

Quantification vectorielle algébrique modulée : codage de source/tatouage conjoints à débit variable

Ludovic GUILLEMOT, Jean-Marie MOUREAUX

Université Henri Poincaré, Nancy 1, CRAN, CNRS UMR 7039,
BP 239, F-54506 Vandœuvre-lès-Nancy Cedex
{ludovic.guillemot, jean-marie.moureaux}@cran.uhp-nancy.fr

Résumé – Nous proposons ici un schéma dit de quantification vectorielle algébrique modulée (QVAM). Il s’agit d’un codeur de source à débit variable permettant l’insertion d’information. Ses bonnes performances en terme de compromis débit distorsion rendent possible son utilisation dans des domaines où la compression avec perte est incontournable, en particulier celui du codage par transformée reconnu comme le plus efficace, notamment pour les images. La QVAM est essentiellement basée sur, d’une part l’utilisation d’un quantificateur vectoriel permettant de diminuer de manière significative la limite de la borne entropique du débit, et d’autre part, sur l’emploi d’une zone morte vectorielle pour exploiter la représentation creuse de la source. En outre, la phase d’estimation de la taille des cellules de quantification est un problème clef des méthodes d’insertion par quantification qui a fait l’objet, à notre connaissance, de travaux uniquement dans le cas scalaire. Nous proposons ici une méthode d’estimation adaptée aux quantificateurs vectoriels algébriques.

Abstract – The scheme we propose here is called modulated lattice vector quantization (MLVQ). It is a variable rate coder dedicated to joint compression and information embedding. Thanks to its good performances in terms of rate distortion trade-off, it can be advantageously used in a transform coding context, for example in image coding. MLVQ is based on one hand, on the use of a vector quantizer permitting to decrease significantly the limit of the entropic bound, and on the other hand, on a vector dead zone which allows to take profit of the sparsity of the source. Moreover, the need to know quantization parameters previous at the extraction step is a key point of quantization based methods. Existing methods are dedicated to scalar quantizers. Here we propose an estimation procedure of the size of the quantization cells adapted to the multidimensional case.

1 L’approche tatouage/compression conjoints

Du fait de l’explosion des communications numériques, la dernière décennie a vu de nombreux travaux de recherche, tant dans le domaine de la compression que dans celui du tatouage. La compression étant une attaque redoutable pour le tatouage, l’approche tatouage/compression conjoints offre de nombreux avantages, parmi lesquels celui pour la compression de ne plus être une source d’interférence lors de l’extraction du tatouage.

Cette approche a été cependant peu abordée dans la littérature ou proposée dans un contexte ad’hoc [5] [9]. Ainsi, par exemple Wu et al. ont développé dans [8] un schéma conjoint effectif basé sur la méthode de quantification par modulation d’index (QIM) [1]. Cependant le codage de source mis en oeuvre est à pas fixe, et par conséquent n’exploite pas les redondances présentes à l’intérieur des signaux. En outre, le cadre scalaire semble peu adapté à ce type d’application en raison de l’existence d’une limite inférieure de la borne entropique (dû au fait que l’on ajoute 1 bit par échantillon). A noter que cette méthode ne permet pas d’exploiter les propriétés de robustesse liées à l’emploi d’un quantificateur vectoriel. Enfin, le problème de l’estimation des paramètres de quantification est crucial dans une approche conjointe qui nécessite de pouvoir faire varier la taille des cellules de quantification, en d’autres termes le taux de compression. A notre connaissance, des solutions à ce problème a été développées dans [2] et [6] mais elles se limitent au cas scalaire.

Le schéma hybride que nous proposons ici est également basé sur le principe de la QIM, mais dans son cadre vectoriel. Il s’agit d’un quantificateur vectoriel algébrique modulé (QVAM) de faible complexité¹ qui permet tatouage et codage à débit variable. L’originalité de nos travaux réside dans le fait que notre méthode permet à la fois d’exploiter les redondances d’une source grâce à l’emploi d’une zone morte vectorielle [7] et d’estimer les paramètres du quantificateur dans le cas vectoriel.

Le plan de cet article est le suivant. Dans le paragraphe 2 nous décrivons l’algorithme QVAM proposé. Le paragraphe 3 est consacré à l’estimation du paramètre de quantification et est suivi par une conclusion.

2 Quantification vectorielle algébrique modulée

Le principe de l’insertion d’information par QIM repose sur le partitionnement d’un dictionnaire en m sous-dictionnaires permettant l’insertion d’un message m -aire². Soit $X \in \mathbb{R}^n$ un vecteur de la source et γ le facteur d’échelle, i.e. le paramètre déterminant la taille des cellules de quantification.

On définit le **quantificateur vectoriel algébrique modulé** Q_i à partir du quantificateur uniforme Q dans $\mathbb{Z}^n$ de la manière

¹A la fois du point de vue du codage de source (grâce à la structure du quantificateur) et du tatouage (effectué conjointement).

²Nous nous limiterons dans cet article à l’insertion d’un message binaire.

suivante :

$$Q_i(X) \triangleq mQ\left(\frac{X - i\frac{\gamma}{m}}{\gamma}\right) + i, \text{ avec } i \in \{0, 1, \dots, m-1\}. \quad (1)$$

L'ensemble des vecteurs obtenus à partir des m quantificateurs Q_i est l'union de m cosets S_i de $\mathbb{Z}^n$: $S_i \triangleq \{m\mathbb{Z}^n + [i]\}$.

Nous appellerons cet ensemble **réseau modulé** $\mathbb{Z}_m^n$:

$$\mathbb{Z}_m^n \triangleq \bigcup_{i=0}^{m-1} S_i, \text{ avec } [i] = \begin{pmatrix} i \\ \vdots \\ i \end{pmatrix}. \quad (2)$$

Afin d'exploiter la structure du réseau nous avons proposé une méthode d'indexage adaptée au réseau modulé [4].

Nous allons à présent expliquer pourquoi coder une source sur ce type de réseau conduit à de faibles performances, en dehors des hauts débits, en terme de compromis débit-distorsion.

Pour des vecteurs de taille n (la population de la couche n du dictionnaire se note $N(n)$) le débit minimal atteignable $\mathcal{H}^\infty$ est donné par la formule suivante :

$$R_C(\gamma) \xrightarrow{\gamma \rightarrow \infty} \mathcal{H}^\infty = 1 + \frac{\log_2(N(n))}{2} \text{ bits par vecteur.} \quad (3)$$

avec R_C le débit entropique de la source.

Preuve. On suppose ici que le message à insérer est équiprobable. Lorsque le facteur d'échelle tend vers l'infini, seulement deux couches du réseau modulé sont utilisées, la couche 0 (correspondant au bit 0) et la couche n (bit 1). D'après la formule du débit du schéma QVA disponible dans [7], on a :

$$\begin{aligned} \mathcal{H}^\infty = & -[P(\|Y\| = 0) \log_2(P(\|Y\| = 0)) + \\ & P(\|Y\| = n) \log_2(P(\|Y\| = n))] + \\ & [P(\|Y\| = 0) \log_2(N(0)) + P(\|Y\| = n) \log_2(N(n))]. \end{aligned}$$

En considérant l'équiprobabilité du message on obtient :

$P(\|Y\| = 0) = P(\|Y\| = n) = \frac{1}{2}$. D'où la relation :

$$\mathcal{H}^\infty = \frac{1}{n} \left\{ 1 + \frac{1}{2} [\log_2(N(0)) + \log_2(N(n))] \right\} \text{ bits/éch.}$$

A noter que l'on a les relation suivantes :

$$\begin{aligned} N(0) &= \text{card}(S_0(0)) = 1 \\ N(n) &= \text{card}(S_0(n)) + \text{card}(S_1(n)) \\ &= \text{card}(N_{\mathbb{Z}^n}(\frac{n}{2})) + \text{card}\{Y = (\pm 1, \dots, \pm 1)\} \\ &= \text{card}(N_{\mathbb{Z}^n}(\frac{n}{2})) + 2^n. \end{aligned}$$

avec $N_{\mathbb{Z}^n}(\frac{n}{2})$ la population de la couche $\frac{n}{2}$ du réseau $\mathbb{Z}^n$. ■

La formule (3) plaide tout d'abord en faveur de l'approche vectorielle, la limite $\mathcal{H}^\infty$ étant de 1,5 bits par éch. pour le cas scalaire contre 0,8 bits/éch., par exemple, pour la QVAM avec des vecteurs de taille 8. Sur la figure 1 sont représentées les courbes distorsion en fonction du débit entropique du quantificateur scalaire uniforme et du quantificateur modulé. Comme on peut le constater, hormis pour les hauts débits, la quantification scalaire modulée est un piètre codeur de source. Un autre facteur entre également en jeu ici : la représentation creuse du signal quantifié est entravée par la modulation.

La contribution majeure de nos travaux réside dans l'introduction du *principe d'exclusion* dans la stratégie d'insertion par QIM : les vecteurs appartenant à une région appelée zone morte vectorielle (ZMV) sont quantifiés par le vecteur nul et exclus du processus d'insertion. De cette manière, la limite $\mathcal{H}^\infty$ est d'une part abaissée et, d'autre part, la représentation creuse du

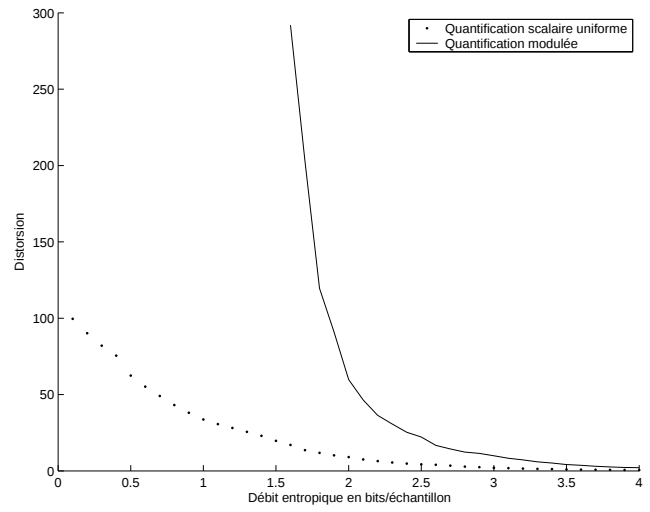


FIG. 1 – Comparaison des courbes débit-distorsion des schémas Quantification Scalaire Uniforme et Quantification Scalaire Modulée : cas d'une source laplacienne centrée d'écart-type $\sigma = 10$.

signal - exploitable par un codeur entropique - est maintenue en sortie du quantificateur.

La figure 2 illustre le principe d'exclusion par zone morte vectorielle, dans le cas de $\mathbb{Z}_2^2$. Les vecteurs de la source appartenant à la zone morte sont quantifiés par zéro et ne contribuent pas à l'insertion du message. Les figures 3 et 4 montrent la supériorité des performances de la QVAM par rapport à la QIM (c'est-à-dire la QVAM sans zone morte), en terme de compromis débit-distorsion. Sur la figure 3 la taille de la zone morte est déterminée de manière à minimiser la distorsion pour un débit donné. Sur la figure 4, la taille est fixée en fonction du taux d'insertion souhaité.

Approche QIM

Approche QVAM

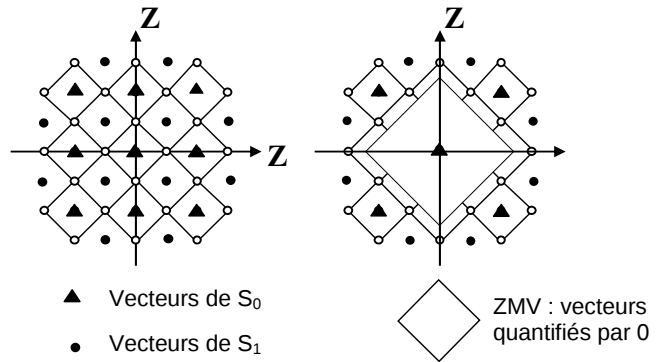


FIG. 2 – Zone morte vectorielle (ZMV) pyramidale sur le réseau $\mathbb{Z}_2^2$.

Les figures 5 et 6 permettent de mieux se rendre compte des performances de la QVAM. On peut y voir deux versions de l'image Lena codées respectivement sans et avec indexage modulé et zone morte vectorielle (l'allocation des débits est la même pour les deux schémas). Le résultat est sans appel : tandis que l'image compressée par QVAM affiche un PSNR correct (33,1 dB), la première montre que le codage sans indexage modulé ni zone morte vectorielle est totalement inadéquat à la

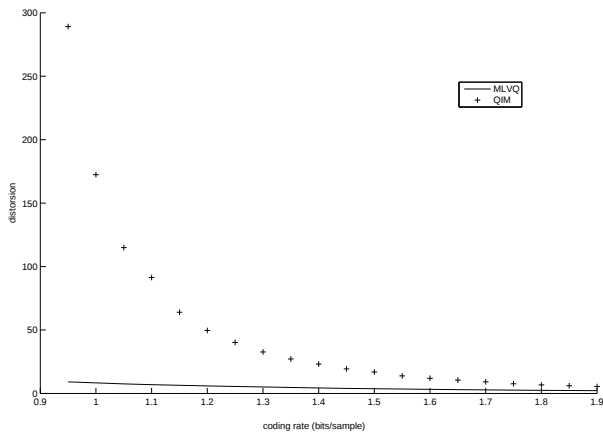


FIG. 3 – Courbes débit-distorsion des schémas QVAM avec zone morte optimale et QIM : cas d'une sous-image d'ondelettes de l'image Lena (filtre 9-7, niveau 3, détails verticaux).

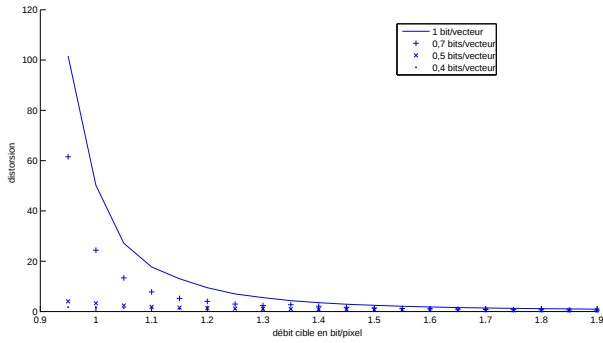


FIG. 4 – Courbe débit-distorsion de l'approche QVAM pour quatre taux d'insertion.

compression.

3 Estimation de la taille de la cellule de quantification

Une des principales limitations des méthodes d'insertion par quantification réside dans le fait que la taille des cellules est un paramètre qui doit être connu à l'extraction du schéma de tatouage. Dans le cadre de l'approche tatouage/compression de source à débit variable il est impossible de fixer cette donnée, puisque de la taille des cellules dépend le débit.

Nous proposons ici une méthode d'estimation de la taille des cellules de quantification (ETCQ) basée sur la requantification de la source tatouée/compressée lors de la phase d'extraction. Cette technique d'estimation est la généralisation aux réseaux de l'approche que nous avons proposé dans [3] pour un quantificateur scalaire uniforme. Soit $\gamma \in \mathbb{R}_+$, on définit $\gamma\Lambda$ comme le réseau dilaté Λ d'un facteur γ . Les deux propriétés du réseau $\gamma\Lambda$ ci-dessous constituent le point central de notre méthode d'estimation :

$$\text{Si } \gamma = k, k \in \mathbb{N}, \text{ alors } \gamma\Lambda \subset \Lambda \quad (4)$$

$$\text{Si } \gamma = \frac{1}{k}, k \in \mathbb{N}, \text{ alors } \gamma\Lambda \supset \Lambda \quad (5)$$



FIG. 5 – Image Lena compressée sans indexage modulé ni zone morte vectorielle, avec un débit de 0,36 bits/pixel et un PSNR de 16,8 dB ; longueur du message inséré : 8,6 kbits.

La fonction distorsion après requantification par un facteur d'échelle $\alpha : \alpha \mapsto d_{Y_{\gamma\Lambda}}(\alpha) = \frac{1}{n} \|Y_{\gamma\Lambda} - Q_{\alpha\Lambda}(Y_{\gamma\Lambda})\|_2^2$ admet alors les deux caractéristiques suivantes :

1. $D_{Y_{\gamma\Lambda}}(\alpha)$ est nulle (et minimale) en $\frac{\gamma}{k}, k \in \mathbb{N}^*$.
2. $D_{Y_{\gamma\Lambda}}(\alpha)$ admet des minima locaux en $\alpha = k\gamma, k \in \mathbb{N}^*$.

L'estimation du facteur d'échelle consiste à déterminer le dernier minimum avant une forte pente. Les figures 7 et 8 illustrent le fonctionnement de l'ETCQ : la source quantifiée/tatouée est requantifiée avec un pas variable. Le facteur d'échelle correspond au dernier minimum local. Comme on peut le voir, ce minimum existe toujours en dépit d'une attaque telle qu'un bruit blanc gaussien additif, ou encore un gain. Dans le cas d'un bruit additif d'écart-type inférieur à 1,5, le taux d'erreur d'extraction est négligeable. Il est de 10% lorsque l'écart-type est de 1,5. Dans le cas du gain, les minima correspondent bien au facteur d'échelle multiplié par le gain utilisé et les erreurs d'extraction sont nulles à chaque fois.

4 Conclusion

Nous avons présenté ici un schéma original de compression et tatouage conjoints : le quantificateur vectoriel algébrique modulé. L'approche conjointe offre les avantages suivants : faible complexité, aucune erreur d'extraction due à la compression avec perte (en l'absence d'autres attaques). Tenant compte des caractéristiques des schémas de compression les plus performants à l'heure actuelle, la QVAM a été mis au point afin de permettre un codage à longueur variable, et surtout de tirer profit de la parcimonie des sources obtenues après transformation.

Cette dernière propriété n'est à notre connaissance pas représentée dans la littérature. Pour ce faire, la QVAM repose d'une part, sur les bonnes propriétés en terme de débit entropique des quantificateurs vectoriels algébriques et, d'autre part sur l'exclusion du processus d'insertion de certains éléments de



FIG. 6 – Image Lena compressée avec indexage modulé et zone morte vectorielle, avec un débit de 0,36 bits/pixel et un PSNR de 33,1 dB ; longueur du message inséré : 3,8 kbits.

la source à travers une zone morte vectorielle. Ce dernier point permettant de maintenir une représentation creuse de la source en sortie du quantificateur.

Enfin, un problème clé relatif à l'approche conjointe a été également abordé : la taille des cellules de quantification ne peut être figée dans l'optique d'un codeur de source performant. L'efficacité de notre méthode d'estimation ETCQ permet de résoudre ce problème, y compris en cas d'attaques telles qu'un bruit additif, ou encore un gain.

Les performances de codage de notre schéma sont encourageantes et démontrent l'intérêt de l'approche conjointe compression / tatouage dans un contexte où la compression constitue la principale attaque. Les questions liées à l'amélioration de la robustesse avec l'emploi de réseaux mieux adaptés ou de codes correcteurs d'erreurs offrent de nombreuses perspectives de recherche. Enfin, la transposition de notre approche à d'autres types de codeurs semble être également un champ d'étude très prometteur.

Références

- [1] B. Chen et G. Wornell, "Quantization Index Modulation : A Class of Provably Good Methods for Digital Watermarking and Information Embedding", IEEE Trans. on Information Theory, vol. 47, NO. 4, May 2001.
- [2] J. J. Eggers, R. Bäuml, R. Tzschoppe et B. Girod, "Scalar Costa Scheme for Information Embedding", IEEE Trans. on Signal Processing, vol. 51, No. 4, pp. 1003-1019, avril 2003.
- [3] L. Guillemot et J. M. Moureaux "Bite-Rate adapted Watermarking algorithm for compressed images", IEEE ICME 2002, Lausanne, août 2002.
- [4] L. Guillemot et J. M. Moureaux, "Hybrid transmission, compression and data hiding : quantisation index modulation as source coding strategy", Electronics letters Vol. 40, No. 17, pp. 1053-1055, août 2004.
- [5] J. M. Moureaux et L. Guillemot, "Image Compression and Watermarking using Lattice Vector Quantization", SPIE security

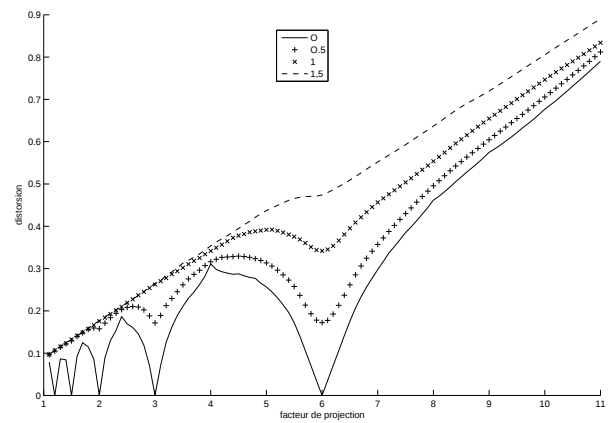


FIG. 7 – Courbes distorsion en fonction du facteur d'échelle en l'absence d'attaques et en présence d'un bruit blanc gaussien additif d'écart-type 0,5 ; 1 et 1,5 (facteur d'échelle initial égal à 6).

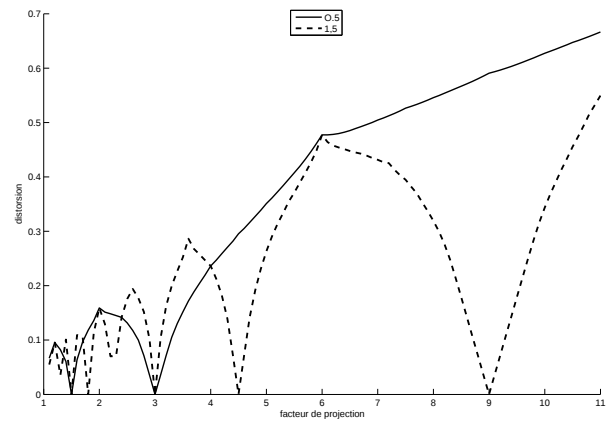


FIG. 8 – Courbes distorsion en fonction du facteur d'échelle en présence d'un gain de 0,5 et 1,5 (facteur d'échelle initial égal à 6).

and watermarking of multimedia contents, pp. 600-610, San Jose, janvier 2002.

- [6] I. D. Shterev, I. L. Lagendijk et R. Heusdens, "Statistical Amplitude Scale Estimation for Quantization-based Watermarking", SPIE security, steganography and watermarking of multimedia contents VII, San José, janvier 2004.
- [7] T. Voinson, L. Guillemot et J.M. Moureaux, "Image compression using Lattice Vector Quantization with code book shape adapted thresholding", IEEE 2002 International Conference on Image Processing, Rochester, septembre 2002.
- [8] G. Wu et E.H. Yang, "Joint watermarking and compression using scalar quantization for maximizing robustness in the presence of additive gaussian attacks", IEEE Trans. on Signal Processing, vol. 53, NO. 2, pp. 834-844, février 2005.
- [9] L. Xie et G.R. Arce, "A Class of Authentication Digital Watermarks for Secure Multimedia Communication," IEEE Trans. on Image Processing, vol. 10, NO. 11, 2001.