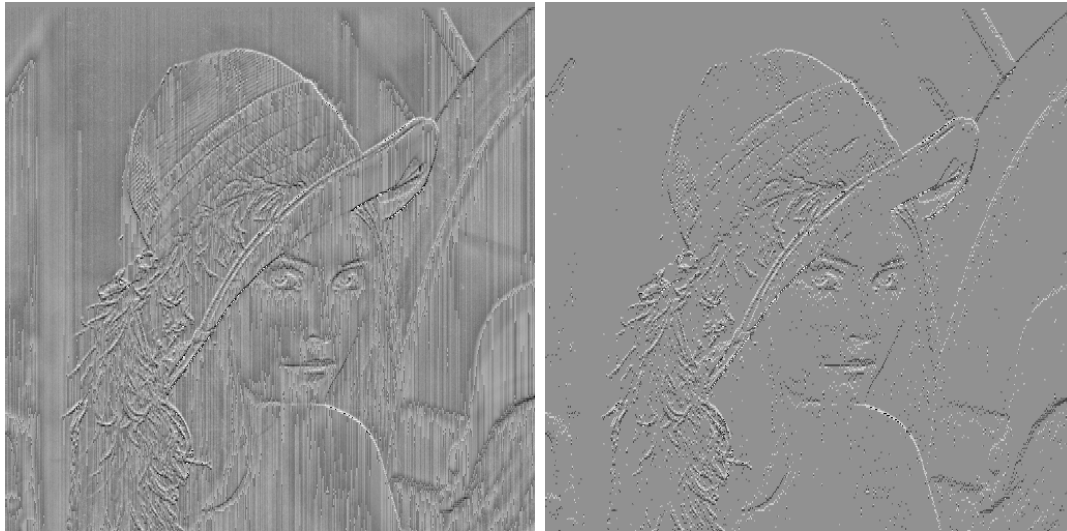


Figure 5: *(a,b) Images des erreurs de prédiction centrées avant (a) et après (b) quantification, la QVAA est réalisée avec Z^4 (cas où $GQ = 3,19$ db et l'entropie des $\mathbf{e}_q$ égale 0.358 bpp).*

6 Perspectives

Ces résultats justifient l'utilisation de dimension vectorielles élevées pour atteindre de bas débits, cependant cela implique la manipulation de SA de tailles imposantes. La solution serait de construire une "souche" au dictionnaire arborescent à partir d'une première SA de très grande taille représentative d'une large gamme de types d'images (ces premiers découpages sont d'ailleurs les plus coûteux en terme de calculs). Ensuite, pour une type d'images données, à partir d'une SA de taille plus restreinte et représentative de cette source, le dictionnaire serait achevé. Cette méthode de quantification, qui sera effectivement rapide, pourra s'inscrire dans un contexte de codage adaptatif.

En seconde perspective, il est envisagé de tester cette approche de QVAA dans un cadre hybride (transformée fréquentielle et prédiction) où la source différentielle serait générée par l'intermédiaire d'une compensation de mouvement.

References

- [1] M. Barlaud, P. Solé, T. Gaidon, M. Antonini, and P. Mathieu. – Pyramidal lattice vector quantization for multiscale image coding. – *IEEE Transactions on Image Processing*, 3(4):367–381, July 1994.
- [2] L. Breiman, J.H. Friedman, R.A. Olshen, and C.J. Stone. – *Classification and regression Trees*. – The Wadsworth Statistics/Probability Series. Wadsworth, Belmont, California, 1984.
- [3] J.H. Conway and Sloane N.J.A. – Fast quantizing and decoding algorithms for lattice quantizers and codes. – *IEEE Transactions on Information Theory*, IT-28(2):227–232, March 1982.
- [4] J.H. Conway and Sloane N.J.A. – *Sphere Packings, Lattices and Groups, 2nd edition*. – A series of Comprehensive Studies in Mathematics. Springer-Verlag, New York, 1993.
- [5] T.R. Fisher. – A pyramid vector quantizer. – *IEEE Transactions on Information Theory*, IT-32(4):568–583, July 1986.
- [6] A. Gersho and R.M. Gray. – *Vector Quantization and Signal Compression*. – Kluwer Academic Publishers, Boston, 1992.
- [7] M. Goldberg, P.R. Boucher, and S. Schlien. – Image compression using adaptative vector quantization. – *IEEE Transactions on Communications*, 34(2):180–187, February 1986.
- [8] Y. Linde, A. Buzo, and R.M. Gray. – An algorithm for vector quantizer design. – *IEEE Transactions on Communications*, 28:84–95, 1980.
- [9] T.D. Lookabaugh and R.M. Gray. – High-resolution quantization theory and the vector quantizer advantage. – *IEEE Transactions on Information Theory*, 35(5):1020–1033, September 1989.
- [10] P. Monet and C. Labit. – Codebook replenishment in classified pruned tree-structured vector quantization of image sequences. – In *Proc. of International Conference on Acoustics, Speech, and Signal Processing*, pages 2285–2288, 1990.
- [11] V. Ricordel and C. Labit. – Vector quantization by hierarchical packing of embedded truncated lattices. – In *Proc. of Visual Communications and Image Processing*. Taiwan, China, May 1995.
- [12] V. Ricordel and C. Labit. – Vector quantization by packing of embedded truncated lattices. – In *Proc. of International Conference on Image Processing*. Washington DC, USA, October 1995.
- [13] E.A. Riskin and R.M. Gray. – A greedy tree growing algorithm for the design of variable rate vector quantizers. – *IEEE Transactions on Signal Processing*, 39(11):2500–2507, November 1991.