

INSTITUT NATIONAL POLYTECHNIQUE DE GRENOBLE

N° attribué par la bibliothèque

THÈSE

pour obtenir le grade de

DOCTEUR DE L'INPG

Spécialité : « Mathématiques appliquées »

préparée au **Laboratoire des Images et des Signaux**

dans le cadre de l'École Doctorale « Mathématiques,
Sciences et Technologies de l'Information, Informatique »

présentée et soutenue publiquement par

Laurent CONDAT

le 18 septembre 2006

Titre :

Méthodes d'approximation pour la reconstruction de signaux et le redimensionnement d'images

Directrice de thèse : Annick MONTANVERT

JURY

Mme Valérie Perrier

Mme Laure Blanc-Féraud Rapporteuse

M. Patrick Flandrin Rapporteur

Mme Annick Montanvert Directrice de thèse

M. Michael Unser

Remerciements

J'adresse tout d'abord mes plus chaleureux remerciements à ma directrice de thèse, Annick Montanvert, pour sa gentillesse, sa bonne humeur, et son esprit éclairé. Malgré ses responsabilités, elle a su se montrer toujours présente et attentive, me montrer la voie tout en me laissant une grande liberté de pensée et d'action. Plus qu'à ma directrice de thèse, c'est à une amie que je témoigne ma plus sincère reconnaissance. Qu'elle sache, au-delà de l'estime que je lui porte, combien je lui suis redevable pour l'accomplissement de ce travail de recherches.

Cette thèse a été réalisée au Laboratoire des Images et des Signaux (LIS) de Grenoble. J'exprime ma sympathie à tous ses membres pour avoir insufflé dans nos murs une atmosphère chaleureuse et conviviale. Le dynamisme et l'ambiance d'un laboratoire tiennent pour une bonne part à ses doctorants. Je témoigne mon amitié à Cyril Kotenkoff, mes joyeux co-bureaux, et beaucoup d'autres, pour avoir partagé mon quotidien de doctorant.

Travailler en collaboration est une aventure stimulante et très enrichissante, sur les plan scientifique et humain. Merci à Thierry Blu, Dimitri Van de Ville, Michael Unser et les autres membres du Laboratoire d'Imagerie Biomédicale (LIB) de Lausanne, pour m'avoir accueilli dans leur équipe. Leur charisme et leur grande maîtrise des arcanes de la recherche ont eu un effet très bénéfique sur mon développement personnel. Je les remercie pour l'influence positive qu'ils ont eu sur l'orientation de ce travail de thèse. Je témoigne aussi ma sympathie à Muthuvel Arigovindan, Akira Hirabayashi, et les doctorants venus d'horizons divers, que j'ai pu rencontrer au LIB.

Les collaborations ne sont qu'une facette de l'interactivité qui règne dans le milieu de la recherche. Merci à toutes les personnes que j'ai pu cotoyer et avec qui j'ai échangé des idées. Je remercie en particulier Monsieur Abdelmounim Belahmidi pour m'avoir envoyé un programme implémentant sa méthode d'agrandissement d'images.

Je souhaite aussi exprimer ma reconnaissance envers les membres de mon jury, en particulier pour leurs avis éclairés sur mon travail. Merci beaucoup à Valérie Perrier, qui m'a fait l'honneur de présider ce jury, à Laure Blanc-Féraud et Patrick Flandrin, pour avoir pris le temps d'étudier et de rapporter ce manuscrit. Je remercie aussi vivement Michael Unser d'avoir pris part à ce jury.

Plus qu'à quiconque, je dédie cette thèse aux membres de ma famille. C'est grâce à leur soutien sans faille que j'ai pu avancer dans la vie. Au travers du magazine "Chasseur d'Images" qu'il m'a fait découvrir, mon père a fait naître en moi l'amour de la photographie et, par voie de conséquence, m'a initié au monde exaltant des images.

Merci à tous pour avoir guidé mes pas vers la connaissance des images, ces mosaïques de pixels aux charmes insoupçonnés.

Sommaire

Table des notations	vii
----------------------------	------------

1 Introduction : contexte et objectifs de l'étude	1
1.1 Motivations	1
1.2 Préliminaires mathématiques et notations	3
1.3 Modèle pour les signaux discrets	6
1.3.1 Considérations sur la nature du processus s	10
1.4 Plan de la thèse	11

I Reconstruction de signaux et images discrets	13
---	-----------

2 Reconstruction 1D à partir de données uniformes	15
2.1 Introduction	15
2.2 Approches optimales	16
2.2.1 Approche variationnelle	16
2.2.2 Approche minimax	18
2.2.3 Approche stochastique	19
2.3 Reconstruction au sens des moindres carrés	20
2.3.1 Interpolation	21
2.3.2 Reconstruction consistante dans le cas non bruité	28
2.3.3 Approches des moindres carrés dans le cas bruité	31
2.4 Estimation de s au sens L_2 dans le cas non bruité	32
2.4.1 Situation du problème	33
2.4.2 Approches minimax	35
2.4.3 Analyse de l'erreur	37
2.5 Reconstruction asymptotiquement L_2 -optimale	39
2.5.1 Conception de préfiltres quasi-interpolants	41
2.6 Application : opérations géométriques sur les signaux discrets	43
2.7 Estimation dans un espace LSI en présence de bruit	51
2.8 Conclusion	55

3	Reconstruction 2D par quasi-projections	57
3.1	Introduction	57
3.1.1	Treillis et signaux 2D	58
3.2	Reconstruction 2D dans un espace LSI	61
3.3	Hex-Splines	62
3.3.1	De l'optimalité du treillis hexagonal	64
3.4	Box-splines pour le treillis hexagonal	71
3.4.1	Les B-splines revisit��es	73
3.4.2	Caract��risation diff��rentielle des box-splines �� trois directions	75
3.4.3	Mise en ��uvre	78
3.5	Quasi-interpolation optimale sur la grille hexagonale	80
3.6	Application au r���chantillonnage hexagonal-vers-orthogonal	85
3.6.1	Principes et r��sultats exp��rimentaux	85
3.6.2	Evaluation des mod��les box-spline et hex-spline	90
3.7	Conclusion	91
4	Reconstruction �� partir d'��chantillons non uniformes	93
4.1	Introduction	93
4.2	Approches traditionnelles	95
4.3	Reconstruction dans un espace LSI	98
4.3.1	Position du probl��me	98
4.3.2	Expression de la solution	99
4.3.3	Reconstruction avec des conditions de bord sym��triques	102
4.3.4	Mise en ��uvre de la m��thode	103
4.3.5	Influence des param��tres	108
4.3.6	Applications	113
4.4	Reconstruction par quasi-projections	115
4.5	Conclusion	119
II	Redimensionnement d'images	121
5	Mod��le pour le redimensionnement	123
5.1	Introduction	123
5.2	Formalisation du redimensionnement	124
5.2.1	Consid��rations sur la fonction $\tilde{\varphi}$	128
5.3	R��duction d'images	131
5.3.1	R��duction de facteur entier	133
5.3.2	R��duction de facteurs non entiers	138
5.4	Conclusion	138

6	Agrandissement d'images	143
6.1	Introduction	143
6.1.1	Contexte de l'étude	144
6.1.2	Objectifs et plan du chapitre	145
6.2	Méthodes linéaires d'agrandissement	146
6.2.1	Interpolation	147
6.2.2	Agrandissement de facteur entier sous contrainte de réduction . . .	149
6.2.3	Agrandissement asymptotiquement optimal	151
6.3	Agrandissement régularisé	154
6.3.1	Régularisation variationnelle	155
6.3.2	Autres méthodes de régularisation	157
6.4	Agrandissement par extrapolation	158
6.5	Agrandissement préservant les structures	161
6.5.1	Méthodes adaptatives	161
6.5.2	Zoom fractal	163
6.5.3	Approches basées contours	163
6.5.4	Méthodes par apprentissage	164
6.6	WiskI : une nouvelle méthode d'interpolation adaptative	165
6.7	Conclusion	168
7	Agrandissement d'images par induction	171
7.1	Introduction	171
7.2	L'agrandissement : un problème inverse	172
7.2.1	L'ensemble induit	172
7.2.2	L'agrandissement : un problème d'extrapolation	174
7.3	L'induction de Didier Calle	176
7.4	L'induction oblique rapide	179
7.4.1	L'induction rapide	179
7.4.2	L'induction oblique	180
7.4.3	Interprétation dans le domaine ondelettes	182
7.5	Extensions de l'induction	184
7.5.1	Induction relaxée	184
7.5.2	Induction de facteur quelconque	185
7.6	Validation expérimentale	188
7.6.1	Expériences de réduction/agrandissement combinées	188
7.6.2	Résultats d'agrandissement au moyen de différentes méthodes . . .	190
7.6.3	Campagnes de tests subjectifs	217
7.7	Conclusion	218

Conclusion générale	221
Résumé	221
Mise en perspective	222
Horizons	224
Annexe 1 : Images classiques de la littérature	227
Bibliographie	233

Table des notations

Notations générales

I	Nombre complexe $\sqrt{-1}$
$.*$	Conjugaison complexe
$.^T$	Transposition matricielle
$\mathbf{I}$	Matrice identité
$.*.$	Convolution (discrète ou continue) de deux signaux
$\hat{\cdot}$	Transformée de Fourier
$\omega, \boldsymbol{\omega}$	Fréquence angulaire
$\mathbb{Z}$	Ensemble des entiers relatifs
$\llbracket a, b \rrbracket$	Intervalle des entiers entre $a \in \mathbb{Z}$ et $b \in \mathbb{Z}$
$\mathbb{R}$	Ensemble des réels
$\mathcal{F}$	Transformée de Fourier (en tant qu'opérateur formel)
$\mathcal{P}_\Omega^\perp$	Projection orthogonale dans l'ensemble Ω
$\mathcal{D}.$	Opérateur de discrétisation produisant v à partir de s
$\mathcal{M}.$	Opérateur de reconstruction produisant une estimée de s à partir de v
$\ \cdot\ _{\mathcal{X}}$	Norme sur l'espace vectoriel $\mathcal{X}$
$\langle \cdot, \cdot \rangle_{\mathcal{X}}$	Produit scalaire dans l'espace de Hilbert $\mathcal{X}$
$O(\cdot)$	Relation de domination
$o(\cdot)$	Relation de négligeabilité
$\cdot \sim \cdot$	Relation d'équivalence
$\binom{\cdot}{\cdot}$	Coefficient binomial
$\mathbf{R}$	Matrice 2×2 caractérisant un treillis 2D
$\Lambda_{\mathbf{R}}$	Treillis 2D, caractérisé par la matrice $\mathbf{R}$
Γ_v	Ensemble de fonctions consistantes avec le signal v
Υ_v	Ensemble d'images agrandies vérifiant la contrainte de réduction avec l'image v
α	Facteur de réduction ou d'agrandissement

Signaux continus et fonctions

$s, s(t), s(\mathbf{x})$	Processus physique défini continûment, engendrant v
φ	Fonction génératrice pour la reconstruction
$\tilde{\varphi}$	Fonction d'analyse servant à la formation de v
β^n	B-spline 1D, ou 2D séparable, de degré n
γ^n	spline cardinale (interpolante) 1D, ou 2D séparable, de degré n
η_L	Hex-spline déployée sur le treillis hexagonal, d'ordre d'approximation L
χ_{Ξ}	Box-spline dont les directions sont les colonnes de la matrice Ξ
χ_{2n}	Box-spline à trois directions sur le treillis hexagonal, d'ordre d'approx. $2n$
$\tilde{\varphi}_T$	Raccourci pour $\frac{1}{T^d} \tilde{\varphi}(\frac{\cdot}{T})$
sinc	Fonction sinus cardinal
L	Ordre d'approximation (généralement, de φ)
L_2	Ensemble des fonctions de carré intégrable
W_2^r	Espace de Sobolev d'ordre r
$f^{(r)}$	Dérivée d'ordre r de f (au sens de Sobolev)
$\ \cdot\ _{\infty}$	Norme infinie
c_f	Autocorrélation d'un processus aléatoire continu dont f est une réalisation
$\mathbb{1}_{\Omega}$	Fonction indicatrice sur le domaine Ω
$\zeta(t)$	Fonction de Riemann-Zeta ($t > 1$)
$\mathcal{L}$	Opérateur de L_2 dans L_2
$\hat{\mathcal{L}}$	Transformée de Fourier de la réponse impulsionnelle de l'opérateur $\mathcal{L}$
$\mathcal{P}_{\phi}$	Projection orthogonale dans l'ensemble $V_T(\phi)$
$\mathcal{Q}_{\cdot}$	Opérateur d'approximation associant à s son estimée à partir de v
$D_{\mathbf{r}}$	Opérateur de différentiation dans la direction $\mathbf{r}$
$h_{\alpha, \tau}$	filtre continu pour l'agrandissement de facteur α non entier
$\tilde{h}_{\alpha, \tau}$	filtre continu pour la réduction de facteur α non entier

Signaux discrets

$v, (v[\mathbf{k}])_{\mathbf{k} \in \mathbb{Z}^2}$	Signal discret ou image discrete formant la donnée initiale
v_0	Version non bruitée de v
μ	Réalisation d'un bruit additif centré stationnaire
c_u	Autocorrélation du signal discret $u \in \ell_2$ ou du processus aléatoire discret dont u est une réalisation
a_ϕ	Autocorrélation de la fonction ϕ discrétisée sur le treillis courant
b_ϕ	Signal discret obtenu en discrétisant la fonction ϕ sur le treillis courant
b_m^n	B-spline de degré n discrétisée avec un pas de $1/m$
$\mathcal{Z}$	Transformée en $\mathbf{z}$ (en tant qu'opérateur formel)
$\mathcal{T} \cdot$	Opérateur de transformation géométrique sur un signal discret
$\mathcal{T}_{\alpha, \tau} \cdot$	Opérateur de redimensionnement d'images
$\mathcal{R}_{\alpha, \tau} \cdot$	Opérateur de réduction d'images
$\mathcal{A}_{\alpha, \tau} \cdot$	Opérateur d'agrandissement d'images
p	Préfiltre pour la reconstruction ou l'agrandissement linéaire
$h_{\alpha, \tau}$	filtre discret pour l'agrandissement de facteur α entier
$\tilde{h}_{\alpha, \tau}$	filtre discret pour la réduction de facteur α entier
u^{-1}	Inverse, au sens de la convolution, de u
$U(\mathbf{z})$	Tranformée en $\mathbf{z}$ de u
$[\cdot] \downarrow a$	sous-échantillonnage/décimation de facteur entier $a \geq 1$
$[\cdot] \uparrow a$	sur-échantillonnage de facteur entier $a \geq 1$

Chapitre 1

Introduction : contexte et objectifs de l'étude

1.1 Motivations

DANS les sociétés modernes, l'activité humaine, tant à des fins professionnelles que de divertissement, se dématérialise, et repose de plus en plus sur des supports médiatiques emmagasinant et véhiculant l'information. Ces vecteurs numériques synthétisent, sous forme d'observations, la perception que nous avons du monde environnant. Ainsi, notre activité est déterminée par une multitude de systèmes mesurant des quantités telles que la température, la pression, la densité d'un milieu, ou tout autre quantité liée à un processus physique.

On parle de signal pour désigner un ensemble de mesures effectuées par un système à différentes dates temporelles ou positions spatiales, sur un processus réel. Un signal est donc un ensemble de mesures, chaque échantillon fournissant une information localisée sur un certain processus physique sous-jacent. Un système analogique délivre un signal défini continûment, alors qu'un système digital fournit un signal discret, défini sur un ensemble discret de positions spatio-temporelles.

Les signaux discrets permettent de synthétiser, stocker, manipuler et traiter efficacement l'information sonore, visuelle, et de bien d'autres modalités, extraite au quotidien du monde qui nous entoure. Plus que cela, les images numériques, qui immortalisent des instantanés visuels complexes, sont des vecteurs d'émotions extrêmement riches. L'explosion récente du marché de la photographie illustre bien le dynamisme du monde des images et signaux, et la nécessité de répondre aux besoins manifestes des applications est plus que jamais d'actualité. Le traitement du signal et des images est un domaine d'étude qui a connu un développement croissant, en particulier depuis l'avancée théorique majeure qu'a constituée, il y a déjà cinquante ans, la théorie de la communication de Shannon.

Dans ce manuscrit, nous nous intéresserons en particulier à la réduction et à l'agrandis-

sement des images numériques fixes. L'agrandissement est une opération courante, à la fois dans les applications grand public (conversion du format TV vers TVHD par exemple) et pour les besoins de certains professionnels, comme en infographie, imagerie satellitaire, où imagerie médicale. Il est souvent nécessaire de « zoomer » sur une partie de l'image pour en scruter le contenu et mieux en comprendre le sens.

Afin d'appréhender le problème de l'agrandissement, on doit se doter d'une méthodologie appropriée. Nous adoptons en effet le point de vue selon lequel il faut d'abord disposer d'un formalisme décrivant l'objet à traiter, avant d'en dériver des méthodes de traitement adéquates. Il est toujours plus satisfaisant de pouvoir justifier l'optimalité d'une approche en regard d'une théorie générique et parcimonieuse, plutôt que de proposer des méthodes *ad hoc*, qui ne sont validées *a posteriori* qu'au moyen de résultats empiriques. En agrandissement d'images, une vaste littérature propose des techniques basées sur des paradigmes très différents, allant du zoom fractal à l'interpolation directionnelle, en passant par l'extrapolation de coefficients d'ondelettes et l'utilisation d'équations aux dérivées partielles. C'est pourquoi nous traitons dans la première partie de cette thèse du problème générique de reconstruction de signaux, fondamental et sous-jacent à une multitude de problématiques et applications, et dont l'agrandissement d'images ne reflète qu'une des multiples facettes. Nous proposons des outils théoriques, ainsi que des méthodes pratiques mettant en œuvre efficacement les principes exposés. Nous nous attachons en effet à proposer des solutions réalisables et implémentables aisément, et s'exécutant avec un temps de calcul raisonnable.

Le problème de reconstruction dont nous traitons consiste à construire un signal défini continûment (une fonction) à partir d'un signal discret, pour modéliser ce dernier et en retenir toute l'information. On est inévitablement confronté à ce problème dès lors que l'on manipule des données numériques, d'une manière qui nécessite de l'information non disponible explicitement dans ces données. La reconstruction, qui établit une passerelle entre le monde des signaux discrets et celui des signaux continus, est fondamentale pour d'innombrables applications. Reconstruire un modèle défini continûment, et pouvant ensuite être évalué en de nouveaux points, est par exemple requis pour changer la fréquence d'échantillonnage d'un signal audio, ou pour effectuer une translation de pas non entier ou une rotation sur une image. L'obtention d'un modèle continu est utile voire nécessaire pour les problèmes de reconnaissance de formes, de vision par ordinateur, ou pour modéliser des données discrètes réelles acquises en imagerie de synthèse. Dans le contexte de l'animation graphique, la reconstruction sert par exemple pour interpoler une séquence d'images entre deux images clés (*inbetweening*) [27]. En conception assistée par ordinateur, la reconstruction sert à définir la forme de surfaces formant les objets en cours de conception (*curve fitting*) [189]. La reconstruction est aussi utilisée pour fusionner des données provenant de différents phénomènes, en imagerie médicale [27]. De plus, avoir un modèle continu des données permet le calcul de quantités différentielles, tâche incontournable pour la détection de contours et la segmentation dans les images. En ce qui concerne le redimensionnement d'images, nous verrons dans la seconde partie de cette thèse les liens forts qui existent entre

ce problème et celui de la reconstruction.

Avant d'aborder en détail les problèmes de reconstruction et de redimensionnement, nous dédions la suite de ce chapitre à l'exposé des notions requises pour leur traitement. Nous présentons d'abord la terminologie et les notations employées. Nous nous dotons ensuite d'un modèle pour les signaux que nous allons étudier, ce qui constitue une étape primordiale pour leur compréhension et la conception de méthodes adéquates. On ne peut en effet manipuler que ce que l'on connaît.

1.2 Préliminaires mathématiques et notations

Tout d'abord, nous notons en gras minuscule les quantités vectorielles. Par exemple $\mathbf{k} = (k_1, \dots, k_d) \in \mathbb{Z}^d$, pour un certain entier $d \geq 1$, est un d -uplet, aussi identifié à un vecteur colonne $\mathbf{k} = [k_1, \dots, k_d]^T \in \mathbb{Z}^{d \times 1}$. Dans le cas scalaire $d = 1$, on utilise des notations non grasses, comme $k \in \mathbb{Z}$. Les majuscules en gras indiquent des matrices, éventuellement de taille infinie, par exemple $\mathbf{A} = (A[k, l])_{k, l \in \mathbb{Z}}$. On indique par T la transposition matricielle.

Un *signal discret* (ou *signal* s'il n'y a pas d'ambiguïté) $u = (u[\mathbf{k}])_{\mathbf{k} \in \mathbb{Z}^d}$ est une suite d'éléments réels, où l'entier $d \geq 1$ est la *dimension* du signal. On ne considèrera que des signaux de taille infinie. Les signaux étant finis en pratique, on les étendra implicitement en des signaux infinis, au moyen de conditions de bords adéquates.

On parle de signaux *unidimensionnels* (ou « 1D ») si $d = 1$, d'*images discrètes* ou simplement *images* si $d = 2$. Chaque élément $u[\mathbf{k}]$ est un *échantillon* de u , mais dans le cas $d = 2$, on parle plus communément de *pixel* (de l'anglais *picture element*), et dans le cas $d = 3$, de *voxel* (*volume element*). Selon ce qui est le plus adéquat, on notera indifféremment $u[\mathbf{k}]$ ou $u[k_1, \dots, k_d]$.

En assimilant un filtre à sa réponse impulsionnelle, un *filtre* sera pour nous un signal discret. On parlera de filtre RIF (pour *à réponse impulsionnelle finie*) s'il possède un nombre fini de coefficients (échantillons) non nuls, de filtre RII sinon.

On définit le produit scalaire ℓ_2 de deux signaux discrets u, v par $\langle u, v \rangle_{\ell_2} = \sum_{\mathbf{k} \in \mathbb{Z}^d} u[\mathbf{k}]v[\mathbf{k}]^*$, où * indique le conjugué complexe (qui n'a pas d'importance pour des signaux réels). La norme ℓ_2 associée est $\|u\|_{\ell_2} = \sqrt{\sum_{\mathbf{k} \in \mathbb{Z}^d} |u[\mathbf{k}]|^2}$. On introduit l'ensemble $\ell_2(\mathbb{Z}^d)$ (ou simplement ℓ_2) des signaux discrets d'énergie finie : $\ell_2(\mathbb{Z}^d) = \{u \in \mathbb{R}^{\mathbb{Z}^d} \mid \|u\|_{\ell_2} < +\infty\}$. On notera $\langle \cdot, \cdot \rangle$ et $\|\cdot\|$ sans l'indice ℓ_2 s'il n'y a pas d'ambiguïté. De plus, on étend le produit scalaire ℓ_2 aux d -uplets et aux vecteurs colonnes, ce qui donne $\langle \mathbf{k}, \mathbf{l} \rangle = \mathbf{k}^T \mathbf{l}^*$.

On définit la *convolution* (discrète) ou *filtrage* (discret) comme l'opération qui à deux signaux discrets (dont l'un est parfois appelé filtre) $u, v \in \ell_2$ associe un nouveau signal discret, noté $u * v$, défini par $u * v[\mathbf{k}] = \sum_{\mathbf{l} \in \mathbb{Z}^d} u[\mathbf{l}]v[\mathbf{k} - \mathbf{l}]$.

Un *signal défini continûment* (ou *signal* s'il n'y a pas d'ambiguïté) est une fonction définie sur l'ensemble $\mathbb{R}^d$, à valeurs réelles. On notera indifféremment f , $f(\mathbf{x})$ ou $f(x_1, \dots, x_d)$ pour indiquer que f dépend de la variable continue $\mathbf{x} \in \mathbb{R}^d$ (qui peut être vue comme une variable spatiale, ou temporelle si $d = 1$).

On définit le produit scalaire L_2 de deux signaux f, g par $\langle f, g \rangle_{L_2} = \int_{\mathbb{R}^d} f(\mathbf{x})g(\mathbf{x})^* d\mathbf{x}$. La norme L_2 associée est $\|u\|_{L_2} = \sqrt{\int_{\mathbb{R}^d} |u(\mathbf{x})|^2 d\mathbf{x}}$. On omettra l'indice « L_2 » s'il n'y a pas d'ambiguïté. On introduit l'ensemble $L_2(\mathbb{R}^d)$ (ou simplement L_2) des signaux d'énergie finie : $L_2(\mathbb{R}^d) = \{f(\mathbf{x}) \mid \|f\|_{L_2} < +\infty\}$.

On notera $\|f\|_\infty$ la *norme infinie* d'un signal f borné : $\|f\|_\infty = \sup_{\mathbb{R}^d} |f(\mathbf{x})|$.

La *convolution* (continue) associe à deux signaux $f, g \in L_2$ le signal noté $f * g$ défini par $f * g(\mathbf{x}) = \int_{\mathbb{R}^d} f(\mathbf{y})g(\mathbf{x} - \mathbf{y})d\mathbf{y}$.

On notera δ l'élément neutre de la convolution.

On définit ainsi la *distribution de Dirac* $\delta(\mathbf{x})$, vue abusivement comme une fonction, par la propriété $\langle f, \delta \rangle = f(\mathbf{0})$ pour toute fonction $f \in L_2$. Il s'ensuit que $f * \delta = f$. De plus, en considérant la fonction f constante égale à 1, on obtient $\int_{\mathbb{R}^d} \delta(\mathbf{x})d\mathbf{x} = 1$. On a aussi la propriété $\delta(\mathbf{x}/a) = |a|^{-d}\delta(\mathbf{x})$.

On introduit aussi le signal discret ($\delta[\mathbf{k}]$) (à ne pas confondre avec la distribution de Dirac) défini par $\delta[\mathbf{0}] = 1$ et $\delta[\mathbf{k}] = 0$ pour tout $\mathbf{k} \neq \mathbf{0}$. On a donc $u * \delta = u$ pour tout signal discret u .

On note $\delta_{a,b}$ le *symbole de Kronecker* qui vaut 1 si $a = b$, 0 sinon. On peut donc écrire : $\delta[\mathbf{k}] = \delta_{\mathbf{k},\mathbf{0}} \forall \mathbf{k} \in \mathbb{Z}^d$.

On introduit l'opérateur $\bar{\cdot}$ tel que $\bar{f}(\mathbf{x}) = f(-\mathbf{x})$ et $\bar{h}[\mathbf{k}] = h[-\mathbf{k}]$. I désignera le nombre complexe racine de -1 .

La *transformée en $\mathbf{z}$* d'un signal discret u est la série formelle de la variable $\mathbf{z} = (z_1, \dots, z_d)$, que nous notons par une majuscule : $U(\mathbf{z}) = \sum_{\mathbf{k} \in \mathbb{Z}^d} u[\mathbf{k}]\mathbf{z}^{-\mathbf{k}}$, où $\mathbf{z}^{-\mathbf{k}}$ est un raccourci pour le scalaire $z_1^{-k_1} z_2^{-k_2} \dots z_d^{-k_d}$. La propriété majeure de la transformée en $\mathbf{z}$ est de transformer les convolutions en produits : si $w = u * v$, alors $W(\mathbf{z}) = U(\mathbf{z})V(\mathbf{z})$.

La *transformée de Fourier* d'un signal discret u est la fonction, que nous notons avec un chapeau, $\hat{u}(\boldsymbol{\omega})$, à valeurs complexes et 2π -périodique, définie par $\hat{u}(\boldsymbol{\omega}) = \sum_{\mathbf{k} \in \mathbb{Z}^d} u[\mathbf{k}]e^{-I\langle \boldsymbol{\omega}, \mathbf{k} \rangle}$. On remarque que $\hat{u}(\boldsymbol{\omega}) = U(e^{I\boldsymbol{\omega}})$.

On définit u^{-1} , l'*inverse* du signal discret u , comme le signal discret tel que $u * u^{-1} = \delta$. On en déduit : $U^{-1}(\mathbf{z}) = 1/U(\mathbf{z})$ et $\widehat{u^{-1}}(\boldsymbol{\omega}) = 1/\hat{u}(\boldsymbol{\omega})$.

Un filtre est dit *inverse* ou *récuratif* s'il est de la forme u^{-1} où u est un filtre RIF. Il est dit *rationnel* s'il est de la forme $u * v^{-1}$ où u et v sont RIF.

La *décimation*, ou *sous-échantillonnage*, de facteur entier a d'un signal discret u donne

le signal discret, noté $[u] \downarrow a$, tel que $[u] \downarrow a [\mathbf{k}] = u[\mathbf{k}]$ pour tout $\mathbf{k} \in \mathbb{Z}^d$. Si $v = [u] \downarrow a$, on a

$$\hat{v}(\boldsymbol{\omega}) = \sum_{\mathbf{k} \in \llbracket 1, a \rrbracket^d} \hat{u}\left(\frac{\boldsymbol{\omega} + 2\pi\mathbf{k}}{a}\right). \quad (1.1)$$

Le *sur-échantillonnage* de facteur entier a d'un signal discret u donne le signal discret $[u] \uparrow a$ tel que $[u] \uparrow a [\mathbf{k}] = u[\mathbf{k}]$ pour tout $\mathbf{k}$, et $[u] \uparrow a [\mathbf{k}] = 0$ partout ailleurs. Cette opération insère donc $a - 1$ zéros entre chaque échantillon de u , dans chaque direction. Si $v = [u] \uparrow a$, on a

$$V(\mathbf{z}) = U(\mathbf{z}^d) \xleftrightarrow{\mathcal{F}} \hat{v}(\boldsymbol{\omega}) = \hat{u}(a\boldsymbol{\omega}). \quad (1.2)$$

Notons que $\llbracket [u] \uparrow a \rrbracket \downarrow a = u$, et les *égalités nobles*

$$[u] \downarrow a * h = [u * [h] \uparrow a] \downarrow a, \quad (1.3)$$

$$[u * h] \uparrow a = [u] \uparrow a * [h] \uparrow a. \quad (1.4)$$

La *transformée de Fourier* de $f \in L_2(\mathbb{R}^d)$ est la fonction, notée avec un chapeau, $\hat{f}(\boldsymbol{\omega}) \in L_2(\mathbb{R}^d)$ à valeurs complexes, définie par $\hat{f}(\boldsymbol{\omega}) = \int_{\mathbb{R}^d} f(\mathbf{x}) e^{-I\langle \boldsymbol{\omega}, \mathbf{x} \rangle} d\mathbf{x}$. Par extension, on définit la transformée de Fourier de la distribution de Dirac par $\hat{\delta}(\boldsymbol{\omega}) = 1$ pour tout $\boldsymbol{\omega}$.

On rappelle la formule de Poisson caractérisant l'échantillonnage d'une fonction $f \in L_2(\mathbb{R}^d)$ en domaine de Fourier : si on construit le signal discret u tel que $u[\mathbf{k}] = f(\mathbf{k}) \forall \mathbf{k} \in \mathbb{Z}^d$, on a

$$\hat{u}(\boldsymbol{\omega}) = \sum_{\mathbf{k} \in \mathbb{Z}^d} f(\mathbf{k}) e^{-I\langle \boldsymbol{\omega}, \mathbf{k} \rangle} = \sum_{\mathbf{k} \in \mathbb{Z}^d} \hat{f}(\boldsymbol{\omega} + 2\pi\mathbf{k}). \quad (1.5)$$

Soit un réel $r \geq 0$. On définit l'espace de Sobolev W_2^r comme l'ensemble de fonctions

$$W_2^r = \left\{ f \in L_2(\mathbb{R}) \mid \int_{\mathbb{R}} (1 + \omega^2)^r |\hat{f}(\omega)|^2 d\omega < +\infty \right\}. \quad (1.6)$$

Par analogie avec cette définition de la régularité, on peut étendre la norme $\|f^{(r)}\|_{L_2}$ à des valeurs non entières de $r \geq 0$ en la posant égale à $\sqrt{\frac{1}{2\pi} \int |\omega|^{2r} |\hat{f}(\omega)|^2 d\omega}$. On peut ainsi caractériser la régularité d'une fonction f par la valeur maximale r telle que $f \in W_2^r$. La dérivée n -ième de s appartient alors à L_2 , pour tout entier $n \leq r$, et elle est de plus continue pour tout entier $n \leq r - \frac{1}{2}$.

En 1D, on définit a_f , l'autocorrélation discrète (séquence de Gram) de $f \in L_2(\mathbf{R})$ par :

$$a_f[k] = \bar{f} * f(k) \xleftrightarrow{\mathcal{F}} \hat{a}_f(\omega) = \sum_{k \in \mathbb{Z}} |\hat{f}(\omega + 2\pi k)|^2 \quad (1.7)$$

Si le signal $(\mu[\mathbf{k}])$ est la réalisation d'un processus aléatoire stationnaire discret de moyenne nulle, on peut le caractériser par son autocorrélation discrète c_μ , définie par

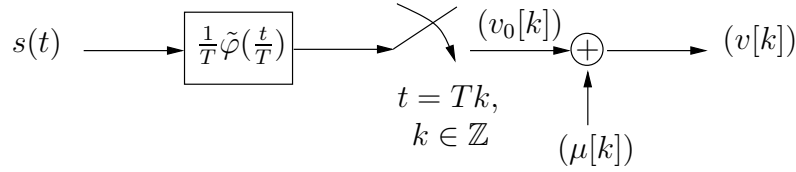


FIG. 1.1 : Modèle de formation des signaux discrets uniformes 1D.

$c_\mu[\mathbf{k}] = \mathcal{E}\{\mu[\mathbf{l}]\mu[\mathbf{k} + \mathbf{l}]\}$, où $\mathcal{E}\{\cdot\}$ est l'espérance mathématique, et qui ne dépend pas de $\mathbf{l} \in \mathbb{Z}^d$. $\hat{c}_\mu$ est alors la *densité spectrale de puissance* de μ .

De même, si $s(\mathbf{x})$ est la réalisation d'un processus aléatoire stationnaire à temps continu de moyenne nulle, on peut le caractériser par son autocorrélation $c_s(\mathbf{x}) = \mathcal{E}\{s(\mathbf{y})s(\mathbf{x} + \mathbf{y})\} \forall \mathbf{y} \in \mathbb{R}^d$ (ne dépendant pas de $\mathbf{y}$ pour un processus stationnaire).

On introduit les notations suivantes :

- $f(\mathbf{x}) = O(g(\mathbf{x}))$ indique que la fonction f est *dominée* par g , c.-à-d. $\limsup_{\|\mathbf{x}\| \rightarrow 0} |f(\mathbf{x})/g(\mathbf{x})| < \infty$.
- $f(\mathbf{x}) = o(g(\mathbf{x}))$ indique que f est *négligeable* devant g , c.-à-d. $|f(\mathbf{x})/g(\mathbf{x})| \rightarrow 0$.
- $f(\mathbf{x}) \sim g(\mathbf{x})$ indique que f est *équivalente* à g , c.-à-d. $\lim_{\|\mathbf{x}\| \rightarrow 0} f(\mathbf{x})/g(\mathbf{x}) = 1$.

On utilisera des fonctions indicatrices : $\mathbb{1}_\Omega(\mathbf{x}) = \{1 \text{ si } \mathbf{x} \in \Omega, 0 \text{ sinon}\}$, pour un certain domaine $\Omega \subset \mathbb{R}^d$.

On définit aussi la fonction de Riemann-Zeta

$$\zeta(x) = \sum_{n=1}^{\infty} \frac{1}{n^x} \quad \forall x > 1. \quad (1.8)$$

D'autre part, on utilisera la notation anglaise pour les coefficients binomiaux : pour tous $a, b \in \mathbb{N}$ tels que $a \geq b$,

$$\binom{a}{b} = \frac{a!}{b!(a-b)!}. \quad (1.9)$$

Les notations principales introduites dans cette section sont rappelées dans la table de notations.

1.3 Modèle pour les signaux discrets

Les signaux et images discrets résultent généralement de la discrétisation d'un processus physique, au moyen d'un dispositif d'acquisition. Le modèle de formation des signaux que nous adoptons tout au long de cette thèse, schématisé sur la figure 1.1, est le suivant : étant donné un signal discret uniforme 1D $v = (v[k])_{k \in \mathbb{Z}}$, nous postulons que l'on peut écrire :

$$\begin{aligned}
\forall k \in \mathbb{Z}, \quad v[k] &= \int_{\mathbb{R}} s(t) \bar{\varphi}\left(\frac{t}{T} - k\right) \frac{dt}{T} + \mu[k] \\
&= \left(s * \frac{1}{T} \bar{\varphi}\left(\frac{\cdot}{T}\right) \right)(Tk) + \mu[k] \\
&= \left\langle s, \frac{1}{T} \bar{\varphi}\left(\frac{\cdot}{T} - k\right) \right\rangle_{L_2} + \mu[k]
\end{aligned} \tag{1.10}$$

Les échantillons sont ainsi des *mesures*, effectuées aux positions uniformes $\{Tk, k \in \mathbb{Z}\}$, sur un processus $s(t)$ défini continûment, et corrompues par un bruit additif μ . On note v_0 la version non bruitée des mesures, c'est-à-dire $v[k] = v_0[k] + \mu[k]$ pour tout k . Nous donnons ici le modèle dans le cas 1D, afin de simplifier les notations, mais nous l'étendrons en 2D dans le chapitre 3, en généralisant la notion d'échantillonnage uniforme à l'échantillonnage sur un *treillis*. Le cas des signaux non uniformes sera traité dans le chapitre 4.

Détaillons les composants de ce modèle de formation :

- $s(t)$ est un processus réel sous-jacent inconnu. Par exemple, comme illustré sur la figure 1.2, dans le cas d'une photographie numérique, s est la scène lumineuse bidimensionnelle telle qu'elle se présente dans le plan de vue de l'appareil, autrement dit ce que voit l'humain qui prend la photo. Plus précisément, s est dans ce cas la puissance d'un flux radiant de photons, intégré sur une certaine plage de longueurs d'ondes, et vue comme une fonction définie dans le plan focal, exprimée initialement en Watts. Des notions de radiométrie sont détaillées dans [245]. Dans le cas général, s modélise un processus physique qui peut être de nature très variée. Dans le cas d'un signal sonore, par exemple, $s(t)$ est la pression de l'air à l'instant t , au niveau du capteur enregistrant le signal. Notons que l'origine $t = 0$ est fixée de manière arbitraire et non restrictive.

- Le réel $T > 0$ est le *pas d'échantillonnage* caractéristique de l'acquisition. On pourrait le fixer arbitrairement à 1, ce qui simplifierait les notations, mais on sera amené à le faire varier afin d'étudier certaines propriétés de convergence. On dit que le signal v est *uniforme*, car on peut voir l'échantillon $v[k]$ comme une mesure localisée à la position Tk . Ainsi il est utile d'assimiler le signal discret ($v[k]$) au peigne de Dirac pondéré

$$v_c(t) = \sum_{k \in \mathbb{Z}} v[k] \delta\left(\frac{t}{T} - k\right). \tag{1.11}$$

On remarque alors que $\widehat{v}_c(\omega) = \widehat{v}(T\omega)$.

- $\bar{\varphi}(t)$, que l'on définit indépendamment de T , modélise la réponse impulsionnelle du dispositif d'acquisition. Ainsi, formellement, s est d'abord dégradée par une convolution avec cette réponse impulsionnelle, puis échantillonnée de manière idéale sur la grille uniforme $(Tk)_{k \in \mathbb{Z}}$. Par exemple, dans le cas 2D d'une photographie numérique, chaque pixel $v[\mathbf{k}]$ est le résultat fourni par un élément du capteur CCD, qui intègre l'énergie lumineuse

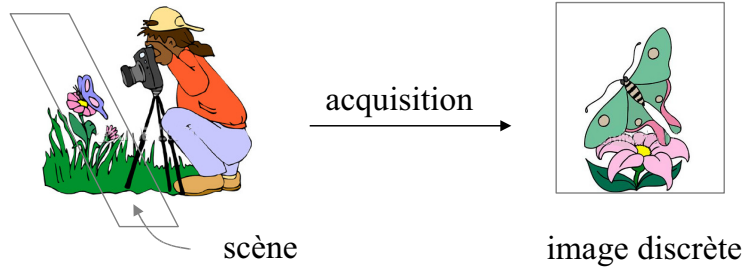


FIG. 1.2 : Une image naturelle discrète provient de l’acquisition d’une scène lumineuse. Les pixels de l’image peuvent donc être interprétés comme des mesures sur la scène.

qu’il reçoit sur sa surface, après que celle-ci ait été dégradée par l’optique. Dans le cas d’un capteur parfait qui effectue l’acquisition sur une grille carrée, et en l’absence de post-filtrage numérique, $\tilde{\varphi}$ est donc la convolution de l’indicatrice $\mathbb{1}_{[k-1/2, k+1/2]^2}$, et d’une fonction de forme gaussienne modélisant les effets introduits par l’optique [47, 243]. En règle générale, $\tilde{\varphi}$ pourra avoir une forme arbitraire. Dans la suite, on définit $\tilde{\varphi}_T(t) = \frac{1}{T}\tilde{\varphi}(\frac{t}{T})$ afin de simplifier les notations dans le modèle d’acquisition. Précisons que $\tilde{\varphi}$ est supposée connue. Son estimation à partir de v , requise dans les problèmes de déconvolution et restauration en aveugle, est réputée difficile, et s’effectue généralement à l’aide de méthodes heuristiques. Dans le cas des images naturelles, nous instancierons $\tilde{\varphi}$ d’une manière que nous jugeons vraisemblable, dans la section 5.2.1.

- Le signal μ est un bruit additif, indépendant de s . On supposera qu’il est la réalisation d’un processus aléatoire stationnaire de moyenne nulle, d’autocorrélation c_μ connue. Le bruit est souvent supposé gaussien [55]. En pratique, il y a toujours du bruit dans les signaux, mais il n’est pas nécessairement additif et indépendant. En pratique, on assimilera les effets de la quantification à un bruit additif indépendant, ce qui est inexact en toute rigueur.

Ce modèle de formation, sur lequel nous nous reposons, est classique, et a été proposé et justifié maintes fois [188, 18, 90, 121, 176, 43, 15]. Des informations supplémentaires sur la formation des images sont données dans [200, 236, 181]. Il s’agit d’un modèle réaliste, modélisant les moyens technologiques actuels : tout processus d’acquisition effectue une série de mesures sur un certain processus énergétique. Le modèle décrit dans [43] est légèrement plus complexe, car il inclut en plus une étape de déformation géométrique sur s , entre la convolution et l’échantillonnage. Bien que cela soit plus réaliste, par exemple pour modéliser la distorsion en barillet introduite par les optiques en photographie, cela rend le modèle non linéaire. Si la déformation n’est pas connue, il est extrêmement difficile de l’estimer, et si elle est connue, on peut la corriger en reformulant le problème à l’aide d’échantillons localisés sur un domaine non uniforme, comme dans la section 4.3.6. Précisons que notre modèle de formation s’applique pour les images et signaux *naturels*, incluant certains types d’images médicales. Il n’est pas représentatif des signaux synthétiques comme les images

de synthèse, ni même des signaux ayant subi des distorsions non linéaires importantes, comme une forte compression.

Certains auteurs considèrent un modèle d'échantillonnage idéal dans lequel $\tilde{\varphi}$ est omise. Cette situation, dans laquelle on dispose d'échantillons ponctuels de s , est un cas particulier de notre modèle obtenu en considérant $\tilde{\varphi}(x) = \delta(x)$. Pour la plupart des problèmes pratiques, ce n'est pas une hypothèse réaliste. D'une part, un capteur qui réalise un échantillonnage ponctuel exact est physiquement irréalisable, car la mesure nécessite une quantité non infinitésimale d'énergie, donc une intégration sur un intervalle de mesure non nulle. D'autre part, le principe d'Heisenberg implique l'impossibilité théorique d'une telle mesure. Cependant, on peut considérer que les $v[k]$ sont des valeurs ponctuelles d'une version dégradée de s , et non de s elle-même, à savoir de $\tilde{s} = s * \tilde{\varphi}_T$. Par exemple, dans le cas des images naturelles acquises par un appareil photo, cette scène dégradée est la scène sur laquelle on aurait fait porter toutes les dégradations du dispositif d'acquisition : flou de mise au point, diffraction optique, et réponse impulsionnelle d'une cellule intégratrice du capteur.

Insistons sur le principe de *mesure* qui est central dans le modèle que nous adoptons : un signal discret est vu comme une suite de mesures linéaires, éventuellement bruitées, sur un processus physique s dont le signal discret n'est qu'une représentation partielle. La complexité de s , en général, ne peut être capturée par le processus de discrétisation qui est donc un processus avec pertes. Bien qu'un signal discret soit vu comme une version appauvrie d'un processus plus riche, il faut cependant garder à l'esprit que seuls les objets discrets sont manipulables et stockables en pratique. L'intérêt principal quant à l'hypothèse d'existence de s est de pouvoir raisonner sur s au lieu de v , en transposant dans le domaine continu, où les outils d'analyse mathématique sont d'un grand secours, des problèmes mal posés dans le domaine discret. En effet, la plupart des problèmes initialement formulés à partir de v consistent en fait à estimer s ou des quantités sur s , comme nous le verrons par la suite. Remarquons que notre modèle n'est pas restrictif, en l'absence d'hypothèse sur s . Il affirme simplement que v n'a pas été créé *ex nihilo*, mais est une *représentation* d'un objet s plus riche, discrétisé par un processus linéaire. Cette abstraction que constitue la modélisation de l'étape de formation est nécessaire pour deux raisons. Elle permet d'une part d'exploiter pleinement la puissance d'expression des théories mathématiques, en faisant porter les manipulations sur des objets certes idéalisés, mais dont l'appréhension est propice à l'approfondissement de la compréhension que nous avons des processus réels. Et d'autre part, elle transcrit plus fidèlement notre conception naturelle des signaux et images, que nous considérons non comme des tableaux de chiffres, mais bien comme des représentations porteuses d'un sens physique.

Nous raisonnerons dans ce travail, hormis au chapitre 4, sur des signaux de taille infinie. Les signaux finis seront étendus implicitement en signaux infinis au moyen de conditions de bords de type symétrie miroir [45]. Ainsi, le signal fini $(v[n])_{n \in \llbracket 0, N-1 \rrbracket}$ sera étendu en un signal infini en posant $u[n] = u[-n]$ si $n < 0$ et $u[n] = u[2N - 2 - n]$ si $n \geq N$.

1.3.1 Considérations sur la nature du processus s

Quelle est la nature de $s(\mathbf{x})$? Bien sûr, cela dépend du contexte, car les processus sous-jacents aux signaux peuvent être de nature très variée. Un signal sonore et une image proviennent certes de processus énergétiques ondulatoires, mais qui n'ont guère de rapport entre eux. Voyons de plus près le cas des scènes lumineuses sous-jacentes aux images naturelles.

De nombreux auteurs ont proposé de modéliser l'ensemble des scènes lumineuses comme un espace mathématique de fonctions. L'espace $BV(\mathbb{R}^2)$ des fonctions à variation bornée est souvent employé [190, 104]. Cela peut sembler correct à première vue, car cet espace contient des fonctions lisses ayant des discontinuités le long de courbes lisses, correspondant aux bords des objets constituant la scène. Cependant, la validité de ce modèle a été contredite en 1999 [12], car $s(\mathbf{x})$ peut localement être très oscillante, par exemple pour certaines textures, ce qui rend l'espace BV impropre à la modélisation des images. Plus généralement, définir un espace des fonctions qui représentent, ou non, des scènes lumineuses, et ce de manière tranchée, semble un objectif hasardeux.

Un paradigme plus souple que le cadre déterministe est l'interprétation stochastique, dans laquelle $s(\mathbf{x})$ est vue comme la réalisation d'un processus aléatoire ayant des caractéristiques connues. Les champs de Markov ont été les premiers modèles aléatoires proposés pour les images, appliqués en conjugaison avec un formalisme d'estimation bayésienne pour des problèmes comme la restauration, la segmentation, ou la reconstruction tomographique [77, 103, 60]. Des extensions multi-échelles plus complexes ont été développées afin de capturer les relations structurelles existant entre les échelles [44, 78, 202]. Les modèles à base de champs de Markov cachés multi-échelles ont montré leur efficacité en traitement d'images [170]. Ils permettent de prendre en compte des dépendances à longue distance, comme celles qui se manifestent dans les processus aléatoires en $1/f$ (c'est-à-dire dont le spectre de puissance se comporte en $1/|\omega|^\gamma$ pour un certain $\gamma > 0$). Or le comportement spectral de type $1/f$ des images naturelles a été avéré à maintes reprises [172, 232, 164, 214]. Les processus de Matérn [115] et les mouvements browniens fractionnaires semblent de bons modèles stochastiques pour $s(\mathbf{x})$ [105, 214, 183]. Cependant, on peut objecter aux modèles aléatoires de ne pas refléter la nature fortement non stationnaire des scènes lumineuses, constituées d'objets variés et complexes. La notion de géométrie et de contours n'est pas prise en compte dans les modèles aléatoires.

Une image est avant tout, pour un observateur humain, composée de zones lisses et de zones texturées. Celles-ci sont séparées par des discontinuités plus ou moins nettes et régulières, correspondant aux contours des éléments de la scène [186, 54, 25]. De nombreuses études psychophysiques ont confirmé cette théorie (dite du *raw primal sketch*) de Marr [151], selon laquelle l'information pertinente (*sketch data*) consiste en la structure géométrique des discontinuités et en les amplitudes des pixels de contours [122][53]. De-

vant la nature fortement non linéaire de la perception visuelle humaine, il est donc difficile de définir mathématiquement des espaces homogènes d'images, ou d'établir des propriétés statistiques vérifiées par les images qu'un humain « de référence » (si tant est qu'il en existe) qualifierait de naturelles. La variété des modèles disponibles permet cependant d'aborder avec des angles différents et complémentaires cette famille d'objets aux propriétés très riches que sont les images naturelles.

Dans cette thèse, on ne fera pas d'hypothèse *a priori* sur la nature de s . On envisagera deux situations :

- la situation *déterministe*, dans laquelle $s \in L_2$. Cette hypothèse de travail ne correspond pas à une réalité physique, mais elle n'est pas restrictive, car les traitements étant essentiellement locaux, le comportement de s à l'infini n'a pas d'influence. Afin que notre modèle d'acquisition soit bien posé, on supposera que s vérifie de plus la contrainte suivante :

$$\int_{\mathbb{R}} |\hat{\varphi}(\omega)|^2 |\hat{s}(\omega)|^2 (1 + |\omega|^2) d\omega < +\infty. \quad (1.12)$$

Cette condition faible implique que $s * \tilde{\varphi}$ est dérivable une fois au sens L_2 , et donc qu'elle est continue. Ainsi, le signal non bruité v_0 appartient à ℓ_2 (voir [36, Appendix C]). Par souci de simplicité, nous noterons abusivement $s \in L_2$ dans tout le manuscrit, tout en nous restreignant implicitement aux fonctions vérifiant (1.12).

- la situation *stochastique*, où s est la réalisation d'un processus aléatoire stationnaire de moyenne nulle, d'autocorrélation $c_s(t)$. On suppose là aussi vérifiée la condition (1.12), où $|\hat{s}(\omega)|^2$ est remplacé par $\hat{c}_s(\omega)$.

1.4 Plan de la thèse

Cette thèse est structurée en deux parties.

Nous abordons, dans la première partie de ce manuscrit, la reconstruction de signaux unidimensionnels et bidimensionnels, à valeurs scalaires, à partir de données uniformes et non uniformes. En exploitant des notions de théorie de l'échantillonnage et de théorie de l'approximation, qui fournissent de puissants outils pour décrire, analyser et synthétiser les signaux, nous établissons des passerelles entre les mondes discret et continu, et nous les mettons à profit pour proposer des nouvelles solutions. Dans le chapitre 2, nous nous intéressons à la reconstruction d'un signal uni-dimensionnel à partir de données uniformes. Cette étude sera étendue au cas des signaux (images) 2D dans le chapitre 3, en introduisant la notion de treillis. La reconstruction à partir de données 1D non uniformes est traitée dans le chapitre 4.

Dans la seconde partie de ce document, nous nous intéressons au thème du redimensionnement d'images, qui recouvre les deux problèmes de réduction et d'agrandissement.

Les concepts mis en œuvre pour la reconstruction trouvent dans le redimensionnement un terrain d'application propice. En effet, redimensionner nécessite d'analyser l'information contenue dans une image, pour la synthétiser non plus continûment, mais de manière discrète à une résolution différente. Dans le chapitre 5, nous présentons un modèle formel original pour le redimensionnement, en cohérence avec le modèle de formation des images que nous avons décrit plus haut. Nous dérivons ensuite de ce modèle des méthodes adéquates et performantes pour la réduction d'images.

Dans le chapitre 6, nous nous focalisons sur le problème plus complexe de l'agrandissement. Nous présentons les méthodes existantes, mises en perspective de notre modèle de redimensionnement, ainsi que de nouveaux procédés, basés là-aussi sur une formalisation précise du problème.

Enfin, nous présentons dans le chapitre 7 l'agrandissement par induction, qui fournit une réponse originale au problème de l'agrandissement. Nous insisterons sur la mise en œuvre efficace de cette méthode, et nous illustrerons la qualité des résultats obtenus à l'aide de comparaisons avec les méthodes existantes.

Première partie

Reconstruction de signaux et images discrets

Chapitre 2

Reconstruction 1D à partir de données uniformes

2.1 Introduction

LE problème de reconstruction s'énonce simplement à partir du modèle de formation des signaux que nous avons adopté à la section 1.3 : il a pour but d'estimer le signal défini continûment $s(t)$, étant données les mesures discrètes $v[k]$. Il s'agit d'un problème générique impliquant à la fois une opération de débruitage, et de déconvolution de la réponse $\tilde{\varphi}$, à la base de nombreuses applications en communications, traitement de la parole et des images. La reconstruction est un problème d'estimation linéaire, vaste domaine reposant sur les travaux classiques de Kolmogorov [134] et Wiener [242]. Notons que la reconstruction telle que nous la définissons est parfois désignée en anglais par le terme *restoration* [43].

Afin de formaliser ce problème, nous introduisons les notations suivantes :

- L'opérateur de *discrétisation* $\mathcal{D} : s \mapsto v$ modélise l'acquisition des mesures $v[k]$ à partir de s .
- L'opérateur de *reconstruction* $\mathcal{M} : v \mapsto f$, qu'il s'agit de concevoir, reconstruit une estimée f de s à partir de v .
- L'opérateur d'*approximation* $\mathcal{Q} = \mathcal{M} \circ \mathcal{D}$ associe à s son estimée à partir de v .

Dans ce chapitre, nous allons nous intéresser à la reconstruction uni-dimensionnelle à partir du signal discret $v = (v[k])_{k \in \mathbb{Z}}$ constituant nos seules données initiales, et supposé provenir d'un processus $s(t)$, comme formalisé par le modèle (1.10). La reconstruction a pour finalité d'obtenir une fonction $f_T = \mathcal{M}v$ qui soit proche de s . Nous adoptons la notation « f_T » afin d'insister sur le fait que la reconstruction dépend du pas d'échantillonnage T .

Nous allons tout d'abord présenter des méthodes de reconstruction qui sont optimales lorsque des connaissances *a priori* sur s sont connues et exploitées. Nous nous placerons ensuite dans le cas général où de telles connaissances ne sont pas disponibles, et on fixera

alors l'espace vectoriel dans lequel chercher la fonction reconstruite.

2.2 Approches optimales

2.2.1 Approche variationnelle

Plaçons-nous dans le cadre déterministe $s \in L_2$. L'approche variationnelle consiste à poser le problème de la reconstruction comme un problème d'optimisation, visant à minimiser un critère qui impose à la solution $f_T(t)$ d'être à la fois proche des données, et suffisamment lisse pour contrebalancer les effets du bruit.

Le critère servant à quantifier la proximité d'une fonction reconstruite aux données $v[k]$ est le critère des moindres carrés :

$$\varepsilon_{\ell_2}(g) = \sum_{k \in \mathbb{Z}} |g * \tilde{\varphi}_T(Tk) - v[k]|^2. \quad (2.1)$$

Selon le contexte dans lequel s'opère la reconstruction, on peut envisager les cas de figure suivants :

- On sait que la fonction s appartient à l'espace $\Omega = \{g \mid \|\mathcal{L}g\| \leq C\}$ où $\mathcal{L} \cdot$ est un opérateur linéaire et invariant par translation. Typiquement, le critère quadratique $\|\mathcal{L}g\|^2$ quantifie le manque de lissitude de la fonction g ; l'exemple le plus courant est l'opérateur de dérivation : $\mathcal{L}g = d^n g / dt^n$. On cherche alors la fonction f_T appartenant à Ω , et minimisant le critère des moindres carrés $\varepsilon_{\ell_2}(g)$. Comme ce dernier est généralement antagoniste de $\|\mathcal{L}g\|$, une fonction sera d'autant plus lisse qu'elle sera éloignée des données, et vice-versa. D'après les conditions de Karush-Kuhn-Tucker [29], le problème se ramène à minimiser $\varepsilon_{\ell_2}(g)$ sous contrainte que $\|\mathcal{L}g\| = C$, ou minimiser $\|\mathcal{L}g\|$ sous contrainte que $\varepsilon_{\ell_2}(g) = 0$.

- On sait que la fonction s est relativement lisse au sens du critère $\mathcal{L}$, mais on ne sait pas dans quelle mesure. Dans ce cas, on va chercher la solution la plus vraisemblable par rapport aux données, vérifiant donc $\varepsilon_{\ell_2}(g) = C$ (pour une certaine valeur C déterminée en fonction de la variance du bruit) et minimisant $\|\mathcal{L}g\|$.

Ces deux problèmes sont en fait équivalents à chercher la fonction f_T minimisant le critère des moindres carrés régularisés (au sens de Tikhonov-Phillips [213, 145]) :

$$\varepsilon_{\text{VAR}}(g) = \sum_{k \in \mathbb{Z}} |g * \tilde{\varphi}_T(Tk) - v[k]|^2 + \lambda \|\mathcal{L}g\|^2 \quad (2.2)$$

où $\lambda > 0$ est un paramètre lagrangien à choisir en fonction de C , afin de rendre ce problème équivalent au problème de minimisation sous contrainte souhaité. Cette formulation variationnelle revient à chercher un compromis entre l'attache aux données au sens des moindres carrés d'une part, et la volonté d'avoir une solution lisse, d'énergie minimale au sens $\|\mathcal{L} \cdot\|$, d'autre part. Cette approche, où deux termes antagonistes sont mis en concurrence, est classique dans les problèmes de modélisation. Par exemple, il y a un lien direct

entre les deux termes ici présents et les notions de *forces externes* (la solution est attirée par les données) et *forces internes* (liées à une fonctionnelle énergétique sur la solution) employées dans la théorie des *snakes* ou *contours actifs*, avec de nombreuses applications en modélisation graphique [120].

Afin d'explicitier la solution de ce problème variationnel, on suppose que l'opérateur $\mathcal{L}$ est *spline-admissible* (voir les conditions données dans [225]) et donc qu'il existe une fonction $\varphi(t) \in L_2$ vérifiant la condition :

$$\mathcal{L}^* \mathcal{L} \varphi(t) = \sum_{k \in \mathbb{Z}} h[k] \tilde{\varphi}(t - k) \quad (2.3)$$

pour un certain signal discret h , où $\mathcal{L}^*$ est l'opérateur adjoint de $\mathcal{L}$ défini par la relation $\langle \mathcal{L}^* f, g \rangle = \langle f, \mathcal{L} g \rangle \forall f, g$. On suppose aussi que les translatés $\{\varphi(t - k), k \in \mathbb{Z}\}$ forment une base de Riesz, ce qui équivaut à dire qu'il existe deux constantes C_1 et C_2 (les bornes de Riesz) telles que $0 < C_1 \leq \hat{a}_\varphi(\omega) \leq C_2 < +\infty$ presque partout.

Alors la fonction f_T minimisant le critère ε_{VAR} dans L_2 s'écrit [225, 182] :

$$f_T(t) = \sum_{k \in \mathbb{Z}} c[k] \varphi\left(\frac{t}{T} - k\right). \quad (2.4)$$

On dit alors que f_T et φ sont des $\mathcal{L}^* \mathcal{L}$ -splines, c'est-à-dire des fonctions que l'on peut écrire sous la forme (2.3). Les coefficients $c[k]$ s'obtiennent par filtrage à partir des données $v[k]$: $c = v * p$, à l'aide du préfiltre p défini par

$$\hat{p}(\omega) = \frac{1}{\sum_{k \in \mathbb{Z}} \hat{\varphi}(\omega + 2k\pi) \tilde{\hat{\varphi}}(\omega + 2k\pi) + \lambda \hat{h}(\omega)} \quad (2.5)$$

Dans le cas non bruité, on posera $\lambda = 0$, et la fonction f_T est alors la solution du problème de minimisation de $\|\mathcal{L} \cdot\|$ sous contrainte de consistance $\varepsilon_{\ell_2}(g) = 0 \Leftrightarrow \mathcal{D}g = v$. L'opérateur de reconstruction $\mathcal{M}$ est, dans ce cas, l'inverse généralisé (inverse de Moore-Penrose) de $\mathcal{D}$, en munissant L_2 de la norme $\|\mathcal{L} \cdot\|$.

Considérons l'exemple $\tilde{\varphi} = \delta$ et $\mathcal{L} = d^n/dt^n$, pour un certain entier $n \geq 1$. On a ainsi $\mathcal{L}^* \mathcal{L} = (-1)^n d^{2n}/dt^{2n}$. La solution f_T est alors une *spline polynomiale* de degré impair $2n-1$, avec ses nœuds localisés aux positions $(Tk)_{k \in \mathbb{Z}}$, c'est-à-dire une fonction de régularité $\mathcal{C}^{2n-2}$ polynomiale de degré $2n-1$ sur chaque segment $[Tk, T(k+1)]$ [84, 218]. L'espace de ces splines admet comme générateur $\varphi = \beta^{2n-1}$, la B-spline centrée de degré $2n-1$ définie comme suit :

$$\beta^0(t) = \begin{cases} 1 & \text{si } t \in]-\frac{1}{2}, \frac{1}{2}[\\ \frac{1}{2} & \text{si } |t| = \frac{1}{2} \\ 0 & \text{sinon} \end{cases} \quad \xleftrightarrow{\mathcal{F}} \quad \hat{\beta}^0(\omega) = \text{sinc}(\omega/2), \quad (2.6)$$

et pour tout $n \geq 1$,

$$\beta^n = \beta^{n-1} * \beta^0 \quad \xleftrightarrow{\mathcal{F}} \quad \hat{\beta}^n(\omega) = \text{sinc}(\omega/2)^{n+1}. \quad (2.7)$$

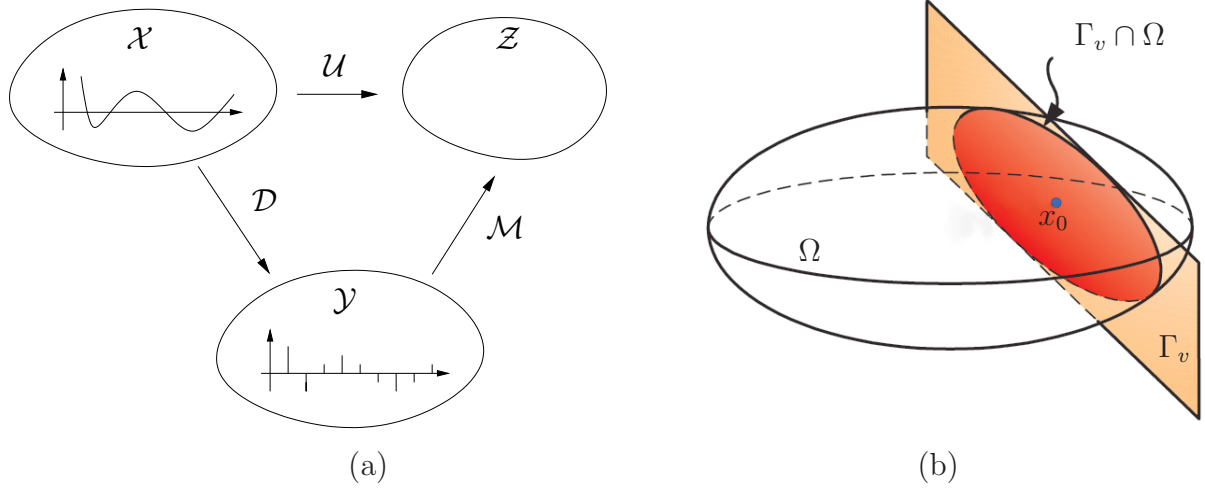


FIG. 2.1 : Estimation selon la théorie de l'*optimal recovery*. (a) : principe général. (b) : La solution proposée est l'image par $\mathcal{U}$ de $x_0 \in \mathcal{X}$, le centre de Chebyshev de $\Gamma_v \cap \Omega$, l'ensemble dans lequel on sait que se trouve l'objet inconnu $s \in \mathcal{X}$.

La propriété (2.3) est bien vérifiée : $(-1)^n d^{2n} \beta^{2n-1} / dt^{2n} = \sum h[k] \delta(t - Tk)$ avec $H(z) = (-z + 2 - z^{-1})^n$.

Si $\mathcal{L}$ est l'identité, ce qui peut être le choix adopté si l'on n'a pas de connaissance *a priori* sur le terme de régularisation approprié, il suffit de choisir $\varphi = \tilde{\varphi}$, qui vérifie bien (2.3) avec $h = a_{\tilde{\varphi}}^{-1}$.

L'approche variationnelle a été étendue dans la littérature à des critères de régularisation non quadratiques, comme la variation totale [191, 10] ou des critères semi-quadratiques [59]. Dans ce cas, le processus de reconstruction $\mathcal{M}$ n'est plus linéaire, et la solution est déterminée à l'aide d'une méthode itérative.

2.2.2 Approche minimax

La théorie de l'*optimal recovery* permet de formaliser de nombreux problèmes d'estimation [157, 201, 165]. Appliquons-la au problème de la reconstruction à partir d'un signal non bruité. Soit $\mathcal{U}$ un opérateur linéaire défini sur un espace vectoriel $\mathcal{X}$ et à valeurs dans un espace vectoriel normé $\mathcal{Z}$. La quantité à estimer est $\mathcal{U}s$, pour un objet inconnu $s \in \mathcal{X}$ dont on ne connaît que $v = \mathcal{D}s$, pour un autre opérateur linéaire $\mathcal{D} : \mathcal{X} \rightarrow \mathcal{Y}$. Afin d'estimer $\mathcal{U}s$, on applique un opérateur $\mathcal{M} : \mathcal{Y} \rightarrow \mathcal{Z}$ à v . Cette situation est schématisée sur la figure 2.1 (a).

En ce qui concerne le problème de reconstruction, on a $\mathcal{X} = \mathcal{Z} = L_2$, et $\mathcal{U}$ est l'identité, car c'est la fonction $s(t)$ elle-même que l'on veut estimer. $\mathcal{D} : \mathcal{X} \rightarrow \mathcal{Y} = \ell_2$ est l'opérateur de discrétisation qui à s associe v , suivant le modèle (1.10) (on suppose qu'il n'y a pas de bruit). $\mathcal{M}$ est l'opérateur de reconstruction à déterminer.

On suppose que $s \in \Omega = \{g \mid \|g\|_{\mathcal{X}} \leq C\}$ où $\|\cdot\|_{\mathcal{X}}$ est une norme qui n'est pas

nécessairement la norme L_2 , et qui modélise la connaissance *a priori* que l'on a sur s . Par exemple, $\|\cdot\|_{\mathcal{X}} = \|\mathcal{L} \cdot\|_{L_2}$ pour un certain opérateur linéaire invariant par translation $\mathcal{L}$.

La donnée v contraint l'ensemble d'instances possibles pour s . En effet, on sait que $s \in \Gamma_v$, que nous appellerons l'**ensemble consistant** :

$$\boxed{\Gamma_v = \{g \in L_2 \mid \mathcal{D}g = v\}}. \quad (2.8)$$

Ainsi, $s \in \Omega \cap \Gamma_v$. Michelli et Rivlin [157] proposent alors de définir $\mathcal{M}v$ comme le centre de Chebyshev z_0 de $\mathcal{U}(\Omega \cap \Gamma_v)$, défini par

$$z_0 = \operatorname{argmin}_{z \in \mathcal{Z}} \max_{w \in \mathcal{U}(\Omega \cap \Gamma_v)} \|z - w\|_{\mathcal{Z}}^2. \quad (2.9)$$

Ainsi, l'objet estimé $\mathcal{M}v$ minimise l'écart au pire cas au véritable objet $\mathcal{U}s$. On montre alors (sous certaines conditions vérifiées dans le cas de la reconstruction) que z_0 est l'image par $\mathcal{U}$ du centre de Chebyshev x_0 de $\Omega \cap \Gamma_v$ [201] :

$$z_0 = \mathcal{U}(x_0) = \mathcal{U}\left(\operatorname{argmin}_{f \in \mathcal{X}} \max_{g \in \Omega \cap \Lambda_v} \|f - g\|_{\mathcal{X}}^2\right). \quad (2.10)$$

De plus, il est connu, d'après les travaux classiques de Golomb et Weinberger [106], que x_0 est l'objet de $\mathcal{X}$ vérifiant $\mathcal{D}x_0 = v$, et de norme $\|\cdot\|_{\mathcal{X}}$ minimale. Nous avons représenté l'élément x_0 sur le schéma 2.1 (b), inspiré de [165]. L'ensemble Γ_v est un hyperplan affine : on peut l'écrire $\Gamma_v = x_0 + \Gamma_0$, où Γ_0 est le noyau de $\mathcal{D}$. Ω est un hyperellipsoïde. L'intersection $\Omega \cap \Gamma_v$ est donc un hyperellipsoïde affine.

En règle générale, on peut expliciter z_0 au moyen des représentants de Riesz des opérateurs $\mathcal{D}$ et $\mathcal{U}$ [201, 165]. Dans notre cas, la fonction reconstruite f_T est simplement la solution du problème de minimisation de $\|\cdot\|_{\mathcal{X}}$ sous contrainte que $\mathcal{D}f_T = v$. Le formalisme adopté ici, cherchant à minimiser la distance au pire cas entre l'objet inconnu et l'objet à estimer, aboutit donc au même résultat que l'approche variationnelle précédente, si $\|\cdot\|_{\mathcal{X}} = \|\mathcal{L} \cdot\|_{L_2}$: f_T est la $\mathcal{L}^*\mathcal{L}$ -spline consistante avec les données, c'est-à-dire obtenue avec (2.4) et (2.5) en posant $\lambda = 0$. Notons que la solution ne dépend pas de la constante C dans la définition de Ω . Ce lien entre l'*optimal recovery* et la théorie variationnelle est discuté plus en détail dans [195, 196] [1, p. 7].

La fonction f_T est donc optimale du point de vue du critère

$$\varepsilon_{\text{MM}}(g) = \max_{f \in \Omega \cap \Gamma_v} \|f - g\|_{\mathcal{Z}}^2, \quad (2.11)$$

où $\|\cdot\|_{\mathcal{Z}}$ peut être une pseudo-norme arbitraire sur L_2 : par exemple, $\|g\|_{\mathcal{Z}} = \|g\|_{L_2}$, $\|g\|_{L_\infty}$, ou $|g(t_0)|$ pour un $t_0 \in \mathbb{R}$ quelconque. Ainsi, cette solution minimax minimise l'erreur au pire cas entre f_T et s , en utilisant simplement le fait que $s \in \Omega \cap \Gamma_v$.

2.2.3 Approche stochastique

On suppose maintenant que s est la réalisation d'un processus aléatoire stationnaire à temps continu de moyenne nulle, de densité spectrale de puissance $\hat{c}_s(\omega) > 0$. On cherche

à minimiser l'erreur quadratique ponctuelle moyenne à un instant $t_0 \in \mathbb{R}$,

$$\varepsilon_{\text{STO},t_0}(g) = \mathcal{E} \{ |s(t_0) - g(t_0)|^2 \}. \quad (2.12)$$

Comme montré dans [182], la solution f_T , qui minimise $\varepsilon_{\text{STO},t_0}$ pour tout t_0 , s'écrit sous la forme (2.4), avec cette fois

$$\varphi = c_s(T \cdot) * \bar{\varphi} \xleftrightarrow{\mathcal{F}} \hat{\varphi}(\omega) = \frac{1}{T} \hat{c}_s\left(\frac{\omega}{T}\right) \hat{\varphi}(\omega)^*. \quad (2.13)$$

Les coefficients $c[k]$ dans l'expression de f_T sont là encore obtenus par filtrage, avec le préfiltre p défini par

$$\hat{p}(\omega) = \frac{1}{\sum_{k \in \mathbb{Z}} \hat{\varphi}(\omega + 2k\pi) \hat{\varphi}(\omega + 2k\pi) + \hat{c}_\mu(\omega)}. \quad (2.14)$$

Il est intéressant de constater que les solutions variationnelle et stochastique sont équivalentes dans le cas où μ est un bruit blanc de variance σ^2 , $\mathcal{L}$ est l'opérateur de blanchiment de s , c'est-à-dire

$$\mathcal{L}^* \mathcal{L}_{c_s}(Tt) = \sigma_0^2 \delta(t), \quad (2.15)$$

$h[k] = \delta[k]$ et $\lambda = \sigma^2/\sigma_0^2$. Cela peut servir d'heuristique pour le choix de $\mathcal{L}$ et λ dans le cas déterministe.

Dans le cas 2D où s est un processus de Matérn, la forme explicite de φ et de la solution f_T ont été étudiées dans [183]. Mentionnons aussi une extension à la reconstruction de mouvements browniens fractionnaires, qui ne sont pas des processus stationnaires, à l'aide de splines polyharmoniques, dans [214].

Si $s(t)$ est un processus aléatoire gaussien TK -périodique et que l'on dispose d'un nombre fini d'observations $v[1], \dots, v[K]$, on peut interpréter la solution f_T comme l'estimée optimale de s d'un point de vue bayésien (solution du maximum *a posteriori*), en voyant (2.2) comme une log-vraisemblance [225]. Plus généralement, l'approche variationnelle peut être ré-interprétée dans le cadre bayésien en associant une densité de probabilité à la fonctionnelle de régularisation, au moyen d'une distribution de Gibbs appropriée.

2.3 Reconstruction au sens des moindres carrés

Nous avons vu dans la section précédente que l'utilisation d'une connaissance *a priori* sur le signal à reconstruire permet de proposer des solutions optimales du point de vue d'un certain critère exploitant cette information. On se place maintenant dans le cas où les mesures $v[k]$ représentent la seule information disponible sur s .

2.3.1 Interpolation

Intéressons-nous d'abord au problème d'interpolation, correspondant au cas le plus simple où l'on dispose d'échantillons ponctuels non bruités $v[k] = s(Tk) \forall k$. Cela revient à poser $\tilde{\varphi} = \delta$ dans notre modèle.

On parle d'*interpolation* lorsque l'on reconstruit une fonction $f_T(t)$ passant par des échantillons connus $v[k] : f_T(Tk) = v[k] \forall k \in \mathbb{Z}$. Cette méthode de reconstruction semble naturelle : puisque les valeurs $s(Tk)$ sont connues, il est judicieux de chercher une fonction ayant les mêmes valeurs. Ainsi la reconstruction est parfaite aux positions $Tk : s(Tk) = f_T(Tk) \forall k$.

Retraçons brièvement l'histoire de l'interpolation, en nous inspirant d'une étude très complète de Meijering [152]. Le mot *interpolation* vient du latin *interpolare*, contraction de *inter* (entre) et *polare* (polire). Ce mot semble avoir été utilisé pour la première fois par Wallis en 1655 dans son livre sur l'arithmétique infinitésimale [238]. Les polynômes furent les premières fonctions utilisées pour définir les fonctions interpolantes. Newton décrivit le premier le polynôme $p(t)$ passant par un nombre fini de valeurs $p(t_1), \dots, p(t_N)$. La formule, due à Waring en 1779, bien qu'attribuée généralement à Lagrange, est la suivante :

$$p(t) = \sum_{i=1..N} p(t_i) \prod_{j \neq i} \frac{t - t_j}{t_i - t_j}. \quad (2.16)$$

Bien que les polynômes n'aient été utilisés au départ que pour leur simplicité de manipulation, le théorème de Weierstrass, en 1885, leur donna leurs lettres de noblesse, en stipulant que toute fonction continue sur un intervalle fermé peut être approchée uniformément par un polynôme à une précision arbitraire. Notons que le polynôme interpolant de Lagrange ne vérifie pas nécessairement cette convergence uniforme : on peut par exemple trouver une fonction $s(t)$ continue sur l'intervalle $[0, 1]$ telle que le polynôme de Lagrange interpolant la fonction aux points $0, 1/N, 2/N \dots 1$ diverge lorsque $N \rightarrow +\infty$. Ce résultat montra l'intérêt limité des polynômes algébriques pour l'interpolation.

Whittaker fit un pas de plus dans son célèbre article de 1915 [241], en s'intéressant à la fonction limite obtenue par l'interpolation de Lagrange aux positions $t_k = Tk, k = -N \dots N$, lorsque N tend vers l'infini. Il montra que cette fonction interpolante $f_T(t)$ peut s'écrire

$$f_T(t) = \sum_{k \in \mathbb{Z}} v[k] \operatorname{sinc}\left(\pi\left(\frac{t}{T} - k\right)\right) \quad \forall t \in \mathbb{R}, \quad (2.17)$$

à l'aide de la fonction *sinus cardinal*

$$\operatorname{sinc}(\pi t) = \frac{\sin(\pi t)}{\pi t} \xleftrightarrow{\mathcal{F}} \widehat{\operatorname{sinc}}(\omega) = \pi \mathbb{1}_{]-1,1[}(\omega). \quad (2.18)$$

Shannon comprit le premier toute l'importance de cette fonction pour la théorie de l'échantillonnage [198, 199]. Son désormais célèbre *théorème de l'échantillonnage* établit que toute fonction $s(t)$ à bande limitée (c'est-à-dire $\hat{s}(\omega) = 0 \forall |\omega| \geq \pi/T$) peut être retrouvée exactement à partir de ses échantillons $v[k] = s(Tk)$, par interpolation avec le sinus cardinal, c'est-à-dire que $f_T = s$ dans (2.17).

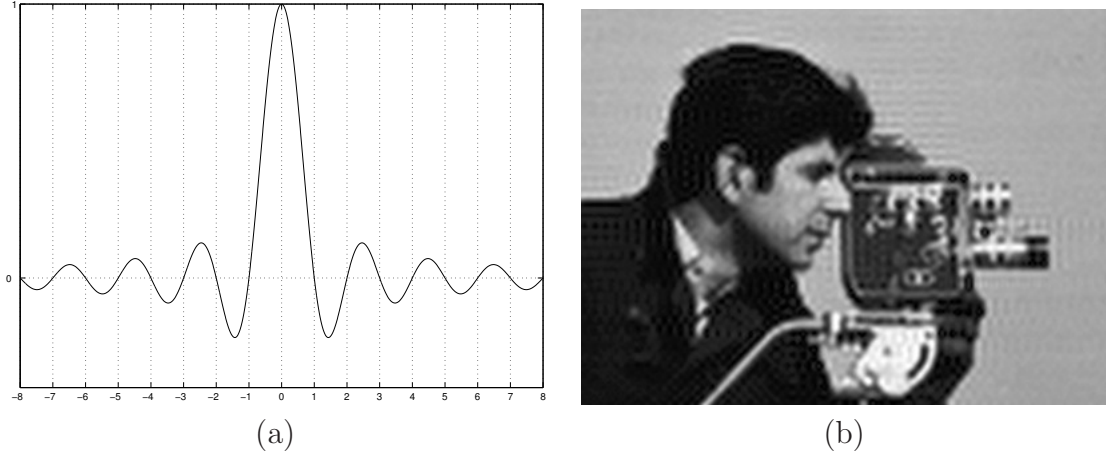


FIG. 2.2 : (a) : la fonction $\text{sinc}(\pi t)$ et (b) : interpolation sinus cardinale $f_T(\mathbf{x})$ de l'image $(v[\mathbf{k}])$ (portion de l'image *Camera*).

Interpolation classique

L'interpolation moderne tire donc ses racines de la théorie de Shannon. Une propriété importante de la fonction $\varphi(t) = \text{sinc}(\pi t)$ est d'être *interpolante* (on dira aussi que φ est un *interpolateur*), c'est-à-dire

$$\varphi(k) = \delta_{k,0} \quad \forall k \in \mathbb{Z}. \quad (2.19)$$

Bien que l'interpolation sinus cardinal permette de reconstruire parfaitement les fonctions à bande limitée, on peut formuler plusieurs critiques à son égard. Tout d'abord, les scènes naturelles et la plupart des signaux ne sont pas à bande limitée [214, 182]. Si l'on reconstruit avec la fonction $\text{sinc}(\pi t)$ un signal qui n'est pas à bande limitée, des oscillations importantes (pseudo-phénomène de Gibbs appelé *ringing*) apparaissent dans la fonction reconstruite, comme illustré sur la figure 2.2. De plus, cette fonction ayant un support infini, la somme dans (2.17) est infinie, ce qui veut dire qu'on ne peut pas implémenter exactement la reconstruction. En outre, la décroissance lente de $|\text{sinc}(t)|$ lorsque $|t| \rightarrow \infty$ rend la reconstruction sensible au bruit et aux erreurs sur les échantillons, ainsi qu'à la troncation éventuelle de la fonction sinc.

Le sinus cardinal est donc généralement remplacé par une fonction φ interpolante à support compact et proche du sinus cardinal, dans (2.17). La méthode de reconstruction par interpolation devient :

$$f_T(t) = \sum_{k \in \mathbb{Z}} v[k] \varphi\left(\frac{t}{T} - k\right) \quad \forall t \in \mathbb{R}. \quad (2.20)$$

Il découle naturellement du fait que φ est interpolante que $f(Tk) = v[k] = s[Tk]$ pour tout k . Il est d'un intérêt pratique certain de choisir φ à support compact, car alors la sommation dans l'équation (2.20) ne porte plus que sur un nombre fini de termes, ce qui rend la reconstruction implémentable.

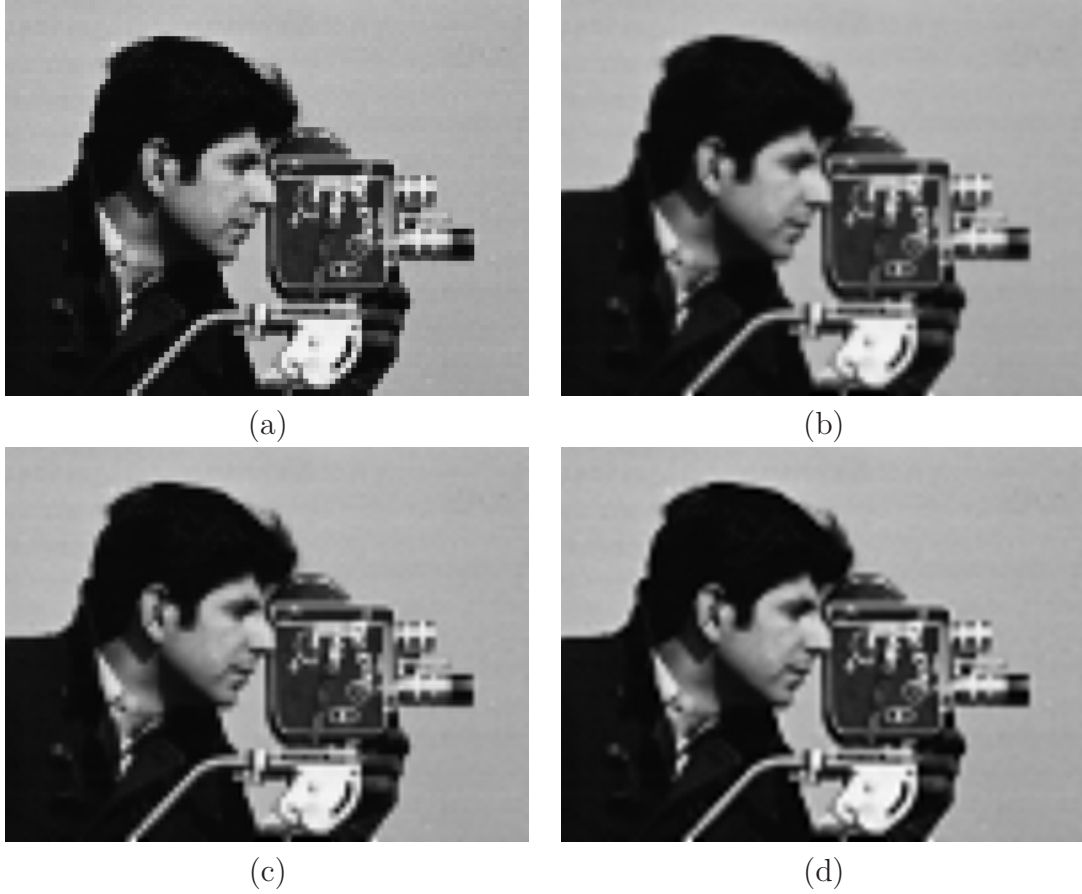


FIG. 2.3 : Portion de l'image *Camera* reconstruite (a) par interpolation au plus proche, (b) par interpolation bilinéaire, (c) par interpolation quadratique, et (d) par interpolation bicubique.

Si l'on reconstruit un signal s dont l'énergie est localisée principalement dans la bande de Nyquist, ce qui est toujours le cas en pratique, il convient de choisir φ la plus proche possible du sinus cardinal, afin de garantir une reconstruction f_T proche de s . Une première idée est d'apodiser (multiplier) le sinus cardinal par une fonction à support compact [153, 154]. Les chercheurs se sont plutôt intéressés à chercher des générateurs polynomiaux par morceaux, pour leur simplicité lors de manipulations comme l'évaluation ponctuelle ou le calcul de dérivées. Les fonctions les plus couramment utilisées sont les quatre ci-après. Nous renvoyons à [155, 154, 152, 212, 211, 35] pour un panorama plus complet des multiples autres générateurs proposés dans la littérature.

- Interpolation par duplication, dite aussi *au plus proche voisin* :

$$\varphi(t) = \begin{cases} 1 & \text{si } t \in]-\frac{1}{2}, \frac{1}{2}[\\ \frac{1}{2} & \text{si } |t| = \frac{1}{2} \\ 0 & \text{sinon} \end{cases} . \quad (2.21)$$

- Interpolation linéaire (terminologie prêtant à confusion), dite aussi bilinéaire en 2D :

$$\begin{cases} \varphi(t) = 1 - |t| & \text{si } |t| \leq 1 \\ \varphi(t) = 0 & \text{sinon.} \end{cases} \quad (2.22)$$

- Interpolation quadratique [88] :

$$\begin{cases} \varphi(t) = -2t^2 + 1 & \text{si } |t| \leq \frac{1}{2} \\ \varphi(t) = t^2 - \frac{5}{2}|t| + \frac{3}{2} & \text{si } \frac{1}{2} \leq |t| \leq \frac{3}{2} \\ \varphi(t) = 0 & \text{sinon.} \end{cases} \quad (2.23)$$

- Interpolation cubique [129], dite aussi bicubique en 2D, utilisée couramment avec le paramètre $A = -\frac{1}{2}$:

$$\begin{cases} \varphi(t) = (A + 2)|t|^3 - (A + 3)t^2 + 1 & \text{si } |t| \leq 1 \\ \varphi(t) = A(|t|^3 - 5t^2 + 8|t| - 4) & \text{si } 1 \leq |t| \leq 2 \\ \varphi(t) = 0 & \text{sinon.} \end{cases} \quad (2.24)$$

La figure 2.4 montre les quatre générateurs ainsi que leurs spectres. Ces fonctions ont été conçues pour approcher au mieux le sinus cardinal en fréquence, pour un support spatial de taille donnée. Toutes présentent une certaine atténuation dans la bande de Nyquist ($\omega \in]-\pi, \pi[$) (donc un effet de flou dans le cas de la reconstruction d'images, visible sur la figure 2.3), et ne sont pas nulles au-delà (ce qui entraîne des effets de repliement de spectre, comme des effets d'escaliers le long des contours obliques en reconstruction d'images, voir la figure 2.3). La transition entre la bande passante et la bande stoppante est suffisamment douce pour que l'effet de Gibbs soit contenu : il est totalement absent pour l'interpolation au plus proche voisin et l'interpolation bilinéaire, grâce à l'unimodalité des générateurs associés (c'est-à-dire $0 \leq \varphi(t') \leq \varphi(t)$ si $|t'| \geq |t|$). Avec l'interpolation bicubique, des oscillations commencent à devenir visibles, comme le montre la figure 2.3 (d).

Interpolation généralisée

L'interpolation à l'aide de la formule (2.20) repose sur l'utilisation d'une fonction φ interpolante et à support compact, afin que l'implémentation soit aisée. Unser et coll. ont montré que l'on pouvait considérer le cadre plus général de l'*interpolation généralisée* [37] : si l'on choisit φ à support compact, non forcément interpolante, on peut interpoler les échantillons $v[k]$ en deux étapes : la reconstruction proprement dite à l'aide de la formule

$$\boxed{f_T(t) = \sum_{k \in \mathbb{Z}} c[k] \varphi\left(\frac{t}{T} - k\right)}, \quad (2.25)$$

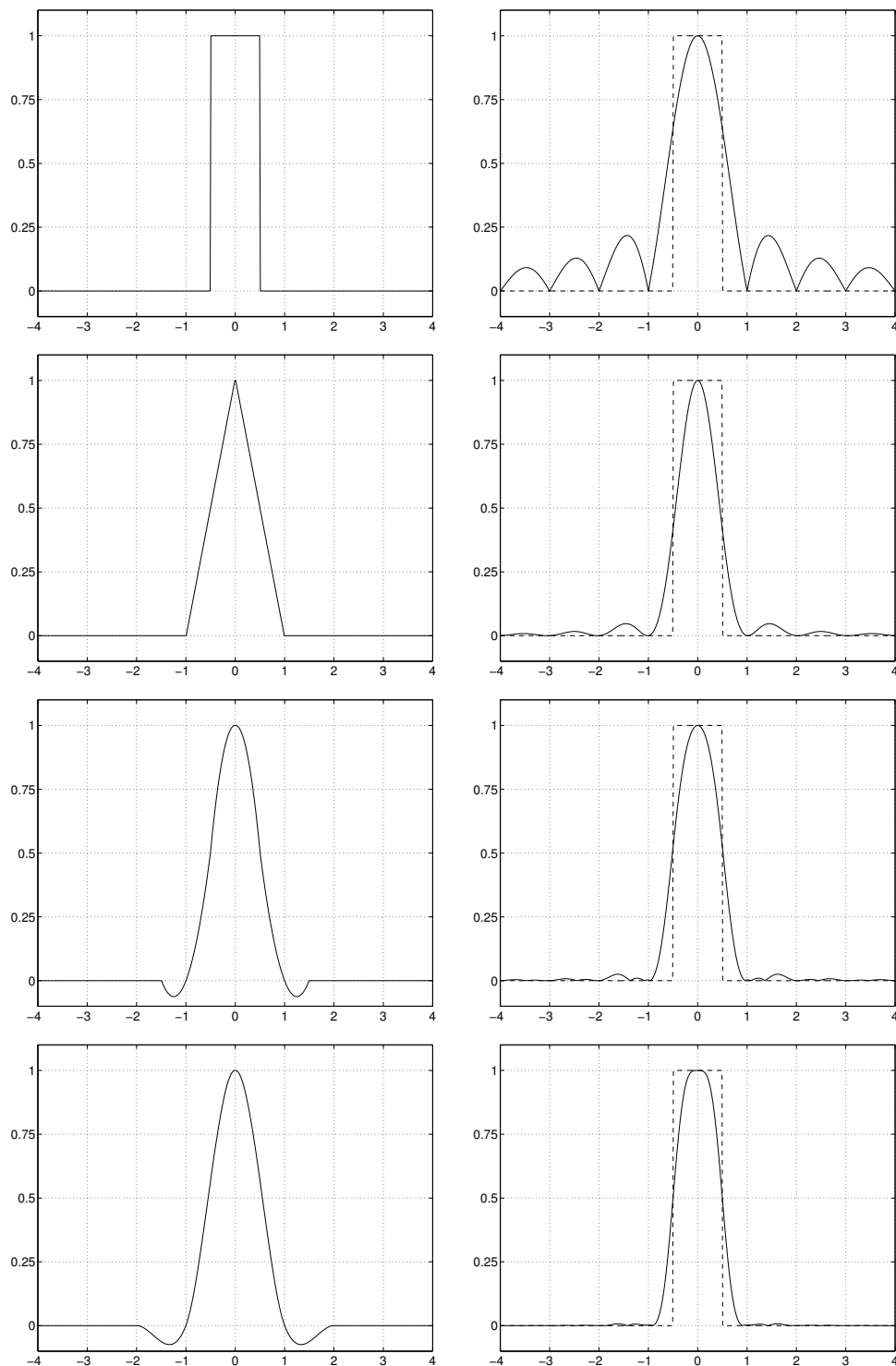


FIG. 2.4 : De haut en bas : générateurs polynomiaux par morceaux $\varphi(t)$ de degrés 0,1,2,3 (à gauche) et leurs spectres respectifs $\hat{\varphi}(\omega)$ (à droite, avec comme variable la fréquence $\nu = \omega/2\pi$).

et un prétraitement qui détermine les coefficients $c[k]$ par filtrage à partir des $v[k]$: $c = v * p_{\text{int}}$, où le préfiltre p_{int} est défini comme l'inverse du filtre b_φ , la version discrétisée de φ : $b_\varphi[k] = \varphi(k) \forall k \in \mathbb{Z}$. On a donc :

$$p_{\text{int}} = b_\varphi^{-1} \xleftrightarrow{\mathcal{Z}} P(z) = \frac{1}{\sum_{k \in \mathbb{Z}} \varphi(k) z^{-k}} \xleftrightarrow{\mathcal{F}} \hat{p}(\omega) = \frac{1}{\sum_{k \in \mathbb{Z}} \hat{\varphi}(\omega + 2k\pi)}. \quad (2.26)$$

Il est intéressant, au moins formellement, de considérer la fonction interpolante, dite aussi fonction *cardinale*, qui vaut

$$\varphi_{\text{int}}(t) = \sum_{k \in \mathbb{Z}} p_{\text{int}}[k] \varphi(t - k). \quad (2.27)$$

L'interpolation généralisée est alors formellement équivalente à l'interpolation classique avec la fonction cardinale, puisque

$$f_T(t) = \sum_{k \in \mathbb{Z}} v[k] \varphi_{\text{int}}\left(\frac{t}{T} - k\right) \quad \forall t \in \mathbb{R}. \quad (2.28)$$

Si φ est interpolante, alors $\varphi_{\text{int}} = \varphi$ et l'on retrouve le cas de l'interpolation classique, car $b_\varphi = b_\varphi^{-1} = \delta$ et aucun préfiltrage n'est requis.

Notons que le filtrage inverse avec b_φ^{-1} nécessite essentiellement le même temps de calcul que la convolution avec b_φ , grâce à l'algorithme de filtrage récursif proposé dans [221]. Lors de l'interpolation, et ce d'autant plus que le facteur est important, le temps de calcul est majoritairement utilisé pour les multiples évaluations ponctuelles de φ lors de la reconstruction (2.25), et le temps est donc proportionnel à la taille du support de φ . Il est donc d'un intérêt certain de choisir φ à support compact et de taille réduite. L'étape de préfiltrage est très rapide et n'augmente pas de manière significative le temps de calcul. Ainsi, seule la taille du support de φ importe, et la fonction équivalente φ_{int} peut avoir un support de taille infini, ce qui constitue la nouveauté de l'interpolation généralisée.

Il n'y a donc pas d'intérêt particulier à imposer à φ d'être interpolante. Au contraire, il a été montré que cela est préjudiciable à la qualité [32]. Ne pas imposer l'interpolation offre des libertés supplémentaires dans le choix de φ . Dans [212] et [211], une approche unifiée de l'interpolation est présentée, et de nombreuses fonctions φ sont comparées des points de vue théorique et expérimental.

En fait, la qualité de l'interpolation n'est pas liée à la régularité du générateur φ comme on pourrait le penser, mais à ses capacités d'approximation. En particulier, on dit que $\varphi(t)$ a un *ordre d'approximation* égal à L si tout polynôme de degré au plus $L - 1$ peut être exprimé comme combinaison linéaire des $\varphi(t - k)$, $k \in \mathbb{Z}$. Strang et Fix ont montré que φ a un ordre d'approximation L si et seulement si [205] :

$$\hat{\varphi}(0) \neq 0 \quad \text{et} \quad \hat{\varphi}^{(r)}(2\pi k) = 0 \quad \forall k \neq 0, \forall r \in \llbracket 0, L - 1 \rrbracket. \quad (2.29)$$

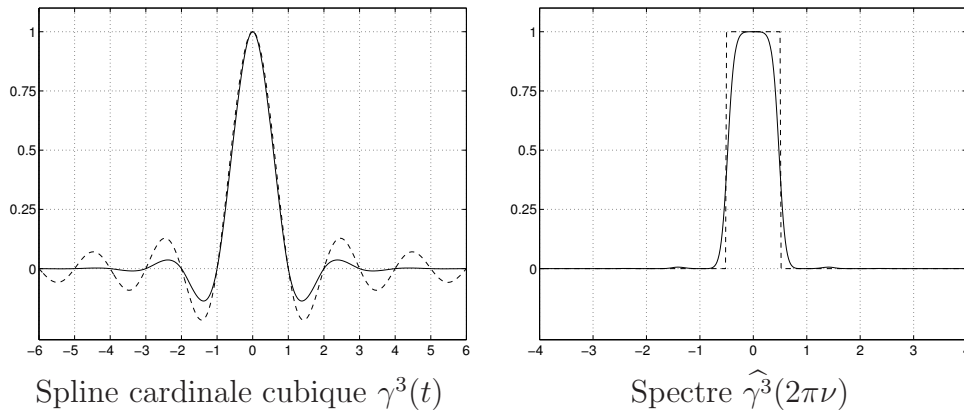


FIG. 2.5 : Spline cardinale cubique et son spectre. En pointillés : le sinus cardinal et son spectre.

La qualité de l'interpolation est directement liée à l'ordre d'approximation L de φ [212, 154, 152, 33]. Nous reviendrons sur l'importance de l'ordre d'approximation pour la qualité de reconstruction à la section 2.5. Dans l'idéal, φ doit donc avoir un support de taille W minimale et un ordre L maximal. Il a été montré que $W \geq L$, et les générateurs optimaux, pour lesquels $W = L$, sont des générateurs polynomiaux par morceaux appelés MOMS [33].

Les B-splines β^n (2.7) sont très utilisées en traitement du signal, et offrent de multiples avantages [221, 218]. La B-spline de degré n $\beta^n(t)$ a un support de taille $n + 1$ et un ordre d'approximation $L = n + 1$, c'est donc une MOMS. On parle alors d'*interpolation spline*, et la fonction f_T est appelée une *spline*. Seules β^0 et β^1 sont interpolantes, et l'interpolation avec ces fonctions est rigoureusement identique à l'interpolation au plus proche voisin et à l'interpolation linéaire, respectivement. La fonction φ_{int} définie par (2.27), lorsque $\varphi = \beta^n$, est appelée *spline cardinale* de degré n . Nous la noterons $\gamma^n(t)$ dans la suite. Le préfiltrage est appelé *transformée B-spline directe*, car il fournit les *coefficients spline* $c[k]$ à partir des échantillons $v[k]$. Comme on le voit sur la figure 2.5, la spline cardinale cubique, de support infini, est plus proche du sinus cardinal que les générateurs à support compact de la figure 2.4. On remarque dans l'exemple de la figure 2.6 que les effets de flou et de crénelage des contours sont bien contenus.

On peut montrer que toutes les MOMS d'ordre L sont des combinaisons linéaires de β^{L-1} et de ses dérivées successives [33]. On peut ainsi construire des fonctions MOMS ayant des propriétés d'approximation légèrement meilleures que les B-splines, appelées O-MOMS (*optimal MOMS*). Les interpolations spline et O-MOMS peuvent être considérées comme les meilleures méthodes pour l'interpolation d'images de nature variée [141, 212, 154]. Cependant, il est recommandé [163] de ne pas aller au-delà du degré $n = 3$, à cause des oscillations qui apparaissent. En effet, lorsque $n \rightarrow \infty$, l'interpolation spline tend vers l'interpolation par sinus cardinal [7].



FIG. 2.6 : Reconstruction par interpolation spline quadratique (a) et cubique (b).

2.3.2 Reconstruction consistante dans le cas non bruité

Dans les exemples de reconstruction optimales de la section précédente, ainsi que pour l'interpolation, nous avons vu que la solution f_T pouvait s'écrire de manière paramétrique comme dans l'équation (2.25). Une telle fonction appartient à l'espace

$$V_T(\varphi) = \left\{ g(t) = \sum_{k \in \mathbb{Z}} c[k] \varphi\left(\frac{t}{T} - k\right) \mid c \in \mathbb{R}^{\mathbb{Z}} \right\}. \quad (2.30)$$

Une fois que φ est choisie, on peut alors reformuler le problème de reconstruction comme visant à chercher la meilleure fonction $f_T(t)$ parmi l'espace $V_T(\varphi)$. Cet espace vectoriel, engendré par les translatés de pas multiple de T de la fonction génératrice $\varphi(t)$, est invariant par translation de pas multiple de T :

$$g(t) \in V_T(\varphi) \Rightarrow g(t - Tk) \in V_T(\varphi) \quad \forall k \in \mathbb{Z}. \quad (2.31)$$

Un tel espace est dit *linear shift-invariant*; on parlera dans la suite d'*espace LSI*. Ces espaces sont très utilisés en théorie de l'approximation [40, 124] et pour les méthodes à éléments finis. Ils sont à la base de l'approximation multi-résolution et de la théorie des ondelettes [149, 83]. En théorie de l'échantillonnage, leur utilisation est de plus en plus répandue [144, 62, 4, 219, 6].

Nous dirons qu'une fonction de $V_T(\varphi)$ a une **résolution** égale à $1/T$. On entend par là qu'une fonction de $V_T(\varphi)$, bien que définie continûment, est paramétrée par des coefficients discrets. Son innovation intrinsèque est de 1 toutes les T unités, c'est-à-dire que $f_T \in V_T(\varphi)$ contient essentiellement la même information qu'un peigne de Dirac de pas T :

$$\sum_{k \in \mathbb{Z}} c[k] \varphi\left(\frac{\cdot}{T} - k\right) = \left(\sum_{k \in \mathbb{Z}} c[k] \delta(\cdot - Tk) \right) * \varphi\left(\frac{\cdot}{T}\right) \quad (2.32)$$

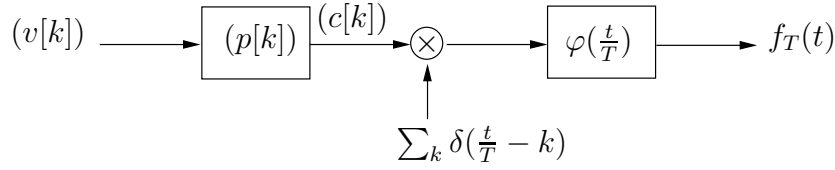


FIG. 2.7 : Formalisation du processus de reconstruction : le signal discret v est d'abord convolué avec le préfiltre p , pour fournir les coefficients $c[k]$ paramétrisant la solution $f_T(t)$ dans l'espace $V_T(\varphi)$.

Le choix de φ est un préalable à la reconstruction elle-même. φ peut être choisie de manière arbitraire, au contraire des situations de la section 2.2 où elle était induite par le critère choisi. Nous avons déjà mentionné les B-splines et autres MOMS pour leurs propriétés attrayantes : elles ont un support compact et un ordre d'approximation maximal. Afin que chaque fonction $g \in V_T(\varphi) \cap L_2(\mathbb{R})$ ait une expression unique et stable de la forme (2.25), on impose que les fonctions $\varphi(\frac{t}{T} - k)$ forment une base de Riesz. On impose aussi une condition supplémentaire peu restrictive sur φ , afin que $g \in V_T(\varphi)$ soit bornée si ses coefficients $c[k]$ le sont :

$$\sup_{t \in [0,1[} \sum_{k \in \mathbb{Z}} |\varphi(t - k)| < +\infty. \quad (2.33)$$

Enfin, on fait l'hypothèse peu restrictive que $V_T(\tilde{\varphi})^\perp \cap V_T(\varphi) = \{0\}$. On supposera dans la suite que la fonction φ vérifie toujours ces trois conditions.

Le problème de reconstruction dans un espace LSI devient donc purement discret, puisqu'à partir d'un signal $(v[k])$, on cherche une suite de coefficients $c = (c[k])_{k \in \mathbb{Z}}$ paramétrisant la fonction reconstruite $f_T(t)$. On souhaite que le processus de reconstruction $\mathcal{M} : v \mapsto f_T$ soit linéaire et invariant par translation de pas multiple de T ; il en résulte que les coefficients $c[k]$ sont déterminés par filtrage discret avec un certain préfiltre p :

$$\boxed{c = v * p}. \quad (2.34)$$

On parle de *préfiltrage* et de *préfiltre* car ce filtrage précède la reconstruction proprement dite, définie par l'équation (2.25). Nous avons schématisé le processus complet de reconstruction sur la figure 2.7. Tout le problème revient donc à chercher un *bon* préfiltre p , indépendant de v . Nous avons déjà donné des exemples de préfiltres : (2.5) (2.14) (2.26).

Lorsque l'on n'a pas de connaissance sur s , on cherche à ce que la reconstruction f_T soit en adéquation avec les données. La solution classique, dans le cas où il n'y a pas de bruit, est de chercher une fonction f_T *consistante* avec les données [220, 226, 219, 94, 93, 96], c'est-à-dire vérifiant

$$f_T * \tilde{\varphi}_T(Tk) = v[k] \quad \forall k \in \mathbb{Z}, \quad (2.35)$$

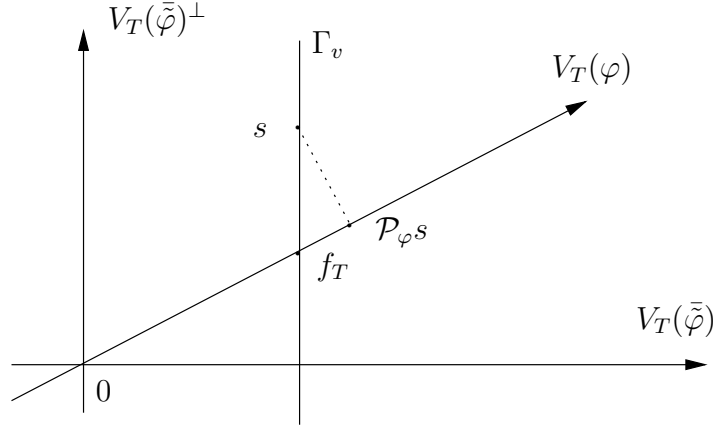


FIG. 2.8 : Illustration de la reconstruction consistante, fournissant la fonction f_T , projection oblique de s dans $V_T(\varphi)$, et unique fonction de $V_T(\varphi)$ consistante avec les données, c'est-à-dire appartenant à Γ_v .

autrement dit $\mathcal{D}f_T = v$, ou $f_T \in \Gamma_v$ défini par (2.8). Ainsi, f_T fournit les mêmes mesures que s si on la place dans le modèle d'acquisition (1.10). On remarque que les solutions présentées jusqu'ici dans le cas non bruité sont toutes consistantes ; l'interpolation dans $V_T(\varphi)$ correspond à la reconstruction consistante dans le cas particulier $\tilde{\varphi} = \delta$.

Il existe une unique fonction f_T consistante dans l'espace $V_T(\varphi)$. Ainsi on peut écrire :

$$\boxed{\{f_T\} = V_T(\varphi) \cap \Gamma_v}. \quad (2.36)$$

Ses coefficients $c[k]$ s'obtiennent par préfiltrage (2.34) à l'aide du préfiltre de reconstruction consistante p_{con} défini par

$$p_{\text{con}}^{-1}[k] = \tilde{\varphi} * \varphi(k) \xleftrightarrow{\mathcal{F}} \hat{p}_{\text{con}}(\omega) = \frac{1}{\sum_{k \in \mathbb{Z}} \hat{\tilde{\varphi}}(\omega + 2k\pi) \hat{\varphi}(\omega + 2k\pi)}. \quad (2.37)$$

La reconstruction consistante a plusieurs propriétés intéressantes :

- La reconstruction consistante est optimale du point de vue du critère des moindres carrés : elle minimise (2.1).
- Si $s \in V_T(\varphi)$, s est reconstruite parfaitement : $f_T = s$. Le principe de reconstruction consistante généralise donc le théorème de Shannon, qui correspond au cas $\varphi(t) = \text{sinc}(\pi t)$ et $\tilde{\varphi} = \delta$. En effet, dire que s est à bande limitée est équivalent à dire que $s \in V_T(\text{sinc}(\pi \cdot))$.
- L'opérateur d'approximation $\mathcal{Q}s \mapsto f_T$ est un projecteur : $\mathcal{Q} \circ \mathcal{Q} = \mathcal{Q}$. f_T est donc la projection *oblique* (c'est-à-dire non orthogonale) de s dans $V_T(\varphi)$. Définissons pour la suite :

$\mathcal{P}_{\phi_1, \phi_2}$, la projection oblique dans $V_T(\phi_1)$ parallèlement à $V_T(\phi_2)$,

$\mathcal{P}_{\phi_1, \phi_2^\perp}$, la projection oblique dans $V_T(\phi_1)$ perpendiculairement à $V_T(\phi_2)$,

c.-à-d. parallèlement à $V_T(\phi_2)^\perp$,

$\mathcal{P}_{\phi_1} = \mathcal{P}_{\phi_1, \phi_1^\perp} = \mathcal{P}_{V_T(\phi_1)}^\perp$, la projection orthogonale dans $V_T(\phi_1)$.

Alors la reconstruction consistante correspond à $\mathcal{Q} = \mathcal{P}_{\varphi, \tilde{\varphi}^\perp}$. On a $f_T = \mathcal{P}_\varphi s$ si et seulement si $\tilde{\varphi} \in V_1(\varphi)$. Cette situation est illustrée sur la figure 2.8.

• L'opérateur de reconstruction $\mathcal{M} : v \mapsto f_T$ est un opérateur pseudo-inverse de l'opérateur de discrétisation $\mathcal{D} : s \mapsto v$. En effet, $\mathcal{D} \circ \mathcal{M}$ est l'identité. $\mathcal{M}$ est même l'inverse généralisée de $\mathcal{D}$ si l'on munit L_2 de la norme $\|\mathcal{L} \cdot\|$, où $\mathcal{L}$ est tel que φ est une $\mathcal{L}^* \mathcal{L}$ -spline, voir (2.3). La reconstruction consistante vérifie donc les propriétés d'optimalité minimax mentionnées en 2.2.2.

Définissons maintenant la fonction $\varphi_{\text{eq}}(t)$ par

$$\boxed{\varphi_{\text{eq}}(t) = \sum_{k \in \mathbb{Z}} p[k] \varphi(t - k)} \xleftrightarrow{\mathcal{F}} \hat{\varphi}_{\text{eq}}(\omega) = \hat{p}(\omega) \hat{\varphi}(\omega). \quad (2.38)$$

Dans le cas où $p = p_{\text{int}}$, on retrouve la fonction cardinale $\varphi_{\text{eq}} = \varphi_{\text{int}}$. On peut alors réécrire le processus de reconstruction comme

$$\boxed{f_T(t) = \sum_{k \in \mathbb{Z}} v[k] \varphi_{\text{eq}}(t - k)}. \quad (2.39)$$

La propriété de consistance est équivalente à la *biorthogonalité* de $\tilde{\varphi}$ et φ_{eq} , c'est-à-dire

$$\tilde{\varphi} * \varphi_{\text{eq}}(k) = \delta_{k,0} \xleftrightarrow{\mathcal{F}} \sum_{k \in \mathbb{Z}} \hat{\tilde{\varphi}}(\omega + 2k\pi) \hat{\varphi}_{\text{eq}}(\omega + 2k\pi) = 1. \quad (2.40)$$

C'est l'étape de préfiltrage qui impose la biorthogonalité du processus, puisque $\tilde{\varphi}$ et φ sont indépendantes l'une de l'autre, et non nécessairement biorthogonales. Notons que la biorthogonalité est à la base de la construction des analyses multi-résolution et de la théorie des ondelettes [83, 149, 2, 228].

2.3.3 Approches des moindres carrés dans le cas bruité

L'approche consistante vise donc à obtenir la fonction la plus proche des données dans l'espace $V_T(\varphi)$, au sens du critère (2.1). Lorsque les données sont bruitées, il faut relativiser ce critère, car il n'est pas nécessairement judicieux d'être exactement fidèle à des données inexactes. La solution généralement adoptée consiste à définir f_T comme la fonction de $V_T(\varphi)$ minimisant le critère des moindres carrés régularisés (2.2), c'est-à-dire que l'on cherche un compromis entre l'attache aux données au sens des moindres carrés, et la lissitude de la solution au sens où $\|\mathcal{L} f_T\|$ doit être faible, pour un certain opérateur $\mathcal{L}$ à choisir, indépendamment de φ .

Notons $\hat{\mathcal{L}}(\omega)$ la réponse fréquentielle telle que $\widehat{\mathcal{L}g}(\omega) = \hat{\mathcal{L}}(\omega) \hat{g}(\omega)$. Par exemple, si $\mathcal{L} = \frac{d^n}{dt^n}$, on a $\hat{\mathcal{L}}(\omega) = (I\omega)^n$. Alors, le préfiltre p fournissant les coefficients $c[k]$ de la solution f_T au problème de minimisation de (2.2) dans $V_T(\varphi)$ est défini par [96] :

$$\hat{p}(\omega) = \frac{\sum_{k \in \mathbb{Z}} \hat{\tilde{\varphi}}(\omega + 2k\pi)^* \hat{\varphi}(\omega + 2k\pi)^*}{|\sum_{k \in \mathbb{Z}} \hat{\tilde{\varphi}}(\omega + 2k\pi) \hat{\varphi}(\omega + 2k\pi)|^2 + \lambda \sum_{k \in \mathbb{Z}} |\hat{\mathcal{L}}(\omega + 2k\pi)|^2 |\hat{\varphi}(\omega + 2k\pi)|^2}. \quad (2.41)$$

On remarque que si φ est une $\mathcal{L}^*\mathcal{L}$ -spline, c'est-à-dire vérifie (2.3), on retrouve le préfiltre optimal (2.5), puisqu'alors $\widehat{\mathcal{L}^*\mathcal{L}\varphi}(\omega) = |\widehat{L}(\omega)|^2 \hat{\varphi}(\omega) = \hat{h}(\omega) \hat{\varphi}(\omega)^*$.

Cette solution régularisée étend donc l'approche consistante au cas bruité, en s'appuyant sur l'approche variationnelle de la section 2.2.1. On peut proposer une nouvelle extension de l'approche consistante au cas bruité, dans le cas stochastique. Supposons que s est la réalisation d'un processus aléatoire stationnaire à temps continu de moyenne nulle, de densité spectrale de puissance $\hat{c}_s(\omega) > 0$. Notons $v_0[k] = v[k] - \mu[k] = s * \tilde{\varphi}_T(Tk)$ la version non bruitée des échantillons $v[k]$. On peut ainsi étendre le critère des moindres carrés (2.1) par le critère :

$$\varepsilon_{\ell_2, \text{STO}}(g) = \mathcal{E}\{|g * \tilde{\varphi}_T(Tk) - v_0[k]|^2\}, \quad (2.42)$$

qui ne dépend pas de $k \in \mathbb{Z}$, afin de chercher une fonction reconstruite qui, en moyenne, fournirait des mesures proches des mesures idéales qu'auraient fournies s en l'absence de bruit. On a tout d'abord

$$\hat{c}_{v_0}(\omega) = \frac{1}{T} \sum_{k \in \mathbb{Z}} |\hat{\varphi}(\omega + 2k\pi)|^2 \hat{c}_s\left(\frac{\omega + 2k\pi}{T}\right). \quad (2.43)$$

Pour une fonction $g \in V_T(\varphi)$ dont les coefficients $c[k]$ ont été obtenus par le préfiltrage $c = v * p$, on a :

$$\begin{aligned} \varepsilon_{\ell_2, \text{STO}}(g) = \frac{1}{2\pi} \int_0^{2\pi} \hat{c}_{v_0}(\omega) & \left| 1 - \hat{p}(\omega) \sum_{k \in \mathbb{Z}} \hat{\varphi}(\omega + 2k\pi) \hat{\varphi}(\omega + 2k\pi) \right|^2 + \\ & \hat{c}_\mu(\omega) \left| \hat{p}(\omega) \sum_{k \in \mathbb{Z}} \hat{\varphi}(\omega + 2k\pi) \hat{\varphi}(\omega + 2k\pi) \right|^2 d\omega. \end{aligned} \quad (2.44)$$

On montre aisément que ce critère est minimal lorsque

$$\hat{p}(\omega) = \frac{\hat{c}_{v_0}(\omega)}{\hat{c}_{v_0}(\omega) + \hat{c}_\mu(\omega)} \frac{1}{\sum_{k \in \mathbb{Z}} \hat{\varphi}(\omega + 2k\pi) \hat{\varphi}(\omega + 2k\pi)}. \quad (2.45)$$

On constate que ce préfiltre se décompose en deux filtres : le filtre de Wiener discret débruitant les échantillons $v[k]$, et le préfiltre de reconstruction consistante. En l'absence de bruit, on se ramène à la reconstruction consistante.

2.4 Estimation de s au sens L_2 dans le cas non bruité

Nous avons vu que le problème de reconstruction pouvait être reformulé comme visant à trouver un *bon* représentant de s dans un certain espace $V_T(\varphi)$, à partir des données $(v[k])$. Nous avons exhibé plusieurs solutions f_T minimisant chacune un certain critère ε . Dans cette section, on se place dans le cas non bruité. On suppose que l'espace de reconstruction

$V_T(\varphi)$ est fixé, indépendamment du fait que l'on dispose ou non de connaissances *a priori* sur s . Plusieurs raisons peuvent justifier une telle situation :

- On ne dispose pas d'information sur s hormis le fait que $s \in \Gamma_v$. Dans ce cas, il faut bien choisir d'une manière ou d'une autre la méthode de reconstruction, et donc fixer $V_T(\varphi)$ arbitrairement.
- On veut choisir $V_T(\varphi)$ pour des raisons liées aux traitements ultérieurs que l'on va appliquer sur f_T ; par exemple, on peut exiger que φ soit différentiable un certain nombre de fois si l'on désire calculer des dérivées de f_T en certains points. On peut aussi être confronté à des contraintes liées à la facilité d'implémentation ou au temps de calcul de la reconstruction. Dans ce cas, effectuer la reconstruction dans un espace spline est un bon choix [218].
- La forme de la fonction reconstruite est imposée, car le processus $\mathcal{M}$ modélise en fait la réponse d'un dispositif de conversion « digital vers analogique », comme un périphérique d'affichage, dont φ est la réponse impulsionnelle. Prenons l'exemple de l'affichage d'une image ($v[\mathbf{k}]$) sur un écran d'ordinateur LCD. En supposant que la réponse impulsionnelle de celui-ci est idéale, et correspond à l'indicatrice d'un pixel, on a $\varphi = \mathbb{1}_{[-\frac{1}{2}, \frac{1}{2}]^2}$. $V_T(\varphi)$ est alors, formellement, l'ensemble des « fonctions » affichables par l'écran.

2.4.1 Situation du problème

Plaçons-nous pour le moment dans le cas déterministe où $s \in L_2$. Le critère le plus naturel pour estimer s est le suivant :

$$\boxed{\varepsilon_{L_2}(g) = \|s - g\|_{L_2}^2} \quad (2.46)$$

Les solutions vues précédemment, optimales pour un certain critère, ne le sont pas vis-à-vis du critère L_2 . En fait, la meilleure reconstruction possible pour ce critère est la projection orthogonale, que nous noterons $\mathcal{P}_T s$ de s dans $V_T(\varphi)$. En effet, pour une fonction quelconque $g \in L_2$, on a l'égalité suivante (théorème de Pythagore) :

$$\|s - g\|_{L_2}^2 = \|s - \mathcal{P}_T s\|_{L_2}^2 + \|\mathcal{P}_T s - g\|_{L_2}^2 \quad (2.47)$$

qui énonce que l'erreur se compose d'un premier terme indépendant de la méthode de reconstruction, et d'un second qui est nul si $g = \mathcal{P}_T s$. Minimiser ε_{L_2} dans $V_T(\varphi)$ est donc équivalent à minimiser le *regret* [96]

$$\varepsilon_{\text{REG}}(g) = \|\mathcal{P}_T s - g\|_{L_2}^2 \quad (2.48)$$

La fonction $\mathcal{P}_T s$ appartient à $V_T(\varphi)$. On peut donc l'écrire sous la forme (2.25) à l'aide de coefficients $c[k]$ valant

$$c[k] = \int_{\mathbb{R}} s(t) \bar{\varphi}_d\left(\frac{t}{T} - k\right) \frac{dt}{T} \quad (2.49)$$

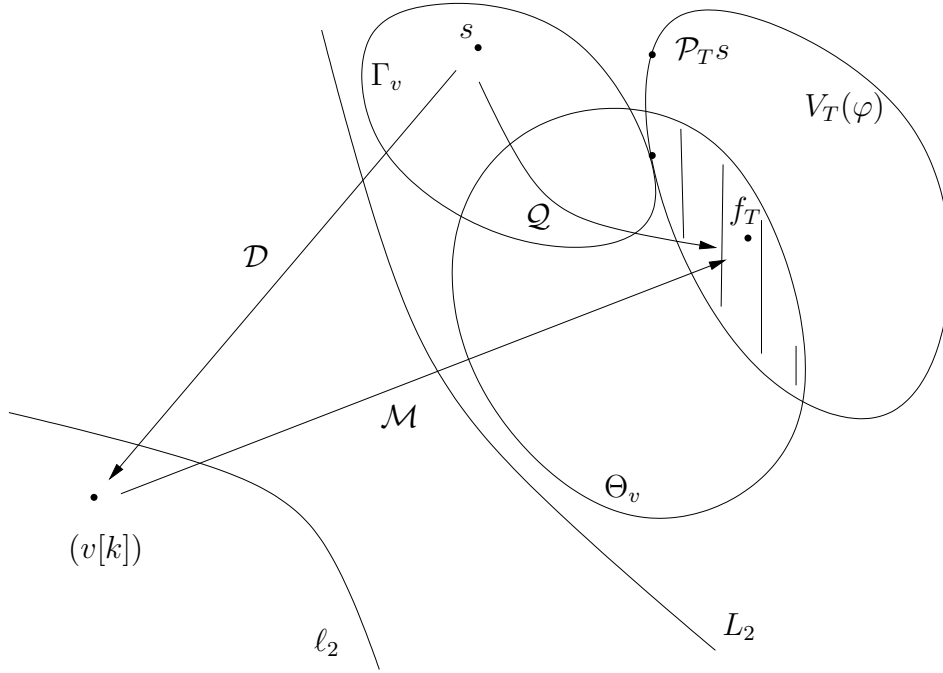


FIG. 2.9 : Formalisation du problème de reconstruction dans l'espace $V_T(\varphi)$. Θ_v représente l'ensemble des fonctions reconstituables linéairement à partir de v . On cherche donc une solution dans $V_T(\varphi) \cap \Theta_v$, ce qui ramène le problème au choix du préfiltre p . On sait juste que $s \in \Gamma_v$. Cependant, l'unique fonction reconstituée consistante de $V_T(\varphi) \cap \Theta_v \cap \Gamma_v$ n'est pas nécessairement la meilleure estimation de s , qui dans l'idéal devrait être $\mathcal{P}_T s$, malheureusement inaccessible.

où φ_d est la *fonction duale* de φ , c'est-à-dire l'unique fonction de $V_T(\bar{\varphi})$ biorthogonale à φ , s'écrivant

$$\varphi_d(t) = \sum_{k \in \mathbb{Z}} a_\varphi^{-1}[k] \bar{\varphi}(t - k) \xleftrightarrow{\mathcal{F}} \hat{\varphi}_d(\omega) = \frac{\hat{\varphi}(\omega)^*}{\sum_{k \in \mathbb{Z}} |\hat{\varphi}(\omega + 2k\pi)|^2} \quad (2.50)$$

Nous avons résumé la situation du problème sur la figure 2.9. La seule connaissance que nous avons de $s(t)$ est qu'elle appartient à l'ensemble consistant Γ_v . Nous avons noté Θ_v l'ensemble des fonctions reconstituables par un processus $\mathcal{M}$ linéaire et invariant par translation, à partir de la donnée v , c'est-à-dire,

$$\Theta_v = \left\{ \sum_{k \in \mathbb{Z}} v[k] \psi\left(\frac{t}{T} - k\right) \mid \psi \in L_2 \right\}. \quad (2.51)$$

Puisque l'on cherche une solution dans l'espace $V_T(\varphi)$, on cherche donc f_T dans $V_T(\varphi) \cap \Theta_v$, la région hachée sur la figure, correspondant à prendre $\psi = \varphi_{\text{eq}}$ dans (2.51); cela revient à ne laisser comme unique degré de liberté que le préfiltre p . L'intersection $V_T(\varphi) \cap \Theta_v \cap \Gamma_v$ est réduite à un seul élément, la fonction consistante obtenue avec le préfiltre p_{con} (2.37).

La solution idéale $\mathcal{P}_T s$ n'est pas accessible à partir des données $v[k]$, sauf si $\varphi_d \in V_1(\tilde{\varphi})$, autrement dit $\tilde{\varphi} \in V_1(\varphi)$, auquel cas la reconstruction consistante fournit $f_T = \mathcal{P}_T s$. Dans le cas général, la solution consistante opère une projection oblique de s dans $V_T(\varphi)$, qui peut être assez éloignée de la projection orthogonale (voir la figure 2.8), et aucun préfiltre ne permet de minimiser le critère (2.46) pour tout s [95, Prop. 2]. Il existe cependant des fonctions dans $V_T(\varphi) \cap \Theta_v$ qui sont plus proches de $\mathcal{P}_T s$ que la reconstruction consistante. Nous verrons par la suite comment construire des préfiltres permettant d'atteindre de telles fonctions. Pour l'instant, on présente l'approche minimax, qui permet de trouver une solution satisfaisante, bien que sous-optimale, au problème de reconstruction.

2.4.2 Approches minimax

Bien que le critère ε_{L_2} ne puisse être minimisé de manière systématique, sans autre information sur s , on peut par contre minimiser ce critère sur une classe restreinte de fonctions. Par exemple, si l'on a la connaissance que $s \in V_T(\varphi)$, la reconstruction consistante est optimale au sens L_2 , puisqu'elle assure la reconstruction parfaite de s , c'est-à-dire $f_T = s$ et $\varepsilon_{L_2}(f_T) = 0$. Plus généralement, si $s \in V_T(\psi)$ pour $\psi \in L_2$ connue, on peut là-aussi trouver un préfiltre minimisant ε_{L_2} : la reconstruction correspondant à $\mathcal{Q} = \mathcal{P}_\varphi \mathcal{P}_{\psi, \tilde{\varphi}^\perp}$ est optimale au sens L_2 [95]. En effet, cet opérateur reconstruit d'abord s parfaitement par reconstruction consistante dans $V_T(\psi)$, puis la projette orthogonalement dans $V_T(\varphi)$, ce qui aboutit bien à $f_T = \mathcal{P}_T s$. Le préfiltre correspondant s'écrit

$$\hat{p}(\omega) = \frac{\sum_{k \in \mathbb{Z}} \hat{\psi}(\omega + 2k\pi) \hat{\varphi}_d(\omega + 2k\pi)}{\sum_{k \in \mathbb{Z}} \hat{\tilde{\varphi}}(\omega + 2k\pi) \hat{\psi}(\omega + 2k\pi)} \quad (2.52)$$

et on retrouve bien la reconstruction consistante comme un cas particulier correspondant à $\psi \in V_1(\varphi)$.

La minimisation de ε_{L_2} sachant $s \in V_T(\psi)$ peut être réécrite comme la minimisation du critère suivant, en remarquant que $V_T(\psi) \cap \Gamma_v = \{s\}$:

$$\varepsilon_{L_2}(g) = \max_{f \in V_T(\psi) \cap \Gamma_v} \|f - g\|_{L_2}^2 \quad (2.53)$$

L'avantage de cette écriture, qui ne vaut que dans ce cas particulier, est qu'elle ne fait plus apparaître s . On voit ainsi apparaître naturellement les méthodes minimax, permettant d'inclure une connaissance sur s dans la minimisation d'un critère indépendant de s .

Considérons maintenant le cas où $s \in \Omega = \{g \mid \|\mathcal{L}g\| \leq C\}$ pour un certain opérateur $\mathcal{L}$ et une certaine constante C . On peut alors définir f_T comme la solution du critère suivant, qui n'est plus équivalent à ε_{L_2} , mais qui minimise l'erreur L_2 au pire cas parmi les instances possibles de s :

$$\varepsilon_{\text{MM}}(g) = \max_{f \in \Omega \cap \Gamma_v} \|f - g\|_{L_2}^2. \quad (2.54)$$

Notons qu'il faut bien faire porter le max sur $\Omega \cap \Gamma_v$ et non Ω , car alors le problème, qui ne tient pas compte du fait que $s \in \Gamma_v$, admet pour solution la fonction nulle [95, section 5.1.1], manifestement peu satisfaisante. Quelle est la solution minimisant ε_{MM} ? Nous avons vu à la section 2.2.2 que la minimisation sur tout L_2 aboutit à la $\mathcal{L}^*\mathcal{L}$ -spline consistante. Afin de déterminer la solution dans l'espace $V_T(\varphi)$, reformulons le problème d'estimation L_2 qui nous préoccupe dans le formalisme utilisé par la théorie de l'*optimal recovery*, présentée à la section 2.2.2. À partir des schémas 2.9 et 2.1, on voit que ce lien s'établit en posant $\mathcal{Z} = V_T(\varphi)$, muni de la norme L_2 , et $\mathcal{U} = \mathcal{P}_\varphi$. Si $g \in V_T(\varphi)$, on peut alors reformuler le critère $\varepsilon_{\text{MM}}(g)$ comme :

$$\varepsilon_{\text{MM}}(g) = \max_{f \in \Omega \cap \Gamma_v} \|\mathcal{P}_T f - g\|_{L_2}^2 = \max_{f \in \mathcal{P}_\varphi(\Omega \cap \Gamma_v)} \|f - g\|_{L_2}^2, \quad (2.55)$$

dont la minimisation dans $\mathcal{Z}$ a pour solution z_0 , le centre de Chebyshev de $\mathcal{P}_\varphi(\Omega \cap \Gamma_v)$. Or $z_0 = \mathcal{P}x_0$, où x_0 est le centre de Chebyshev de $\Omega \cap \Gamma_v$, à savoir la $\mathcal{L}^*\mathcal{L}$ -spline consistante. Soit $\psi \in L_2$, une $\mathcal{L}^*\mathcal{L}$ -spline vérifiant (2.3). Alors $\mathcal{Q} = \mathcal{P}_\varphi \mathcal{P}_{\psi, \tilde{\varphi}^\perp}$ est optimale au sens du critère ε_{MM} , et on retrouve le préfiltre (2.52), que l'on peut réécrire :

$$\hat{p}(\omega) = \frac{\sum_{k \in \mathbb{Z}} \hat{\tilde{\varphi}}(\omega + 2k\pi)^* \hat{\varphi}_d(\omega + 2k\pi) / |\hat{\mathbf{L}}(\omega + 2k\pi)|^2}{\sum_{k \in \mathbb{Z}} |\hat{\tilde{\varphi}}(\omega + 2k\pi)|^2 / |\hat{\mathbf{L}}(\omega + 2k\pi)|^2}. \quad (2.56)$$

Ce nouveau résultat généralise les théorèmes 2 et 3 de [95], qui énoncent le résultat dans le cas où $\mathcal{L}$ est l'identité (auquel cas on peut poser $\psi = \tilde{\varphi}$), ce qui donne $\mathcal{Q} = \mathcal{P}_\varphi \mathcal{P}_{\tilde{\varphi}}$.

Cette méthode de reconstruction est donc optimale à la fois lorsque s appartient à l'espace linéaire $V_T(\psi)$, et lorsque s appartient à l'ensemble quadratique Ω , quelle que soit la constante C . Cette solution est en fait naturelle : si l'on n'impose pas de contrainte sur l'espace de reconstruction, la $\mathcal{L}^*\mathcal{L}$ -spline consistante est une bonne solution, et si l'on impose l'espace de reconstruction $V_T(\varphi)$, il suffit de considérer la meilleure représentation de cette solution dans $V_T(\varphi)$, à savoir son projeté orthogonal.

Nous avons donc vu que le critère L_2 , pour lequel on ne peut pas trouver de méthode de reconstruction optimale dans le cas général, pouvait être modifié en un critère minimax afin de prendre en compte une information *a priori* sur le signal s . Cependant, une telle information n'est généralement pas disponible. De plus, la modification du critère L_2 en un critère minimax ne garantit plus la proximité systématique de f_T et s au sens L_2 . D'autre part, nous avons présenté plusieurs méthodes de reconstruction, minimisant un certain critère, mais nous ne disposons pas pour l'instant d'un moyen de les comparer quantitativement. Nous présentons dans la suite un outil fondamental, le noyau d'erreur fréquentiel, particulièrement utile pour comparer les méthodes d'approximation.

2.4.3 Analyse de l'erreur

Afin d'évaluer quantitativement la qualité d'une méthode d'approximation linéaire $\mathcal{Q} : s \mapsto f_T \in V_T(\varphi)$, définissons la quantité

$$\eta_s(T) = \sqrt{\frac{1}{2\pi} \int |\hat{s}(\omega)|^2 E(T\omega) d\omega}, \quad (2.57)$$

au moyen du *noyau d'erreur fréquentiel*

$$E(\omega) = \underbrace{1 - \hat{\varphi}(\omega)\hat{\varphi}_d(\omega)}_{E_{\min}(\omega)} + \underbrace{\hat{a}_\varphi(\omega)|\hat{p}(\omega)\hat{\varphi}(\omega) - \hat{\varphi}_d(\omega)|^2}_{E_{\text{res}}(\omega)}. \quad (2.58)$$

Alors, suivant un résultat remarquable de théorie de l'approximation dû à Thierry Blu en 1999 [36, 37], l'erreur d'approximation L_2 peut être prédite très précisément par $\eta_s(T)$, autrement dit

$$\|f_T - s\|_{L_2} \approx \eta_s(T). \quad (2.59)$$

L'estimation de l'erreur (2.59) peut être justifiée de plusieurs points de vue [37]. Tout d'abord, l'égalité dans (2.59) est exacte si s est à bande limitée dans $[-\frac{(k+1)\pi}{T}, -\frac{k\pi}{T}] \cup [\frac{k\pi}{T}, \frac{(k+1)\pi}{T}]$ pour un certain entier $k \geq 0$, ou si $\tilde{\varphi}$ et φ sont à bande limitée dans $[-\pi, \pi]$. D'autre part, une propriété importante à laquelle s'intéresse la théorie de l'approximation est la convergence de f_T vers s lorsque $T \rightarrow 0$. Introduisons l'énergie relative hors bande de Nyquist d'une fonction $g \in L_2$ par

$$e_g(T) = \frac{\int_{|\omega| \geq \frac{\pi}{T}} |\hat{g}(\omega)|^2 d\omega}{\int_{\mathbb{R}} |\hat{g}(\omega)|^2 d\omega}. \quad (2.60)$$

On a $0 \leq e_g(T) \leq 1$ et $e_g(T) \rightarrow 0$ quand $T \rightarrow 0$. Alors, si s appartient à l'espace de Sobolev W_2^r avec $r > \frac{1}{2}$, on peut écrire

$$\|s - f_T\|_{L_2} = \eta_s(T) + \gamma e_{s^{(r)}}(T) T^r \|s^{(r)}\|_{L_2}, \quad (2.61)$$

où

$$|\gamma| \leq \frac{2}{\pi^r} \sqrt{\zeta(2r) \|E\|_\infty}. \quad (2.62)$$

Ainsi, le terme $\eta_s(T)$ est le terme dominant de l'erreur $\|s - f_T\|_{L_2}$. Le second terme se comporte en $o(T^r)$, et son influence est négligeable.

Le noyau d'erreur permet aussi de calculer exactement l'erreur moyenne dans le cadre stochastique. Définissons le critère suivant, qui généralise naturellement l'erreur L_2 pour des signaux aléatoires stationnaires :

$$\varepsilon_{\text{STO}}(g) = \lim_{A \rightarrow +\infty} \frac{1}{2A} \int_{-A}^A \mathcal{E}\{|g(t) - s(t)|^2\} dt \quad (2.63)$$

$$= \frac{1}{T} \int_0^T \mathcal{E}\{|g(t) - s(t)|^2\} dt, \quad (2.64)$$

correspondant à une version moyennée sur t_0 de l'erreur quadratique ponctuelle moyenne $\varepsilon_{\text{STO},t_0}$ de (2.12), et qui est T -périodique pour une fonction reconstruite linéairement à partir de v . Notons que la minimisation directe de $\varepsilon_{\text{STO},t_0}$ dans l'espace $V_T(\varphi)$ donne une solution dépendant de t_0 ; c'est pourquoi on définit une version moyennée de l'erreur quadratique, afin d'obtenir une solution proche de s en tout point. Alors on peut montrer [37, Appendix 3] que

$$\varepsilon_{\text{STO}}(f_T) = \frac{1}{2\pi} \int_{\mathbb{R}} \hat{c}_s(\omega) E(T\omega) d\omega, \quad (2.65)$$

ce qui correspond à la même formule que pour l'erreur L_2 , en remplaçant $|\hat{s}(\omega)|^2$ par $\hat{c}_s(\omega)$ dans (2.57).

Dans le cas stochastique, le noyau d'erreur permet de calculer exactement l'erreur moyenne, et dans le cas déterministe, de calculer l'erreur L_2 de manière approchée mais précise. Etudier le noyau d'erreur associé à une méthode de reconstruction donnée est donc d'un grand intérêt. La valeur $E(\omega)$ doit être interprétée comme l'erreur relative engendrée par la méthode d'approximation à la fréquence ω . Une valeur proche de 1 signifie que l'information est perdue à cette fréquence, et une valeur supérieure à 1 signifie que non seulement l'information est perdue, mais qu'en plus l'énergie à cette fréquence va apparaître (généralement après repliement de spectre) à une autre fréquence, introduisant une distorsion supplémentaire.

Le noyau $E_{\min}(\omega)$ est une borne inférieure pour $E(\omega)$, indépendante du préfiltre p , et qui caractérise la projection orthogonale dans $V_T(\varphi)$. On peut donc quantifier la qualité d'approximation intrinsèque d'un espace $V_T(\varphi)$ au moyen du noyau $E_{\min}(T\omega)$ associé. Une fois choisie φ , le rôle du préfiltre sur la qualité d'approximation peut être quantifié au moyen de $E_{\text{res}}(T\omega)$. Le noyau d'erreur offre donc un moyen privilégié pour évaluer et comparer des méthodes de reconstruction, c'est-à-dire des couples (φ, p) .

Nous avons représenté sur la figure 2.10 les noyaux d'erreur associés à l'approximation spline par projection orthogonale ($\tilde{\varphi} = \varphi_d$) et par interpolation ($\tilde{\varphi} = \delta$). Nous avons aussi représenté le noyau d'erreur associé à l'interpolation cubique (2.24) obtenue avec $\varphi(t) = 3\beta^3(t) - \beta^2(t + \frac{1}{2}) - \beta^2(t - \frac{1}{2})$.

Comme on l'a déjà mentionné à la section 2.3.1, la qualité de reconstruction est directement liée à l'ordre d'approximation de φ (égal à $n + 1$ pour la B-spline β^n , et à 3 pour le générateur cubique). Il s'agit donc du premier paramètre à prendre en compte lors du choix de φ . On peut voir aussi que l'écart entre l'interpolation et la projection orthogonale n'est pas négligeable; en dehors de la zone de Nyquist, le contenu spectral de s se replie et entraîne la présence d'*aliasing* dans le signal reconstruit, et ce pour toutes les méthodes de reconstruction. En l'absence de fonction d'analyse, on est en effet très éloigné des conditions d'échantillonnage idéal de Shannon.

Dans la suite de ce chapitre, nous allons concevoir des préfiltres exploitant au mieux les propriétés d'approximation de l'espace de reconstruction $V_T(\varphi)$. L'idée sous-jacente

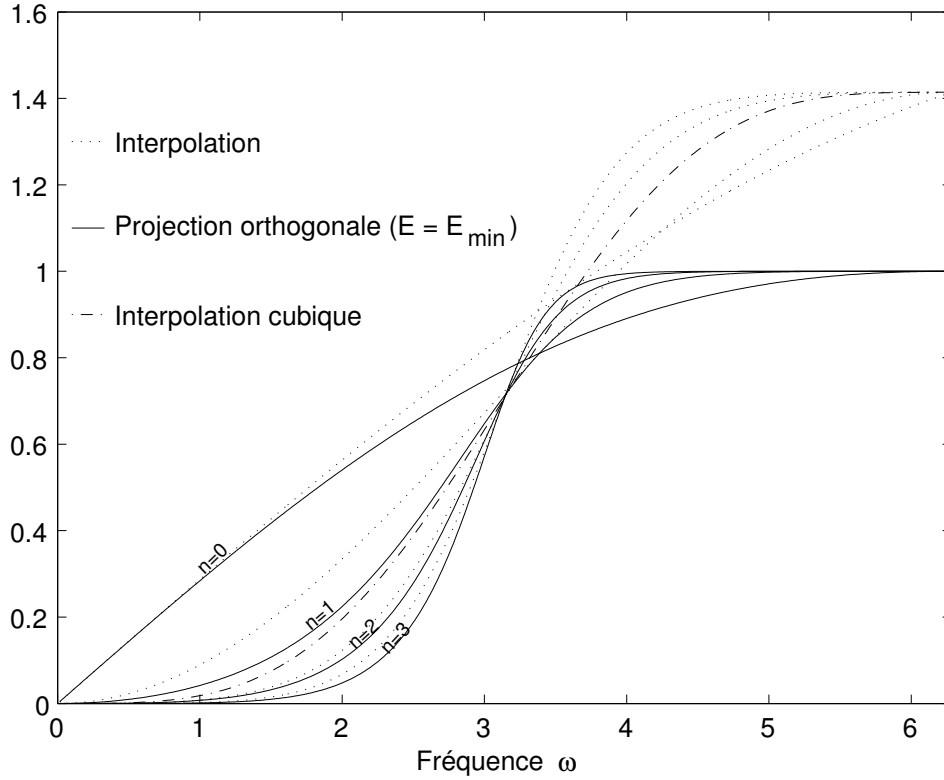


FIG. 2.10 : $\sqrt{E(\omega)}$ pour la projection orthogonale et l'interpolation spline de degré $n = 0$ à 3 (c.-à-d. $\varphi = \beta^n$) et l'interpolation cubique avec le générateur (2.24).

consiste à minimiser le noyau d'erreur $E_{\text{res}}(\omega)$.

2.5 Reconstruction asymptotiquement L_2 -optimale

Nous avons vu à la section précédente que la reconstruction optimale au sens L_2 ne pouvait être atteinte directement. Bien que l'on puisse contourner la difficulté en modifiant le critère L_2 afin de le rendre indépendant de s , en se plaçant dans le paradigme minimax, cette solution ne nous semble pas totalement satisfaisante. Minimiser une erreur au pire cas n'implique pas, en effet, que la qualité de reconstruction soit bonne en règle générale. Indépendamment de cette étude, nous avons présenté le noyau d'erreur, qui fournit un excellent moyen d'évaluer la qualité de telle ou telle méthode de reconstruction. Plus le noyau est faible, meilleure est la qualité de reconstruction, et cette analyse peut être faite à une fréquence fixée ou sur une plage de fréquences.

Dans cette section, on se place encore dans le cas non bruité. Le cas bruité sera abordé à la section 2.7. Nous présentons une stratégie nouvelle pour la reconstruction, adaptée aux signaux essentiellement passe-bas, c'est-à-dire ayant leur énergie localisée dans les basses fréquences. Cela est le cas pour de nombreux signaux acquis depuis le monde environ-

nant, et en particulier pour les images naturelles, qui ont un comportement en $1/f$ avéré [214, 182, 183]. Nous allons donc chercher un préfiltre tel que le noyau d'erreur $E(\omega)$ associé à la reconstruction soit aussi proche que possible de zéro, en particulier au voisinage de $\omega = 0$.

Précisons d'abord les propriétés asymptotiques d'approximation de la méthode de reconstruction $\mathcal{M} : v \mapsto f_T$ définie par le couple (φ, p) . Dans la suite, on note L l'ordre d'approximation de φ . Un développement de Taylor dans (2.61) donne, pour une fonction $s \in W_2^r$,

$$\|s - f_T\|_{L_2}^2 = \sum_{k=0}^r \frac{E^{(2k)}(0)}{(2k)!} \|s^{(k)}\|_{L_2}^2 T^{2k} + o(T^{2r}). \quad (2.66)$$

On va donc chercher les conditions sur p permettant d'avoir la vitesse de convergence maximale, ou de manière équivalente le noyau d'erreur le plus plat possible au voisinage de l'origine, c'est-à-dire

$$\boxed{E(\omega) \sim E_{\min}(\omega)}, \quad (2.67)$$

puisque alors, en supposant que $r \geq L$, on a la vitesse de convergence maximale

$$\|s - f_T\|_{L_2} \sim \|s - \mathcal{P}_T s\|_{L_2} \sim C_{\min} \|s^{(L)}\|_{L_2} T^L \quad \text{lorsque } T \rightarrow 0. \quad (2.68)$$

On voit donc pourquoi l'ordre d'approximation (dont le nom se trouve ici justifié) de φ est primordial : c'est lui qui détermine la vitesse de convergence en $O(T^L)$ de l'erreur d'approximation (si p est correctement choisi). En raisonnant sur le noyau d'erreur, on constate que

$$\sqrt{E_{\min}(\omega)} \sim C_{\min} \omega^L. \quad (2.69)$$

On doit donc choisir p afin que $E(\omega)$ ait ce même développement de Taylor. Dans l'idéal, on aimerait avoir $E(\omega) = E_{\min}(\omega)$ sur toute la bande de Nyquist $]-\pi, \pi[$. Cela est possible, en choisissant le préfiltre idéal

$$\hat{p}^{\text{id}}(\omega) = \frac{\hat{\varphi}_d(\omega)}{\hat{\hat{\varphi}}(\omega)} \quad \forall \omega \in [-\pi, \pi]. \quad (2.70)$$

En particulier, ce filtre assure que $f_T = \mathcal{P}_T s$ dans le cas où s est à bande limitée dans $]-\frac{\pi}{T}, \frac{\pi}{T}[$. Le principal inconvénient de ce filtre est de ne pas être implémentable en domaine spatial. C'est pourquoi on préfère chercher un filtre rationnel, de la forme $P(z) = H_1(z)/H_2(z)$, où h_1 et h_2 sont des filtres RIF, car on dispose d'algorithmes rapides pour implémenter le préfiltrage avec un tel filtre [218].

La contrainte (2.67) est équivalente à l'égalité suivante, où N est un entier supérieur ou égal à $L + 1$:

$$\hat{p}(\omega) = \frac{\hat{\varphi}_d(\omega)}{\hat{\hat{\varphi}}(\omega)} + O(\omega^N). \quad (2.71)$$

Dès lors que $N \geq L$ dans (2.71), le processus d'approximation $\mathcal{Q} : s \mapsto f_T$ est un **quasi-projecteur** d'ordre L dans $V_T(\varphi)$ [85, 217], c'est-à-dire que $\mathcal{Q}s = s$ si s est un polynôme de degré au plus $L - 1$. Ainsi un quasi-projecteur n'effectue une projection que pour un espace restreint de polynômes, et non pour toute fonction. Alors que $\mathcal{Q}$ est un projecteur si et seulement si $\tilde{\varphi}$ et φ_{eq} , définie par (2.38), sont biorthogonales (2.40), $\mathcal{Q}$ est un quasi-projecteur d'ordre L si et seulement si $\tilde{\varphi}$ et φ_{eq} sont *quasi-biorthogonales* d'ordre L [36], c'est-à-dire

$$\hat{\tilde{\varphi}}(\omega)\hat{\varphi}_{\text{eq}}(\omega + 2k\pi) = \delta_{k,0} + O(\omega^L) \quad \forall k \in \mathbb{Z}. \quad (2.72)$$

Dans le cas particulier $\tilde{\varphi} = \delta$, on parle de *quasi-interpolation* au lieu de quasi-projection : on dit que φ_{eq} est un quasi-interpolateur d'ordre L si

$$\hat{\varphi}_{\text{eq}}(\omega + 2k\pi) = \delta_{k,0} + O(\omega^L) \quad \forall k \in \mathbb{Z}. \quad (2.73)$$

Dans ce cas, $\mathcal{Q}s = s$ si les $v[k]$ sont les valeurs ponctuelles en les Tk d'un polynôme de degré au plus $L - 1$.

On cherche donc un préfiltre réalisant une quasi-projection d'ordre L dans $V_T(\varphi)$, afin d'avoir la vitesse de convergence de l'erreur en $O(T^L)$ dans (2.68). La reconstruction consistante réalise une projection oblique, et donc une quasi-projection d'ordre L , de s dans $V_T(\varphi)$. La quasi-projection est en effet une contrainte moins forte que la projection. Cependant, on n'impose pas seulement la propriété de quasi-projection, équivalente à $N = L$ dans (2.71), mais $N \geq L + 1$, ce qui n'est généralement pas vérifié pour la solution consistante. Autrement dit, on ne veut pas seulement une vitesse de convergence en $C(T^L)$, mais on demande aussi à avoir la constante asymptotique optimale $C = C_{\min}$, afin d'avoir la convergence optimale de f_T vers s au sens L_2 lorsque $T \rightarrow 0$. C'est pourquoi on dira du préfiltre p , si (2.67) est vérifiée, qu'il est asymptotiquement optimal au sens L_2 .

2.5.1 Conception de préfiltres quasi-interpolants

Etudions en détail la conception de préfiltres quasi-interpolants optimaux ($\tilde{\varphi} = \delta$), en nous intéressant plus particulièrement à la reconstruction spline, *i.e* $\varphi = \beta^n$. On cherche un préfiltre p de faible complexité, vérifiant (2.71) avec $N \geq L + 1 = n + 2$. Thierry Blu a déterminé les filtre RIF de taille minimale ($N = L + 1$) vérifiant cette condition d'optimalité [37]. Nous allons montrer que des filtres inverses, c'est-à-dire de la forme $P(z) = 1/Q(z)$ avec q un filtre RIF, donnent de meilleurs résultats. Cela est intuitif, puisque l'on cherche $\hat{p}(\omega)$ approchant $\hat{\varphi}_d(\omega) \approx 1/\hat{\varphi}(\omega)$, si φ est passe-bas.

Cherchons donc un filtre RIF q tel que $p = q^{-1}$ vérifie (2.71). En substituant les conditions de Strang-Fix (2.29) dans la formule de Poisson (1.7), nous obtenons

$$\hat{a}_\varphi(\omega) = |\hat{\varphi}(\omega)|^2 + O(\omega^{2L}). \quad (2.74)$$

Par conséquent, si $N \leq 2L$, (2.71) est équivalente à

$$\frac{1}{\hat{p}(\omega)} = \hat{\varphi}(\omega) + O(\omega^N). \quad (2.75)$$

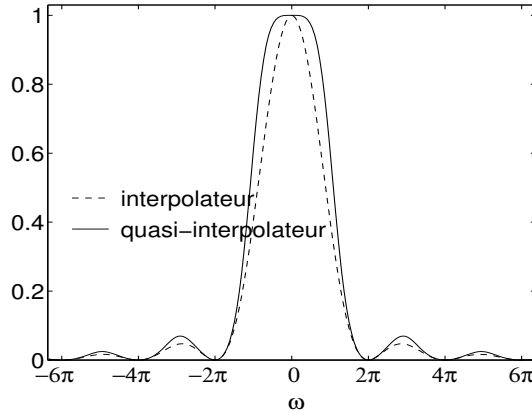


FIG. 2.11 : spectre $\hat{\varphi}_{\text{eq}}(\omega)$ du quasi-interpolateur correspondant à $\varphi = \beta^1$ combinée au préfiltre du tableau 2.1.

Cette dernière égalité simplifie encore la conception du filtre inverse $q = p^{-1}$, puisqu'il suffit de choisir ce filtre de telle manière que son développement de Taylor soit égal à celui de φ jusqu'à l'ordre désiré. Seul le calcul de ce développement de Taylor peut nécessiter l'usage d'un logiciel de calcul formel. Notons qu'il n'est pas nécessaire de connaître l'autocorrélation a_φ pour concevoir le préfiltre.

Détaillons le calcul du préfiltre dans le cas d'une reconstruction linéaire par morceaux, obtenue avec $\varphi = \beta^1$, d'ordre $L = 2$. φ étant symétrique, on cherche un préfiltre symétrique afin que le processus d'approximation soit lui aussi symétrique (c.-à-d. $\overline{Qs} = Q\bar{s}$). On a le développement de Taylor :

$$\hat{\varphi}(\omega) = \text{sinc}(\omega/2)^2 = 1 - \frac{1}{12}\omega^2 + O(\omega^4). \quad (2.76)$$

Un filtre q symétrique de taille 3, de la forme $[\frac{b}{2}, a, \frac{b}{2}]$, vérifie $\hat{q}(\omega) = a + b\cos(\omega) = a + b(1 - \frac{1}{2}\omega^2 + O(\omega^4))$. En imposant $\hat{q}(\omega) = \hat{\varphi}(\omega) + O(\omega^4)$, on obtient par identification $a = 1 - b$ et $b = \frac{1}{6}$, ce qui donne le filtre

$$Q(z) = \frac{1}{12}z + \frac{5}{6} + \frac{1}{12}z^{-1}. \quad (2.77)$$

La figure (2.11) montre le spectre du quasi-interpolateur φ_{eq} correspondant (voir (2.38)). On voit qu'il est plus proche du filtre passe-bas idéal que l'interpolateur β^1 . Cela indique que la reconstruction d'image avec la méthode de quasi-interpolation proposée introduit nettement moins de flou que l'interpolation correspondante, comme nous allons le confirmer expérimentalement par la suite. Notons que tous les quasi-interpolateurs splines que nous proposons sont passifs, c'est-à-dire que $|\hat{\varphi}_{\text{eq}}(\omega)| \leq 1$. Il n'y a donc jamais de réhaussement artificiel du contenu fréquentiel lors de la reconstruction.

$\varphi(x)$	$Q(z) = 1/P(z)$
$\beta^0(x)$	$-\frac{1}{24}z + \frac{13}{12} - \frac{1}{24}z^{-1}$
$\beta^1(x)$	$\frac{1}{12}z + \frac{5}{6} + \frac{1}{12}z^{-1}$
$\beta^2(x)$	$-\frac{7}{1920}z^2 + \frac{67}{480}z + \frac{233}{320} + \frac{67}{480}z^{-1} - \frac{7}{1920}z^{-2}$
$\beta^3(x)$	$-\frac{1}{720}z^2 + \frac{31}{180}z + \frac{79}{120} + \frac{31}{180}z^{-1} - \frac{1}{720}z^{-2}$

TAB. 2.1 – Préfiltres inverses de la forme $p = q^{-1}$ proposés pour la quasi-interpolation spline.

La conception de préfiltres quasi-interpolants est donc particulièrement aisée. Nous donnons dans le tableau 2.1 les préfiltres symétriques obtenus pour la quasi-interpolation spline de degré 0 à 3. Lorsque n est pair, l'interpolation vérifie déjà la contrainte d'optimalité (2.67). On a donc choisi $N = L + 2$ lorsque n est pair, et $N = L + 1$ lorsque n est impair, dans (2.71), afin de proposer les préfiltres les plus courts, tout en étant différents des préfiltres d'interpolation. On peut vérifier, en visualisant sur la figure 2.12 les noyaux d'erreur associés aux préfiltres, que la qualité d'approximation est bien meilleure que celle obtenue par interpolation. Le noyau $E(\omega)$ associé à la quasi-interpolation avec nos préfiltres est non seulement optimal à l'origine, mais se révèle être très proche de $E_{\min}$ sur toute la bande de Nyquist, ce qui valide notre approche. D'ailleurs, il apparaît inutile d'augmenter la valeur de N dans (2.71) car, du moins pour la quasi-interpolation spline, la plupart du gain potentiel par rapport à l'interpolation est déjà atteint avec les préfiltres que nous avons proposés.

2.6 Application : opérations géométriques sur les signaux discrets

Il est courant d'avoir à réaliser une transformation géométrique sur un signal ou une image, par exemple une translation ou une rotation. Cependant, ces opérations, qui sont bien définies dans le domaine continu, ne le sont pas dans le domaine discret. Ainsi, la rotation d'images est une opération qui, à partir d'une image discrète, fournit une autre image discrète, mais la relation entre les pixels de ces deux images n'est pas explicite *a priori*. Il est possible d'adopter un point de vue purement discret et d'utiliser des notions de géométrie discrète pour formaliser les opérations géométriques sur les images (voir par exemple [169] pour un exposé de ce type d'approches).

Nous préférons adopter une démarche différente, mettant à profit les liens entre les mondes discret et continu qui sont apparus en arrière-plan de ce travail sur la reconstruction. Ainsi, puisque nous nous sommes donnés le modèle (1.10) pour définir le signal initial $v = (v[k])$, il nous faut aussi proposer un modèle pour le signal transformé $\mathcal{T}v$, qui soit cohérent avec le processus de formation de v , et donc avec le signal sous-jacent s . Nous pro-

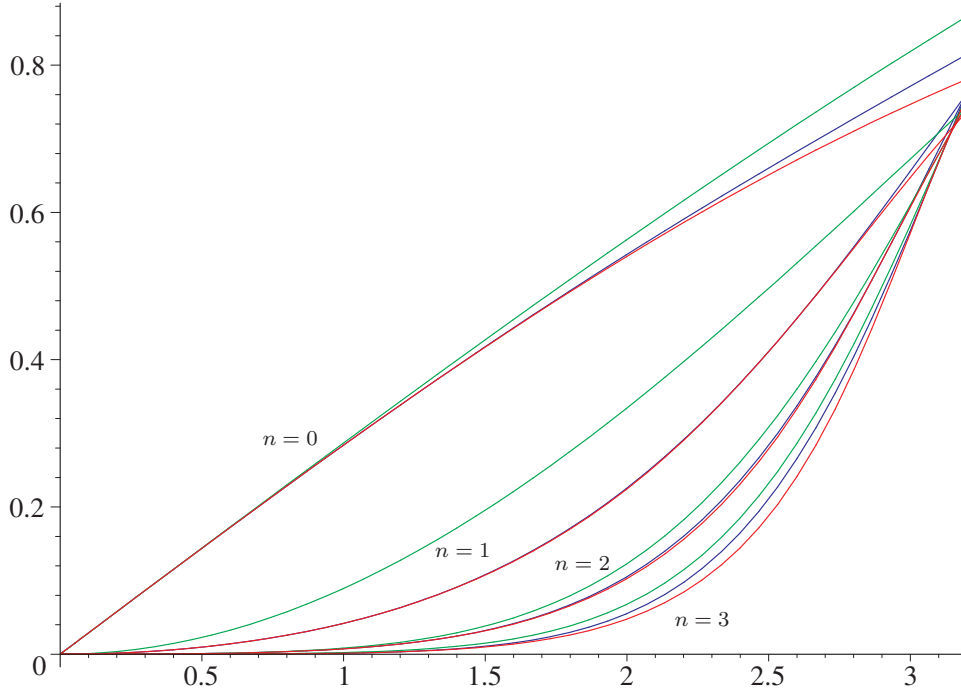


FIG. 2.12 : Noyaux d'erreur associés à la projection orthogonale (en rouge), l'interpolation (en vert), et la quasi-interpolation (en bleu) spline de degré $n = 0$ à 3.

posons un formalisme générique permettant de définir rigoureusement de telles opérations géométriques, et de proposer des solutions concrètes pour les mettre en œuvre.

Nous cherchons à définir un opérateur idéal discret $\mathcal{T}^{\text{id}} : \ell_2 \rightarrow \ell_2$ (dit « idéal », car c'est l'opérateur dont on aimerait disposer mais qu'en pratique on ne pourra qu'approcher). Prenons l'exemple de la rotation d'image discrète d'angle θ . On suppose qu'il existe un opérateur $\mathfrak{T} : L_2 \rightarrow L_2$ effectuant la même transformation dans le domaine continu. Dans notre exemple, il s'agit de la rotation 2D d'angle θ définie par

$$\mathfrak{T}s(\mathbf{x}) = s(\mathbf{R}_{-\theta}\mathbf{x}) \text{ où } \mathbf{R}_{-\theta} = \begin{bmatrix} \cos(\theta) & \sin(\theta) \\ -\sin(\theta) & \cos(\theta) \end{bmatrix}. \quad (2.78)$$

On définit $\mathcal{T}v$ comme le signal discret qui aurait été obtenu si l'on avait fait porter l'opérateur $\mathfrak{T}$ sur s , avant l'acquisition, c'est-à-dire remplacé s par $\mathfrak{T}s$ dans notre modèle (1.10), et en l'absence de bruit. Appliquer un opérateur à un signal discret, c'est donc obtenir le signal qu'aurait produit le dispositif d'acquisition, en l'absence de bruit, si le processus sous-jacent avait subi l'équivalent continu de l'opérateur. En d'autres termes, $\mathcal{T}^{\text{id}}$ est défini par :

$$\mathcal{T}^{\text{id}} : v = \mathcal{D}s \mapsto \mathcal{D}_0 \mathfrak{T}s \quad \forall s \in L_2. \quad (2.79)$$

où $\mathcal{D}_0 : s \mapsto (s * \tilde{\varphi}_T(Tk))_{k \in \mathbb{Z}}$ est la version non bruitée de $\mathcal{D}$.

Les opérations géométriques les plus courantes sont les suivantes :

$$\text{rotation 2D d'angle } \theta \leftrightarrow \mathfrak{T} : s(\mathbf{x}) \mapsto s(\mathbf{R}_{-\theta}\mathbf{x}) \quad (2.80)$$

$$\text{translation de pas } \boldsymbol{\tau} \leftrightarrow \mathfrak{T} : s(\mathbf{x}) \mapsto s(\mathbf{x} - \boldsymbol{\tau}) \quad (2.81)$$

$$\text{réduction de facteur } \alpha \leftrightarrow \mathfrak{T} : s(\mathbf{x}) \mapsto s(\alpha\mathbf{x}) \quad (2.82)$$

$$\text{agrandissement de facteur } \alpha \leftrightarrow \mathfrak{T} : s(\mathbf{x}) \mapsto s(\mathbf{x}/\alpha) \quad (2.83)$$

Comme pour le problème de reconstruction, où l'on cherche un opérateur $\mathcal{M}$ proche de $\mathcal{M}^{\text{id}} : \mathcal{D}_s \mapsto s$, il nous faut proposer un opérateur $\mathcal{T}$ proche de $\mathcal{T}^{\text{id}}$. Une solution est d'estimer s à l'aide d'un opérateur de reconstruction $\mathcal{M}$, puis d'appliquer l'opérateur de discrétisation $\mathcal{D}_0$:

$$\mathcal{T} : v \mapsto \mathcal{D}_0 \mathfrak{T} \mathcal{M} v \quad (2.84)$$

Nous montrerons dans la seconde partie de cette thèse qu'il est possible de trouver des solutions plus satisfaisantes pour la réduction et l'agrandissement. Intéressons-nous pour l'instant à l'opération de rotation. On suppose dans ce paragraphe que le pas d'échantillonnage T vaut 1 afin de simplifier les notations. On cherche donc pour tout $\mathbf{k} \in \mathbb{Z}^2$,

$$\mathcal{T}^{\text{id}} v[\mathbf{k}] = (s(\mathbf{R}_{-\theta} \cdot) * \tilde{\varphi})(\mathbf{k}) \quad (2.85)$$

$$= (s * \tilde{\varphi}(\mathbf{R}_{\theta} \cdot))(\mathbf{R}_{-\theta} \mathbf{k}). \quad (2.86)$$

Notons $\tilde{s} = s * \tilde{\varphi}$, et supposons que $\tilde{\varphi}$ est isotrope, c.-à-d. $\tilde{\varphi}(\mathbf{R}_{\theta} \mathbf{x}) = \tilde{\varphi}(\mathbf{x})$ pour tous $\mathbf{x}$ et θ . Alors le problème de rotation revient à évaluer des échantillons de la fonction $\tilde{s}$ à de nouvelles positions $\mathbf{R}_{-\theta} \mathbf{k}$, connaissant les $v[\mathbf{k}] = \tilde{s}(\mathbf{k}) + \mu[\mathbf{k}]$, $\mathbf{k} \in \mathbb{Z}^2$. On est donc confronté à un problème de reconstruction de $\tilde{s}$: on va chercher à évaluer cette fonction connaissant ses échantillons bruités uniformes $v[\mathbf{k}]$, puis échantillonner la fonction reconstruite $\mathcal{M}v$ en de nouvelles positions pour former le signal transformé. Notons que la connaissance de $\tilde{\varphi}$ n'est pas requise, et qu'il s'agit d'un problème d'interpolation dans le cas non bruité. On définit donc notre opérateur de rotation par

$$\mathcal{T} v[\mathbf{k}] = \mathcal{M} v(\mathbf{R}_{-\theta} \mathbf{k}) \quad \forall \mathbf{k} \in \mathbb{Z}^2 \quad (2.87)$$

pour un certain opérateur de reconstruction $\mathcal{M}$ adapté au cas $\tilde{\varphi} = \delta$ dans le modèle (1.10) (puisque l'on raisonne sur $\tilde{s}$ et non s pour effectuer la reconstruction).

Ainsi, on peut valider expérimentalement notre méthode de reconstruction par quasi-interpolation. Nous proposons l'expérience suivante : 17 rotations successives d'angle $2\pi/17$ sont effectuées sur plusieurs images (représentées dans l'annexe 7.7), en utilisant différentes méthodes de reconstruction. L'image finale est ensuite comparée à l'image initiale, puisque $\mathcal{T}^{\text{id}}$ itéré 17 fois est l'identité. Puisque l'on est en 2D, on effectue la reconstruction à l'aide de la fonction produit tensoriel $\varphi(\mathbf{x}) = \varphi(x_1)\varphi(x_2)$ et du préfiltre $P(\mathbf{z}) = P(z_1)P(z_2)$ appliqué de manière séparable le long des lignes et colonnes de l'image. Nous justifierons dans le chapitre 3 que les préfiltres asymptotiquement optimaux en 1D que nous avons développés

φ p	β^1 int.	β^1 int. shift.	β^1 quasi. RIF	β^1 quasi. RII	bicubique int.	β^3 int.	β^3 quasi. RII
L	2	2	2	2	3	4	4
<i>Lena</i>	29.40	34.87	35.47	36.81	35.29	38.69	39.81
<i>Barbara</i>	23.84	25.96	26.13	27.32	25.95	28.99	30.55
<i>Baboon</i>	21.98	24.40	25.17	26.23	25.01	27.63	28.64
<i>Lighthouse</i>	22.88	27.56	27.57	29.04	27.35	31.04	32.57
<i>Goldhill</i>	28.35	31.65	32.63	33.73	32.47	35.24	36.31
<i>Boat</i>	26.33	30.59	31.37	32.52	31.21	34.07	35.00
<i>Camera</i>	23.29	26.54	27.49	28.64	27.33	30.18	31.23
<i>Peppers</i>	28.89	33.08	34.36	35.41	34.22	36.83	37.64
temps	1 U	1.05 U	1.1 U	1.1 U	2.3 U	2.6 U	2.7 U

TAB. 2.2 – PSNR après 17 rotations d’angle $2\pi/17$ opérées sur les images de l’annexe 7.7, en utilisant différentes méthodes de reconstruction par interpolation et quasi-interpolation. Nous avons rappelé l’ordre d’approximation L de chaque méthode, et mentionné le temps de calcul moyen pour effectuer une rotation, en unité normalisée relative au temps requis pour l’interpolation bilinéaire.

ont un produit tensoriel lui aussi asymptotiquement optimal en 2D.

On définit la distance PSNR entre deux images u et v de taille $N \times N$ par :

$$\text{PSNR}(u, v) = 10 \log_{10} \left(\frac{255^2 N^2}{\sum_{\mathbf{k} \in [1, N]^2} |u[\mathbf{k}] - v[\mathbf{k}]|^2} \right). \quad (2.88)$$

Le tableau 2.2 fournit les valeurs PSNR obtenues, et le temps de calcul moyen d’une opération de rotation. Les différentes méthodes comparées sont :

- l’interpolation bilinéaire ($\varphi = \beta^1$, $P(z) = 1$)
- l’interpolation spline cubique ($\varphi = \beta^3$, $1/P(z) = \frac{1}{6}z + \frac{2}{3} + \frac{1}{6}z^{-1}$)
- la quasi-interpolation spline linéaire et cubique, obtenue avec les préfiltres du tableau 2.1 (notée « quasi. RII »)
- la quasi-interpolation spline linéaire, notée « quasi. RIF », obtenue avec le préfiltre RIF proposé par Thierry Blu [37] : $P(z) = \frac{-1}{12}z + \frac{7}{6} + \frac{-1}{12}z^{-1}$
- l’interpolation bicubique (φ donnée par (2.24), $P(z) = 1$)
- l’interpolation bilinéaire *shiftée*, une méthode récente proposée par Thierry Blu en 2004 [34] : $\varphi(t) = \beta^1(t - \tau)$ où $\tau = \frac{1}{2}(1 - \frac{\sqrt{3}}{3})$, $P(z) = 1/(1 - \tau + \tau z^{-1})$.

Comme on peut le voir, numériquement dans le tableau 2.2 et visuellement sur la figure 2.13, les méthodes de reconstruction par quasi-interpolation fournissent une amélioration significative de la qualité de reconstruction, par rapport à leurs équivalents interpolants. Nous avons indiqué en bas du tableau 2.2 les temps de calcul moyens d’une opération de rotation, pour chaque méthode. On voit que l’étape de préfiltrage n’est pas la plus coûteuse

du processus de reconstruction, la majorité du temps de calcul étant dévolue au rééchantillonnage du modèle reconstruit $f_T(\mathbf{x})$. L'usage de méthodes de quasi-interpolation n'entraîne donc qu'un accroissement marginal du temps de calcul. Ainsi, la quasi-interpolation bilinéaire donne de meilleurs résultats que l'interpolation bicubique, pour un temps de calcul deux fois moindre, ce qui est un résultat important. Ces expériences valident l'étude théorique des propriétés d'approximation des filtres, et notre conception de filtres asymptotiquement optimaux. De plus, nos nouveaux préfiltres RII sont plus performants que les préfiltres RIF proposés dans [37], pour un temps de calcul identique.

La figure 2.14 montre les images obtenues après plusieurs rotations d'angle $2\pi/5$. Les distorsions s'accumulent à chaque rotation, donc on peut estimer que la différence entre l'image obtenue après une seule rotation et l'image inconnue que l'on cherche à estimer est ici égale à $1/5$ de la différence entre l'image de référence et l'image obtenue après 5 rotations successives. La faible distorsion occasionnée par la quasi-interpolation bilinéaire avec notre filtre est donc remarquable.



image initiale



interpolation bicubique



interpolation bilinéaire



quasi-interpolation bilinéaire



interpolation spline cubique



quasi-interpolation spline cubique

FIG. 2.13 : Extraits de l'image Boat après 17 rotations d'angle $2\pi/17$ effectuées avec différentes méthodes d'interpolation et quasi-interpolation.

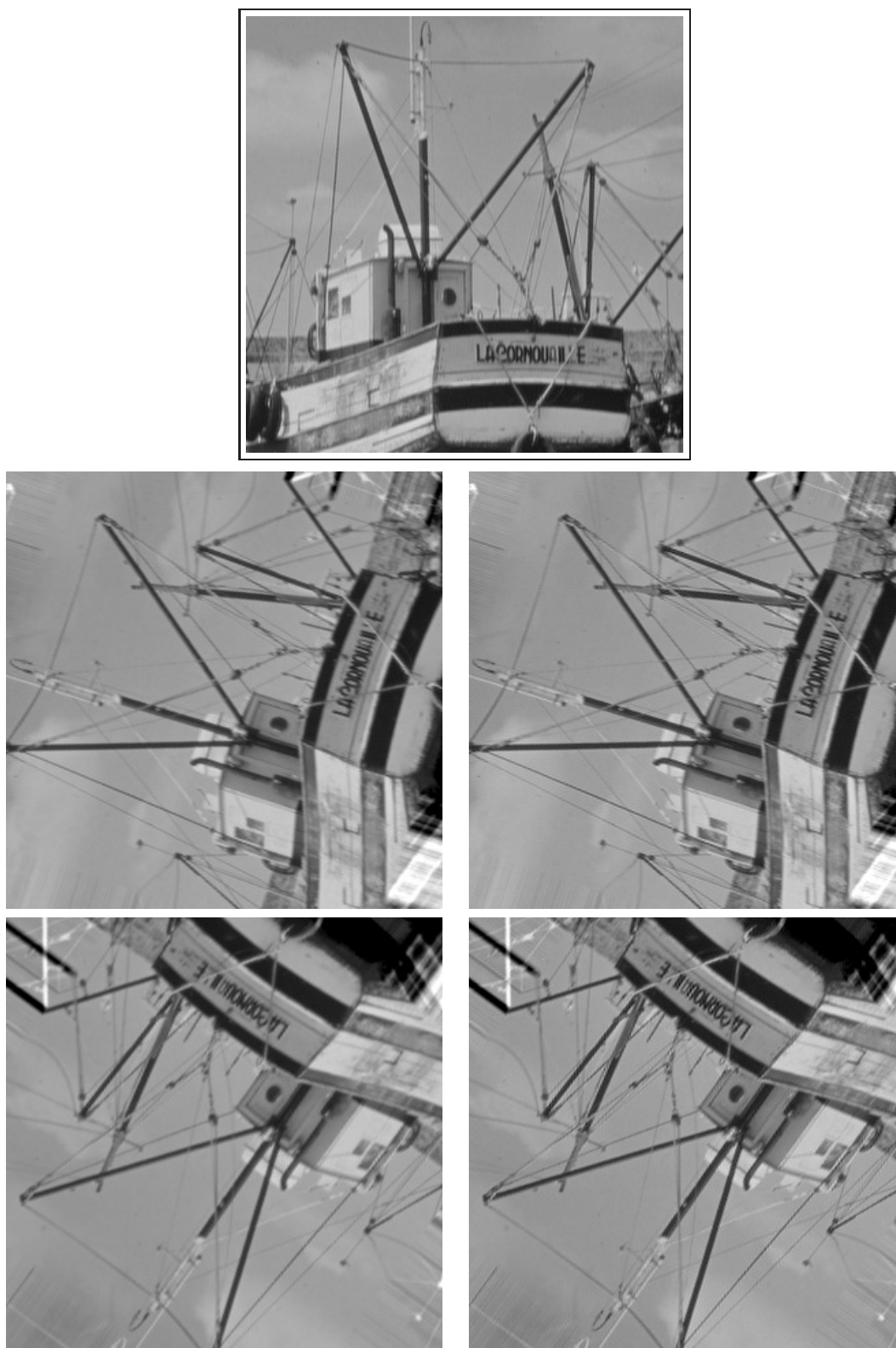


FIG. 2.14 : Résultats obtenus après rotations successives d'angle $2\pi/5$ sur une partie de l'image Boat (montrée en haut). À gauche : par interpolation bilinéaire. À droite : par quasi-interpolation bilinéaire. (Figure continuée sur la page suivante)

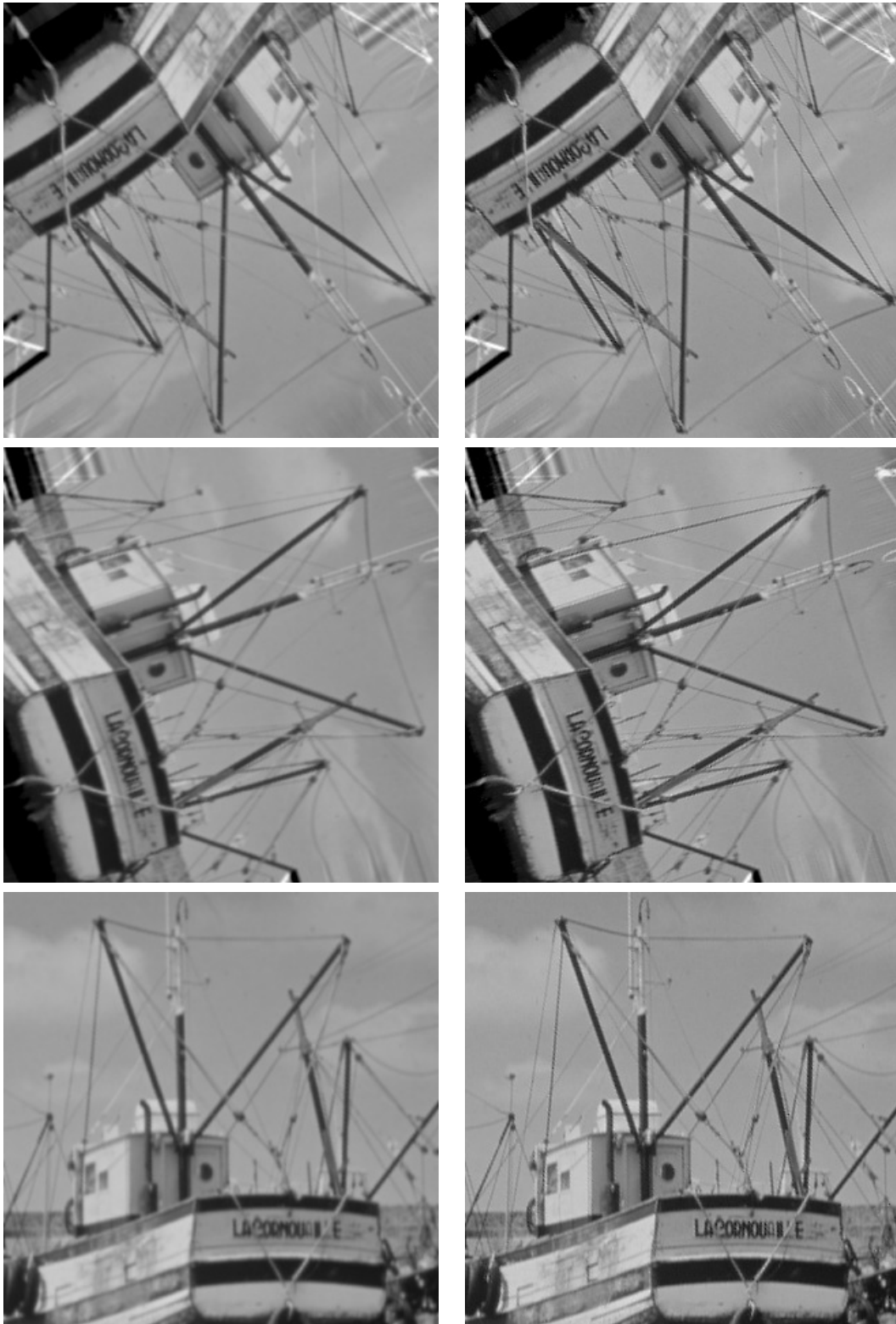


FIG. 2.14 : Suite. Résultats après 3, 4, 5 rotations successives. Le PSNR entre les images de la ligne du bas et l'image initiale est 31.17 à gauche, 37.26 à droite.

2.7 Estimation dans un espace LSI en présence de bruit

Dans les sections précédentes, nous nous sommes concentrés sur la reconstruction au sens L_2 , en l'absence de bruit. Dans le cas où les mesures sont contaminées par un bruit stationnaire μ , le signal v n'appartient plus à ℓ_2 , et donc la fonction reconstruite f_T n'appartient pas à L_2 . On ne peut plus dès lors définir l'erreur de reconstruction comme la distance $\|s - f_T\|_{L_2}$, et il nous faut trouver un nouveau critère généralisant celui-ci.

Intéressons-nous tout d'abord à la situation déterministe dans laquelle $s \in L_2$. Une extension naturelle du critère L_2 est l'erreur quadratique moyennée définie par (2.63). Cependant, seul le bruit contribue à ce critère : $\varepsilon_{\text{STO}}(f_T)$ vaut 0 en l'absence de bruit, quel que soit le préfiltre choisi pour la reconstruction, car alors $|f_T(t) - s(t)|^2$ est de somme sur $\mathbb{R}$ finie, donc de moyenne sur $\mathbb{R}$ nulle. Dans le cas bruité, la solution minimisant ce critère est donc la fonction nulle obtenue avec le préfiltre nul. Eldar et Unser ont proposé de modifier ce critère afin d'obtenir une solution non triviale [95]. Faisons l'hypothèse que $s \in \Omega = \{g \mid \|\mathcal{L}g\| \leq C\}$ pour une certaine constante C . Alors la minimisation des deux critères suivants :

$$\varepsilon_{\text{HYB1}}(g) = \max_{f \in \Omega} \mathcal{E}\{|g(t_0) - \mathcal{P}_T f(t_0)|^2\} \quad (2.89)$$

$$\varepsilon_{\text{HYB2}}(g) = \lim_{A \rightarrow +\infty} \frac{1}{2A} \int_{-A}^A \max_{f \in \Omega} \mathcal{E}\{|g(t) - f(t)|^2\} dt \quad (2.90)$$

fournit la même solution, indépendante de t_0 , et obtenue avec le préfiltre [95] :

$$\hat{p}(\omega) = \frac{\sum_{k \in \mathbb{Z}} \hat{\varphi}(\omega + 2k\pi)^* \hat{\varphi}_d(\omega + 2k\pi) / |\hat{L}(\omega + 2k\pi)|^2}{\hat{c}_\mu(\omega) / C^2 + \sum_{k \in \mathbb{Z}} |\hat{\varphi}(\omega + 2k\pi)|^2 / |\hat{L}(\omega + 2k\pi)|^2}. \quad (2.91)$$

Notons que dans le cas non bruité, on retrouve le préfiltre (2.56) optimal au sens minimax. De plus, si la seule connaissance que l'on a est que $s \in L_2$, donc que l'on fait tendre C vers l'infini, on retrouve aussi ce même préfiltre. On peut donc objecter aux critères $\varepsilon_{\text{HYB1}}$ et $\varepsilon_{\text{HYB2}}$ de ne pas être naturels, et d'être trop conservatifs : le préfiltre solution dépend de la constante C caractérisant l'*a priori* que l'on fait sur s , et f_T tend vers la solution minimax lorsque $C \rightarrow +\infty$, autrement dit une solution qui ne dépend pas du niveau de bruit.

En fait, cette situation hybride mi-déterministe, mi-stochastique n'est pas naturelle, et il est malaisé de formuler la reconstruction dans ce cadre comme un problème d'optimisation ayant une solution non triviale satisfaisante. Il semble plus judicieux de se placer dans le cadre pleinement stochastique, dans lequel on suppose que le signal s comme le bruit sont des réalisations de processus stochastiques stationnaires. Il est alors naturel de formuler l'erreur de reconstruction comme l'erreur quadratique moyennée (2.64). Avant d'explicitier le préfiltre minimisant ce critère, nous allons généraliser la formule (2.65) qui caractérise l'erreur d'approximation dans le cas non bruité.

PROPOSITION : Soit $f_T \in V_T(\varphi)$ la fonction reconstruite dont les coefficients $c[k]$ sont obtenus à partir du signal v par convolution avec le préfiltre p . Alors l'erreur $\varepsilon_{\text{STO}}(f_T)$ s'écrit exactement au moyen du noyau d'erreur fréquentiel, et d'un second noyau d'erreur de bruit :

$$\frac{1}{T} \int_0^T \mathcal{E}\{|f_T(t) - s(t)|^2\} dt = \frac{1}{2\pi} \int_{\mathbb{R}} \hat{c}_s(\omega) E(T\omega) + T \hat{c}_\mu(T\omega) E_b(T\omega) d\omega, \quad (2.92)$$

où nous définissons le *noyau d'erreur de bruit* :

$$E_b(\omega) = |\hat{p}(\omega)|^2 |\hat{\varphi}(\omega)|^2. \quad (2.93)$$

PREUVE : Notons $f_{T,0}$ la fonction qui aurait été reconstruite à partir des données non bruitées $v_0[k]$ (rappelons que $v = v_0 + \mu$). Nous utilisons aussi la fonction φ_{eq} définie par (2.38). On peut écrire :

$$\frac{1}{T} \int_0^T \mathcal{E}\{|f_T(t) - s(t)|^2\} dt = \frac{1}{T} \int_0^T \mathcal{E}\{|f_{T,0}(t) - s(t)|^2\} dt \quad (2.94)$$

$$- \frac{2}{T} \int_0^T \mathcal{E}\left\{s(t) \sum_{k \in \mathbb{Z}} \mu[k] \varphi_{\text{eq}}\left(\frac{t}{T} - k\right)\right\} dt \quad (2.95)$$

$$+ \frac{2}{T} \int_0^T \mathcal{E}\left\{f_{T,0}(t) \sum_{k \in \mathbb{Z}} \mu[k] \varphi_{\text{eq}}\left(\frac{t}{T} - k\right)\right\} dt \quad (2.96)$$

$$+ \frac{1}{T} \int_0^T \mathcal{E}\left\{\left|\sum_{k \in \mathbb{Z}} \mu[k] \varphi_{\text{eq}}\left(\frac{t}{T} - k\right)\right|^2\right\} dt. \quad (2.97)$$

Les termes (2.95) et (2.96) sont nuls car μ et s sont indépendants l'un de l'autre. Le terme (2.94) caractérise l'erreur sur le signal, indépendamment du bruit, et s'écrit avec le noyau d'erreur fréquentiel comme dans (2.65). Le dernier terme (2.97) caractérise l'erreur

occasionnée par le bruit, et fournit la seconde partie de (2.92) :

$$\begin{aligned} \frac{1}{T} \int_0^T \mathcal{E} \left\{ \left| \sum_{k \in \mathbb{Z}} \mu[k] \varphi_{\text{eq}}\left(\frac{t}{T} - k\right) \right|^2 \right\} dt &= \sum_{k, k' \in \mathbb{Z}} \mathcal{E} \{ \mu[k] \mu[k'] \} \int_0^T \varphi_{\text{eq}}\left(\frac{t}{T} - k\right) \varphi_{\text{eq}}\left(\frac{t}{T} - k'\right) \frac{dt}{T} \\ &= \sum_{k, k' \in \mathbb{Z}} c_\mu[k' - k] \int_k^{k+1} \varphi_{\text{eq}}(t) \varphi_{\text{eq}}(t - (k' - k)) dt \\ &= \sum_{l \in \mathbb{Z}} c_\mu[l] \int_{\mathbb{R}} \varphi_{\text{eq}}(t) \varphi_{\text{eq}}(t - l) dt \end{aligned} \quad (2.98)$$

$$= \sum_{l \in \mathbb{Z}} c_\mu[l] (p * \bar{p} * a_\varphi)[l] \quad (2.99)$$

$$= \frac{1}{2\pi} \int_0^{2\pi} \hat{c}_\mu(\omega) |\hat{p}(\omega)|^2 \hat{a}_\varphi(\omega) d\omega \quad (2.100)$$

$$= \frac{1}{2\pi} \int_{\mathbb{R}} \hat{c}_\mu(\omega) |\hat{p}(\omega)|^2 |\hat{\varphi}(\omega)|^2 d\omega \quad (2.101)$$

$$= \frac{1}{2\pi} \int_{\mathbb{R}} T \hat{c}_\mu(T\omega) E_b(T\omega) d\omega. \quad (2.102)$$

□

L'expression (2.92) permet de dissocier l'erreur intrinsèque à la méthode d'approximation de l'erreur due au bruit. Plutôt que faire la part des choses entre ces deux termes, il peut être intéressant de les recombinaer afin de faire apparaître le rôle des paramètres de la reconstruction. Faisons d'abord apparaître explicitement la dépendance de l'erreur par rapport à φ_{eq} :

$$\begin{aligned} \varepsilon_{\text{STO}}(g) &= \frac{1}{2\pi} \int_{\mathbb{R}} \hat{c}_s(\omega) - 2\hat{c}_s(\omega) \Re(\hat{\varphi}(T\omega) \hat{\varphi}_{\text{eq}}(T\omega)) + \hat{c}_s(\omega) |\hat{\varphi}(T\omega)|^2 \hat{a}_{\varphi_{\text{eq}}}(T\omega) \\ &\quad + T \hat{c}_\mu(T\omega) |\hat{\varphi}_{\text{eq}}(T\omega)|^2 d\omega \end{aligned} \quad (2.103)$$

$$\begin{aligned} &= \frac{1}{2\pi} \int_{\mathbb{R}} \hat{c}_s(\omega) - 2\Re(\hat{c}_s(\omega) \hat{\varphi}(T\omega) \hat{\varphi}_{\text{eq}}(T\omega)) \\ &\quad + |\hat{\varphi}_{\text{eq}}(T\omega)|^2 \left(\sum_{k \in \mathbb{Z}} \hat{c}_s(\omega + \frac{2k\pi}{T}) |\hat{\varphi}(T\omega + 2k\pi)|^2 + T \hat{c}_\mu(T\omega) \right) d\omega \end{aligned} \quad (2.104)$$

$$= \frac{1}{2\pi} \int_{\mathbb{R}} A(T\omega) |\hat{\varphi}_{\text{eq}}(T\omega) - B(T\omega)|^2 + C(\omega) d\omega, \quad (2.105)$$

où $A(\omega) = T \hat{c}_v(\omega) \geq 0$, $B(\omega)$, et $C(\omega)$ ne dépendent pas de φ ni de p . L'erreur est donc minimale lorsque $\hat{\varphi}_{\text{eq}}(\omega) = B(\omega)$ pour tout $\omega \in \mathbb{R}$, c'est-à-dire

$$\boxed{\hat{\varphi}_{\text{eq}}(\omega) = \frac{\frac{1}{T} \hat{c}_s(\frac{\omega}{T}) \hat{\varphi}(\omega)^*}{\sum_{k \in \mathbb{Z}} \frac{1}{T} \hat{c}_s(\frac{\omega + 2k\pi}{T}) |\hat{\varphi}(\omega + 2k\pi)|^2 + \hat{c}_\mu(\omega)}}. \quad (2.106)$$

On retrouve ainsi $\hat{\varphi}_{\text{eq}} = \hat{p} \hat{\varphi}$, avec φ et $p = c_v^{-1}$ donnés par (2.13) et (2.14). La minimisation de ε_{STO} fournit donc la même solution que la minimisation de $\varepsilon_{\text{STO}, t_0}$ pour tout t_0 , étudiée

à la section 2.2.3 ; ce résultat était attendu, car le critère ε_{STO} étant une version moyennée sur t_0 de $\varepsilon_{\text{STO},t_0}$, la solution optimale pour le premier critère l'est aussi pour le second. On retrouve ainsi, grâce au noyau d'erreur, le résultat montré dans [182], mais sous une forme plus faible (optimalité en moyenne sur t , alors que [182] montre l'optimalité ponctuelle partout). La solution f_T est donc une $\mathcal{L}^*\mathcal{L}$ -spline, en choisissant $\mathcal{L}$ comme l'opérateur de blanchiment de s . De plus, il s'agit de la $\mathcal{L}^*\mathcal{L}$ -spline consistante dans le cas non bruité.

Le problème qui nous intéresse est plutôt de trouver la fonction $f_T \in V_T(\varphi)$ minimisant ε_{STO} , pour φ fixée indépendamment de c_s . Réécrivant le critère ε_{STO} sous une autre forme, afin de faire apparaître explicitement le préfiltre p :

$$\begin{aligned} \varepsilon_{\text{STO}}(g) = & \frac{1}{2\pi} \int_0^{2\pi} \sum_{k \in \mathbb{Z}} \frac{1}{T} \hat{c}_s\left(\frac{\omega + 2k\pi}{T}\right) \\ & - 2\Re\left(\hat{p}(\omega) \sum_{k \in \mathbb{Z}} \frac{1}{T} \hat{c}_s\left(\frac{\omega + 2k\pi}{T}\right) \hat{\varphi}(\omega + 2k\pi) \hat{\varphi}(\omega + 2k\pi)\right) \\ & + |\hat{p}(\omega)|^2 \hat{a}_\varphi(\omega) \left(\sum_{k \in \mathbb{Z}} \frac{1}{T} \hat{c}_s\left(\frac{\omega + 2k\pi}{T}\right) |\hat{\varphi}(\omega + 2k\pi)|^2 + \hat{c}_\mu(\omega) \right) d\omega. \end{aligned} \quad (2.107)$$

On a donc l'intégrale d'une forme quadratique en $\hat{p}(\omega)$, minimale, lorsque pour tout $\omega \in \mathbb{R}$,

$$\boxed{\hat{p}(\omega) = \frac{\sum_{k \in \mathbb{Z}} \frac{1}{T} \hat{c}_s\left(\frac{\omega + 2k\pi}{T}\right) \hat{\varphi}(\omega + 2k\pi)^* \hat{\varphi}_d(\omega + 2k\pi)}{\sum_{k \in \mathbb{Z}} \frac{1}{T} \hat{c}_s\left(\frac{\omega + 2k\pi}{T}\right) |\hat{\varphi}(\omega + 2k\pi)|^2 + \hat{c}_\mu(\omega)}}. \quad (2.108)$$

En minimisant le critère

$$\varepsilon_{\text{REG},t_0}(g) = \mathcal{E}\{|g(t_0) - \mathcal{P}_T s(t_0)|^2\}, \quad (2.109)$$

Eldar et Unser ont montré que la solution ne dépend pas de t_0 [96], et le préfiltre qu'ils obtiennent est exactement celui de (2.108). Cependant, la projection orthogonale n'est pas définie dans le cas stochastique, et $\mathcal{P}_T s$ est donc défini par extension du cas déterministe, au moyen de (2.49).

D'autre part, si l'on cherche un préfiltre optimal de forme fixée, dépendant linéairement de paramètres à déterminer, par exemple un filtre RIF symétrique de taille 3 de la forme $\hat{p}(\omega) = u + v \cos(\omega)$, le critère ε_{STO} s'exprime comme une forme quadratique positive en les paramètres. On peut donc toujours expliciter les valeurs optimales des paramètres en fonction de $\tilde{\varphi}$, φ , c_s et c_μ .

L'approche asymptotiquement optimale présentée à la section 2.5, adaptée au cas où c_s n'est pas connue et $c_\mu = 0$, n'est pas transposable directement au cas bruité. On peut penser trouver les conditions telles que f_T soit un estimateur optimal de s au sens de l'erreur quadratique moyennée (2.63), lorsque s est un polynôme de degré au plus N .

Malheureusement, un polynôme croissant à l'infini, la notion de rapport signal-sur-bruit n'est pas bien définie pour un polynôme, et on trouvera alors des conditions définissant une quasi-projection, et ne dépendant pas du bruit.

Par contre, il est intéressant de remarquer que le préfiltre (2.108) peut être approché par le préfiltre suivant, avec égalité lorsque c_s , $\tilde{\varphi}$ ou φ sont à bandes limitées :

$$\hat{p}(\omega) = \frac{\hat{\varphi}_d(\omega) \sum_{k \in \mathbb{Z}} \frac{1}{T} \hat{c}_s\left(\frac{\omega+2k\pi}{T}\right) |\hat{\varphi}(\omega + 2k\pi)|^2}{\hat{\varphi}(\omega) \sum_{k \in \mathbb{Z}} \frac{1}{T} \hat{c}_s\left(\frac{\omega+2k\pi}{T}\right) |\hat{\varphi}(\omega + 2k\pi)|^2 + \hat{c}_\mu(\omega)} \quad (2.110)$$

$$= \frac{\hat{c}_{v_0}(\omega)}{\hat{c}_{v_0}(\omega) + \hat{c}_\mu(\omega)} \frac{\hat{\varphi}_d(\omega)}{\hat{\varphi}(\omega)}, \quad (2.111)$$

où l'on rappelle que v_0 est la version non bruitée de v . Ce filtre se décompose donc en deux filtres : le filtre de Wiener discret débruitant les mesures $v[k]$, puis le préfiltre p^{id} idéal dans le cas non bruité (2.70). Cela signifie que dans le cas bruité, il suffit de débruiter les échantillons dans un premier temps, puis d'effectuer la reconstruction comme dans le cas non bruité, à l'aide d'un préfiltre de quasi-projection optimale. De plus, pour ce faire, seule la densité spectrale $\hat{c}_\mu$ doit être connue, puisque le filtre de Wiener estime $\hat{c}_{v_0} = \hat{c}_v - \hat{c}_\mu$ directement à partir des mesures $v[k]$.

2.8 Conclusion

Dans ce chapitre, nous avons présenté le problème de reconstruction 1D d'un signal $f_T(t)$ défini continûment, à partir de données discrètes $(v[k])_{k \in \mathbb{Z}}$ interprétées comme des mesures uniformes, aux positions Tk , effectuées sur un signal inconnu $s(t)$. Nous avons formulé ce problème de modélisation comme un problème d'estimation de s , donc comme un problème mal posé, inverse de l'opération de discrétisation (formalisant le processus d'acquisition des données).

Les contributions nouvelles apportées dans ce chapitre peuvent être résumées comme suit :

- Nous avons dressé un panorama des méthodes existantes, et mis en lumière les relations qui existent entre elles, en les présentant dans un formalisme commun. Cela nous a permis par exemple d'étendre certains résultats minimax de Eldar et coll. dans la section 2.4.2, au moyen de la théorie de l'*optimal recovery*.

- Nous avons proposé une approche originale pour la reconstruction, basée sur l'approximation asymptotique optimale des basses fréquences du signal inconnu s . Cette idée avait été mentionnée par Thierry Blu dans son article [37] sur le noyau d'erreur fréquentiel. En l'exploitant plus en avant, nous avons proposé des préfiltres inverses particulièrement performants. Le noyau d'erreur fréquentiel s'est révélé être un outil de choix pour concevoir, au moyen de simples développements de Taylor, les préfiltres désirés. Nous avons illustré la pertinence de notre démarche par des expériences de rotations d'images, validant les

méthodes de quasi-projection obtenues, et montrant le gain significatif par rapport à l'approche consistante classique. De plus, ce gain n'est pas obtenu au détriment de la vitesse d'exécution. Ce travail a fait l'objet de la publication [65].

◦ Nous avons étendu la définition du noyau d'erreur, en lui adjoignant un terme quantifiant la perte de qualité due à la présence de bruit dans les données. L'approche classique, dans le cas bruité, est de régulariser la reconstruction consistante au moyen d'un terme variationnel. En fait, nous avons exhibé à la section 2.3.3 une solution plus satisfaisante, qui consiste à d'abord appliquer un filtre de Wiener débruitant les données, avant d'effectuer la reconstruction elle-même. En étudiant le noyau d'erreur bruité, nous avons vu que cette décomposition est en fait proche de la solution optimale, si la reconstruction est effectuée avec notre approche plutôt que de manière consistante (voir (2.111)). Si les caractéristiques spectrales du signal et du bruit sont connues, nous avons aussi déterminé le préfiltre fournissant la reconstruction optimale, au sens de l'espérance moyennée (formule (2.108)). Cette étude du cas bruité nécessiterait une validation expérimentale, non présentée ici. Nous avons cependant obtenu des résultats préliminaires prometteurs, en modélisant les images naturelles comme des processus de Matérn, suivant la méthode développée dans [214],[183].

Dans le chapitre suivant, nous étendons la problématique de la reconstruction au cas des signaux 2D uniformes, définis sur des treillis.

Chapitre 3

Reconstruction 2D par quasi-projections

3.1 Introduction

LA plupart des images, comme d'autres types de données 2D, sont définies et traitées sur la grille cartésienne (le treillis orthogonal), qui est le domaine discret régulier 2D le plus simple. Il existe cependant une infinité d'autres domaines discrets réguliers, appelés treillis, sur lesquels peuvent être localisés des signaux discrets uniformes. Parmi ceux-ci, le treillis hexagonal présente un intérêt certain. Ses avantages théoriques sont bien connus ; ainsi, l'échantillonnage hexagonal est optimal dans le cas de signaux 2D isotropes à bande limitée [156, 175]. Les propriétés d'isotropie de la grille hexagonale, comme la 12-symétrie ou la 6-connexité [142] ont montré leurs atouts dans de nombreux champs d'applications, comme la détection de contours [204, 244, 158] et la reconnaissance de formes [173, 101]. De plus, il existe des capteurs permettant d'acquérir des données directement sur ce type de grille [126, 216]. Cela justifie donc que l'on s'intéresse à la représentation des signaux 2D non pas seulement sur le treillis orthogonal, mais sous une forme plus générale. Afin de présenter des solutions concrètes, nous nous concentrerons sur le treillis hexagonal, mais les concepts et méthodes proposés pourront être étendus sans difficulté à un treillis quelconque. Simplement, le treillis hexagonal a certaines spécificités dont il est intéressant de tirer parti. Nous exhiberons en outre de nouvelles propriétés avantageuses du treillis hexagonal du point de vue de l'approximation.

Dans ce chapitre, nous allons donc traiter de la reconstruction de signaux à partir de mesures uniformes localisées sur un treillis 2D. L'étude du chapitre précédent peut être transposée aisément en 2D pour des images définies sur le treillis orthogonal, en appliquant les traitements de manière séparable, le long des lignes et colonnes de l'image. Cette simplicité explique certainement la popularité du treillis orthogonal. Nous montrerons que la mise en œuvre de traitements non séparables n'est pas si complexe que l'on pourrait le croire au premier abord, et permet en contrepartie de profiter pleinement des propriétés

d'approximation des treillis.

Nous avons insisté dans le chapitre précédent sur l'utilisation de splines, à la fois pour leur simplicité d'utilisation et leur optimalité du point de vue de l'approximation. Le concept de spline, reposant sur des fonctions polynomiales par morceaux, est suffisamment versatile pour être étendu dans un cadre multi-dimensionnel. Par exemple, les splines polyharmoniques apparaissent naturellement en 2D, car elles minimisent des critères de régularité 2D naturels, et sont optimales pour certains signaux en $1/f$ [182, 214]. La construction des B-splines peut être étendue de plusieurs manières pour la définition de splines sur un treillis arbitraire. Sur le treillis orthogonal, il suffit d'utiliser le produit tensoriel de B-splines 1D pour construire la B-spline séparable $\beta(\mathbf{x}) = \beta(x_1) \cdots \beta(x_d)$. Sur le treillis hexagonal, deux extensions majeures ont été proposées : les box-splines et les hex-splines. Ce sont ces deux familles de splines que nous étudierons pour la reconstruction 2D, en les déployant sur le treillis hexagonal.

Après avoir étudié les propriétés d'approximation de ces splines, qui tirent pleinement parti des avantages du treillis hexagonal, nous développerons des préfiltres quasi-interpolants optimaux permettant d'exploiter au mieux ses propriétés d'approximation. La faisabilité et l'efficacité de notre approche sont validées par l'étude du problème de rééchantillonnage du treillis hexagonal vers le treillis orthogonal.

3.1.1 Treillis et signaux 2D

Un *treillis* (dit aussi *réseau*) 2D $\Lambda_{\mathbf{R}}$ est un ensemble régulier, géométriquement périodique, de points du plan $\mathbb{R}^2$, appelés *sites* du treillis. Un treillis est caractérisé par deux vecteurs linéairement indépendants $\mathbf{r}_1$ et $\mathbf{r}_2$, regroupés en une matrice $\mathbf{R} = [\mathbf{r}_1 \ \mathbf{r}_2]$ telle que

$$\Lambda_{\mathbf{R}} = \{\mathbf{R}\mathbf{k} \mid \mathbf{k} \in \mathbb{Z}^2\}, \quad (3.1)$$

c'est-à-dire que les sites de $\Lambda_{\mathbf{R}}$ sont les points de coordonnées $\mathbf{R}\mathbf{k}$. La surface du parallélogramme engendré par $\mathbf{r}_1$ et $\mathbf{r}_2$ est $T^2 \triangleq |\det(\mathbf{R})|$. Ainsi, le treillis a une *densité* de $1/T^2$, exprimée en nombre de sites par unité de surface (voir [90] pour plus de détails sur les treillis).

On parle de treillis orthogonal (cartésien) si $\mathbf{R} = T\mathbf{I}$, où $\mathbf{I}$ est la matrice identité. Le treillis hexagonal uniforme (du premier type) Λ_{hex} et le treillis orthogonal $\Lambda_{\text{orth}} = \mathbb{Z}^2$, de densité 1, sont représentés sur la figure 3.1. Ils sont décrits par leurs matrices respectives :

$$\mathbf{R}_{\text{hex}} = \sqrt{\frac{2}{\sqrt{3}}} \begin{bmatrix} 1 & 1/2 \\ 0 & \sqrt{3}/2 \end{bmatrix}, \quad \mathbf{R}_{\text{orth}} = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}. \quad (3.2)$$

On définit aussi $\widehat{\Lambda}_{\mathbf{R}}$, le treillis *dual* ou *réciroque* de $\Lambda_{\mathbf{R}}$, comme le treillis de matrice $2\pi\widehat{\mathbf{R}} \triangleq 2\pi(\mathbf{R}^{-1})^T$. Ainsi, pour les treillis hexagonal et orthogonal, on a :

$$\widehat{\mathbf{R}}_{\text{hex}} = \sqrt{\frac{2}{\sqrt{3}}} \begin{bmatrix} \sqrt{3}/2 & 0 \\ -1/2 & 1 \end{bmatrix}, \quad \widehat{\mathbf{R}}_{\text{orth}} = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}. \quad (3.3)$$

Lorsque l'on échantillonne une fonction $s(\mathbf{x})$ sur $\Lambda_{\mathbf{R}}$, on obtient un signal discret 2D uniforme localisé sur le treillis $\Lambda_{\mathbf{R}}$. Le signal $v = (v[\mathbf{k}])_{\mathbf{k} \in \mathbb{Z}^2}$, tel que $v[\mathbf{k}] = s(\mathbf{R}\mathbf{k})$ pour tout $\mathbf{k}$, peut donc être interprété comme un peigne de Dirac 2D pondéré

$$v_c(\mathbf{x}) = \sum_{\mathbf{k} \in \mathbb{Z}^2} v[\mathbf{k}] \delta(\mathbf{x} - \mathbf{R}\mathbf{k}). \quad (3.4)$$

On remarque que

$$\hat{v}_c(\boldsymbol{\omega}) = \hat{v}(\mathbf{R}^T \boldsymbol{\omega}). \quad (3.5)$$

L'échantillonnage de $s(\mathbf{x})$ sur $\Lambda_{\mathbf{R}}$ a pour effet, en domaine de Fourier, de dupliquer son spectre à chaque site $2\pi\hat{\mathbf{R}}\mathbf{k}$, comme décrit par la formule de Poisson (prouvée dans [229]) :

$$\hat{v}_c(\boldsymbol{\omega}) = \sum_{\mathbf{k} \in \mathbb{Z}^2} s(\mathbf{R}\mathbf{k}) e^{-I\langle \boldsymbol{\omega}, \mathbf{R}\mathbf{k} \rangle} = \sum_{\mathbf{k} \in \mathbb{Z}^2} \hat{s}(\boldsymbol{\omega} - 2\pi\hat{\mathbf{R}}\mathbf{k}). \quad (3.6)$$

Notons l'analogie avec le cas 1D, où un signal uniforme dont les échantillons sont localisés aux Tk , $k \in \mathbb{Z}$ peut être vu comme défini sur le treillis $T\mathbb{Z}$.

On étend la définition de l'autocorrélation discrète à une fonction $\varphi \in L_2(\mathbb{R}^2)$, en posant $a_\varphi[\mathbf{k}] = (\bar{\varphi} * \varphi)(\mathbf{R}\mathbf{k})$. En utilisant la formule de Poisson, on obtient donc

$$\hat{a}_\varphi(\mathbf{R}^T \boldsymbol{\omega}) = \sum_{\mathbf{k} \in \mathbb{Z}^2} |\hat{\varphi}(\boldsymbol{\omega} - 2\pi\hat{\mathbf{R}}\mathbf{k})|^2. \quad (3.7)$$

On définit aussi la fonction duale de φ par [17] :

$$\hat{\varphi}_d(\boldsymbol{\omega}) = \frac{\hat{\varphi}(\boldsymbol{\omega})^*}{\hat{a}_\varphi(\mathbf{R}^T \boldsymbol{\omega})}. \quad (3.8)$$

Chaque treillis a une unique *région de Voronoï*, de surface T , qui est le domaine du plan composé de tous les points plus proches de l'origine $\mathbf{0}$ que de tout autre site. Les régions de Voronoï des treillis orthogonal et hexagonal sont représentées sur la figure 3.1. On peut définir la fonction indicatrice sur la région de Voronoï par

$$\mathbb{1}_{\mathbf{R}}(\mathbf{x}) = \begin{cases} 1 & \text{si } \mathbf{x} \in \text{région de Voronoï} \\ 1/m & \text{si } \mathbf{x} \in \text{frontière de la région de Voronoï} \\ 0 & \text{si } \mathbf{x} \notin \text{région de Voronoï} \end{cases}, \quad (3.9)$$

où m est le nombre de sites à égale distance de $\mathbf{x}$. Par définition, cette fonction pave le plan, lorsqu'on la réplique à chaque site du treillis, ce qui signifie que $\mathbb{1}_{\mathbf{R}}(\mathbf{x})$ vérifie la partition de l'unité :

$$\sum_{\mathbf{k} \in \mathbb{Z}^2} \mathbb{1}_{\mathbf{R}}(\mathbf{x} - \mathbf{R}\mathbf{k}) = 1. \quad (3.10)$$

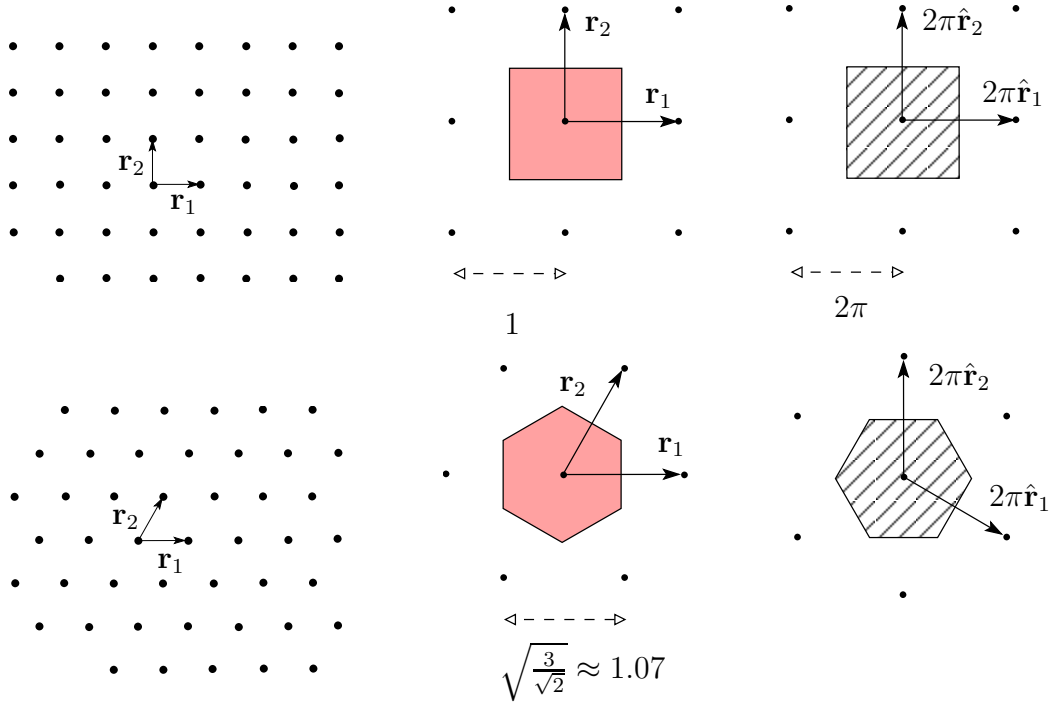


FIG. 3.1 : En haut : treillis orthogonal de densité 1, sa région de Voronoï, et sa région de Nyquist (région de Voronoï du treillis dual). En bas : même chose pour le treillis hexagonal du premier type, de densité 1.

Autrement dit, $\mathbb{1}_{\mathbf{R}}$ a un ordre d'approximation égal à 1 ; on dit que $\varphi(\mathbf{x})$ a un ordre d'approximation égal à L en 2D si tout polynôme de degré¹ au plus $L - 1$ s'exprime comme combinaison linéaire des fonctions $\varphi(\mathbf{x} - \mathbf{R}\mathbf{k})$. Cette propriété est équivalente aux conditions de Strang-Fix [205] :

$$\hat{\varphi}(\mathbf{0}) \neq 0 \quad \text{et} \quad \hat{\varphi}(\boldsymbol{\omega} - 2\pi\hat{\mathbf{R}}\mathbf{k}) = O(\|\boldsymbol{\omega}\|^L) \quad \forall \mathbf{k} \neq \mathbf{0}. \quad (3.11)$$

La cellule de Voronoï du treillis $\widehat{\Lambda}_{\mathbf{R}}$ est la *région de Nyquist* associée à $\Lambda_{\mathbf{R}}$: si l'on échantillonne sur $\Lambda_{\mathbf{R}}$ une fonction $s(\mathbf{x})$ telle que $\hat{s}$ est nulle en dehors de cette région, l'échantillonnage ne crée pas de repliement de spectre, car les répliques $\hat{s}(\boldsymbol{\omega} - 2\pi\hat{\mathbf{R}}\mathbf{k})$ ne se recouvrent pas entre elles. La région de Nyquist pour le treillis orthogonal est le carré $]-\pi, \pi[^2$, alors que pour le treillis hexagonal, il s'agit de l'hexagone de second type, de même surface $4\pi^2$, comme illustré sur la figure 3.1. Notons que cet hexagone s'étend au-delà du carré de même surface, comme illustré sur la figure 3.2, le long des directions horizontale et verticale. Cela signifie qu'une structure variant le long de ces directions sera mieux représentée sur le treillis hexagonal, alors que le treillis orthogonal est mieux adapté pour la représentation d'un motif diagonal à hautes fréquences, de type échiquier.

1. Le degré d'un polynôme à plusieurs variables $p(x_1, \dots, x_d)$ est défini comme le degré du polynôme à une variable $p(t, \dots, t)$.

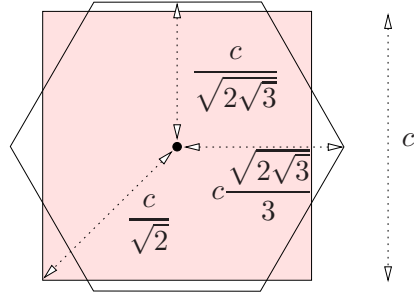


FIG. 3.2 : Dimensions de l'hexagone du second type et du carré, de même surface c^2 .

Dans toute la suite, on fera l'hypothèse que $\mathbf{R}$ a un déterminant égal à 1. Le treillis de densité $1/T^2$ correspondant sera donc $\Lambda_{T\mathbf{R}}$.

3.2 Reconstruction 2D dans un espace LSI

Définissons maintenant le problème de reconstruction en 2D, à partir de mesures discrètes localisées sur un treillis 2D. Le modèle 1D (1.10) se généralise comme suit : on dispose d'un signal $(v[\mathbf{k}])_{\mathbf{k} \in \mathbb{Z}^2}$ tel que

$$v[\mathbf{k}] = \int_{\mathbb{R}^2} s(\mathbf{x}) \tilde{\varphi}\left(\frac{\mathbf{x}}{T} - \mathbf{R}\mathbf{k}\right) \frac{d\mathbf{x}}{T^2} + \mu[\mathbf{k}] \quad \forall \mathbf{k} \in \mathbb{Z}^2, \quad (3.12)$$

c'est-à-dire que le signal v provient de la discrétisation de la fonction inconnue $s(\mathbf{x})$, à l'aide d'un dispositif d'acquisition de réponse impulsionnelle $\tilde{\varphi}_T(\mathbf{x}) = \frac{1}{T^2} \tilde{\varphi}(\frac{\mathbf{x}}{T})$, fournissant des mesures corrompues par un bruit additif μ , et localisées sur le treillis $\Lambda_{T\mathbf{R}}$.

On définit l'espace LSI $V_T(\varphi)$, engendré par les translatés sur le treillis $\Lambda_{T\mathbf{R}}$ de la fonction $\varphi \in L_2$, par

$$V_T(\varphi) = \left\{ g(\mathbf{x}) = \sum c[\mathbf{k}] \varphi\left(\frac{\mathbf{x}}{T} - \mathbf{R}\mathbf{k}\right) \mid c \in \mathbb{R}^{\mathbb{Z}^2} \right\}. \quad (3.13)$$

Nous dirons que les fonctions appartenant à un tel espace ont une résolution de $1/T$. Autrement dit, la notion de résolution est directement liée à la densité du treillis auquel les fonctions sont attachées.

On cherche à reconstruire une fonction $f_T(\mathbf{x}) \in V_T(\varphi)$ proche de s , par un processus de reconstruction $\mathcal{M} : v \mapsto f_T$ linéaire et invariant par translation sur $\Lambda_{T\mathbf{R}}$. f_T s'écrit donc sous la forme

$$f_T(\mathbf{x}) = \sum_{\mathbf{k} \in \mathbb{Z}^2} c[\mathbf{k}] \varphi\left(\frac{\mathbf{x}}{T} - \mathbf{R}\mathbf{k}\right), \quad (3.14)$$

où les coefficients $c[\mathbf{k}]$ sont obtenus par préfiltrage 2D à partir des données $v[\mathbf{k}]$:

$$c = v * p, \quad (3.15)$$

pour un certain préfiltre p à déterminer.

Les approches classiques vues au chapitre précédent, comme la reconstruction consistante, s'étendent sans difficulté au cas multi-dimensionnel. Par exemple, le préfiltre d'interpolation s'écrit

$$P_{\text{int}}(\mathbf{z}) = \frac{1}{\sum_{\mathbf{k} \in \mathbb{Z}^2} \varphi(\mathbf{R}\mathbf{k}) \mathbf{z}^{-\mathbf{k}}}. \quad (3.16)$$

Afin d'évaluer la qualité des méthodes de reconstruction, on étend en 2D le noyau d'erreur fréquentiel présenté à la section (2.4.3). Les résultats d'approximation valides en 1D se généralisent aisément au cadre multi-dimensionnel. On peut donc approcher très précisément l'erreur d'approximation par la formule

$$\|s - f_T\|_{L_2}^2 \approx \frac{1}{2\pi^2} \int_{\mathbb{R}^2} |s(\boldsymbol{\omega})|^2 E(T\boldsymbol{\omega}) d\boldsymbol{\omega}, \quad (3.17)$$

où le noyau d'erreur fréquentiel $E(\boldsymbol{\omega})$ s'écrit

$$E(\boldsymbol{\omega}) = 1 - \underbrace{\hat{\varphi}_d(\boldsymbol{\omega}) \hat{\varphi}(\boldsymbol{\omega})}_{E_{\min}(\boldsymbol{\omega})} + \hat{a}_\varphi(\mathbf{R}^T \boldsymbol{\omega}) \left| \hat{p}(\mathbf{R}^T \boldsymbol{\omega}) \hat{\varphi}(\boldsymbol{\omega}) - \hat{\varphi}_d(\boldsymbol{\omega}) \right|^2. \quad (3.18)$$

Nous étendrons dans la section 3.5 la démarche de reconstruction asymptotiquement optimale présentée au chapitre précédent. Avant cela, nous allons étudier le choix de la fonction φ définissant l'espace de reconstruction. Pour cela, nous présentons deux familles de splines multi-dimensionnelles, les hex-splines et les box-splines. Sauf mention contraire, on se place donc sur le treillis Λ_{hex} , et on pose $\mathbf{R} = \mathbf{R}_{\text{hex}}$ définie par (3.2).

3.3 Hex-Splines

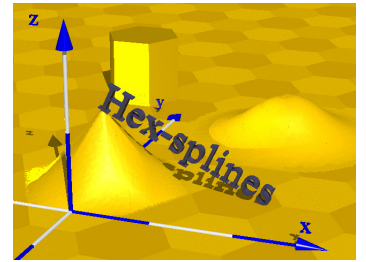
Une nouvelle famille de splines multi-dimensionnelles a été récemment conçue [229], afin détendre la construction des B-splines à des treillis multi-dimensionnels quelconques. Ces « hex-splines » ont depuis été utilisées avec succès pour des applications d'impression [231].

Le principe sous-jacent aux hex-splines est d'interpréter la B-spline centrée de degré 0 $\beta^0 = \mathbb{1}_{]-\frac{1}{2}, \frac{1}{2}[}$ comme l'indicatrice de Voronoï associée au treillis 1D $\mathbb{Z}$. On peut donc étendre cette construction à n'importe quel treillis $\Lambda_{\mathbf{R}}$, en définissant la hex-spline d'ordre 1 comme l'indicatrice de Voronoï du treillis :

$$\eta_1(\mathbf{x}) = \mathbb{1}_{\mathbf{R}}(\mathbf{x}). \quad (3.19)$$

Les hex-splines d'ordres plus élevés sont construites par simples convolutions successives : pour $L > 1$,

$$\eta_L = \eta_{L-1} * \eta_1 \xrightarrow{\mathcal{F}} \hat{\eta}_L(\boldsymbol{\omega}) = \hat{\eta}_1(\boldsymbol{\omega})^L. \quad (3.20)$$



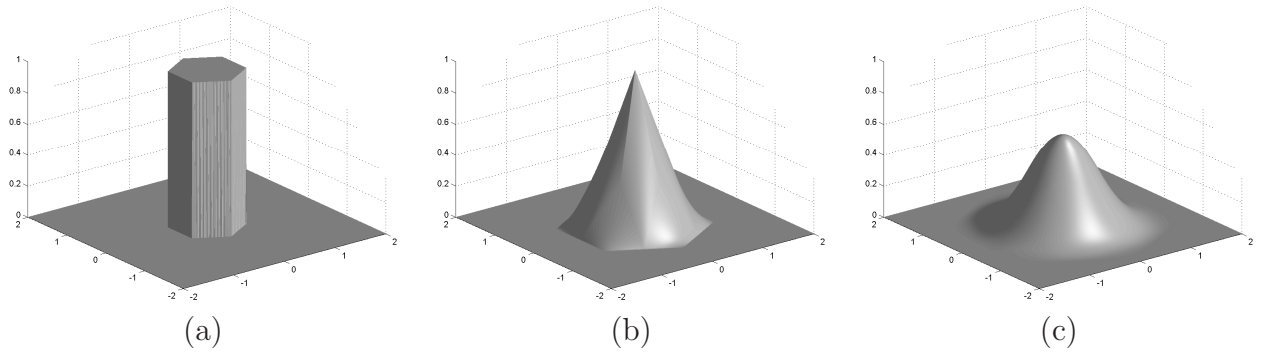


FIG. 3.3 : Hex-splines (a) $\eta_1(\mathbf{x})$, (b) $\eta_2(\mathbf{x})$, (c) $\eta_3(\mathbf{x})$ déployées sur le treillis hexagonal Λ_{hex} .

Notons que le suffixe L indique l'ordre d'approximation de η_L .

Sur le treillis orthogonal, on retrouve $\eta_L = \beta^{L+1}$. Sur le treillis hexagonal, on peut expliciter la forme des hex-splines en domaine de Fourier, en introduisant les vecteurs suivants, représentés sur la figure 3.8 :

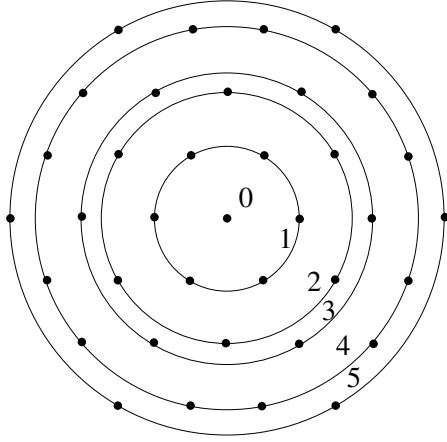
$$\mathbf{u}_1 = \sqrt{\frac{2}{\sqrt{3}}} \begin{bmatrix} 1/2 \\ \sqrt{3}/6 \end{bmatrix}, \quad \mathbf{u}_2 = \sqrt{\frac{2}{\sqrt{3}}} \begin{bmatrix} -1/2 \\ \sqrt{3}/6 \end{bmatrix}, \quad \mathbf{u}_3 = \sqrt{\frac{2}{\sqrt{3}}} \begin{bmatrix} 0 \\ -1/\sqrt{3} \end{bmatrix}. \quad (3.21)$$

Ainsi, pour tout $L \geq 1$,

$$\begin{aligned} \hat{\eta}_L(\boldsymbol{\omega}) = & \left(\frac{1}{3} \operatorname{sinc}\left(\frac{\langle \boldsymbol{\omega}, \mathbf{u}_1 \rangle}{2}\right) \operatorname{sinc}\left(\frac{\langle \boldsymbol{\omega}, \mathbf{u}_2 \rangle}{2}\right) \cos\left(\frac{\langle \boldsymbol{\omega}, \mathbf{u}_3 \rangle}{2}\right) \right. \\ & + \frac{1}{3} \operatorname{sinc}\left(\frac{\langle \boldsymbol{\omega}, \mathbf{u}_2 \rangle}{2}\right) \operatorname{sinc}\left(\frac{\langle \boldsymbol{\omega}, \mathbf{u}_3 \rangle}{2}\right) \cos\left(\frac{\langle \boldsymbol{\omega}, \mathbf{u}_1 \rangle}{2}\right) \\ & \left. + \frac{1}{3} \operatorname{sinc}\left(\frac{\langle \boldsymbol{\omega}, \mathbf{u}_3 \rangle}{2}\right) \operatorname{sinc}\left(\frac{\langle \boldsymbol{\omega}, \mathbf{u}_1 \rangle}{2}\right) \cos\left(\frac{\langle \boldsymbol{\omega}, \mathbf{u}_2 \rangle}{2}\right) \right)^L. \end{aligned} \quad (3.22)$$

$\eta_L(\mathbf{x})$ a la même symétrie d'ordre 12 que le treillis, provenant de la forme hexagonale de la cellule de Voronoï. Nous avons représenté les 3 premières hex-splines sur la figure 3.3. Les hex-splines ont un support compact de forme hexagonale, comme illustré sur les figures 3.8 et 3.3. Nous renvoyons à [229] pour d'autres propriétés des hex-splines (construction analytique, base de Riesz...), ainsi que pour les aspects pratiques liés à leur implémentation.

On définit la hex-spline discrète b_{η_L} par $b_{\eta_L}[\mathbf{k}] = \eta_L(\mathbf{Rk})$. En utilisant les programmes en langage Maple décrits dans [229], on peut déterminer ces hex-splines discrètes. Grâce à la symétrie d'ordre 12, les coefficients situés à égale distance de l'origine ont la même valeur. Il est donc pratique d'exprimer ces filtres à l'aide des fonctions $\text{ring}_n(\mathbf{z})$ définies par la figure 3.4. Ainsi, $B_{\eta_3}(\mathbf{z}) = \frac{42}{72} + \frac{5}{72} \text{ring}_1(\mathbf{z})$. η_1 et η_2 sont interpolantes, ce qui signifie que $b_{\eta_L} = \delta$ pour $L = 1$ et 2. Rappelons que l'interpolation (généralisée) à l'aide de hex-splines se fait au moyen du préfiltre $p_{\text{int}} = b_{\eta_L}^{-1}$.



n	$\text{ring}_n(\mathbf{z})$
0	1
1	$z_1 + z_2 + z_1^{-1} + z_2^{-1} + z_1 z_2^{-1} + z_2 z_1^{-1}$
2	$z_1 z_2 + z_1^{-1} z_2^{-1} + z_1^2 z_2^{-1} + z_1^{-2} z_2 + z_1^{-1} z_2^2 + z_1 z_2^{-2}$
3	$z_1^2 + z_2^2 + z_1^{-2} + z_2^{-2} + z_1^2 z_2^{-2} + z_2^2 z_1^{-2}$
4	$z_1^2 z_2 + z_1^{-2} z_2^{-1} + z_1^3 z_2^{-1} + z_1^{-3} z_2 + z_1 z_2^2 + z_1^{-1} z_2^{-2} + z_1^{-1} z_2^3 + z_1 z_2^{-3} + z_1^{-2} z_2^3 + z_1^2 z_2^{-3} + z_1^{-3} z_2^2 + z_1^3 z_2^{-2}$
5	$z_1^3 + z_2^3 + z_1^{-3} + z_2^{-3} + z_1^3 z_2^{-3} + z_1^{-3} z_2^3$

FIG. 3.4 : On définit des anneaux (à gauche) collectant les sites du treillis Λ_{hex} à égale distance de l'origine. On peut ainsi exprimer aisément, et sous une forme compacte, les transformées en $\mathbf{z}$ de filtres discrets isotropes, au moyen des fonctions $\text{ring}_n(\mathbf{z})$ (à droite), qui sont les transformées en $\mathbf{z}$ de filtres concentriques dont les coefficients sont 1 sur le n -ième anneau, et 0 ailleurs.

3.3.1 De l'optimalité du treillis hexagonal

Dans ce paragraphe, nous nous intéressons aux propriétés d'approximation des treillis orthogonal et hexagonal. Pour cela, nous considérons une base hex-spline « naturelle » sur chaque treillis, c'est-à-dire les splines séparables sur le treillis orthogonal, et les hex-splines hexagonales sur le treillis hexagonal. Nous allons évaluer, pour un ordre d'approximation donné, la qualité d'approximation par projection orthogonale hex-spline sur chacun des treillis, en comparant les noyaux d'erreur $E_{\min}$ associés.

Notons $\mathcal{P}_T s$ la projection orthogonale de $s(\mathbf{x})$ dans l'espace $V_T(\varphi)$ où $\varphi = \eta_L$ est la hex-spline d'ordre L , et ce pour un treillis quelconque. Alors, au moyen de développements de Taylor dans (3.18) et (3.17), on montre aisément que $E_{\min}(\boldsymbol{\omega}) = O(\|\boldsymbol{\omega}\|^{2L})$ et $\|s - \mathcal{P}_T s\|_{L_2} = O(T^L)$. Plus précisément, en faisant le changement de coordonnées polaires

$$\boldsymbol{\omega} = (\omega_1, \omega_2) = (\|\boldsymbol{\omega}\| \cos(\theta), \|\boldsymbol{\omega}\| \sin(\theta)), \quad (3.23)$$

on peut définir la *constante asymptotique angulaire* $C_L(\theta)$ telle que

$$E(\boldsymbol{\omega}) \sim C_L(\theta)^2 \|\boldsymbol{\omega}\|^{2L} \quad \text{quand } \|\boldsymbol{\omega}\| \rightarrow 0, \quad (3.24)$$

ou de manière équivalente,

$$\|s - \mathcal{P}_T s\|_{L_2}^2 \sim \frac{T^{2L}}{4\pi^2} \int_0^{2\pi} \int_0^\infty C_L(\theta)^2 \|\boldsymbol{\omega}\|^{2L+1} |\hat{s}(\boldsymbol{\omega})|^2 d\theta d\|\boldsymbol{\omega}\| \quad \text{quand } T \rightarrow 0. \quad (3.25)$$

Pour un ordre d'approximation donné, la constante asymptotique $C_L(\theta)$ fournit un bon indicateur de la qualité d'approximation, puisqu'elle détermine la tendance du noyau

d'erreur au voisinage de l'origine, et donc la constante dans la vitesse de convergence de l'erreur d'approximation lorsque le pas T du treillis tend vers 0. Ainsi, on peut comparer les capacités d'approximation des treillis, au moyen de leurs constantes asymptotiques associées. Déterminons donc $C_L(\theta)$, sur les treillis orthogonal et hexagonal.

Calcul de $C_L(\theta)$ sur le treillis orthogonal

On se place sur le treillis orthogonal Λ_{orth} . La hex-spline d'ordre L , déployée sur ce treillis, est simplement la B-spline séparable $\varphi(\mathbf{x}) = \beta^{L-1}(\mathbf{x}) = \beta^{L-1}(x_1)\beta^{L-1}(x_2)$. On a :

$$\hat{\varphi}(\boldsymbol{\omega}) = \left(\frac{\sin(\frac{\omega_1}{2})}{\frac{\omega_1}{2}} \frac{\sin(\frac{\omega_2}{2})}{\frac{\omega_2}{2}} \right)^L \quad (3.26)$$

et

$$\hat{a}_\varphi(\boldsymbol{\omega}) = \sum_{\mathbf{k} \in \mathbb{Z}^2} \left(\frac{\sin(\frac{\omega_1}{2})}{\frac{\omega_1 + 2k_1\pi}{2}} \frac{\sin(\frac{\omega_2}{2})}{\frac{\omega_2 + 2k_2\pi}{2}} \right)^{2L}, \quad (3.27)$$

d'où

$$E_{\min}(\boldsymbol{\omega}) = 1 - \frac{1}{\omega_1^{2L} \omega_2^{2L} \sum_{\mathbf{k} \in \mathbb{Z}^2} \left(\frac{\omega_1}{\omega_1 + 2k_1\pi} \frac{\omega_2}{\omega_2 + 2k_2\pi} \right)^{2L}} \quad (3.28)$$

$$= \omega_1^{2L} \sum_{k_2 \neq 0} \frac{1}{(2k_2\pi)^{2L}} + \omega_2^{2L} \sum_{k_1 \neq 0} \frac{1}{(2k_1\pi)^{2L}} + o(\|\boldsymbol{\omega}\|^{2L}) \quad (3.29)$$

$$= \frac{2\zeta(2L)}{(2\pi)^{2L}} (\omega_1^{2L} + \omega_2^{2L}) + o(\|\boldsymbol{\omega}\|^{2L}), \quad (3.30)$$

où la fonction zeta de Riemann est définie par (1.8). On a donc :

$$C_L(\theta)^2 = \frac{2\zeta(2L)}{(2\pi)^{2L}} (\cos(\theta)^{2L} + \sin(\theta)^{2L}). \quad (3.31)$$

Cette constante est $\frac{\pi}{2}$ -périodique. Elle est minimale pour $\theta = \frac{\pi}{4}$, car les fréquences diagonales sont les mieux préservées sur le treillis orthogonal. À l'opposé, $C_L(\theta)$ est maximale pour $\theta = 0$. On peut donc définir

$$C_L^{\min 2} = C_L(\frac{\pi}{4})^2 = \frac{4\zeta(2L)}{(2\sqrt{2}\pi)^{2L}} \quad (3.32)$$

et

$$C_L^{\max 2} = C_L(0)^2 = \frac{2\zeta(2L)}{(2\pi)^{2L}}. \quad (3.33)$$

On peut aussi définir l'anisotropie du treillis, par $\sqrt[2]{\frac{C_L^{\max}}{C_L^{\min}}}$, qui tend ici vers $\sqrt{2}$ quand $L \rightarrow +\infty$. Cela signifie que l'on peut représenter avec la même qualité d'approximation

un motif oscillant dans la direction diagonale de fréquence $\sqrt{2}$ fois plus élevée qu'un motif horizontal ou vertical (asymptotiquement, dans les basses fréquences). Les ordres faibles ne permettent pas de tirer parti de cette spécificité du treillis orthogonal : pour $L = 1$, $C_L(\theta)$ est isotrope, donc les motifs diagonaux ne sont pas mieux représentés que les autres dans l'espace de reconstruction constant par morceaux.

Calcul de $C_L(\theta)$ sur le treillis hexagonal

Plaçons-nous maintenant sur le treillis Λ_{hex} . Afin de calculer la constante asymptotique angulaire $C_L(\theta)$, on déduit tout d'abord de (3.20) que

$$\hat{a}_{\eta_L}(\omega) = \sum_{\mathbf{k} \in \mathbb{Z}^2} \left| \hat{\eta}_1(\omega - 2\pi \hat{\mathbf{R}}\mathbf{k}) \right|^{2L}. \quad (3.34)$$

En utilisant (3.11), on peut simplifier le noyau d'erreur en

$$E(\omega) = \frac{\sum_{\mathbf{k} \neq \mathbf{0}} \left| \hat{\eta}_1(\omega - 2\pi \hat{\mathbf{R}}\mathbf{k}) \right|^{2L}}{\sum_{\mathbf{k} \in \mathbb{Z}^2} \left| \hat{\eta}_1(\omega - 2\pi \hat{\mathbf{R}}\mathbf{k}) \right|^{2L}} = \sum_{\mathbf{k} \neq \mathbf{0}} \left| \hat{\eta}_1(\omega - 2\pi \hat{\mathbf{R}}\mathbf{k}) \right|^{2L} + O(\|\omega\|^{4L}). \quad (3.35)$$

Par conséquent,

$$C_L(\theta)^2 = \lim_{\omega \rightarrow \mathbf{0}} \frac{\sum_{\mathbf{k} \neq \mathbf{0}} \left| \hat{\eta}_1(\omega - 2\pi \hat{\mathbf{R}}\mathbf{k}) \right|^{2L}}{\|\omega\|^{2L}} = \sum_{\mathbf{k} \neq \mathbf{0}} |A(\mathbf{k}, \theta)|^{2L}, \quad (3.36)$$

où

$$A(\mathbf{k}, \theta) = \lim_{\|\omega\| \rightarrow \mathbf{0}} \hat{\eta}_1(\omega - 2\pi \hat{\mathbf{R}}\mathbf{k}) / \|\omega\|. \quad (3.37)$$

À l'aide du logiciel de calcul formel Maple, nous avons obtenu le résultat suivant, en effectuant le développement de Taylor de $\hat{\eta}_1(\omega - 2\pi \hat{\mathbf{R}}^T \mathbf{k})$ à l'origine, après être passés en coordonnées polaires :

$$A(\mathbf{k}, \theta) = \sqrt{\frac{2}{\sqrt{3}}} \frac{3\sqrt{3}}{4\pi^2} \frac{S(k_1 + k_2)}{k_1 + k_2} \left(\frac{\cos(\theta)}{2k_2 - k_1} + \frac{\cos(\theta + \frac{2\pi}{3})}{2k_1 - k_2} \right) \quad (3.38)$$

avec

$$S(k_1 + k_2) = \begin{cases} 0 & \text{si } k_1 + k_2 \equiv 0[3] \\ -1 & \text{si } k_1 + k_2 \equiv 1[3] \\ 1 & \text{si } k_1 + k_2 \equiv 2[3] \end{cases}. \quad (3.39)$$

La formule (3.38) n'est valide que lorsque son dénominateur ne s'annule pas. Sinon, on a :

$$\text{si } k_2 = -k_1, \quad A(\mathbf{k}, \theta) = \sqrt{\frac{2}{\sqrt{3}}} \frac{-1}{6\pi k_1} \cos(\theta + \frac{\pi}{3}), \quad (3.40)$$

$$\text{si } k_2 = 2k_1, \quad A(\mathbf{k}, \theta) = \sqrt{\frac{2}{\sqrt{3}}} \frac{-1}{6\pi k_1} \cos(\theta - \frac{\pi}{3}), \quad (3.41)$$

$$\text{si } k_1 = 2k_2, \quad A(\mathbf{k}, \theta) = \sqrt{\frac{2}{\sqrt{3}}} \frac{-1}{6\pi k_2} \cos(\theta). \quad (3.42)$$

En calculant la contribution $C_L^1(\theta)^2$ de ces trois derniers termes à $C_L(\theta)^2$, on trouve

$$C_L^1(\theta)^2 = \sum_{\substack{\mathbf{k} \neq \mathbf{0} \text{ et} \\ (k_2 = -k_1 \text{ ou } k_2 = 2k_1 \text{ ou } k_1 = 2k_2)}} |A(\mathbf{k}, \theta)|^{2L} = \sum_{n \neq 0} \left(\sqrt{\frac{2}{\sqrt{3}}} \frac{1}{6\pi n} \right)^{2L} \left(\cos(\theta)^{2L} + \cos(\theta + \frac{\pi}{3})^{2L} + \cos(\theta - \frac{\pi}{3})^{2L} \right) \quad (3.43)$$

$$= \frac{2\zeta(2L)}{(18\pi^2\sqrt{3})^L} \left(\cos(\theta)^{2L} + \cos(\theta + \frac{\pi}{3})^{2L} + \cos(\theta - \frac{\pi}{3})^{2L} \right) \quad (3.44)$$

$$= \frac{6\zeta(2L)}{(72\pi^2\sqrt{3})^L} \sum_{n=-\lfloor L/3 \rfloor}^{\lfloor L/3 \rfloor} \binom{2L}{3n+L} \cos(6n\theta). \quad (3.45)$$

Le reste des termes $C_L^2(\theta)^2 = C_L(\theta)^2 - C_L^1(\theta)^2$ vaut

$$C_L^2(\theta)^2 = \sum_{k_2 \neq -k_1 \text{ et } k_2 \neq 2k_1 \text{ et } k_1 \neq 2k_2} |A(\mathbf{k}, \theta)|^{2L} \quad (3.46)$$

$$= \left(\frac{27}{8\sqrt{3}\pi^4} \right)^L \sum_{k_2+k_1 \not\equiv 0[3]} \left[\frac{1}{k_1+k_2} \left(\frac{\cos(\theta)}{2k_2-k_1} + \frac{\cos(\theta + \frac{2\pi}{3})}{2k_1-k_2} \right) \right]^{2L}. \quad (3.47)$$

On a donc finalement

$$C_L(\theta)^2 = \frac{6\zeta(2L)}{(72\pi^2\sqrt{3})^L} \sum_{n=-\lfloor L/3 \rfloor}^{\lfloor L/3 \rfloor} \binom{2L}{3n+L} \cos(6n\theta) \quad (3.48)$$

$$+ \left(\frac{27}{8\sqrt{3}\pi^4} \right)^L \sum_{k_2+k_1 \not\equiv 0[3]} \left[\frac{1}{k_1+k_2} \left(\frac{\cos(\theta)}{2k_2-k_1} + \frac{\cos(\theta + \frac{2\pi}{3})}{2k_1-k_2} \right) \right]^{2L}. \quad (3.49)$$

$C_L(\theta)$ est $\frac{\pi}{3}$ -périodique, et son maximum est atteint pour $\theta = \frac{\pi}{2}$. On définit donc $C_L^{\max} = C_L(\frac{\pi}{2})$, caractérisant le comportement au pire cas. Calculons cette constante de manière explicite. On a d'abord

$$C_L^1(\frac{\pi}{2})^2 = \frac{4\zeta(2L)}{(18\pi^2\sqrt{3})^L} \left(\frac{\sqrt{3}}{2} \right)^{2L} = \frac{4\zeta(2L)}{(24\pi^2\sqrt{3})^L}. \quad (3.50)$$

La deuxième partie de la constante vaut :

$$C_L^2(\frac{\pi}{2})^2 = \left(\frac{27\sqrt{3}}{32\pi^4}\right)^L \sum_{k_2+k_1 \not\equiv 0[3]} \frac{1}{(k_1+k_2)^{2L}(k_2-2k_1)^{2L}} \quad (3.51)$$

$$= \left(\frac{27\sqrt{3}}{32\pi^4}\right)^L \left[\sum_{(n_1, k_1) \in \mathbb{Z}^2} \frac{1}{(3n_1+1)^{2L}(3n_1+1-3k_1)^{2L}} \right. \quad (3.52)$$

$$\left. + \sum_{(n_1, k_1) \in \mathbb{Z}^2} \frac{1}{(3n_1+2)^{2L}(3n_1+2-3k_1)^{2L}} \right] \quad (3.53)$$

$$= \left(\frac{27\sqrt{3}}{32\pi^4}\right)^L \left[\sum_{(n_1, n_2) \in \mathbb{Z}^2} \frac{1}{(3n_1+1)^{2L}(3n_2+1)^{2L}} + \sum_{(n_1, n_2) \in \mathbb{Z}^2} \frac{1}{(3n_1+2)^{2L}(3n_2+2)^{2L}} \right]. \quad (3.54)$$

où nous avons coupé la somme en deux en posant respectivement $k_1+k_2=3n_1+1$ et $k_1+k_2=3n_1+2$, avant d'effectuer le changement de variable $n_2=n_1-k_1$. En remarquant que les sommes sont séparables en n_1, n_2 et que les deux sommes sont égales, on obtient :

$$C_L^2(\frac{\pi}{2})^2 = \left(\frac{27\sqrt{3}}{32\pi^4}\right)^L 2 \left(\sum_{n \in \mathbb{Z}} \frac{1}{(3n+1)^{2L}} \right)^2 \quad (3.55)$$

$$= \left(\frac{27\sqrt{3}}{32\pi^4}\right)^L 2 \left(\frac{1}{2} \sum_{n \in \mathbb{Z}} \frac{1}{(3n+1)^{2L}} + \frac{1}{2} \sum_{n \in \mathbb{Z}} \frac{1}{(3n+2)^{2L}} \right)^2 \quad (3.56)$$

$$= \left(\frac{27\sqrt{3}}{32\pi^4}\right)^L 2 \left(\frac{1}{2} \sum_{m \neq 0} \frac{1}{m^{2L}} - \frac{1}{2} \sum_{n \neq 0} \frac{1}{(3n)^{2L}} \right)^2 \quad (3.57)$$

$$= \left(\frac{27\sqrt{3}}{32\pi^4}\right)^L 2 \left(\zeta(2L) \left(1 - \frac{1}{3^{2L}} \right) \right)^2. \quad (3.58)$$

D'où finalement

$$C_L^{\max 2} = \frac{2\zeta(2L)^2 (9^L - 1)^2}{(32\sqrt{3}\pi^4)^L} + \frac{4\zeta(2L)}{(24\sqrt{3}\pi^2)^L}. \quad (3.59)$$

Lorsque $L \rightarrow \infty$, on a l'équivalent

$$C_L^{\max 2} \sim \left(\frac{27\sqrt{3}}{32\pi^4}\right)^L \quad (3.60)$$

qui est assez précis : pour $L=4$, l'équivalent fournit la valeur de $C_L^{\max}$ avec une erreur relative de 0.9%.

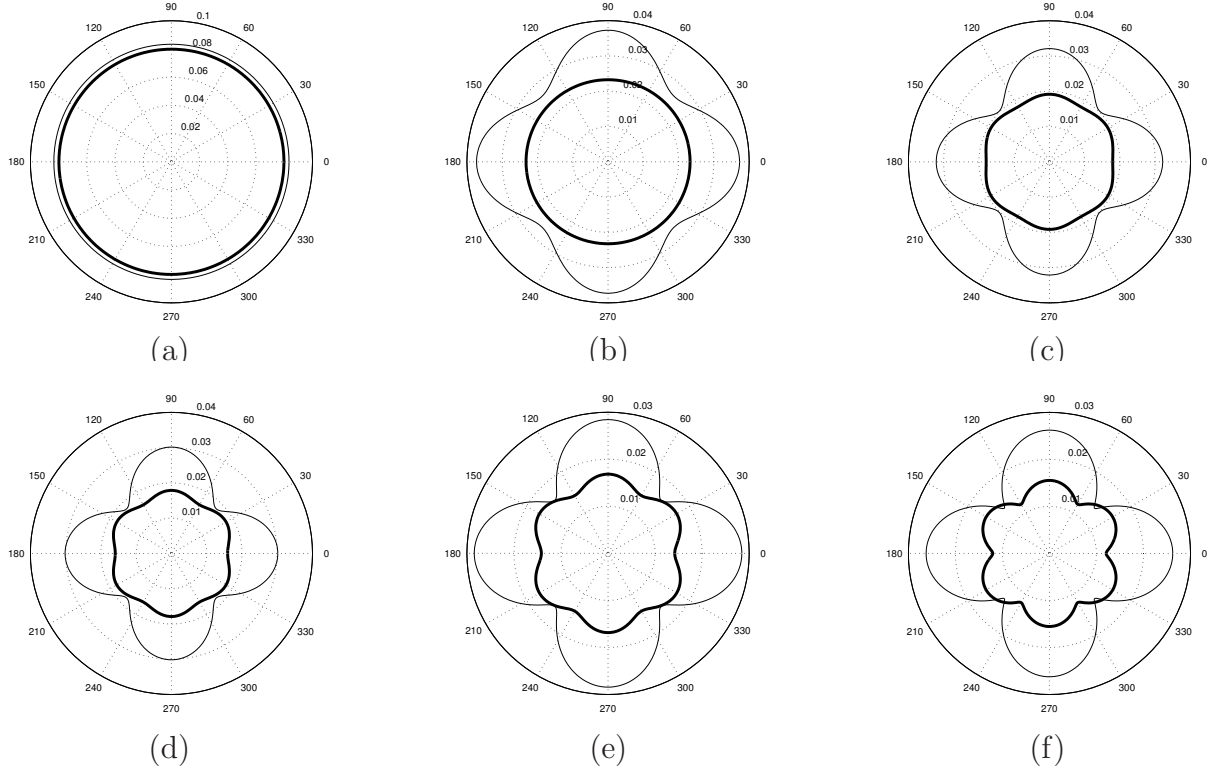


FIG. 3.5 : Représentation des constantes angulaires $\sqrt[2]{C_L(\theta)^2}$, pour L valant, de (a) à (e), 1,2,3,4,6,20. Le trait fin et le trait épais sont respectivement associés aux constantes des treillis orthogonal et hexagonal.

Comparaison des constantes asymptotiques

Lorsque l'on considère le treillis $\Lambda_{T\mathbf{R}}$, de densité $1/T^2$, au lieu du treillis normalisé $\Lambda_{\mathbf{R}}$, le noyau d'erreur associé devient $E_{\min}(T\boldsymbol{\omega})$. Par conséquent, la constante asymptotique $C_L(\theta)^2$ associée est multipliée par T^{2L} . On peut ainsi comparer les treillis orthogonal et hexagonal en cherchant le facteur T tel que $T\Lambda_{\text{hex}}$ et Λ_{orth} aient des constantes asymptotiques $C_L(\theta)$ similaires. Afin de rendre la comparaison indépendante de l'angle θ , nous considérons les constantes moyennes

$$\bar{C}_L^2 = \frac{1}{2\pi} \int_0^{2\pi} C_L(\theta)^2 d\theta. \quad (3.61)$$

L'intérêt d'étudier cette constante moyenne apparaît en particulier pour l'approximation de fonctions isotropes. En effet, dans ce cas, on peut réécrire (3.25) comme :

$$\|s - \mathcal{P}_T s\|_{L_2}^2 \sim \frac{T^{2L}}{4\pi^2} \int_0^{2\pi} C_L(\theta)^2 d\theta \int_0^\infty |\hat{s}(\|\boldsymbol{\omega}\|)|^2 \|\boldsymbol{\omega}\|^{2L+1} d\|\boldsymbol{\omega}\| \quad (3.62)$$

$$\sim \bar{C}_L^2 T^{2L} \|\Delta^L s\|_{L_2}^2, \quad (3.63)$$

ordre d'approx. L	gain surfaccique T^2	amélioration de Λ_{hex} par rapport à Λ_{orth} (%)
1	$\frac{3\sqrt{3}}{5} \approx 1.039$	4
2	$\frac{9}{\sqrt{42}} \approx 1.389$	28
3	≈ 1.478	32
4	≈ 1.53	35
6	≈ 1.58	37
20	≈ 1.65	40

TAB. 3.1 – Gain surfaccique T^2 tel que le treillis hexagonal de densité $1/T^2$ ait le même comportement asymptotique de l'erreur d'approximation que le treillis orthogonal de densité 1.

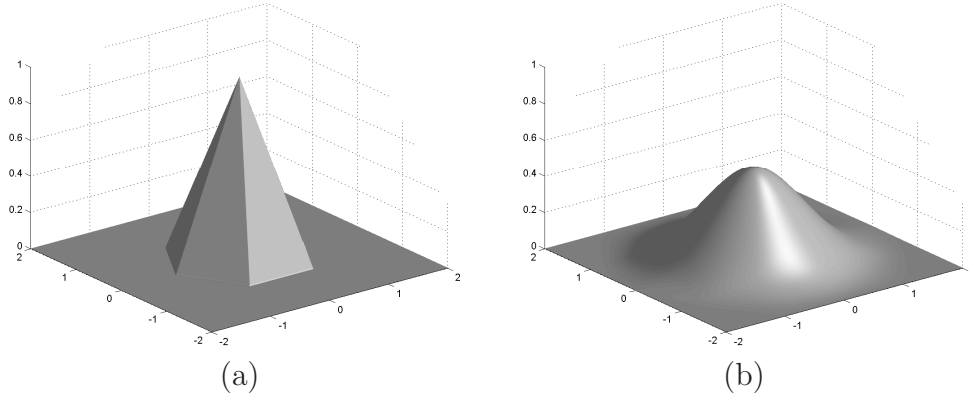
où Δ^L est l'opérateur Laplacien itéré d'ordre L . On a donc ici l'équivalent 2D de la formule asymptotique (2.68), avec une vitesse de convergence égale à $\bar{C}_L$ fois la puissance L -ième du pas T .

On considère donc le facteur de gain surfaccique T^2 rendant égales les constantes asymptotiques moyennes sur les treillis $T\Lambda_{\text{hex}}$ et Λ_{orth} , c'est-à-dire

$$T^2 = \sqrt[L]{\frac{\bar{C}_L^{\text{orth } 2}}{\bar{C}_L^{\text{hex } 2}}}. \quad (3.64)$$

Les valeurs de ce gain pour les premiers ordres sont reportées dans le tableau 3.1. On voit que, asymptotiquement (dans les basses fréquences), on a la même qualité d'approximation sur un treillis orthogonal et sur un treillis hexagonal ayant une densité d'échantillons près de 40% moindre. Il faut cependant garder à l'esprit qu'un tel gain n'est possible que pour l'approximation des basses fréquences des signaux. Les hautes fréquences ne s'accrochent pas de la réduction de la bande de Nyquist occasionnée par un tel changement de densité. Ainsi, le gain effectif que l'on peut espérer en pratique dépend essentiellement de la proportion dans laquelle l'énergie du signal est concentrée dans les basses fréquences. On remarque que l'usage de splines d'ordre 1 (approximation constante par morceaux) ne permet pas de tirer parti de cet avantage du treillis hexagonal. Cette pathologie disparaît dès l'ordre 2.

Notre étude complète le résultat classique selon lequel le treillis hexagonal permet une meilleure représentation des signaux à bande limitée isotropes [156]. Plus précisément, pour représenter de tels signaux dans les conditions du théorème de l'échantillonnage, un treillis hexagonal de densité $\frac{2}{\sqrt{3}}$ fois moindre que celle d'un treillis orthogonal est requis. Cela correspond à 13% d'échantillons en moins. Notre résultat est cependant plus général, au sens où l'hypothèse de bande limitée n'est pas requise, ni sur les signaux, ni sur l'espace de reconstruction, sachant que l'on tend vers l'approximation à bande limitée ($\varphi = \text{sinc}$) lorsque $L \rightarrow \infty$. Si l'on veut étudier de manière plus précise l'approximation de signaux non

FIG. 3.6 : Box-splines à trois directions (a) $\chi_2(\mathbf{x})$ et (b) $\chi_4(\mathbf{x})$.

isotropes, variant selon une direction privilégiée, il faut comparer les constantes angulaires $C(\theta)$ et non plus seulement les constantes moyennes $\bar{C}$. La figure 3.5 montre les constantes sur les treillis orthogonal et hexagonal, pour différentes valeurs de L . Il est intéressant de constater que, contrairement à ce que laisse penser l'étude des régions de Nyquist (voir figure 3.2), les basses fréquences diagonales (d'angle $\frac{\pi}{4}$) sont aussi bien représentées sur les deux treillis. Le treillis orthogonal ne prend l'avantage que pour l'approximation de hautes fréquences diagonales, par exemple pour la reconstruction d'une texture de type échiquier.

3.4 Box-splines pour le treillis hexagonal

Les box-splines sont une autre généralisation multi-dimensionnelle des B-splines, elles aussi polynomiales par morceaux, proposée par De Boor [86]. Les box-splines ont de multiples applications en modélisation géométrique, représentation multi-échelles, et de nombreux autres champs d'études [177, 146, 87]. Une box-spline centrée $\chi_{\Xi}(\mathbf{x})$ dépend de N vecteurs, arrangés en une matrice $\Xi = [\mathbf{v}_1 \cdots \mathbf{v}_N]$ ($N \geq 2$), et peut être définie en domaine de Fourier par

$$\hat{\chi}_{\Xi}(\boldsymbol{\omega}) = \prod_{i=1}^N \text{sinc}\left(\frac{\langle \boldsymbol{\omega}, \mathbf{v}_i \rangle}{2}\right). \quad (3.65)$$

Ainsi, on a la propriété de convolution $\chi_{\Xi_1 \cup \Xi_2} = \chi_{\Xi_1} * \chi_{\Xi_2}$. On peut aussi construire les box-splines en domaine spatial, comme suit : si $\mathbf{v}_1$ et $\mathbf{v}_2$ sont deux vecteurs linéairement indépendants,

$$\chi_{[\mathbf{v}_1 \ \mathbf{v}_2]}(\mathbf{x}) = \begin{cases} 1/|\det(\Xi)| & \text{si } \Xi^{-1}\mathbf{x} \in]-\frac{1}{2}, \frac{1}{2}]^2, \\ 0 & \text{sinon,} \end{cases} \quad (3.66)$$

et récursivement,

$$\chi_{\Xi \cup [\mathbf{v}]}(\mathbf{x}) = \int_{-\frac{1}{2}}^{\frac{1}{2}} \chi_{\Xi}(\mathbf{x} - t\mathbf{v}) dt. \quad (3.67)$$

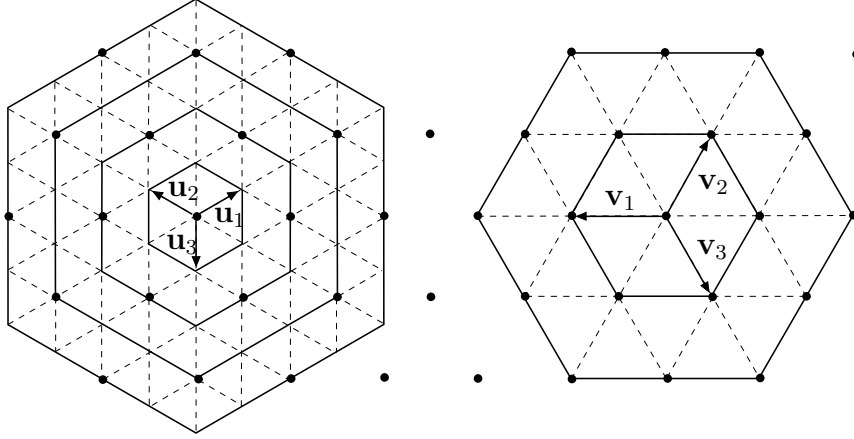


FIG. 3.8 : Les hex-splines η_L et les box-splines χ_{2n} sont polynomiales par morceaux sur des triangles (en pointillés). Ces triangles sont de surface trois fois plus grande dans le cas des box-splines (figure de droite). Les supports de η_L et χ_{2n} , indiqués par des hexagones en trait plein, ont une surface égale respectivement à L^2 et $3n^2$.

En considérant les vecteurs $\mathbf{v}_1 = [1 \ 0]^T$ et $\mathbf{v}_2 = [0 \ 1]^T$ définissant le treillis Λ_{orth} , on obtient les B-splines séparables : $\chi_{[\mathbf{v}_1 \ \mathbf{v}_2]} = \beta^0$.

Les box-splines à trois directions, adaptées au treillis hexagonal, sont définies à l'aide des vecteurs suivants (ou de manière équivalente, à l'aide des vecteurs de (3.77)), représentés sur la figure 3.8 :

$$\mathbf{v}_1 = \sqrt{\frac{2}{\sqrt{3}}} \begin{bmatrix} -1 \\ 0 \end{bmatrix}, \quad \mathbf{v}_2 = \sqrt{\frac{2}{\sqrt{3}}} \begin{bmatrix} 1/2 \\ \sqrt{3}/2 \end{bmatrix}, \quad \mathbf{v}_3 = \sqrt{\frac{2}{\sqrt{3}}} \begin{bmatrix} 1/2 \\ -\sqrt{3}/2 \end{bmatrix}. \quad (3.68)$$

Dans la suite, on définit $\chi_2(\mathbf{x}) = \chi_{[\mathbf{v}_1, \mathbf{v}_2, \mathbf{v}_3]}$ (connu sous le nom de *Courant element* [75]) et on définit, pour tout $n > 1$, la box-spline d'ordre d'approximation $2n$ par

$$\chi_{2n}(\mathbf{x}) = \chi_{2n-2} * \chi_2. \quad (3.69)$$

Les box-splines $\chi_{2n}(\mathbf{x})$ ont une symétrie hexagonale et un support compact, comme le montrent les figures 3.8 et 3.6. Lorsque n augmente, χ_{2n} tend vers une gaussienne. Les box-splines ont été utilisées avec succès dans de nombreux problèmes impliquant des données définies sur un treillis hexagonal [146, 177, 171].

On définit la box-spline discrète $b_{\chi_{2n}}$ par $b_{\chi_{2n}}[\mathbf{k}] = \chi_{2n}(\mathbf{R}\mathbf{k})$. Grâce à la propriété $\chi_{2n} * \chi_{2n} = \chi_{4n}$, l'autocorrélation discrète $a_{\chi_{2n}}$ est égale à $b_{\chi_{4n}}$. Les deux premières box-splines discrètes, calculées grâce à notre formule (3.99), sont (en positionnant $b_{\chi_{2n}}[\mathbf{k}]$ au site $\mathbf{R}\mathbf{k}$) :

$$\begin{array}{ccccccc}
& & & & 1/12 & & 1/12 \\
& & & & & & \\
b_{\chi_2} = 1, & & b_{\chi_4} = & 1/12 & 1/2 & 1/12 & .
\end{array} \quad (3.70)$$

Grâce à la symétrie d'ordre 12, les coefficients situés à la même distance de l'origine ont la même valeur. Comme pour les hex-splines discrètes, il est donc pratique d'utiliser les fonctions $\text{ring}_n(\mathbf{z})$ de la figure 3.4 pour exprimer les box-splines discrètes. Ainsi, $B_{\chi_4}(\mathbf{z}) = \frac{1}{2} + \frac{1}{12} \text{ring}_1(\mathbf{z})$. Notons que χ_2 est interpolante, autrement dit $B_{\chi_2}(\mathbf{z}) = 1$.

Dans la suite de cette section, nous proposons une nouvelle caractérisation des box-splines χ_{2n} . Cette formulation différentielle fournit une formule analytique explicite et originale pour les box-splines, ainsi qu'une nouvelle implémentation particulièrement efficace. En effet, les algorithmes disponibles dans la littérature reposent tous sur des propriétés récursives des box-splines, et sont donc très coûteux en temps et en mémoire [64, 81, 82, 63, 138, 39, 133]. Nous allons nous inspirer d'une construction déjà appliquée avec succès pour les splines exponentielles [224] et les $\mathcal{L}$ -splines [225] : nous décomposons une box-spline comme la convolution d'un filtre discret, jouant le rôle d'un opérateur de localisation, et d'une fonction génératrice, définie comme une fonction de Green d'un opérateur différentiel approprié. Ce paradigme est très général, et permet de définir de nombreuses fonctions. Avant de construire les box-splines à trois directions, voyons comment on peut définir de manière similaire les B-splines 1D.

3.4.1 Les B-splines revisitées

Nous avons précédemment défini les B-splines 1D récursivement à partir de la B-spline de degré 0, cf. (2.7). Cette définition, qui est généralement adoptée, sert d'ailleurs de base pour l'extension que constituent les hex-splines, en interprétant β^0 comme l'indicatrice de Voronoï du treillis $\mathbb{Z}$. Une autre construction est cependant possible, en remarquant que les B-splines 1D sont des $\mathcal{L}^n$ -splines pour les opérateurs de dérivation $\mathcal{L}^n = d^n/dt^n$, comme on l'a remarqué à la section 2.2.1.

Soit $n \geq 1$. On considère la B-spline causale 1D $\tilde{\beta}_n$ de degré $n-1$, notée avec en suffixe son ordre d'approximation n , et définie par $\tilde{\beta}_n(t) \triangleq \beta^{n-1}(t - \frac{n}{2})$. On peut alors écrire cette fonction sous la forme

$$\tilde{\beta}_n(t) = \Delta_c^n * \rho^n(t). \quad (3.71)$$

où :

- la fonction génératrice $\rho^n(t) = (t)_+^{n-1}/(n-1)!$ est la fonction de Green causale de l'opérateur différentiel $\mathcal{L}^n = d^n/dt^n$, c'est-à-dire que $\mathcal{L}^n \rho^n(t) = \delta(t)$, où l'on définit pour

tout $m \in \mathbb{N}$ la fonction de puissance causale

$$(t)_+^m \triangleq \begin{cases} t^m & \text{si } t > 0 \\ 0 & \text{sinon} \end{cases}. \quad (3.72)$$

• Δ^n est un filtre discret, interprété dans le domaine continu comme un peigne de Dirac pondéré sur le treillis $\mathbb{Z}$: $\Delta_c(t) = \sum_{k \in \mathbb{Z}} \Delta[k] \delta(t - k)$. Il s'agit de l'opérateur de différence finie d'ordre n , dont la transformée en $\mathbf{z}$ est $\Delta^n(z) = (1 - z^{-1})^n$.

Ainsi, une spline polynomiale de degré $n - 1$ avec ses nœuds positionnés sur le treillis $\mathbb{Z}$ est une fonction qui, si on la dérive n fois, donne un peigne de Dirac pondéré uniforme. L'opérateur Δ^n agit comme un opérateur de localisation sur la fonction génératrice ρ^n , puisque leur combinaison fournit une B-spline à support compact.

En domaine de Fourier, la décomposition (3.71) apparaît naturellement et distinctement :

$$\hat{\beta}^{n-1}(\omega) = \frac{(1 - e^{-I\omega})^n}{(I\omega)^n} = \frac{\hat{\Delta}_c^n(\omega)}{\hat{\mathbf{L}}^n(\omega)} = \hat{\Delta}_c^n(\omega) \hat{\rho}^n(\omega). \quad (3.73)$$

En effet, d'après la théorie des distributions,

$$\rho^n(t) = \frac{(t)_+^{n-1}}{(n-1)!} \xleftrightarrow{\mathcal{F}} \hat{\rho}^n(\omega) = \frac{1}{(I\omega)^n} + \frac{\pi I^{n-1} \delta^{(n-1)}(\omega)}{(n-1)!}, \quad (3.74)$$

où $\delta^{(n)}$ est la distribution dérivée n -ième de Dirac. La distribution dont la transformée de Fourier est $\pi I^{n-1} \delta^{(n-1)} / (n-1)!$ ne joue aucun rôle, puisqu'elle est dans le noyau des opérateurs $\mathcal{L}^n$ et Δ^n .

Sur le treillis 2D orthogonal, cette construction se généralise aisément, et permet de définir les B-splines séparables causales $\tilde{\beta}^n(\mathbf{x})$ en utilisant les fonctions génératrices séparables $(\mathbf{x})_+^n / (n!)^2 = (x_1)_+^n (x_2)_+^n / (n!)^2$ et les opérateurs de localisation séparables $\Delta^n(\mathbf{z}) = \Delta^n(z_1) \Delta^n(z_2)$. L'opérateur différentiel associé à $\tilde{\beta}^{n-1}$ est alors

$$\mathcal{L}^n = \frac{\partial^{2n}}{\partial x_1^n \partial x_2^n} = \mathbf{D}_{\mathbf{r}_1}^n \mathbf{D}_{\mathbf{r}_2}^n \xleftrightarrow{\mathcal{F}} \hat{\mathbf{L}}^n(\omega) = (I\langle \omega, \mathbf{r}_1 \rangle)^n (I\langle \omega, \mathbf{r}_2 \rangle)^n, \quad (3.75)$$

où l'on définit l'opérateur de différentiation directionnelle

$$\mathbf{D}_{\mathbf{v}} f(\mathbf{x}) = \lim_{t \rightarrow 0} \frac{f(\mathbf{x} + t\mathbf{v}) - f(\mathbf{x})}{t}, \quad (3.76)$$

et où les vecteurs $\mathbf{r}_1 = [1 \ 0]^T$ et $\mathbf{r}_2 = [0 \ 1]^T$ sont les vecteurs définissant le treillis Λ_{orth} .

Nous allons maintenant étendre cette construction aux box-splines à trois directions, en montrant que χ_{2n} est une $\mathcal{L}^n$ -spline pour un opérateur différentiel $\mathcal{L}^n$ à trois directions.

Remarquons que nous avons pris soin de construire de manière différentielle les B-splines causales, et non pas centrées. En effet, une B-spline centrée de degré n pair n'est pas une $\mathcal{L}^n$ -spline, car si on lui applique $\mathcal{L}^n$, on trouve bien un peigne de Dirac pondéré, mais pas positionné sur le treillis $\mathbb{Z}$. Cette pathologie disparaît avec les box-splines, et on peut donc considérer directement les box-splines centrées.

3.4.2 Caractérisation différentielle des box-splines à trois directions

En nous inspirant de la construction des B-splines utilisant des fonctions de Green, nous proposons de l'étendre aux box-splines sur le treillis hexagonal. Définissons les trois vecteurs suivants, le long des trois directions naturelles du treillis hexagonal :

$$\mathbf{r}_1 = \sqrt{\frac{2}{\sqrt{3}}} \begin{bmatrix} 1 \\ 0 \end{bmatrix}, \quad \mathbf{r}_2 = \sqrt{\frac{2}{\sqrt{3}}} \begin{bmatrix} 1/2 \\ \sqrt{3}/2 \end{bmatrix}, \quad \mathbf{r}_3 = \sqrt{\frac{2}{\sqrt{3}}} \begin{bmatrix} 1/2 \\ -\sqrt{3}/2 \end{bmatrix}. \quad (3.77)$$

On peut alors définir l'opérateur différentiel à trois directions, pour tout $n \geq 1$,

$$\mathcal{L}^n = D_{\mathbf{r}_1}^n D_{\mathbf{r}_2}^n D_{\mathbf{r}_3}^n. \quad (3.78)$$

Sa transformée de Fourier, au sens des distributions, est

$$\widehat{\mathcal{L}}^n(\boldsymbol{\omega}) = (I\langle \boldsymbol{\omega}, \mathbf{r}_1 \rangle)^n (I\langle \boldsymbol{\omega}, \mathbf{r}_2 \rangle)^n (I\langle \boldsymbol{\omega}, \mathbf{r}_3 \rangle)^n. \quad (3.79)$$

En utilisant les vecteurs $\mathbf{r}_1, \mathbf{r}_2, \mathbf{r}_3$, on peut réécrire la transformée de Fourier des box-splines à trois directions comme :

$$\widehat{\chi}_{2n}(\boldsymbol{\omega}) = \left(\prod_{i=1}^3 \text{sinc}(\langle \boldsymbol{\omega}, \mathbf{r}_i \rangle / 2) \right)^n \quad (3.80)$$

$$= \frac{(e^{I\langle \boldsymbol{\omega}, \mathbf{r}_1 \rangle} \prod_{i=1}^3 1 - e^{-I\langle \boldsymbol{\omega}, \mathbf{r}_i \rangle})^n}{(I\langle \boldsymbol{\omega}, \mathbf{r}_1 \rangle)^n (I\langle \boldsymbol{\omega}, \mathbf{r}_2 \rangle)^n (I\langle \boldsymbol{\omega}, \mathbf{r}_3 \rangle)^n}. \quad (3.81)$$

On reconnaît $\widehat{\mathcal{L}}^n(\boldsymbol{\omega})$ au dénominateur de (3.81). Le numérateur correspond, quant à lui, à la transformée de Fourier $\widehat{\Delta}_c^n(\boldsymbol{\omega})$ d'un peigne de Dirac pondéré localisé sur le treillis Λ_{hex} , que l'on peut identifier avec le filtre discret Δ^n défini par

$$\Delta^n(\mathbf{z}) = (z_1 - 1)^n (1 - z_2^{-1})^n (1 - z_1 z_2^{-1})^n. \quad (3.82)$$

En effet, du fait que $\mathbf{r}_1 = \mathbf{r}_2 + \mathbf{r}_3$,

$$\widehat{\Delta}_c^n(\boldsymbol{\omega}) = \widehat{\Delta}^n(\mathbf{R}^T \boldsymbol{\omega}) = \Delta^n(e^{I\langle \boldsymbol{\omega}, \mathbf{r}_1 \rangle}, e^{I\langle \boldsymbol{\omega}, \mathbf{r}_2 \rangle})^n \quad (3.83)$$

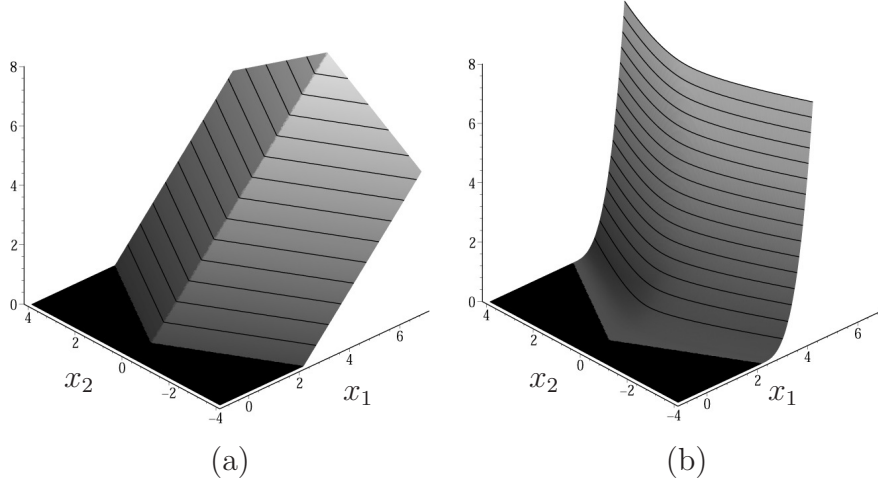


FIG. 3.9 : Les fonctions de Green $\rho^1 = \mu^{1,0}$ (a), et $\rho^2 = \mu^{3,1} + 2\mu^{4,0}$ (b), servant à générer les box-splines χ_2 et χ_4 , respectivement.

fournit exactement le numérateur de (3.81). On peut expliciter les coefficients du filtre Δ^n en développant (3.82) et collectant les coefficients devant les termes $z_1^{-k_1} z_2^{-k_2}$. On obtient, pour tout $k_1, k_2 \in \mathbb{Z}$,

$$\Delta^n[k_1, k_2 - k_1] = \sum_{i=\max(k_1, k_2, 0)}^{\min(n+k_1, n+k_2, n)} (-1)^{k_1+k_2+i} \binom{n}{i-k_1} \binom{n}{i-k_2} \binom{n}{i}. \quad (3.84)$$

En positionnant les coefficients $\Delta^n[\mathbf{k}]$ aux sites $\mathbf{Rk}$, on peut représenter les deux premiers filtres de localisation par

$$\Delta = \begin{pmatrix} & & 1 & -2 & 1 \\ & -1 & 1 & & \\ 1 & 0 & -1 & & \\ & -1 & 1 & & \end{pmatrix}, \quad \Delta^2 = \begin{pmatrix} & & 1 & -2 & 1 \\ & -2 & 2 & 2 & -2 \\ 1 & 2 & -6 & 2 & 1 \\ & -2 & 2 & 2 & -2 \\ & 1 & -2 & 1 & \end{pmatrix}. \quad (3.85)$$

Il nous faut maintenant trouver une fonction de Green $\rho^n(\mathbf{x})$ de l'opérateur $\mathcal{L}^n$, afin d'avoir la caractérisation en domaine spatial de χ_{2n} :

$$\chi^n(\mathbf{x}) = \Delta_c^n * \rho^n(\mathbf{x}) = \sum_{\mathbf{k} \in \mathbb{Z}^2} \Delta^n[\mathbf{k}] \rho^n(\mathbf{x} - \mathbf{Rk}) \xleftrightarrow{\mathcal{F}} \widehat{\chi}_{2n}(\boldsymbol{\omega}) = \frac{\widehat{\Delta}^n(\mathbf{R}^T \boldsymbol{\omega})}{\widehat{L}^n(\boldsymbol{\omega})} = \widehat{\Delta}^n(\mathbf{R}^T \boldsymbol{\omega}) \widehat{\rho}^n(\boldsymbol{\omega}). \quad (3.86)$$

PROPOSITION : Une fonction de Green $\rho^n(\mathbf{x})$ de l'opérateur $\mathcal{L}^n$, $n \geq 1$, est donnée par

$$\rho^n(\mathbf{x}) = \sum_{i=0}^{n-1} \binom{n-1+i}{i} \mu^{n-1-i, 2n-1+i}(\mathbf{x}), \quad (3.87)$$

où l'on définit pour tous entiers n_1, n_2 ,

$$\mu^{n_1, n_2}(\mathbf{x}) = \frac{1}{n_1! n_2!} \left(\sqrt{\frac{2}{\sqrt{3}}} |x_2| \right)^{n_2} \left(\sqrt{\frac{\sqrt{3}}{2}} x_1 - \frac{|x_2|}{\sqrt{2\sqrt{3}}} \right)_+^{n_1}. \quad (3.88)$$

Les fonctions μ^{n_1, n_2} et ρ^n ont toutes le même support en forme de coin. Elles sont causales selon x_1 et symétriques selon x_2 , comme illustré sur la figure 3.9.

PREUVE : Vérifions que ρ^n est bien une fonction de Green de $\mathcal{L}^n$, c.-à-d. $\mathcal{L}^n \rho^n(\mathbf{x}) = \delta(\mathbf{x})$. Pour cela, introduisons les vecteurs

$$\mathbf{r}_1^\perp = \sqrt{\frac{2}{\sqrt{3}}} \begin{bmatrix} 0 \\ 1 \end{bmatrix}, \quad \mathbf{r}_2^\perp = \sqrt{\frac{2}{\sqrt{3}}} \begin{bmatrix} \sqrt{3}/2 \\ -1/2 \end{bmatrix}, \quad \mathbf{r}_3^\perp = \sqrt{\frac{2}{\sqrt{3}}} \begin{bmatrix} \sqrt{3}/2 \\ 1/2 \end{bmatrix}, \quad (3.89)$$

qui permettent d'exprimer les bases duales de $(\mathbf{r}_1, \mathbf{r}_2)$ et $(\mathbf{r}_1, \mathbf{r}_3)$ comme $(\mathbf{r}_2^\perp, \mathbf{r}_1^\perp)$ et $(\mathbf{r}_3^\perp, -\mathbf{r}_1^\perp)$, respectivement. Par exemple, les coordonnées de $\mathbf{x}$ dans la base $(\mathbf{r}_1, \mathbf{r}_2)$ sont $(\langle \mathbf{x}, \mathbf{r}_2^\perp \rangle, \langle \mathbf{x}, \mathbf{r}_1^\perp \rangle)$.

À l'aide de ces vecteurs, on peut écrire $\mu^{n_1, n_2}(\mathbf{x})$ plus facilement :

$$\mu^{n_1, n_2}(x_1, x_2) = (x_2)_+^0 \mu^{n_1, n_2}(x_1, x_2) + (-x_2)_+^0 \mu^{n_1, n_2}(x_1, x_2) \quad (3.90)$$

$$= (\langle \mathbf{x}, \mathbf{r}_2^\perp \rangle)_+^{n_1} (\langle \mathbf{x}, \mathbf{r}_1^\perp \rangle)_+^{n_2} / n_1! n_2! + (\langle \mathbf{x}, \mathbf{r}_3^\perp \rangle)_+^{n_1} (\langle \mathbf{x}, -\mathbf{r}_1^\perp \rangle)_+^{n_2} / n_1! n_2!. \quad (3.91)$$

En utilisant un produit tensoriel et un changement de base (de Jacobien 1) de la base canonique vers la base $(\mathbf{r}_2, \mathbf{r}_1)$, on obtient à partir de (3.74)

$$(x_2)_+^0 \mu^{n_1, n_2}(x_1, x_2) \xleftrightarrow{\mathcal{F}} \frac{1}{(I\langle \boldsymbol{\omega}, \mathbf{r}_1 \rangle)^{n_1+1} (I\langle \boldsymbol{\omega}, \mathbf{r}_2 \rangle)^{n_2+1}} + \mathcal{D}(\boldsymbol{\omega}), \quad (3.92)$$

où $\mathcal{D}$ est une distribution basée sur des dérivées de Dirac. On omet (abusivement) ce terme dans les formules suivantes, car il s'annule lorsqu'on lui applique l'opérateur $\mathcal{L}^n$ ou Δ^n (voir pour plus de détails [229, Appendix C]). Il n'a donc aucune influence sur la démonstration.

De manière similaire, on obtient la transformée de Fourier de $(-x_2)_+^0 \mu^{n_1, n_2}$, ce qui donne finalement

$$\mu^{n_1, n_2} \xleftrightarrow{\mathcal{F}} \frac{1}{(I\langle \boldsymbol{\omega}, \mathbf{r}_1 \rangle)^{n_1+1} (I\langle \boldsymbol{\omega}, \mathbf{r}_2 \rangle)^{n_2+1}} + \frac{1}{(I\langle \boldsymbol{\omega}, \mathbf{r}_1 \rangle)^{n_1+1} (I\langle \boldsymbol{\omega}, \mathbf{r}_3 \rangle)^{n_2+1}}. \quad (3.93)$$

Définissons maintenant les fonctions γ^{n_1, n_2, n_3} , pour tous entiers n_1, n_2, n_3 , par

$$\gamma^{n_1, n_2, n_3} \xleftrightarrow{\mathcal{F}} \frac{1}{(I\langle \boldsymbol{\omega}, \mathbf{r}_1 \rangle)^{n_1} (I\langle \boldsymbol{\omega}, \mathbf{r}_2 \rangle)^{n_2} (I\langle \boldsymbol{\omega}, \mathbf{r}_3 \rangle)^{n_3}}. \quad (3.94)$$

Définissons $\rho^n = \gamma^{n, n, n}$, $\forall n \geq 1$, qui est bien une fonction de Green de $\mathcal{L}^n$, et exprimons cette fonction à partir des μ^{n_1, n_2} .

À partir de la propriété $\mathbf{r}_1 = \mathbf{r}_2 + \mathbf{r}_3$, on montre aisément la relation de récurrence suivante : pour tous $n_1 \geq 0, n_2 \geq 1, n_3 \geq 1$,

$$\gamma^{n_1, n_2, n_3} = \gamma^{n_1+1, n_2, n_3-1} + \gamma^{n_1+1, n_2-1, n_3}. \quad (3.95)$$

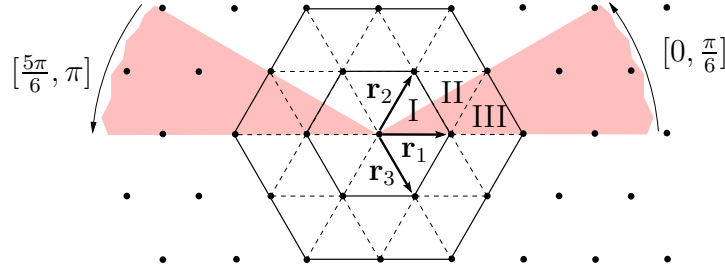


FIG. 3.10 : Grâce à la symétrie d'ordre 12 des box-splines, il suffit de savoir évaluer χ_{2n} dans un secteur d'angle $\pi/6$. Ainsi, il suffit de connaître la forme polynomiale de χ_4 à l'intérieur des triangles I, II, III, qui intersectent le secteur $[0, \frac{\pi}{6}]$.

À l'aide de cette relation, et par récurrence sur $n_2 + n_3$, on peut aussi montrer que

$$\gamma^{n_1, n_2, n_3} = \sum_{i=0}^{n_3-1} \binom{n_2-1+i}{i} \gamma^{n_1+n_2+i, 0, n_3-i} + \sum_{i=0}^{n_2-1} \binom{n_3-1+i}{i} \gamma^{n_1+n_3+i, n_2-i, 0}. \quad (3.96)$$

On a donc

$$\rho^n = \sum_{i=0}^{n-1} \binom{n-1+i}{i} (\gamma^{2n+i, 0, n-i} + \gamma^{2n+1, n-i, 0}). \quad (3.97)$$

En remarquant que (3.93) peut s'écrire

$$\mu^{n_1, n_2} = \gamma^{n_1+1, n_2+1, 0} + \gamma^{n_1+1, 0, n_2+1}, \quad (3.98)$$

et en identifiant (3.98) avec (3.97), on aboutit bien à (3.87). $\square$.

L'expression analytique complète de $\chi_{2n}(\mathbf{x})$, $n \geq 1$, peut ainsi être explicitée :

$$\begin{aligned} \chi_{2n}(x_1, x_2) &= \sum_{k_1, k_2=-n}^n \sum_{i=\max(k_1, k_2, 0)}^{\min(n+k_1, n+k_2, n)} (-1)^{k_1+k_2+i} \binom{n}{i-k_1} \binom{n}{i-k_2} \binom{n}{i} \\ &\times \sum_{d=0}^{n-1} \binom{n-1+d}{d} \frac{1}{(2n-1+d)! (n-1-d)!} \\ &\times \left| \sqrt{\frac{2}{\sqrt{3}}} x_2 + k_1 - k_2 \right|^{n-1-d} \left(\sqrt{\frac{\sqrt{3}}{2}} x_1 - \frac{k_1+k_2}{2} - \left| \frac{x_2}{\sqrt{2\sqrt{3}}} + \frac{k_1-k_2}{2} \right| \right)_+^{2n-1+d}. \end{aligned} \quad (3.99)$$

3.4.3 Mise en œuvre

L'équation (3.99) fournit un moyen réellement efficace pour effectuer l'évaluation ponctuelle en un point $\mathbf{x}$ d'une box-spline χ_{2n} . Remarquons que les fonctions puissances croissent rapidement en s'éloignant de l'origine, comme le montre la figure 3.9. Une implémentation naïve pourrait donc amener à des problèmes de stabilité numérique. Un remède

simple consiste à n'évaluer χ_{2n} que pour $x_1 \leq 0$, afin d'exploiter la causalité de ρ^n en x_1 et la symétrie $\chi_{2n}(x_1, x_2) = \chi_{2n}(-x_1, x_2)$. Il est préférable d'implémenter efficacement la fonction $t^n/n!$, par exemple avec une boucle (en langage Matlab, `val=1; for i=n:-1:1, val=val*(t/i); end`). Remarquons que cette fonction est au pire évaluée en $t = 2n$. Or la fonction $(2n)^n/n!$ croît de manière raisonnable : elle vaut environ 10^5 pour $n = 8$. Il est peu probable que des box-splines d'ordres si élevés aient à être utilisées.

Le code Matlab suivant effectue une série d'évaluations d'une box-spline, à une suite de points $(x[m], y[m])$, indexée par m . La symétrie d'ordre 12 est utilisée afin de replier les coordonnées des points dans le secteur $[5\pi/6, \pi]$, illustré sur la figure 3.10, car le nombre d'évaluations des fonctions puissances est minimal dans cette partie du plan. On utilise les coordonnées (u, v) dans la base $(\mathbf{r}_3, \mathbf{r}_2)$, au lieu des coordonnées (x, y) dans la base canonique. `nchoosek(n, k)` est la commande Matlab calculant le coefficient binomial $\binom{n}{k}$. Remarquons que ce code pourrait être amélioré car l'implémentation des fonctions du type $t^n/n!$ est ici naïve.

```
function val=boxspline(x,y,n)
d=sqrt(sqrt(3)/2);
x=-abs(x)*d;    y=abs(y)*d;
u=x-y/sqrt(3);  v=x+y/sqrt(3);
id=find(v>0);   v(id)=-v(id); u(id)=u(id)+v(id);
id=find(v>u/2); v(id)=u(id)-v(id);
val=zeros(size(x));
for K=-n:ceil(max(max(u)))-1,
    for L=-n:ceil(max(max(v)))-1,
        for i=0:min(n+K,n+L),
            coeff=(-1)^(K+L+i)*nchoosek(n,i-K)*nchoosek(n,i-L)*nchoosek(n,i);
            for d=0:n-1,
                aux=abs(v-L-u+K);
                aux2=(u-K+v-L-aux)/2;
                aux2(find(aux2<0))=0;
                val=val+coeff*nchoosek(n-1+d,d)/factorial(2*n-1+d)/factorial(n-1-d)*...
                    aux.^(n-1-d).*aux2.^(2*n-1+d);
            end,
        end,
    end,
end
```

Nous avons généré les images de la figure 3.6 au moyen de cet algorithme. Son temps de calcul est polynomial en n , et non pas exponentiel comme avec les méthodes récursives décrites dans la littérature [63, 138, 39, 133]. Par exemple, l'évaluation `boxspline(1,1,3)` a nécessité 0.002s sur notre PC portable à 1.6GHz, alors que 47s ont été requises pour la même opération avec le code Matlab proposé par De Boor [39] (qui est cependant plus général que le nôtre, puisqu'il permet d'évaluer n'importe quelle box-spline, et pas seulement celles à trois directions).

Il est aussi possible d'envisager une implémentation hybride, mi-formelle mi-numérique, sur le modèle des algorithmes d'évaluation des hex-splines proposé dans [229], pour évaluer

$\chi_{2n}(\mathbf{x})$ en un point $\mathbf{x}_0$. La forme polynomiale $p(\mathbf{x})$ de la box-spline $\chi_{2n}(\mathbf{x})$, à l'intérieur du triangle contenant $\mathbf{x}_0$, est d'abord déterminée en opérant des simplifications dans (3.99). Il suffit ensuite d'évaluer $p(\mathbf{x}_0)$.

Lorsque l'on désire effectuer de multiples évaluations de χ_{2n} , pour n fixé, il est intéressant de disposer d'un algorithme spécifiquement dédié au calcul de cette fonction. Pour cela, la forme polynomiale de χ_{2n} dans chacun des triangles composant son support est précalculée et stockée. L'algorithme ci-dessous a été écrit suivant ce principe, pour l'évaluation de χ_4 . Les coordonnées sont d'abord repliées dans le secteur angulaire $[0, \pi/2]$, puis dans $[0, \pi/3]$, et enfin dans $[0, \pi/6]$. Les coordonnées (u, v) dans la base $(\mathbf{r}_3, \mathbf{r}_2)$ permettent d'utiliser aisément ces symétries du treillis. Les coordonnées $(g = u - v/2, v)$ dans la base orthogonale $(\mathbf{r}_3, (\mathbf{r}_1 + \mathbf{r}_2)/2)$ sont les plus appropriées pour avoir des expressions polynomiales simples avec des coefficients rationnels.

```
float boxspline2(float x, float y) {
    float d=sqrt(sqrt(3.0)/2.0);
    float u=(fabs(x)-fabs(y)/sqrt(3.0))*d;
    float v=(fabs(x)+fabs(y)/sqrt(3.0))*d;
    if (u<0) { u=-u; v=v+u; } /*symétrie par rapport à r2*/
    if (2*u<v) u=v-u; /*symétrie par rapport à r1+r2*/
    float g=u-v/2.0;
    if (v>2.0) return 0.0; /*extérieur du support*/
    if (v<1.0) return 0.5+((5/3.0-v/8.0)*v-3)*v*v/4.0+
        ((1-v/4.0)*v+g*g/6.0-1)*g*g; /*triangle I*/
    if (u>1.0) return (v-2)*(v-2)*(v-2)*(g-1)/6.0; /*triangle III*/
    return 5/6.0+((1+(1/3.0-v/8.0)*v)*v/4.0-1)*v+
        ((1-v/4.0)*v+g*g/6.0-1)*g*g; /*triangle II*/
}
```

Précisons que nos algorithmes permettent d'évaluer des box-splines à trois directions $\chi'_{2n}(\mathbf{x})$ déployées sur n'importe quel treillis de matrice $\mathbf{R}'$, c'est-à-dire formées à partir de trois vecteurs $\mathbf{r}'_1, \mathbf{r}'_2, \mathbf{r}'_2 - \mathbf{r}'_1$. Il suffit pour cela d'effectuer un changement de base $\chi'_{2n}(\mathbf{x}) = \chi_{2n}(\mathbf{R}\mathbf{R}'^{-1}\mathbf{x})$.

3.5 Quasi-interpolation optimale sur la grille hexagonale

Nous avons présenté dans les deux sections précédentes deux familles de splines multidimensionnelles pouvant être utilisées sur n'importe quel treillis, et en particulier sur le treillis hexagonal, où leurs propriétés (symétrie, support compact, expression polynomiale par morceaux relativement simple, ordre d'approximation arbitraire...) sont des plus intéressantes. En suivant les mêmes principes qu'au chapitre précédent, nous allons maintenant nous intéresser au choix du préfiltre p , intervenant dans la reconstruction d'une fonction

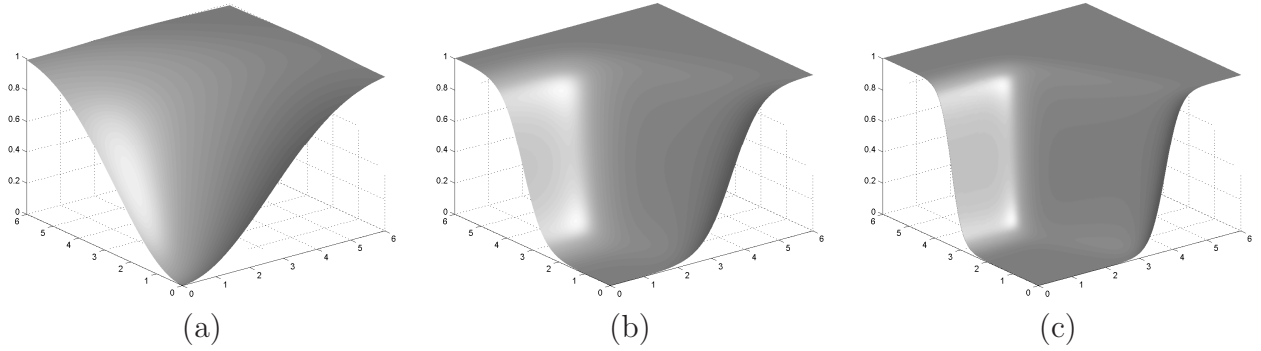


FIG. 3.11 : Noyaux d'erreur $E_{\min}(\omega)$ associés à la projection orthogonale spline sur le treillis hexagonal. (a) $\varphi = \eta_1$, (b) $\varphi = \eta_2$, (c) $\varphi = \chi_4$. Lorsque l'ordre d'approximation de φ augmente, le noyau d'erreur s'annule de plus en plus dans la région de Nyquist, de forme hexagonale.

$f_T \in V_T(\varphi)$ à partir de mesures sur le treillis hexagonal, lorsque φ est une hex-spline ou une box-spline à trois directions.

Au chapitre précédent, nous nous sommes concentrés sur les propriétés asymptotiques du schéma d'approximation $\mathcal{Q} : s \mapsto f_T$, afin d'avoir une reconstruction optimale des basses fréquences de s . Cette prédominance des basses fréquences dans les signaux rencontrés en pratique est encore plus vraie dans le cas des images naturelles ou biomédicales. Il est donc opportun d'étendre l'approche asymptotique du chapitre précédent au cas multi-dimensionnel. On cherche ainsi à avoir un noyau d'erreur (3.18) le plus plat possible à l'origine, c'est-à-dire que l'on cherche un préfiltre p tel que

$$E(\omega) \sim E_{\min}(\omega) \sim C_{\min}(\theta)^2 \|\omega\|^{2L}, \quad (3.100)$$

où L est l'ordre d'approximation de φ . Nous avons représenté sur la figure 3.11 des noyaux d'erreur $E_{\min}(\omega)$ hex-spline et box-spline. Comme en 1D, l'ordre d'approximation détermine la qualité de reconstruction dans toute la bande de Nyquist, et on peut donc s'attendre à ce que la minimisation de $E(\omega)$ à l'origine ait un effet favorable non seulement dans les basses fréquences, mais aussi dans toute la bande de Nyquist.

La contrainte (3.100) se réécrit simplement :

$$\hat{p}(\mathbf{R}^T \omega) \hat{\varphi}(\omega) = \frac{\hat{\varphi}(\omega)^*}{\hat{a}_{\varphi}(\mathbf{R}^T \omega)} + O(\|\omega\|^N), \quad (3.101)$$

pour un certain $N \geq L + 1$.

Avant d'étudier la reconstruction sur treillis hexagonal, montrons que sur le treillis orthogonal, l'utilisation du produit tensoriel des filtres 1D du chapitre 2 (comme on l'a fait dans la section 2.6 pour la rotation d'images) est licite. On suppose que $\mathbf{R}$ est l'identité, et que $\tilde{\varphi}$ et φ sont séparables. On considère le préfiltre séparable $P(\mathbf{z}) = P(z_1)P(z_2)$, où p

est un filtre 1D asymptotiquement optimal, vérifiant (2.71). On a alors

$$\hat{p}(\boldsymbol{\omega})\hat{\varphi}(\boldsymbol{\omega}) = \hat{p}(\omega_1)\hat{\varphi}(\omega_1)\hat{p}(\omega_2)\hat{\varphi}(\omega_2) = \varphi_d(\omega_1)\varphi_d(\omega_2) + O(w_1^N + w_2^N) = \varphi_d(\boldsymbol{\omega}) + O(\|\boldsymbol{\omega}\|^N). \quad (3.102)$$

donc le préfiltre séparable est bien asymptotiquement optimal au sens de (3.100).

Intéressons-nous maintenant en détail à la reconstruction sur treillis hexagonal, lorsque $\tilde{\varphi} = \delta$. Comme $\hat{a}_\varphi(\boldsymbol{\omega}) = |\hat{\varphi}(\boldsymbol{\omega})|^2 + O(\|\boldsymbol{\omega}\|^{2L})$ d'après la formule de Poisson (3.7) et les conditions de Strang-Fix (3.11), la conception d'un préfiltre de quasi-interpolation asymptotiquement optimal se ramène à la condition particulièrement simple, portant sur le développement de Taylor du filtre : pour un certain $N \geq L + 1$, on veut

$$\hat{p}(\boldsymbol{\omega}) = \frac{1}{\hat{\varphi}(\boldsymbol{\omega})} + O(\|\boldsymbol{\omega}\|^N) \quad (3.103)$$

$$\Leftrightarrow \frac{1}{\hat{p}(\boldsymbol{\omega})} = \hat{\varphi}(\boldsymbol{\omega}) + O(\|\boldsymbol{\omega}\|^N). \quad (3.104)$$

Puisque l'on peut choisir la forme du préfiltre, nous proposons trois formes différentes. Dans tous les cas, nous cherchons un filtre ayant la même symétrie d'ordre 12 que le treillis et φ . Nous allons utiliser les fonctions $\text{ring}_n(\mathbf{z})$ définies par la figure 3.4, afin d'exprimer aisément les filtres dans le domaine en $\mathbf{z}$.

- Nous cherchons en premier lieu des filtres RIF, notés p_{RIF} , puisque l'implémentation du préfiltrage avec de tels filtres est particulièrement simple et rapide. De plus, ces filtres agissent sur les données de manière locale ; cela est très intéressant lorsque l'on manipule des données de taille importante, car le traitement peut alors être effectué en accédant localement aux données, sans avoir à charger en mémoire toute l'image. Il faut donc déterminer les coefficients $h[k]$ tels que

$$P_{\text{RIF}}(\mathbf{z}) = \sum_{k=0}^K h[k] \text{ring}_k(\mathbf{z}). \quad (3.105)$$

- Le second type de filtres auquel nous nous intéressons est le pendant 2D des filtres RII inverses, dont nous avons montré la supériorité en 1D. On cherche donc les $h[k]$ tels que p_{RII} soit un filtre RII inverse défini par

$$P_{\text{RII}}(\mathbf{z}) = \frac{1}{\sum_{k=0}^K h[k] \text{ring}_k(\mathbf{z})}. \quad (3.106)$$

Un tel préfiltre n'a pas d'implémentation directe en domaine spatial, car il n'y a pas d'équivalent à l'algorithme récursif rapide disponible dans le cas 1D. On peut donc soit se tourner vers des méthodes itératives de type descente de gradient [23], soit implémenter le préfiltrage en domaine de Fourier, en effectuant une FFT des données puis une multiplication

φ	η_1			χ_2			η_2			η_3			χ_4		
p	RIF	RII1	RII2	RIF	RII1	RII2	RIF	RII1	RII2	RIF	RII1	RII2	RIF	RII1	RII2
$h[0]$	$\frac{31}{36}$	$\frac{41}{36}$	$\frac{113}{108}$	$\frac{5}{4}$	$\frac{3}{4}$	$\frac{11}{12}$	$\frac{23}{18}$	$\frac{13}{18}$	$\frac{49}{54}$	$\frac{682}{405}$	$\frac{887}{1620}$	$\frac{21713}{25920}$	$\frac{37}{20}$	$\frac{29}{60}$	$\frac{97}{120}$
$h[1]$	$\frac{5}{216}$	$\frac{-5}{216}$	$\frac{-5}{216}$	$\frac{-1}{24}$	$\frac{1}{24}$	$\frac{1}{24}$	$\frac{-5}{108}$	$\frac{5}{108}$	$\frac{5}{108}$	$\frac{-883}{6480}$	$\frac{127}{1620}$	$\frac{3307}{38880}$	$\frac{-41}{240}$	$\frac{7}{80}$	$\frac{1}{10}$
$h[2]$										$\frac{433}{19440}$	$\frac{-29}{9720}$	$\frac{-607}{155520}$	$\frac{7}{240}$	$\frac{-1}{720}$	$\frac{-1}{240}$

TAB. 3.2 – Coefficients $h[k]$ de (3.105), (3.106), (3.108) pour les préfiltres de quasi-interpolation hex-spline et box-spline.

par $\hat{p}(\omega)$. C'est cette seconde solution que nous avons retenue. Sa mise en œuvre, détaillée dans [229, Appendix E], se fait au moyen de FFTs rectangulaires classiques, le long des deux directions $\mathbf{r}_1$ et $\mathbf{r}_2$ du treillis.

- On s'intéresse enfin à une troisième forme de filtre RII inverse, définie par

$$P_{\text{RII2}}(\mathbf{z}) = \frac{1}{H(z_1)H(z_2)H(z_1z_2^{-1})}, \quad (3.107)$$

où $H(z)$ la transformée en $\mathbf{z}$ d'un filtre 1D symétrique. On cherche donc les coefficients $h[k]$ tels que

$$H(z) = h[0] + \sum_{k=1}^K h[k](z^k + z^{-k}). \quad (3.108)$$

Cette structure particulière du filtre est choisie afin que le filtre soit factorisable en trois filtres 1D le long des trois directions naturelles du treillis hexagonal. Ainsi, on peut implémenter le préfiltrage avec un tel filtre de manière efficace, en effectuant de manière séparable le filtrage inverse de h , au moyen de l'algorithme récursif proposé par Unser et coll. [218]. Cette forme de filtres a déjà été proposée pour les hex-splines [80], mais sans l'optimalité au sens de (3.101).

Nous choisissons donc les coefficients $h[k]$ des filtres afin que la condition d'optimalité (3.103) ou (3.104) soit vérifiée. On choisit des filtres de taille minimale, donc avec $K = (N - 1)/2$ dans (3.105), (3.106) et (3.108). De plus, on choisit $N = L + 1$ dans (3.103) et (3.104) dans le cas des box-splines et des hex-splines d'ordre pair. Pour les hex-splines d'ordre impair, on pose $N = L + 2$, car $N = L + 1$ conduit au préfiltre d'interpolation dans le cas RII1, ainsi que dans le cas RIF pour la hex-spline d'ordre 1. Notons qu'il est tout à fait possible de prendre des valeurs de N supérieures, afin d'obtenir une qualité de reconstruction encore meilleure. Les coefficients $h[k]$ des filtres proposés sont résumés dans le tableau 3.2. Par exemple, le filtre p_{RII2}^{-1} pour la box-spline d'ordre 2 est :

$$\begin{array}{cccc}
\frac{1}{13824} & \frac{11}{6912} & \frac{1}{13824} & \\
\frac{11}{6912} & \frac{253}{6912} & \frac{253}{6912} & \frac{11}{6912} \\
\frac{1}{13824} & \frac{253}{6912} & \frac{1775}{2304} & \frac{253}{6912} & \frac{1}{13824} \\
\frac{11}{6912} & \frac{253}{6912} & \frac{253}{6912} & \frac{11}{6912} & \\
\frac{1}{13824} & \frac{11}{6912} & \frac{1}{13824} & &
\end{array} = \begin{array}{ccccccc}
& & \frac{1}{24} & & & & \\
& \frac{11}{12} & * & \frac{1}{24} & \frac{11}{12} & \frac{1}{24} & * & \frac{11}{12} \\
& & \frac{1}{24} & & & & & \frac{1}{24}
\end{array}$$

Notons que les préfiltres box-spline proposés peuvent être utilisés sur tous les treillis, et pas seulement le treillis hexagonal. En effet, les box-splines à trois directions peuvent être déployées sur n'importe quel treillis par un simple changement de base. Les hex-splines sont, elles, construites spécifiquement pour un treillis. Par conséquent, les préfiltres construits ici sont spécifiques au treillis hexagonal. Cependant, la méthode proposée est générique, et permet de construire aisément des préfiltres sur n'importe quel treillis.

Afin d'illustrer la facilité de conception des préfiltres, détaillons le calcul du filtre RIF pour la hex-spline η_2 , en utilisant (3.103) avec $N = L + 1 = 3$. Le développement de Taylor de (3.20) à l'origine donne

$$\frac{1}{\hat{\eta}_2(\boldsymbol{\omega})} = 1 + \frac{5\sqrt{3}}{108}(\omega_1^2 + \omega_2^2) + O(\|\boldsymbol{\omega}\|^4). \quad (3.109)$$

Le préfiltre recherché a la forme $P_{\text{RIF}}(\mathbf{z}) = h[0] + h[1] \text{ring}_1(\mathbf{z})$, ou de manière équivalente

$$\hat{p}_{\text{RIF}}(\boldsymbol{\omega}) = (h[0] + 6h[1]) - h[1]\sqrt{3}(\omega_1^2 + \omega_2^2) + O(\|\boldsymbol{\omega}\|^4). \quad (3.110)$$

Il suffit d'identifier les deux développements (3.109) et (3.110), ce qui donne $h[1] = -\frac{5}{108}$ et $h[0] = \frac{23}{18}$.

Notre choix de filtres asymptotiquement optimaux peut être validé en examinant les noyaux d'erreur associés. Par exemple, pour la hex-spline d'ordre 2, la figure 3.12 (a) montre le noyau d'erreur associé à l'interpolation. On voit que la différence avec $E_{\min}$, représenté sur la figure 3.11 (b), n'est pas négligeable, ce qui signifie que l'interpolation n'exploite pas pleinement les possibilités d'approximation de l'espace de reconstruction. La figure 3.12 (b) montre le noyau d'erreur associé à la quasi-interpolation avec le préfiltre RIF développé au paragraphe précédent. Celui-ci est meilleur dans toute la bande de Nyquist, comme le montre la différence entre les deux noyaux, représentée sur la figure 3.12 (c).

La figure 3.13 montre un exemple de surfaces box-spline et hex-spline $f_T(\mathbf{x})$ reconstruites par quasi-interpolation. L'utilisation d'un espace de reconstruction d'ordre d'approximation au moins égal à trois est nécessaire si l'on veut obtenir une reconstruction

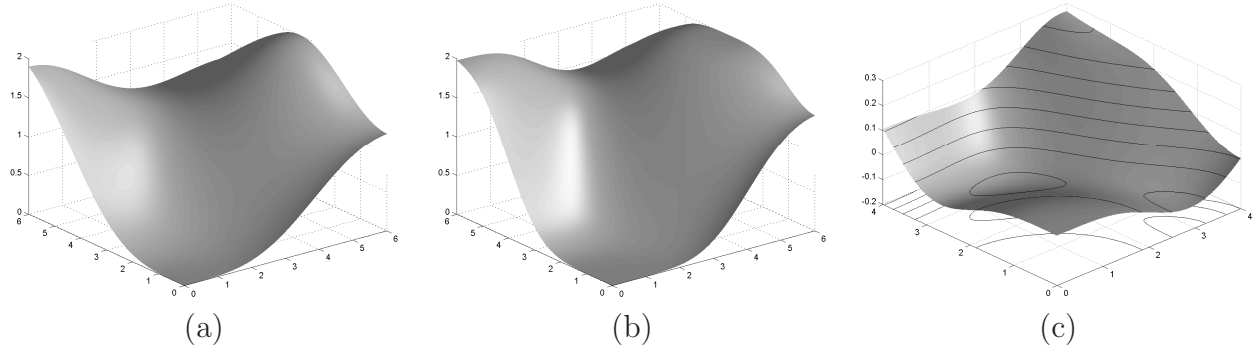


FIG. 3.12 : Noyaux d'erreur pour la reconstruction hex-spline d'ordre 2 : (a) $E_{\text{int}}(\omega)$ associé à l'interpolation (donc sans préfiltre car η_2 est interpolante); (b) $E(\omega)$ pour la quasi-interpolation avec le filtre RIF du tableau 3.2; (c) la différence $E - E_{\text{int}}$ montrant clairement l'avantage de la quasi-interpolation, qui offre une meilleure qualité d'approximation dans toute la bande de Nyquist.

lisse, sans crénelage le long des contours obliques. Les modèles d'ordre inférieur sont trop simples, et la structure du treillis apparaît en filigrane : $\varphi = \eta_1$ donne un modèle constant par morceaux, et χ_2 donne une surface plane sur chaque domaine triangulaire formé par trois sites voisins. $f_T(\mathbf{x})$ n'est continûment différentiable qu'à partir de l'ordre 3, ce qui fait de η_3 et χ_4 des candidats de choix pour les applications où le calcul de dérivées est requis.

Dans la section suivante, des expériences de rééchantillonnage sont proposées, afin de confirmer expérimentalement les bonnes propriétés des préfiltres de quasi-interpolation que nous avons conçus.

3.6 Application au rééchantillonnage hexagonal-vers-orthogonal

3.6.1 Principes et résultats expérimentaux

Dans cette section, on se concentre sur le problème du rééchantillonnage d'un signal v localisé sur le treillis hexagonal, vers le treillis orthogonal de même densité (on suppose que $T = 1$). Afin d'évaluer la qualité des différentes méthodes de reconstruction, nous proposons la procédure suivante : une image de référence v_0 est d'abord rééchantillonnée du treillis Λ_{orth} vers le treillis Λ_{hex} , par interpolation O-MOMS, c'est-à-dire en prenant pour φ la O-MOMS cubique séparable proposée dans [33]. Nous pourrions aussi utiliser la quasi-interpolation spline cubique séparable développée au chapitre 2, pour un résultat quasiment identique. La surface reconstruite par interpolation O-MOMS joue le rôle de $s(\mathbf{x})$, et est échantillonnée sur Λ_{hex} , fournissant ainsi les échantillons $v[\mathbf{k}]$. On cherche à rééchantillonner l'image v sur Λ_{orth} , c'est-à-dire que l'on cherche à estimer v_0 à partir de v . On est donc bien dans les conditions d'un problème de reconstruction, puisque connaissant

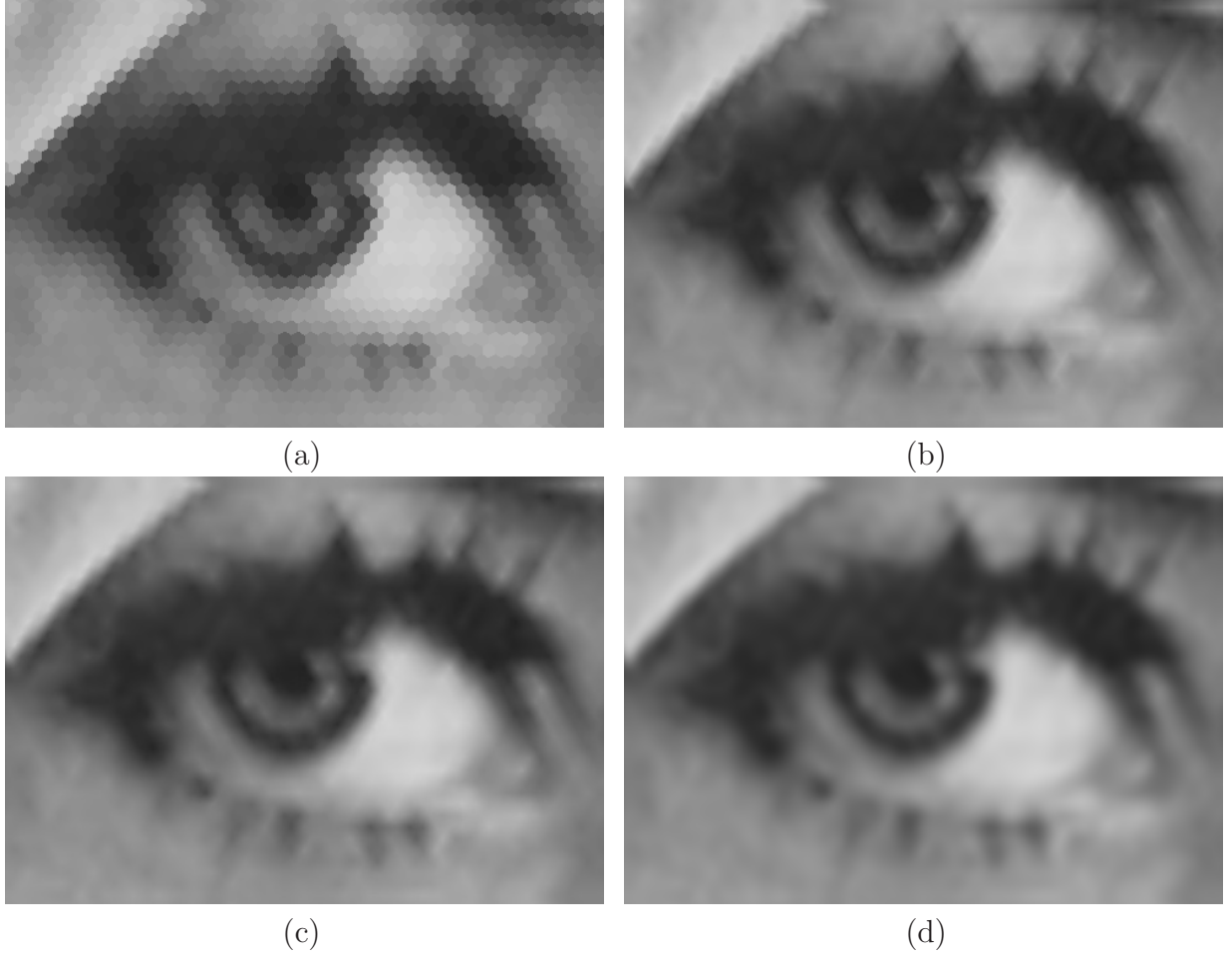


FIG. 3.13 : Reconstruction $f_T(\mathbf{x})$ de l'œil de *Lena* à partir d'échantillons localisés sur le treillis hexagonal, en prenant comme fonction génératrice φ , (a) η_1 , (b) χ_2 , (c) η_2 , (d) χ_4 . Le préfiltre de quasi-interpolation RII2 est utilisé dans tous les cas.

les valeurs ponctuelles $s(\mathbf{R}_{\text{hex}}\mathbf{k})$, on cherche à estimer les $s(\mathbf{R}_{\text{orth}}\mathbf{k})$, ce qui nécessite de reconstruire une estimée f_T de s puis de l'échantillonner sur Λ_{orth} . L'image finale rééchantillonnée peut alors être comparée à l'image de référence v_0 . Cette mesure d'erreur discrète est un bon indicateur qualitatif de l'erreur de reconstruction $\|s - f_T\|$, que l'on ne peut pas calculer directement.

Nous avons appliqué cette procédure de test à 7 images de taille 512×512 , représentées dans l'annexe 7.7. Des conditions symétriques miroir ont été utilisées aux bords des images. Les mesures de PSNR sont reportées dans le tableau 3.3, pour les différentes combinaisons (φ, p) testées pour la reconstruction. On remarque tout d'abord que pour les générateurs d'ordre 2, le gain de la quasi-interpolation sur l'interpolation est important, même en utilisant les filtres RIF de faible complexité. Pour les ordres d'approximation plus élevés,

φ	η_1				χ_2				η_2			
p	int.	RIF	RII1	RII2	int.	RIF	RII1	RII2	int.	RIF	RII1	RII2
<i>Lena</i>	35.86	35.97	35.97	35.97	42.16	44.46	44.74	44.67	41.87	44.73	45.26	45.10
<i>Barbara</i>	29.24	29.46	29.46	29.46	33.60	35.99	36.35	36.27	33.31	36.35	37.07	36.86
<i>Baboon</i>	28.09	28.22	28.23	28.23	31.88	33.87	34.11	34.06	31.57	34.13	34.69	34.53
<i>Lighthouse</i>	29.50	29.64	29.64	29.64	34.92	37.59	37.95	37.87	34.65	38.08	38.83	38.61
<i>Goldhill</i>	34.61	34.72	34.73	34.73	39.39	41.57	41.90	41.81	39.10	41.84	42.48	42.28
<i>Boat</i>	32.82	32.91	32.92	32.92	37.75	39.74	40.00	39.94	37.48	39.95	40.46	40.31
<i>Peppers</i>	35.33	35.28	35.31	35.31	39.66	41.43	41.50	41.50	39.12	41.59	42.07	41.93
temps	1	2	12	2	1	2	12	2	3	4	13	3

φ	η_3				χ_4			
p	int.	RIF	RII1	RII2	int.	RIF	RII1	RII2
<i>Lena</i>	46.75	45.93	47.37	46.65	47.16	45.18	47.57	46.34
<i>Barbara</i>	39.64	38.33	40.66	39.73	40.77	37.47	41.50	39.80
<i>Baboon</i>	36.31	35.30	37.12	36.20	36.88	34.40	37.44	35.86
<i>Lighthouse</i>	41.44	40.18	42.59	41.33	42.43	39.23	43.22	41.08
<i>Goldhill</i>	44.15	43.16	44.99	44.01	44.74	42.30	45.30	43.69
<i>Boat</i>	41.63	40.89	42.20	41.50	41.91	40.13	42.26	41.14
<i>Peppers</i>	42.95	42.20	43.47	42.78	43.09	41.33	43.40	42.33
temps	14	5	14	4	14	5	14	4

TAB. 3.3 – PSNR obtenus après rééchantillonnage de différentes images du treillis orthogonal vers hexagonal, puis inversement, en utilisant différentes méthodes d'interpolation et quasi-interpolation.

la reconstruction avec les filtres RII non séparables est meilleure que l'interpolation, mais les filtres RIF et RII séparables sont par contre moins performants. Dans tous les cas, les résultats obtenus avec le filtre séparable RII2 se situent entre ceux des filtres RIF et RII1. On note aussi que la hiérarchie entre les générateurs φ est respectée, c'est-à-dire que l'ordre d'approximation est bien le critère le plus déterminant de la qualité de reconstruction. La différence entre les box-splines et les hex-splines dépend du préfiltre : ainsi, pour l'ordre 2, les hex-splines sont légèrement meilleures dans le cas de la quasi-interpolation, alors que la tendance est inversée pour l'interpolation.

La figure 3.14 montre une partie de l'image *Barbara* après rééchantillonnage avec plusieurs méthodes. La reconstruction hex-spline d'ordre 1 souffre de distorsions importantes de type « effet de blocs », clairement visibles le long des structures obliques. Les générateurs d'ordre 2 sont meilleurs de ce point de vue, mais l'interpolation avec ces fonctions introduit un effet de flou non négligeable. Ce défaut est corrigé avec la quasi-interpolation. En augmentant l'ordre d'approximation, on obtient des images de plus en plus fidèles à l'original. L'effet de flou qui subsiste avec l'interpolation est systématiquement corrigé avec la quasi-interpolation, du moins avec les filtres RII1.

Ces résultats empiriques sont confirmés par l'examen des noyaux d'erreur associés, dont l'étude prédit bien les résultats obtenus. Pour η_1 , η_3 et χ_4 , le noyau d'erreur E_{int} est déjà très proche de E_{min} , ce qui signifie que l'on ne peut pas espérer un gain important en changeant le préfiltre. La forme inverse non séparable RII1, qui est celle du préfiltre d'interpolation, est la seule permettant une amélioration systématique par rapport à l'interpolation. Pour

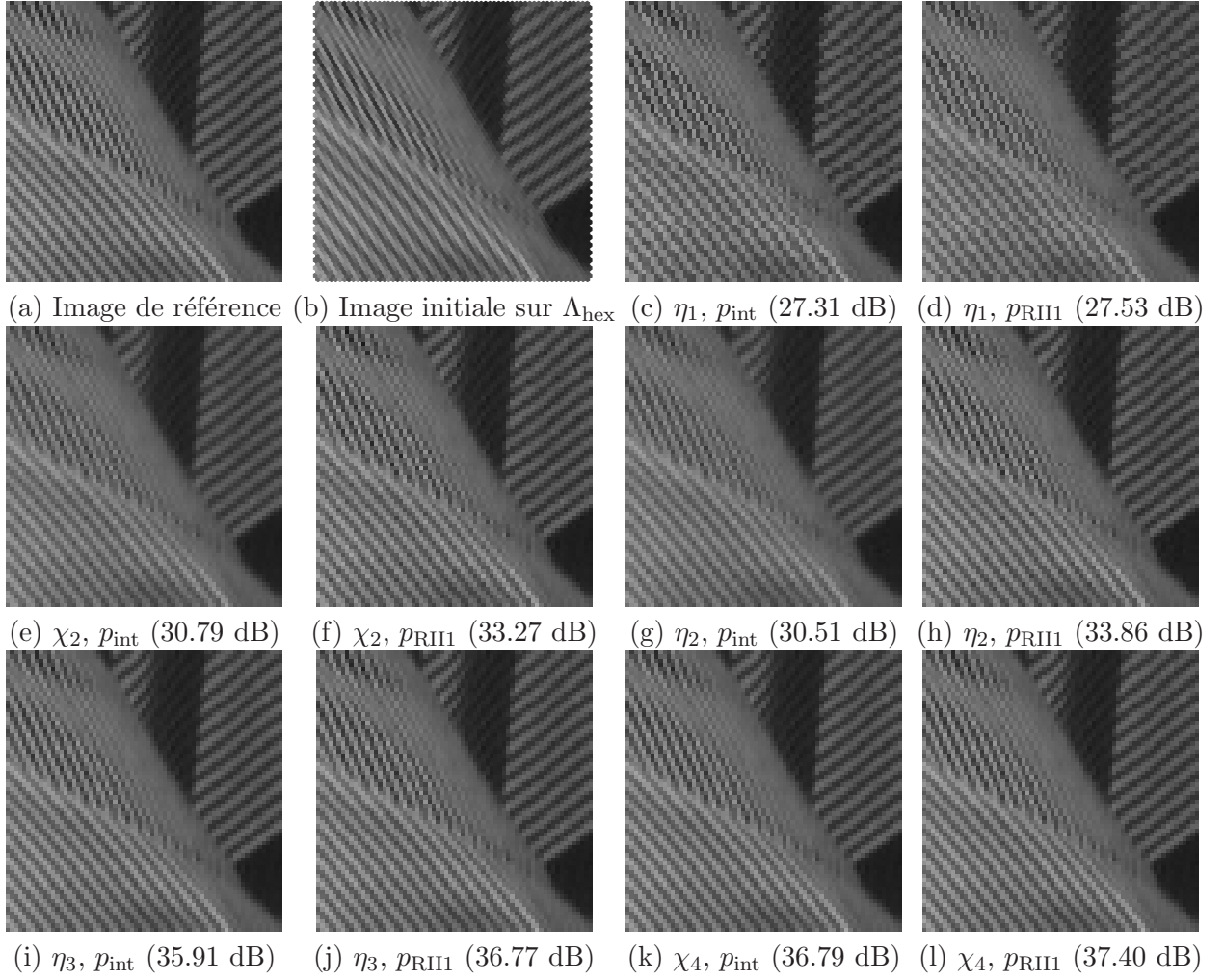


FIG. 3.14 : Partie de l'image *Barbara* après rééchantillonnage de l'image (b) sur le treillis orthogonal au moyen de différentes méthodes d'interpolation et quasi-interpolation. L'image obtenue est comparée à l'image de référence (a) ayant servi à générer l'image (b).

les filtres RIF et RII2 associés à η_3 et χ_4 , le noyau d'erreur est maximale plat au voisinage de $\omega = \mathbf{0}$, mais le comportement se dégrade dans le reste de la bande de Nyquist. Cela montre que la condition d'optimalité asymptotique n'implique pas nécessairement un noyau d'erreur meilleur dans toute la bande de Nyquist. Il vaut donc parfois mieux chercher à minimiser le noyau d'erreur globalement, et pas seulement à l'origine.

Par exemple, cherchons un filtre RIF pour la quasi-interpolation avec η_3 , avec trois coefficients $h[0], h[1], h[2]$ dans (3.105). Deux degrés de liberté sont utilisés pour assurer l'optimalité $N = L + 1$ dans (3.103) (et pas $N = L + 2$ comme pour le filtre du tableau 3.2); le troisième est optimisé manuellement afin d'avoir un noyau d'erreur minimal dans

la bande de Nyquist. Nous proposons ainsi le préfiltre

$$P(\mathbf{z}) = \frac{121}{60} - \frac{79}{360} \text{ring}_1(\mathbf{z}) + \frac{1}{20} \text{ring}_2(\mathbf{z}), \quad (3.111)$$

qui donne une amélioration importante de 1.5 dB en moyenne sur les images en comparaison avec le filtre du tableau 3.2.

La dernière ligne du tableau 3.3 donne les temps de calcul du rééchantillonnage, comparativement au temps de l'algorithme le plus rapide, qui est l'interpolation au plus proche voisin (1 unité correspond approximativement à 0.1s pour un code C exécuté sur un PC récent). Le préfiltrage avec les filtres inverses non séparables (préfiltres RII1 et d'interpolation) a été implémenté en domaine de Fourier, en calculant la FFT de v et de p . Avec une telle implémentation, la taille du préfiltre n'entre pas en compte, mais, pour une image de taille $N \times N$, le temps de calcul est en $O(N^2 \log(N))$, au lieu de $O(N^2)$ pour un préfiltrage avec un filtre RIF ou RII2. Ces deux dernières formes de filtres sont donc les plus intéressantes du point de vue du temps de calcul, qui est réduit d'un facteur important avec les splines d'ordre élevé. Associés aux générateurs splines interpolants d'ordre 1 ou 2, ces filtres fournissent une amélioration nette de la qualité pour une augmentation faible du temps de calcul. Le choix d'un préfiltre de type RII2 semble donc particulièrement pertinent dans tous les cas, surtout en tenant compte du fait que le préfiltrage avec un tel filtre n'implique que des traitements 1D, et se révèle donc encore plus rapide que le filtrage avec un filtre RIF non séparable. L'implémentation de la reconstruction, après l'étape de préfiltrage, est détaillée dans la section 3.6.2.

Remarquons que nous avons effectué le rééchantillonnage d'un treillis vers un autre de même densité. En effet, nous avons formulé le rééchantillonnage comme un problème d'interpolation, au même titre que la rotation d'image dans la section 2.6. En fait, si l'on rééchantillonne vers un treillis plus grossier, des problèmes de repliement de spectre peuvent se manifester avec cette approche. Cela s'explique par le fait que dans ce cas, ce ne sont pas des échantillons de $s(\mathbf{x})$ que l'on veut estimer (et donc $s(\mathbf{x})$ que l'on cherche à reconstruire), mais des échantillons d'une version préfiltrée de $s(\mathbf{x})$ à l'aide d'un filtre de fréquence de coupure adaptée à la zone de Nyquist du treillis d'arrivée. Si l'on désire effectuer le rééchantillonnage vers un treillis de densité plus faible, on peut donc préfiltrer v avant d'effectuer la reconstruction. Nous formaliserons de manière rigoureuse au chapitre 5 le problème de réduction de la résolution.

Notre approche est, bien entendu, aussi appropriée pour effectuer le rééchantillonnage vers un treillis de plus grande densité. Les images de la figure 3.13 sont en fait rééchantillonnées du treillis Λ_{hex} vers un treillis orthogonal de haute densité. Enfin, notons que les box-splines et hex-splines sont basées sur des principes génériques, qu'il est possible d'étendre sans mal aux dimensions supérieures. Par exemple, la visualisation de données 3D, qui constitue un champ de recherche actif [97], pourrait bénéficier de nos méthodes rapides de quasi-interpolation. L'interpolation exacte est en effet prohibitive en 3D.

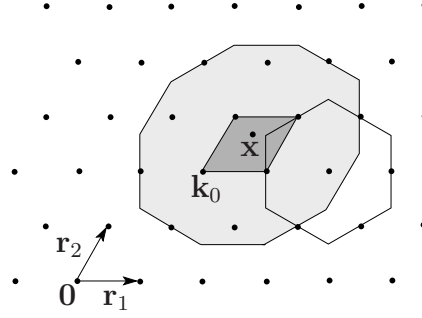


FIG. 3.15 : Pour évaluer $f_T(\mathbf{x})$ à un certain point $\mathbf{x}$, à l'aide de la formule (3.112), il faut connaître les sites $\mathbf{Rk}$ tels que $\varphi(\mathbf{x} - \mathbf{Rk}) \neq 0$. Ils appartiennent au domaine gris clair obtenu par dilatation morphologique du losange contenant $\mathbf{x}$ (en gris foncé) avec le support de φ comme élément structurant. Dans cet exemple, $\varphi = \eta_2$, dont le support est l'hexagone représenté. Celui-ci est ici centré en l'un des 8 sites susceptibles de contribuer à la valeur de $f_T(\mathbf{x})$, si l'on sait seulement de $\mathbf{x}$ qu'il est dans le losange.

3.6.2 Evaluation des modèles box-spline et hex-spline

Une fois déterminés les coefficients $c[\mathbf{k}]$ du modèle spline $f_T(\mathbf{x})$, il faut être en mesure d'évaluer efficacement cette fonction si l'on veut pouvoir l'exploiter en pratique. Puisque

$$f_T(\mathbf{x}) = \sum_{\mathbf{k} \in \mathbb{Z}^2} c[\mathbf{k}] \varphi\left(\frac{\mathbf{x}}{T} - \mathbf{Rk}\right), \quad (3.112)$$

la plupart du temps de calcul est consacré aux multiples évaluations de $\varphi(\mathbf{x})$. L'évaluation des hex-splines est discutée en détail dans [229], et nous avons proposé des méthodes efficaces pour l'évaluation des box-splines à trois directions dans la section 3.4.3. Pour l'évaluation de $\varphi = \eta_1, \chi_2, \eta_2, \eta_3, \chi_4$, le temps de calcul dépend directement du nombre de triangles élémentaires composant le support de φ (c.-à-d. 6, 6, 24, 54, 24 respectivement, voir la figure 3.8), et du degré de l'expression polynomiale de φ dans chacun des triangles (c.-à-d. 0, 1, 2, 4, 4 respectivement). Nous nous sommes limités aux cinq générateurs mentionnés car l'implémentation et le temps de calcul associés restent raisonnables, et un ordre d'approximation ≤ 4 semble suffisant pour la plupart des applications potentielles.

En fait, le point important, qu'il s'agit de détailler, dans l'évaluation de $f_T(\mathbf{x})$ à l'aide de (3.112), est de limiter le nombre d'appels à φ . Pour cela, il faut connaître les indices $\mathbf{k}$ tels que $\varphi(\mathbf{x} - \mathbf{Rk}) \neq 0$. Contrairement au cas cartésien, il n'est pas aisé de déterminer ces indices directement, lorsque φ a un support hexagonal. Nous proposons donc la stratégie suivante, qui est sous-optimale (au sens où elle ne fournit pas uniquement les $\mathbf{k}$ tels que $\varphi(\mathbf{x} - \mathbf{Rk}) \neq 0$) mais facile à mettre en œuvre.

On calcule d'abord les coordonnées (u, v) de $\mathbf{x}$ dans la base $(\mathbf{r}_1, \mathbf{r}_2)$: $[u \ v]^T = \mathbf{R}^{-1}\mathbf{x}$. En prenant les parties entières de ces coordonnées, on obtient $\mathbf{k}_0 = [\lfloor u \rfloor \ \lfloor v \rfloor]^T \in \mathbb{Z}^2$. Ainsi, le point $\mathbf{x}$ appartient au losange dont un des sommets est $\mathbf{Rk}_0$, comme illustré sur

la figure 3.15. Alors, la somme (3.112) se réduit à une somme finie portant sur les quelques indices $\mathbf{k}$ pour lesquels le support de $\varphi(\cdot - \mathbf{R}\mathbf{k})$ intersecte le losange. Ces valeurs $\mathbf{k}$ sont telles que les sites $\mathbf{R}\mathbf{k}$ sont compris dans le domaine obtenu par dilatation morphologique du losange, avec le support de φ comme élément structurant [197]. Ainsi, ce nombre de termes à calculer dans la somme, qui est directement proportionnel au temps de calcul de la valeur $f_T(\mathbf{x})$, est respectivement de 4, 4, 8, 14, 14 pour nos cinq générateurs. Par exemple, si $\varphi = \chi_2$, on peut utiliser la formule suivante quel que soit $\mathbf{x}$, une fois calculé le $\mathbf{k}_0$ associé :

$$\begin{aligned} f_T(\mathbf{x}) = & \chi_2(\mathbf{x} - \mathbf{R}\mathbf{k}_0) + \chi_2(\mathbf{x} - \mathbf{R}\mathbf{k}_0 - \mathbf{r}_1) \\ & + \chi_2(\mathbf{x} - \mathbf{R}\mathbf{k}_0 - \mathbf{r}_2) + \chi_2(\mathbf{x} - \mathbf{R}\mathbf{k}_0 - \mathbf{r}_1 - \mathbf{r}_2). \end{aligned} \quad (3.113)$$

3.7 Conclusion

Dans ce chapitre, nous avons développé des méthodes de reconstruction pour des données 2D localisées sur un treillis arbitraire, en nous concentrant sur le cas concret du treillis hexagonal. Nous avons utilisé des hex-splines et des box-splines à trois directions, pour leurs bonnes propriétés d'approximation et leur mise en œuvre relativement aisée. L'idée d'effectuer une quasi-projection asymptotiquement optimale est tout à fait nouvelle dans ce contexte. Les méthodes de quasi-interpolation spline sur grille hexagonale que nous avons conçues donnent de bien meilleurs résultats que l'interpolation, qui est pourtant systématiquement adoptée par les praticiens du traitement d'images. En outre, le préfiltre peut être choisi afin d'avoir une implémentation exacte et rapide, au contraire de l'interpolation spline qui pose problème, puisque le préfiltre d'interpolation n'est pas séparable. La quasi-interpolation est donc une alternative sérieuse à l'interpolation pour de nombreux traitements portant sur les images, et nous sommes convaincus que ce type d'approches va être amené à se développer dans un avenir proche.

Résumons les contributions nouvelles apportées dans ce chapitre :

- Nous avons étendu le noyau d'erreur fréquentiel dans le cadre multi-dimensionnel, sur un treillis quelconque. Cette extension depuis le cas 1D, qui ne pose pas de difficulté théorique particulière, devrait être justifiée par une publication ([230], en préparation).
- Nous avons mis en exergue les propriétés d'approximation avantageuses du treillis hexagonal. Ainsi, nos résultats montrent qu'un signal sur-échantillonné (c'est-à-dire dont l'énergie est concentrée dans les basses fréquences par rapport à la fréquence d'échantillonnage) pouvait être représenté sur un treillis hexagonal de densité 40% inférieure à celle requise sur un treillis cartésien. Cela fournit un argument supplémentaire en faveur du treillis hexagonal, qui malgré des avantages indéniables et connus, reste pourtant peu utilisé en pratique. Ce travail a fait l'objet de la publication [73].
- Nous avons découvert une nouvelle caractérisation des box-splines à trois directions, utilisant une localisation discrète de fonctions de Green associées à des opérateurs différentiels. Cela nous a permis de développer une implémentation très efficace, et bien plus

rapide que les méthodes récursives disponibles actuellement dans la littérature. Ce travail a été publié dans [72, 74]. Les box-splines deviennent ainsi des fonctions de choix pour la reconstruction sur treillis hexagonal, principalement grâce à leurs bonnes propriétés d'approximation, combinées à un support compact de taille réduite.

◦ Nous avons proposé des méthodes de quasi-interpolation, à la fois plus performantes et plus rapides que l'interpolation. Les avantages pour les applications de rééchantillonnage sont réels. Ce travail a fait l'objet des articles [74, 71]. Des chercheurs travaillant dans le domaine de la visualisation se sont montrés intéressés par l'extension en 3D de notre approche, qui au demeurant ne pose pas de problème. Des splines d'ordre élevé sont en effet nécessaires afin de garantir des propriétés de lissitude des modèles 3D reconstruits, mais la mise en œuvre de l'interpolation est alors très difficile. En ce qui nous concerne, nous avons obtenu des résultats préliminaires très encourageants de démosaïquage 2D : en simulant une image mosaïquée sur grille hexagonale, que l'on démosaïque puis rééchantillonne sur grille orthogonale, on obtient de meilleurs résultats qu'en démosaïquant une image acquise directement sur la grille orthogonale, comme le font tous les appareils photos numériques.

Dans le chapitre suivant, nous nous tournons vers le problème de reconstruction à partir d'échantillons localisés de manière non uniforme, et non plus sur un treillis uniforme comme dans ce chapitre et le précédent.

Chapitre 4

Reconstruction à partir d'échantillons non uniformes

4.1 Introduction

DANS les deux chapitres précédents, les signaux que nous avons traités étaient uniformes, c'est-à-dire localisés sur des treillis. Cependant, on est souvent amené à manipuler des données définies non uniformément : on parle de signal non uniforme lorsque les échantillons qui le constituent sont des mesures localisées de manière quelconque, et non sur un domaine régulier. Par exemple, on rencontre des signaux non uniformes dans les cas suivants :

- En géophysique, des données comme le potentiel magnétique ne peuvent être collectées qu'à certaines positions sur le sol ou en sous-sol, auxquelles l'observateur a accès. Il en va de même pour les données météorologiques.
- En observation astronomique, la mesure de luminosité d'une étoile ne peut être faite qu'à des moments particuliers, dépendant du cycle de rotation quotidien et annuel de la terre, et de la météo.
- En imagerie tomographique, les observations sont réalisées par un ensemble de capteurs généralement localisés avec une géométrie cylindrique ou hélicoïdale, autour de l'objet observé [246, 203]. En effet, il n'est pas possible de disposer des capteurs n'importe où, en particulier à l'intérieur du corps humain ! On rencontre aussi ce genre de problème en imagerie satellitaire, à cause des effets de parallaxe dus à l'angle formé entre le satellite et la surface terrestre.

On peut aussi penser à des systèmes qui enregistrent des événements dont l'amplitude dépasse un certain seuil de sensibilité, comme les battements cardiaques, qui sont généralement non uniformes. D'autre part, un ensemble de mesures peut comporter des échantillons manquants (perdus lors de la transmission sur un canal imparfait) ou non valides (corrompus par un bruit important, ou suite à une défaillance d'un élément du capteur), ce qui rend le flux total de mesures non uniforme.

Dans les deux chapitres précédents, les signaux étaient uniformes, et nous avons effectué la reconstruction sur le même treillis que les mesures, c'est-à-dire cherché une fonction reconstruite $f_T(\mathbf{x})$ s'écrivant comme une combinaison linéaire des fonctions $\varphi(\mathbf{x}/T - \mathbf{R}\mathbf{k})$ localisées aux sites $T\mathbf{R}\mathbf{k}$ du treillis. Dans le présent chapitre, on généralise le problème de reconstruction à des données localisées à des positions arbitraires. Que l'échantillonnage non uniforme soit voulu ou indésirable, il rend en tout état de cause le traitement des données, en particulier la reconstruction, plus difficile, car les outils d'analyse fréquentielle, pourtant si utiles, ne sont plus utilisables.

Bien que les échantillons puissent être localisés de manière quelconque, on impose à la fonction reconstruite d'être uniforme, dans un espace LSI fixé $V_T(\varphi)$. Par exemple, cet espace peut modéliser l'ensemble des fonctions pouvant être générées par un dispositif de synthèse analogique de résolution $1/T$, de réponse impulsionnelle $\varphi(t/T)$, comme un écran d'ordinateur. Les données peuvent aussi être localisées sur un treillis uniforme, mais dissocié de celui de reconstruction. Par exemple, étant donnés des échantillons 1D $v[k] = s(T_0 k)$, $k \in \mathbb{Z}$, où T_0 est le pas d'échantillonnage caractéristique de l'acquisition, on peut vouloir reconstruire une fonction $f_T \in V_T(\varphi)$ ayant un pas de résolution T différent de T_0 , puisque les systèmes d'acquisition et de synthèse n'ont pas de raison particulière d'avoir la même résolution. En 2D, le treillis peut être différent, par exemple si l'on souhaite visualiser des données acquises sur un treillis hexagonal, sur un écran d'ordinateur dont les pixels forment une grille carrée. Du point de vue du treillis de reconstruction, les données sont alors à des positions non uniformes, car le treillis d'acquisition n'a pas de sites communs avec celui de reconstruction, hormis l'origine.

On s'intéresse donc au problème de reconstruction à partir de données non uniformes, sous contrainte que la fonction reconstruite appartienne à un espace $V_T(\varphi)$ fixé indépendamment des données. Il ne s'agit donc plus de reconstruction pure, pour laquelle on cherche à synthétiser toute l'information contenue dans les données, mais plutôt du problème visant à obtenir le meilleur modèle pour les données, représentable dans l'espace $V_T(\varphi)$ des reconstructions réalisables. Cette approche est donc orientée par l'utilisation que l'on veut faire de la fonction f_T par la suite.

Dans ce chapitre, nous traitons de la reconstruction de données uni-dimensionnelles. L'extension au cadre multi-dimensionnel est possible, bien que des problèmes d'implémentation se posent, auxquels on n'est pas confronté en 1D. Nous renvoyons à [21] pour le traitement du problème en 2D et 3D. Nous nous limitons aussi au cas d'échantillons ponctuels, c'est-à-dire $\tilde{\varphi} = \delta$ dans notre modèle d'acquisition (1.10). Ce choix, qui permet de simplifier les notations, n'est pas restrictif : il n'y a aucun obstacle particulier à l'extension de l'étude dans le cas général $\tilde{\varphi} \neq \delta$.

Le problème que nous abordons est le suivant : on dispose d'échantillons ponctuels $v[n]$ d'une fonction inconnue $s(t)$, éventuellement bruités, localisés en des positions arbitraires

t_n , c'est-à-dire

$$\boxed{v[n] = s(t_n) + \mu[n]}. \quad (4.1)$$

Les échantillons $\mu[n]$ sont des réalisations d'un bruit stationnaire indépendant de s . Le but est de reconstruire une fonction $f_T(t)$ proche de $s(t)$, et appartenant à l'espace de reconstruction $V_T(\varphi)$.

Dans ce chapitre, nous présentons d'abord brièvement le contexte de la reconstruction à partir de données non uniformes, en insistant plus particulièrement sur l'approche variationnelle. Nous présentons ensuite notre approche, qui consiste à minimiser dans $V_T(\varphi)$ un critère des moindres carrés régularisé. Nous détaillons les aspects pratiques liés à la mise en œuvre de cette solution, et présentons quelques applications. Nous discutons ensuite de l'extension au cas non uniforme des méthodes par quasi-projections vues au chapitre 2.

4.2 Approches traditionnelles

De nombreuses méthodes ont été proposées dans la littérature pour la reconstruction de signaux à partir d'échantillons non uniformes (voir [6] pour un panorama). Depuis le travail pionnier de Shannon [199], il est connu qu'une fonction à bande limitée $s(t) \in L_2(\mathbb{R})$ (telle que $\hat{s}(\omega) = 0$ en dehors de $]-\pi, \pi[$) peut être reconstruite parfaitement à partir de ses échantillons aux entiers $s(n), n \in \mathbb{Z}$. Cela reste vrai si l'on dispose d'échantillons $s(t_n)$ à des positions arbitraires t_n , si certaines hypothèses fortes sur l'ensemble $\{t_n \mid n \in \mathbb{Z}\}$ sont vérifiées [110]. Cette possibilité de reconstruction parfaite d'une fonction à partir d'échantillons non uniformes est aussi assurée si la fonction appartient à un certain espace LSI, non nécessairement celui des fonctions à bande limitée [62], là aussi avec des hypothèses fortes sur les positions t_n [5]. Des algorithmes itératifs ont été proposés afin de mettre en œuvre cette reconstruction parfaite, dans le cas où s est à bande limitée [99, 109, 100, 206, 207], et dans le cas où s appartient à un espace LSI [144, 6]. Des conditions encore plus fortes sur les positions t_n doivent être satisfaites pour que ces algorithmes convergent [61].

Une grande partie de la littérature traitant de la reconstruction à partir de données non uniformes est donc consacrée à l'étude des conditions théoriques et pratiques sous lesquelles la reconstruction parfaite d'un signal $s(t)$ est possible, à partir d'échantillons discrets de s [110, 5]. Bien sûr, il est nécessaire pour cela que les échantillons ne soient pas bruités, et que l'ensemble $\{t_n\}$ soit suffisamment dense pour que toute l'information de s soit contenue dans les valeurs $s(t_n)$. En pratique, les applications sont nombreuses pour lesquelles les positions t_n sont arbitraires, et où le signal inconnu $s(t)$ n'a pas de raison particulière d'appartenir à un certain espace LSI. Il faut alors renoncer à la reconstruction parfaite, et s'intéresser plutôt à modéliser les données avec une fonction $f(t)$ satisfaisante.

Une approche classique consiste à modéliser les données par un modèle paramétrique, au sens des moindres carrés [6]. Dans nos notations, on cherche la fonction f_T appartenant à un certain espace $V_T(\varphi)$ minimisant le critère $\sum_n |f_T(t_n) - v[n]|^2$. Là encore, des hypothèses fortes doivent être faites pour que l'ensemble $\{t_n\}$ soit suffisamment dense. La

seconde approche, qui est couramment adoptée lorsque les échantillons sont localisés de manière arbitraire, est l'approche variationnelle. $f(t)$ est alors définie comme la solution d'un problème d'optimisation : $f(t)$ minimise un critère garantissant à la fois une reconstruction lisse et proche des données [237]. Nous détaillons cette approche dans la suite de cette section. Pour un panorama des autres méthodes de reconstruction à partir de données non uniformes, nous renvoyons à [240, 21].

L'approche variationnelle, dans le cas d'échantillons non bruités, vise à reconstruire la fonction $f(t)$ qui soit consistante avec les données, et qui minimise une fonctionnelle de régularisation $J(f)$ [46, 42] :

$$f = \operatorname{argmin}_{g \in L_2} J(g) \quad \text{sous contrainte que } g(t_n) = v[n] \quad \forall n. \quad (4.2)$$

Cette définition, où $J(\cdot)$ est une fonctionnelle convexe, rend le problème bien posé et assure que les régions vides d'échantillons sont « remplies » de manière lisse. Les fonctionnelles les plus utilisées proviennent de la famille des semi-normes de Duchon [91]. Dans le cas 1D, cela correspond à l'énergie de la dérivée d'ordre entier r . On pose donc, pour $r \geq 1$,

$$J(g) = \int_{\mathbb{R}} |g^{(r)}(t)|^2 dt. \quad (4.3)$$

Notons que pour que $J(g)$ soit bien défini, on impose que $g(t)$ appartienne à l'espace de Sobolev d'ordre r , W_2^r .

Dans le cas général où les mesures sont bruitées, il faut relâcher la contrainte de consistance, et la moduler avec la minimisation du critère variationnel [91, 237, 98]. On cherche donc la solution au problème

$$f = \operatorname{argmin}_{g \in W_2^r} \left(\sum_n |g(t_n) - v[n]|^2 + \lambda \int_{\mathbb{R}} |g^{(r)}(t)|^2 dt \right), \quad (4.4)$$

où $\lambda > 0$ est un paramètre lagrangien de régularisation contrôlant le compromis entre l'attache aux données et la régularisation, qui sont deux critères antagonistes. Il faut en effet trouver une fonction qui soit en adéquation avec les données $v[n]$, mais qui soit dans le même temps suffisamment lisse. On retrouve donc le critère des moindres carrés régularisé (solution de Tikhonov) présenté pour le cas uniforme dans la section 2.2.1. Lorsque λ tend vers 0, la solution f tend vers la solution du problème de minimisation sous contrainte (4.2). Dans la suite, on dira que f est solution de (4.4) avec $\lambda = 0^+$, pour indiquer que f est solution de (4.2).

L'entier r contrôle la lissitude de la reconstruction. Les valeurs $r = 1$ et $r = 2$ sont les plus fréquemment utilisées ; elles correspondent à chercher respectivement une fonction la plus plate possible, ou ayant une courbure minimale.

La solution au problème (4.4) peut être explicitée comme suit :

$$f(t) = \sum_n c[n] \rho(|t - t_n|) + \sum_{i=0}^{r-1} q[i] t^i. \quad (4.5)$$

Elle est constituée de deux termes. Le premier est une combinaison linéaire de *fonctions de base radiales* (*radial basis functions*) symétriques et localisées aux positions t_n . Dans le cas de la fonctionnelle choisie, on a $\rho(t) = |t|^{2r-1}$. Le second terme est un polynôme de degré au plus $r - 1$, qui est dans le noyau de $J(\cdot)$. Les coefficients $c[n]$ et $q[i]$ sont la solution du système linéaire

$$\begin{bmatrix} \mathbf{A} + \lambda \mathbf{I} & \mathbf{B} \\ \mathbf{B}^T & \mathbf{0} \end{bmatrix} \begin{bmatrix} \mathbf{c} \\ \mathbf{q} \end{bmatrix} = \begin{bmatrix} \mathbf{v} \\ \mathbf{0} \end{bmatrix}, \quad (4.6)$$

où $\mathbf{c} = [\cdots c[i] \cdots]^T$, $\mathbf{v} = [\cdots v[i] \cdots]^T$, $\mathbf{q} = [q[0] \cdots q[r-1]]^T$, $\mathbf{A}[n, m] = \rho(|t_n - t_m|)$ et $\mathbf{B}[n, m] = (t_n)^m$. Notons que la solution dans le cas $\lambda = 0^+$ s'obtient en résolvant ce système pour $\lambda = 0$ exactement.

La solution $f(t)$ est une spline non uniforme, avec ses nœuds aux t_n [84]. Par exemple, si $r = 1$, et $\lambda = 0^+$, $f(t)$ est simplement la fonction interpolant les valeurs $v[n]$, linéaire sur chaque segment $[t_n, t_{n+1}]$ (en supposant les t_n ordonnés par ordre croissant).

Le système linéaire à résoudre est très mal conditionné, car les fonctions radiales croissent vers l'infini en s'éloignant de l'origine. De plus, le système n'est pas creux, et la résolution est donc coûteuse. Cependant, il existe une autre base que celle des fonctions radiales pour l'espace de reconstruction $\{\sum_n c[n] |t - t_n|^{2r-1} + \sum_{i=0}^{r-1} q[i] t^i \mid c[n], q[i] \in \mathbb{R}\}$. Il s'agit de la base des B-splines non uniformes [192, 84, 98]. Ces fonctions sont à support compact, mais elles ne s'expriment plus simplement comme les translatés d'un unique prototype. Il faut donc exprimer une nouvelle B-spline en chaque point t_n . Cela rend aussi coûteuse l'évaluation de $f(t)$ en une valeur t quelconque, car f n'a pas l'expression simple d'une fonction appartenant à un espace LSI. Notons que les B-splines non uniformes deviennent les B-splines classiques $\beta^{2r-1}(x/T - n)$ dans le cas uniforme $t_n = Tn$.

L'avantage principal de l'approche variationnelle que nous venons d'aborder, est de ne pas imposer de conditions sur les positions t_n . De plus, cette approche est robuste au bruit sur les données, grâce au paramètre λ qui peut être ajusté. Les méthodes variationnelles sont donc très utilisées, en particulier en statistiques [237]. Notons que la solution $f(t)$ est invariante par dilatation et translation : si les positions t_n sont dilatées ou translatées — remplacées par αt_n et $t_n + \tau$, respectivement —, la solution correspondant au nouveau problème est la dilatée ou translatée de f correspondante — c.-à-d. $f(t/\alpha)$ et $f(t - \tau)$, respectivement.

4.3 Reconstruction dans un espace LSI

4.3.1 Position du problème

Nous proposons maintenant une nouvelle approche pour la reconstruction à partir d'échantillons non uniformes. Comme pour l'approche variationnelle, on ne fait aucune hypothèse sur les positions t_n . On minimise le même critère que dans (4.4), avec la différence majeure que l'on contraint la solution à appartenir à un certain espace LSI, comme dans les chapitres précédents.

On traite de la reconstruction d'un nombre fini N d'échantillons $v[n]$ localisés aux positions t_n , pour $n \in \llbracket 1, N \rrbracket$. On effectue la reconstruction sur un intervalle fini $\mathcal{I}$ contenant les échantillons, et pouvant être arbitrairement grand. En effet, dans les applications pratiques, le nombre d'échantillons est toujours fini. Cela permet en outre de s'intéresser au problème du traitement aux bords.

On cherche une fonction appartenant à l'espace de reconstruction $V_T(\varphi)$, s'écrivant donc sous la forme

$$f_T(t) = \sum_k c[k] \varphi\left(\frac{t}{T} - k\right), \quad (4.7)$$

où les coefficients $c[k]$ sont à déterminer en fonction des données $(t_n, v[n])$, afin que f_T soit solution du problème :

$$f_T = \operatorname{argmin}_{g \in V_T(\varphi)} \left(\sum_{n=1}^N |g(t_n) - v[n]|^2 + \lambda \int_{\mathcal{I}} |g^{(r)}(t)|^2 dt \right), \quad (4.8)$$

pour un certain $\lambda > 0$. Notons que l'intégrale porte sur $\mathcal{I}$ et non plus sur $\mathbb{R}$.

Dans notre approche, φ est arbitraire, comme dans les chapitres précédents, et n'est pas induite par le choix de la fonctionnelle $J(\cdot)$ comme avec l'approche variationnelle, dans laquelle la solution est une spline non uniforme de degré impair. Cependant, si l'on a le choix de φ , il semble judicieux de choisir comme espace de reconstruction l'espace spline de degré $2r - 1$: $\varphi = \beta^{2r-1}$. Ainsi notre solution coïncide exactement avec la solution variationnelle dans le cas où les positions t_n sont uniformes : $t_n = Tn$.

Avant d'étudier la solution de notre problème, mentionnons une approche alternative proposée dans [162] : f_T est définie comme la projection orthogonale dans $V_T(\varphi)$ de la spline non uniforme $f(t)$ solution du problème variationnel (4.4). Cette approche est plus compliquée que la nôtre, puisqu'il faut d'abord reconstruire $f(t)$ dans une base de B-splines non uniformes, puis effectuer l'étape de projection, qui n'est pas triviale. Notre formulation a l'avantage d'exprimer directement la solution $f_T(t)$ à partir des données discrètes $v[n]$, sans passer par une fonction intermédiaire.

Les coefficients $c[k]$ dans l'équation (4.7), qui déterminent la solution $f_T(t)$ de notre problème, vont être calculés à partir des données $v[n]$ par un algorithme direct (non itératif)

et rapide, réalisant l'équivalent d'un filtrage inverse non stationnaire. Pour qu'une telle implémentation soit possible, il nous faut faire une hypothèse importante : φ doit avoir un **support compact**. De manière non-restrictive, on supposera ce support inclus dans l'intervalle $] -W, W[$, pour un certain entier W .

Détaillons les autres hypothèses requises. On suppose que φ est bornée, et que les fonctions $\varphi(\frac{t}{T} - k)$ forment une base de Riesz. De plus, on impose que φ appartienne à l'espace de Sobolev W_2^r ; par exemple, si l'on choisit une B-spline de degré d , $\varphi = \beta^d$, il est nécessaire d'avoir $d \geq r$. On cherche à effectuer la reconstruction sur l'intervalle $\mathcal{I} = [0, S]$, avec $S = KT$ pour un certain $K \in \mathbb{N}$. On requiert que $t_n \in \mathcal{I}$ pour tout $n \in \llbracket 1, N \rrbracket$, et qu'il y ait au moins r positions t_n distinctes. La fonction reconstruite f_T est entièrement décrite sur $\mathcal{I}$ par un nombre fini de coefficients $c[k]$, pour $k \in \llbracket -W + 1, K + W - 1 \rrbracket$. Dans la suite, on notera $k \in \mathbb{Z}$, mais on supposera qu'on ne manipule que des fonctions de $V_T(\varphi)$ telles que $c[k] = 0$ pour tout $k \notin \llbracket -W + 1, K + W - 1 \rrbracket$.

4.3.2 Expression de la solution

Déterminer la fonction $f_T(t)$ solution de (4.8) revient à calculer la séquence $c = (c[k])_{k \in \mathbb{Z}}$ de (4.7) afin que la fonction coût

$$\varepsilon(c) = \sum_{n=1}^N |f_T(t_n) - v[n]|^2 + \lambda \int_{\mathcal{I}} |f_T^{(r)}(t)|^2 dt \quad (4.9)$$

soit minimale. Premièrement, réécrivons le terme d'attache aux données en fonction de c :

$$\sum_{n=1}^N |f_T(t_n) - v[n]|^2 = \sum_{n=1}^N \left(v[n] - \sum_{k \in \mathbb{Z}} c[k] \varphi\left(\frac{t_n}{T} - k\right) \right)^2. \quad (4.10)$$

Deuxièmement, réécrivons le terme variationnel en fonction de c , mais avec l'intégrale portant sur $\mathbb{R}$ et non sur $\mathcal{I}$ dans un premier temps. Introduisons la séquence discrète $q_{\varphi,r}$ définie par

$$q_{\varphi,r}[k] = \frac{(-1)^r}{T} a_{\varphi^{(r)}}[k]. \quad (4.11)$$

Puisque $\varphi^{(r)} = (-1)^r \bar{\varphi}^{(r)}$, on a :

$$\int_{\mathbb{R}} |f_T^{(r)}(t)|^2 dt = \frac{1}{T^2} \int_{\mathbb{R}} \left(\sum_{k \in \mathbb{Z}} c[k] \varphi^{(r)}\left(\frac{t}{T} - k\right) \right)^2 dt \quad (4.12)$$

$$= \frac{1}{T} \sum_{k, l \in \mathbb{Z}} c[k] c[l] \int_{\mathbb{R}} \varphi^{(r)}(x - (k - l)) \varphi^{(r)}(x) dx \quad (4.13)$$

$$= \frac{(-1)^r}{T} \sum_{k, l \in \mathbb{Z}} c[k] c[l] (\bar{\varphi}^{(r)} * \varphi^{(r)})(k - l) \quad (4.14)$$

$$= \frac{(-1)^r}{T} \sum_{k, l \in \mathbb{Z}} c[k] c[l] a_{\varphi^{(r)}}[k - l] \quad (4.15)$$

$$= \sum_{k \in \mathbb{Z}} c[k] (c * q_{\varphi, r})[k]. \quad (4.16)$$

Si l'on effectue une reconstruction spline ($\varphi = \beta^d$), on peut expliciter la forme générale du filtre $q_{\varphi, r}$. En effet, les B-splines vérifient la relation simple $\bar{\beta}^d * \beta^d = \beta^{2d+1}$. De plus, la dérivée d'une spline est aussi une spline de degré inférieur : $\beta^{d(1)}(t) = \beta^{d-1}(t + \frac{1}{2}) - \beta^{d-1}(t - \frac{1}{2})$. Pour tout $d \in \mathbb{N}$, on définit b^d comme la B-spline discrète centrée de degré d : $b^d[k] = \beta^d(k)$ pour tout $k \in \mathbb{Z}$. Rappelons l'expression des premières B-splines discrètes dans le domaine en z : $B^0(z) = 1$, $B^1(z) = 1$, $B^2(z) = \frac{1}{8}z + \frac{3}{4} + \frac{1}{8}z^{-1}$, $B^3(z) = \frac{1}{6}z + \frac{2}{3} + \frac{1}{6}z^{-1}$. Ainsi, le filtre $q_{\varphi, r}$ a la forme suivante :

$$Q_{\beta^d, r}(z) = \frac{1}{T} (-z + 2 - z^{-1})^r B^{2d+1-2r}(z). \quad (4.17)$$

Par exemple, si $\varphi = \beta^1$ et $r = 1$, alors $Q_{\varphi, r}(z) = \frac{1}{T} (-z + 2 - z^{-1})$. Si $\varphi = \beta^3$ et $r = 2$, alors $Q_{\varphi, r}(z) = \frac{1}{6T} (z^3 - 9z + 16 - 9z^{-1} + z^{-3})$.

Afin d'exprimer le coût $\varepsilon(c)$ en termes de matrices et vecteurs, on introduit les quantités matricielles suivantes :

$$\begin{aligned} \mathbf{c} &= [c[-W + 1] \ \cdots \ c[K + W - 1]]^T, \\ \mathbf{v} &= [v[1] \ \cdots \ v[N]]^T, \\ \mathbf{M} &= [M[n, k]]_{n, k} \text{ avec } M[n, k] = \varphi\left(\frac{t_n}{T} - k\right) \ \forall n \in [1, N], k \in [-W + 1, K + W - 1], \\ \mathbf{Q} &= [Q[k, l]]_{k, l} \ \forall k, l \in [-W + 1, K + W - 1]. \end{aligned} \quad (4.18)$$

En utilisant ces matrices, le coût $\varepsilon(c)$ peut être réécrit :

$$\varepsilon(\mathbf{c}) = \|\mathbf{M}\mathbf{c} - \mathbf{v}\|^2 + \lambda \mathbf{c}^T \mathbf{Q} \mathbf{c}, \quad (4.19)$$

où les valeurs de la matrice $\mathbf{Q}$ sont données par :

$$Q[k, l] = q_{\varphi, r}[k - l] = \frac{(-1)^r}{T} a_{\varphi^{(r)}}[k - l], \quad (4.20)$$

sauf pour les premières et dernières lignes de $\mathbf{Q}$ qui contiennent des valeurs particulières, car $\varepsilon(c)$ est défini avec une intégrale portant sur $\mathcal{I}$, et non sur $\mathbb{R}$ comme dans (4.16). Ces valeurs particulières sont contenues dans des domaines carrés de taille $(2W - 1)^2$, dans les coins en haut à gauche et en bas à droite de la matrice. Afin de calculer ces valeurs pour le bord gauche (la méthode serait similaire pour le bord droit), on doit développer le membre gauche de l'égalité suivante, et l'identifier avec le membre droit :

$$\int_0^{2W-1} \left| \sum_{k=-W+1}^{W-1} c[k] \varphi^{(r)}\left(\frac{t}{T} - k\right) \right|^2 dt = \sum_{k,l=-W+1}^{W-1} Q[k,l] c[k] c[l]. \quad (4.21)$$

Par exemple, dans les deux cas $\varphi = \beta^1, r = 1$ et $\varphi = \beta^3, r = 2$, on a respectivement :

$$\mathbf{Q} = \frac{1}{T} \begin{bmatrix} \mathbf{1} & -\mathbf{1} & 0 & 0 & 0 & \cdots \\ -\mathbf{1} & \mathbf{2} & -\mathbf{1} & 0 & 0 & \ddots \\ 0 & -\mathbf{1} & \mathbf{2} & -\mathbf{1} & 0 & \ddots \\ 0 & 0 & -\mathbf{1} & \mathbf{2} & -\mathbf{1} & \ddots \\ 0 & 0 & 0 & -\mathbf{1} & \mathbf{2} & \ddots \\ \vdots & \ddots & \ddots & \ddots & \ddots & \ddots \end{bmatrix}, \quad \mathbf{Q} = \frac{1}{6T} \begin{bmatrix} \mathbf{2} & -\mathbf{3} & \mathbf{0} & \mathbf{1} & 0 & \cdots \\ -\mathbf{3} & \mathbf{8} & -\mathbf{6} & 0 & \mathbf{1} & \ddots \\ \mathbf{0} & -\mathbf{6} & \mathbf{14} & -\mathbf{9} & 0 & \ddots \\ \mathbf{1} & 0 & -\mathbf{9} & \mathbf{16} & -\mathbf{9} & \ddots \\ 0 & \mathbf{1} & 0 & -\mathbf{9} & \mathbf{16} & \ddots \\ \vdots & \ddots & \ddots & \ddots & \ddots & \ddots \end{bmatrix}. \quad (4.22)$$

Maintenant, définissons $\mathbf{A} = \mathbf{M}^T \mathbf{M} + \lambda \mathbf{Q}$ et $\mathbf{y} = \mathbf{M}^T \mathbf{v}$. Minimiser le coût $\varepsilon(\mathbf{c})$ revient à résoudre le système linéaire

$$\mathbf{A} \mathbf{c} = \mathbf{y}. \quad (4.23)$$

En considérant les lignes de ce système, on obtient un ensemble d'équations, que l'on peut retrouver directement à partir de (4.10) et (4.16), en mettant à zéro les dérivées partielles $\partial \varepsilon / \partial c[k]$, pour chaque $k \in \llbracket -W + 1, K + W - 1 \rrbracket$:

$$\sum_{l \in \mathbb{Z}} \left[\sum_{n=1}^N \varphi\left(\frac{t_n}{T} - k\right) \varphi\left(\frac{t_n}{T} - l\right) + \lambda Q[k,l] \right] c[l] = \sum_{n=1}^N \varphi\left(\frac{t_n}{T} - k\right) v[n]. \quad (4.24)$$

On doit encore vérifier que la solution du système linéaire (4.23) est bien définie. Tout d'abord, $\mathbf{M}^T \mathbf{M}$ est positive, c.-à-d. $\mathbf{x}^T \mathbf{M}^T \mathbf{M} \mathbf{x} \geq 0$ pour tout vecteur $\mathbf{x}$, car $\mathbf{x}^T \mathbf{M}^T \mathbf{M} \mathbf{x} = \|\mathbf{M} \mathbf{x}\|_{\ell_2}^2 \geq 0$. $\mathbf{Q}$ est aussi positive : si l'on choisit un vecteur $\mathbf{x}$ et que l'on définit

$$g(t) = \sum_{k=-W+1}^{K+W-1} x[k] \varphi\left(\frac{t}{T} - k\right), \quad (4.25)$$

alors $\mathbf{x}^T \mathbf{Q} \mathbf{x} = \int_{\mathcal{I}} |g^{(r)}(t)|^2 dt$ par construction de $\mathbf{Q}$, et cette intégrale est positive. Prouvons maintenant que $\mathbf{A}$ est définie positive, c.-à-d. $(\mathbf{x}^T \mathbf{A} \mathbf{x} = 0 \Rightarrow \mathbf{x} = 0)$. On suppose que $\mathbf{x}^T \mathbf{A} \mathbf{x} = 0$. Alors $\mathbf{x}^T \mathbf{M}^T \mathbf{M} \mathbf{x} = 0$ et $\mathbf{x}^T \mathbf{Q} \mathbf{x} = 0$. Cela nous donne $g(t_n) = 0$ pour tout $n \in \llbracket 1, N \rrbracket$ et $\int_{\mathcal{I}} |g^{(r)}(t)|^2 dt = 0$. Par conséquent, $g^{(r)}(t) = 0$ dans $\mathcal{I}$, et $g(t)$ est un polynôme

de degré strictement inférieur à r dans cet intervalle. Ce polynôme a N racines aux t_n (dont au moins r distinctes par hypothèse). Donc $g = 0$ et $x[k] = 0$ pour tout k . Par conséquent, le système linéaire a une solution unique.

De plus, grâce à l'hypothèse sur le support compact de φ , $\mathbf{A}$ est diagonale par bandes avec seulement $4W - 1$ diagonales contenant des valeurs non nulles ; plus précisément, $A[k, l] = 0$ si $|k - l| \geq 2W$. Ce sera l'élément clé pour résoudre efficacement ce système linéaire, comme détaillé dans la section 4.3.4. Notons qu'effectuer un produit matrice/vecteur comme $\mathbf{A}\mathbf{c}$ est équivalent à appliquer un filtre non stationnaire au signal c . Résoudre un système linéaire revient donc à effectuer un filtrage inverse non stationnaire sur le signal c .

4.3.3 Reconstruction avec des conditions de bord symétriques

La solution de notre problème de reconstruction, sur un intervalle fini de taille K , est paramétrée par $K + 2W - 1$ coefficients. Dans cette section, nous envisageons une solution alternative, qui permet d'éviter l'excès de coefficients décrivant la solution aux bords de l'intervalle. Pour cela, nous adoptons des conditions de bords symétriques, c'est-à-dire que nous cherchons une fonction f_T symétrique par rapport à $t = 0$ et $t = TK$, et donc $2K$ -périodique [45]. Le principal avantage d'une telle approche est que la solution est alors entièrement caractérisée par les coefficients $c[k]$ pour k allant de 0 à K . Dans cette section seulement, nous faisons l'hypothèse que φ est symétrique ($\varphi = \bar{\varphi}$).

Nous résolvons maintenant le problème suivant :

$$f_T = \underset{g \in V_T(\varphi) \cap \mathcal{S}}{\operatorname{argmin}} \left(\sum_{n=0}^{N-1} |g(t_n) - v[n]|^2 + \lambda \int_{\mathcal{I}} |g^{(r)}(t)|^2 dt \right). \quad (4.26)$$

où $\mathcal{S}$ est l'ensemble des fonctions symétriques par rapport à $t = 0$ et $t = TK$. Avec ces conditions de bords, les coefficients c de f_T vérifient aussi des propriétés de symétrie et périodicité $c[-k] = c[k]$ and $c[k + 2pK] = c[k]$ pour tout $k, p \in \mathbb{Z}$. f_T est donc bien complètement déterminée par les $c[k], k \in \llbracket 0, K \rrbracket$. Une interprétation de (4.26) est de considérer que l'on veut reconstruire une fonction à partir d'un ensemble infini d'échantillons, symétrique et périodique, comme illustré sur la figure 4.1. Ainsi, la symétrie et la périodicité de f_T sont « héritées » de la structure de ces données.

Les coefficients $c[k], k \in \llbracket 0, K \rrbracket$, paramétrant la solution du problème (4.26), sont solutions d'un système linéaire $\tilde{\mathbf{A}}\tilde{\mathbf{c}} = \tilde{\mathbf{y}}$ à déterminer — les tildes sont pour distinguer les matrices de celles précédemment définies. Etudions la structure du problème au bord gauche seulement, puisque le traitement du bord droit est identique. Comme dans l'étude précédente, le problème présent revient à minimiser le critère $\varepsilon(\mathbf{c})$ de (4.19). Le vecteur $\mathbf{c}$ a la forme $\mathbf{c} = [c[W-1] \ \cdots \ c[1] \ c[0] \ c[1] \ \cdots]^T$ à cause des conditions de bords $c[-k] = c[k]$. Chaque ligne de $\mathbf{M}\mathbf{c}$ peut être développée en

$$\cdots + M[k, -1]c[-1] + M[k, 0]c[0] + M[k, 1]c[1] + \cdots, \quad (4.27)$$

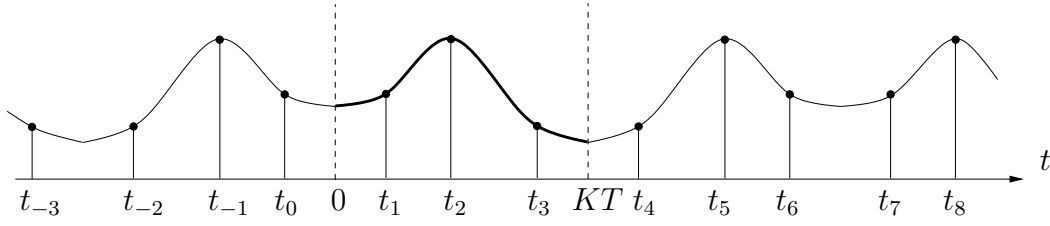


FIG. 4.1 : Un ensemble fini de $N = 3$ échantillons étendu en un ensemble infini en utilisant des conditions de bords symétriques.

que l'on peut réécrire

$$M[k, 0]c[0] + (M[k, -1] + M[k, 1])c[1] + \dots \quad (4.28)$$

On a donc $\widetilde{\mathbf{M}}\mathbf{c} = \mathbf{M}\mathbf{c}$, où $\widetilde{\mathbf{M}}$ est obtenue à partir de $\mathbf{M}$ en supprimant ses $W - 1$ premières colonnes, et en ajoutant leurs valeurs à celles des colonnes miroirs par rapport à la W -ième colonne (avec la même manipulation au bord droit). Intuitivement, cela revient à « replier » la matrice $\mathbf{M}$, comme illustré dans l'exemple suivant avec $W = 2$, où la première colonne est ajoutée à la troisième et supprimée :

$$\mathbf{M} = \begin{bmatrix} a & b & c & d & \dots \\ e & f & g & h & \dots \\ \vdots & \vdots & \vdots & \vdots & \ddots \end{bmatrix} \longrightarrow \widetilde{\mathbf{M}} = \begin{bmatrix} b & c+a & d & \dots \\ f & g+e & h & \dots \\ \vdots & \vdots & \vdots & \ddots \end{bmatrix}. \quad (4.29)$$

Cette opération de pliage doit être effectuée à la fois sur les premières et dernières lignes et colonnes des matrices $\mathbf{Q}$ et $\mathbf{A}$, afin d'obtenir $\widetilde{\mathbf{Q}}$ et $\widetilde{\mathbf{A}}$. De même, les premières et dernières valeurs de $\mathbf{y}$ doivent être repliées. Par exemple, si $\varphi = \beta^3$ et $r = 2$ (donc $W = 2$), cela donne :

$$\mathbf{Q} = \frac{1}{6T} \begin{bmatrix} \mathbf{2} & \mathbf{-3} & \mathbf{0} & \mathbf{1} & \mathbf{0} & \dots \\ \mathbf{-3} & \mathbf{8} & \mathbf{-6} & \mathbf{0} & \mathbf{1} & \ddots \\ \mathbf{0} & \mathbf{-6} & \mathbf{14} & \mathbf{-9} & \mathbf{0} & \ddots \\ \mathbf{1} & \mathbf{0} & \mathbf{-9} & \mathbf{16} & \mathbf{-9} & \ddots \\ \mathbf{0} & \mathbf{1} & \mathbf{0} & \mathbf{-9} & \mathbf{16} & \ddots \\ \vdots & \ddots & \ddots & \ddots & \ddots & \ddots \end{bmatrix} \longrightarrow \widetilde{\mathbf{Q}} = \frac{1}{6T} \begin{bmatrix} \mathbf{8} & \mathbf{-9} & \mathbf{0} & \mathbf{1} & \ddots \\ \mathbf{-9} & \mathbf{16} & \mathbf{-8} & \mathbf{0} & \ddots \\ \mathbf{0} & \mathbf{-8} & \mathbf{16} & \mathbf{-9} & \ddots \\ \mathbf{1} & \mathbf{0} & \mathbf{-9} & \mathbf{16} & \ddots \\ \ddots & \ddots & \ddots & \ddots & \ddots \end{bmatrix}. \quad (4.30)$$

4.3.4 Mise en œuvre de la méthode

Nous avons vu que le problème posé se ramène à la résolution du système linéaire $\mathbf{A}\mathbf{c} = \mathbf{y}$, symétrique, défini positif, et diagonal par bandes. La meilleure stratégie pour résoudre un tel système repose sur la factorisation de Cholesky de la matrice $\mathbf{A}$, et consiste en trois étapes :

1. La décomposition de Cholesky de $\mathbf{A}$ est effectuée ; cette opération consiste à calculer l'unique matrice triangulaire inférieure $\mathbf{L}$ telle que $\mathbf{A} = \mathbf{L}\mathbf{L}^T$ [178, 107]. $\mathbf{L}$ est aussi

diagonale par bandes, avec $2W$ diagonales contenant des éléments non nuls : $L[k, l] \neq 0$ seulement si $-2W < k - l \leq 0$. La décomposition de Cholesky peut être effectuée ligne par ligne par une version rapide de l'algorithme de Crout [178], qui tire avantage de la structure diagonale par bandes de $\mathbf{A}$.

2. Le système triangulaire inférieur $\mathbf{L}^{\circ} \mathbf{c} = \mathbf{y}$ est résolu par substitution : pour k allant de $k_{\min}$ à $k_{\max}$,

$$\mathring{c}[k] = \frac{1}{L[k, k]} \left(y[k] - \sum_{i=-2W+1}^{-1} L[k, k+i] \mathring{c}[k+i] \right). \quad (4.31)$$

3. Le système triangulaire supérieur $\mathbf{L}^T \mathbf{c} = \mathring{\mathbf{c}}$ est finalement résolu par substitution en sens inverse : pour k allant de $k_{\max}$ à $k_{\min}$,

$$c[k] = \frac{1}{L[k, k]} \left(\mathring{c}[k] - \sum_{i=1}^{2W-1} L[k+i, k] c[k+i] \right). \quad (4.32)$$

Il est possible d'effectuer ces opérations l'une après l'autre, en utilisant des algorithmes standards d'algèbre linéaire, que l'on trouve dans les bibliothèques comme LAPACK implémentant de nombreuses routines génériques de factorisation matricielle et de résolution de systèmes linéaires. De telles bibliothèques sont hautement optimisées pour tirer au mieux parti de la plateforme sur laquelle le programme est exécuté, mais elles ne contiennent que des algorithmes généraux, non dédiés à la résolution d'un problème particulier. Au contraire, l'algorithme que nous proposons tire pleinement parti de la structure spécifique du système à résoudre. Expliciter notre algorithme nous semble un bon moyen d'en comprendre les mécanismes, et le développeur désirant l'implémenter aura loisir de l'optimiser à plus ou moins haut niveau.

On s'intéresse donc à l'implémentation pratique de l'algorithme qui calcule les coefficients $c[k]$, $k \in \llbracket -W + 1, K + W - 1 \rrbracket$. On traite de l'approche sans conditions de bords, mais l'approche avec conditions de bords symétriques est très similaire, bien qu'un peu plus longue à écrire. L'algorithme consiste en deux passes : la première implémente la décomposition de Cholesky et résout le premier système linéaire (4.31), tout cela en un seul balayage des données. La seconde passe résout le second système linéaire (4.32), lors d'un parcours en sens inverse des coefficients précédemment calculés.

Définissons les variables auxiliaires $a[i] = A[k, k+i]$ et $u[k, i] = L[k+i, k]$. Alors notre algorithme peut s'écrire en pseudo-code comme suit :

- Première passe :
 - pour k de $-W + 1$ à $K + W - 1$ {
 - $i_{\min} := \max(-2W + 1, -W + 1 - k)$;
 - $i_{\max} := \min(2W - 1, K + W - 1 - k)$;

$$\begin{aligned}
& \text{pour } i \text{ de } 0 \text{ à } i_{\max}, \\
& \quad a[i] := \sum_{n=1}^N \varphi\left(\frac{t_n}{T} - k\right) \varphi\left(\frac{t_n}{T} - k - i\right) + \lambda Q[k, k+i]; \\
& \quad u[k, 0] := \left(a[0] - \sum_{i=i_{\min}}^{-1} u[k+i, -i]^2 \right)^{1/2}; \\
& \quad \circ c[k] := \frac{1}{u[k, 0]} \left(\sum_{n=1}^N \varphi\left(\frac{t_n}{T} - k\right) v[n] - \sum_{i=i_{\min}}^{-1} u[k+i, -i] \circ c[k+i] \right); \\
& \quad \text{pour } i \text{ de } 1 \text{ à } i_{\max}, \\
& \quad \quad u[k, i] := \frac{1}{u[k, 0]} \left(a[i] - \sum_{j=\max(i-2W+1, -W+1-k)}^{-1} u[k+j, -j] u[k+j, i-j] \right); \\
& \quad \} \\
& - \text{Seconde passe :} \\
& \quad \text{pour } k \text{ de } K+W-1 \text{ à } -W+1 \{ \\
& \quad \quad i_{\max} := \min(2W-1, K+W-1-k); \\
& \quad \quad c[k] := \frac{1}{u[k, 0]} \left(\circ c[k] - \sum_{i=1}^{i_{\max}} u[k, i] c[k+i] \right); \\
& \quad \}
\end{aligned}$$

En fait, les sommes indexées par n ne contiennent qu'un petit nombre de valeurs non nulles correspondant aux échantillons situés localement dans l'intervalle $]T(k-W), T(k+W)[$. Si les données sont ordonnées de telle manière que $t_{n+1} \geq t_n$ pour tout n , on peut accéder à l'ensemble de données $(t_n, v[n])$ de manière progressive. Ainsi la première passe peut être exécutée au fil de l'eau, au fur et à mesure que les données sont rendues disponibles. Par exemple, si t_n est la date d'acquisition de la mesure $v[n]$, cette première passe peut être réalisée en temps réel, avec un retard de W unités de temps. De plus, la seconde passe, qui parcourt les coefficients calculés en sens inverse, peut être réalisée en place : les valeurs $c[k]$ remplacent les valeurs intermédiaires $\circ c[k]$ calculées durant la première passe. Ainsi, la seconde passe peut être vue comme un post-traitement mettant à jour les coefficients. Cet algorithme réalise donc un double filtrage récursif, avec un filtre dont les coefficients varient en fonction du temps, et qui sont calculés à la volée par la factorisation de Cholesky.

Si les positions t_n ne sont pas triées, ou si N est très grand devant K , une grande partie du temps de calcul sera consommée pour l'évaluation des sommes indexées par n . Dans ce cas, ou si le traitement au fil de l'eau des données n'est pas requis, il est plus approprié d'utiliser la variante suivante de l'algorithme, composée de trois passes. Durant la première passe, la partie triangulaire supérieure de $\mathbf{A}$ est calculée et stockée provisoirement dans les coefficients $u[k, i]$. Les valeurs $y[k]$ sont elles aussi calculées et stockées dans les coefficients $\circ c[k]$. Les vraies valeurs de ces coefficients sont ensuite calculées en place durant la seconde

pas, comme précédemment.

– Première passe :

pour k de $-W + 1$ à $K + W - 1$ {
 $\circ c[k] = 0$;
 pour i de 0 à $\min(2W - 1, K + W - 1 - k)$,
 $u[k, i] := \lambda Q[k, k + i]$;
 }
 pour n de 1 à N ,
 pour k de $\left\lfloor \frac{t_n}{T} + 1 - W \right\rfloor$ à $\left\lfloor \frac{t_n}{T} - 1 + W \right\rfloor$ {
 pour i de 0 à $\min(2W - 1, K + W - 1 - k)$,
 $u[k, i] := u[k, i] + \varphi\left(\frac{t_n}{T} - k\right) \varphi\left(\frac{t_n}{T} - k - i\right)$;
 $\circ c[k] := \circ c[k] + \varphi\left(\frac{t_n}{T} - k\right) v[n]$;
 }

– Seconde passe :

pour k de $-W + 1$ à $K + W - 1$ {
 $i_{\min} := \max(-2W + 1, -W + 1 - k)$;
 $i_{\max} := \min(2W - 1, K + W - 1 - k)$;
 $u[k, 0] := \left(u[k, 0] - \sum_{i=i_{\min}}^{-1} u[k + i, -i]^2 \right)^{1/2}$;
 $\circ c[k] := \frac{1}{u[k, 0]} \left(\circ c[k] - \sum_{i=i_{\min}}^{-1} u[k + i, -i] \circ c[k + i] \right)$;
 pour i de 1 à $i_{\max}$,
 $u[k, i] := \frac{1}{u[k, 0]} \left(u[k, i] - \sum_{j=\max(i-2W+1, -W+1-k)}^{-1} u[k + j, -j] u[k + j, i - j] \right)$;
 }

– Troisième passe :

pour k de $K + W - 1$ à $-W + 1$ {
 $i_{\max} := \min(2W - 1, K + W - 1 - k)$;
 $c[k] := \frac{1}{u[k, 0]} \left(\circ c[k] - \sum_{i=1}^{i_{\max}} u[k, i] c[k + i] \right)$;
 }

L'implémentation dans le cas de l'utilisation de conditions de bords symétriques est très similaire. L'opération de « pliage » des matrices est alors implémentée en affectant

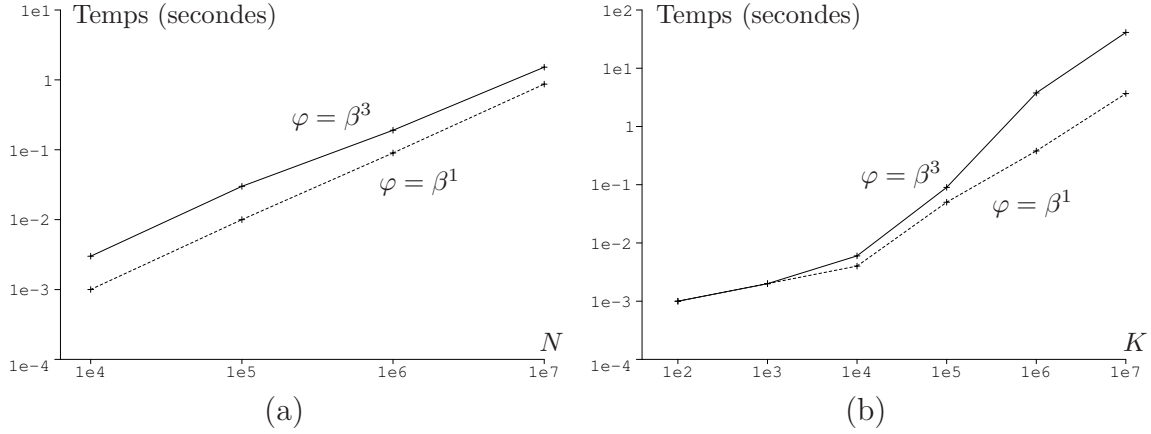


FIG. 4.2 : Temps de calcul des coefficients $c[k], k \in \llbracket 0, K \rrbracket$, en échelle log-log, pour un problème de reconstruction de N échantillons localisés à des positions choisies aléatoirement dans l'intervalle $[0, 100]$, comme sur la figure 4.6 ($T = 100/K$, $\lambda = 0.01$). (a) : K est fixé à 100 et N varie. (b) : N est fixé à 10000 et K varie. La ligne en pointillés est pour la reconstruction spline linéaire ($\varphi = \beta^1$, $r = 1$), alors que la ligne solide est pour la reconstruction spline cubique ($\varphi = \beta^3$, $r = 2$).

directement chaque contribution $\varphi(\frac{t_n}{T} - k)\varphi(\frac{t_n}{T} - k - i)$ à sa place après repliement, par exemple à $u[-k, -i]$ au lieu de $u[k, i]$ si $k < 0$.

Le temps de calcul de notre algorithme peut être modélisé comme suit : $O(W^2N)$, $O(W^2K)$, $O(WN)$ pour le calcul des éléments de $\mathbf{A}$, $\mathbf{L}$ et $\mathbf{y}$ respectivement, et $O(KW)$ pour les substitutions. Ainsi, le temps total s'écrit $O(W^2(N + K))$; il est linéaire en N et K , ce qui est le mieux que l'on puisse espérer. Si la reconstruction est effectuée sur un intervalle $\mathcal{I}$ de taille fixée $S = KT$, le temps total peut être réécrit $O(W^2(N + S/T))$ pour faire apparaître la dépendance linéaire en $1/T$. Des temps de calcul expérimentaux sont reportés sur la figure 4.2 pour une implémentation en langage C de la deuxième version de l'algorithme que nous avons donnée, exécutée sur un PC portable à 1.6GHz. Le temps de calcul est asymptotiquement linéaire en K et N comme prédit.

Hormis la mémoire requise pour stocker les coefficients $c[k]$, un emplacement mémoire auxiliaire de taille $2W(K + 2W - 1)$ unités (ou $2W(K + 1)$ si des conditions de bords symétriques sont utilisées) est nécessaire pour stocker les valeurs $u[k, i]$, qui sont générées durant la première passe et utilisées lors de la seconde.

Enfin, notre algorithme repose sur la décomposition de Cholesky, réputée « extrêmement stable numériquement » [178]. Une divergence de l'algorithme impliquerait que la matrice n'est pas définie positive, ce qui est exclu. Cependant, même si le système linéaire est défini positif, son conditionnement dépend des positions t_n . En fait, notre méthode

revient à effectuer une déconvolution avec un filtre non stationnaire. Dans une région sans échantillon, ce filtre se réduit à $q_{\varphi,r}$, qui a des racines sur le cercle unité complexe. Un filtre avec des pôles sur le cercle unité est dit *marginalelement stable*, car la réponse impulsionnelle correspondant à ces pôles ne décroît pas, mais ne croît pas non plus. Donc si une erreur de calcul, de l'ordre de la précision de la machine, survient durant l'exécution sur un coefficient $c[k]$, elle se propagera aux coefficients voisins dans une région sans échantillon, mais sans être amplifiée. Ce n'est donc pas un point préoccupant.

Notons qu'au lieu de la décomposition de Cholesky, il est possible d'utiliser une factorisation LU, qui n'exploite pas la symétrie de la matrice. Bien que cette décomposition prenne deux fois plus de temps, l'utilisation de la fonction racine carrée n'est pas requise, et une diagonale en moins est nécessaire pour le stockage de $\mathbf{L}$; c'est-à-dire, $(2W-1)(K+2W-1)$ au lieu de $2W(K+2W-1)$ unités mémoire sont utilisées. Une autre variante est la factorisation $\mathbf{LDL}^T$, qui nécessite une passe de plus de parcours des coefficients, mais évite aussi l'utilisation de la racine carrée. L'opportunité d'utiliser telle ou telle factorisation doit être étudiée au cas par cas.

La mise en œuvre rapide et non itérative que nous avons proposée dans cette section est nouvelle. Gröchenig et coll. se sont intéressés au cas particulier que constitue le problème de la reconstruction dans $V_T(\varphi)$ au sens des moindres carrés, donc sans critère variationnel, autrement dit $\lambda = 0$ dans (4.8) [111]. Il faut pour cela que l'ensemble $\{t_n\}$ soit suffisamment dense afin que le problème soit bien posé, c'est-à-dire que la solution dans $V_T(\varphi)$ soit unique, sans que la régularisation soit nécessaire. Ces auteurs aboutissent à une formulation matricielle similaire à la nôtre, mais ils proposent d'utiliser des algorithmes d'algèbre classiques, pour résoudre le système linéaire obtenu. En entrant dans le détail de cette résolution, nous avons obtenu un algorithme original effectuant un filtrage récursif non stationnaire en deux passes, l'une causale, l'autre anti-causale. Cet algorithme peut être vu comme une généralisation de l'algorithme récursif stationnaire de Unser et coll. permettant d'effectuer les filtrages inverses [221] .

On gardera cependant à l'esprit que notre mise en œuvre rapide n'est disponible que dans le cas 1D. Récemment, deux équipes indépendantes ont étudié l'équivalent 2D du problème abordé dans ce chapitre (à savoir, la reconstruction régularisée dans un espace LSI à partir de données non-uniformes) [22][233]. Les algorithmes qu'ils proposent sont itératifs, et bien que la stratégie multi-résolution adoptée dans [22] soit relativement efficace, son équivalent 1D reste bien plus lent que notre implémentation directe.

4.3.5 Influence des paramètres

Afin d'illustrer l'influence de φ et r sur la reconstruction, nous avons représenté trois configurations avec un exemple de données synthétiques, sur la figure 4.3. En (a), le choix $\varphi = \beta^1, r = 1$ produit une reconstruction linéaire par morceaux, avec ses nœuds uniformes aux $Tk, k \in \mathbb{Z}$. Notons qu'avec $\varphi = \beta^1$, le seul choix possible est $r = 1$. De plus, dans ce cas,

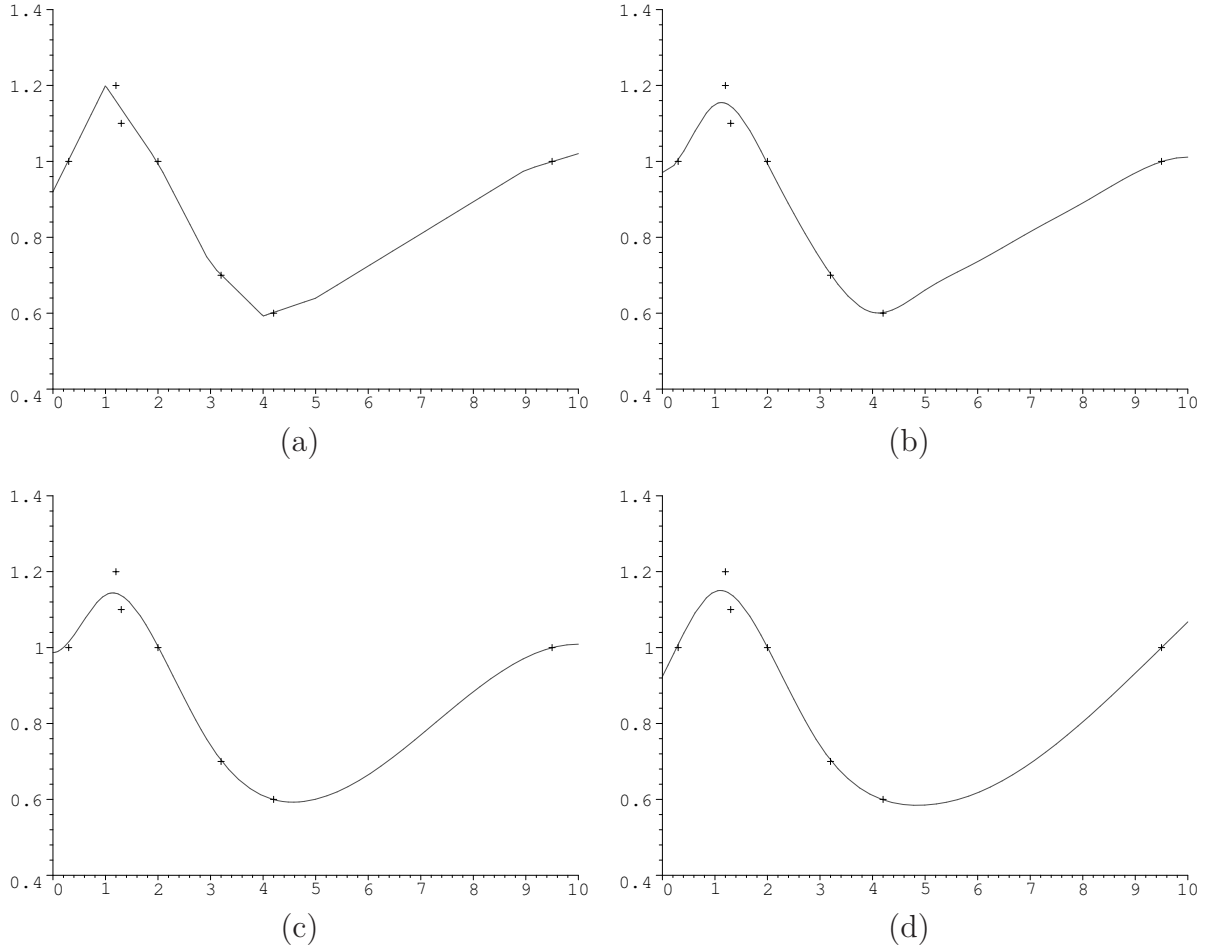


FIG. 4.3 : Splines uniformes avec nœuds aux entiers ($T = 1$, $\lambda = 0.01$) reconstruites à partir de 7 échantillons localisés dans l'intervalle $[0, 10]$, avec différents degrés et différentes valeurs du paramètre r . (a) : $\varphi = \beta^1$, $r = 1$. (b) : $\varphi = \beta^3$, $r = 1$. (c), (d) : $\varphi = \beta^3$, $r = 2$. Des conditions de bords symétriques sont utilisées en (a), (b), (c), pas en (d).

il est équivalent d'appliquer la stratégie avec et sans conditions de bords symétriques. En (b), (c), (d), $\varphi = \beta^3$ fournit une reconstruction plus lisse, de régularité $\mathcal{C}^2$. Si $r = 1$, comme en (b), la solution se comporte comme une ligne droite dans les régions vides d'échantillons, alors que si $r = 2$, comme en (c), f_T se comporte comme un polynôme de degré 3 ; dans ce deuxième cas, f_T peut s'étendre au-delà de la dynamique initiale des échantillons.

Le choix des conditions de bords est illustré sur la figure 4.3 (c) et (d) : des conditions de bords symétriques donnent une reconstruction dont la dérivée première est nulle aux frontières de l'intervalle $\mathcal{I}$. Dans ce cas, f_T est paramétrée par 11 coefficients $c[k]$, $k \in \llbracket 0, 10 \rrbracket$. En (d), où l'on n'impose pas de conditions aux bords, la solution est paramétrée par 13 coefficients $c[k]$, $k \in [-1, 11]$.

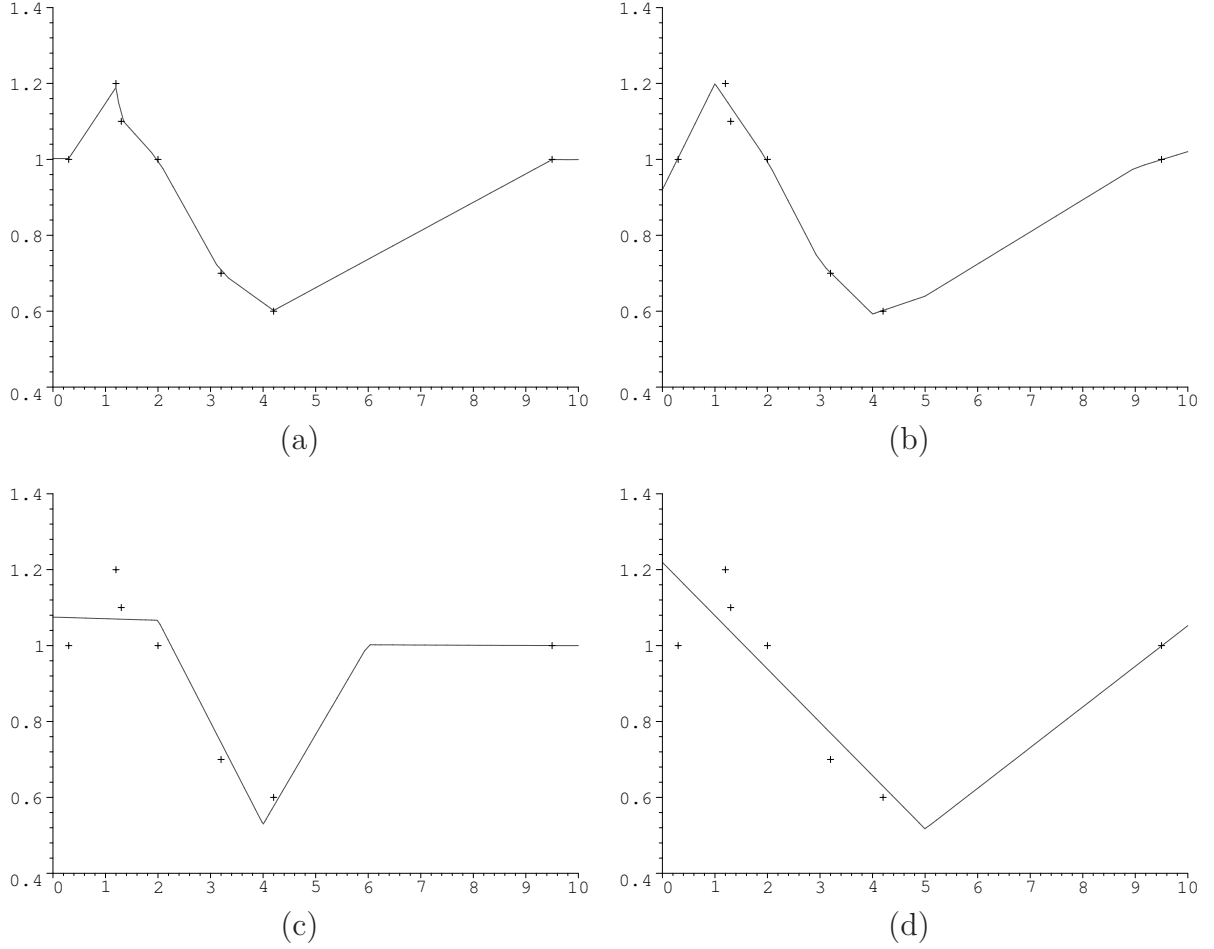


FIG. 4.4 : Splines linéaires uniformes de différentes résolutions ($\varphi = \beta^1$, $r = 1$, $\lambda = 0.01$) reconstruites à partir de 7 échantillons localisés dans l'intervalle $[0, 10]$. (a) : $T = 0.1$. (b) : $T = 1$. (c) : $T = 2$. (d) : $T = 5$. Les splines ont leurs nœuds aux Tk , $k \in [0, 10/T]$.

Le paramètre T permet d'ajuster la résolution de la représentation f_T des données. Si l'on effectue la reconstruction sur un intervalle fixé $[0, S]$, on obtient une solution paramétrique avec $K + 1 = S/T + 1$ degrés de liberté. On a donc une représentation d'autant plus grossière que T est grand, mais aussi d'autant plus parcimonieuse, ce qui peut avoir son importance dans des applications de codage. Choisir une valeur importante de T permet aussi de réduire le temps de calcul. Inversement, lorsque $T \rightarrow 0$, f_T converge vers la solution variationnelle $f(t)$ de (4.4), car l'espace $V_T(\varphi)$ devient dense dans l'espace de Sobolev W_2^r . L'influence de T est illustrée par le figure 4.4, avec $\varphi = \beta^1$, $r = 1$. La fonction f_T est paramétrée par $10/T + 1$ coefficients $c[k]$, $k \in \llbracket 0, 10/T \rrbracket$. Quand $T \rightarrow 0$, f_T tend vers une spline linéaire non uniforme, ayant ses nœuds aux positions t_n .

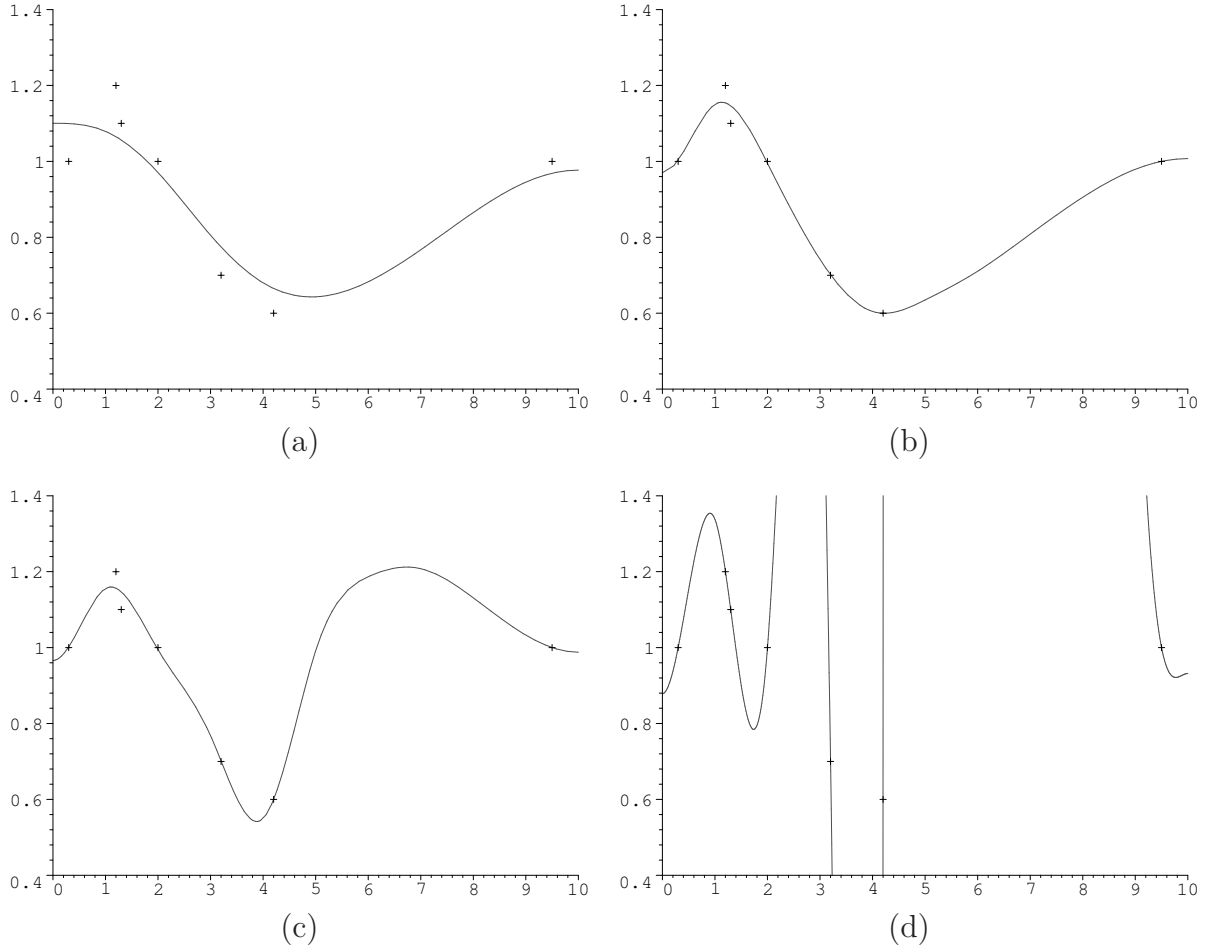


FIG. 4.5 : Splines uniformes cubiques avec leurs nœuds aux entiers ($\varphi = \beta^3$, $r = 2$, $T = 1$, conditions aux bords symétriques) reconstruites à partir de 7 échantillons localisés dans l'intervalle $[0, 10]$, pour différentes valeurs du paramètre de lissage λ . (a) : $\lambda = 1.0$. (b) : $\lambda = 0.001$. (c) : $\lambda = 0.0001$. (d) : cas limite $\lambda = 0^+$.

Le paramètre de régularisation λ est, en pratique, le seul paramètre à ajuster dans notre méthode, les paramètres φ , T , r étant généralement imposés par des considérations liées à l'application concernée. Le choix de λ est important : une valeur excessive rendra la solution excessivement lisse, alors qu'une valeur faible fournira une solution proche des données, mais avec potentiellement de larges variations. Considérons le comportement de f_T dans le cas limite $\lambda = 0^+$. La solution $\mathbf{c} = (\mathbf{M}^T \mathbf{M} + \lambda \mathbf{Q})^{-1} \mathbf{M}^T \mathbf{s}$ a une limite bien définie lorsque $\lambda \rightarrow 0$. Si $\mathbf{M}^T \mathbf{M}$ est inversible, cette limite correspond à ce que l'on obtient en posant $\lambda = 0$ exactement dans les équations. Cela est le cas si l'ensemble des t_n est suffisamment dense, par exemple si $0 < t_{n+1} - t_n < T$ pour tout n . Dans le cas général,

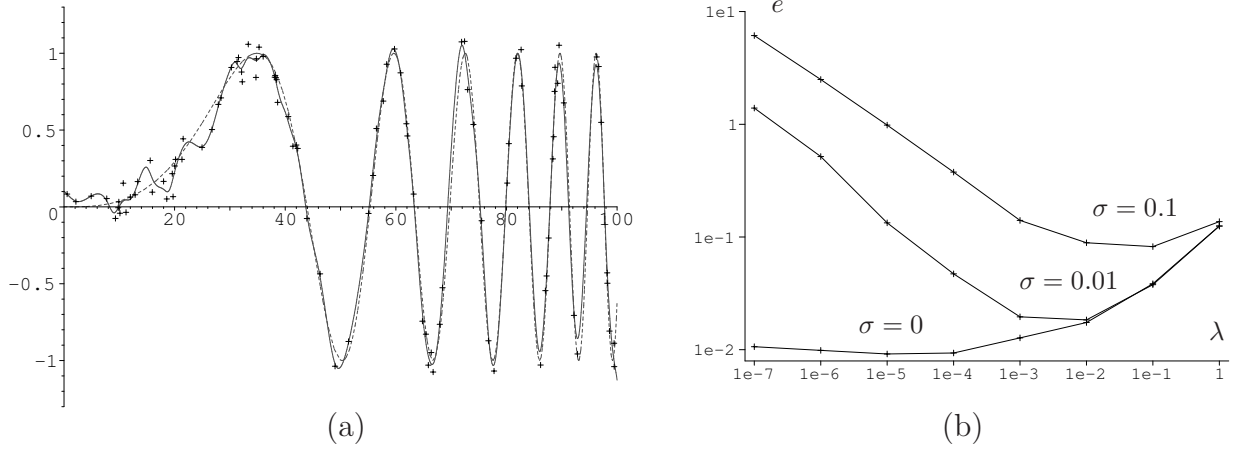


FIG. 4.6 : Reconstruction de la fonction $s(t) = \sin((\frac{t}{30})^3)$ (en pointillés) à partir de 100 échantillons contaminés par un bruit blanc additif gaussien d'écart-type σ à des positions aléatoires localisées dans l'intervalle $[0, 100]$. (a) : spline uniforme cubique ayant ses nœuds aux entiers ($\varphi = \beta^3$, $r = 2$, $T = 1$) reconstruite à partir des données bruitées ($\sigma = 0.1$, $\lambda = 0.1$). (b) : erreur $e = (\frac{1}{100} \int_0^{100} |s(t) - f_T(t)|^2 dt)^{1/2}$ en fonction du niveau de bruit σ et du paramètre λ , en échelle log-log.

la solution pour $\lambda = 0^+$ peut être interprétée comme suit : le terme de régularisation dans (4.8) devient négligeable en comparaison du terme d'attache aux données. Donc f_T minimise le critère des moindres carrés $\sum_n |f_T(t_n) - v[n]|^2$, et si la solution n'est pas unique, les degrés de liberté restants sont mis à profit pour minimiser $\int_T |f_T^{(r)}(t)|^2$.

Le comportement de f_T est donc le suivant : si λ est grand, f_T est lisse, quelles que soient les positions t_n , comme illustré sur la figure 4.5 (a). Si λ est très petit, on peut distinguer deux cas : si localement on a peu de données (c.-à-d. $|t_{n+1} - t_n|$ est grand devant T pour quelques valeurs successives de n), alors f_T interpole presque exactement les échantillons, et les régions vides sont reconstruites de manière lisse. (figure 4.5 (b), (c), (d), pour $t > 3$). Inversement, dans une région dense en échantillons, f_T approche les données au mieux au sens des moindres carrés. (figure 4.5 (b), (c), (d), pour $t < 2$).

En pratique, les échantillons sont souvent bruités, et il apparaît souhaitable de trouver le bon degré de régularisation par l'intermédiaire de λ . Cependant, il n'y a pas de règle absolue fournissant une valeur de λ optimale, et l'ajustement empirique au cas par cas est encore la méthode la plus fiable. Dans certaines applications, il est possible d'apprendre λ à partir des données en utilisant des techniques de validation croisée [76, 118]. De plus, dans le cas d'échantillons uniformes et dans un cadre stochastique, où les données et le bruit sont des réalisations de processus aléatoires stationnaires de caractéristiques spectrales connues, il est suggéré dans [96] de choisir λ inversement proportionnel au rapport signal sur bruit. Cela peut servir d'heuristique dans le cadre déterministe présenté. Cette règle

semble confirmée pour l'exemple montré sur la figure 4.6, où $\lambda = \sigma$ donne l'erreur de reconstruction minimale, lors de l'approximation d'un signal à partir de ses échantillons contaminés par un bruit additif d'écart-type σ . Ainsi, il est tentant, dans le cas où il n'y a pas de bruit, de choisir λ très petit. Cela peut cependant conduire à de larges oscillations imprévues, s'étendant bien au-delà de la dynamique du signal, comme montré sur la figure 4.5 (c), (d) : l'interpolation exacte est donc une contrainte trop forte dans de nombreux cas.

4.3.6 Applications

La méthode de reconstruction que nous avons proposée dans ce chapitre modélise, dans un espace LSI, des échantillons localisés à des positions arbitraires. Le fonction reconstruite $f_T(t)$ dépend d'un paramètre de résolution T , et modélise au mieux les données à cette résolution. Les applications sont nombreuses, où il est nécessaire de modéliser des données par un modèle de résolution fixée. On pense par exemple aux représentations multi-échelles/multi-résolutions, pour laquelle est considérée non pas une fonction f_T , mais toute une collection de fonctions f_T pour différentes valeurs de T . On a ainsi plusieurs représentations, plus ou moins grossières, des données, ce qui est très avantageux algorithmiquement pour des problèmes comme la détection de contours ou la mise en correspondance d'images [210] : une solution grossière est déterminée rapidement et de manière robuste à partir d'une reconstruction de résolution faible, puis elle est progressivement affinée pour prendre en compte les détails à des résolutions plus fines. Cela évite par exemple de se perdre dans des optima locaux en manipulant directement la reconstruction de résolution la plus fine. Dans cette optique, notre approche permet de choisir un pas T quelconque, et non pas seulement une puissance de deux comme dans les approches dyadiques classiques. Détaillons d'autres applications, pour lesquelles notre méthode peut s'avérer utile.

Rééchantillonnage de données non uniformes vers un signal uniforme

Supposons que l'on veuille obtenir un signal uniforme discret u , dont les échantillons $u[k]$ soient localisés aux positions Tk , et ce à partir d'échantillons non uniformes $v[n]$ aux positions t_n . On peut penser par exemple à la visualisation de données non uniformes sur un écran d'ordinateur de pas inter-pixel T . u est alors l'image qui sera affichée sur l'écran. À cette fin, il suffit de reconstruire la fonction $f_T(t) \in V_T(\varphi)$ de résolution $1/T$ voulue, pour une fonction φ choisie arbitrairement ou imposée par le dispositif de sortie (réponse impulsionnelle de l'écran, par exemple). On pose alors $u[k] = f_T(Tk)$ pour tout $k \in \mathbb{Z}$. Ainsi, au contraire de la situation du chapitre 2 où la résolution de f_T dépendait de celle des données, T est maintenant choisi en fonction de l'utilisation que l'on veut faire de f_T : ici, l'échantillonner pour avoir un signal discret de résolution $1/T$.

Le signal u peut être obtenu directement par filtrage discret à partir des coefficients $c[k]$ de f_T . En effet, si l'on définit $b[k] = \varphi(k)$ pour tout $k \in \mathbb{Z}$, on a simplement $u = c * b$.

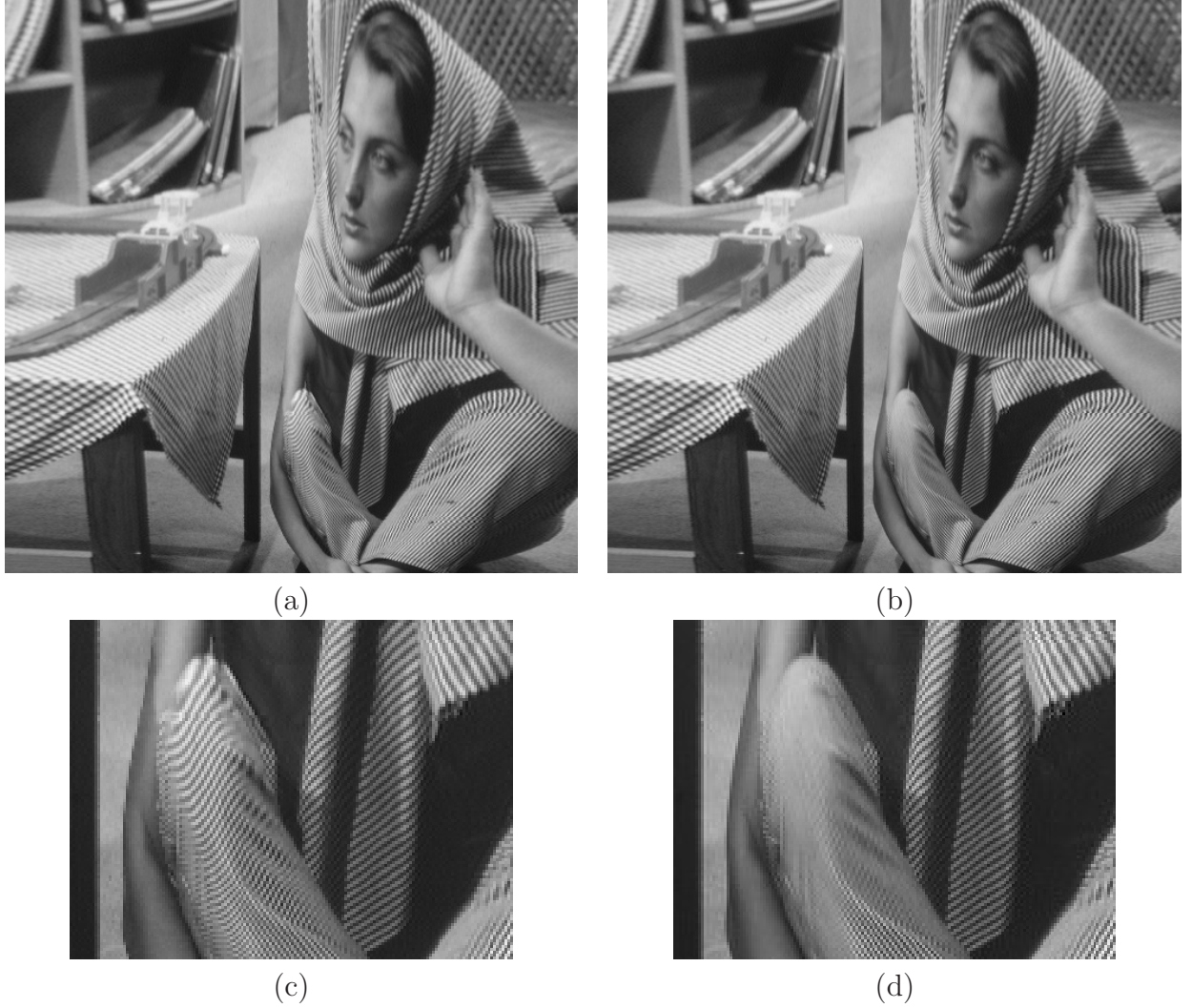


FIG. 4.7 : Image *Barbara* (voir annexe 1) après déformation géométrique. (a) : *backward mapping* avec interpolation spline cubique. (b) : notre approche de reconstruction spline cubique, avec $\lambda = 10^{-4}$. En (c) et (d), détails des images (a) et (b).

Le problème se traite donc de manière entièrement discrète : c est obtenu à partir de v avec notre algorithme, puis u par une simple convolution sur c .

Remarquons que l'approche décrite dans ce paragraphe revient à formuler le rééchantillonnage comme un problème inverse : on cherche le signal u qui, s'il était interpolé par f_T puis rééchantillonné aux positions non uniformes t_n , serait le plus proche des données $v[n]$.

Déformations géométriques

La déformation géométrique (*warping*) d'une image a de nombreuses applications pratiques, par exemple en imagerie de synthèse pour le plaquage de textures [243]. Supposons

que l'on dispose d'une image uniforme, localisée sur le treillis orthogonal de pas T , telle que $v[\mathbf{k}] = s(T\mathbf{k})$ pour une scène s inconnue. On veut obtenir les valeurs $u[\mathbf{k}] = \mathcal{T}v[\mathbf{k}]$ de la scène déformée $\mathfrak{T}s(\mathbf{x}) = s(w^{-1}(\mathbf{x}))$, sur le même treillis, comme on l'a déjà formalisé dans la section 2.6. w est une fonction réversible d'inverse w^{-1} caractérisant la déformation, par exemple $w(\mathbf{x}) = \mathbf{x} + \boldsymbol{\tau}$ pour une translation de pas $\boldsymbol{\tau}$. En d'autres termes, on cherche à évaluer les valeurs $s(w^{-1}(T\mathbf{k}))$ à partir des $v[\mathbf{k}] = s(T\mathbf{k})$.

La méthode classique, appelée *backward mapping*, consiste à reconstruire $f_T = \mathcal{M}v$, généralement par interpolation, puis à échantillonner $\mathfrak{T}f_T$ pour obtenir l'image déformée :

$$u[\mathbf{k}] = \mathfrak{T}\mathcal{M}v(T\mathbf{k}) = f_T(w^{-1}(T\mathbf{k})) \quad \forall \mathbf{k} \in \mathbb{Z}^2. \quad (4.33)$$

Des artéfacts dus au repliement de spectre peuvent apparaître avec cette méthode, puisque le spectre du signal déformé $s(w^{-1}(\mathbf{x}))$ s'étend potentiellement au-delà du seuil de Nyquist du treillis de reconstruction, même si ce n'est pas le cas de s . Cette approche, qui est en fait celle que nous avons adoptée dans la section 2.6 pour la rotation d'image, est donc applicable pour une transformation comme la rotation ou la translation, mais pas pour une déformation arbitraire.

De fait, on connaît les échantillons $v[\mathbf{k}] = s(w^{-1}(w(T\mathbf{k})))$ de la fonction déformée $s(w^{-1}(\mathbf{x}))$ aux positions non uniformes $w(T\mathbf{k})$. On est donc ramené à un problème de rééchantillonnage de données non uniformes vers un signal uniforme, ce dont on a parlé au paragraphe précédent. Si la fonction $w(\mathbf{x})$ est séparable, on peut donc appliquer notre algorithme, le long des lignes puis des colonnes de l'image, pour obtenir le signal $u = \mathcal{T}v$. Il est cependant plus rigoureux d'appliquer l'approche 2D proposée dans [22, 233], impliquant un critère variationnel 2D, bien que l'on perde dans ce cas le bénéfice de notre algorithme rapide.

Afin de donner une illustration concrète de ce problème, nous avons déformé l'image *Barbara* avec $w(x_1/511, x_2/511)/511 = (2.2x_1 - 3.6x_1^2 + 2.4x_1^3, 0.5x_2^2 + 0.5x_2)$, pour $(x_1, x_2) \in [0, 511]^2$. Comme on le voit sur la figure 4.7 en (b) et (e), le motif hautes fréquences du pantalon entraîne du repliement de spectre avec la méthode *backward mapping*, car ce motif, une fois déformé, n'est plus représentable sur le treillis de reconstruction. Notre méthode n'a pas cet inconvénient, et produit une région plane sans artéfact, puisqu'il s'agit localement de la meilleure représentation dans $V_T(\varphi)$ de cette texture déformée, au sens des moindres carrés.

4.4 Reconstruction par quasi-projections

Au chapitre 2, nous avons montré qu'il était possible d'effectuer une meilleure reconstruction qu'avec l'approche des moindres carrés, en effectuant une quasi-projection asymptotiquement optimale de s dans $V_T(\varphi)$. Or, dans ce chapitre, la méthode de reconstruction proposée repose sur une approche des moindres carrés, et coïncide avec l'approche de la section 2.3 lorsque les échantillons sont uniformes aux positions $T\mathbf{k}$. On peut donc espérer

obtenir de meilleurs résultats au moyen d'une approche par quasi-projection. Cependant, les outils d'analyse fréquentielle que nous avons utilisés, en particulier le noyau d'erreur, ne sont plus applicables dans le cas non uniforme. Il faut donc renoncer à l'espoir de concevoir, sur le même modèle que dans le chapitre 2, une méthode asymptotiquement optimale. On remarque tout de même que si φ a un ordre d'approximation égal à L et si $2r \geq L$, notre approche effectue bien une quasi-projection (bien que non asymptotiquement optimale) de s dans $V_T(\varphi)$, dans le cas $\lambda = 0^+$.

Toutefois, il y a un cas particulier dans lequel il est possible d'étendre l'approche du chapitre 2 : on suppose que les données sont uniformes de pas T_0 et non bruitées, c'est-à-dire

$$v[k] = \left\langle s(t), \frac{1}{T_0} \bar{\varphi}\left(\frac{t}{T_0} - k\right) \right\rangle_{L_2}. \quad (4.34)$$

Cela correspond exactement au cas non bruité du modèle (1.10), à l'exception du pas d'échantillonnage $T_0 > 0$ qui n'est plus égal au pas de reconstruction $T > 0$. On souhaite effectuer la reconstruction dans l'espace $V_T(\varphi)$. On est donc confronté à un problème de reconstruction uniforme à partir de données uniformes, mais à une résolution différente. Cette situation a de nombreuses applications en analyse multi-résolution. Les approches usuelles reposent toutes sur des approximations par projection. Une première approche est celle des moindres carrés, comme dans le début de ce chapitre ou dans [3]. Une seconde approche est d'effectuer la reconstruction consistante dans $V_{T_0}(\varphi)$, puis à projeter orthogonalement la fonction obtenue dans $V_T(\varphi)$ [222, 223]. L'utilisation de quasi-projections dans ce contexte est donc nouvelle.

Définissons le facteur de changement de résolution

$$\alpha = \frac{T}{T_0}. \quad (4.35)$$

Nous nous contentons d'étudier le cas où α est un entier. L'étude qui suit pourrait être étendue au cas où α est un nombre rationnel quelconque, puis un réel quelconque par passage à la limite, en utilisant la version multi-ondelettes du noyau d'erreur [36]. On se place en outre dans le cas 1D afin de simplifier les notations, mais l'extension multi-dimensionnelle sur un treillis quelconque ne pose aucun problème.

On cherche donc à reconstruire une fonction de la forme

$$f_T(t) = \sum c[k] \varphi\left(\frac{t}{T} - k\right) \quad (4.36)$$

approchant la fonction inconnue $s(t)$ dont on connaît les mesures $v[k]$ acquises selon le modèle (4.34), et ce pour un pas de reconstruction $T = \alpha T_0$ avec α entier. En imposant à l'opérateur d'approximation $\mathcal{Q} : s \mapsto f_T$ d'être linéaire et invariant par translation de pas multiple de T , on montre aisément que les coefficients $c[k]$ s'obtiennent à partir des données $v[k]$ par filtrage puis décimation :

$$\boxed{c = [v * p_\alpha] \downarrow \alpha}. \quad (4.37)$$

On retrouve bien $c = v * p_1$ dans le cas $\alpha = 1$, étudié au chapitre 2.

De la même manière qu'au chapitre 2, on peut prédire l'erreur d'approximation au moyen d'un noyau d'erreur approprié. Pour définir celui-ci, il suffit de remarquer que le problème ici présent peut être rendu formellement équivalent à celui du chapitre 2. En effet, définissons

$$\tilde{\varphi}_{\text{eq}}(t) = \sum_{k \in \mathbb{Z}} p_\alpha[k] \tilde{\varphi}(t - k) \xleftrightarrow{\mathcal{F}} \hat{\varphi}_{\text{eq}}(\omega) = \hat{p}_\alpha(\omega) \hat{\varphi}(\omega). \quad (4.38)$$

Alors

$$c[k] = \left(s * \frac{1}{T_0} \tilde{\varphi}_{\text{eq}}\left(\frac{\cdot}{T_0}\right) \right)(Tk). \quad (4.39)$$

On est donc bien dans les conditions du modèle (1.10), en remplaçant dans celui-ci la fonction d'analyse $\tilde{\varphi}(t)$ par $\alpha \tilde{\varphi}_{\text{eq}}(\alpha t)$. On peut donc réécrire (2.59) comme suit :

$$\|f_T - s\|_{L_2}^2 \approx \frac{1}{2\pi} \int_{\mathbb{R}} |\hat{s}(\omega)|^2 E(T\omega) d\omega, \quad (4.40)$$

où $E(\omega)$ est le noyau d'erreur :

$$E(\omega) = \underbrace{1 - \hat{\varphi}(\omega) \hat{\varphi}_d(\omega)}_{E_{\min}(\omega) \geq 0} + \underbrace{\hat{a}_\varphi(\omega) \left| \hat{p}_\alpha\left(\frac{\omega}{\alpha}\right) \hat{\varphi}\left(\frac{\omega}{\alpha}\right) - \hat{\varphi}_d(\omega) \right|^2}_{E_{\text{res}}(\omega) \geq 0}. \quad (4.41)$$

On peut dès lors chercher un filtre réalisant une quasi-projection optimale, en imposant comme précédemment

$$E(\omega) \sim E_{\min}(\omega), \quad (4.42)$$

ce qui est maintenant équivalent à avoir

$$p_\alpha(\omega) = \frac{\hat{\varphi}_d(\alpha\omega)}{\hat{\varphi}(\omega)} + O(\omega^N). \quad (4.43)$$

avec $N \geq L + 1$, en notant L l'ordre d'approximation de φ .

Dans le chapitre 2, nous avons opté pour des préfiltres inverses, fournissant de meilleurs résultats que les filtres RIF. Nous proposons d'étendre la forme inverse $P_1(z) = 1/Q_1(z)$ au cas $\alpha \neq 1$, en cherchant un filtre rationnel de la forme $p_\alpha = h_\alpha * [q_\alpha^{-1}] \uparrow \alpha$, autrement dit,

$$P_\alpha(z) = \frac{H_\alpha(z)}{Q_\alpha(z^\alpha)}, \quad (4.44)$$

où

$$H_\alpha(z) = \frac{1}{\alpha^2} \left(\frac{z^{-\frac{\alpha}{2}} - z^{\frac{\alpha}{2}}}{z^{-\frac{1}{2}} - z^{\frac{1}{2}}} \right)^2 = \frac{z^{-\alpha-1}}{\alpha^2} (z + \dots + z^\alpha)^2 \quad (4.45)$$

est un facteur de régularité permettant d'éviter le repliement de spectre dans la bande fréquentielle $\omega \in [\pi/T, \pi/T_0]$.

φ	$Q_2(z)$ tel que $H_2(z) = P_2(z)/Q_2(z^2)$
β^0	1
β^1	$\frac{17}{48}z + \frac{7}{24} + \frac{17}{48}z^{-1}$
β^2	$\frac{3}{16}z + \frac{5}{8} + \frac{3}{16}z^{-1}$
β^3	$\frac{59}{11520}z^2 + \frac{601}{2880}z + \frac{1099}{1920} + \frac{601}{2880}z^{-1} + \frac{59}{11520}z^{-2}$

TAB. 4.1 – Préfiltres proposés pour l'approximation spline à une résolution deux fois plus faible.

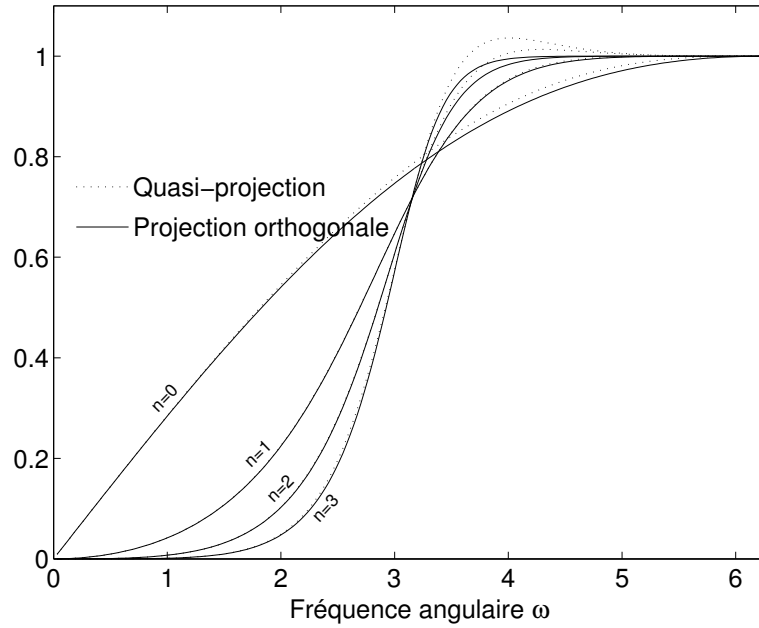


FIG. 4.8 : $\sqrt{E(\omega)}$ lorsque $\varphi = \beta^n$, $\tilde{\varphi} = \delta$, et $\alpha = 2$.

En utilisant l'identité noble [227], permettant d'échanger un filtrage sur-échantillonné et une décimation, on obtient l'implémentation efficace suivante pour l'étape de préfiltrage :

$$c = [v * h_\alpha] \downarrow \alpha * q_\alpha^{-1}. \quad (4.46)$$

Le tableau 4.1 donne les filtres de taille minimale vérifiant (4.43) (donc avec $N = L + 1$) dans le cas $\alpha = 2$ et $\tilde{\varphi} = \delta$. Les noyaux d'erreur $E(\omega)$ associés, représentés sur la figure (4.8), sont très proches des noyaux $E_{\min}(\omega)$; cela indique l'excellente qualité du processus de reconstruction.

4.5 Conclusion

Dans ce chapitre, nous nous sommes intéressés à la reconstruction à partir de données non uniformes, et ce dans un espace de reconstruction uniforme de résolution fixée. La formulation variationnelle proposée est originale et a plusieurs avantages. Elle permet d'abord de lever les hypothèses habituelles sur les positions des échantillons ; ainsi, celles-ci peuvent être complètement arbitraires, sans que cela n'influence le temps de calcul de la reconstruction. Ensuite, le fait de contrôler la résolution de la fonction reconstruite permet de s'adapter aux contraintes d'utilisation et de représentation de la fonction reconstruite, selon l'application envisagée. Notre approche est nouvelle, puisque les méthodes habituelles reconstruisent soit une fonction non-uniforme, soit une fonction uniforme mais avec des hypothèses fortes sur la répartition des échantillons. En exploitant la structure spécifique du problème, nous avons proposé un algorithme efficace pour le calcul de la solution, qui effectue de manière implicite, sans nécessiter la manipulation de matrices, la résolution d'un système linéaire diagonal par bandes. Ce travail est présenté dans les publications [69, 70].

L'extension de l'approche par quasi-projection asymptotiquement optimale des chapitres précédents au cas non uniforme est une question ouverte. L'étude que nous avons menée dans la section 4.4, et publiée dans l'article [67], suggère que des méthodes nouvelles efficaces et rapides pourraient être développées, pour la représentation de données uniformes à des résolutions différentes.

Dans les chapitres suivants, nous nous tournons vers la réduction et l'agrandissement des images naturelles. Les concepts exposés dans ce chapitre nous seront utiles par la suite, en particulier pour l'agrandissement par induction de facteur quelconque (section 7.5.2).

Deuxième partie

Redimensionnement d'images



Chapitre 5

Modèle pour le redimensionnement

5.1 Introduction

Redimensionner une image consiste à en modifier la taille, c'est-à-dire le nombre de pixels, pour obtenir une image qui soit la représentation de la même scène, mais à une résolution différente. Le redimensionnement recouvre deux opérations distinctes : l'*agrandissement*, qui vise à augmenter la résolution de l'image (chaque élément de la scène sera représenté avec plus de pixels), et la *réduction*, qui vise à la réduire. Nous allons traiter du redimensionnement des images, mais le problème peut être transposé aisément à d'autres types de signaux uni-dimensionnels ou multi-dimensionnels. Seulement, les images naturelles ont ceci de particulier que l'information visuellement pertinente est de nature principalement géométrique, et repose sur les contours des objets. Par conséquent, il est primordial de conserver la cohérence de ces structures lors du redimensionnement. Nous verrons que dans le contexte de l'agrandissement, cela impose l'usage de méthodes plus évoluées que les simples méthodes linéaires, alors que ces dernières se montrent généralement suffisantes pour l'agrandissement de signaux 1D, et tout à fait appropriées pour la réduction. L'opération de redimensionnement doit retenir au mieux l'information de l'image initiale, tout en fournissant une image dépourvue d'artéfacts. Le système visuel est très sensible à la présence d'une signature synthétique montrant qu'un traitement artificiel a été opéré ; il faut donc veiller à ce que le redimensionnement laisse une trace la moins visible possible dans l'image résultante. Plus encore en imagerie médicale, où l'image a un impact décisionnel sur le diagnostic du praticien, il importe de maîtriser les distorsions introduites, comme l'introduction de fausses structures ou la disparition d'éléments éventuellement pertinents.

Le redimensionnement a de multiples applications. On pense en premier lieu à la visualisation d'une photographie numérique sur un périphérique d'affichage donné, par exemple un écran d'ordinateur. Les appareils photos numériques récents génèrent des images bien plus grandes que les écrans, et il faut donc les réduire pour les visualiser. De plus, à l'heure

où les applications multimédias occupent une place prépondérante dans tous les secteurs d'activité, les besoins sont croissants en vidéophonie et dans les domaines du stockage et de la transmission des images. Malgré l'avènement des réseaux haut débit, les technologies de compression restent souvent insuffisantes pour permettre la transmission en temps réel d'une vidéo. La réduction apparaît alors comme un moyen de pallier à ce problème. Inversement, il est parfois requis d'afficher ou d'imprimer une petite image avec une plus grande résolution, que l'image ait été précédemment réduite, ou qu'elle soit issue d'un processus d'acquisition de résolution insuffisante. L'agrandissement de l'image est alors nécessaire. Cette opération permet aussi de « zoomer », comme avec une loupe virtuelle, sur une partie d'une image, pour en scruter un détail et mieux comprendre la nature de l'objet représenté.

Les applications du redimensionnement vont au-delà du rendu et de la visualisation à une résolution différente de celle de l'image initiale. En imagerie de synthèse animée, on peut simuler le mouvement en changeant la taille effective d'un objet dans le plan de vue, et donc son éloignement perçu par l'observateur. Les techniques dites *multi-résolution*, qui sont utilisées dans des champs d'applications très vastes, nécessitent aussi des méthodes de changement de résolution, c'est-à-dire de redimensionnement. Le fait de disposer de différentes versions, à des résolutions différentes, d'une même image, offre tout un éventail de possibilités pour l'analyse. Ainsi, de nombreux algorithmes, en reconnaissance de formes, détection de contour, ou segmentation, utilisent une stratégie pyramidale, en commençant par agir à une résolution grossière, puis en affinant leurs résultats en progressant au fur et à mesure vers les niveaux de résolution les plus fins. La théorie des ondelettes repose aussi sur l'analyse multi-résolution [149] ; elle est à la base des méthodes modernes de compression d'images, domaine lié à celui du redimensionnement [208].

5.2 Formalisation du redimensionnement

Afin de formuler précisément le problème du redimensionnement, il est important de garder à l'esprit la notion de localisation associée à une image. Nous avons en effet insisté dans la première partie de cette thèse sur l'importance d'interpréter les échantillons comme des mesures localisées spatialement. Chaque pixel d'une image se trouve donc positionné à un site du treillis associé à l'image. Avant de définir le redimensionnement, nous étendons la notion de treillis 2D à celle de *treillis translaté*, défini par

$$\Lambda_{\mathbf{R},\boldsymbol{\tau}} = \Lambda_{\mathbf{R}} + \mathbf{R}\boldsymbol{\tau} = \{\mathbf{R}(\mathbf{k} + \boldsymbol{\tau}) \mid \mathbf{k} \in \mathbb{Z}^2\}. \quad (5.1)$$

Cela permet de localiser une image sur un treillis décalé par rapport à l'origine. On ne s'intéressera qu'au redimensionnement d'images localisées sur des treillis orthogonaux, même si l'extension à des treillis quelconques ne pose pas de problème particulier. Ainsi, on écrira dans la suite $\Lambda_{T,\boldsymbol{\tau}}$ au lieu de $\Lambda_{T\mathbf{I},\boldsymbol{\tau}}$, où $\mathbf{I}$ est la matrice identité.

On définit maintenant le redimensionnement $\mathcal{T}_{\alpha,\boldsymbol{\tau}}$ de facteur α et d'offset $\boldsymbol{\tau}$. Cette opération fournit, à partir d'une image localisée sur le treillis $\Lambda_{T,\boldsymbol{\tau}_0}$, une image localisée sur

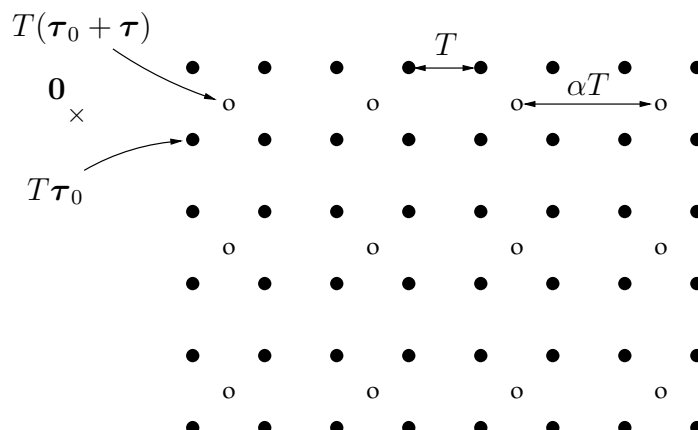


FIG. 5.1 : En redimensionnant une image définie sur le treillis noir Λ_{T, τ_0} avec un facteur $\alpha = 2$ et un offset $\tau = (\frac{1}{2}, \frac{1}{2})$, on obtient une image localisée sur le treillis blanc $\Lambda_{\alpha T, (\tau_0 + \tau)/\alpha}$. Inversement, on passe d'une image localisée sur le treillis blanc à une autre localisée sur le treillis noir par redimensionnement de facteur $\frac{1}{2}$, d'offset $(-\frac{1}{4}, -\frac{1}{4})$.

le treillis $\Lambda_{\alpha T, (\tau_0 + \tau)/\alpha}$. Un exemple est représenté sur la figure 5.1. Afin de préciser la nature de l'image redimensionnée, et donc la valeur attendue de ses pixels, on définit l'opérateur de redimensionnement idéal $\mathcal{T}_{\alpha, \tau}^{\text{id}}$, fournissant l'image redimensionnée que l'on souhaiterait avoir dans l'idéal. Celle-ci n'étant généralement pas accessible, on ne pourra qu'approcher cet opérateur en pratique, hormis dans des cas particuliers que nous détaillerons. Nous nous basons sur le modèle de formation d'images qui nous a servi jusqu'alors : on suppose que l'image initiale v est telle que

$$v[\mathbf{k}] = \int_{\mathbb{R}^2} s(\mathbf{x}) \tilde{\varphi}\left(\mathbf{k} + \tau_0 - \frac{\mathbf{x}}{T}\right) \frac{d\mathbf{x}}{T^2} \quad \forall \mathbf{k} \in \mathbb{Z}^2, \quad (5.2)$$

pour un certain processus inconnu sous-jacent $s(\mathbf{x})$ et une fonction d'analyse $\tilde{\varphi}(\mathbf{x})$. On suppose donc que l'on veut redimensionner une image non bruitée. Nous n'aborderons pas le problème du débruitage, qui représente à lui seul un vaste champ d'études en traitement d'images, et repose sur des méthodes non-linéaires et/ou adaptatives évoluées. De même que pour la reconstruction, le mieux à faire dans le cas bruité est de débruiter l'image dans un premier temps, puis d'effectuer le redimensionnement ensuite.

Remarquons que pour que la notion de localisation des échantillons ait un sens, il faut que $\tilde{\varphi}(\mathbf{x})$ soit essentiellement localisée en $\mathbf{x} = \mathbf{0}$, typiquement, symétrique en $\mathbf{0}$.

L'image v redimensionnée de facteur α et d'offset τ est alors définie comme l'image qui aurait été obtenue avec le même processus de formation, mais déployé sur le treillis $\Lambda_{\alpha T, (\tau_0 + \tau)/\alpha}$ au lieu de Λ_{T, τ_0} :

$$\boxed{\mathcal{T}_{\alpha, \tau}^{\text{id}} v[\mathbf{k}] = \int_{\mathbb{R}^2} s(\mathbf{x}) \tilde{\varphi}\left(\mathbf{k} + \frac{\tau_0 + \tau}{\alpha} - \frac{\mathbf{x}}{\alpha T}\right) \frac{d\mathbf{x}}{(\alpha T)^2} \quad \forall \mathbf{k} \in \mathbb{Z}^2.} \quad (5.3)$$

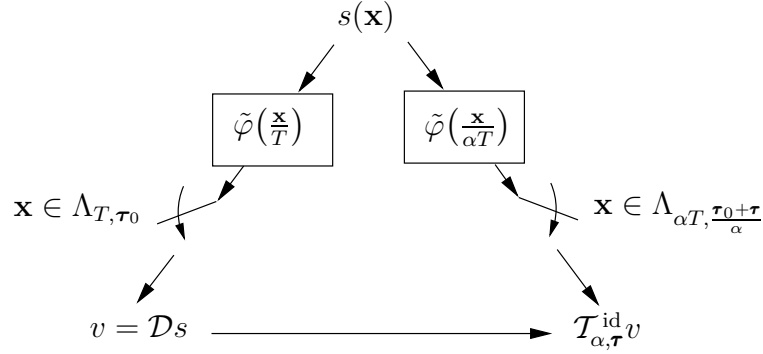


FIG. 5.2 : En redimensionnant une image définie sur le treillis Λ_{T, τ_0} avec un facteur α et un offset τ , on vise à obtenir l'image qui aurait été acquise sur le treillis $\Lambda_{\alpha T, (\tau_0 + \tau)/\alpha}$, avec le même dispositif, mais dilaté de facteur α .

Hormis l'introduction d'un offset, cela revient simplement à remplacer T par αT dans (5.2). Autrement dit, l'image redimensionnée est celle qui aurait été obtenue, si le dispositif d'acquisition avait été parfaitement dilaté de facteur α . On peut aussi interpréter l'image redimensionnée comme **l'image qui aurait été obtenue si la scène $s(\mathbf{x})$ avait été remplacée par $s(\alpha \mathbf{x})$, et acquise avec le même dispositif**. Dans les notations de la section 2.6, le redimensionnement est donc l'équivalent discret de la contraction continue $\mathfrak{T} : s(\mathbf{x}) \mapsto s(\alpha \mathbf{x})$. Ce processus de redimensionnement idéal est schématisé sur la figure 5.2. Intuitivement, dans le cas d'une photographie, cela revient à considérer, en négligeant les changements de perspective, que le photographe s'approche ou s'éloigne de la scène avant d'effectuer sa prise de vue.

Notons que nous avons introduit la notion d'offset, car il n'y a pas de raison de redimensionner des images en calant spatialement leurs origines plutôt que tout autre pixel. La notion de position absolue n'a pas vraiment de sens pour les images, et le choix de l'origine est arbitraire. Remarquons aussi que l'opération de redimensionnement ne dépend pas intrinsèquement de T et τ_0 , pour le calcul des valeurs des pixels de l'image redimensionnée. Ces deux paramètres permettent simplement de définir le treillis de localisation de l'image redimensionnée, de manière cohérente avec celui de l'image initiale. Par conséquent, étant donnée une image d'origine inconnue, on peut de manière non restrictive la supposer localisée sur le treillis $\Lambda_{1,0}$. Associer un treillis à une image n'est pas nécessaire en soi pour concevoir des méthodes de redimensionnement, mais cela nous semble important pour définir formellement le redimensionnement et en comprendre la nature.

Signalons qu'avec notre définition, pour $\alpha = 1$, le redimensionnement est simplement une translation lorsque $\tau \neq \mathbf{0}$. La translation peut donc être vue comme un cas particulier du redimensionnement, et il est intéressant de garder à l'esprit que certaines méthodes présentées pourront être utilisées pour effectuer une translation. Enfin, $\mathcal{T}_{1,0}^{\text{id}}$ est l'identité.

Le redimensionnement idéal que nous avons défini ne peut être réalisé en pratique, car le signal sous-jacent s est inconnu. Cependant, il est intéressant d'étudier certaines propriétés vérifiées par le redimensionnement idéal, et susceptibles d'être vérifiées par un opérateur concret $\mathcal{T}_{\alpha,\tau}$.

- On dira que $\mathcal{T}_{\alpha,\tau}$ est linéaire si

$$\mathcal{T}_{\alpha,\tau}(v_1 + \lambda v_2) = \mathcal{T}_{\alpha,\tau}(v_1) + \lambda \mathcal{T}_{\alpha,\tau}(v_2) \quad \forall v_1, v_2 \in \ell_2(\mathbb{Z}^2), \lambda \in \mathbb{R}. \quad (5.4)$$

$\mathcal{T}_{\alpha,\tau}$ est linéaire si et seulement si il existe un ensemble de filtres $\{h_{\mathbf{k}} \in \ell_2 \mid \mathbf{k} \in \mathbb{Z}^2\}$ tels que

$$\mathcal{T}_{\alpha,\tau} v[\mathbf{k}] = \sum_{\mathbf{l} \in \mathbb{Z}^2} h_{\mathbf{k}}[-\mathbf{l}] v[\mathbf{l}] \quad \forall \mathbf{k} \in \mathbb{Z}^2. \quad (5.5)$$

- $\mathcal{T}_{\alpha,\tau}$ est dit *invariant par translation* si on peut permuter une translation et un redimensionnement : soit v_1 l'image translatée de $\mathbf{l} \in \mathbb{Z}^2$ à partir de v : $v_1[\mathbf{k}] = v[\mathbf{k} - \mathbf{l}]$ pour tout $\mathbf{k} \in \mathbb{Z}^2$. Alors $\mathcal{T}_{\alpha,\tau}$ est *invariant par translation* si pour tout $\mathbf{k}, \mathbf{l} \in \mathbb{Z}^2$ tels que $\frac{1}{\alpha} \in \mathbb{Z}^2$, on a :

$$\mathcal{T}_{\alpha,\tau} v_1[\mathbf{k}] = \mathcal{T}_{\alpha,\tau} v[\mathbf{k} - \frac{1}{\alpha}] \quad \forall \mathbf{k} \in \mathbb{Z}^2. \quad (5.6)$$

- $\mathcal{T}_{\alpha,\tau}$ est linéaire et invariant par translation si et seulement si il existe un noyau $h_{\alpha,\tau}(\mathbf{x})$ tel que

$$\mathcal{T}_{\alpha,\tau} v[\mathbf{k}] = \sum_{\mathbf{l} \in \mathbb{Z}^2} h_{\alpha,\tau}(\alpha \mathbf{k} + \boldsymbol{\tau} - \mathbf{l}) v[\mathbf{l}] \quad \forall \mathbf{k} \in \mathbb{Z}^2. \quad (5.7)$$

Dans le cas où $\alpha = \frac{p}{q}$ est rationnel, seul un nombre fini de valeurs de $h_{\alpha,\tau}(\mathbf{x})$ interviennent dans (5.7). On peut donc montrer [31] que pour α rationnel, $\mathcal{T}_{\alpha,\tau}$ est linéaire et invariant par translation si et seulement si il existe un filtre discret $(h_{\alpha,\tau}[\mathbf{k}])_{\mathbf{k} \in \mathbb{Z}^2}$ tel que

$$\mathcal{T}_{\alpha,\tau} v = [[v] \uparrow q * h_{\alpha,\tau}] \downarrow p. \quad (5.8)$$

Si α n'est pas rationnel, l'invariance par translation n'est pas définie, et alors tout opérateur $\mathcal{T}_{\alpha,\tau}$ linéaire peut se mettre sous la forme (5.7).

Si maintenant on s'intéresse à la cohérence des opérateurs $\mathcal{T}_{\alpha,\tau}$ entre eux, pour différentes valeurs de α et $\boldsymbol{\tau}$, on aimerait avoir la propriété hiérarchique

$$\mathcal{T}_{\alpha_2,\boldsymbol{\tau}_2} \circ \mathcal{T}_{\alpha_1,\boldsymbol{\tau}_1} = \mathcal{T}_{\alpha_1\alpha_2,\boldsymbol{\tau}_1+\alpha_1\boldsymbol{\tau}_2} \quad \forall \alpha_1, \alpha_2 \in \mathbb{R}, \boldsymbol{\tau}_1, \boldsymbol{\tau}_2 \in \mathbb{R}^2. \quad (5.9)$$

Dans le cas d'opérateurs linéaires et invariants par translation, cela réduit le choix de $h_{\alpha,\tau}(\mathbf{x})$ dans (5.7) à la seule solution intéressante

$$h_{\alpha,\tau}(\mathbf{x}) = \begin{cases} \frac{1}{\alpha^2} \text{sinc}\left(\frac{\pi \mathbf{x}}{\alpha}\right) & \text{si } \alpha < 1 \\ \text{sinc}(\pi \mathbf{x}) & \text{sinon} \end{cases}. \quad (5.10)$$

Nous allons voir par la suite que le redimensionnement à l'aide de la fonction sinus cardinal donne des résultats insatisfaisants visuellement, à cause des artéfacts de *ringing* qui

apparaissent. Il faudra donc, en pratique, renoncer à la propriété hiérarchique, ce qui est fort dommage du point de vue théorique.

Dans la suite de cette thèse, on distinguera la réduction et l'agrandissement, et on étudiera ces deux problèmes séparément. Si $\alpha > 1$, on a affaire à une réduction de facteur α , et on notera $\mathcal{R}_{\alpha,\tau}$ au lieu de $\mathcal{T}_{\alpha,\tau}$. Si $\alpha < 1$, on parlera d'agrandissement de facteur $1/\alpha$, et on notera $\mathcal{A}_{\frac{1}{\alpha},\tau}$ au lieu de $\mathcal{T}_{\alpha,\tau}$.

5.2.1 Considérations sur la fonction $\tilde{\varphi}$

Dans toute la première partie de cette thèse, consacrée à la reconstruction, nous n'avons pas fait d'hypothèse sur la fonction $\tilde{\varphi}$, que nous avons laissée comme paramètre libre du modèle. De nombreux problèmes, comme la rotation et la translation, se ramènent au cas $\tilde{\varphi} = \delta$. En règle générale cependant, $\tilde{\varphi}$ est souvent inconnue en pratique, et le redimensionnement, comme la reconstruction, ne peuvent alors pas être effectués. Il faut donc faire au préalable une hypothèse vraisemblable sur la fonction $\tilde{\varphi}$ ayant permis la formation de l'image v dont on dispose.

Du point de vue physique, la réponse impulsionnelle d'un capteur ne peut être que positive. Cependant, comme le signal est généralement corrigé par un post-filtrage discret, le dispositif d'acquisition dans son ensemble peut avoir une réponse impulsionnelle $\tilde{\varphi}$ arbitraire. En ce qui concerne les appareils photos numériques, $\tilde{\varphi}(\mathbf{x})$ est la convolution d'une fonction isotrope gaussienne englobant les distorsions optiques [193, 47], d'une fonction porte $\mathbb{1}_{]-\frac{1}{2}, \frac{1}{2}[}^2$ modélisant l'intégration lumineuse sur la surface d'un élément carré du capteur CCD [243], et d'un post-filtre discret effectuant un réhaussement fréquentiel de l'image. Cependant, chaque appareil photo a une réponse qui lui est propre, et l'estimation de $\tilde{\varphi}$ à partir de l'image v est très difficile. Malgré cela, nous sommes assez peu sensibles aux variations de $\tilde{\varphi}$, et une estimation grossière de cette fonction suffit.

La question qu'il faut en fait se poser est la suivante : si on maîtrisait le processus d'acquisition, quelle fonction $\tilde{\varphi}^{\text{id}}$ faudrait-il choisir, dans l'idéal ? Afin d'étudier l'influence de tel ou tel choix de $\tilde{\varphi}^{\text{id}}$, il est possible de simuler une acquisition de la manière suivante : on part d'une image de grande taille v_0 , supposée acquise sur $\Lambda_{1,0}$ avec une fonction d'analyse $\tilde{\varphi}_0$ inconnue, et on la réduit N fois d'un facteur 2 à l'aide d'un filtre discret h symétrique, pour obtenir

$$v = \left[\cdots \left[[v_0 * h] \downarrow 2 * h \right] \downarrow 2 \cdots * h \right] \downarrow 2 \quad (5.11)$$

$$= [v_0 * h_N] \downarrow 2^N, \quad (5.12)$$

où $h_N = h * [h] \uparrow 2 * \cdots * [h] \uparrow 2^{N-1}$. Alors de

$$v_0[\mathbf{k}] = \int_{\mathbb{R}^2} s(\mathbf{x}) \tilde{\varphi}_0(\mathbf{k} - \mathbf{x}) d\mathbf{x} \quad \forall \mathbf{k} \in \mathbb{Z}^2, \quad (5.13)$$

on déduit

$$v[\mathbf{k}] = \int_{\mathbb{R}^2} s(\mathbf{x}) \tilde{\varphi}(\mathbf{k} - \frac{\mathbf{x}}{2^N}) \frac{d\mathbf{x}}{2^{2N}} \quad \forall \mathbf{k} \in \mathbb{Z}^2, \quad (5.14)$$

où $\tilde{\varphi}$ est définie par

$$\tilde{\varphi}(\mathbf{x}) = 2^{2N} \sum_{\mathbf{k} \in \mathbb{Z}^2} h_N[\mathbf{k}] \tilde{\varphi}_0(2^N \mathbf{x} - \mathbf{k}) \xleftrightarrow{\mathcal{F}} \hat{\tilde{\varphi}}(\boldsymbol{\omega}) = \hat{\varphi}_0\left(\frac{\boldsymbol{\omega}}{2^N}\right) \prod_{i=1}^N \hat{h}\left(\frac{\boldsymbol{\omega}}{2^i}\right). \quad (5.15)$$

Pour N grand ($N = 2$ suffit en pratique), on peut donc faire l'approximation

$$\hat{\tilde{\varphi}}(\boldsymbol{\omega}) = \prod_{i=1}^{+\infty} \hat{h}\left(\frac{\boldsymbol{\omega}}{2^i}\right), \quad (5.16)$$

et $\tilde{\varphi}$ est alors entièrement caractérisée par le filtre h , que l'on appelle son **filtre d'échelle** (*refinement filter*), tel que

$$H(\mathbf{z}) = \frac{\sum_{\mathbf{k} \in \mathbb{Z}^2} \tilde{\varphi}(\frac{\mathbf{k}}{2}) \mathbf{z}^{-\mathbf{k}}}{2^2 \sum_{\mathbf{k} \in \mathbb{Z}^2} \tilde{\varphi}(\mathbf{k}) \mathbf{z}^{-\mathbf{k}}}. \quad (5.17)$$

$\tilde{\varphi}$ vérifie alors la *relation à deux échelles* :

$$\tilde{\varphi}(\mathbf{x}) = 2^2 \sum_{\mathbf{k} \in \mathbb{Z}^2} h[\mathbf{k}] \tilde{\varphi}(2\mathbf{x} - \mathbf{k}). \quad (5.18)$$

Ce type de relation est à la base de la théorie des ondelettes [149, 83].

On peut ainsi, par cette simple expérience de réduction, et en choisissant un certain filtre h , simuler une acquisition avec la fonction d'analyse $\tilde{\varphi}$ ayant h pour filtre d'échelle. Par exemple, en prenant $H(\mathbf{z}) = (\frac{1}{4}z_1 + \frac{1}{2} + \frac{1}{4}z_1^{-1})(\frac{1}{4}z_2 + \frac{1}{2} + \frac{1}{4}z_2^{-1})$, on peut simuler une acquisition avec la fonction B-spline séparable de degré 1 :

$$\tilde{\varphi}(\mathbf{x}) = \begin{cases} (1 - |x_1|)(1 - |x_2|) & \text{si } \mathbf{x} \in [-1, 1]^2 \\ 0 & \text{sinon} \end{cases}. \quad (5.19)$$

Si on effectue la décimation sans préfiltrage, c'est-à-dire avec $H(\mathbf{z}) = 1$, on simule une acquisition avec $\tilde{\varphi} = \delta$.

Après de nombreux essais sur des images variées, nous avons abouti aux constatations suivantes, par ailleurs bien connues :

- Si $\hat{\tilde{\varphi}}(\boldsymbol{\omega})$ n'est pas proche de 0 en dehors de la zone de Nyquist, il en résulte des effets de moiré et des effets d'escaliers sur les contours obliques dans l'image, dus au repliement de spectre (*aliasing*) lors de l'échantillonnage. Ces effets sont visibles sur la figure 5.3.
- Si $\hat{\tilde{\varphi}}(\boldsymbol{\omega})$ n'est pas proche de 1 dans la zone de Nyquist, il en résulte un effet de flou dans l'image.
- Si la transition de $\hat{\tilde{\varphi}}(\boldsymbol{\omega})$ de 1 à 0 entre l'intérieur et l'extérieur de la zone de Nyquist est trop brutale, il en résulte des oscillations dans toute l'image (effet dit de *ringing*), comme le montre la figure 5.3 (c).

Il faut donc trouver une fonction $\tilde{\varphi}^{\text{id}}$ qui offre un compromis entre les effets de flou, de repliement de spectre, et d'oscillations, sachant que ces effets sont antagonistes, et que l'on ne peut pas les annuler tous les trois. Nous avons abouti empiriquement à la conclusion que la fonction offrant le meilleur compromis, du point de vue de la qualité visuelle des images générées, est $\boxed{\tilde{\varphi}^{\text{id}}(\mathbf{x}) = \gamma^3(\mathbf{x})}$, la fonction spline cardinale cubique séparable (représentée en 1D, ainsi que son spectre, sur la figure 2.5), que l'on peut définir par :

$$\hat{\gamma}^3(\omega) = \frac{1}{(\frac{2}{3} + \frac{1}{3} \cos(\omega_1))(\frac{2}{3} + \frac{1}{3} \cos(\omega_2))} \frac{4 \sin(\frac{\omega_1}{2}) \sin(\frac{\omega_2}{2})}{\omega_1 \omega_2}. \quad (5.20)$$

Cette fonction vérifie une relation à deux échelles dyadique (5.18) avec le filtre d'échelle $H(\mathbf{z}) = H(z_1)H(z_2)$, où

$$H(z) = \frac{\frac{1}{96}(z^3 + z^{-3}) + \frac{1}{12}(z^2 + z^{-2}) + \frac{23}{96}(z + z^{-1}) + \frac{1}{3}}{\frac{1}{6}(z^2 + z^{-2}) + \frac{2}{3}}. \quad (5.21)$$

Etant donnée une image v , il est difficile d'identifier la fonction $\tilde{\varphi}$ caractérisant le dispositif d'acquisition qui l'a générée. Cependant, si l'image est de bonne qualité et ne présente pas de défaut particulier (flou ou moiré excessif), on en déduit que $\tilde{\varphi}$ est proche de $\tilde{\varphi}^{\text{id}}$. Par conséquent, étant donnée une image naturelle de bonne qualité visuelle, on fera l'hypothèse qu'elle a été engendrée à partir du modèle (5.2), avec $\boxed{\tilde{\varphi} = \gamma^3}$.

Le modèle de redimensionnement que nous avons décrit dans cette section présente des similarités avec celui proposé par Hussein Aly [16, 15]. Cet auteur distingue cependant le dispositif d'acquisition ayant engendré v de celui générant l'image redimensionnée. Cette sophistication permet de prendre en compte le cas où l'image v a été générée avec un dispositif de réponse $\tilde{\varphi}$ connue et différente de la réponse idéale $\tilde{\varphi}^{\text{id}}$. Nous ne faisons pas une telle distinction, et pour nous, $\tilde{\varphi} = \tilde{\varphi}^{\text{id}}$, en l'absence d'éléments laissant supposer le contraire dans l'image v . D'autre part, Aly formule le redimensionnement comme un problème de reconstruction : il cherche l'image redimensionnée qui, affichée sur un périphérique de réponse φ (typiquement, un écran CRT), donne une fonction f_T proche de s . Cela l'amène à faire une hypothèse qu'il considère vraisemblable sur φ et à choisir $\tilde{\varphi}^{\text{id}} = \varphi_d$, la fonction duale de φ . On remarque que ce contexte est exactement celui que nous avons abordé à la section 4.4 : la reconstruction d'une fonction de résolution différente de celle des mesures $v[\mathbf{k}]$. En ce qui nous concerne, nous préférons formuler le redimensionnement d'une image indépendamment de la manière dont elle est visualisée. En pratique, les périphériques d'affichage et d'impression sont optimisés et étalonnés pour donner de bons résultats à partir d'images « types ». Il n'y a donc pas lieu de se demander comment préfiltrer une image pour optimiser son rendu, à chaque fois qu'on l'affiche.

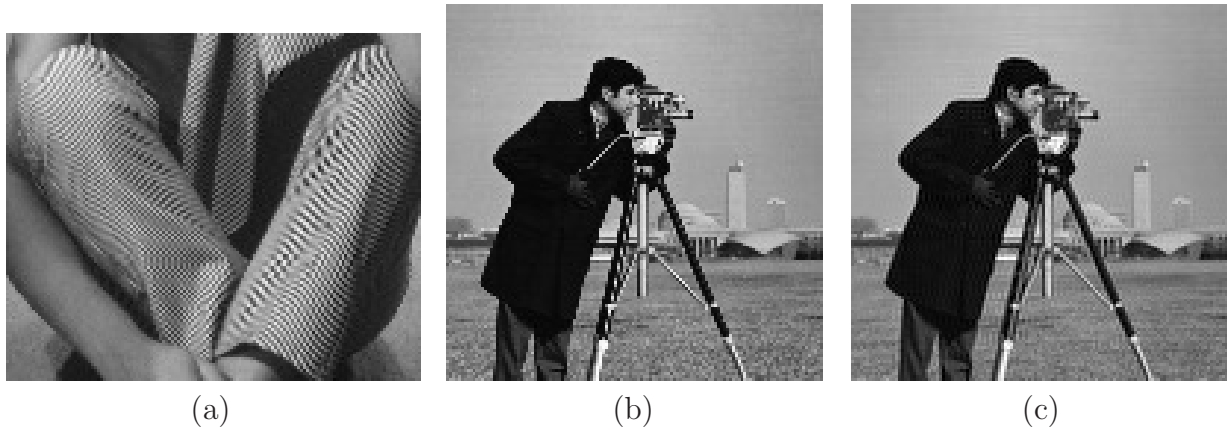


FIG. 5.3 : En réduisant de facteur 2 le pantalon de Barbara et le cameraman par décimation sans préfiltrage, le repliement de spectre se manifeste de manière très visible par du moiré (a) et des effets d'escaliers sur les contours obliques (b). À l'opposé, la réduction de facteur 2 à l'aide du filtre sinus cardinal introduit de nombreuses oscillations se propageant dans toute l'image à partir des contours (c).

5.3 Réduction d'images

La réduction d'une image consiste à diminuer le nombre de pixels qui la composent. Dans le cas d'une réduction d'un facteur 2, l'image réduite contiendra ainsi quatre fois moins de pixels que l'image initiale. Cela conduit à une disparition de structures contenues dans l'image initiale et, si l'on n'y prend garde, à l'apparition de distorsions indésirables, comme des effets de moiré ou de crénelage. La réduction est un processus avec pertes, et il faudra s'efforcer de conserver le maximum de l'information contenue dans l'image initiale, tout en maintenant une bonne qualité visuelle.

La méthode généralement employée pour réduire une image est décrite par les formules (5.7) et (5.8) : l'image est filtrée puis échantillonnée sur le treillis d'arrivée (ce qui équivaut à un filtrage discret suivi d'une décimation dans la cas où α est entier). Cela revient à chercher un processus de réduction linéaire et invariant par translation, ce qui est intuitif d'une part, et conduit aux méthodes pratiques les plus simples d'autre part. Nous justifierons ce procédé à l'aide de notre modèle de redimensionnement dans la suite de cette section. Le problème se ramène donc au choix du filtre $h_{\alpha,\tau}$.

On remarque tout d'abord, comme le montre la figure 5.3, que l'absence de filtrage avant le rééchantillonnage occasionne des artéfacts très visibles dus au repliement de spectre. En réécrivant (5.7) dans le domaine fréquentiel, on obtient

$$\widehat{\mathcal{R}_{\alpha,\tau}v}(\omega) = \sum_{\mathbf{k} \in \mathbb{Z}^2} e^{I\langle \omega + 2\pi\mathbf{k}, \alpha\tau \rangle} \hat{v}\left(\frac{\omega + 2\pi\mathbf{k}}{\alpha}\right) \hat{h}_{\alpha,\tau}\left(\frac{\omega + 2\pi\mathbf{k}}{\alpha}\right). \quad (5.22)$$

Les fréquences au-delà du seuil de Nyquist $\frac{\pi}{\alpha}$ sont perdues lors de l'échantillonnage, et se re-

plient spectralement sous forme de basses fréquences parasites, créant de fausses structures. Il faut donc à la fois supprimer ces répliques, et conserver intactes les basses fréquences. La suppression de tout repliement de spectre est équivalente à la condition $\hat{h}(\omega) = 0$ en dehors de la bande $]-\frac{\pi}{\alpha}, \frac{\pi}{\alpha}[^2$. De plus, les distorsions du contenu fréquentiel conservé, typiquement un effet de flou dû à une atténuation, sont évitées si et seulement si $\hat{h}(\omega) = 1$ dans la bande $]-\frac{\pi}{\alpha}, \frac{\pi}{\alpha}[^2$. Ces deux contraintes définissent le filtre de réduction sinus cardinal :

$$h_{\alpha,\tau}(\mathbf{x}) = \frac{1}{\alpha^2} \text{sinc}\left(\frac{\pi\mathbf{x}}{\alpha}\right). \quad (5.23)$$

Bien que ce filtre n'entraîne ni repliement de spectre, ni effet de flou dans l'image réduite, des oscillations, visibles sur la figure 5.3 (c), apparaissent dans toute l'image. La réduction à l'aide de la fonction sinus cardinal n'est donc pas satisfaisante et il faudra, pour atténuer le *ringing*, choisir une fonction $h_{\alpha,\tau}(\mathbf{x})$ dont le spectre présente une transition moins abrupte au niveau de la frontière de la région de Nyquist. Cela implique nécessairement la présence d'un peu de flou et/ou de repliement de spectre dans l'image réduite. Il y a donc un compromis à trouver entre les effets de flou, *ringing*, et *aliasing*, sachant qu'il est possible d'en annuler deux parmi les trois, mais jamais les trois à la fois. Blanchet et coll. ont ainsi proposé une classe de filtres ne présentant aucun *aliasing*, et paramétrée par un facteur de compromis entre flou et *ringing* [30].

Notons que la réduction peut être effectuée en domaine fréquentiel, en multipliant la FFT de l'image par la réponse fréquentielle du filtre, avant de la tronquer à la taille voulue et d'effectuer une FFT inverse. C'est généralement la stratégie adoptée pour la réduction avec le sinus cardinal, car son support infini et sa décroissance lente empêchent une implémentation exacte dans le domaine spatial.

De nombreux auteurs ont proposé des filtres pour la réduction, que ce soit en domaine spatial ou fréquentiel (filtres de Burt [48], de Meer, ...). Nous renvoyons à [49] pour un panorama de leurs propriétés.

Puisque la réduction est un processus où une partie de l'information contenue dans l'image initiale est perdue, une classe de méthodes de réduction proposées dans la littérature vise à minimiser cette perte d'information. On peut mentionner les deux approches suivantes (on suppose que l'image initiale v a une résolution de $1/T$).

- Unser et coll. associent implicitement à toute image un modèle défini continûment, par interpolation spline. Le redimensionnement consiste alors à chercher l'image qui minimise la perte au sens L_2 , entre les modèles associés aux images initiale et redimensionnée [222, 223, 139]. Ainsi, la fonction spline $f_T(\mathbf{x}) \in V_T(\beta^n)$ interpolant l'image initiale est projetée orthogonalement dans l'espace spline $V_{\alpha T}(\beta^n)$, ce qui fournit le modèle $f_{\alpha T}(\mathbf{x})$ dont les échantillons $f_{\alpha T}(\alpha T\mathbf{k})$ sur le treillis $\Lambda_{\alpha T,0}$ forment l'image redimensionnée. On remarque que cette méthode part d'une formalisation inverse du problème de réduction : on cherche l'image réduite qui, interpolée, donne une fonction $f_{\alpha T}$ la plus proche possible de l'image initiale, au travers de son modèle f_T . Cette approche ne repose pas sur un modèle de

l'opération de redimensionnement. De plus, la modélisation par interpolation spline (dont le degré est un paramètre libre) est arbitraire et criticable.

- Une seconde approche, dans le même esprit, consiste à chercher l'image réduite à la résolution $1/(\alpha T)$ qui, interpolée, donne une fonction $f_{\alpha T} \in V_{\alpha T}(\varphi)$ la plus proche possible de l'image initiale, au sens ℓ_2 cette fois-ci [3]. On cherche donc la fonction $f_{\alpha T}$ minimisant le critère des moindres carrés $\sum_{\mathbf{k}} |f_{\alpha T}(T\mathbf{k}) - v[\mathbf{k}]|^2$. Il s'agit en fait d'un cas particulier du problème de modélisation que nous avons traité au chapitre 4 : on cherche le meilleur modèle pour l'image initiale, mais ayant la résolution de l'image réduite. Notons que cette méthode ne vérifie pas la propriété de hiérarchie : deux réductions successives de facteur 2 ne sont pas équivalentes à une réduction de facteur 4. De plus, cette approche non plus n'est pas cohérente avec notre modèle de redimensionnement. Par contre, elle est pertinente si le but est effectivement de minimiser la distorsion introduite par la réduction. On pense en particulier à une application de compression d'images, reposant sur une réduction pour la compression, et un agrandissement à la décompression, ce dernier étant fixé et s'effectuant par interpolation. Cette problématique n'est pas l'objet de notre étude.

En fait, minimiser un critère de perte d'information revient à poser la réduction comme un problème pseudo-inverse de l'interpolation. Le modèle que nous adoptons pour la réduction ne vise pas le même but : nous cherchons à d'obtenir une image réduite qui soit réaliste et satisfaisante visuellement, c'est-à-dire proche de ce qu'aurait fourni un bon dispositif réel d'acquisition. Nous ne cherchons pas une image réduite qui, si on la ré-agrandit ensuite de manière connue, soit la plus proche possible de l'image initiale. Ce problème est différent du problème de réduction proprement dit, puisqu'il faut alors optimiser le processus de réduction en tenant compte de la méthode d'agrandissement. Ce n'est pas, pour nous, l'objectif de la réduction en tant que telle, que nous abordons indépendamment de tout *a priori* sur l'utilisation de l'image réduite.

Dans la suite de cette section, nous présentons de nouvelles méthodes linéaires de réduction, que nous dérivons directement de notre modèle de redimensionnement.

5.3.1 Réduction de facteur entier

Nous nous intéressons en premier lieu à la réduction de facteur α entier. Ce cas est le plus simple, car la réduction se fait, de manière linéaire, à l'aide d'un filtre discret ($h_{\alpha, \tau}[\mathbf{k}]$), et il n'est pas nécessaire de disposer d'un filtre $h_{\alpha, \tau}(\mathbf{x})$ défini continûment. On a alors

$$\mathcal{R}_{\alpha, \tau} v = [v * h_{\alpha, \tau}] \downarrow \alpha. \quad (5.24)$$

En utilisant le modèle de formation de v (5.2), on peut écrire

$$\mathcal{R}_{\alpha, \tau} v[\mathbf{k}] = \int_{\mathbb{R}^2} s(\mathbf{x}) \tilde{\varphi}_{\text{eq}}\left(\mathbf{k} + \frac{\boldsymbol{\tau} + \boldsymbol{\tau}_0}{\alpha} - \frac{\mathbf{x}}{\alpha T}\right) \frac{d\mathbf{x}}{(\alpha T)^2} \quad \forall \mathbf{k} \in \mathbb{Z}^2, \quad (5.25)$$

où

$$\tilde{\varphi}_{\text{eq}}(\mathbf{x}) = \alpha^2 \sum_{\mathbf{k} \in \mathbb{Z}^2} h_{\alpha, \tau}[\mathbf{k}] \tilde{\varphi}(\alpha \mathbf{x} - \mathbf{k} - \boldsymbol{\tau}). \quad (5.26)$$

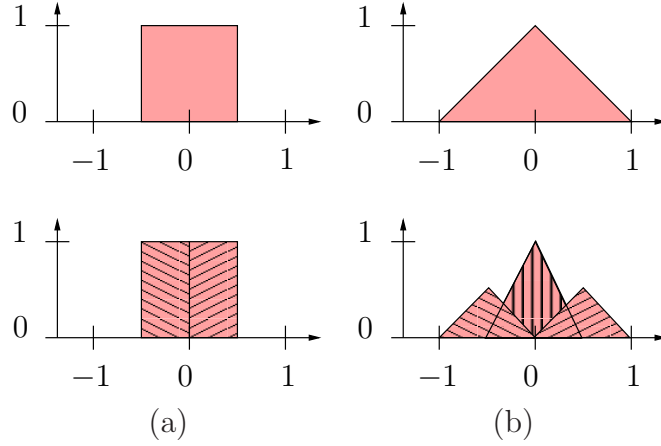


FIG. 5.4 : Les B-splines centrées 1D $\beta^n(t)$ vérifient toutes une relation à deux échelles généralisée de facteur $\alpha = 2$: elles s'écrivent comme une somme pondérée des translatés sur le treillis $\Lambda_{\frac{1}{\alpha}, \tau}$ de leurs α -contractées. En (a), le filtre d'échelle de β^0 est $h = [\frac{1}{2}, \frac{1}{2}, 0]$, avec $\tau = \frac{1}{2}$, alors qu'en (b), le filtre associé à β^1 est $h = [\frac{1}{4}, \frac{1}{2}, \frac{1}{4}]$, avec $\tau = 0$.

Par conséquent, on a exactement l'image réduite désirée ($\mathcal{R}_{\alpha, \tau} v = \mathcal{R}_{\alpha, \tau}^{\text{id}} v$) si et seulement si $\tilde{\varphi}$ vérifie une *relation à deux échelles généralisée* :

$$\tilde{\varphi}(\mathbf{x}) = \alpha^2 \sum_{\mathbf{k} \in \mathbb{Z}^2} h_{\alpha, \tau}[\mathbf{k}] \tilde{\varphi}(\alpha \mathbf{x} - \mathbf{k} - \boldsymbol{\tau}). \quad (5.27)$$

On a donc la propriété suivante :

Si $\tilde{\varphi}$ vérifie une relation à deux échelles généralisée de facteur α entier et d'offset $\boldsymbol{\tau}$, on peut effectuer la réduction $\mathcal{R}_{\alpha, \tau}$ de manière parfaite d'après notre modèle, au moyen du processus de filtrage décimation (5.24), avec pour filtre de réduction $h_{\alpha, \tau}$ le filtre d'échelle de $\tilde{\varphi}$.

Les B-splines et splines cardinales séparables vérifient toutes une relation à deux échelles généralisée, et ce pour tout α entier. La figure 5.4 illustre cette propriété pour les B-splines 1D de degré 0 et 1. Notons que $\boldsymbol{\tau} = \mathbf{0}$ pour les splines de degré impair, et $\boldsymbol{\tau} = (\frac{1}{2}, \frac{1}{2})$ pour les splines de degré pair. Une famille récente de splines, appelées les « $\alpha - \tau$ splines » [38] (α et τ n'ont pas la même signification que dans ce document), si on les centre en $\mathbf{0}$, vérifient aussi une relation à deux échelles généralisée pour tout α entier, avec un offset $\boldsymbol{\tau}$ qui les caractérise.

Ainsi, dans l'hypothèse où $\tilde{\varphi} = \gamma^3$ est la spline cardinale cubique, on peut effectuer la réduction de facteur α entier et d'offset $\boldsymbol{\tau} = \mathbf{0}$ à l'aide de son filtre d'échelle

$$h_{\alpha, \mathbf{0}}[\mathbf{k}] = \frac{1}{\alpha^2} \gamma^3\left(\frac{\mathbf{k}}{\alpha}\right) \Leftrightarrow h_{\alpha, \mathbf{0}} = \frac{1}{\alpha^2} b_{\alpha}^3 * [(b_1^3)^{-1}] \uparrow \alpha \xleftrightarrow{\mathcal{Z}} H_{\alpha, \mathbf{0}}(\mathbf{z}) = \frac{1}{\alpha^2} \frac{B_{\alpha}^3(\mathbf{z})}{B_1^3(\mathbf{z}^{\alpha})}, \quad (5.28)$$

où l'on définit pour tous entiers $n \geq 1, m$, le filtre B-spline discret (ou *B-spline dilatée discrète*)

$$b_m^n[k] = \beta^n\left(\frac{k}{m}\right) \quad \forall k \in \mathbb{Z}, \quad (5.29)$$

ainsi que le filtre séparable 2D $(b_m^n[\mathbf{k}])_{\mathbf{k} \in \mathbb{Z}^2}$ par produit tensoriel de $(b_m^n[k])_{k \in \mathbb{Z}}$.

La réduction s'implémente alors efficacement en deux étapes :

$$\mathcal{R}_{\alpha, \mathbf{0}} v = \frac{1}{\alpha^2} [v * b_\alpha^3] \downarrow \alpha * (b_1^3)^{-1}. \quad (5.30)$$

Remarquons que (5.30) et (5.7) sont équivalentes, en posant dans (5.7)

$$h_{\alpha, \tau}(\mathbf{x}) = \frac{1}{\alpha^2} \gamma^3 \left(\frac{\mathbf{x}}{\alpha} \right). \quad (5.31)$$

Enfin, on peut noter qu'on a, avec cette méthode de réduction, la propriété hiérarchique

$$\mathcal{R}_{\alpha_1, \mathbf{0}} \circ \mathcal{R}_{\alpha_2, \mathbf{0}} = \mathcal{R}_{\alpha_1 \alpha_2, \mathbf{0}} \quad \forall \alpha_1, \alpha_2 \in \mathbb{N} \setminus \{0\}. \quad (5.32)$$

Si maintenant on souhaite effectuer une réduction de facteur α et d'offset τ sans que $\tilde{\varphi}$ ne vérifie une relation à deux échelles avec ces paramètres, on ne peut plus trouver de filtre $h_{\alpha, \tau}$ assurant $\mathcal{R}_{\alpha, \tau} v = \mathcal{R}_{\alpha, \tau}^{\text{id}} v$. On va donc plutôt chercher le filtre qui, s'il n'annule pas, du moins minimise, dans un cadre stochastique, l'erreur

$$\varepsilon_{\text{RED}} = \mathcal{E} \{ |\mathcal{R}_{\alpha, \tau}^{\text{id}} v[\mathbf{k}] - \mathcal{R}_{\alpha, \tau} v[\mathbf{k}]|^2 \} \quad (5.33)$$

$$= \mathcal{E} \left\{ \left| s(T(\cdot + \tau_0)) * \tilde{\phi}(\alpha \mathbf{k}) \right|^2 \right\}, \quad (5.34)$$

qui ne dépend ni de $\mathbf{k}$, ni de τ_0 , et où on a posé

$$\tilde{\phi}(\mathbf{x}) = \frac{1}{\alpha^2} \tilde{\varphi} \left(\frac{\mathbf{x} + \tau}{\alpha} \right) - \sum_{\mathbf{k} \in \mathbb{Z}^2} h_{\alpha, \tau}[\mathbf{k}] \tilde{\varphi}(\mathbf{x} - \mathbf{k}). \quad (5.35)$$

En domaine de Fourier, on obtient :

$$\varepsilon_{\text{RED}} = \frac{1}{(2\pi)^2} \int_{\mathbb{R}^2} \frac{1}{T^2} \hat{c}_s \left(\frac{\omega}{T} \right) \left| \hat{\tilde{\varphi}}(\omega) \hat{h}_{\alpha, \tau}(\omega) - \hat{\tilde{\varphi}}(\alpha \omega) e^{I\langle \omega, \tau \rangle} \right|^2 d\omega \quad (5.36)$$

$$= \frac{1}{(2\pi)^2} \int_{\mathbb{R}^2} \hat{c}_s(\omega) E(T\omega) d\omega, \quad (5.37)$$

où nous définissons le *noyau d'erreur fréquentiel de réduction entière* $E(\omega)$ par

$$E(\omega) = \left| \hat{\tilde{\varphi}}(\omega) \hat{h}_{\alpha, \tau}(\omega) - \hat{\tilde{\varphi}}(\alpha \omega) e^{I\langle \omega, \tau \rangle} \right|^2. \quad (5.38)$$

• En utilisant ces formules, on peut trouver le filtre $h_{\alpha, \tau}^{\text{id}}$ minimisant l'erreur ε_{RED} . En effet, on peut écrire

$$\varepsilon_{\text{RED}} = \frac{1}{(2\pi)^2} \int_{[0, 2\pi]^2} \sum_{\mathbf{k} \in \mathbb{Z}^2} \frac{1}{T^2} \hat{c}_s \left(\frac{\omega + 2\pi \mathbf{k}}{T} \right) |\hat{\tilde{\varphi}}(\omega + 2\pi \mathbf{k})|^2 \left| \hat{h}_{\alpha, \tau}(\omega) - \hat{h}_{\alpha, \tau}^{\text{id}}(\omega) \right|^2 + C(\omega) d\omega, \quad (5.39)$$

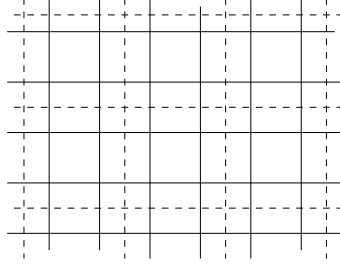


FIG. 5.5 : Le problème de réduction de facteur 2 d'offset $\mathbf{0}$ illustré dans le cas où $\tilde{\varphi} = \mathbb{1}_{]-\frac{1}{2}, \frac{1}{2}[}^2$: connaissant les moyennes dans les carrés solides de la fonction inconnue $s(\mathbf{x})$, on cherche les moyennes de s dans les carrés en pointillés. Ce cas est représentatif de l'acquisition d'un signal à l'aide d'un capteur CCD idéal, dont chaque cellule réalise l'intégration lumineuse sur sa surface carrée. Notons que les sites des treillis sont identifiés aux centres des carrés.

où $C(\omega)$ est un terme constant indépendant de $h_{\alpha, \tau}$, et

$$\hat{h}_{\alpha, \tau}^{\text{id}}(\omega) = \frac{\sum_{\mathbf{k} \in \mathbb{Z}^2} \frac{1}{T^2} \hat{c}_s\left(\frac{\omega + 2\pi \mathbf{k}}{T}\right) \hat{\varphi}(\omega + 2\pi \mathbf{k})^* \hat{\varphi}(\alpha(\omega + 2\pi \mathbf{k})) e^{I(\omega, \tau)}}{\sum_{\mathbf{k} \in \mathbb{Z}^2} \frac{1}{T^2} \hat{c}_s\left(\frac{\omega + 2\pi \mathbf{k}}{T}\right) |\hat{\varphi}(\omega + 2\pi \mathbf{k})|^2}. \quad (5.40)$$

• Ce filtre optimal dépendant de la densité spectrale de s , qui est généralement inconnue, on peut adopter une approche similaire à celle que nous avons développée pour la reconstruction : on cherche un filtre assurant une réduction de bonne qualité lorsque s a son énergie localisée dans les basses fréquences, ce qui revient à dire que les basses fréquences de $\mathcal{R}_{\alpha, \tau} v$ seront celles de $\mathcal{R}_{\alpha, \tau}^{\text{id}} v$. Autrement dit, on cherche h tel que $E(\omega)$ soit le plus plat possible à l'origine, c'est-à-dire

$$\hat{h}(\omega) = \frac{\hat{\varphi}(\alpha\omega) e^{I(\omega, \tau)}}{\hat{\varphi}(\omega)} + O(\|\omega\|^N), \quad (5.41)$$

avec N le plus grand possible. Remarquons aussi que cette dernière propriété équivaut à la propriété de convergence

$$\varepsilon_{\text{RED}} = O(T^N), \quad T \rightarrow 0. \quad (5.42)$$

Prenons l'exemple concret suivant, schématisé sur la figure 5.5 : $\tilde{\varphi} = \beta^0$ est la fonction porte, et on cherche à effectuer au mieux la réduction de facteur 2 et d'offset $\mathbf{0}$. Cherchons un filtre $h_{2, \mathbf{0}}$ séparable symétrique de taille 3, de la forme

$$H_{2, \mathbf{0}}(\mathbf{z}) = (a + b(z_1 + z_1^{-1}))(a + b(z_2 + z_2^{-1})). \quad (5.43)$$

Ses deux degrés de liberté a et b peuvent être choisis pour que (5.41) soit vérifiée avec $N = 4$, puisque (5.41) se ramène dans ce cas à

$$a + 2b - b\omega^2 + O(\omega^4) = 1 - \frac{1}{8}\omega^2 + O(\omega^4), \quad (5.44)$$

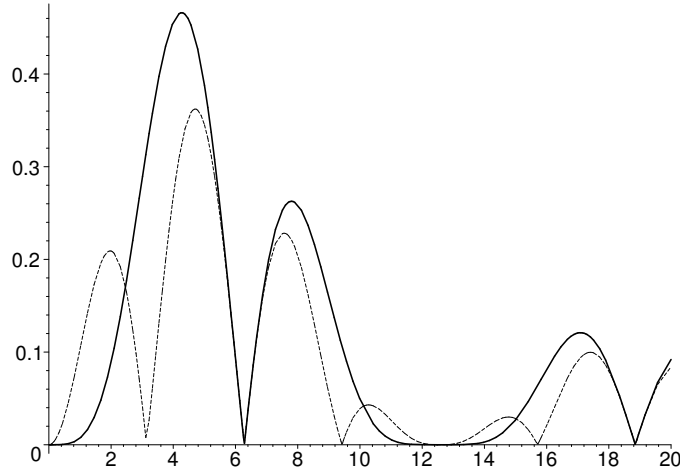


FIG. 5.6 : Noyau d'erreur de réduction $\sqrt{E(\omega)}$ associé aux filtres $h_{2,0} = [\frac{1}{8}, \frac{3}{4}, \frac{1}{8}]$ (trait plein) et $h_{2,0}^{\text{Price}} = [\frac{1}{4}, \frac{1}{2}, \frac{1}{4}]$ (pointillés) pour la réduction de facteur 2 d'offset 0, lorsque $\tilde{\varphi} = \beta^0$ (dans le cas 1D). Notre filtre est plus plat à l'origine que celui de Price et coll. [179], et sera donc meilleur pour la réduction d'images.

d'où l'on identifie $a = \frac{3}{4}$, $b = \frac{1}{8}$. On peut comparer ce filtre à celui proposé par Price et coll. dans [179] pour résoudre le même problème. Les auteurs proposent de choisir le filtre assurant $\mathcal{R}_{\alpha,\tau}v = \mathcal{R}_{\alpha,\tau}^{\text{id}}v$ dans le cas particulier où $s \in V_T(\beta^0)$, c'est-à-dire est constante par morceaux. On voit aisément d'après la figure 5.5 que si $s(\mathbf{x})$ est constante dans chaque carré solide, on obtient ses moyennes dans les carrés pointillés en prenant le filtre

$$H_{2,0}^{\text{Price}}(\mathbf{z}) = \left(\frac{1}{2} + \frac{1}{4}(z_1 + z_1^{-1})\right)\left(\frac{1}{2} + \frac{1}{4}(z_2 + z_2^{-1})\right). \quad (5.45)$$

En comparant les noyaux d'erreur des deux filtres (figure 5.6), on voit que notre filtre est meilleur que $h_{2,0}^{\text{Price}}$ au voisinage de l'origine.

Dans le cas où $\tilde{\varphi}$ est la spline cardinale cubique, plutôt que de concevoir un filtre optimal pour chaque τ , la méthode que nous conseillons pour la réduction de facteur entier α et d'offset τ consiste simplement à utiliser pour tout τ la formule (5.7) avec la fonction spline cardinale cubique dilatée donnée par (5.31). Cette méthode n'est optimale que pour $\tau = \mathbf{0}$, mais elle donne entière satisfaction en pratique. L'implémentation correspondante est la suivante :

$$\mathcal{R}_{\alpha,\tau}v = c * (b_1^3)^{-1}, \quad (5.46)$$

où b_1^3 , donné par (5.29), est le filtre $[\frac{1}{6}, \frac{2}{3}, \frac{1}{6}]^2$, et où

$$c[\mathbf{k}] = \frac{1}{\alpha^2} \sum_{\mathbf{l} \in \mathbb{Z}^2} v[\mathbf{l}] \beta^3\left(\mathbf{k} + \frac{\tau - \mathbf{l}}{\alpha}\right). \quad (5.47)$$

5.3.2 Réduction de facteurs non entiers

Pour effectuer la réduction avec un facteur α non entier, et pour un offset $\boldsymbol{\tau}$ quelconque, il faut disposer d'une fonction $h_{\alpha,\boldsymbol{\tau}}(\mathbf{x})$, afin d'appliquer la formule (5.7). La réduction consiste donc formellement à construire la fonction

$$g_{\alpha,\boldsymbol{\tau}}(\mathbf{x}) = \sum_{\mathbf{k} \in \mathbb{Z}^2} v[\mathbf{k}] h_{\alpha,\boldsymbol{\tau}}(\mathbf{x} - \mathbf{k}), \quad (5.48)$$

avant de l'échantillonner :

$$\mathcal{R}_{\alpha,\boldsymbol{\tau}} v[\mathbf{k}] = g_{\alpha,\boldsymbol{\tau}}(\alpha \mathbf{k} + \boldsymbol{\tau}). \quad (5.49)$$

Notons qu'il n'y a aucune raison pour que $h_{\alpha,\boldsymbol{\tau}}$ et $g_{\alpha,\boldsymbol{\tau}}$ dépende de $\boldsymbol{\tau}$, puisque ce paramètre n'intervient que lors de l'échantillonnage de g , et pas dans sa définition.

Une première méthode pour choisir $h_{\alpha,\boldsymbol{\tau}}$ consiste à étendre la méthode utilisée pour les facteurs entiers en choisissant la fonction (5.31). Ce n'est pas une bonne idée : après dilatation, γ^3 perd ses propriétés d'approximation. Par exemple, si v est un signal constant, $g_{\alpha,\boldsymbol{\tau}}(\mathbf{x})$ n'est pas une fonction constante, et des oscillations néfastes apparaîtront alors dans l'image réduite.

La solution pratique que nous adoptons consiste à renormaliser localement la fonction $\beta^3(\frac{\cdot}{\alpha})$ afin de forcer la partition de l'unité dans l'implémentation (5.46),(5.47). On calcule donc

$$c[\mathbf{k}] = \frac{\sum_{\mathbf{l} \in \mathbb{Z}^2} v[\mathbf{l}] \beta^3(\mathbf{k} + \frac{\boldsymbol{\tau} - \mathbf{l}}{\alpha})}{\sum_{\mathbf{l} \in \mathbb{Z}^2} \beta^3(\mathbf{k} + \frac{\boldsymbol{\tau} - \mathbf{l}}{\alpha})}, \quad (5.50)$$

avant d'obtenir l'image réduite par le post-filtrage

$$\mathcal{R}_{\alpha,\boldsymbol{\tau}} v = c * (b_1^3)^{-1}. \quad (5.51)$$

Notons que cette méthode est bien une généralisation de celle employée dans la situation où α est entier. En effet, puisque les B-splines vérifient la partition de l'unité, on a (uniquement si α est entier) :

$$\sum_{\mathbf{l} \in \mathbb{Z}^2} \beta^3(\mathbf{k} + \frac{\boldsymbol{\tau} - \mathbf{l}}{\alpha}) = \alpha^2 \quad (5.52)$$

et (5.50) est alors équivalente à (5.47).

La figure 5.7 montre des images réduites avec notre méthode, pour différents facteurs et offsets. De notre point de vue, la qualité de cette méthode est excellente. Le flou, l'*aliasing* et le *ringing* sont tous les trois maîtrisés. Ce compromis très avantageux est à créditer aux propriétés de la fonction spline cardinale cubique, proche du sinus cardinal, mais avec une transition fréquentielle douce à la fréquence de Nyquist.

5.4 Conclusion

Dans ce chapitre, nous avons proposé un nouveau formalisme pour le redimensionnement : l'image redimensionnée doit être la plus proche possible de celle qui aurait été

produite par le dispositif d'acquisition ayant engendré l'image initiale, si la scène lumineuse avait été contractée ou dilatée.

Cela nous a permis de proposer dans un premier temps des méthodes originales et efficaces pour la réduction d'images de facteur entier, reposant sur la préservation asymptotique des basses fréquences. Nous avons aussi défini un noyau d'erreur de réduction, permettant de visualiser à chaque fréquence l'erreur relative introduite par un filtre de réduction donné.

Nous avons ensuite étendu notre méthode de réduction dans le cas général de facteurs non entiers. Bien que l'optimalité théorique soit perdue, les images réduites sont, en pratique, de très bonne qualité. La formulation de solutions optimales dans le cas de facteurs non entiers est une voie de développement qui reste à explorer.

Dans le chapitre suivant, nous allons appliquer notre formalisme au problème de l'agrandissement d'images. Celui-ci est de nature différente de la réduction : au contraire de cette dernière qui supprime de l'information pour obtenir une représentation plus grossière, l'agrandissement nécessite l'ajout d'information dans les hautes fréquences, dont l'extrapolation à partir de l'image initiale est la clé pour l'obtention d'images de bonne qualité.



(a)

(b) $\alpha = \pi$, $\tau = (0.0, 0.0)$ (c) $\alpha = 5$, $\tau = (0.0, 0.0)$ (d) $\alpha = 9.3$, $\tau = (0.8, 0.8)$

FIG. 5.7 : Image *Barbara* (a) réduite avec la fonction spline cardinale cubique, avec différents facteurs et offsets. (Figure continuée sur la page suivante)

(e) $\alpha = 1.1$, $\tau = (0.5, 0.5)$ (f) $\alpha = 1.9$, $\tau = (0.2, 0.2)$

FIG. 5.7 : Suite.

Chapitre 6

Agrandissement d'images

6.1 Introduction

LE problème de l'agrandissement occupe une place importante dans l'univers du traitement d'images, car il est à la fois difficile et essentiel. L'essor des applications multimédias implique en effet un besoin croissant en images et vidéos de grandes tailles et de bonne qualité. Or, les dispositifs d'acquisition d'images ont des limites physiques, et les systèmes ayant une haute résolution ont une sensibilité moindre (la densité surfacique de photons reçus diminuant), et sont donc plus sujets aux problèmes de bruit. En acquisition vidéo, la vitesse de capture impose aussi une résolution limitée, et les caméras haute résolution sont lourdes et encombrantes, sans parler de leur coût prohibitif. En imagerie satellitaire, ce sont les perturbations de l'air et le mouvement du satellite qui imposent actuellement une limite inférieure à la résolution d'environ 1pixel/60cm. De plus, les capacités de stockage et les débits de transmission restent limités, et obligent à manipuler des images de résolution plus faible que souhaitée, souvent parce qu'elles ont été précédemment réduites en même temps que compressées.

L'agrandissement, qui vise à augmenter artificiellement la résolution d'une image, permet de pallier ces contraintes. L'augmentation de taille et de résolution des périphériques d'affichage et d'impression rendent l'agrandissement nécessaire, par exemple pour la conversion de vidéos au format TV vers le format émergent TVHD, ou pour l'impression d'affiches publicitaires de très grande taille. L'agrandissement permet aussi de scruter une partie d'image, afin de pouvoir analyser et comprendre le sens d'un détail, avec un confort visuel et une acuité meilleurs. Les enjeux sont importants, en particulier dans le cas où une image médicale agrandie sert de support à un praticien pour le guider dans son diagnostic.

Le problème de l'agrandissement est essentiellement mal posé : dans le but d'augmenter la résolution d'une image, on ne peut en pratique qu'analyser l'information contenue dans cette image, pour la synthétiser au mieux à une résolution supérieure. Il faut donc interpréter les images à un niveau sémantique suffisant pour agrandir de manière cohérente et visuellement satisfaisante les contours des objets, qui sont porteurs de l'information

géométrique, à laquelle nous sommes le plus sensibles.

6.1.1 Contexte de l'étude

L'agrandissement vise avant tout à obtenir une image agrandie de bonne qualité visuelle. Etant donné que l'on ne dispose pas de référence avec laquelle comparer l'image agrandie, l'estimation quantitative de la qualité d'une méthode d'agrandissement donnée est problématique. L'évaluation de la qualité d'une image redimensionnée, et d'une image en règle générale, est une tâche très difficile [239]. L'inspection visuelle reste en fin de compte le meilleur juge de qualité.

Le problème de l'agrandissement est tributaire de la nature de l'image initiale. On s'intéressera aux images naturelles non bruitées, c'est-à-dire aux images acquises à partir d'une scène lumineuse à l'aide d'un dispositif d'acquisition linéaire, comme formalisé par (5.2). Cela inclut certains types d'images médicales ou satellitaires, et pas seulement les photographies numériques. Au contraire, les images synthétiques ou les images de texte contiennent des contours abrupts de nature particulière, qui nécessitent des méthodes spécifiques. On ne traitera pas de l'agrandissement de séquences vidéos, car dans ce domaine, la qualité est déterminée par l'exploitation plus ou moins poussée de la redondance temporelle entre images successives. Ce problème est similaire au problème d'agrandissement de plusieurs images légèrement translatées l'une par rapport à l'autre, pour obtenir une seule image agrandie fusionnant l'information de toutes les petites images. Dans ces situations, on ne parle pas d'agrandissement mais de *super-résolution* [41]. La particularité de ce problème est qu'il existe une source d'information supplémentaire, provenant de la redondance spatiale ou temporelle, que l'on peut exploiter pour améliorer significativement la qualité des images agrandies.

Les artéfacts les plus courants pouvant apparaître dans les images agrandies sont les suivants. Ce sont les mêmes que ceux qui se manifestent après la réduction, ou la formation d'images : ils sont intrinsèques à tout processus fournissant une image discrète.

- Apparition d'oscillations (*ringing*), au voisinage des contours. Un contour, dans le domaine continu, n'est pas à bande limitée; par conséquent, en le représentant sur un treillis discret, on tronque la partie hautes fréquences de son énergie, ce qui fait apparaître des oscillations dans l'image. Cet effet est très visible sur la figure 6.1 (a), et dans une moindre mesure sur la figure 6.1 (d).
- Effet de flou : si la méthode d'agrandissement ne crée pas de hautes fréquences et/ou atténue le contenu fréquentiel de l'image initiale, un effet de flou global apparaît. Le système visuel humain est très sensible au flou. Il préfère la présence de ringing, repliement de spectre, ou bruit, plutôt qu'un effet de flou. On peut voir cet effet sur la figure 6.1 (c).
- Repliement de spectre (*aliasing*) : cet effet se manifeste lorsque le signal qui a été échantillonné pour former l'image contient de l'énergie au-delà de la bande de Nyquist associée au treillis de l'image. Ces hautes fréquences se « replient » sous forme de basses

fréquences, et de fausses structures apparaissent dans l'image. Dans le cas de l'agrandissement, cet effet apparaît essentiellement sous forme d'un effet « d'escaliers » le long des contours obliques, par exemple sur le pied de la caméra des figures 6.1 (b) et (c). Etant localisé, de même que le ringing, cet effet est moins gênant que le flou, hormis dans le cas extrême de l'interpolation au plus proche voisin (figure 6.1 (b)).

D'autres artéfacts peuvent apparaître, ou être renforcés, à cause de traitements faits sur l'image, et non plus à cause de la nature discrète de l'image elle-même. Les distorsions les plus fréquentes sont la présence de bruit, d'effets de blocs dus à la compression JPEG, ou d'un effet d'« à plat » (*painting effect, clustering*) si les contours sont bien rendus mais que l'image est floue entre ceux-ci.

Jusqu'à maintenant, nous n'avons parlé que de signaux et images à une composante. En pratique, on est plus souvent amené à manipuler des images couleur à trois composantes, représentant l'intensité lumineuse dans les bandes centrées autour du rouge, vert, et bleu. Le système visuel humain n'est pas sensible à ces informations de manière uniforme, et l'information de luminance, sorte de moyenne des trois composantes, est beaucoup plus importante que l'information de chrominance, qui caractérise les différences d'amplitude entre les trois canaux. Par conséquent, on s'attachera particulièrement à l'agrandissement de l'information de luminance, ce qui ramène le problème à l'agrandissement des images en niveaux de gris. L'agrandissement de la chrominance sera réalisé au moyen d'une des méthodes linéaires, simples et rapides, présentées à la section 6.2, la différence de qualité occasionnée par l'emploi d'une méthode plus complexe n'étant pas perceptible par un observateur humain.

6.1.2 Objectifs et plan du chapitre

Le problème de l'agrandissement, tel que nous l'avons formulé à la section 5.2, est le suivant : étant donnée l'image initiale v acquise selon le modèle (5.2), on cherche à estimer l'image agrandie idéale $\mathcal{A}_{\alpha, \tau}^{\text{id}} v$ définie par

$$\mathcal{A}_{\alpha, \tau}^{\text{id}} v[\mathbf{k}] = \int_{\mathbb{R}^2} s(\mathbf{x}) \tilde{\varphi}(\mathbf{k} + \alpha(\boldsymbol{\tau} + \boldsymbol{\tau}_0) - \frac{\alpha \mathbf{x}}{T}) \frac{\alpha^2 d\mathbf{x}}{T^2} \quad \forall \mathbf{k} \in \mathbb{Z}^2. \quad (6.1)$$

Cette manière de formuler l'objectif de l'agrandissement est nouvelle. Elle est générique, puisque la fonction $\tilde{\varphi}$ peut être choisie arbitrairement, en fonction des connaissances *a priori* dont on dispose. Ainsi, le problème d'interpolation ($\tilde{\varphi} = \delta$) à une résolution supérieure est un cas particulier d'agrandissement. Pour l'agrandissement des images naturelles, et en l'absence de connaissances supplémentaires, nous faisons l'hypothèse, vraisemblable *a posteriori*, que la réponse $\tilde{\varphi}$ peut être bien modélisée par la spline cardinale γ^3 .

Ainsi, notre modèle d'agrandissement diffère de celui adopté par Aly [16, 13], qui cherche l'image agrandie s'affichant le mieux sur un périphérique donné, de réponse impulsionnelle φ connue. Il s'agit alors d'un problème de reconstruction, que l'on peut traiter en suivant

les principes exposés dans la première partie de cette thèse. Pour nous, l'agrandissement doit être dissocié de l'affichage proprement dit, puisqu'on ne peut espérer avoir une image agrandie ayant un rendu satisfaisant sur tous les périphériques d'affichage. Cependant, cet ajustement peut être réalisé après coup au moyen d'un simple filtrage linéaire.

Dans ce chapitre, nous présentons d'abord dans la section 6.2 les méthodes linéaires d'agrandissement existantes, ainsi que de nouvelles méthodes linéaires plus appropriées. Ces méthodes sont globales, car l'image est traitée identiquement en tout point, et elles ne tiennent donc pas compte des non-stationnarités des images. Or une image est composée de zones homogènes, de textures, et de contours. Une méthode globale ne peut pas tirer parti des spécificités de ces zones pour les agrandir de manière appropriée.

Dans les sections 6.3 à 6.5, nous présentons un panorama des différentes classes de méthodes non linéaires proposées dans la littérature. Nous verrons que les méthodes d'agrandissement les plus satisfaisantes sont basées sur un traitement approprié des contours. Pour chaque méthode, les images servant d'illustration sont extraites de l'article donné en référence, sauf mention contraire.

Enfin, dans la section 6.6, nous proposons une nouvelle méthode d'interpolation non linéaire particulièrement performante.

6.2 Méthodes linéaires d'agrandissement

Si l'on impose au processus d'agrandissement d'être linéaire et invariant par translation, alors il existe une certaine fonction $\varphi_{\alpha,\tau}^{\text{eq}}(\mathbf{x})$ telle que

$$\mathcal{A}_{\alpha,\tau} v[\mathbf{k}] = \sum_{\mathbf{l} \in \mathbb{Z}^2} \varphi_{\alpha,\tau}^{\text{eq}}\left(\frac{\mathbf{k}}{\alpha} + \boldsymbol{\tau} - \mathbf{l}\right) v[\mathbf{l}] \quad \forall \mathbf{k} \in \mathbb{Z}^2. \quad (6.2)$$

Formellement, l'agrandissement est donc équivalent à reconstruire la fonction

$$g_{\alpha,\tau}(\mathbf{x}) = \sum_{\mathbf{k} \in \mathbb{Z}^2} v[\mathbf{k}] \varphi_{\alpha,\tau}^{\text{eq}}\left(\frac{\mathbf{x}}{T} - \boldsymbol{\tau}_0 - \mathbf{k}\right) \quad (6.3)$$

avant de l'échantillonner sur le treillis $\Lambda_{\frac{T}{\alpha}, \alpha(\boldsymbol{\tau}_0 + \boldsymbol{\tau})}$. Notons qu'il n'y a pas de raison pour que $\varphi_{\alpha,\tau}^{\text{eq}}$ et $g_{\alpha,\tau}$ dépendent de $\boldsymbol{\tau}$, puisque ce paramètre n'intervient que dans la localisation des échantillons. On note donc dans la suite $\varphi_{\alpha}^{\text{eq}}$ et g_{α} , pour indiquer que ces fonctions sont indépendantes de $\boldsymbol{\tau}$.

On voit donc que l'on peut formuler l'agrandissement linéaire comme un problème de reconstruction, suivi d'un échantillonnage. La différence majeure est que l'on cherche, connaissant des mesures linéaires uniformes sur s , non pas s elle-même, mais $s * \frac{\alpha^2}{T^2} \tilde{\varphi}(\frac{\alpha \cdot}{T})$.

Dans le cas d'un facteur α entier, seules les valeurs ponctuelles $h_{\alpha,\tau}[\mathbf{k}] = \varphi_{\alpha}^{\text{eq}}(\frac{\mathbf{k}}{\alpha} + \boldsymbol{\tau})$ entrent en compte dans le calcul (6.2). L'agrandissement linéaire et invariant par translation

de facteur entier revient donc à définir le filtre d'agrandissement $h_{\alpha,\tau}[\mathbf{k}]$, tel que :

$$\mathcal{A}_{\alpha,\tau}v[\mathbf{k}] = \sum_{\mathbf{l} \in \mathbb{Z}^2} h_{\alpha,\tau}[\mathbf{k} - \alpha\mathbf{l}]v[\mathbf{l}] \quad \forall \mathbf{k} \quad \Leftrightarrow \quad \mathcal{A}_{\alpha,\tau} = [v] \uparrow \alpha * h_{\alpha,\tau}. \quad (6.4)$$

Remarquons que l'agrandissement classique de facteur 2 au plus proche voisin, au moyen du filtre $[0, \frac{1}{2}, \frac{1}{2}]^2$, correspond pour nous à un offset $\tau = (-\frac{1}{4}, -\frac{1}{4})$. Si on compare l'image agrandie avec ce filtre à celle agrandie bilinéairement avec le filtre $[\frac{1}{4}, \frac{1}{2}, \frac{1}{4}]^2$ (offset $\tau = \mathbf{0}$), en les affichant sur le même treillis, on observera que la première image est translatée d'un demi-pixel vers la droite et le haut par rapport à la seconde. C'est pourquoi il faut toujours garder à l'esprit le treillis sur lequel est localisée une image. En particulier, calculer une distance PSNR entre une image traitée et une image de référence n'a de sens que si elles sont définies sur le même treillis.

6.2.1 Interpolation

La méthode usuelle pour effectuer l'agrandissement d'une image consiste à utiliser une méthode linéaire d'interpolation. Cela correspond à considérer, dans nos notations, que $\tilde{\varphi} = \delta$. Ce modèle n'est pas correct : les systèmes d'acquisition n'effectuent pas un échantillonnage idéal sur la scène lumineuse s , et si cela était le cas, les images seraient polluées par des artéfacts de repliement de spectre, puisque la scène s n'est pas à bande limitée. On en déduit *a posteriori* que la fonction $\hat{\tilde{\varphi}}(\mathbf{x})$ réalise une forte atténuation en dehors de la bande de Nyquist.

L'autre situation qui puisse justifier un agrandissement par interpolation est le cas $\tilde{\varphi}(\mathbf{x}) = \text{sinc}(\pi\mathbf{x})$. Là aussi, ce modèle n'est pas réaliste, car alors les images seraient contaminées par de nombreuses oscillations. $\tilde{\varphi}$ présente donc aussi une atténuation à l'intérieur de la zone de Nyquist. Par conséquent, l'agrandissement devra corriger dans une certaine mesure cette atténuation, par un réhaussement du contenu fréquentiel de l'image initiale. Les méthodes d'interpolation, linéaires ou non, ne peuvent que reconstruire la fonction dont les échantillons forment l'image v , c'est-à-dire $s * \frac{1}{T^2}\tilde{\varphi}(\frac{\cdot}{T})$. Or l'image agrandie que l'on cherche consiste en des échantillons de la fonction $s * \frac{\alpha^2}{T^2}\tilde{\varphi}(\frac{\alpha\cdot}{T})$. **Il y a donc un problème de déconvolution derrière celui de l'agrandissement.**

En fait, le paradigme sous-jacent au choix habituel de $\varphi_{\alpha}^{\text{eq}}$ est la théorie de l'échantillonnage de Whittaker et Shannon, qui stipule qu'un signal continu peut être parfaitement reconstruit à partir de ses échantillons, si toutefois il a été échantillonné à une fréquence au moins deux fois supérieure à sa plus haute fréquence (fréquence de Nyquist). Ainsi, si s est à bande limitée dans $]-\frac{\pi}{T}, \frac{\pi}{T}[$ et que $\tilde{\varphi} = \delta$, l'agrandissement revient à effectuer la reconstruction de s , car $s * \frac{\alpha^2}{T^2}\tilde{\varphi}(\frac{\alpha\cdot}{T}) = s$. Cette reconstruction peut être réalisée exactement, au moyen de la fonction sinus cardinal $\varphi_{\alpha}^{\text{eq}}(\mathbf{x}) = \text{sinc}(\pi\mathbf{x})$.

Bien que le choix de cette fonction parte d'un principe louable de conservation du contenu fréquentiel de l'image initiale, il repose sur un modèle erroné. La scène lumineuse



FIG. 6.1 : Agrandissement $\times 4$ par interpolation (a) sinus cardinal (b) au plus proche (c) bilinéaire (d) spline cubique, d'une portion de l'image *Camera*.

a, certes, son énergie concentrée dans les basses fréquences, mais celle-ci ne s'annule pas en dehors d'une bande fréquentielle donnée. L'énergie dans les hautes fréquences correspond essentiellement à la présence de contours, qui sont localisés spatialement de manière précise. La non reconstruction de hautes fréquences entraîne donc un rendu insatisfaisant des contours, avec une transition transverse apparaissant floue.

Malgré ces objections, l'interpolation sinus cardinal est toujours présentée comme une méthode « idéale ». Si elle n'est pas utilisée, c'est parce que d'une part son implémentation n'est pas aisée, le sinus cardinal ayant un support infini, et d'autre part elle entraîne un effet de *ringing* très important (voir figure 6.1 (a)). Les méthodes usuelles consistent donc à choisir $\varphi_\alpha^{\text{eq}}$ interpolante (et indépendante de α), dans le but d'approcher la fonction sinus cardinal. Les fonctions classiques sont celles dont on a discuté à la section 2.3.1 : interpolation au plus proche voisin et bilinéaire avec $\varphi_\alpha^{\text{eq}} = \beta^0$ et β^1 respectivement, interpolation

bicubique avec la fonction de Keys (2.24), et interpolation spline cubique avec $\varphi_\alpha^{\text{eq}} = \gamma^3$, la spline cardinale cubique. Dans ce dernier cas, on a une étape de préfiltrage, suivie de la reconstruction avec la B-spline cubique :

$$g_{\alpha, \tau}(\mathbf{x}) = \sum_{\mathbf{k} \in \mathbb{Z}^2} c[\mathbf{k}] \varphi\left(\frac{\mathbf{x}}{T} - \tau_0 - \mathbf{k}\right), \quad (6.5)$$

avec

$$c = p_\alpha * v, \quad (6.6)$$

en prenant $\varphi = \beta^3$ et le préfiltre d'interpolation $p_\alpha = p_{\text{int}}$. Ainsi, l'interpolation spline de degré n de facteur α entier et d'offset $\mathbf{0}$ s'implémente simplement par

$$c = v * (b_1^n)^{-1} \quad \text{puis} \quad \mathcal{A}_{\alpha, \mathbf{0}} v = [v] \uparrow \alpha * b_\alpha^n, \quad (6.7)$$

où b_α^n est la B-spline dilatée discrète définie par (5.29).

La figure 6.1 montre un exemple d'image agrandie avec les méthodes d'interpolation usuelles. Bien que l'effet de flou soit contenu avec l'interpolation spline cubique et sinus cardinal, toutes les méthodes, hormis cette dernière, sont sujettes au repliement de spectre qui, bien que faible dans l'absolu, prend la forme d'un effet de crénelage particulièrement gênant le long des contours. L'interpolation sinus cardinal, qui ne présente ni flou, ni *aliasing*, paye ses propriétés par l'introduction d'oscillations nuisibles. Aucune fonction interpolante n'offre donc un compromis satisfaisant.

6.2.2 Agrandissement de facteur entier sous contrainte de réduction

Plaçons-nous temporairement dans le cas où le facteur d'agrandissement α est entier. Faisons aussi l'hypothèse que $\tilde{\varphi}$ vérifie une relation à deux échelles généralisée avec le filtre $\tilde{h}_{\alpha, -\alpha\tau}$ (propriété définie par (5.27)). Notons que l'on désignera dans la suite par $\tilde{h}$ un filtre associé au processus de réduction, et par h un filtre intervenant dans l'agrandissement.

Nous avons vu au chapitre précédent que l'on pouvait alors effectuer la réduction de facteur α et d'offset $-\alpha\tau$ de manière exacte, dans le cadre de notre modèle de redimensionnement :

$$\mathcal{R}_{\alpha, -\alpha\tau}^{\text{id}} u = [u * \tilde{h}_{\alpha, -\alpha\tau}] \downarrow \alpha. \quad (6.8)$$

De plus, d'après notre modèle,

$$\mathcal{R}_{\alpha, -\alpha\tau}^{\text{id}} \circ \mathcal{A}_{\alpha, \tau}^{\text{id}} = \mathcal{I}, \quad (6.9)$$

où $\mathcal{I}$ est l'identité. Par conséquent, l'image agrandie $y^{\text{id}} = \mathcal{A}_{\alpha, \tau}^{\text{id}} v$ que l'on souhaite estimer vérifie la propriété :

$$\mathcal{R}_{\alpha, -\alpha\tau}^{\text{id}} y^{\text{id}} = v \Leftrightarrow [y^{\text{id}} * \tilde{h}_{\alpha, -\alpha\tau}] \downarrow \alpha = v. \quad (6.10)$$

Autrement dit, **l'image agrandie que l'on cherche redonne l'image initiale v lorsqu'on la réduit**. Nous dirons qu'une image agrandie pour laquelle cette propriété est satisfaite vérifie la *contrainte de réduction*.

Il paraît donc judicieux de chercher une méthode d'agrandissement $\mathcal{A}_{\alpha,\tau}$ vérifiant la contrainte de réduction. Cette contrainte se substitue alors à la contrainte d'interpolation, qui n'est pas conforme à notre modèle. En fait, la contrainte d'interpolation est la contrainte de réduction avec le filtre trivial $\tilde{h}_{\alpha,-\alpha\tau} = \delta$, qui n'est clairement pas un bon filtre de réduction, à cause de l'*aliasing* introduit (voir figure 5.3).

De nombreux auteurs ont cherché des méthodes d'agrandissement non linéaires vérifiant la contrainte de réduction [194, 92, 179, 9, 180, 215, 26]. Plutôt que d'agrandissement sous contrainte de réduction, certains auteurs parlent d'agrandissement *linéairement réversible* [112, 148, 28]. Tous ces auteurs s'appuient cependant sur des modèles de formation simplistes, où $\tilde{\varphi}$ est généralement la fonction indicatrice $\mathbb{I}_{[-\frac{1}{2}, \frac{1}{2}]^2}$. Cela revient à considérer comme filtre de réduction de facteur 2 et d'offset $1/2$ le filtre moyennneur $[\frac{1}{2}, \frac{1}{2}, 0]$. Dans le prochain chapitre, nous montrerons comment effectuer l'agrandissement de manière non linéaire, tout en vérifiant exactement la contrainte de réduction avec un filtre arbitraire.

Dans cette section, nous nous intéressons aux méthodes d'agrandissement linéaires et invariantes par translation, ce qui nous ramène à trouver le filtre d'agrandissement $(h_{\alpha,\tau}[\mathbf{k}])_{\mathbf{k} \in \mathbb{Z}^2}$ dans l'équation (6.4). Alors, la contrainte de réduction s'écrit

$$[v] \uparrow \alpha * h_{\alpha,\tau} * \tilde{h}_{\alpha,-\alpha\tau} = v \downarrow \alpha \quad (6.11)$$

$$\Leftrightarrow [h_{\alpha,\tau} * \tilde{h}_{\alpha,-\alpha\tau}] \downarrow \alpha = \delta. \quad (6.12)$$

Lorsque cette dernière condition est vérifiée, on dit que les filtres $h_{\alpha,\tau}$ et $\tilde{h}_{\alpha,-\alpha\tau}$ sont des *partenaires biorthogonaux* de facteur α [228]. Remarquons que l'on effectue un agrandissement par interpolation lorsque le filtre $h_{\alpha,\tau}$ est interpolant :

$$h_{\alpha,\tau}[\mathbf{k}] = \begin{cases} 1 & \text{si } \mathbf{k} \in \alpha\mathbb{Z}^2 \\ 0 & \text{sinon} \end{cases}, \quad (6.13)$$

ce qui revient exactement à dire que $h_{\alpha,\tau}$ est biorthogonal à δ .

On peut maintenant trouver l'équivalent de l'interpolation spline, vérifiant la contrainte de réduction. Il suffit pour cela de remplacer le préfiltre d'interpolation $(b_1^n)^{-1}$ dans (6.7) par

$$c = v * ([b_\alpha^n * \tilde{h}_{\alpha,0}] \downarrow \alpha)^{-1}. \quad (6.14)$$

Par exemple, dans le cas de la réduction spline cubique donnée par (5.30), qui nous sert de référence pour la réduction, l'agrandissement spline cubique vérifiant la contrainte de réduction s'obtient par

$$c = v * ([b_\alpha^3 * b_\alpha^3] \downarrow \alpha)^{-1} * b_1^3 \text{ puis } \mathcal{A}_{\alpha,0}v = [v] \uparrow \alpha * b_\alpha^3. \quad (6.15)$$



FIG. 6.2 : Agrandissement $\times 4$ spline cubique. (a) : par interpolation. (b) : sous contrainte de réduction.

La figure 6.2 illustre le résultat obtenu. On voit que par rapport à l'interpolation spline cubique, l'image est moins floue, mais les effets de *ringing* et *aliasing* sont renforcés. L'absence de hautes fréquences se fait donc encore plus cruellement sentir que dans le cas de l'interpolation. Cependant, le contraste global est meilleur, car en accord avec notre modèle. On retiendra que les méthodes linéaires vérifiant la contrainte de réduction ne donnent pas des résultats visuellement plus satisfaisants que ceux obtenus par interpolation.

6.2.3 Agrandissement asymptotiquement optimal

Puisque l'agrandissement linéaire peut être vu comme un problème de reconstruction d'une fonction $g_{\alpha, \tau}$ suivi d'un échantillonnage, comme rappelé par (6.2) et (6.2), on peut, au lieu de s'intéresser directement au critère d'erreur $\|\mathcal{A}_{\alpha, \tau} v - \mathcal{A}_{\alpha, \tau}^{\text{id}} v\|_{\ell_2}$, définir $g_{\alpha, \tau}$ comme la fonction minimisant le critère suivant :

$$\varepsilon_{\text{AG}}(g) = \|g - \tilde{s}\|_{L_2}, \quad (6.16)$$

où l'on définit $\tilde{s}$, la fonction que l'on cherche à approcher, par

$$\tilde{s} = s * \frac{\alpha^2}{T^2} \tilde{\varphi}\left(\frac{\alpha \cdot}{T}\right) \xleftrightarrow{\mathcal{F}} \hat{\tilde{s}}(\omega) = \hat{s}(\omega) \hat{\tilde{\varphi}}\left(\frac{T}{\alpha} \omega\right). \quad (6.17)$$

Nous substituons donc à l'erreur discrète sur l'image agrandie, l'erreur, avant échantillonnage, entre la fonction reconstruite et ce qu'elle devrait être dans l'idéal. Puisque l'on est maintenant confronté à un problème de reconstruction, on peut utiliser le formalisme adopté à la section 2.4.3, afin d'analyser l'erreur ε_{AG} . On remarque d'abord que les échantillons $v[\mathbf{k}]$, qui sont des mesures sur s , peuvent aussi être interprétés comme des mesures

sur $\tilde{s}$:

$$\hat{v}(\omega) = \frac{1}{T^2} \sum_{\mathbf{k} \in \mathbb{Z}^2} \hat{s}\left(\frac{\omega + 2\pi\mathbf{k}}{T}\right) e^{I(\omega, \tau)} \hat{\varphi}(\omega + 2\pi\mathbf{k}) = \frac{1}{T^2} \sum_{\mathbf{k} \in \mathbb{Z}^2} \hat{s}\left(\frac{\omega + 2\pi\mathbf{k}}{T}\right) e^{I(\omega, \tau)} \frac{\hat{\varphi}(\omega + 2\pi\mathbf{k})}{\hat{\varphi}\left(\frac{\omega + 2\pi\mathbf{k}}{\alpha}\right)}. \quad (6.18)$$

On peut donc raisonner sur $\tilde{s}$ pour exprimer l'erreur d'approximation, au moyen du noyau d'erreur fréquentiel de la section 2.4.3 : pour une fonction g reconstruite linéairement selon le procédé décrit par (6.5), (6.6), on a

$$\varepsilon_{\text{AG}}(g)^2 = \frac{1}{(2\pi)^2} \int_{\mathbb{R}^2} |\hat{s}(\omega)|^2 \left[1 - \frac{|\hat{\varphi}(T\omega)|^2}{\hat{a}_{\varphi}(T\omega)} + \hat{a}_{\varphi}(T\omega) \left| \hat{p}_{\alpha}(T\omega) \frac{\hat{\varphi}(T\omega)}{\hat{\varphi}\left(\frac{T}{\alpha}\omega\right)} - \hat{\varphi}_d(T\omega) \right|^2 \right] d\omega + o(T^{2L}) \quad (6.19)$$

$$= \frac{1}{(2\pi)^2} \int_{\mathbb{R}^2} |\hat{s}(\omega)|^2 E(T\omega) d\omega + o(T^{2L}), \quad (6.20)$$

où L est l'ordre d'approximation de φ , et où l'on définit le *noyau d'erreur fréquentiel d'agrandissement* :

$$E(\omega) = \underbrace{\left(1 - \frac{|\hat{\varphi}(\omega)|^2}{\hat{a}_{\varphi}(\omega)} \right) |\hat{\varphi}\left(\frac{\omega}{\alpha}\right)|^2}_{E_{\min}(\omega)} + \hat{a}_{\varphi}(\omega) \left| \hat{p}_{\alpha}(\omega) \hat{\varphi}(\omega) - \hat{\varphi}_d(\omega) \hat{\varphi}\left(\frac{\omega}{\alpha}\right) \right|^2. \quad (6.21)$$

En utilisant ce noyau d'erreur, nous pouvons suivre la méthodologie adoptée à la section 2.5 : on fixe φ , et on choisit le préfiltre p_{α} afin que l'erreur ε_{AG} ait la décroissance la plus rapide possible lorsque T tend vers 0. Ainsi, l'agrandissement sera le meilleur possible pour le contenu basses fréquences de s . Notons $\mathcal{P}_T$ la projection orthogonale dans l'espace $V_T(\varphi)$. On cherche donc p_{α} tel que

$$\|g_{\alpha, \tau} - \tilde{s}\|_{L_2} \sim \|g - \mathcal{P}_T \tilde{s}\|_{L_2}, \quad T \rightarrow 0 \quad (6.22)$$

$$\Leftrightarrow E(\omega) \sim E_{\min}(\omega), \quad \omega \rightarrow 0 \quad (6.23)$$

$$\Leftrightarrow E(\omega) = E_{\min}(\omega) + O(\|\omega\|^{2L+2}) \quad (6.24)$$

$$\Leftrightarrow \hat{p}_{\alpha}(\omega) = \frac{\hat{\varphi}_d(\omega) \hat{\varphi}\left(\frac{\omega}{\alpha}\right)}{\hat{\varphi}(\omega)} + O(\|\omega\|^{L+1}). \quad (6.25)$$

Illustrons notre approche, en cherchant tout d'abord un filtre inverse symétrique séparable de taille 5×5 pour l'agrandissement spline cubique. p_{α} , $\tilde{\varphi}$ et φ étant séparables, on peut se ramener à des calculs 1D. On cherche donc un filtre de la forme

$$P_{\alpha}(z) = \frac{1}{a + b(z + z^{-1}) + c(z^2 + z^{-2})} \quad (6.26)$$

Le développement de Taylor (6.25) nous donne alors directement

$$a = \frac{78\alpha^4 + 1}{120\alpha^4}, \quad b = \frac{32\alpha^4 - 1}{180\alpha^4}, \quad c = \frac{1 - 2\alpha^4}{720\alpha^4}. \quad (6.27)$$

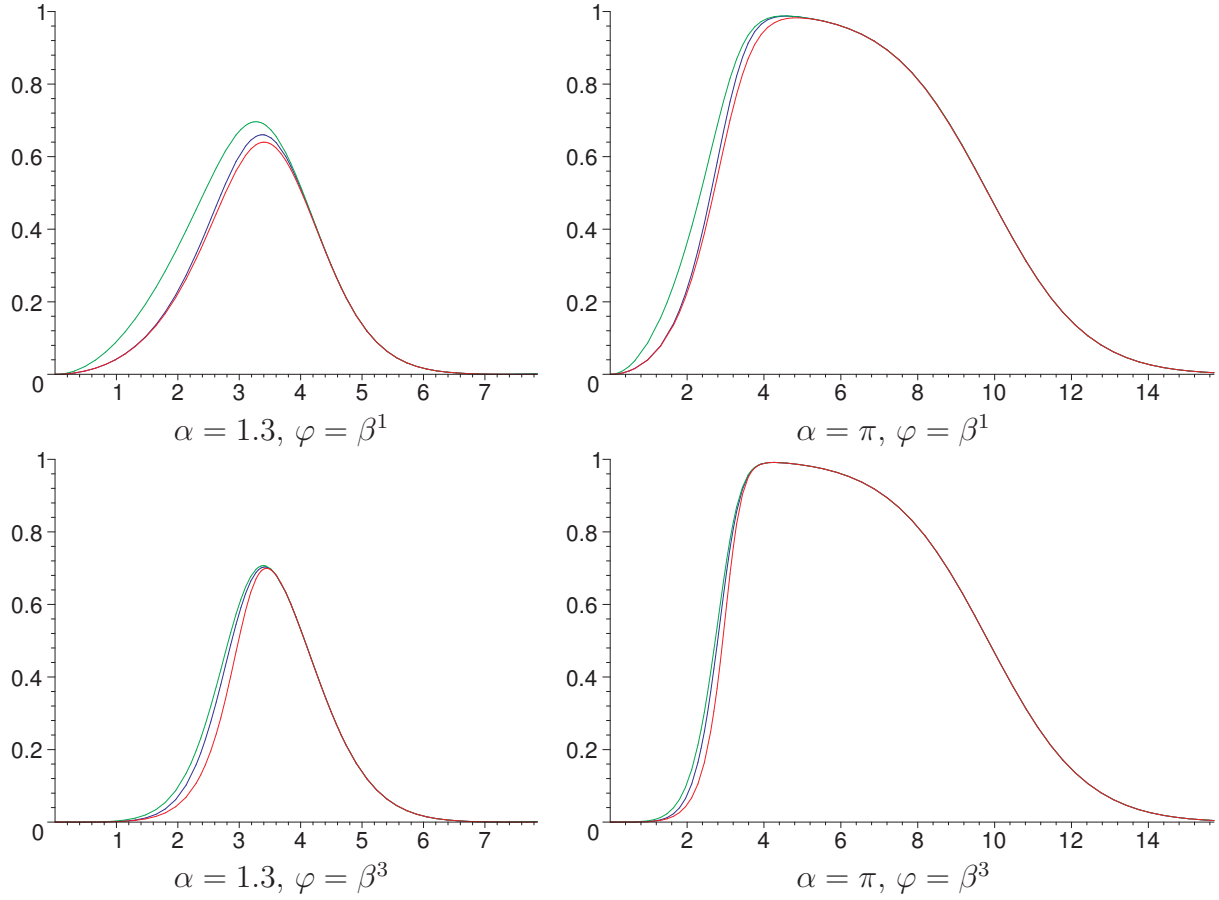


FIG. 6.3 : Noyaux d'erreur d'agrandissement, en fixant $\tilde{\varphi} = \gamma^3$ dans le modèle d'acquisition. en rouge : $\sqrt{E_{\min}(\omega)}$; en vert : $\sqrt{E_{\text{int}}(\omega)}$ associé à l'agrandissement par interpolation; en bleu : $\sqrt{E(\omega)}$ associé à nos nouveaux préfiltres.

Si maintenant on cherche un filtre inverse de taille 3 pour l'agrandissement spline linéaire, on obtient le préfiltre

$$P_{\alpha}(z) = \frac{1}{\frac{5}{6} + \frac{1}{6}(z + z^{-1})} \quad (6.28)$$

qui ne dépend pas de α , et qui se trouve être le même filtre que le préfiltre de quasi-interpolation présenté à la section 2.5.1.

Nous avons tracé sur la figure 6.3 les noyaux d'erreur pour une reconstruction spline de degré 1 et 3, et pour deux facteurs d'agrandissement différents. Les noyaux d'erreur montrent que l'erreur d'approximation est due en grande partie au fait que le contenu fréquentiel créé dans la bande $[\frac{\pi}{\alpha}, \pi]^2$ de l'image agrandie est très faible. On voit que les deux filtres que nous proposons sont meilleurs, et ce à toutes les fréquences, que les préfiltres d'interpolation. Cependant, le choix du préfiltre joue peu, au sens où le noyau E est proche dans tous les cas de sa borne inférieure $E_{\min}$. On constate donc que les méthodes d'agrandissement linéaires sont intrinsèquement limitées. Visuellement, les préfiltres asymptoti-

quement optimaux conduisent à des images très similaires à celles vérifiant la contrainte de réduction : l'effet de flou est réduit par rapport à l'interpolation, mais l'*aliasing* et le *ringing* sont renforcés.

Dans la suite de ce chapitre, nous nous tournons vers les méthodes d'agrandissement non linéaires présentées dans la littérature, dans le but de pallier le défaut des méthodes linéaires que constitue l'absence de création significative de contenu hautes fréquences, pourtant requis pour un rendu visuel satisfaisant des contours.

6.3 Agrandissement régularisé

Les méthodes visant à minimiser un certain critère, sous contrainte d'interpolation ou sous contrainte de réduction, forment une classe importante en agrandissement d'images. Ces méthodes partent d'un principe de parcimonie selon lequel une image lisse, au sens d'une certaine fonctionnelle énergétique, est plus vraisemblable qu'une image chaotique. Puisque le problème de l'agrandissement est un problème mal posé, la régularisation est une solution pour rendre la solution unique et bien définie. En assouplissant le problème de minimisation sous contrainte à l'aide d'un paramètre lagrangien (afin de tenir compte de la présence de bruit dans l'image v , ou à cause d'une incertitude ou inexactitude sur le modèle de réduction), on obtient la formulation générale suivante :

$$\mathcal{A}_{\alpha, \tau} v = \operatorname{argmin}_{u \in \ell_2(\mathbb{Z}^2)} \|\mathcal{R}_{\alpha, -\alpha\tau} u - v\|_{\ell_2}^2 + \lambda \mathcal{C}u, \quad (6.29)$$

pour une méthode de réduction $\mathcal{R}_{\alpha, -\alpha\tau}$ fixée *à priori*. Les méthodes interpolantes, dans le cas α entier, reviennent à choisir $\mathcal{R}_{\alpha, 0} v = [v] \downarrow \alpha$. La quasi-totalité des auteurs formulent l'attache aux données comme la contrainte d'interpolation (par exemple [125, 161, 160]), ou comme la contrainte de réduction avec le filtre $[\frac{1}{2}, \frac{1}{2}, 0]^2$, dans le cas $\alpha = 2$, $\tau = (-\frac{1}{4}, -\frac{1}{4})$, en partant d'un modèle d'acquisition simpliste avec $\tilde{\varphi} = \mathbb{1}_{[-\frac{1}{2}, \frac{1}{2}]^2}$ (par exemple [194, 8]). Ces deux modèles ne sont pas réalistes, et les résultats ne sont donc pas satisfaisants : en particulier, les méthodes d'interpolation donnent des textures (quand ce ne sont pas aussi des contours) flous. Notons aussi que ces deux formes d'attache aux données sont les plus simples, car elles amènent des contraintes sur les pixels agrandis orthogonales entre elles. La plupart des méthodes ne peuvent donc pas être généralisées à des filtres de réduction quelconques. De plus, toutes ces méthodes sont aussi limitées à des facteurs de réduction entiers (la plupart des auteurs ne s'intéressant d'ailleurs qu'au cas $\alpha = 2$). À notre connaissance, le seul auteur formulant l'agrandissement comme un problème de régularisation avec un modèle de réduction générique, utilisé en pratique avec un filtre déduit d'un formalisme rigoureux, est Hussein Aly [16, 15, 14].

Si le critère $\mathcal{C}$ est quadratique, on retombe sur les méthodes linéaires de la section précédente, dont on connaît les limites. Une étude systématique des critères quadratiques présentant de bonnes propriétés, comme l'invariance par rotation, a été menée dans [136, 137]. On

retrouve des critères basés sur les semi-normes de Duchon. D'autres critères quadratiques, et leur mise en œuvre numérique, ont été étudiés dans [127].

Le choix d'un critère non quadratique ouvre des possibilités beaucoup plus vastes pour améliorer la qualité visuelle de l'image agrandie. La solution du problème de minimisation ne dépend alors plus linéairement des données. On ne dispose généralement pour calculer l'image agrandie solution, que de méthodes itératives de descente de gradient, qui sont lentes, peuvent présenter des problèmes de stabilité et de convergence, et sont sensibles à la condition initiale. De plus, l'attache aux données peut être difficile à intégrer dans le schéma numérique sans en dégrader les propriétés théoriques.

6.3.1 Régularisation variationnelle

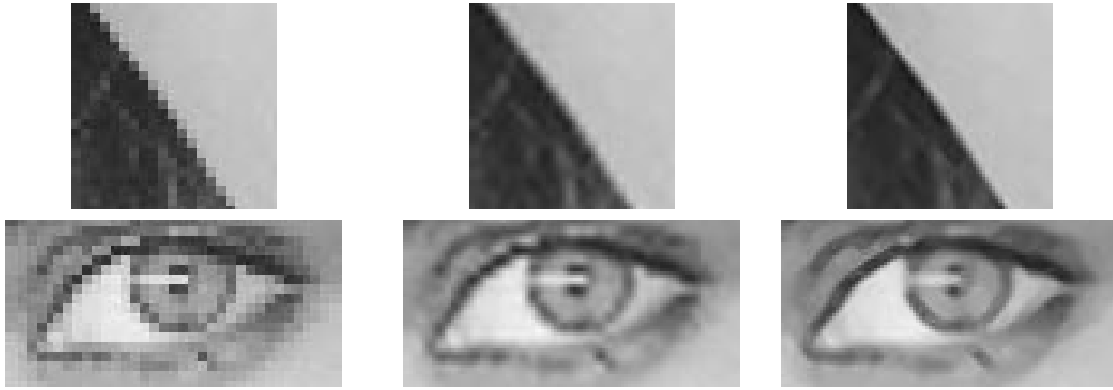
Plusieurs auteurs ont proposé de minimiser la variation totale de l'image agrandie, sous contrainte d'interpolation. La variation totale d'une fonction g s'exprime dans le domaine continu, comme l'intégrale du module du gradient :

$$\int_{\mathbb{R}^2} |\nabla g(\mathbf{x})| d\mathbf{x} = \int_{\mathbb{R}^2} \sqrt{\left(\frac{\partial g}{\partial x_1}\right)^2 + \left(\frac{\partial g}{\partial x_2}\right)^2} d\mathbf{x} \quad (6.30)$$

L'introduction de la variation totale comme mesure de régularité est due à Rudin et coll. [191], pour le débruitage des images. Elle a ensuite été relaxée et étudiée par Chambolle et Lions [56], puis adaptée pour l'interpolation d'images par Guichard et Malgouyres [112, 147, 148]. Almansa propose aussi une approche originale pour la restauration d'images floues et bruitées, basée sur la minimisation de la variation totale combinée avec un filtre de Wiener sur un domaine fréquentiel adapté au signal [10, 11].

Perceptuellement, nous sommes plus sensibles à des oscillations dans les zones uniformes que dans les zones texturées, mais la minimisation de la variation totale ne prend pas en compte ce phénomène. C'est pourquoi Moisan [159] propose de pondérer la variation totale par une fonction poids dépendant de l'inverse du module du gradient. Cela permet d'outrepasser l'inconvénient principal des méthodes à base de variation totale, dont la solution tend à être constante par morceaux avec des transitions franches, ce qui donne un effet d'« à plat » désagréable.

D'autres variantes du critère de variation totale ont été proposées, utilisant généralement le lien qui existe entre ce critère et la minimisation de la courbure des lignes de niveaux (*level sets*). Par exemple, la méthode proposée par Morse et coll. [160, 161] (voir figure 6.4) utilise une première estimation de l'image agrandie, qui est ensuite régularisée afin de supprimer les oscillations le long des lignes de niveaux, au moyen d'un lissage le long de celles-ci. Une autre méthode, dont la mise en œuvre n'est pas itérative, a été proposée par Jiang et coll. [125]. La discrétisation du problème conduit à une formule d'interpolation non linéaire assez simple. Les auteurs revendiquent des résultats aussi bons que ceux obtenus par les méthodes itératives [160, 161] pour une complexité algorithmique nettement réduite. Cependant, nous n'avons pas été en mesure de reproduire ces résultats.



Interpolation par duplication Interpolation bicubique Méthode de Morse et coll.

FIG. 6.4 : Illustration de la méthode d'agrandissement de Morse et coll. [160].

Généralement, le critère variationnel est formulé dans le domaine continu, et porte formellement sur la fonction dont les valeurs ponctuelles sur le treillis $\Lambda_{\frac{T}{\alpha}, \alpha(\tau_0 + \tau)}$ formeront l'image agrandie. La minimisation analytique du critère n'est généralement pas accessible. La mise en œuvre pratique consiste donc à discrétiser le critère sur le treillis (en approchant les quantités différentielles par des différences finies), afin de faire porter le problème directement sur les valeurs de l'image agrandie. L'image agrandie n'est donc qu'une solution approximative du problème posé dans le domaine continu. Il y a aussi souvent un amalgame entre la fonction (g_α dans nos notations) interpolant l'image agrandie, sur laquelle porte le problème, et la scène s : la contrainte de réduction est assimilée, à tort, à la contrainte de consistance, autrement dit, l'agrandissement est formalisé comme un problème de reconstruction suivi d'un simple échantillonnage (par exemple dans [148]).

Un autre moyen de formuler l'agrandissement régularisé est d'utiliser des équations aux dérivées partielles (EDPs). Un panorama de leurs applications en traitement d'image est présenté dans [113]. Leur utilisation en agrandissement d'images a émergé récemment [54]. Plutôt que de chercher l'image minimisant un certain critère, la solution est formulée comme la solution d'une EDP, sous contrainte de réduction. On peut par exemple chercher une fonction de classe C^2 dont le Laplacien est nul : $\Delta g = 0$. Une telle fonction minimise l'intégrale de Dirichlet $\int_{\mathbb{R}^2} |\nabla g(x, y)|^2 dx dy$. Ainsi, pour certains problèmes de minimisation d'un critère, on peut trouver une EDP équivalente. Généralement, celle-ci n'a pas de solution analytique. Le problème est donc transformé en introduisant un paramètre t supplémentaire, le temps, et on s'intéresse à la fonction $g(\mathbf{x}, t)$ dont l'évolution au cours du temps est décrite par l'EDP, par exemple $\frac{\partial g}{\partial t} = \Delta(g)$. La solution du problème est alors définie comme la fonction stationnaire obtenue après un temps infini. L'utilisation d'EDPs fournit donc un formalisme puissant et esthétique pour l'utilisation de méthodes itératives régularisant une image agrandie initiale : le caractère itératif de la mise en œuvre n'est plus *ad hoc*, mais résulte du formalisme adopté. Le travail de Belahmidi et coll. [28] est un bon

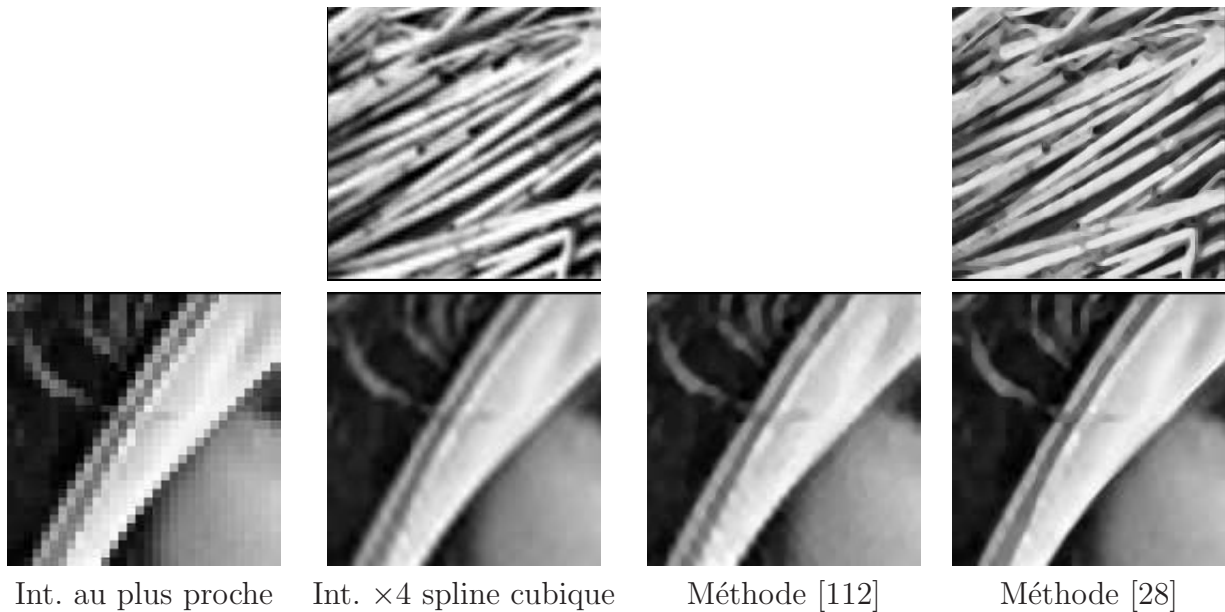


FIG. 6.5 : Illustration de la méthode d'agrandissement de Belahmidi et coll. [28].

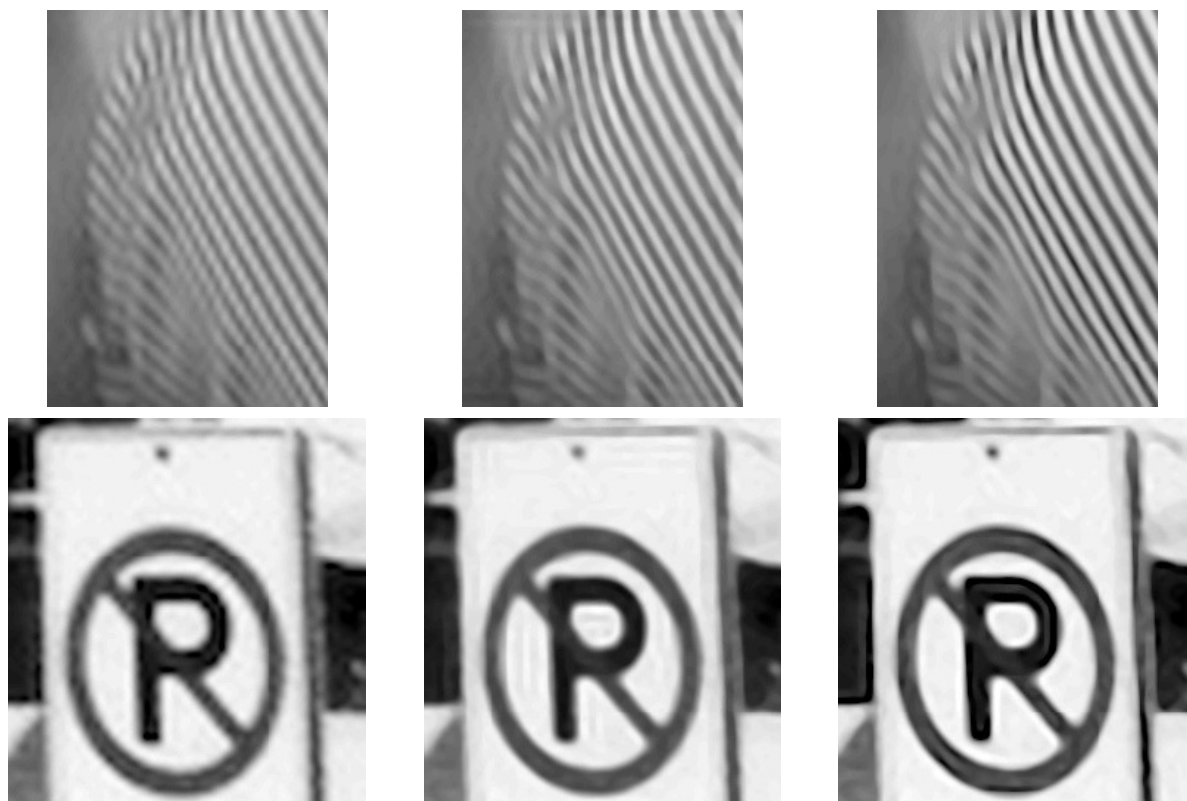
exemple de ce type de méthodes. Cet auteur utilise une EDP issue du travail pionnier de Perona et Malik [174]. Il obtient ainsi de meilleurs résultats que ceux reportés dans [194] et [112].

Nous renvoyons à la thèse de Hussein Aly [16] pour un panorama des méthodes d'agrandissement par régularisation variationnelle. Cet auteur propose une méthode basée sur la variation totale, avec une contrainte de réduction formulée précisément [14]. Il obtient ainsi des images moins floues que celles issues des méthodes variationnelles classiques (voir figure 6.6).

6.3.2 Autres méthodes de régularisation

Une des premières méthodes d'agrandissement par régularisation, proposée par Schultz et Stevenson [194], repose sur un formalisme statistique. L'image agrandie est modélisée par un champ de Markov, et définie comme la solution d'un problème de minimisation impliquant un critère qui dépend localement de la présence ou non d'un contour. Il a été souligné dans [28] que les propriétés de cette méthode ne sont pas clairement identifiables, et que l'énergie est difficile à minimiser car la descente de gradient est instable.

En s'appuyant sur la théorie de l'*optimal recovery* et ses propriétés d'optimalité minimax, Muresan a proposé la méthode AQua [167], qui construit l'image agrandie minimisant localement un critère quadratique, celui-ci étant appris dans une voisinage centré en chaque point de l'image initiale [167]. Une variante (AQua2), plus simple et plus rapide, a récemment été proposée [166]. Comme le montre la figure 6.7, la méthode AQua2 est particulièrement efficace pour fournir des images dépourvues d'effets d'escaliers le long des contours



Interp. $\times 5$ spline cubique Méthode de Mal. et coll. [147] Méthode de Aly et coll.

FIG. 6.6 : Illustration de la méthode d'agrandissement de Aly et coll. [14].

obliques, mais il s'agit d'une méthode d'interpolation, et les contours apparaissent flous malgré tout. De plus, un effet désagréable de moutonnement apparaît dans les textures.

6.4 Agrandissement par extrapolation

L'agrandissement sous contrainte de réduction est un problème inverse mal posé. Pour une image v de taille N^2 à agrandir de facteur α , on cherche $(\alpha N)^2$ nouvelles valeurs de pixels, étant données N^2 contraintes linéaires sur ceux-ci. La contrainte de réduction porte plutôt sur le contenu basses fréquences de l'image agrandie, alors que les fréquences au-delà de la fréquence de Nyquist associée à v ne sont presque pas contraintes. Plutôt que de mettre à zéro la partie hautes fréquences, comme le font les méthodes linéaires, on peut chercher à extrapoler à la résolution de l'image agrandie les fréquences nécessaires à la synthèse des structures géométriques détectées dans l'image initiale.

En remarquant que dans la pyramide Laplacienne [48] d'une image, des structures apparaissent à tous les niveaux de la pyramide, on peut chercher comment ces structures sont reliées à travers les niveaux, afin d'extrapoler un niveau de pyramide supplémentaire. Celui-ci est sensé contenir les détails manquants lorsqu'on agrandit v par une méthode li-

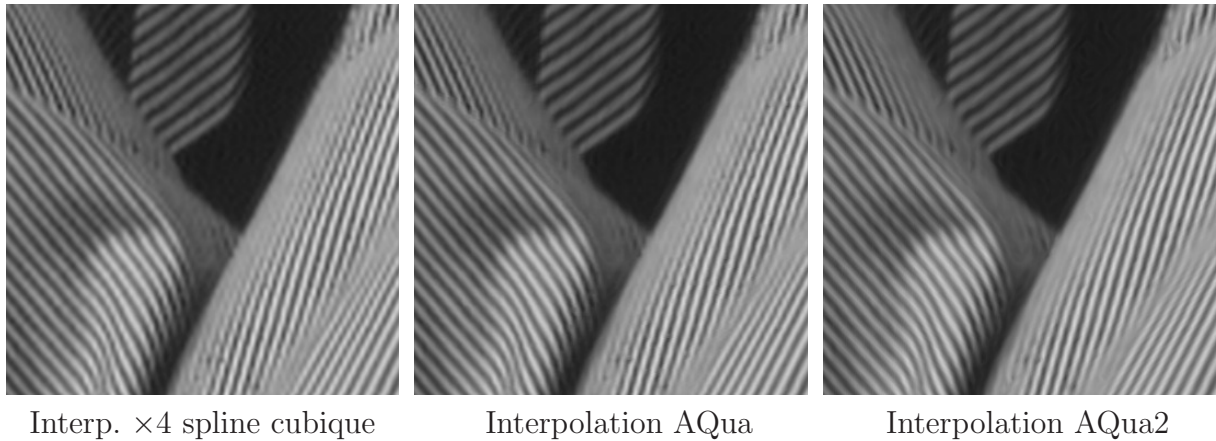
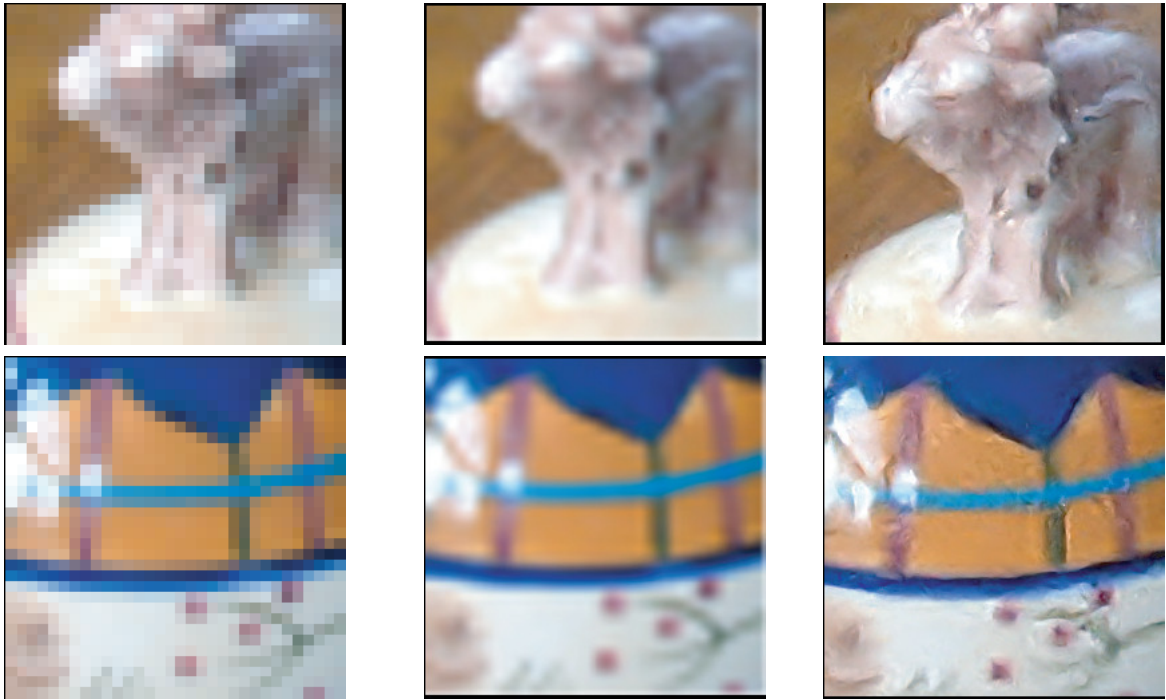


FIG. 6.7 : Illustration, sur une partie du pantalon de Barbara, des méthodes d'interpolation AQua et AQua2 de Muresan et coll. [167]. L'image du milieu a été obtenue à l'adresse <http://dsplab.ece.cornell.edu/papers/results/orinterp/zipped/> et l'image de droite générée à l'aide du logiciel libre Pictura disponible à l'adresse http://www.dmmd.net/products/pictura_downloads.htm

néaire d'agrandissement, et sera donc ajouté à une première image agrandie linéairement à partir de v . La méthode de Greenspan et coll. [108] repose sur cette idée, mais le caractère non linéaire de la dépendance entre les niveaux rend difficile l'extrapolation, en particulier à cause des problèmes de signe des coefficients. La méthode de Freeman et coll. [102] consiste aussi à chercher cette image de différence qui doit être ajoutée à l'image agrandie par interpolation spline cubique, pour obtenir une image agrandie avec des détails nets. La méthode repose sur un apprentissage à partir duquel un dictionnaire de blocs d'images est constitué, avec pour chaque bloc, le bloc de différence apparié à ajouter dans l'image agrandie. L'image à agrandir est ensuite décomposée en blocs, qui sont comparés à ceux du dictionnaire pour former l'image de différence adéquate. On peut objecter que les détails haute résolution pertinents dans une image ne sont pas forcément les mêmes pour une autre image ayant localement une structure similaire à la résolution de l'image initiale. On peut voir sur les exemples de la figure 6.8 que les contours dans l'image agrandie obtenue sont francs, mais pas droits, et l'image est polluée par de nombreux artéfacts prenant la forme de mini structures courbées, donnant un aspect de moutonnement artificiel déplaisant à l'image agrandie.

La plupart des méthodes d'extrapolation se placent dans le domaine de la transformée en ondelettes, et posent le problème de l'agrandissement comme celui de l'extrapolation des coefficients d'ondelettes à la résolution de l'image agrandie. Nous verrons au chapitre 7 que ce formalisme est rigoureux, si les filtres associés à la transformée sont choisis de manière adéquate.

Une première possibilité pour extrapoler les coefficients d'ondelettes est de s'intéres-



Interpolation par duplication Interpolation spline cubique Méthode de Freeman[102]

FIG. 6.8 : Illustration de la méthode d'extrapolation de Freeman et coll. [102].

ser à leur évolution à travers les résolutions. En effet, comme montré par plusieurs auteurs [150][79], les maxima d'amplitude des coefficients d'ondelettes qui correspondent à des transitions franches dans l'image, susceptibles d'être des contours, ont une certaine régularité (appelée Hölderienne ou Lipschitzienne) que l'on peut caractériser. La méthode proposée dans [58, 57] utilise cette propriété d'évolution des maxima locaux avec une transformée en ondelette non décimée, ce qui rend l'approche similaire à celles basées sur une pyramide de différences. Des méthodes d'extrapolation basées sur la transformée en ondelettes décimée, utilisant la régularité de Hölder, sont proposées dans [131, 132] et [140]. Malheureusement, avec la transformée décimée, la loi d'évolution des coefficients à travers les échelles n'est plus vérifiée.

Nicolier et coll. [168] utilisent, quant à eux, une propriété portant sur les passages par zéro (*zero-crossings*) de la transformée en ondelettes, plutôt que les maxima des coefficients. Dans [130], des arbres de Markov cachés sont utilisés pour modéliser la dépendance des coefficients à trouver. Dans [249], les contours sont détectés dans l'image initiale et un modèle de contour idéal paramétrique (amplitude, angle...) est appliqué sur chacun, pour en déduire les coefficients d'ondelettes appropriés à la synthèse du contour dans l'image agrandie.

Les méthodes à base d'extrapolation explicites de coefficients d'ondelettes ne donnent pas de résultats probants. Même si l'on peut prédire l'amplitude des coefficients à extrapoler, il n'en va pas de même de leur signe. Enfin, la théorie valide en 1D s'étend difficilement

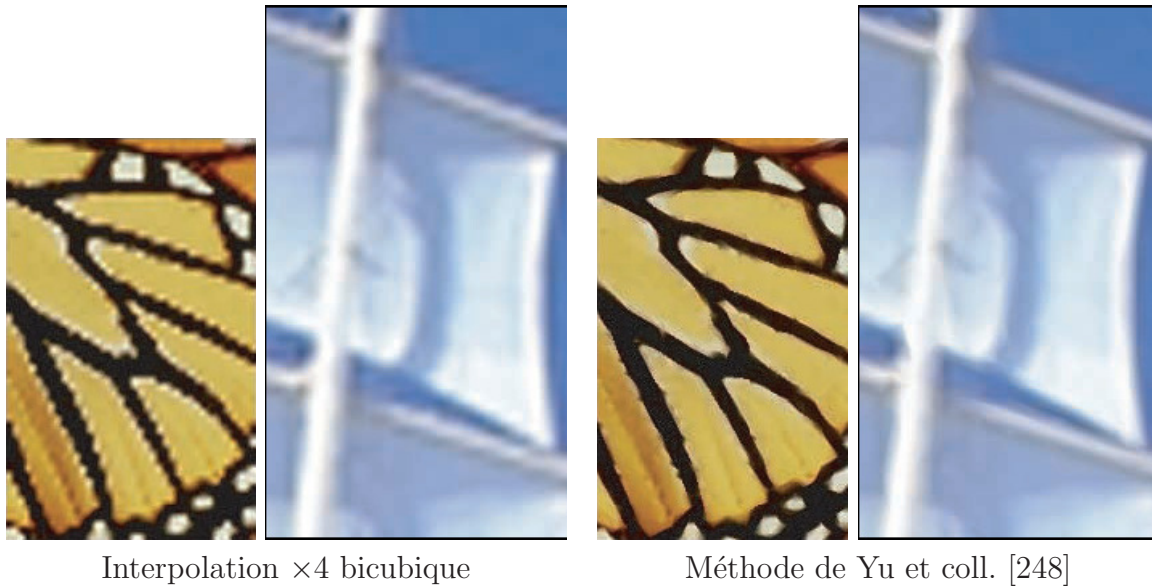


FIG. 6.9 : Illustration de la méthode d'agrandissement par triangulation de Yu et coll. [248].

en 2D, où la transformée en ondelettes séparable ne tient pas compte de la structure particulière des contours obliques, dont la représentation dans le domaine ondelettes est très difficile à caractériser. Généralement, il y a donc apparition d'effets d'escaliers importants dans l'image agrandie. Nous verrons au chapitre 7 que l'on peut formuler l'agrandissement comme un problème d'extrapolation de coefficients d'ondelettes, sans entreprendre explicitement cette extrapolation en pratique.

6.5 Agrandissement préservant les structures

Plutôt que de préserver les fréquences (qui n'ont pas de caractère local), de calculer une image lisse (la prise en compte des discontinuités étant difficile), ou d'extrapoler des coefficients dans un certain domaine (où les structures de l'image initiale sont transformées et rendues inexploitable), les méthodes les plus efficaces agissent directement dans le domaine spatial, en prenant en compte les caractéristiques locales de l'image initiale. Ces méthodes, plus ou moins heuristiques, visent principalement à opérer un traitement spécifique au voisinage d'un contour, afin de l'agrandir en limitant l'effet d'escalier.

6.5.1 Méthodes adaptatives

Une approche simple pour traiter les contours de manière adéquate, visant à diminuer les effets d'escaliers par rapport aux méthodes linéaires, consiste à découper l'image en blocs, qui sont classés dans des catégories, par exemple : zone homogène, zone texturée, contour, et d'utiliser un filtre d'agrandissement approprié à chaque catégorie. C'est la méthode

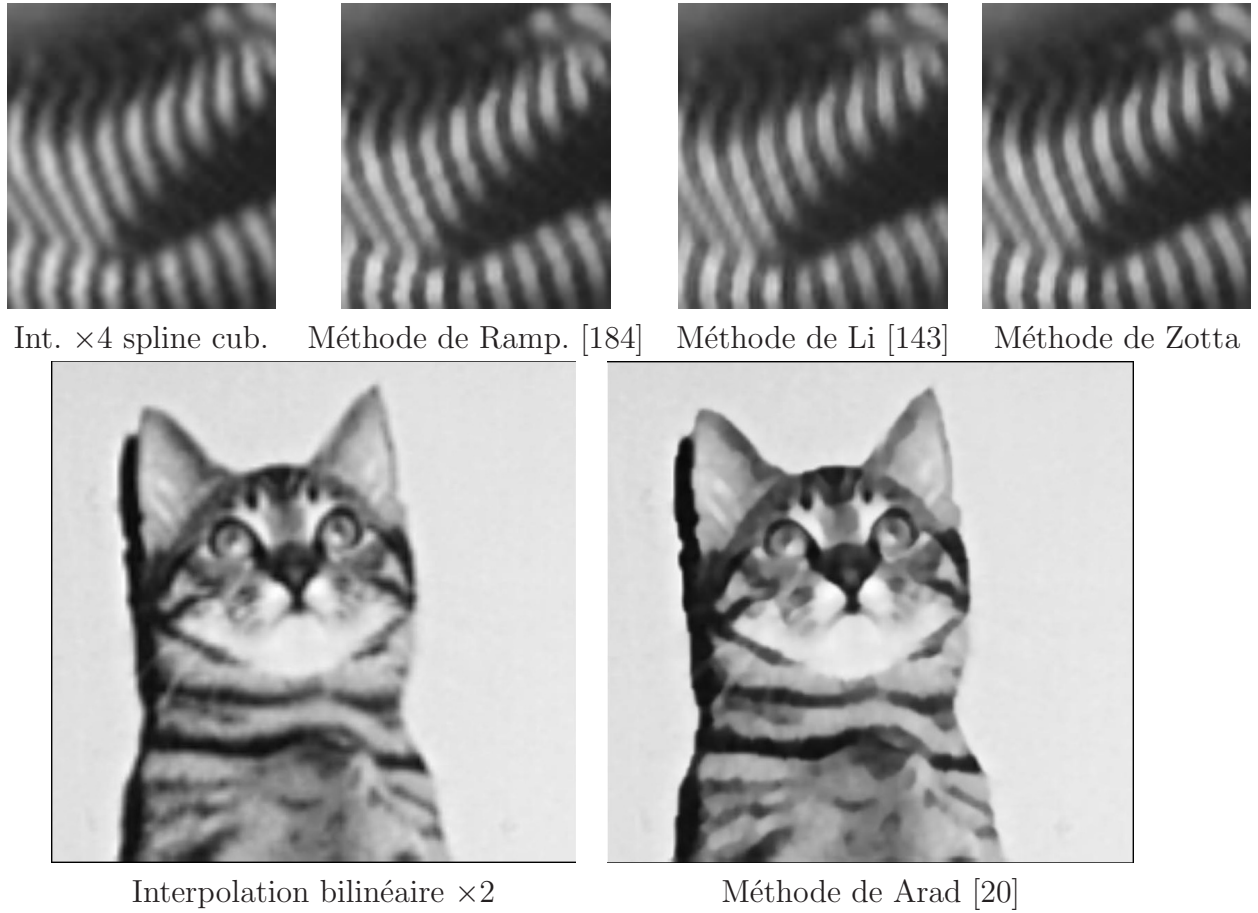


FIG. 6.10 : Illustration des méthodes de *distance warping* de Zotta et coll. [250], et de Arad et coll. [20].

retenue dans [116] : la classification est réalisée sur des blocs 4×4 et une interpolation directionnelle adaptée est ensuite effectuée. Ce type de méthodes présente un inconvénient majeur : si les blocs sont trop grands, le caractère local des contours n'est pas bien exploité, et si les blocs sont trop petits, la décision est sensible au bruit, et une mauvaise décision entraîne des artefacts très visibles. De plus, le nombre de classes étant fini, tous les contours obliques d'orientations arbitraires ne peuvent pas être synthétisés.

Dans un tout autre registre, la méthode proposée par Yu et coll. [248] consiste à chercher une triangulation de l'image initiale telle que les contours soient le long des bords des triangles, et ne les traversent pas, tout en minimisant une certaine fonctionnelle. La fonction reconstruite est continue, interpolante, et linéaire à l'intérieur de chaque triangle. La recherche de la triangulation optimale est assez longue, mais les contours obliques semblent exempts d'effet d'escaliers (voir figure 6.9).

Les méthodes plus sophistiquées mettent en œuvre une décision locale suivant qu'il y a vraisemblablement un contour ou non. Ramponi et coll. ont proposé plusieurs méthodes en ce sens [51, 185, 52, 209]. Une autre approche, toujours de la même équipe, est basée

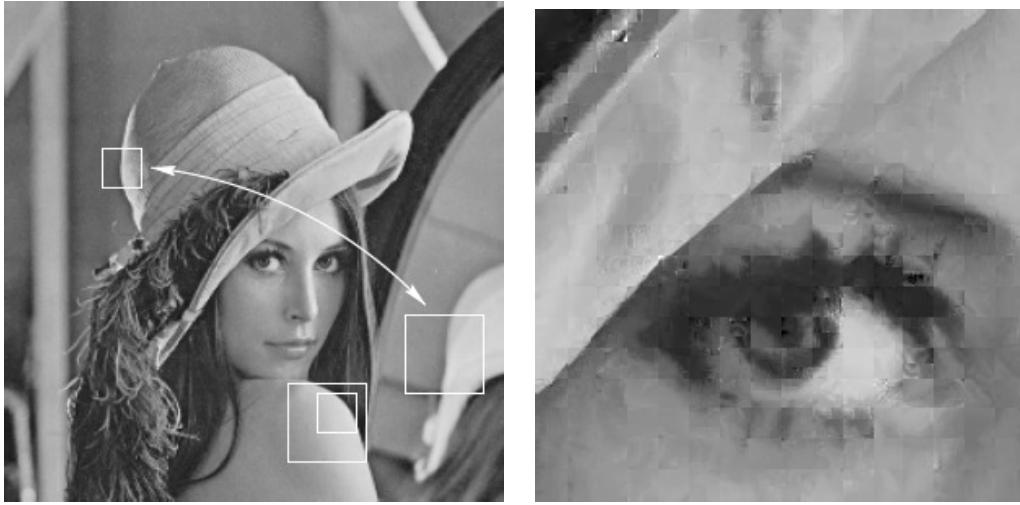


FIG. 6.11 : Exemples d'auto-similarités locales dans *Lena* et, à droite, un détail de son agrandissement par zoom fractal $\times 4$.

sur le principe de déformation de distance (*distance warping*) [184, 250] : les contours sont synthétisés avec un profil plus franc dans l'image agrandie, sans être détectés explicitement. Une autre technique utilisant un champ de déformation 2D est proposée par Arad et coll. [20]. La déformation tient compte du gradient de l'image initiale, qui est calculé par le détecteur de Canny-Deriche. Les contours obtenus sont francs, mais l'image agrandie souffre là aussi d'un effet d'« à plat » (voir figure 6.10).

6.5.2 Zoom fractal

L'agrandissement par zoom fractal est basé sur l'existence de similarités structurales entre les résolutions. Cette propriété d'invariance par changement de résolution des contours et de certaines textures est aussi utilisée par les méthodes d'extrapolation utilisant la régularité Lipschitzienne. Plutôt que de constituer une base de blocs d'images servant de référence, le zoom fractal consiste à chercher des auto-similarités à différentes échelles dans l'image à agrandir. Une étude détaillée du zoom fractal est donnée dans [49]. La méthode fournit des images avec des contours francs, mais le découpage en blocs est souvent visible dans l'image agrandie, avec des structures parasites. Cela découle du principe d'auto-similarité qui n'est jamais parfaitement vérifié dans les images naturelles.

6.5.3 Approches basées contours

Les méthodes directionnelles cherchent localement si un contour est présent et, le cas échéant, quelles sont son orientation et son amplitude. Ainsi chaque pixel de l'image agrandie est calculé en appliquant localement sur l'image initiale un filtre aligné le long du contour, afin de ne pas introduire de flou, qui résulte habituellement du moyennage de

Interpolation $\times 4$ bicubique

Méthode de Allebach [8]

Méthode de Li [143]

FIG. 6.12 : Illustration des méthodes d'interpolation directionnelle de Allebach et coll. [8], et Li et coll. [143].

pixels d'objets différents. Allebach et coll. [8] ont proposé une méthode allant en ce sens : une interpolation bilinéaire est réalisée, sauf si un contour est détecté, auquel cas seuls certains pixels du voisinage servent au calcul. L'approche a été améliorée par Li et coll. [143] (voir figure 6.12) : l'interpolation est réalisée sur le treillis quinquonce à l'aide des quatre plus proches voisins du pixel à interpoler, avec des poids appris dans un voisinage local, en faisant une hypothèse d'auto-similarité aux différentes résolutions.

La méthode de Jensen et coll. [123] utilise une détection de contours. Pour chaque pixel de l'image initiale, le bloc 3×3 avoisinant est considéré, et divers critères sont utilisés pour décider si l'on est en présence d'un contour. Si c'est le cas, les paramètres d'un contour sigmoïdal idéal sont calculés, et c'est ce contour qui est échantillonné pour donner les pixels de l'image agrandie. S'il n'y a pas de contour, une simple interpolation bilinéaire est effectuée. Malgré sa simplicité, cette méthode se révèle particulièrement efficace pour le rendu des contours, qui sont francs et présentent peu de crénelage. Par contre, on a une atténuation forte des objets de petite taille, et les coins et intersections donnent parfois lieu à des artéfacts. L'exemple de la figure 6.13 a été obtenu avec notre propre implémentation de cette méthode.

6.5.4 Méthodes par apprentissage

Afin de trouver le meilleur agrandissement vérifiant une contrainte de réduction avec $\mathcal{R}_{\alpha, -\alpha\tau}$, il est possible d'apprendre l'agrandissement optimal à partir d'une base d'images et de leurs versions réduites avec $\mathcal{R}_{\alpha, -\alpha\tau}$. Si l'on cherche un filtre d'agrandissement linéaire, on obtient un filtre de Wiener, l'autocorrélation des images étant apprise plutôt que modélisée (si tant est qu'existe une telle autocorrélation générique pour l'ensemble des images naturelles). Atkins et coll. [24] proposent une version améliorée de cette méthode, agissant localement sur des blocs de l'image (figure 6.14). Lors d'une phase d'apprentissage, les blocs d'une base d'images sont rangés dans des classes, et le filtre d'agrandissement linéaire

Interpolation $\times 2$ spline cubique

Méthode de Jensen [123]

FIG. 6.13 : Illustration de la méthode d'agrandissement de Jensen et coll. [123].

optimal est calculé pour chaque classe. Les probabilités *a priori* d'appartenance d'un bloc arbitraire à chaque classe sont aussi calculées. L'agrandissement d'un bloc donné se fait ensuite suivant un formalisme bayésien.

Des approches plus ou moins similaires ont été proposées par Kondo [135] et par Hu et coll. [117]. Ce dernier utilise un réseau de neurones optimisé pour chaque classe au lieu d'un agrandissement linéaire. On peut opposer aux méthodes reposant sur un apprentissage plusieurs critiques : l'apprentissage est coûteux en temps et en stockage de la base d'images, et empirique car reposant sur une base d'images arbitraire. Il faut ensuite maintenir en mémoire les résultats de l'apprentissage. De plus, le fait de raisonner sur des blocs locaux ne permet pas d'assurer une contrainte de réduction globale. Enfin, une erreur de classification, lors de l'agrandissement, entraînera un résultat imprévisible avec généralement création de structures parasites.

6.6 WiskI : une nouvelle méthode d'interpolation adaptative

Nous avons présenté dans les sections précédentes plusieurs classes de méthodes, reposant sur des intuitions différentes. Cependant, toutes cherchent à prendre en compte la structure locale de l'image, afin d'effectuer un traitement approprié en présence d'un contour. Nous proposons dans cette section une nouvelle méthode d'interpolation, spécifique au cas $\alpha = 2$, $\tau = \mathbf{0}$. L'extension à d'autres facteurs ainsi que la prise en compte d'une contrainte de réduction n'ont pas encore été développées, et sont une voie d'étude future. Cependant, il nous semble important de présenter cette approche originale, les résultats préliminaires obtenus étant très satisfaisants.



Agrandissement bicubique

Méthode proposée par Atkins et coll.

FIG. 6.14 : Illustration de la méthode par apprentissage de Atkins et coll. [24]

Si l'on suppose que l'image agrandie y^{id} que l'on cherche à estimer est la réalisation d'un processus aléatoire stationnaire d'autocorrélation connue c_y , il est possible de déterminer le filtre optimal $h_{\alpha,\tau}$ d'agrandissement de facteur entier α , minimisant l'erreur

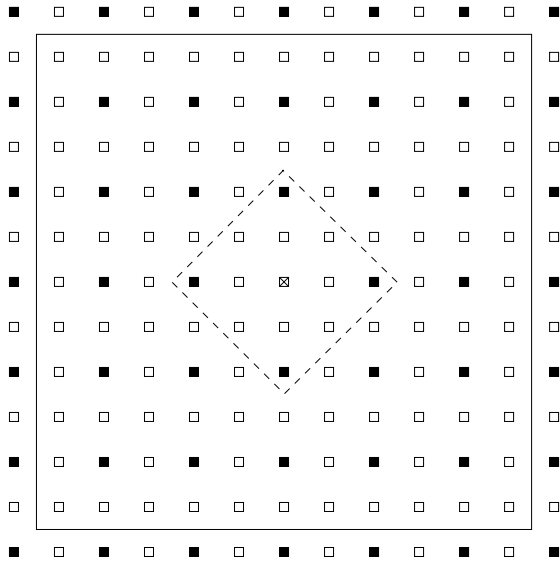
$$\varepsilon_{\text{STO},\mathbf{k}} = \mathcal{E} \left\{ \left| [v] \uparrow \alpha * h_{\alpha,\tau}[\mathbf{k}] - y^{\text{id}}[\mathbf{k}] \right|^2 \right\} \quad (6.31)$$

pour tout $\mathbf{k} \in \mathbb{Z}^2$. Ce filtre, version discrétisée du filtre de reconstruction optimal présenté au chapitre 2 (équations (2.13),(2.14)), est le suivant :

$$h_{\alpha,\tau} = c_y * \tilde{h}_{\alpha,-\alpha\tau} * \left([c_y * \tilde{h}_{\alpha,-\alpha\tau} * \tilde{h}_{\alpha,-\alpha\tau}] \downarrow \alpha \uparrow \alpha \right)^{-1}. \quad (6.32)$$

On remarque que ce filtre est bien biorthogonal à $\tilde{h}_{\alpha,-\alpha\tau}$, et la contrainte de réduction est donc assurée.

Quand bien même on connaîtrait l'autocorrélation c_y , cette méthode, comme toutes les méthodes d'agrandissement linéaires, n'est pas satisfaisante, car l'hypothèse de stationnarité des images est fausse. Par contre, les images naturelles sont structurées, et il est vraisemblable de les modéliser comme des processus « localement stationnaires » : chaque pixel de l'image agrandie n'est pas indépendant de ses voisins, et que ce soit dans une zone homogène, une zone texturée, ou au niveau d'un contour, la corrélation d'un pixel avec ses voisins varie peu d'un pixel à l'autre. Cependant, cela ne résout pas le problème de l'agrandissement, puisque l'autocorrélation de l'image agrandie, et *a fortiori* son autocorrélation locale, ne sont pas connues. Li, partant de ce principe de « stationnarité locale », a proposé une méthode d'interpolation, en faisant une hypothèse d'invariance à travers les échelles pour estimer l'autocorrélation locale de l'image agrandie à partir de celle de l'image initiale [143]. Sa méthode, comme toutes les méthodes d'interpolation directionnelle, donne un résultat mitigé : certes, les contours ne sont pas crénelés, mais un effet de traînée désagréable est présent (voir figure 6.12).



Les carrés noirs sont les sites du treillis de l'image initiale v , les carrés blancs sont les sites du treillis de l'image agrandie y . Pour chaque site $\alpha\mathbf{k}$ (par exemple, le site avec une croix), on calcule le filtre $a_{\mathbf{k}}$, qui représente l'autocorrélation locale autour du pixel $y[\alpha\mathbf{k}]$. Le support $\aleph$ de $a_{\mathbf{k}}$ a une forme de diamant (carré quinconce en pointillés, représenté centré en $\alpha\mathbf{k}$) contenant 13 sites. La fenêtre d'apprentissage $\mathcal{U} + \alpha\mathbf{k}$ servant au calcul de $a_{\mathbf{k}}$ contient 121 pixels autour de $y[\alpha\mathbf{k}]$.

FIG. 6.15 : Procédé de calcul de l'autocorrélation locale pour la méthode d'interpolation WiskI.

Muresan propose, lui, de modéliser localement les pixels de l'image agrandie comme appartenant à une classe quadratique qu'il s'agit d'estimer. Il applique ensuite le formalisme d'estimation minimax au sens de l'*optimal recovery*, pour en déduire le filtre d'agrandissement à appliquer localement [167]. Bien que son interprétation déterministe minimax diffère de l'approche stochastique pour laquelle nous optons, des similitudes apparaissent dans l'objectif à atteindre. La mise en œuvre adoptée par Muresan est intéressante : il montre que l'on peut estimer avec une bonne qualité la classe quadratique locale de l'image agrandie, en la calculant à partir d'une version interpolée de l'image initiale, par exemple par interpolation bicubique. Ce résultat est empirique, et n'est pas expliqué par Muresan, mais les résultats obtenus avec cette technique sont bons *a posteriori*.

Nous nous inspirons de l'astuce employée par Muresan pour sa méthode AQua, et nous proposons ainsi une nouvelle méthode d'interpolation, que nous appelons WiskI¹. Afin d'estimer localement l'autocorrélation de l'image agrandie, nous utilisons l'interpolée bicubique de v , que nous notons y^0 . Comme pour toute méthode adaptative, les paramètres importants sont la taille du filtre d'autocorrélation d'une part (puisque l'on suppose celle-ci locale donc pertinente uniquement dans un voisinage de taille finie), et la taille de la fenêtre servant pour son apprentissage d'autre part (une grande fenêtre donnant un résultat plus fiable mais aussi moins local). Pour chaque pixel d'indice $\alpha\mathbf{k}$ de l'image agrandie y à construire (c'est-à-dire aux positions des pixels de l'image initiale), nous calculons le filtre d'autocorrélation $a_{\mathbf{k}}$, dont le support $\aleph$ est un diamant de taille 13 (voir la figure 6.15).

1. pour « Interpolation par estimation de **W**iener **l**ocale »

Nous choisissons une fenêtre d'apprentissage $\mathcal{U} + \alpha \mathbf{k}$ carrée de taille 11×11 , centrée autour du point $\alpha \mathbf{k}$. On calcule donc, pour chaque indice $\mathbf{k}$ de l'image initiale, l'autocorrélation empirique

$$a_{\mathbf{k}}[\mathbf{l}] = \frac{1}{\#\mathcal{U}} \sum_{\mathbf{j} \in \mathcal{U}} y^0[\alpha \mathbf{k} + \mathbf{l} + \mathbf{j}] y^0[\alpha \mathbf{k} + \mathbf{j}] \quad \forall \mathbf{l} \in \mathbb{N}. \quad (6.33)$$

Une fois déterminés les filtres $a_{\mathbf{k}}$, l'image agrandie s'écrit

$$y[\mathbf{l}] = \sum_{\mathbf{k} \in \mathbb{Z}^2} c[\mathbf{k}] a_{\mathbf{k}}[\mathbf{l} - \alpha \mathbf{k}] \quad \forall \mathbf{l} \in \mathbb{Z}^2, \quad (6.34)$$

pour des coefficients $c[\mathbf{k}]$ à déterminer afin que la contrainte d'interpolation $y[\alpha \mathbf{k}] = v[\mathbf{k}]$ soit vérifiée. Si les $a_{\mathbf{k}}$ étaient tous identiques, c s'obtiendrait par le filtrage inverse $c = v * ([a_{\mathbf{k}}] \downarrow \alpha)^{-1}$. Ici, il faudrait effectuer un filtrage inverse spatialement variant, qui ne peut être réalisé en 2D qu'au moyen d'une méthode itérative. Nous avons mis au point une méthode qui effectue ce filtrage de manière directe mais approchée. Nous ne détaillons pas ce procédé expérimental qui demande encore plusieurs ajustements.

Si l'on applique la méthode que nous venons de décrire, on obtient une image y meilleure que l'estimée initiale y^0 , mais les résultats sont encore meilleurs si on itère la méthode : on utilise y pour estimer à nouveau l'autocorrélation, qui sert à générer l'image y^2 , etc. Deux itérations sont suffisantes en pratique. La figure 6.16 montre le résultat obtenu sur le pantalon de Barbara. Notre méthode fournit des images interpolées sans aucun effet de crénelage le long des contours obliques. Cela confirme le bien-fondé du paradigme de stationnarité locale. Cela étant, nous ne sommes pas encore en mesure de justifier pourquoi chaque itération améliore l'estimation de l'autocorrélation, et donc la qualité de l'image agrandie.

6.7 Conclusion

Dans ce chapitre, nous avons d'abord montré comment concevoir des méthodes d'agrandissement cohérentes avec notre modèle de redimensionnement, et ayant des propriétés d'approximation optimales. Malgré le formalisme rigoureux adopté, les résultats des méthodes linéaires sont intrinsèquement limités. Poser l'agrandissement comme un problème pseudo-inverse de la réduction ne garantit que la synthèse cohérente des basses fréquences. Il faut donc se détourner de la linéarité si l'on veut synthétiser correctement l'information géométrique dans l'image agrandie. Les méthodes linéaires pèchent sur ce point, car elles étendent essentiellement à zéro le spectre de l'image initiale.

Nous avons dans la suite de ce chapitre présenté un état de l'art synthétique de l'agrandissement d'images, en mettant en perspective par rapport à notre modèle les principes sous-jacents aux différentes approches. Les méthodes d'interpolation sont basées avant tout sur le simple paradigme de conservation des fréquences, insuffisant pour obtenir une impression de netteté satisfaisante dans l'image agrandie. Les méthodes d'extrapolation partent

d'un modèle plus évolué, mais échouent car les relations unissant les coefficients à extrapoler aux coefficients disponibles dans l'image initiale sont trop complexes. Les méthodes régularisant la contrainte de réduction sont basées sur la minimisation d'un coût, qui doit être suffisamment complexe pour autoriser les discontinuités. Les mises en œuvre itératives lentes et les difficultés d'implémentation desservent ce type d'approches. Il semble difficile de capturer dans un critère analytique explicite la complexité des images, et la minimisation d'un critère générique ne semble pas à même de fournir une image agrandie systématiquement satisfaisante. Les méthodes préservant les structures se fixent des objectifs plus modestes, et utilisent des heuristiques pour fournir des images agrandies visuellement satisfaisantes, en faisant preuve d'un effort particulier pour le rendu des contours obliques. Les méthodes adaptatives, basées sur une détection et un traitement approprié des contours, donnent les résultats les plus satisfaisants.

Nous avons proposé une méthode d'interpolation, appelée WiskI, qui évite particulièrement bien l'écueil de la plupart des méthodes, à savoir le crénelage des contours obliques. Il s'agit cependant d'une méthode d'interpolation, et l'extension à une contrainte de réduction générique, ainsi qu'à des facteurs autres que 2, est en cours d'étude.

On peut être déçu par l'idée qu'aucune méthode d'agrandissement ne combine à l'heure actuelle les avantages particuliers constatés avec telle ou telle approche. Nous sommes optimistes en ce qui nous concerne, et l'obtention d'une méthode vérifiant une contrainte de réduction réaliste, produisant des contours à la fois nets transversalement et dénués de crénelage longitudinalement, nous semble possible. Une telle méthode devrait fonctionner avec un facteur arbitraire, et avoir une implémentation directe. Dans le chapitre suivant, nous proposons l'agrandissement par induction, qui réalise un pas en avant vers un tel compromis, non encore atteint à ce jour.

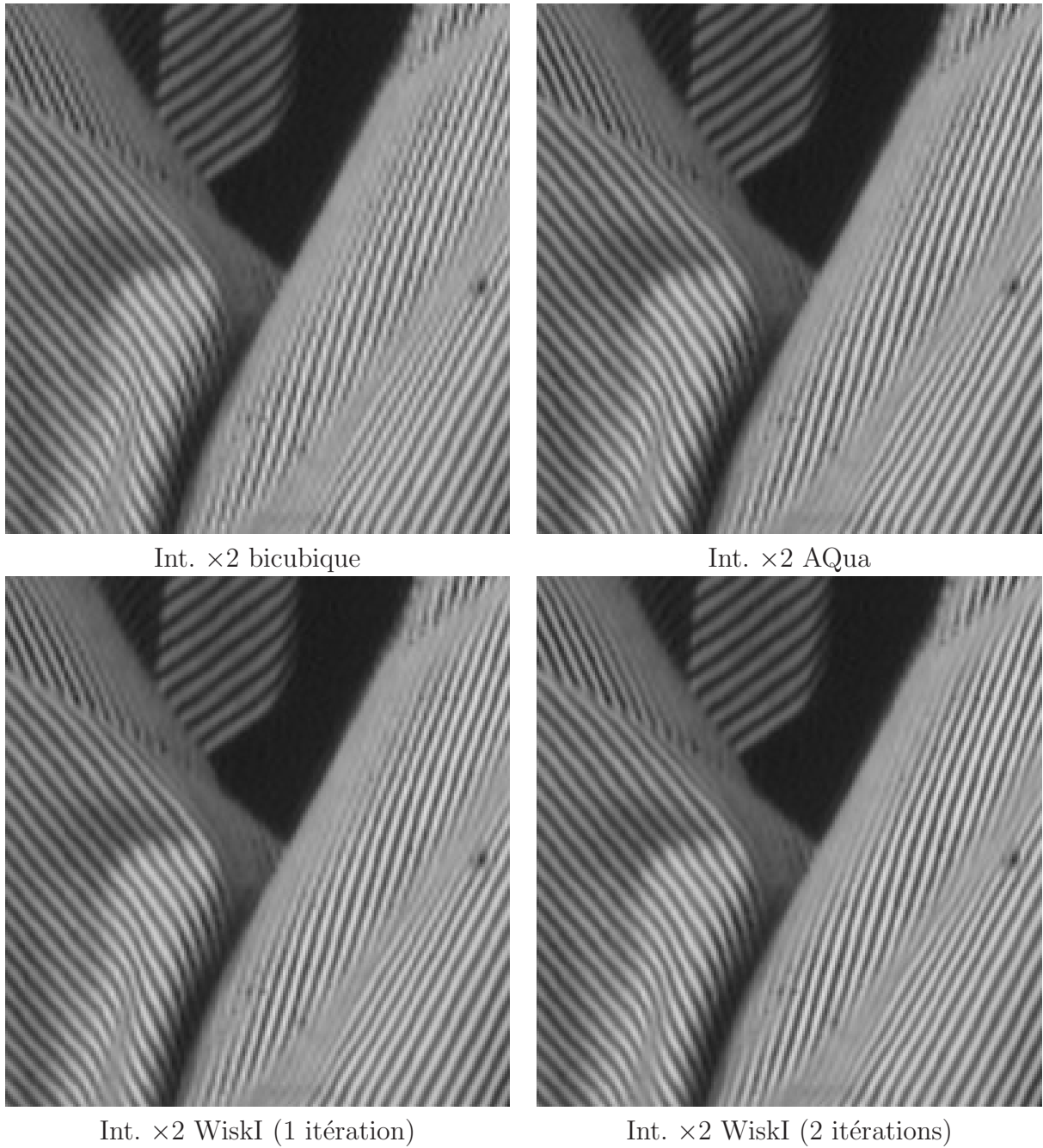


FIG. 6.16 : Illustration sur le pantalon de Barbara de notre méthode d'interpolation WiskI. Noter l'absence totale d'effets d'escaliers sur les contours obliques, dès la deuxième itération. Même avec une itération, le résultat est meilleur que celui de la méthode AQua [167].

Chapitre 7

Agrandissement d'images par induction

7.1 Introduction

LES scènes lumineuses sous-jacentes aux images sont composées de manière essentiellement non linéaire, car elles reposent fondamentalement sur le phénomène d'occlusion : les contours, qui jouent un rôle capital dans le processus de vision, déterminent les frontières d'objets qui se masquent les uns les autres. Les méthodes d'agrandissement linéaires échouent dans le but de capturer cette spécificité des images, car elles ne créent pas les hautes fréquences requises pour la synthèse des contours. Elles reposent sur un postulat de stationnarité qui n'est pas vérifié dans les images naturelles. À l'inverse, les méthodes d'agrandissement non linéaires sont plus efficaces sur ce point, mais elles sont principalement basées sur des heuristiques. La cohérence entre l'image initiale et l'image agrandie n'est assurée généralement que par la contrainte d'interpolation, qui n'est pas appropriée du point de vue de notre modèle de redimensionnement, et qui fournit en pratique des images d'aspect flou ou lissé.

Dans ce chapitre, nous présentons une réponse originale au problème de l'agrandissement sous contrainte de réduction. Nous allons d'abord analyser plus en détail cette contrainte, qui pose le problème de l'agrandissement comme un problème inverse mal posé de la réduction : l'image agrandie doit, si on la réduit, redonner l'image initiale. Les méthodes de la littérature imposent pour la plupart une contrainte de réduction avec un filtre élémentaire : le filtre identité, ou moyennneur 2×2 pour la réduction de facteur 2. De plus, cette contrainte est souvent vérifiée de manière approximative, car elle est difficile à combiner avec les critères de vraisemblance (régularité, cohérence...) antagonistes, imposés par ailleurs. La méthode que nous proposons, appelée *agrandissement par induction*, vérifie la contrainte de réduction avec un filtre arbitraire et de manière exacte, tout en donnant de bons résultats visuels. Cette méthode originale est basée sur un formalisme ensembliste, et marie le meilleur de deux mondes : celui des méthodes linéaires pour la synthèse cohérente

des basses fréquences, et celui des méthodes non linéaires pour le rendu des contours. La philosophie sous-jacente est la suivante : plutôt que de chercher une image solution d'un problème compliqué combinant la contrainte de réduction et un critère de régularisation, on va d'abord agrandir l'image de manière visuellement satisfaisante, avant de forcer *après coup* la contrainte de réduction.

7.2 L'agrandissement : un problème inverse

Nous allons nous placer dans les conditions suivantes, pour l'étude du problème d'agrandissement. Ces trois hypothèses ne seront relaxées que dans la section 7.5 :

- L'image initiale v à agrandir n'est pas bruitée ;
- Le facteur d'agrandissement α est entier ;
- La fonction $\tilde{\varphi}$ vérifie la relation à deux échelles généralisée (5.27) avec le filtre d'échelle $\tilde{h}_{\alpha, -\alpha\tau}$.

Sous ces conditions, comme nous l'avons décrit dans la section 6.2.2, l'image agrandie $y^{\text{id}} = \mathcal{A}_{\alpha, \tau}^{\text{id}} v$, que l'on cherche à estimer, vérifie la contrainte de réduction avec le filtre de réduction $\tilde{h}_{\alpha, -\alpha\tau}$. Notons $\mathcal{R}_{\alpha, -\alpha\tau}$ l'opérateur de réduction avec ce filtre. On a :

$$\mathcal{R}_{\alpha, -\alpha\tau} y^{\text{id}} = v \Leftrightarrow [y^{\text{id}} * \tilde{h}_{\alpha, -\alpha\tau}] \downarrow \alpha = v. \quad (7.1)$$

Nous allons voir, dans la suite de cette section, en quoi cette contrainte implique une connaissance sur l'image agrandie idéale y^{id} , et en quoi elle nous aide à l'estimer.

7.2.1 L'ensemble induit

Etant donnée l'image v , il y a tout un ensemble d'images agrandies, que nous appellerons **l'ensemble induit** (sous-entendu, par v), qui vérifie la contrainte de réduction :

$$\Upsilon_v = \{y \mid \mathcal{R}_{\alpha, -\alpha\tau} y = v\}. \quad (7.2)$$

Cet ensemble est un espace affine : étant donnée une image $y_0 \in \Upsilon_v$, on a

$$\Upsilon_v = y_0 + \Upsilon_0 = \{y + y_0 \mid \mathcal{R}_{\alpha, -\alpha\tau} y = 0\}, \quad (7.3)$$

où l'espace vectoriel Υ_0 est le noyau de l'opérateur $\mathcal{R}_{\alpha, -\alpha\tau}$. L'ensemble Υ_v capture toute l'information que l'on a sur l'image y^{id} : on sait juste de y^{id} qu'elle appartient à cet ensemble. Le problème de l'agrandissement est donc un problème inverse mal posé, au sens de Hadamard [114] : il y a une infinité d'images agrandies donnant v après réduction, et nous n'avons aucune information sur la composante de y^{id} se trouvant dans le noyau de l'opérateur de réduction.

Comme nous l'avons détaillé dans la section (6.2.2), une méthode d'agrandissement linéaire vérifie la contrainte de réduction si et seulement si le filtre d'agrandissement $h_{\alpha,\tau}$ est biorthogonal au filtre de réduction $\tilde{h}_{\alpha,-\alpha\tau}$. Afin d'interpréter ce procédé d'agrandissement, définissons le sous-espace vectoriel de ℓ_2 associé à un filtre quelconque $h \in \ell_2$, et défini par :

$$V_\alpha(h) = \left\{ \sum_{\mathbf{l} \in \mathbb{Z}^2} c[\mathbf{l}] (h[\mathbf{k} - \alpha\mathbf{l}])_{\mathbf{k} \in \mathbb{Z}^2} \mid c \in \ell_2 \right\}. \quad (7.4)$$

$V_\alpha(h)$ est l'espace des combinaisons linéaires des α -translatés de h . Il s'agit de l'équivalent discret des espaces LSI de fonctions que nous avons manipulés pour le problème de reconstruction. On peut par exemple construire les espaces splines discrets $V_\alpha(b_\alpha^n)$, où b_α^n est la B-spline discrète définie par (5.29). $V_\alpha(b_\alpha^n)$ est donc l'espace de tous les signaux discrets que l'on peut obtenir en échantillonnant aux entiers une fonction de $V_\alpha(\beta^n)$, c'est-à-dire une surface spline de degré n avec ses nœuds aux positions $\alpha\mathbf{k}$.

Lorsque l'on définit l'image agrandie par

$$\mathcal{A}_{\alpha,\tau}v = [c] \uparrow \alpha * h_{\alpha,\tau} \quad (7.5)$$

pour une certaine séquence c de coefficients, on cherche une image appartenant à l'espace $V_\alpha(h_{\alpha,\tau})$. L'agrandissement linéaire sous contrainte de réduction revient alors à calculer la **projection oblique** de l'image agrandie idéale y^{id} dans l'espace $V_\alpha(h_{\alpha,\tau})$, parallèlement à l'espace $V_\alpha(\tilde{h}_{\alpha,-\alpha\tau})^\perp$. On voit donc les fortes similitudes qui existent entre les problèmes d'agrandissement et de reconstruction. Le premier est en fait une version discrétisée du second, et la contrainte de réduction est au problème de l'agrandissement ce que la consistance est au problème de la reconstruction. Dans le cas de méthodes linéaires, ces deux contraintes sont équivalentes à la propriété de biorthogonalité : sur les filtres $\tilde{h}$ et h dans le premier cas, sur les fonctions $\tilde{\varphi}$ et φ dans le second cas, comme exprimé respectivement par les formules (6.12), (2.40). On peut aussi faire le lien entre ces deux problèmes en voyant la reconstruction comme un cas limite d'agrandissement : si on itère à l'infini à partir de v l'agrandissement linéaire de filtre $h_{\alpha,\tau}$ vérifiant la contrainte de réduction, on obtient exactement la fonction consistante de l'espace $V_T(\varphi)$, où φ est la fonction limite ayant $h_{\alpha,\tau}$ comme filtre d'échelle :

$$\hat{\varphi}(\omega) = \prod_{i=1}^{+\infty} \frac{1}{\alpha^2} \hat{h}_{\alpha,\tau}\left(\frac{\omega}{\alpha^i}\right). \quad (7.6)$$

Les méthodes asymptotiquement optimales de la section 6.2.3 reviennent, quant à elles, à chercher une meilleure représentation que cette projection oblique, de l'image inconnue y^{id} dans l'espace $V_\alpha(h_{\alpha,\tau})$. En effet, la meilleure représentation de y^{id} dans cet espace est sa projection orthogonale, définie par (preuve dans [221]) :

$$\mathcal{P}_{V_\alpha(h_{\alpha,\tau})}^\perp y^{\text{id}} = [c] \uparrow \alpha * h_{\alpha,\tau} \quad \text{avec } c = [y^{\text{id}} * h_d] \downarrow \alpha, \quad (7.7)$$

où h_d est le dual de $h_{\alpha, \tau}$:

$$h_d = \bar{h}_{\alpha, \tau} * ([h_{\alpha, \tau} * \bar{h}_{\alpha, \tau}] \downarrow \alpha \uparrow \alpha)^{-1}. \quad (7.8)$$

Bien sûr, cette projection orthogonale de y^{id} n'est pas accessible à partir de l'image v . Par contre, les méthodes développées dans la section 6.2.3 effectuent une quasi-projection de y^{id} dans l'espace $V_\alpha(h_{\alpha, \tau})$, asymptotiquement plus proche de la projection orthogonale que la projection oblique atteinte par l'agrandissement sous contrainte de réduction. Cela étant, on ne peut surpasser l'agrandissement sous contrainte de réduction que parce que l'on se restreint à chercher une image appartenant à un espace $V_\alpha(h_{\alpha, \tau})$ fixé, par exemple l'espace spline cubique. Si l'on s'autorise à produire des images agrandies de manière non linéaire, imposer la contrainte de réduction est à nouveau le mieux que l'on puisse faire. Cela se justifie par l'argument suivant : si une image agrandie y^0 ne vérifie pas la contrainte de réduction, on peut obtenir de manière systématique une image plus proche de y^{id} que y^0 , en projetant orthogonalement y^0 dans l'ensemble induit Υ_v . Cela résulte de l'égalité de Pythagore :

$$\|y^{\text{id}} - y^0\|_{\ell_2}^2 = \|y^{\text{id}} - \mathcal{P}_{\Upsilon_v}^\perp y^0\|_{\ell_2}^2 + \|y^0 - \mathcal{P}_{\Upsilon_v}^\perp y^0\|_{\ell_2}^2 \geq \|y^{\text{id}} - \mathcal{P}_{\Upsilon_v}^\perp y^0\|_{\ell_2}^2. \quad (7.9)$$

Cette formalisation de l'agrandissement en termes de projections dans des espaces nous semble intéressante pour mettre en lumière les mécanismes qui rentrent en jeu. Il ne s'agit pas que d'une simple interprétation : le principe même de la méthode d'induction, que nous allons détailler dans la suite de ce chapitre, est d'effectuer une projection orthogonale d'une image agrandie y^0 dans l'ensemble induit, suivant le constat effectué au paragraphe précédent. Nous montrerons le rôle effectif, dans la conception de l'image agrandie, de ces projections, qui peuvent paraître abstraites au premier abord.

7.2.2 L'agrandissement : un problème d'extrapolation

Afin de simplifier les notations, adoptons les écritures h pour $h_{\alpha, \tau}$ et $\tilde{h}$ pour $\tilde{h}_{\alpha, -\alpha\tau}$. Il est intéressant d'écrire la contrainte de réduction en domaine fréquentiel :

$$\hat{v}(\omega) = \sum_{\mathbf{n} \in \llbracket 1, \alpha \rrbracket^2} \hat{y}^{\text{id}}\left(\frac{\omega + 2\pi\mathbf{n}}{\alpha}\right) \hat{h}\left(\frac{\omega + 2\pi\mathbf{n}}{\alpha}\right). \quad (7.10)$$

Ainsi, à chaque fréquence $\omega \in [0, 2\pi/\alpha]^2$, on ne dispose que d'une moyenne pondérée des valeurs $\hat{y}^{\text{id}}(\omega + 2\pi\mathbf{n}/\alpha)$. L'agrandissement consiste donc à re-séparer correctement cette énergie dans les différentes bandes $[2\pi\mathbf{n}/\alpha, 2\pi(\mathbf{n} + \mathbf{1})/\alpha]$.

L'interpolation au moyen du filtre sinc revient à affecter toute l'énergie dans les basses-fréquences :

$$\hat{y}(\omega) = \begin{cases} \hat{v}(\alpha\omega) & \omega \in [0, 2\pi/\alpha]^2 \\ 0 & \omega \in [2\pi\mathbf{n}/\alpha, 2\pi(\mathbf{n} + \mathbf{1})/\alpha] \quad \forall \mathbf{n} \in \llbracket 0, \alpha - 1 \rrbracket^2 \setminus \{\mathbf{0}\} \end{cases} \quad (7.11)$$

On voit que ce choix n'est pas correct : il faudrait, pour avoir $y = y^{\text{id}}$, d'une part que y^{id} soit à bande limitée dans $[0, 2\pi/\alpha]^2$, et d'autre part que le filtre de réduction soit un sinus cardinal discrétisé. Or, c'est justement la présence d'*aliasing* dans l'image v qui rend le problème de l'agrandissement si spécifique : il faut à la fois « désaliaser » l'image v , et déconvoluer le filtre $\tilde{h}$.

Plutôt que d'écrire la contrainte de réduction dans le domaine fréquentiel, où l'agrandissement n'apparaît pas comme un problème d'extrapolation, nous allons l'expliciter dans le domaine *ondelettes*. Remarquons tout d'abord que le noyau de l'opérateur de réduction est $V_\alpha(\tilde{h})^\perp$, l'orthogonal de $V_\alpha(\tilde{h})$ dans $\ell_2(\mathbb{Z}^2)$. Il est possible de déterminer un ensemble de filtres passe-haut $g_{\mathbf{n}}$, $\mathbf{n} \in \llbracket 0, \alpha - 1 \rrbracket^2 \setminus \{\mathbf{0}\}$ tels que l'on ait la somme directe [234, 128, 187]

$$V_\alpha(\tilde{h})^\perp = \bigoplus_{\mathbf{n} \in \llbracket 0, \alpha - 1 \rrbracket^2 \setminus \{\mathbf{0}\}} V_\alpha(g_{\mathbf{n}}). \quad (7.12)$$

Ainsi, on a existence et unicité de séquences de coefficients d'ondelettes $c_{\mathbf{n}}$, $\mathbf{n} \in \llbracket 0, \alpha - 1 \rrbracket^2 \setminus \{\mathbf{0}\}$ telles que

$$y^{\text{id}} = [v] \uparrow \alpha * h + \sum_{\mathbf{n} \in \llbracket 0, \alpha - 1 \rrbracket^2 \setminus \{\mathbf{0}\}} [c_{\mathbf{n}}] \uparrow \alpha * g_{\mathbf{n}}. \quad (7.13)$$

Si par exemple l'image v est de taille N^2 , l'image agrandie de facteur 2 y^{id} est de taille $4N^2$. Elle peut être décomposée en N^2 coefficients d'échelle, qui sont en fait les pixels $v[\mathbf{k}]$, et $3N^2$ coefficients d'ondelettes $c_{\mathbf{n}}[\mathbf{k}]$ qu'il s'agit d'estimer. Sous cette forme, l'agrandissement apparaît bien comme un problème d'extrapolation : la composante basses fréquences v de y^{id} est connue, et on cherche les coefficients d'ondelettes $c_{\mathbf{n}}[\mathbf{k}]$ permettant de reconstruire le reste du contenu fréquentiel de y^{id} . Lorsque l'on fait l'agrandissement linéaire $y = [v] \uparrow \alpha * h$, cela revient à mettre à zéro tous ces coefficients d'ondelettes, ce qui n'est pas satisfaisant. L'agrandissement consiste donc à extrapoler les coefficients d'ondelettes $c_{\mathbf{n}}[\mathbf{k}]$ à partir des coefficients d'échelle $v[\mathbf{k}]$. Cela n'est possible que si ces coefficients ne sont pas indépendants. En calculant la transformée en ondelettes d'images naturelles, on constate effectivement que des structures similaires apparaissent aux différentes échelles. Il est important de souligner que ces dépendances inter-échelles sont essentiellement non linéaires. S'il subsiste une corrélation linéaire entre les échelles, c'est que le filtre d'agrandissement h a été mal choisi. Nous ne nous sommes pas attardés, au chapitre précédent, sur le choix du filtre h , ou plus exactement de l'espace $V_\alpha(h)$. Comme pour le problème de reconstruction, on peut simplement faire la constatation suivante : si l'on considère que y^{id} est la réalisation d'un processus d'autocorrélation connue, on peut exhiber le filtre d'agrandissement optimal (6.32). Il est avéré que les images naturelles ont un comportement spectral en $1/f$, proche de celui d'un mouvement brownien fractionnaire [214]. Dans ce cas, le filtre d'agrandissement le meilleur est le filtre d'échelle de la spline polyharmonique [235] consistante avec v . En pratique, nous choisissons $V_\alpha(h)$ comme étant un espace spline séparable cubique discret (c.-à-d. $V_\alpha(b_\alpha^3)$ lorsque $\boldsymbol{\tau} = \mathbf{0}$). Nous considérons, de manière empirique, qu'il s'agit

du meilleur choix pour effectuer un agrandissement linéaire, et exploiter au mieux *de manière linéaire* l'information contenue dans l'image v . Ainsi, on peut donner l'interprétation suivante pour la formule (7.13) : l'image agrandie $y = [v] \uparrow \alpha * h$ représente la partie de y^{id} qu'il est possible de synthétiser linéairement à partir de v . Le résidu, dissocié de y , contient la partie manquante de y^{id} , intrinsèque à la résolution de l'image agrandie, et qui ne peut plus être prédit linéairement. La synthèse des coefficients d'ondelettes de y^{id} est donc un problème essentiellement **non linéaire**.

Bien que l'agrandissement (de facteur entier) puisse être vu comme un problème d'extrapolation de coefficients d'ondelettes, la nature non linéaire du problème rend délicate la synthèse de ces coefficients à partir de v . La méthode d'agrandissement par induction que nous allons développer dans la suite de ce chapitre effectue cette extrapolation de manière implicite, en la « déléguant » à une autre méthode d'agrandissement non linéaire.

7.3 L'induction de Didier Calle

L'agrandissement par induction est une méthode qui régularise une image agrandie visuellement satisfaisante, obtenue au moyen d'une méthode préservant les structures, mais qui ne vérifie pas la contrainte de réduction. Cette image n'est pas cohérente avec l'image initiale au sens de notre modèle; l'induction vise donc à lui faire subir un traitement permettant à la contrainte de réduction d'être satisfaite, tout en conservant les qualités visuelles de l'image.

L'image agrandie à régulariser est appelée *image inductrice*. On la notera y^0 . La méthode d'induction proposée par Didier Calle en 1999 dans son travail de thèse [50, 49], consiste à définir l'image agrandie par induction, appelée *image induite* et notée y^∞ , comme la projetée orthogonale de y^0 dans l'ensemble induit Υ_v :

$$\boxed{y^\infty = \mathcal{P}_{\Upsilon_v}^\perp y^0}. \quad (7.14)$$

L'équation (7.9) justifie cette définition : on est ainsi sûr d'obtenir après induction une image plus proche de l'image idéale y^{id} , quelle que soit l'image inductrice de départ y^0 . Ainsi, l'image induite est l'image de l'ensemble induit la plus proche de l'image inductrice, au sens ℓ_2 :

$$y^\infty = \underset{y \in \Upsilon_v}{\operatorname{argmin}} \|y - y^0\|_{\ell_2}. \quad (7.15)$$

La distance ℓ_2 étant, faute de mieux, un indicateur acceptable de la distance visuelle entre les images¹, on peut considérer que l'image induite est l'image la plus proche visuellement de l'image inductrice, vérifiant la contrainte de réduction.

1. Cette affirmation doit être relativisée, car la distance ℓ_1 semble plus appropriée que la distance ℓ_2 pour comparer les images [163]. Cependant, seule cette dernière donne lieu à un problème linéaire lorsqu'on la minimise. Elle reste donc la seule utilisable en pratique.

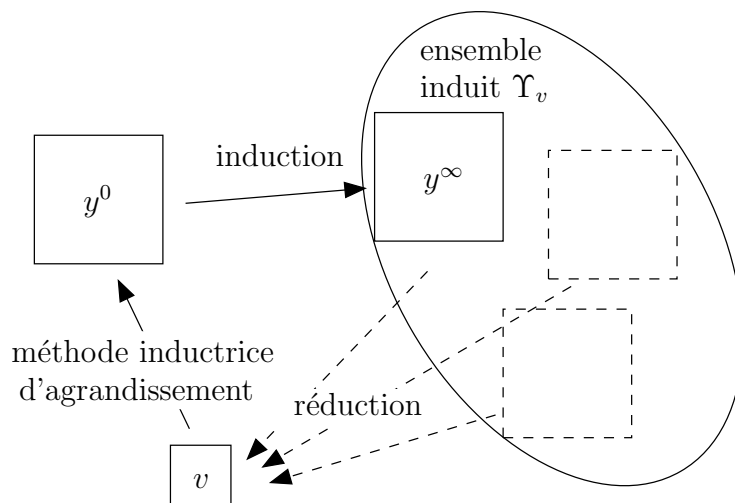


FIG. 7.1 : Principe de l'agrandissement par induction d'une image v . L'induction consiste à projeter orthogonalement une image agrandie y^0 ne vérifiant pas la contrainte de réduction, dans l'ensemble induit Υ_v , défini comme l'ensemble de toutes les images vérifiant cette contrainte, c'est-à-dire l'ensemble des images qui, lorsqu'on les réduit par le procédé $\mathcal{R}_{\alpha, -\alpha\tau}$, produisent l'image v .

L'induction agit comme un post-traitement sur l'image inductrice y^0 afin de la régulariser, pour la rendre cohérente avec notre modèle de redimensionnement, au travers de la contrainte de réduction qui lui est imposée. Ce procédé, qui peut être appliqué à n'importe quelle image inductrice, est illustré sur la figure 7.1. La figure 7.2 montre les images induites y^∞ que l'on obtient lorsque l'on agrandit l'image *Camera*, en prenant différentes images inductrices y^0 choisies de manière naïve. Les hautes fréquences de l'image inductrice sont préservées dans l'image induite ; si elles ne sont pas en phase avec le contenu de l'image v que l'on veut agrandir, elle apparaissent en filigrane dans l'image induite, de manière bien visible. Ainsi, le choix de l'image inductrice est crucial, comme on va le détailler par la suite, pour la qualité de l'image induite obtenue. Naturellement, le procédé d'induction n'est intéressant² que si l'image inductrice a été obtenue à partir de l'image initiale v au moyen d'une certaine méthode d'agrandissement, que nous appellerons *méthode inductrice*. Dans le cas contraire, il n'y a aucun espoir d'approcher correctement l'image y^{id} que l'on cherche à estimer.

Au départ, Didier Calle a développé l'induction afin de régulariser sa propre méthode inductrice, de type zoom fractal [49].

Intéressons-nous à la mise en œuvre pratique de l'induction. Afin d'effectuer cette projection orthogonale dans l'ensemble induit, Didier Calle utilise des méthodes itératives de

2. Du moins pour l'agrandissement. On pourrait imaginer des applications, comme le tatouage, où l'ajout en filigrane d'information hautes fréquences indépendante du contenu de l'image puisse être intéressante.

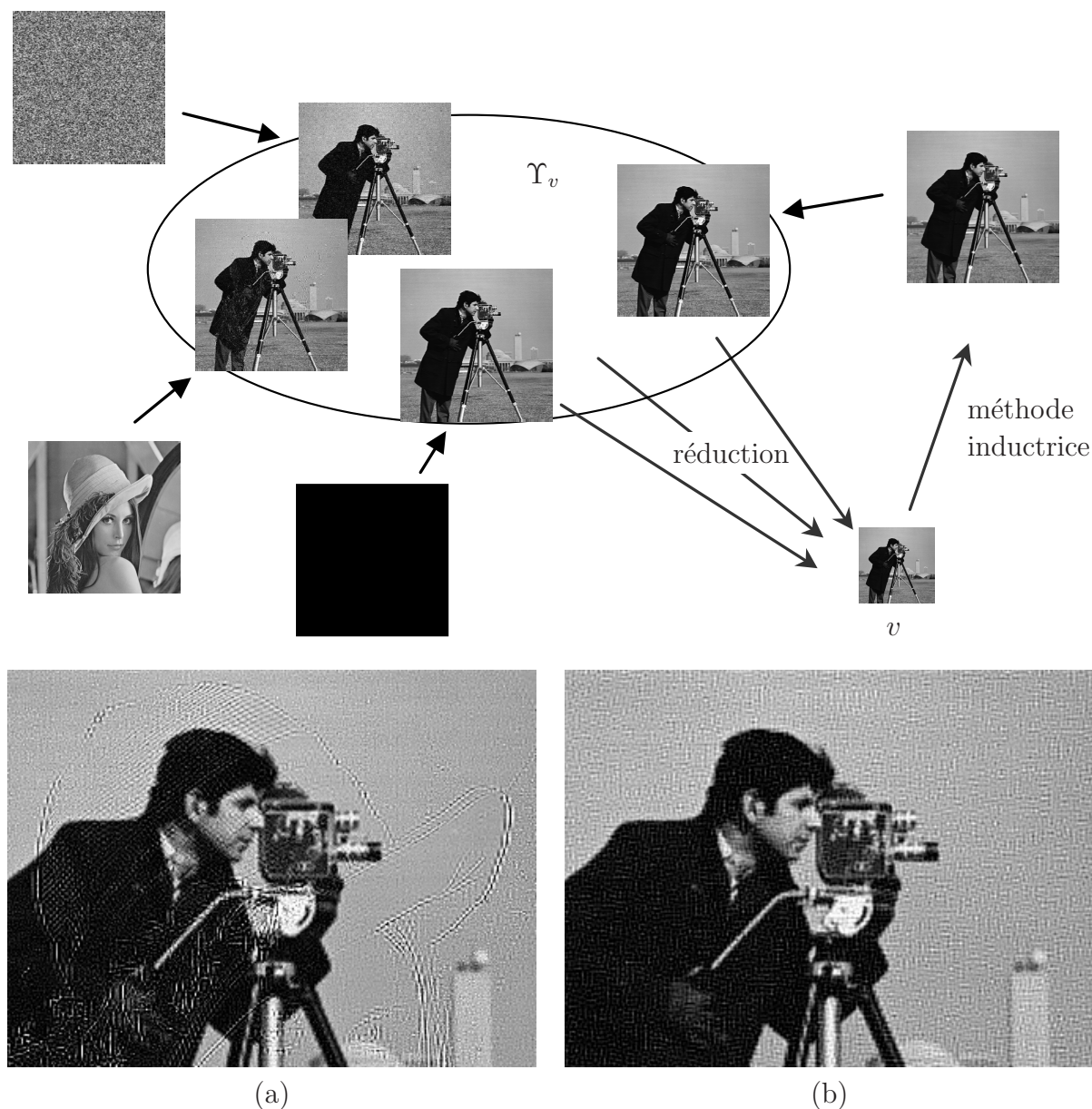


FIG. 7.2 : L'ensemble induit contient une infinité d'images, pouvant être assez éloignées de l'image idéale y^{id} . Il va nous falloir trouver une « bonne » méthode inductrice, en un sens à préciser, fournissant une image inductrice y^0 qui, après induction, soit visuellement satisfaisante. Un choix naïf d'image inductrice donne des résultats peu convaincants. Par exemple, si l'image v est *Camera* et que l'on considère *Lena* comme inductrice, l'image induite contient bien le cameraman, mais avec les contours de Lena en filigrane (a). En (b), l'image induite en prenant comme inductrice une image de bruit : le cameraman est bien présent, mais avec du bruit en surimpression.

projection dans des ensembles convexes (POCS) [247]. En effet, chaque pixel de l'image initiale définit une contrainte linéaire sur l'image agrandie : on peut écrire

$$\Upsilon_v = \bigcap_{\mathbf{k}} \Upsilon_{v,\mathbf{k}}, \quad (7.16)$$

où

$$\Upsilon_{v,\mathbf{k}} = \{y \mid \mathcal{R}_{\alpha, -\alpha\tau} y[\mathbf{k}] = v[\mathbf{k}]\} \quad (7.17)$$

est un hyperplan affine, donc un ensemble convexe élémentaire, sur lequel la projection orthogonale s'exprime simplement :

$$\mathcal{P}_{\Upsilon_{v,\mathbf{k}}}^\perp u[\mathbf{l}] = u + \frac{v[\mathbf{k}] - \sum_{\mathbf{n} \in \mathbb{Z}^2} u[\mathbf{n}] \tilde{h}[\alpha\mathbf{k} - \mathbf{n}]}{\sum_{\mathbf{n} \in \mathbb{Z}^2} \tilde{h}[\mathbf{n}]^2} \tilde{h}[\alpha\mathbf{k} - \mathbf{l}] \quad \forall \mathbf{l} \in \mathbb{Z}^2. \quad (7.18)$$

Ainsi, sous sa forme originelle, l'induction est une méthode itérative. La méthode itérative la plus simple, dite des *projections alternées* est la suivante : lors de la $n + 1$ -ième itération l'image agrandie y^n est mise à jour par projections successives sur chaque ensemble $\Upsilon_{v,\mathbf{k}}$: pour une image de taille $N \times N$, et à partir de l'image inductrice y^0 ,

$$y^{n+1} = \mathcal{P}_{\Upsilon_{v,[N,N]}}^\perp \circ \cdots \circ \mathcal{P}_{\Upsilon_{v,[1,1]}}^\perp y^n. \quad (7.19)$$

On peut montrer que cette méthode converge, lorsque $n \rightarrow +\infty$, vers l'image induite y^∞ , projection orthogonale de y^0 dans l'ensemble induit. La convergence peut être significativement améliorée en remplaçant les projections orthogonales successives par des projections obliques plus élaborées, avec des facteurs de *damping* [50].

La première amélioration que nous avons apportée à l'induction est une mise en œuvre directe, non itérative, qui fournit l'image induite de manière exacte, et beaucoup plus rapidement que les méthodes itératives proposées par Didier Calle [50, 49]. Nous décrivons maintenant cette *induction rapide*.

7.4 L'induction oblique rapide

7.4.1 L'induction rapide

En étudiant les formules (7.18) et (7.19), on constate que l'image induite y^∞ est obtenue par ajout sur y^0 d'une succession de corrections $e^n = \underline{y}^{n+1} - y^n$, qui s'expriment toutes comme combinaisons linéaires des $\alpha\mathbf{k}$ -translatés du filtre $\tilde{h}$, le « retourné » de $\tilde{h}$. En d'autres termes, la différence $e = y^\infty - y^0$ appartient à l'espace $V_\alpha(\tilde{h})$. Il existe donc une séquence de coefficients $(c[\mathbf{k}])_{\mathbf{k} \in \mathbb{Z}^2}$ tels que

$$y^\infty - y^0 = [c] \uparrow \alpha * \tilde{h}. \quad (7.20)$$

Afin d'expliciter c , il suffit d'utiliser le fait que y^∞ vérifie la contrainte de réduction :

$$[y^\infty * \tilde{h}] \downarrow \alpha = v \quad (7.21)$$

$$\Leftrightarrow [(y^0 + [c] \uparrow \alpha * \tilde{h}) * \tilde{h}] \downarrow \alpha = v \quad (7.22)$$

$$\Leftrightarrow [y^0 * \tilde{h}] \downarrow \alpha + c * [\tilde{h} * \tilde{h}] \downarrow \alpha = v \quad (7.23)$$

$$\Leftrightarrow c = (v - [y^0 * \tilde{h}] \downarrow \alpha) * ([\tilde{h} * \tilde{h}] \downarrow \alpha)^{-1}. \quad (7.24)$$

On a donc la forme explicite de l'image induite y^∞ :

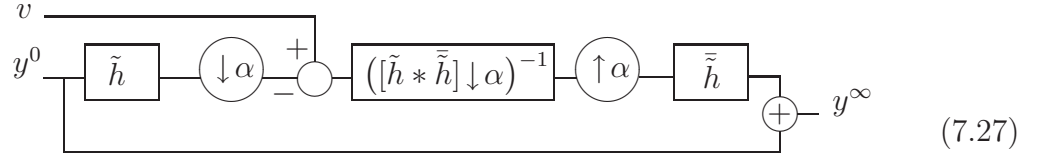
$$y^\infty = y^0 + [v - [y^0 * \tilde{h}] \downarrow \alpha] \uparrow \alpha * \tilde{h}_d, \quad (7.25)$$

où $\tilde{h}_d$ est le filtre dual de $\tilde{h}$:

$$\tilde{h}_d = \tilde{\tilde{h}} * ([\tilde{h} * \tilde{\tilde{h}}] \downarrow \alpha \uparrow \alpha)^{-1}. \quad (7.26)$$

Ainsi, l'induction consiste à rétro-propager dans l'inductrice y^0 la différence entre v et la version réduite de y^0 . La clé de la méthode, qui permet d'obtenir l'image induite en une seule étape, est le filtrage inverse $([\tilde{h} * \tilde{\tilde{h}}] \downarrow \alpha)^{-1}$, qui compense lors du processus de correction la corrélation entre $\tilde{h}$ et ses $\alpha \mathbf{k}$ -translatés.

L'implémentation directe de l'induction se fait donc comme indiqué par le schéma (*flow-graph*) suivant :



7.4.2 L'induction oblique

L'induction peut être interprétée de deux manières différentes. On peut en premier lieu voir ce processus comme un traitement sur l'image inductrice y^0 pour la rendre cohérente avec l'image v . Si maintenant on écrit (7.20) sous la forme

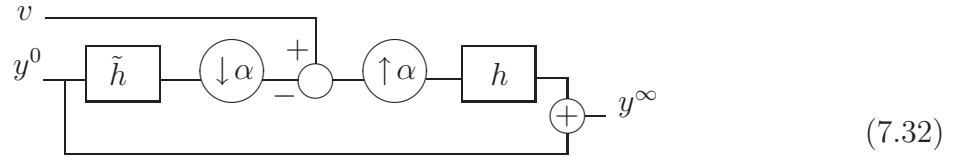
$$y^\infty = \underbrace{([v] \uparrow \alpha * \tilde{h}_d)}_{y_L^\infty} + \underbrace{(y^0 - [y^0 * \tilde{h}] \downarrow \alpha \uparrow \alpha * \tilde{h}_d)}_{y_H^\infty} \quad (7.28)$$

on voit que l'image induite est la somme de deux images apportant chacune une contribution spécifique :

- L'image $[v] \uparrow \alpha * \tilde{h}_d$, que nous notons y_L^∞ , dépend de v mais pas de y^0 . Elle fait donc partie de toute image induite. C'est par cette image que s'exprime la contrainte de réduction. On peut ainsi caractériser l'ensemble induit par

$$\Upsilon_v = y_L^\infty + \Upsilon_0, \quad (7.29)$$

où Υ_0 est le noyau de l'opérateur de réduction, qui ne dépend pas de v .



Par exemple, si l'on choisit un espace spline de degré n pour $V_\alpha(h)$, cela correspond formellement à

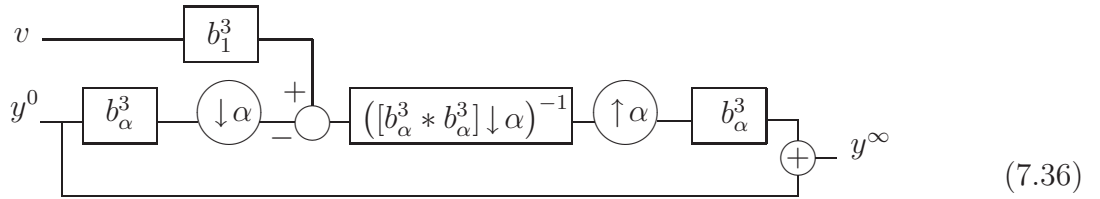
$$h = b_\alpha^n * ([b_\alpha^n * \tilde{h}] \downarrow \alpha \uparrow \alpha)^{-1}. \quad (7.33)$$

De fait, si $\tilde{h}$ est le filtre spline cardinal cubique discret (5.28) préconisé pour la réduction, et que l'on choisit l'espace d'agrandissement spline cubique, l'induction est en fait orthogonale. On a alors

$$\tilde{h} = b_\alpha^3 * [(b_1^3)^{-1}] \uparrow \alpha \quad (7.34)$$

$$h = [b_1^3 * ([b_\alpha^3 * b_\alpha^3] \downarrow \alpha)^{-1}] \uparrow \alpha * b_\alpha^3. \quad (7.35)$$

Ce sont ces filtres que nous utilisons en pratique pour l'induction de facteur entier, d'offset $\tau = \mathbf{0}$. L'implémentation correspondante s'effectue exactement comme suit :



Si l'on veut effectuer l'agrandissement avec un offset différent de $\mathbf{0}$, il suffit de considérer les filtres B-splines discrets translatés, et de substituer à $b_\alpha^3[\mathbf{k}]$ la valeur $b_{\alpha,\tau}^3[\mathbf{k}] = \beta^3(\frac{\mathbf{k}}{\alpha} + \tau)$.

D'autre part, la complexité algorithmique est réduite si les filtres $\tilde{h}$ et h sont tous les deux RIF, ce qui permet de se passer de filtrage inverse. On peut ainsi prendre $\tilde{h} = D9$ et $h = D7$, la paire de filtres de Cohen-Daubechies-Fauveau [83] rendue populaire par ses performances en compression [19] et son emploi dans le standard JPEG2000 [208]. On obtient alors des résultats visuellement identiques à ceux obtenus avec les filtres splines précédents.

7.4.3 Interprétation dans le domaine ondelettes

L'idée sous-jacente à l'induction, que nous avons mise en lumière dans la section précédente, est la suivante : l'image v est d'abord agrandie linéairement du mieux que l'on peut, en respectant la contrainte de réduction, puis l'on fait appel à une méthode inductrice extérieure afin de créer le contenu hautes fréquences y_H^∞ manquant pour un bon rendu

de l'image agrandie. Si l'on regarde à nouveau la figure 7.2, on voit bien que les images induites comportent toutes comme composante principale l'image y_L^∞ , version agrandie de l'image v Camera, ainsi que y_H^0 , qui contient les hautes-fréquences de l'inductrice y^0 . Ainsi, l'image y_H^0 est la partie intrinsèquement non-linéaire des détails de y^0 , qui ne peut être prédite linéairement à partir de la résolution inférieure. L'induction repose sur le postulat suivant : si l'image y^0 est visuellement satisfaisante, ses hautes-fréquences estiment bien celles recherchées de y^{id} , c'est-à-dire $y_H^{\text{id}} \approx y_H^0$.

Plutôt que de raisonner en termes de basses et hautes fréquences, ce qui ne fournit qu'une interprétation qualitative, il vaut mieux considérer les décompositions en ondelettes associées aux images agrandies. Ainsi, le procédé d'induction consiste exactement à considérer que les coefficients d'ondelettes à la résolution de l'image agrandie, que l'on cherche à extrapoler, sont ceux de l'image inductrice. Dans les notations de la section 7.2.2, si l'on décompose l'image inductrice en

$$y^0 = [c_L] \uparrow \alpha * h + \sum_{\mathbf{n} \in \llbracket 0, \alpha-1 \rrbracket^2 \setminus \{\mathbf{0}\}} [c_{\mathbf{n}}] \uparrow \alpha * g_{\mathbf{n}}, \quad (7.37)$$

alors l'image induite est définie comme

$$y^\infty = [v] \uparrow \alpha * h + \sum_{\mathbf{n} \in \llbracket 0, \alpha-1 \rrbracket^2 \setminus \{\mathbf{0}\}} [c_{\mathbf{n}}] \uparrow \alpha * g_{\mathbf{n}}. \quad (7.38)$$

On peut donc soit interpréter l'induction comme régularisant l'image inductrice y^0 , en remplaçant ses coefficients d'échelle $c_L[\mathbf{k}]$ par les valeurs $v[\mathbf{k}]$ qui forment l'image initiale ; soit voir l'image induite comme la somme de $y_L^\infty = [v] \uparrow \alpha * h$ et de l'image y_H^∞ synthétisée à partir des coefficients d'ondelettes $c_{\mathbf{n}}[\mathbf{k}]$ de l'image inductrice. Cette vision de l'induction en termes de coefficients d'ondelettes est schématisée sur la figure 7.4.

Notons qu'en pratique, il n'y a pas lieu d'effectuer réellement une transformée en ondelettes pour mettre en œuvre l'induction, et les implémentations par filtrage que nous avons proposées sont les plus efficaces.

Ainsi, l'agrandissement par induction peut être vu comme une méthode d'extrapolation implicite de coefficients d'ondelettes, ceux-ci étant extraits d'une image agrandie par une méthode d'agrandissement inductrice. On demande avant tout à cette dernière d'effectuer un traitement satisfaisant des contours, avec un rendu franc et sans crénelage. La figure 7.5 illustre la création de coefficients d'ondelettes lors de l'agrandissement par induction. Le rôle de la méthode inductrice est primordial, et il faut la choisir avec soin. On peut faire les remarques suivantes sur l'information apportée par l'image inductrice :

- Si l'image inductrice est uniforme (par exemple complètement noire), aucune haute fréquence n'est apportée et $y^\infty = y_L^\infty$.
- Si l'image inductrice est agrandie linéairement à partir de v : $y^0 = [v] \uparrow \alpha * h^0$, le procédé d'induction devient complètement linéaire, et l'image induite vaut alors $y^\infty = [v] \uparrow \alpha * h^\infty$, pour un certain filtre y^∞ combinaison de h et h^0 .

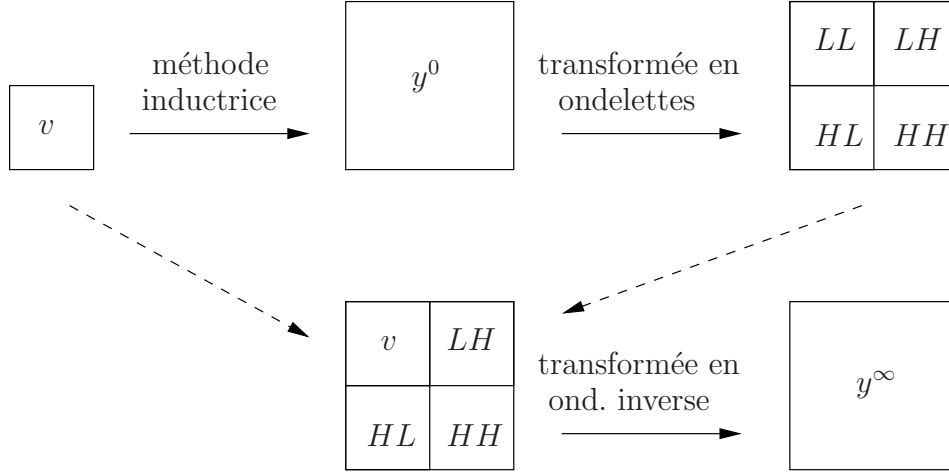


FIG. 7.4 : Interprétation du processus d'induction en termes d'ondelettes : l'image induite contient v comme coefficients d'échelle, et les coefficients d'ondelettes de l'image inductrice.

L'induction ne présente donc d'intérêt que si la méthode inductrice est non linéaire. Nous montrerons dans la section 7.6 que la méthode de Jensen et coll. [123] donne de bons résultats en pratique.

7.5 Extensions de l'induction

7.5.1 Induction relaxée

La méthode d'agrandissement par induction que nous avons présentée vérifie parfaitement la contrainte de réduction, avec un filtre de réduction arbitraire. Il y a cependant de nombreux cas où imposer rigoureusement cette contrainte peut sembler inadéquat. Le cas le plus typique est celui où l'image v contient du bruit. Si la méthode inductrice prend en compte la présence de bruit et fait son possible pour l'éliminer, l'induction va annuler en partie ce travail en forçant l'image induite à être cohérente avec l'image initiale bruitée, et donc en réinjectant du bruit dans l'image induite. Notons qu'il y a toujours du bruit dans les images, ne serait-ce que le bruit de quantification aux entiers des valeurs des pixels. Dans un autre registre, la contrainte de réduction peut être relaxée s'il règne une incertitude sur cette contrainte. Ainsi, le filtre $\tilde{h}$ n'est généralement pas connu. Même si la fonction $\tilde{\varphi}$ caractérisant le processus d'acquisition de v est connue, il y a peu de chances qu'elle vérifie exactement une relation à deux échelles de filtre $\tilde{h}$. Dans tous ces cas, on peut chercher un compromis entre la contrainte de réduction d'une part (proximité de y^∞ à Υ_v), et la proximité de l'image induite à l'image inductrice. Il suffit pour cela d'introduire un facteur $\lambda \in [0, 1]$ de relaxation, et de définir l'image induite relaxée par

$$y_\lambda^\infty = \lambda y^\infty + (1 - \lambda) y^0. \quad (7.39)$$

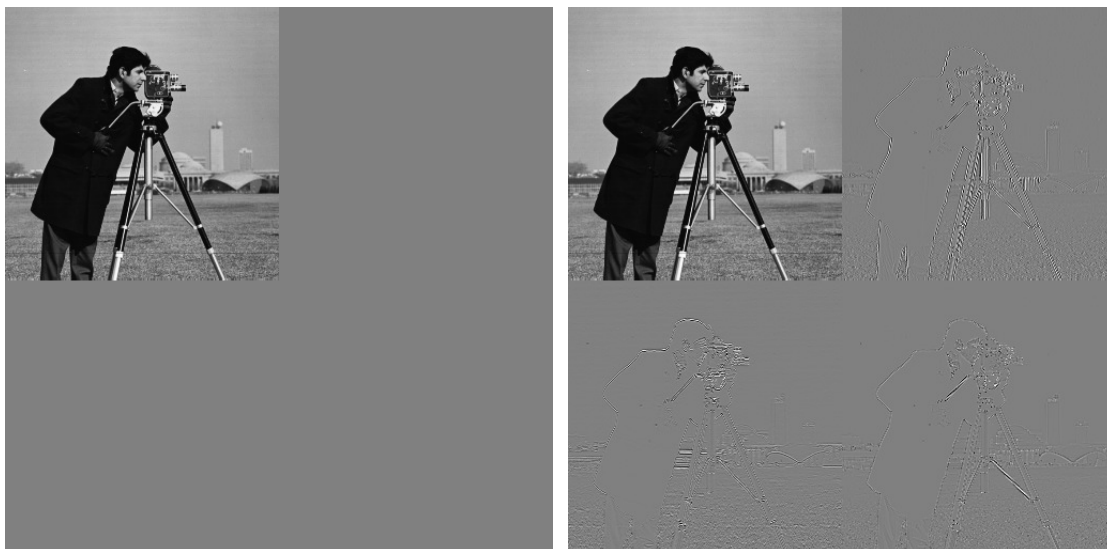


FIG. 7.5 : Décomposition en ondelettes de l'image *Camera* agrandie d'un facteur 2 par induction. L'inductrice est : à gauche, nulle, à droite, obtenue par la méthode de Ramponi et coll. [185].

Ainsi, lorsque $\lambda = 1$, on a confiance dans la contrainte de réduction et on l'impose exactement. Lorsque $\lambda = 0$, l'induction ne fait rien et on s'en remet à la méthode inductrice pour effectuer l'agrandissement. L'implémentation se fait tout simplement en multipliant par λ le terme correctif apporté à l'image inductrice :

$$\begin{array}{c}
 v \\
 \downarrow \\
 y^0 \rightarrow \boxed{\tilde{h}} \rightarrow \textcircled{\downarrow \alpha} \rightarrow \textcircled{+} \rightarrow \textcircled{\uparrow \alpha} \rightarrow \boxed{\lambda h} \rightarrow \textcircled{+} \rightarrow y^\infty
 \end{array}
 \quad (7.40)$$

7.5.2 Induction de facteur quelconque

Nous avons présenté l'induction avec un facteur α entier, par soucis de simplicité, mais la méthode peut être étendue à un facteur α quelconque. La mise en œuvre devient cependant plus complexe. Comme nous l'avons décrit au chapitre 5, la réduction de facteur quelconque peut être effectuée de manière linéaire, bien que cette méthode ne trouve plus de justification théorique dans notre modèle de redimensionnement. La méthode que nous avons proposée est la suivante : pour une certaine fonction $\tilde{h}(\mathbf{x})$,

$$\mathcal{R}_{\alpha, -\alpha\tau} v[\mathbf{k}] = \sum_{\mathbf{l} \in \mathbb{Z}^2} \tilde{h}\left(\mathbf{k} - \tau - \frac{\mathbf{l}}{\alpha}\right) v[\mathbf{l}] \quad \forall \mathbf{k} \in \mathbb{Z}^2. \quad (7.41)$$

De même, l'agrandissement linéaire de facteur α s'écrit, pour une certaine fonction $h(\mathbf{x})$ et une suite de coefficients $c[\mathbf{l}]$ obtenus linéairement à partir des pixels $v[\mathbf{l}]$,

$$\mathcal{A}_{\alpha,\tau}v[\mathbf{k}] = \sum_{\mathbf{l} \in \mathbb{Z}^2} h\left(\frac{\mathbf{k}}{\alpha} + \tau - \mathbf{l}\right) c[\mathbf{l}] \quad \forall \mathbf{k} \in \mathbb{Z}^2. \quad (7.42)$$

Définissons maintenant l'espace discret $\bigvee_{\alpha,\tau}(g) \subset \ell_2$ associé à une fonction $g(\mathbf{x})$, et défini comme l'ensemble des signaux $u \in \ell_2$ s'écrivant sous la forme

$$u[\mathbf{k}] = \sum_{\mathbf{l} \in \mathbb{Z}^2} c[\mathbf{l}] g\left(\frac{\mathbf{k}}{\alpha} + \tau - \mathbf{l}\right) \quad \forall \mathbf{k} \in \mathbb{Z}^2, \quad (7.43)$$

pour une certaine séquence de coefficients $(c[\mathbf{l}])_{\mathbf{l} \in \mathbb{Z}^2} \in \ell_2$. $\bigvee_{\alpha,\tau}(g)$ est la version discrétisée de l'espace LSI de fonctions $V_\alpha(g)$. Contrairement au cas α entier, où l'espace discret $V_\alpha(h)$ s'exprime à partir des $\alpha\mathbf{k}$ translatés du filtre discret h , $\bigvee_{\alpha,\tau}(h)$ associé à la fonction $h(\mathbf{x})$ est ici l'ensemble des combinaisons linéaires de filtres discrets tous distincts, obtenus en échantillonnant $h(\mathbf{x})$ avec un pas de $1/\alpha$, mais des phases distinctes.

Si l'on effectue l'agrandissement linéaire décrit par (7.42), on obtient une image $y = \mathcal{A}_{\alpha,\tau}v$ qui appartient à l'espace $\bigvee_{\alpha,\tau}(h)$. Quel que soit α , on peut définir l'ensemble induit comme précédemment :

$$\Upsilon_v = \{y \mid \mathcal{R}_{\alpha,-\alpha\tau}y = v\}. \quad (7.44)$$

Dès lors, l'image induite y^∞ à partir de l'image inductrice y^0 est définie comme **la projection oblique de l'image inductrice dans l'ensemble induit, parallèlement à $\bigvee_{\alpha,\tau}(h)$** . Si les coefficients $c[\mathbf{l}]$ dans (7.42) sont définis de manière à ce que l'agrandissement $\mathcal{A}_{\alpha,\tau}$ vérifie la contrainte de réduction, on a, comme dans le cas où α est entier :

$$y^\infty = \underbrace{\mathcal{A}_{\alpha,\tau}v}_{y_L^\infty} + \underbrace{(y^0 - \mathcal{A}_{\alpha,\tau}\mathcal{R}_{\alpha,-\alpha\tau}y^0)}_{y_H^\infty = y_H^0} \quad (7.45)$$

$$= y^0 + \mathcal{A}_{\alpha,\tau}(v - \mathcal{R}_{\alpha,-\alpha\tau}y^0). \quad (7.46)$$

L'image $y_L^\infty = y_L^{\text{id}}$ est la *projection oblique* de y^{id} dans $\bigvee_{\alpha,\tau}(h)$ parallèlement à l'espace $\bigvee_{\alpha,\tau}(\tilde{h})^\perp$. Toute la difficulté consiste à déterminer cette projection oblique, autrement dit à calculer les coefficients $c[\mathbf{l}]$ dans (7.42) pour que l'opérateur d'agrandissement $\mathcal{A}_{\alpha,\tau}$ vérifie la contrainte de réduction. Ces coefficients sont obtenus par filtrage inverse à partir de v . Lorsque α est entier, ce filtrage est simplement l'inverse de la corrélation entre les filtres discrets h et $\tilde{h}$, afin de forcer la biorthogonalité de la combinaison réduction/agrandissement : $c = v * ([h * \tilde{h}] \downarrow \alpha)^{-1}$. Dans le cas où α n'est pas un entier, il s'agit d'un filtrage non stationnaire, *spatialement variant*. En effet,

$$\mathcal{R}_{\alpha,-\alpha\tau}\mathcal{A}_{\alpha,\tau}u = u \quad \forall u \in \ell_2 \quad (7.47)$$

$$\Leftrightarrow \sum_{\mathbf{j} \in \mathbb{Z}^2} \tilde{h}(\mathbf{k} - \tau - \frac{\mathbf{j}}{\alpha}) \sum_{\mathbf{l} \in \mathbb{Z}^2} h\left(\frac{\mathbf{j}}{\alpha} + \tau - \mathbf{l}\right) c[\mathbf{l}] = v[\mathbf{k}] \quad \forall \mathbf{k} \in \mathbb{Z}^2. \quad (7.48)$$

Plaçons-nous dans le cas où les fonctions h et $\tilde{h}$ sont séparables. Nous sommes ramenés à effectuer des filtrages inverses 1D non stationnaires sur les lignes et les colonnes de l'image à agrandir. Chaque filtrage revient à résoudre un système linéaire de la forme

$$\mathbf{A}\mathbf{c} = \mathbf{v}, \quad (7.49)$$

où

$$\mathbf{A}[k, l] = \sum_{j \in \mathbb{Z}} \tilde{h}(k - \tau - \frac{j}{\alpha}) h(\frac{j}{\alpha} + \tau - l) \quad \forall k, l \in \mathbb{Z}. \quad (7.50)$$

L'algorithme effectuant un tel filtrage est exactement celui que nous avons développé au chapitre 4. Il faut simplement que les fonctions h et $\tilde{h}$ soient à support compact.

Détaillons l'implémentation de l'induction, dans le cas d'un facteur α non entier. Nous considérons que $h = \beta^3$ est la B-spline cubique et que $\tilde{h}$ est, formellement, la spline cardinale cubique renormalisée, que nous avons proposée pour la réduction dans la section (5.3.2). L'induction (ici orthogonale puisque $\bigvee_{\alpha, \tau}(h) = \bigvee_{\alpha, \tau}(\tilde{h})$) consiste ainsi en les étapes suivantes (voir l'analogie avec le schéma (7.36)) :

1. On réduit l'image inductrice y^0 à l'aide de la formule (5.50). On obtient ainsi les coefficients $u[\mathbf{k}]$ définis par

$$u[\mathbf{k}] = \frac{\sum_{\mathbf{l} \in \mathbb{Z}^2} v[\mathbf{l}] \beta^3(\mathbf{k} - \boldsymbol{\tau} - \frac{\mathbf{l}}{\alpha})}{\sum_{\mathbf{l} \in \mathbb{Z}^2} \beta^3(\mathbf{k} - \boldsymbol{\tau} - \frac{\mathbf{l}}{\alpha})}. \quad (7.51)$$

2. On convolue v par b_1^3 .
3. On calcule l'image de résidu r par soustraction pixel à pixel des deux images précédentes :

$$r[\mathbf{k}] = v * b_1^3[\mathbf{k}] - u[\mathbf{k}] \quad \forall \mathbf{k} \in \mathbb{Z}^2. \quad (7.52)$$

4. On effectue l'agrandissement $\mathcal{A}_{\alpha, \tau}$ de r . La première étape consiste à effectuer le filtrage inverse de décorrélation, qui équivaut formellement à résoudre le système linéaire

$$\mathbf{A}\mathbf{c} = \mathbf{r}, \quad (7.53)$$

avec

$$\mathbf{A}[k, l] = \frac{\sum_{j \in \mathbb{Z}} \beta^3(k - \tau - \frac{j}{\alpha}) \beta^3(\frac{j}{\alpha} + \tau - l)}{\sum_{j \in \mathbb{Z}} \beta^3(k - \tau - \frac{j}{\alpha})} \quad \forall k, l \in \mathbb{Z}. \quad (7.54)$$

On utilise pour ce faire l'algorithme développé au chapitre 4, appliqué successivement sur les lignes et les colonnes de e .

5. On effectue la seconde partie de l'agrandissement de r , en appliquant la formule (7.42). On obtient ainsi l'image de résidu agrandie e définie par

$$e[\mathbf{k}] = \sum_{\mathbf{l} \in \mathbb{Z}^2} \beta^3(\frac{\mathbf{k}}{\alpha} + \boldsymbol{\tau} - \mathbf{l}) c[\mathbf{l}] \quad \forall \mathbf{k} \in \mathbb{Z}^2. \quad (7.55)$$

6. On ajoute ce résidu $e = y_L^\infty - y_L^0$ à l'image inductrice, pour finalement obtenir l'image induite :

$$y^\infty = y^0 + \lambda e. \quad (7.56)$$

où λ est le coefficient de relaxation proposé dans la section 7.5.1, que l'on peut mettre à 1 dans la plupart des situations.

Afin de pouvoir effectuer l'agrandissement par induction de facteur quelconque, il faut disposer d'une méthode d'agrandissement inductrice non linéaire, permettant l'agrandissement avec n'importe quel facteur. À cette fin, nous avons développé une extension de la méthode de Jensen et coll., initialement proposée pour le facteur 2 [123]. Cette méthode consiste à effectuer une interpolation spline cubique de l'image initiale, puis à modifier localement cette image agrandie en fonction des contours détectés dans l'image initiale. C'est la méthode inductrice que nous utilisons en pratique pour effectuer l'induction.

7.6 Validation expérimentale

7.6.1 Expériences de réduction/agrandissement combinées

L'agrandissement d'images est une opération délivrant une image agrandie dont la qualité ne peut pas être évaluée directement, puisque l'on ne dispose pas de référence (« vérité terrain ») permettant une comparaison directe. On dispose donc des solutions suivantes :

- évaluer la qualité de l'image agrandie directement, en lui mettant une « note ». Malheureusement, il n'existe pas de méthode systématique et fiable pour quantifier la qualité d'une image donnée. Le meilleur juge reste l'observateur humain, mais même en ce domaine, de fortes disparités se manifestent d'un individu à l'autre.
- effectuer des comparaisons visuelles entre des images agrandies par différentes méthodes. Cela permet d'estimer qualitativement la valeur relative d'une méthode par rapport à une autre. C'est la méthode de validation qui nous semble la plus appropriée.
- créer artificiellement une image de référence, en effectuant l'agrandissement non pas sur l'image initiale v , mais sur sa version précédemment réduite $r = \mathcal{R}_{\alpha, -\alpha\tau}v$. On peut alors quantifier, au moyen d'une distance PSNR entre l'image agrandie et la référence v , la qualité de l'agrandissement. Cette méthode est généralement adoptée dans la littérature. Elle constitue le seul moyen pratique de quantifier la qualité d'une méthode d'agrandissement.

Nous émettons des réserves quant à la validation de l'agrandissement suivant cette dernière méthode, qui consiste à effectuer une combinaison réduction/agrandissement. Le fait de réduire une image avant de l'agrandir ajoute, de fait, une connaissance *a priori* sur l'image v que l'on veut estimer, à savoir le fait qu'elle appartient à l'ensemble induit Υ_r , c'est-à-dire que l'on sait de v qu'elle engendre r par réduction avec une méthode connue $\mathcal{R}$. Ce type d'expérimentation introduit donc un biais en faveur des méthodes qui utilisent

	Nulle	Zeng		Ramponi		Jensen	
	$y^\infty = v_L$	y^0	y^∞	y^0	y^∞	y^0	y^∞
<i>Lena</i>	35.77	34.26	35.52	34.40	35.29	33.14	35.70
<i>Barbara</i>	25.71	25.24	25.67	25.40	25.80	24.92	25.72
<i>Baboon</i>	24.68	23.81	24.56	24.02	24.59	23.26	24.82
<i>Lighthouse</i>	27.07	26.29	27.17	25.91	26.89	25.65	27.27
<i>Goldhill</i>	32.28	31.37	32.16	31.52	32.12	30.49	32.28
<i>Boat</i>	31.07	30.05	30.97	30.17	30.89	29.13	31.21
<i>Camera</i>	27.27	26.48	27.32	26.62	27.30	25.95	27.68
<i>Peppers</i>	31.63	31.77	32.49	32.11	32.80	31.47	32.84

TAB. 7.1 – Résultats numériques après une réduction suivie d’un agrandissement par induction, de facteur $\alpha = 2$, d’offset $\tau = \mathbf{0}$. Cette expérience permet de comparer différentes méthodes inductrices. On voit que la méthode de Jensen et coll. donne les meilleurs résultats, ce qui indique que son utilisation est appropriée pour créer les coefficients d’ondelettes requis par le processus d’induction.

cette information disponible. En particulier, il n’est pas correct de comparer une méthode d’agrandissement vérifiant la contrainte de réduction par rapport à $\mathcal{R}$, à une autre ne la vérifiant pas. En d’autres termes, l’agrandissement de r sous contrainte de réduction fournit une estimée *oracle* de v . Les expériences de réduction/agrandissement n’apportent donc qu’une réponse relative au problème de l’évaluation de la qualité d’agrandissement. En particulier, ce procédé de validation dépend de la méthode de réduction employée.

En gardant à l’esprit ces remarques, nous proposons une expérience de réduction/agrandissement, dont le tableau 7.1 présente les résultats. Plusieurs images ont été réduites d’un facteur 2, par la méthode spline cubique décrite par (5.30). Nous avons implémenté et testé trois méthodes proposées dans la littérature pour l’agrandissement de facteur 2. Les méthodes de Zeng et coll. et Ramponi et coll. effectuent une interpolation non linéaire avec respectivement un filtre de type médian et un filtre rationnel. La troisième méthode, proposée par Jensen et coll., réalise quant à elle une interpolation spline cubique, à l’exception des zones où un contour est détecté, auquel cas les pixels agrandis sont calculés à partir du modèle idéal représentant au mieux le contour. Les PSNRs entre l’image initiale v et l’image réduite puis ré-agrandie avec ces méthodes sont donnés dans les colonnes « y^0 ». Nous avons appliqué le procédé d’induction, en utilisant comme inductrice chacune de ces images agrandies. Le modèle de réduction utilisé par l’induction est celui effectivement employé pour réduire les images avant agrandissement. Les PSNRs entre les images induites et l’image initiale sont reportés dans les colonnes « y^∞ ». Nous avons aussi indiqué les résultats de l’induction avec une inductrice nulle (tous ses pixels à zéro). Dans ce dernier cas, l’image induite est simplement l’image v réduite, puis ré-agrandie linéairement selon la méthode (6.15). Cette image représente la partie basses fréquences de v , commune à toutes les images induites : l’induction ajoute à cette image le résidu hautes fréquences y_H^0 .

extrait de l'image inductrice y^0 .

Cette expérience de réduction/agrandissement permet de valider le procédé d'induction : un PSNR associé à l'image induite supérieur à celui obtenu avec l'inductrice nulle indique que les coefficients d'ondelettes de l'image inductrice sont de bons prédictors des coefficients inconnus de l'image v , que l'on cherche à estimer. Seule la méthode de Jensen réussit dans ce but. C'est ainsi cette méthode que nous avons adoptée en pratique comme méthode inductrice à combiner avec l'induction. Nous voyons que même combinée à la méthode de Jensen, l'induction offre un gain numérique peu important par rapport à l'agrandissement linéaire sous contrainte de réduction (revenant formellement à effectuer l'induction avec une inductrice nulle). Cependant, les différences se manifestent au niveau des contours, dont nous sommes visuellement très sensibles à la qualité. Le gain qualitatif obtenu avec l'induction est ainsi important dans les faits, comme le montreront les images dans la suite de cette section. Cela illustre d'ailleurs une des limites du PSNR, qui est une mesure globale ne tenant pas compte de la prévalence de l'information géométrique.

Les PSNRs entre les images inductrices et l'image de référence v ne sont donnés qu'à titre indicatif. On voit la nette différence de résultats entre les images inductrices, qui ne vérifient pas la contrainte de réduction, et les images induites. Comme nous le disions, cela traduit le biais existant dans ce type d'évaluation en faveur des méthodes vérifiant la contrainte de réduction. Seule la comparaison des PSNRs associés aux images induites est donc pertinente. On remarque ainsi que l'agrandissement linéaire sous contrainte de réduction surpasse toutes les méthodes non linéaires ne vérifiant pas cette contrainte, malgré leur complexité beaucoup plus grande.

7.6.2 Résultats d'agrandissement au moyen de différentes méthodes

Nous illustrons maintenant, au moyen d'images agrandies, les performances de différentes méthodes d'agrandissement. Nous proposons une image test, représentée sur la figure 7.6, composée d'une scène complexe, et particulièrement difficile pour l'agrandissement. Cela permet d'exacerber les comportements des différentes méthodes testées, et de mettre en exergue les défauts de chacune. Les résultats sont représentés sur les figures 7.7 à 7.22. Les différentes méthodes testées sont l'interpolation au plus proche voisin, l'interpolation spline cubique, notre méthode d'interpolation WiskI, la méthode de Jensen, l'induction (utilisant la méthode de Jensen comme méthode inductrice). La méthode AQua2 de Muresan [166] est aussi testée³. La dernière méthode prise en compte⁴ est celle de Belahmidi et coll. [28], représentative des approches variationnelles fondées sur des équations aux dérivées partielles. Les méthodes de Muresan et Belahmidi ne fonctionnent que pour des

3. Le logiciel Pictura, qui implémente la méthode AQua2, est disponible librement à l'adresse http://www.dmmnd.net/products/pictura_downloads.htm.

4. Nous remercions A. Belahmidi de nous avoir gracieusement fourni un exécutable implémentant sa méthode.

facteurs entiers, et seul le facteur 2 peut être employé pour le moment avec notre méthode WiskI. Nous avons étendu la méthode de Jensen, initialement proposée par ses auteurs uniquement pour le facteur 2, afin de pouvoir l'utiliser avec un facteur quelconque. Afin de pouvoir comparer toutes les méthodes mentionnées, nous avons opté pour un facteur d'agrandissement égal à 4, le facteur 2 ne nous semblant pas suffisamment discriminant. Pour la méthode WiskI, deux agrandissements successifs de facteur 2 sont réalisés. Notons que toutes les images de cette section sont reproduites avec la même résolution, qui n'est pas nécessairement la même que celle des autres images incluses dans ce manuscrit.

Les figures qui suivent montrent notre image de test ainsi que des extraits d'agrandissements de facteur 4 de cette image.



FIG. 7.6 : Image test représentant une scène complexe. Les figures suivantes montrent des extraits des images agrandies de facteur 4 à partir de cette image.

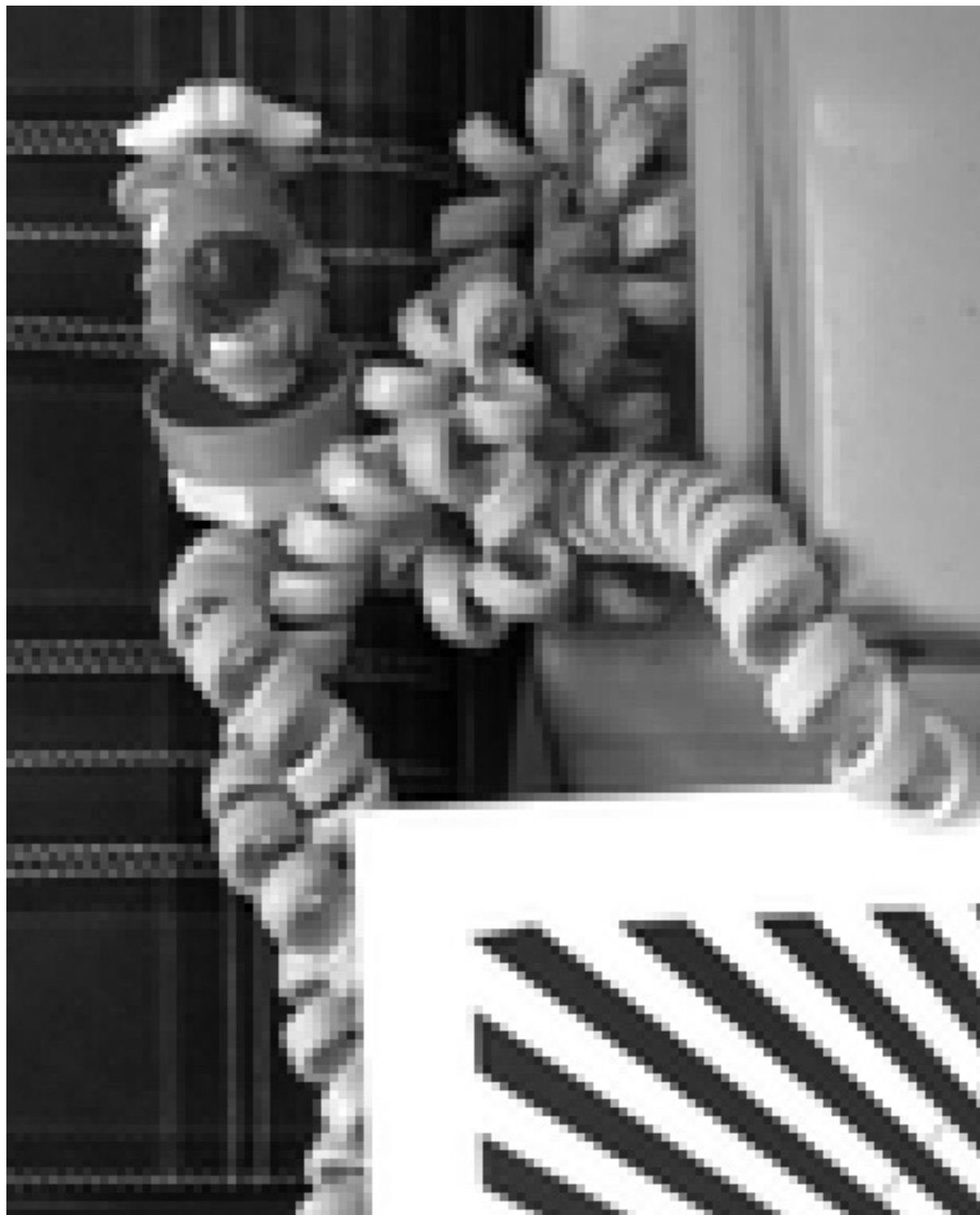


FIG. 7.7 : Agrandissement ($\alpha = 4$, $\tau = [-3/8, -3/8]$) par interpolation au plus proche.



FIG. 7.8 : Agrandissement ($\alpha = 4$, $\tau = [0, 0]$) par interpolation spline cubique.



FIG. 7.9 : Agrandissement ($\alpha = 4$, $\tau = [0, 0]$) par la méthode AQua2 de Muresan.



FIG. 7.10 : Agrandissement ($\alpha = 4$, $\tau = [0, 0]$) par interpolation WiskI (2 itérations).



FIG. 7.11 : Agrandissement ($\alpha = 4$, $\tau = [-3/8, -3/8]$) par la méthode de Belahmidi.



FIG. 7.12 : Agrandissement ($\alpha = 4$, $\tau = [0, 0]$) par la méthode de Jensen.



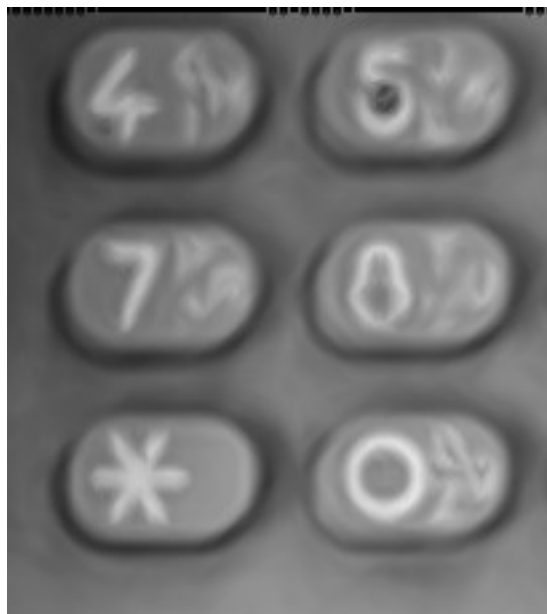
FIG. 7.13 : Agrandissement ($\alpha = 4$, $\tau = [0, 0]$) par induction.



FIG. 7.14 : Agrandissement ($\alpha = 4$, $\tau = [0, 0]$) par induction relaxée ($\lambda = 0.6$).



(a) Interpolation spline cubique



(b) Méthode Aqua2.



(c) Méthode de Jensen

(d) Induction relaxée, $\lambda = 0.6$.FIG. 7.15 : Agrandissement ($\alpha = 4$, $\tau = [0, 0]$) par différentes méthodes.

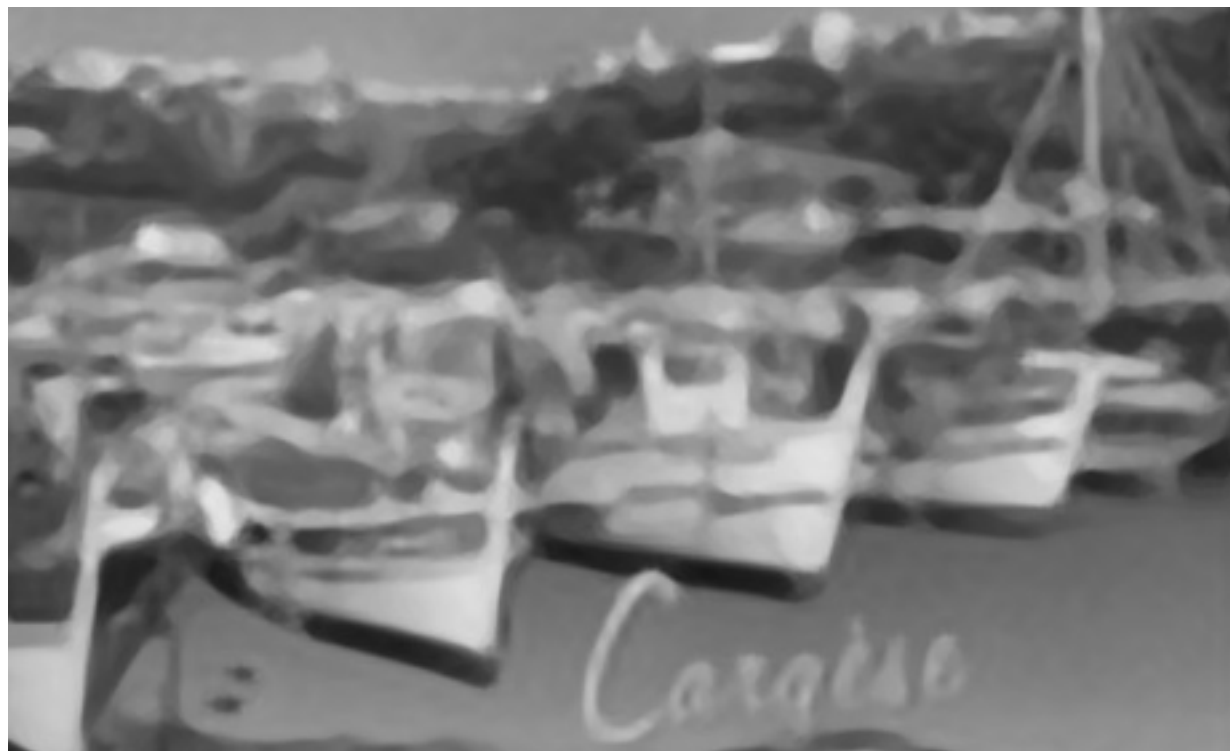


(a) Interpolation spline cubique.



(b) Induction.

FIG. 7.16 : Agrandissement ($\alpha = 4$, $\tau = [0, 0]$ sauf en (f)) par différentes méthodes.

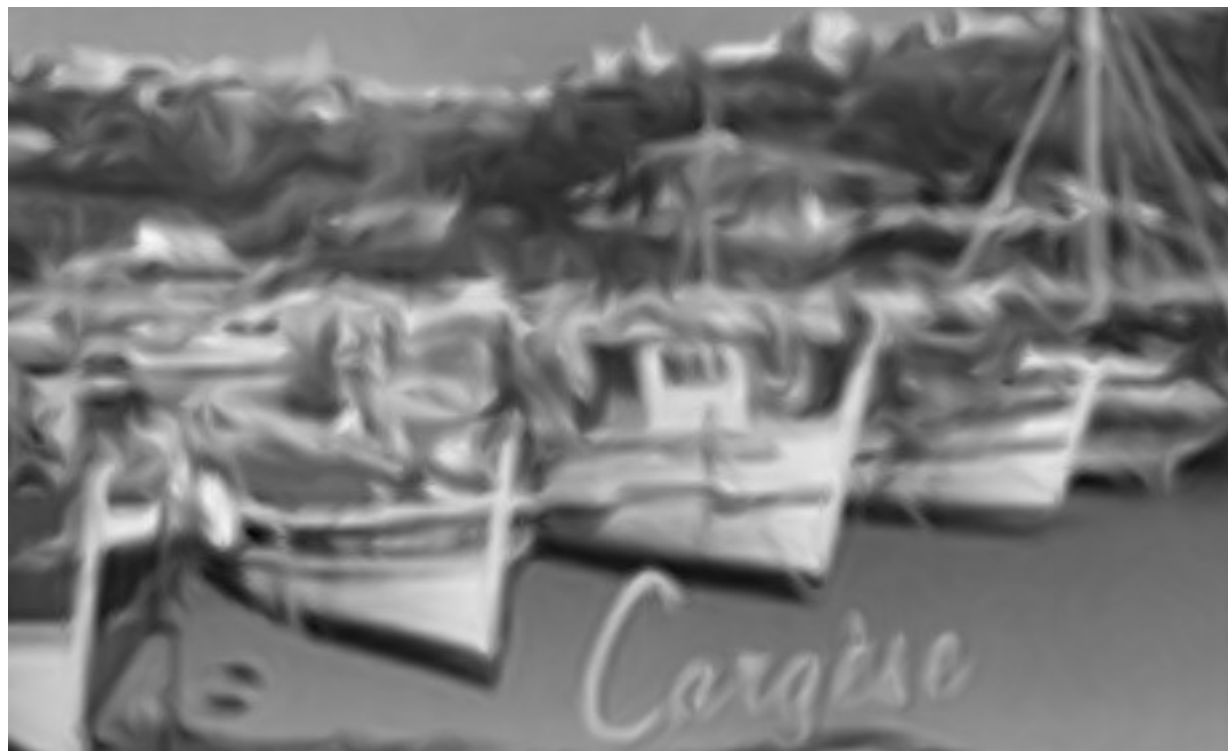


(c) Méthode de Jensen.



(d) Interpolation WiskI, 2 itérations.

FIG. 7.16 : Suite.



(e) Méthode AQua2.

(f) Méthode de Belahmidi ($\tau = [-3/8, -3/8]$).

FIG. 7.16 : Suite.



FIG. 7.17 : Agrandissement ($\alpha = 4$, $\tau = [0, 0]$) par interpolation spline cubique.



FIG. 7.18 : Agrandissement ($\alpha = 4$, $\tau = [0, 0]$) par la méthode de Jensen.



FIG. 7.19 : Agrandissement ($\alpha = 4$, $\tau = [0, 0]$) par la méthode AQUA2 de Muresan.



FIG. 7.20 : Agrandissement ($\alpha = 4$, $\tau = [0, 0]$) par la méthode de Belahmidi.



FIG. 7.21 : Agrandissement ($\alpha = 4$, $\tau = [0, 0]$) par interpolation WiskI, 10 itérations.



FIG. 7.22 : Agrandissement ($\alpha = 4$, $\tau = [0, 0]$) par induction.

À la vue des images précédentes, plusieurs constatations peuvent être faites.

L'image 7.7 agrandie au plus proche voisin montre le niveau de détails disponibles dans l'image initiale. L'interpolation spline cubique produit un effet de flou systématique, que l'on retrouve aussi avec la méthode d'interpolation WiskI, présentant sur ce point un rendu similaire à la méthode précédente. Cet effet de flou apparaît aussi avec la méthode AQua2. Il est moins présent avec la méthode de Belahmidi et l'induction. La méthode de Jensen présente des contours nets, mais souffre d'un effet de flou important ailleurs, dû à la disparition des petits détails (voir les yeux du personnage dans l'image 7.12). La méthode de Jensen produit ainsi des images à l'allure de peintures, avec un effet d'« à plat » (voir les images 7.15 (c) et 7.16 (c)) dû à la dichotomie entre des contours principaux nets d'une part, et le reste de l'image qui est atténué d'autre part.

L'effet de crénelage au niveau des contours obliques est bien visible sur les lignes de la mire, par exemple avec l'interpolation spline cubique en bas à droite de l'image 7.8, ou sur l'image 7.17. Cet effet est bien maîtrisé par toutes les autres méthodes ; il s'agit là du point fort des méthodes d'agrandissement non linéaires. Par contre, la méthode Belahmidi tend à « quantifier » les contours d'orientation proche de la verticale ou l'horizontale, par exemple le bord de la feuille blanche dans l'image 7.11. Cela donne un effet de « cubisme » désagréable et systématique, particulièrement visible au niveau des textures, comme le montre l'image 7.16 (f). La méthode AQua2 de Muresan présente aussi une importante contre-partie à l'absence de crénelage : des micro-contours sont créés dans toute l'image, comme si celle-ci était localement lissée dans la direction de plus grande régularité. Un effet « grains de riz » en résulte au niveau des textures, comme le montre l'image 7.16 (e).

La méthode de Jensen et l'induction sont les seules à fournir des contours francs, c'est-à-dire sans transition floue dans la direction transverse. En effet, l'agrandissement correct des contours nécessite d'une part une transition dépourvue d'oscillations dans la direction longitudinale, mais aussi une transition franche dans la direction transversale. Sur ce premier point, la méthode WiskI est de loin la plus performante, alors que la méthode de Jensen et par voie de conséquence l'induction surpassent les autres méthodes sur le second point. Le crénelage est une des deux formes que peut prendre l'*aliasing* dans les images agrandies, les effets de moiré étant la seconde manifestation de ce phénomène de repliement de spectre. La capacité remarquable de WiskI à contrecarrer l'*aliasing* présent dans l'image initiale est illustrée par l'image 7.21. Aucune autre méthode n'est capable de reconstituer à la résolution de l'image agrandie le centre de la mire, où les effets de moiré sont conséquents avec toutes les autres méthodes. L'*aliasing* est aussi visible en bas à droite du centre de la mire, sous forme de vagues concentriques perpendiculaires aux lignes de la mire. Cet effet, visible dans l'image interpolée spline cubique 7.17, l'est encore plus dans l'image induite 7.22. Par contre, les lignes de la mire sont dénuées d'effets de crénelage dans l'image induite, ce qui ne peut pas être le cas avec une méthode d'agrandissement linéaire (hormis

si l'on choisit l'espace d'agrandissement $V_\alpha(\text{sinc})$.

Les effets de l'*aliasing* se mêlent à ceux du *ringing* pour créer des oscillations au voisinage des contours, apparaissant sous forme de vagues ou d'échos. Cet effet est plus visible avec l'induction qu'avec les autres méthodes (AQua2 étant la meilleure sur ce point) : ces dernières souffrent d'un effet de flou qui atténue certes ces artéfacts, mais aussi toutes les autres structures de l'image. La plus grande visibilité des échos dans les images induites, (voir le bord de touches dans l'image 7.15 (d), le C de Cargèse dans l'image 7.16 (b), l'intérieur des lignes noires dans l'image 7.13) est la contre-partie à accepter pour la diminution du flou. Cela montre la limite de l'induction : si le contenu hautes fréquences extrait de l'image induite n'est pas localement suffisant, les effets d'*aliasing* et de *ringing* prennent le pas visuellement sur la netteté supplémentaire des contours. Ainsi, bien que les effets d'échos soient visuellement désagréables dans l'image induite 7.16 (b), la lisibilité de l'image par rapport aux autres méthodes est meilleure : le bord des bateaux, par exemple, est bien mieux rendu.

La seconde image que nous présentons afin d'illustrer les particularités des différentes méthodes est l'image *Camera*. La figure 7.23 montre les résultats obtenus sur la partie la plus discriminante de cette image, dont le fort contraste rend là aussi l'agrandissement particulièrement difficile.

Tout comme dans le précédent exemple, les caractéristiques des méthodes se manifestent clairement avec cette image. L'interpolation spline cubique souffre des trois défauts que sont le flou, l'*aliasing* et le *ringing*. La méthode de Jensen corrige ces trois défauts, au prix d'une atténuation voire disparition des textures et petites structures, comme on peut le voir dans la caméra. Avec la méthode AQua2, les contours sont moins francs par rapport à ceux produits avec la méthode de Jensen, et des effets de lissage directionnel désagréables apparaissent. L'induction bénéficie des bonnes propriétés de la méthode de Jensen, et l'intérieur de la caméra ainsi que les autres structures sont, cette fois, bien reproduites. Le *ringing* apparaît tout de même avec l'induction, dans une proportion similaire à ce qu'introduit l'interpolation spline cubique. Les effets de crénelage au niveau des contours sont toutefois bien mieux maîtrisés.



FIG. 7.23 : Agrandissement ($\alpha = 4$, $\tau = [0, 0]$). En haut : interpolation spline cubique ; en bas : méthode de Jensen.



FIG. 7.23 : Suite. En haut : méthode AQua2 ; en bas : induction relaxée, $\lambda = 0.6$.

Ces deux premières images ont permis de montrer les défauts inhérents à chaque méthode, et l'on constate qu'aucune n'est complètement satisfaisante, au sens où aucune méthode ne combine tous les avantages spécifiques à chacune dans certains cas. En pratique, sur la plupart des images, l'induction tire nettement son épingle du jeu, en fournissant des images dépourvues d'artéfacts systématiques. L'introduction d'une « signature » synthétique et caractéristique d'une méthode nous semble en effet un défaut rédhibitoire. Il en va ainsi de l'effet d'« à plat » introduit par la méthode de Jensen, de l'effet de « cubisme » introduit par la méthode de Belahmidi, ou de l'effet de lissage directionnel des textures introduit par la méthode AQua2 de Muresan. L'induction est exempte d'un tel biais systématique.

La figure 7.24 montre un exemple typique d'image fournie par l'induction, alors que les deux exemples précédents ne sont pas représentatifs des images couramment rencontrées en pratique. L'induction produit des images au contraste naturel et avec des bords francs, ce qui en fait une méthode d'agrandissement de choix, d'autant plus qu'elle s'accommode de n'importe quel facteur d'agrandissement, par exemple le nombre irrationnel π dans ce dernier exemple.



FIG. 7.24 : Agrandissement ($\alpha = \pi$, $\tau = [0, 0]$). (a) : interpolation spline cubique.



FIG. 7.24 : Suite. (b) : induction.

Image 1	0.5	Image 5	1.3	2	beaucoup mieux
Image 2	1.3	Image 6	0.2	1	un peu mieux
Image 3	0.8	Image 7	0.8	0	pareil
Image 4	1.0	Image 8	0.5	-1	un peu moins bien
				-2	beaucoup moins bien

TAB. 7.2 – Résultats de la campagne de tests subjectifs menée par Philips Applied Technologies auprès de 6 observateurs. L'induction est comparée à l'interpolation bicubique de facteur 2, sur 8 images. À gauche : les notes moyennées sur les observateurs. À droite : signification des notes.

7.6.3 Campagnes de tests subjectifs

Nous présentons maintenant deux campagnes de tests menées séparément auprès de plusieurs observateurs, afin de comparer différentes méthodes d'agrandissement.

La première série de tests a été menée au sein de la société Philips Applied Technologies, dans le cadre d'un contrat de collaboration. Le service recherche et développement de cette entreprise s'intéresse aux méthodes d'agrandissement d'images, pour afficher un même *medium* sur des périphériques de résolutions différentes (téléphone, PDA, téléviseur, PC...), et en particulier pour la conversion du format TV vers TVHD. Nous avons délivré à Philips un exécutable implémentant la méthode de Jensen, ainsi qu'une version modifiée de l'induction reposant entre autres sur une relaxation locale de l'induction. Cette option « défensive » permet d'augmenter la robustesse vis-à-vis des effets d'*aliasing* et de *ringing*. Cependant, les résultats de cette sophistication ne varient pas sensiblement par rapport à l'induction classique.

Cette méthode d'induction modifiée a été évaluée par Philips pour un agrandissement de facteur 2, en comparaison avec l'interpolation bicubique, sur une base de 8 images. Les tests ont été effectués par Cécile Dufour et Jean Gobert, auprès de 6 observateurs non experts. Les images étaient affichées l'une après l'autre, dans un ordre aléatoire, et séparée par un écran gris. L'observateur devait alors comparer les deux méthodes selon le barème donné dans le tableau 7.2. Les tests étaient réalisés sur un écran LCD de 19 pouces, avec une luminosité ambiante maintenue basse dans la pièce. Les résultats, reportés dans le tableau 7.2, montrent que notre méthode est majoritairement préférée à l'interpolation bicubique. Le gain qualitatif n'est pas très élevé, mais les différences sont effectivement peu visibles dans le cas d'un facteur 2, et ce quelles que soient les méthodes employées.

Philips n'a pas choisi d'opter pour une diffusion de l'induction dans ses périphériques d'affichage grand public, malgré l'intérêt manifesté pour notre méthode. En effet, les facteurs d'agrandissement requis dépassent rarement 2, et une interpolation suivie d'un réhaussement, pour obtenir des images dépourvues de flou (*sharp*) est jugé suffisant par Philips pour le type d'applications visées.

	Spline3	Jensen	Induction	AQua2
Spline3		15	98	25
Jensen	85		90	90
Induction	2	10		9
AQua2	75	10	91	

TAB. 7.3 – Résultats de la campagne de tests subjectifs menée par Thierry Dolmière auprès de 21 sujets, sur 9 images de scènes naturelles. Les valeurs indiquent les pourcentages de fois où une méthode (indiquée en ordonnée) est préférée à une autre (indiquée en abscisse). Lorsque l'induction est comparée à une des trois autres méthodes, elle est préférée dans 93% des cas (moyenne de la troisième colonne).

La seconde campagne de tests, de plus grande envergure, a été menée par Thierry Dolmière durant son projet de fin d'études⁵ portant sur l'évaluation subjective de la qualité d'images [89]. 9 images de scènes naturelles (paysages, scènes urbaines. . .) ont été comparées par 21 sujets. L'expérience s'est faite dans le respect des recommandations de l'Union Internationale des Télécommunications (UIT) pour ce type de tests psycho-cognitifs [119] : pièce entièrement tapissée en gris neutre, luminosité ambiante contrôlée, moniteur CRT étalonné. . . Quatre méthodes d'agrandissement étaient comparées deux à deux, pour un facteur d'agrandissement égal à 4. L'ordre des comparaisons ainsi que l'ordre des images dans chaque comparaison étaient choisis aléatoirement. Les méthodes comparées étaient : l'interpolation spline cubique, l'induction, la méthode AQua2, et la méthode de Jensen. Le tableau 7.3 synthétise une partie des résultats obtenus par T. Dolmière. L'induction est nettement préférée aux trois autres méthodes d'agrandissement : elle est plébiscitée dans 93% des cas.

Cette campagne de tests valide, du point de vue de la qualité visuelle ressentie, les principes théoriques sur lesquels reposent l'induction : à la fois un rendu des basses fréquences fidèle à celui de l'image initiale, par l'intermédiaire de la contrainte de réduction, et l'ajout de hautes fréquences appropriées au niveau des contours, pour un rendu cohérent de ces derniers et une impression de netteté globale grandement améliorée, du fait de la prédominance des contours et de l'information géométrique dans le processus de vision et de perception sémantique des images.

7.7 Conclusion

Dans ce chapitre, nous avons présenté l'induction, approche puissante, efficace et rapide pour l'agrandissement sous contrainte de réduction. Nous avons revisité en profondeur cette méthode, initialement formulée par Didier Calle dans son travail de thèse [50, 49] comme un procédé de régularisation, visant à restaurer dans l'image inductrice l'information contenue dans l'image initiale, mais perdue lors de l'agrandissement. Cette méthode se trouvait

5. Stage de Master 2 Signal Image Parole Télécoms, École Doctorale EEATS, Grenoble.

pénalisée, sous sa forme originelle, par une implémentation complexe due aux méthodes POCS sous-jacentes et par un temps de calcul élevé dû à sa forme itérative. En reformulant et ré-interprétant entièrement l'induction, nous avons pu non seulement l'étendre (induction oblique, relaxation, facteur d'agrandissement quelconque), mais aussi mettre en lumière ses propriétés (caractérisation précise de son action en termes de coefficients d'ondelettes). Nous avons fourni dans tous les cas des implémentations directes, non itératives. Cette refonte profonde de la méthode d'induction la rend plus que jamais attrayante pour l'agrandissement d'images.

L'étude menée dans ce chapitre montre qu'il est possible d'obtenir une image vérifiant les propriétés de cohérence souhaitées (contrainte de réduction), tout en ayant de bonnes qualités visuelles. Cela est rendu possible par l'extrapolation implicite de coefficients d'ondelettes que réalise l'induction, considérée comme un processus global. En effet, il nous semble réducteur de dissocier la méthode inductrice de l'induction, vue alors comme un post-traitement. La méthode inductrice et l'induction agissent de concert pour former l'image induite finale.

Ce travail sur l'induction a fait l'objet des publications [68, 66].

L'application de l'induction au problème de fusion d'images satellitaires constitue une perspective intéressante. En effet, le processus de fusion d'une image panchromatique haute résolution et d'une image multispectrale basse résolution consiste à insérer de manière pertinente les hautes fréquences de l'image panchromatique dans la version agrandie de l'image multispectrale. Le parallèle avec l'agrandissement apparaît alors, l'image panchromatique jouant le rôle de l'image inductrice.

Conclusion générale

Résumé

CE travail de recherche a été focalisé sur la reconstruction de signaux, ainsi que la réduction et l'agrandissement d'images fixes.

Dans une première partie, nous avons abordé le problème de reconstruction, qui vise à modéliser un signal (ou une image) discret par une fonction définie continûment. Nous avons formulé ce problème comme un problème d'estimation inverse mal posé, en interprétant les échantillons du signal comme des mesures linéaires, éventuellement bruitées, sur une fonction inconnue qu'il s'agit d'approcher. Dans les trois chapitres composant cette première partie, nous avons traité respectivement les cas 1D uniforme, 2D uniforme, et 1D non uniforme. Dans tous les cas, nous avons choisi d'opérer la reconstruction dans un espace fonctionnel linéaire et invariant par translation, typiquement un espace spline, paramétré par un ensemble de coefficients formant un signal discret uniforme. Cette approche a deux avantages : d'une part la fonction reconstruite uniforme est facilement manipulable, au moyen de traitements opérant directement sur les coefficients qui la caractérisent, et d'autre part, la fonction reconstruite a une résolution intrinsèque qui peut être choisie selon les applications envisagées, et ce indépendamment de la densité de mesures disponibles.

En maximisant la décroissance asymptotique de l'erreur de reconstruction lorsque le pas d'échantillonnage tend vers 0, nous avons proposé des méthodes de reconstruction originales, basées sur la notion de quasi-projection. Bien que connu et utilisé dans la communauté mathématique, ce concept n'a eu que peu d'écho en traitement du signal. La quasi-interpolation est généralement utilisée lorsque l'interpolation exacte pose des difficultés de mise en œuvre, ou présente un temps de calcul trop important. L'utilisation que nous avons faite de la quasi-interpolation, et des quasi-projections en général, est nouvelle, puisqu'en mettant à profit les degrés de liberté supplémentaires offerts par rapport à la reconstruction consistante, nous avons été en mesure de proposer des solutions meilleures du point de vue des propriétés asymptotiques d'approximation. Dans le cas bi-dimensionnel, nous nous sommes concentrés sur le treillis hexagonal, afin de montrer que nos solutions permettent de tirer parti des avantages théoriques de ce treillis, et ce sans pénalité de temps de calcul, par rapport aux solutions cartésiennes classiques. Pour toutes les méthodes proposées, nous nous sommes attachés à la réalisabilité de l'implémentation, et à la rapidité

de l'exécution. Les différentes facettes de ce travail sur la reconstruction ont fait l'objet des publications [65, 73, 72, 74, 69, 67].

Une idée maîtresse nous semble émerger de ce travail sur la reconstruction : lorsque l'on est confronté à un problème inverse d'un processus connu, il ne faut pas nécessairement chercher une méthode pseudo-inverse dudit processus. L'information disponible sur l'objet à construire, au travers des données disponibles, forment généralement un ensemble de contraintes *obliques* sur cet objet. L'idée générale et pourtant intuitive de *consistance* doit donc être dépassée dans certaines applications.

Dans la seconde partie de cette thèse, nous avons appliqué la même démarche méthodologique, pour traiter les problèmes de réduction et agrandissement d'images. Après avoir proposé un nouveau formalisme pour le redimensionnement et des noyaux d'erreurs permettant de quantifier en domaine fréquentiel les distorsions introduites par rapport au résultat idéal escompté, nous avons proposé des méthodes linéaires efficaces pour la réduction d'images. Le problème de l'agrandissement est quant à lui plus complexe. Nous avons montré les limites des méthodes linéaires pour l'agrandissement, et dressé un panorama des méthodes non linéaires proposées dans la littérature pour dépasser ces limitations. En effet, il nous faut opter pour des méthodes d'agrandissement non linéaires, seules à même de synthétiser sans artéfacts les contours des objets de l'image, qui représentent l'information géométrique à laquelle le système visuel est le plus sensible. Nous avons présenté l'agrandissement par induction, une approche originale initialement proposée par Didier Calle en 1999. En revisitant en profondeur cette méthode, et en lui apportant des améliorations notables, nous avons montré que l'induction est une réponse pertinente au problème de l'agrandissement. Les résultats obtenus, avec une mise en œuvre pourtant relativement simple et rapide, sont meilleurs que ceux des méthodes existantes. Les différents volets de l'induction ont été publiés dans [66, 68]. Des collaborations ont été entreprises, durant lesquelles des campagnes de tests subjectifs ont été menées, qui ont confirmé de manière indépendante les qualités de l'induction.

Mise en perspective

Ce travail a été réalisé avec en toile de fond la volonté d'investiguer les lois mathématiques sous-jacentes aux signaux et images discrets. Beaucoup d'algorithmes existent dans la littérature, visant à exécuter des tâches diverses, et basés en général sur un mélange d'intuitions se manifestant sous forme d'heuristiques, et d'exigences informatiques. Les principes formels, théories ou axiomatiques recouvrant ces méthodes sont peu nombreux, voire absents, en particulier dans le domaine de l'agrandissement d'images. Sans prétendre proposer des principes théoriques généraux gouvernant le monde des signaux et des images, nous avons tenu à adopter une démarche méthodologique scrupuleuse.

Tout d'abord, nous nous sommes dotés d'un modèle décrivant la nature des objets

que nous entendions manipuler. En ce sens, nous nous sommes appuyés sur un modèle de formation des signaux linéaires, prenant en compte la non-idéalité des dispositifs d'acquisition (au moyen de la fonction de base $\tilde{\varphi}$), ainsi que la présence éventuelle de bruit sur les mesures. Tout au long de cette étude, nous avons interprété un signal discret comme un ensemble de mesures *localisées*, de manière uniforme ou non uniforme, et formant une représentation discrète d'une fonction inconnue $s(\mathbf{x})$. Cette formalisation du processus de discrétisation, donnant naissance à un signal discret à partir d'une fonction définie continûment, nous semble essentielle. Le signal discret hérite en effet, au travers de cette étape d'acquisition, de propriétés sur s qui peuvent fournir des connaissances *a priori* importantes, selon la nature du processus physique dont s est une représentation. Nous avons en particulier exploité la concentration de l'énergie, pour beaucoup de signaux d'intérêt comme les images naturelles, dans les basses fréquences.

Dans un second temps, nous avons formalisé les problèmes auxquels nous nous sommes ensuite confrontés, à savoir la reconstruction de signaux et le redimensionnement d'images. Le premier problème a été défini comme un problème inverse mal posé d'estimation de la fonction inconnue $s(\mathbf{x})$. Le redimensionnement, quant à lui, a été modélisé comme un problème d'estimation de mesures discrètes sur s , à partir de mesures discrètes différentes, et formant une image de résolution différente de celle désirée. Ce travail se situe ainsi à la confluence des mondes des signaux discrets et continus. Les passerelles entre ces deux univers sont trop souvent occultées, et résumées par un raccourci inadéquat au simple théorème de Shannon, qui ne décrit que le monde restreint des signaux à bande limitée.

Ce n'est qu'après avoir défini précisément à la fois les objets manipulés, les problèmes posés, et les résultats escomptés dans l'idéal, que nous avons développé des outils mathématiques permettant de proposer des solutions pratiques et performantes. Pour ainsi dire, l'intuition doit être formalisée par un modèle dont on dérive le bon outil. Différents noyaux d'erreur ont été proposés, permettant de quantifier la distorsion entre la solution atteinte en pratique, et la solution idéale inaccessible sans la présence d'un « *oracle* ». C'est en minimisant cette distorsion de manière asymptotique, dans les basses fréquences, que de nouvelles méthodes linéaires de reconstruction et redimensionnement ont été proposées, reposant sur la notion de quasi-projection. La philosophie qui sous-tend cette approche peut se résumer comme suit : puisque l'on compte opérer un traitement sur un objet inconnu à partir d'informations partielles sur celui-ci, on conçoit le traitement pour fournir le résultat idéal sur une classe restreinte de signaux représentatifs des signaux d'intérêt. Intuitivement, les polynômes sont des fonctions dont toute l'énergie est localisée à la fréquence nulle. Les méthodes de quasi-projection asymptotiquement optimales, qui fournissent le résultat idéal lorsque s est un polynôme, sont donc tout indiquées pour le traitement des signaux basses fréquences comme les images naturelles. Le critère d'optimalité choisi est donc fondamentalement pragmatique, et notre approche est versatile et extensible à d'autres classes de signaux : en remplaçant par exemple les polynômes par des sinusoides à différents harmoniques, on pourrait se focaliser sur d'autres types de signaux oscillants, comme des signaux

sonores ou des bruits mécaniques.

Ainsi, à l'aide d'outils comme les noyaux d'erreurs ou des méthodes de projections formelles dans des espaces de signaux, nous avons développé des solutions pratiques, efficaces, cohérentes avec nos modèles et optimales du point de vue de critères pragmatiques. De plus, les principes mis en avant sont génériques, à l'inverse des traitements *ad hoc* généralement proposés dans la littérature pour un problème spécifique. Nous avons mis l'accent sur les méthodes linéaires, car il nous semble nécessaire de pousser ces méthodes dans leurs derniers retranchements avant de passer à des approches non linéaires. L'expérimentation a été menée sans biais en faveur de nos méthodes ; en particulier, les tests d'agrandissement ont été effectués en double aveugle. Notons aussi que toutes nos solutions sont implémentables relativement aisément, sans artifices occultés. Pour notre méthode de reconstruction à partir d'échantillons non uniformes, nous avons mené l'étude algorithmique de bout en bout. Il nous semble en effet toujours opportun d'exploiter au maximum les spécificités du problème dans la mise en œuvre de la solution, plutôt que d'utiliser des outils « boîtes noires », ne serait-ce que pour mettre en lumière les mécanismes sous-tendant l'approche.

Horizons

Concernant tout d'abord le problème de l'agrandissement, notre étude mérite d'être poursuivie. Il faut garder à l'esprit la difficulté de ce problème : par exemple, dans le cas du facteur 2, les trois quarts de l'information brute sont perdus dans l'image agrandie que l'on cherche à estimer. On ne peut que chercher à synthétiser au mieux l'information géométrique, mais en aucun cas un algorithme ne peut retrouver des textures et microstructures qui ne sont pas présentes dans l'image initiale. Cependant, nous avons vu qu'il se trouve généralement, parmi les différentes méthodes exposées, une méthode fournissant un résultat très satisfaisant. Le *challenge* est de combiner les avantages de ces méthodes en une seule, afin de disposer d'une méthode unique d'agrandissement fournissant systématiquement le résultat atteignable le meilleur. En effet, il n'y a création d'artéfacts que lorsque la méthode ne tire pas parti correctement de l'information disponible dans l'image initiale, ou prend une mauvaise décision. Ainsi, nous aimerions combiner les avantages des méthodes de Jensen et WiskI, afin d'obtenir en toute situation des contours à la fois dépourvu d'*aliasing*, et présentant une transition franche.

Les principes et méthodes proposés dans cette thèse peuvent trouver un terrain d'application propice pour de nombreux problèmes. Nous comptons appliquer notre méthodologie au démosaïquage d'images couleurs, ainsi qu'à la fusion d'images multi-spectrales. Du point de vue théorique, ces deux problèmes consistent à estimer des mesures discrètes étant données d'autres mesures discrètes, et en ce sens peuvent bénéficier de l'approche asymptotique que nous avons proposée. Concernant l'aspect pratique, l'application de l'induction

au traitement de ces problèmes est une piste à poursuivre.

La compression d'images ou de sons fournissent aussi des horizons d'élargissement, où la recherche de représentations creuses nous semble passer par une bonne compréhension préalable de la notion de résolution, qui est au cœur de ce travail sur la reconstruction et le redimensionnement.

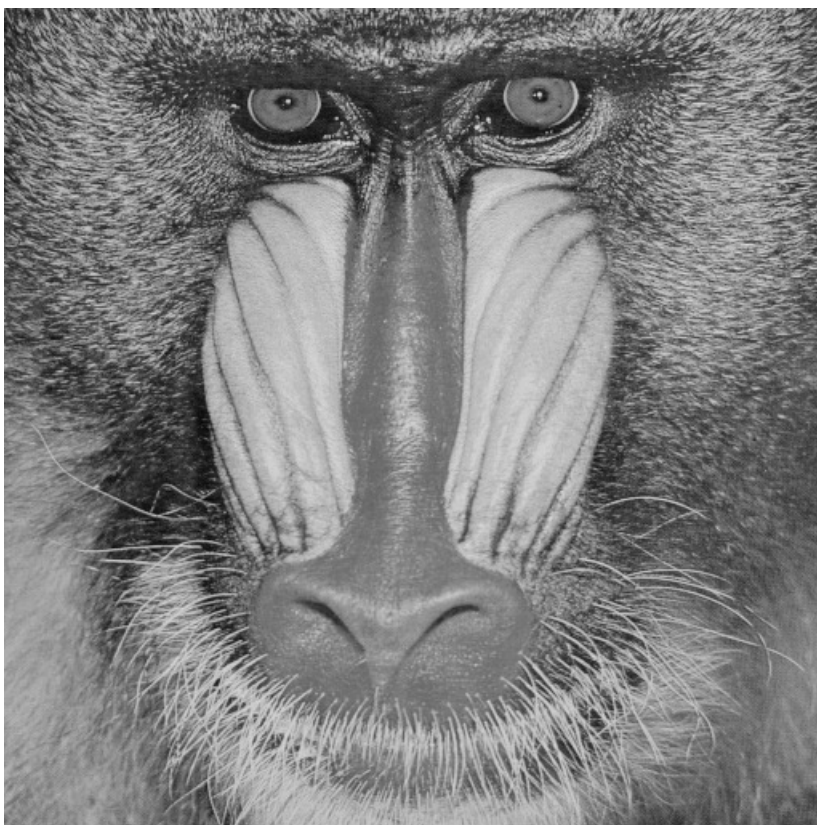
Annexe 1 : Images classiques de la littérature

Nous représentons dans cette annexe quelques images classiques de la littérature du traitement d'images, utilisées pour les expériences dans cette thèse. Elles sont toutes de taille 512×512 , sauf *Camera* de taille 256×256 .



Barbara

Précisons que l'image *Barbara* existe en plusieurs versions dans la littérature. Elles ont la même résolution et la même taille, mais n'ont pas été scannées de manière identique. Il faudra prendre garde à ne pas les confondre, et à ne pas tirer des conclusions hâtives quant aux mesures PSNR annoncées sur l'image *Barbara* dans tel ou tel article.

*Lena**Baboon*



Goldhill



Lighthouse

*Boat**Peppers*



Camera

Bibliographie

- [1] J. H. AHLBERG, E. N. WILSON et J. L. WALSH, *The Theory of Splines and Their Applications*, Academic Press, New York, 1967. 19
- [2] A. ALDROUBI, P. ABRY et M. UNSER, « Construction of biorthogonal wavelets starting from any two multiresolutions », *IEEE Trans. Signal Processing*, vol. 46, n° 4, p. 1130-1133, avr. 1998. 31
- [3] A. ALDROUBI, M. EDEN et M. UNSER, « Discrete spline filters for multiresolution and wavelets of ℓ_2 », *SIAM J. Math. Anal.*, vol. 25, n° 5, p. 1412-1432, sept. 1994. 116, 133
- [4] A. ALDROUBI et H. FEICHTINGER, « Exact iterative reconstruction algorithm for multivariate irregularly sampled functions in spline-like spaces: The L_p -theory », *Proc. Am. Math. Soc.*, vol. 126, n° 9, p. 2677-2686, sept. 1998. 28
- [5] A. ALDROUBI et K. GRÖCHENIG, « Beurling-Landau-type theorems for non-uniform sampling in shift invariant spline spaces », *J. Fourier Anal. Appl.*, vol. 6, n° 1, p. 93-103, 2000. 95
- [6] A. ALDROUBI et K. GRÖCHENIG, « Nonuniform sampling and reconstruction in shift-invariant spaces », *SIAM Rev.*, vol. 43, n° 4, p. 585-620, 2001. 28, 95
- [7] A. ALDROUBI, M. UNSER et M. EDEN, « Cardinal spline filters: Stability and convergence to the ideal sinc interpolator », *Signal Processing*, vol. 28, n° 2, p. 127-138, août 1992. 27
- [8] J. ALLEBACH et P. W. WONG, « Edge-directed interpolation », *Proc. of IEEE ICIP*, vol. 3, p. 707-710, 1996. 154, 164
- [9] J. P. ALLEBACH, « Image scanning, sampling, and interpolation », A. BOVIK, éd., *Handbook of Image and Video processing*, p. 629-643, Academic Press, San Diego, CA, 2000. 150
- [10] A. ALMANSA, *Échantillonnage, Interpolation et Détection. Applications en Imagerie Satellitaire*, Thèse de doctorat, École Normale Supérieure de Cachan, 2002. 18, 155
- [11] A. ALMANSA, S. DURAND et B. ROUGÉ, « Measuring and improving image resolution by adaptation of the reciprocal cell », *Journal of Mathematical Imaging and Vision*, vol. 21, n° 3, p. 235-279, nov. 2004. 155

- [12] L. ALVAREZ, Y. GOUSSEAU et J.-M. MOREL, « The size of objects in natural and artificial images », *Advances in Imaging and Electron Physics*, vol. 111, p. 167-242, Academic Press, 1999. 10
- [13] H. ALY et E. DUBOIS, « Crafting the observation model for regularized image up-sampling », *Proc. of IEEE ICASSP*, p. 101-104, avr. 2003. 145
- [14] H. ALY et E. DUBOIS, « Image up-sampling using total-variation regularization with a new observation model », *IEEE Trans. Image Processing*, vol. 14, n° 10, p. 1647-1659, oct. 2005. 154, 157, 158
- [15] H. ALY et E. DUBOIS, « Specification of the observation model for regularized image up-sampling », *IEEE Trans. Image Processing*, vol. 14, n° 5, p. 567-576, mai 2005. 8, 130, 154
- [16] H. A. ALY, *Regularized Image Up-sampling*, PhD thesis, University of Ottawa, Ottawa, Ontario, Canada, nov. 2004. 130, 145, 154, 157
- [17] H. A. ALY et E. DUBOIS, « Design of optimal camera apertures adapted to display devices over arbitrary sampling lattices », *IEEE Signal Processing Lett.*, vol. 11, n° 4, p. 443-445, avr. 2004. 59
- [18] H. C. ANDREWS et B. R. HUNT, *Digital Image Restoration*, Prentice Hall, Englewood Cliffs, NJ, 1977. 8
- [19] M. ANTONINI, M. BARLAUD, P. MATHIEU et I. DAUBECHIES, « Image coding using wavelet transformation », *IEEE Trans. Image Processing*, vol. 1, n° 2, p. 205-220, avr. 1992. 182
- [20] N. ARAD et C. GOTSMAN, « Enhancement by image dependent warping », *IEEE Trans. Image Processing*, vol. 8, n° 8, p. 1063-1074, août 1999. 162, 163
- [21] M. ARIGOVINDAN, *Variational Reconstruction of Vector and Scalar Images from Non-Uniform Samples*, Thèse de doctorat, École Polytechnique Fédérale de Lausanne (EPFL), 2005. 94, 96
- [22] M. ARIGOVINDAN, M. SÜHLING, P. HUNZIKER et M. UNSER, « Variational image reconstruction from arbitrarily spaced samples: A fast multiresolution spline solution », *IEEE Trans. Image Processing*, vol. 14, n° 4, p. 450-460, avr. 2005. 108, 115
- [23] T. ASAH, K. ICHIGE et R. ISHII, « An efficient algorithm for decomposition and reconstruction of images by box splines », *IEICE Trans. Fundamentals*, vol. E-84A, n° 8, p. 1883-1891, août 2001. 82
- [24] C. B. ATKINS, C. A. BOUMAN et J. P. ALLEBACH, « Optimal image scaling using pixel classification », *Proc. of IEEE ICIP*, vol. 3, p. 864-867, 2001. 164, 166
- [25] J.-F. AUJOL, *Contribution à l'analyse de textures en traitement d'images par méthodes variationnelles et équations aux dérivées partielles*, Thèse de doctorat, Université de Nice-Sophia Antipolis, 2004. 10

- [26] S. BAKER et T. KANADE, « Limits on super-resolution and how to break them », *IEEE Trans. Pattern Anal. Machine Intell.*, vol. 24, p. 1167-1183, sept. 2002. 150
- [27] R. BARTELS, J. BEATTY et B. BARSKY, *An Introduction to Splines for Use in Computer Graphics and Geometric Modeling*, Morgan Kaufmann Publishers, Los Altos, CA, 1986. 2
- [28] A. BELAHMIDI et F. GUICHARD, « A partial differential equation approach to image zoom », *Proc. of IEEE ICIP*, p. 649-652, oct. 2004. 150, 156, 157, 190
- [29] D. P. BERTSEKAS, *Nonlinear Programming*, Athena Scientific, Belmont, MA, 2^{de} édition, 1999. 16
- [30] G. BLANCHET, L. MOISAN et B. ROUGÉ, « A linear prefilter for image sampling with ringing artifact control », *Proc. of IEEE ICIP*, vol. 3, p. 577-580, 2005. 132
- [31] T. BLU, « Iterated filter banks with rational rate changes. Connection with discrete wavelet transforms », *IEEE Trans. Signal Processing*, vol. 41, n° 12, p. 3232-3244, déc. 1993. 127
- [32] T. BLU, P. THÉVENAZ et M. UNSER, « Generalized interpolation: Higher quality at no additional cost », *Proc. of IEEE ICIP*, vol. 3, p. 667-671, oct. 1999. 26
- [33] T. BLU, P. THÉVENAZ et M. UNSER, « MOMS: Maximal-order interpolation of minimal support », *IEEE Trans. Image Processing*, vol. 10, n° 7, p. 1069-1080, juil. 2001. 27, 85
- [34] T. BLU, P. THÉVENAZ et M. UNSER, « Linear interpolation revitalized », *IEEE Trans. Image Processing*, vol. 13, n° 5, p. 710-719, mai 2004. 46
- [35] T. BLU, P. THÉVENAZ et M. UNSER, « Complete parametrization of piecewise-polynomial interpolation kernels », *IEEE Trans. Image Processing*, vol. 12, n° 11, p. 1297-1309, sept. 2003. 23
- [36] T. BLU et M. UNSER, « Approximation error for quasi-interpolators and (multi-) wavelet expansions », *Applied and Computational Harmonic Analysis*, vol. 6, n° 2, p. 219-251, mars 1999. 11, 37, 41, 116
- [37] T. BLU et M. UNSER, « Quantitative Fourier analysis of approximation techniques: Part I—Interpolators and projectors — Part II—Wavelets », *IEEE Trans. Signal Processing*, vol. 47, n° 10, p. 2783-2806, oct. 1999. 24, 37, 38, 41, 46, 47, 55
- [38] T. BLU et M. UNSER, « A complete family of scaling functions: The (α, τ) -fractional splines », *Proc. of IEEE ICASSP*, vol. VI, p. 421-424, 2003. 134
- [39] C. De BOOR, « On the evaluation of box splines », *Numerical Algorithms*, vol. 5, p. 5-23, mars 1993. 73, 79
- [40] C. De BOOR, R. DEVORE et A. RON, « The structure of finitely generated shift-invariant spaces in $L_2(\mathbb{R}^d)$ », *J. Funct. Anal.* 119, p. 37-78, 1994. 28

- [41] S. BORMAN et R. L. STEVENSON, « Spatial resolution enhancement of low-resolution image sequences: A comprehensive review with directions for future research », Technical report, Department of Electrical Engineering, University of Notre Dame, Notre Dame, Indiana, juil. 1998. 144
- [42] A. BOUHAMIDI et A. MÉHAUTÉ, « Radial basis functions under tension », *Journal of Approximation Theory*, vol. 127, p. 135-154, 2004. 96
- [43] T. E. BOULT et G. WOLBERG, « Local image reconstruction and sub-pixel restoration algorithms », *Computer Graphics and Image Processing: Graphical Models and Image Processing (CVGIP:GMIP)*, vol. 55, n° 5, p. 63-77, jan. 1993. 8, 15
- [44] C. BOUMAN et M. SHAPIRO, « A multiscale random field model for Bayesian image segmentation », *IEEE Trans. Image Processing*, vol. 3, n° 2, p. 162-177, 1994. 10
- [45] C. M. BRISLAWN, « Classification of nonexpansive symmetric extension transforms for multirate filter banks », *Applied and Comp. Harmonic Anal.*, vol. 3, p. 337-357, 1996. 9, 102
- [46] M. BUHMANN, *Radial Basis Functions: Theory and Implementations*, Cambridge University Press, Cambridge, 2003. 96
- [47] P. D. BURNS, « Signal-to-noise ratio analysis of charge-coupled device imagers », *Proc. of SPIE*, vol. 1242, 1990. 8, 128
- [48] P. J. BURT et E. H. ADELSON, « The Laplacian pyramid as a compact image code », *IEEE Trans. Commun.*, vol. 31, n° 4, p. 532-540, avr. 1983. 132, 158
- [49] D. CALLE, *Agrandissement d'images par synthèse de similarités et par induction sur un ensemble*, Thèse de doctorat, Université J. Fourier, Grenoble, France, nov. 1999. 132, 163, 176, 177, 179, 218
- [50] D. CALLE et A. MONTANVERT, « Super-resolution inducing of an image », *Proc. of IEEE ICIP*, p. 232-236, 1998. 176, 179, 218
- [51] S. CARRATO, G. RAMPONI et S. MARSI, « A simple edge-sensitive image interpolation filter », *Proc. of IEEE ICIP*, sept. 1996. 162
- [52] S. CARRATO et L. TENZE, « A high quality 2x image interpolator », *IEEE Signal Processing Lett.*, vol. 7, p. 132-134, juin 2000. 162
- [53] J. R. CASAS, *Image Compression based on Perceptual Coding Techniques*, PhD thesis, Dept. Signal Theory Commun., UPC, Barcelona, Spain, mars 1996. 10
- [54] V. CASELLES, J.-M. MOREL et C. SBERT, « An axiomatic approach to image interpolation », *IEEE Trans. Image Processing*, vol. 7, n° 3, p. 376-386, mars 1998. 10, 156
- [55] K. CASTLEMAN, *Digital Image Processing*, Prentice-Hall, Englewood Cliffs, New Jersey, 1996. 8
- [56] A. CHAMBOLLE et P. L. LIONS, « Image recovery via total variation minimisation and related problems », *Numer. Math.*, vol. 76, p. 167-188, 1997. 155

- [57] S. G. CHANG, Z. CVETKOVIC et M. VETTERLI, « Locally adaptive wavelet-based image interpolation », *IEEE Trans. Image Processing*, vol. 15, n° 6, p. 1471-1485, juin 2006. 160
- [58] S. Grace CHANG, Z. CVETKOVIC et M. VETTERLI, « Resolution enhancement of images using wavelet transform extrema interpolation », *Proc. of IEEE ICASSP*, p. 2379-2382, mai 1995. 160
- [59] P. CHARBONNIER, L. BLANC-FÉRAUD et M. BARLAUD, « Deterministic edge-preserving regularization in computed imaging », *IEEE Trans. Image Processing*, vol. 6, n° 2, p. 298-311, fév. 1997. 18
- [60] R. CHELLAPPA et A. JAIN, *Markov Random Fields: Theory and Applications*, Academic Press, San Diego, CA, 1993. 10
- [61] W. CHEN, B. HAN et R.-Q. JIA, « Maximal gap of a sampling set for the exact iterative reconstruction algorithm in shift-invariant spaces », *IEEE Signal Processing Lett.*, vol. 11, n° 8, p. 655-658, août 2004. 95
- [62] W. CHEN et S. ITOH, « A sampling theorem for shift-invariant subspaces », *IEEE Trans. Signal Processing*, vol. 46, n° 10, p. 2822-2824, oct. 1998. 28, 95
- [63] C. K. CHUI et M.-J. LAI, « Algorithms for generating B-nets and graphically displaying spline surfaces on three- and four-directional meshes », *Comput. Aided Geom. Design*, vol. 8, n° 6, p. 479-493, 1991. 73, 79
- [64] E. COHEN, T. LYCHE et R. RIESENFELD, « Discrete box splines and refinement algorithms », *Comput. Aided Geom. Design*, vol. 1, p. 131-141, 1984. 73
- [65] L. CONDAT, T. BLU et M. UNSER, « Beyond interpolation: Optimal reconstruction by quasi-interpolation », *Proc. of IEEE ICIP*, vol. 1, p. 33-36, sept. 2005. 56, 222
- [66] L. CONDAT et A. MONTANVERT, « Agrandissement et compression d'images par induction », *Actes de CORESA*, p. 121-124, mai 2004. 219, 222
- [67] L. CONDAT et A. MONTANVERT, « Analyse multirésolution L_2 -optimale : Estimation par quasi-projections », *Actes du GRETSI*, 2005. 119, 222
- [68] L. CONDAT et A. MONTANVERT, « A framework for image magnification: Induction revisited », *Proc. of IEEE ICASSP*, vol. 2, p. 845-848, mars 2005. 219, 222
- [69] L. CONDAT et A. MONTANVERT, « Fast reconstruction from non-uniform samples in shift-invariant spaces », *Proc. of EUSIPCO*, 2006. 119, 222
- [70] L. CONDAT et A. MONTANVERT, « Reconstruction from non-uniform samples: A direct, variational approach in shift-invariant spaces », *IEEE Trans. Signal Processing*, 2006, en préparation. 119
- [71] L. CONDAT et D. VAN DE VILLE, « Hexagonal to cartesian grid conversion using multi-dimensional splines », *IEEE Trans. Image Processing*, 2006, soumis, en révision. 92

- [72] L. CONDAT et D. VAN DE VILLE, « Three-directional box-splines: Characterization and efficient evaluation », *IEEE Signal Processing Lett.*, vol. 13, n° 7, p. 417-420, juil. 2006. 92, 222
- [73] L. CONDAT, D. VAN DE VILLE et T. BLU, « Hexagonal versus orthogonal lattices: A new comparison using approximation theory », *Proc. of IEEE ICIP*, vol. 3, p. 1116-1119, sept. 2005. 91, 222
- [74] L. CONDAT, D. VAN DE VILLE et M. UNSER, « Efficient reconstruction of hexagonally sampled data using three-directional box-splines », *Proc. of IEEE ICIP*, 2006. 92, 222
- [75] R. COURANT, « Variational methods for the solution of problems in equilibrium and vibrations », *Bulletin of the American Mathematical Society*, vol. 49, p. 1-23, 1943. 72
- [76] P. CRAVEN et G. WAHBA, « Smoothing noisy data with spline functions: Estimating the correct degree of smoothing by the method of generalized cross-validation », *Numerical Mathematics*, vol. 31, p. 377-403, 1979. 112
- [77] G. CROSS et A. JAIN, « Markov random field texture models », *IEEE Trans. Pattern Anal. Machine Intell.*, vol. 5, p. 25-39, 1983. 10
- [78] M. CROUSE, R. NOWAK et R. BARANIUK, « wavelet-based statistical signal processing using hidden Markov models », *IEEE Trans. Signal Processing*, vol. 46, p. 886-902, 1998. 10
- [79] Z. CVETKOVIC et M. VETTERLI, « Discrete-time wavelet extrema representation: Design and consistent reconstruction », *IEEE Trans. Signal Processing*, vol. 43, n° 3, p. 681-693, mars 1995. 160
- [80] D. VAN DE VILLE, T. BLU et M. UNSER, « Recursive filtering for splines on hexagonal lattices », *Proc. of IEEE ICASSP*, vol. 3, p. 301-304, 2003. 83
- [81] W. DAHMEN et C. A. MICHELLI, « Subdivision algorithms for the generation of box-spline surfaces », *Comput. Aided Geom. Design*, vol. 1, p. 115-129, 1984. 73
- [82] W. DAHMEN et C. A. MICHELLI, « Line average algorithm: A method for the computer generation of smooth surfaces », *Comput. Aided Geom. Design*, vol. 2, p. 77-85, 1985. 73
- [83] I. DAUBECHIES, *Ten Lectures on Wavelets*, Society for Industrial and Applied Mathematics (SIAM), Philadelphia, PA, 1992. 28, 31, 129, 182
- [84] C. de BOOR, *A Practical Guide to Splines*, Springer-Verlag, New York, 1978. 17, 97
- [85] C. de BOOR et G. FIX, « Spline approximation by quasi-interpolants », *J. Approx. Theory*, vol. 8, p. 19-45, 1973. 41
- [86] C. de BOOR, K. HÖLLIG et S. RIEMENSCHNEIDER, *Box Splines*, vol. Applied Mathematical Sciences, vol. 98, Springer-Verlag, Berlin, 1993. 71

- [87] R. DEVORE et A. RON, « Developing a computation-friendly mathematical foundation for spline functions », *SIAM News*, vol. 38, n° 4, mai 2005. 71
- [88] N. A. DODGSON, « Quadratic interpolation for image resampling », *IEEE Trans. Image Processing*, vol. 6, n° 9, p. 1322-1326, sept. 1997. 24
- [89] T. DOLMIÈRE, « évaluation subjective de la qualité des images », Mémoire de Master, École Nationale Supérieure de Physique de Grenoble, 2006. 218
- [90] E. DUBOIS, « The sampling and reconstruction of time-varying imagery with application in video systems », *Proc. IEEE*, vol. 73, p. 502-522, avr. 1985. 8, 58
- [91] J. DUCHON, « Splines minimizing rotation-invariant semi-norms in Sobolev spaces », W. SCHEMPP et K. ZELLER, éd., *Constructive Theory of Functions of Several Variables*, p. 85-100, Springer-Verlag, Berlin-Heidelberg, 1977. 96
- [92] M. ELAD et A. FEUER, « Restoration of single superresolution image from several blurred, noisy, and undersampled measured images », *IEEE Trans. Image Processing*, vol. 6, p. 1646-1658, déc. 1997. 150
- [93] Y. C. ELDAR, « Sampling without input constraints: Consistent reconstruction in arbitrary spaces », A. I. ZAYED et J. J. BENEDETTO, éd., *Sampling, Wavelets and Tomography*, p. 33-60, Birkhauser, Boston, MA, 1998. 29
- [94] Y. C. ELDAR, « Sampling and reconstruction in arbitrary spaces and oblique dual frame vectors », *J. Fourier Analys. Appl.*, vol. 1, n° 9, p. 77-96, jan. 2003. 29
- [95] Y. C. ELDAR et T. G. DVORKIND, « A minimum squared-error framework for generalized sampling », *IEEE Trans. Signal Processing*, vol. 54, n° 6, p. 2155-2167, juin 2006. 35, 36, 51
- [96] Y. C. ELDAR et M. UNSER, « Nonideal sampling and interpolation from noisy observations in shift-invariant spaces », *IEEE Trans. Image Processing*, vol. 54, n° 7, p. 2636-2651, juil. 2006. 29, 31, 33, 54, 112
- [97] A. ENTEZARI, R. DYER et T. MÖLLER, « Linear and cubic box splines for the body centered cubic lattice », *Proc. of IEEE Visualization*, p. 11-18, Austin, TX, oct. 2004. 89
- [98] R. L. EUBANK, *Nonparametric Regression and Spline Smoothing*, Marcel Dekker, New York, 1999. 96, 97
- [99] H. G. FEICHTINGER et K. GRÖCHENIG, « Iterative reconstruction of multivariate bandlimited functions from irregular sampling values », *SIAM J. Math. Anal.*, vol. 231, p. 244-261, 1992. 95
- [100] H. G. FEICHTINGER, K. GRÖCHENIG et T. STROHMER, « Efficient numerical methods in non-uniform sampling theory », *Numer. Math.*, vol. 69, n° 4, p. 423-440, 1995. 95
- [101] A. P. FITZ et R. J. GREEN, « Fingerprint classification using a hexagonal fast Fourier transform », *Pattern Recognition*, vol. 29, n° 10, p. 1587-1597, 1996. 57

- [102] W. T. FREEMAN, T. R. JONES et E. C. PASZTOR, « Example-based super-resolution », *IEEE Comp. Graph. and Appl.*, vol. 22, n° 2, p. 56-65, mars-avril 2002. 159, 160
- [103] S. GEMAN et D. GEMAN, « Stochastic relaxation, Gibbs distribution and the Bayesian restoration of images », *IEEE Trans. Pattern Anal. Machine Intell.*, vol. 6, n° 6, p. 712-741, 1984. 10
- [104] E. D. GIORGI et L. AMBROSIO, « Un nuovo tipo di funzionale del calcolo delle variazioni », *Atti. Accad. Naz. Lincei*, vol. 8, p. 199-210, 1988. 10
- [105] C. A. GLASBEY, « Optimal linear interpolation of images with known point spread function », *Proc. of Scandinavian Image Analysis Conference (SCIA)*, p. 161-168, 2001. 10
- [106] M. GOLOMB et H. F. WEINBERGER, « Optimal approximation and error bounds », R. E. LANGER, éd., *On numerical approximation*, p. 117-190, University of Wisconsin Press, Madison, WI, 1959. 19
- [107] G. H. GOLUB et C. F. Van LOAN, *Matrix Computations*, Johns Hopkins Univ. Press, Baltimore, MD, 3^e édition, 1996. 103
- [108] H. GREENSPAN, C. H. ANDERSON et S. AKBER, « Image enhancement by nonlinear extrapolation in frequency space », *IEEE Trans. Image Processing*, vol. 9, n° 6, p. 1035-1048, 2000. 159
- [109] K. GRÖCHENIG, « Reconstruction algorithms in irregular sampling », *Math. Comput.*, vol. 59, n° 199, p. 181-194, 1992. 95
- [110] K. GRÖCHENIG et H. RAZAFINJATOVO, « On Landau's necessary density conditions for sampling and interpolation of band-limited functions », *J. Lond. Math. Soc.*, vol. 54, n° 3, p. 557-567, 1996. 95
- [111] K. GRÖCHENIG et H. SCHWAB, « Fast local reconstruction methods for nonuniform sampling in shift invariant spaces », *SIAM J. Matrix Anal. and Appl.*, vol. 24, n° 4, p. 899-913, 2003. 108
- [112] F. GUICHARD et F. MALGOUYRES, « Total variation based interpolation », *Proc. of EUSIPCO*, vol. 3, p. 1741-1744, 1998. 150, 155, 157
- [113] F. GUICHARD, L. MOISAN et J.-M. MOREL, « A review of P.D.E. models in image processing and image analysis », *Journal de Physique IV*, vol. 12, p. 137-154, 2002. 156
- [114] J. HADAMARD, *Lectures on Cauchy's Problem in Linear Partial Differential Equations*, Dover Pub. Inc., New York, 1952. 172
- [115] M. S. HANDCOCK et J. R. WALLIS, « An approach to statistical spatial-modeling of meteorological fields (with discussion) », *Journal of the American Statistical Association*, vol. 89, p. 368-390, 1994. 10

- [116] K. P. HONG, J. K. PAIK, H. J. KIM et C. H. LEE, « An edge preserving image interpolation system for a digital camcorder », *IEEE Trans. Cons. Elec.*, vol. 41, n° 3, p. 279-284, août 1996. 162
- [117] H. HU, P. M. HOFMAN et G. de HAAN, « Image interpolation using classification-based neural networks », *Proc. of ISCE*, p. 133-137, sept. 2004. 165
- [118] M. F. HUTCHINSON et F. R. De HOOG, « Smoothing noisy data with spline functions », *Numer. Math.*, vol. 4, n° 7, p. 99-106, 1985. 112
- [119] ITU-R, « Methodology for the subjective assessment of the quality of television pictures », 2000, Recommendation ITU-R BT.500-7. 218
- [120] M. JACOB, T. BLU et M. UNSER, « Efficient energies and algorithms for parametric snakes », *IEEE Trans. Image Processing*, vol. 13, n° 9, p. 1231-1244, sept. 2004. 17
- [121] A. K. JAIN, *Fundamentals of Digital Image Processing*, Prentice Hall, Englewood Cliffs, NJ, 1989. 8
- [122] N. S. JAYANT, J. D. JOHNSTON et R. J. SAFRANEK, « Signal compression based on models of human perception », *Proc. IEEE*, vol. 81, n° 10, p. 1385-1422, oct. 1993. 10
- [123] K. JENSEN et D. ANASTASSIOU, « Subpixel edge localization and the interpolation of still images », *IEEE Trans. Image Processing*, vol. 4, n° 3, p. 265-295, mars 1995. 164, 165, 184, 188
- [124] R.-Q. JIA, « Shift-invariant spaces and linear operator equations », *Israel J. Math.*, vol. 41, p. 259-288, 1998. 28
- [125] H. JIANG et C. MOLONEY, « A new direction adaptive scheme for image interpolation », *Proc. of IEEE ICIP*, vol. 3, p. 369-372, 2002. 154, 155
- [126] S. JUNG, R. THEWES, T. SCHEITER, K. F. GOSER et W. WEBER, « Low-power and high-performance CMOS fingerprint sensing and encoding architecture », *IEEE J. Solid-State Circuits*, vol. 34, p. 978-984, juil. 1999. 57
- [127] N. B. KARAYIANNIS et A. N. VENETSANOPOULOS, « Image interpolation based on variational principles », *Signal Processing*, vol. 25, p. 259-288, 1992. 155
- [128] G. KARLSSON et M. VETTERLI, « Theory of two-dimensional multirate filter banks », *IEEE Trans. Acoust., Speech, Signal Processing*, vol. 38, p. 925-937, juin 1990. 175
- [129] R. G. KEYS, « Cubic convolution interpolation for digital image processing », *IEEE Trans. Acoust., Speech, Signal Processing*, vol. 29, p. 1153-1160, 1981. 24
- [130] K. KINEBUCHI, D. D. MURESAN et T. W. PARKS, « Image interpolation using wavelet-based hidden Markov trees », *Proc. of IEEE ICASSP*, 2001. 160
- [131] W. KNOX CAREY, D. B. CHUANG et S. S. HEMAMI, « Regularity-preserving image interpolation », *Proc. of IEEE ICIP*, p. 901-908, 1997. 160

- [132] W. KNOX CAREY, D. B. CHUANG et S. S. HEMAMI, « Regularity-preserving image interpolation », *IEEE Trans. Image Processing*, vol. 8, n° 9, p. 1293-1297, sept. 1999. 160
- [133] L. KOBBELT, « Stable evaluation of box-splines », *Numerical Algorithms*, vol. 14, n° 4, p. 377-382, 1997. 73, 79
- [134] A. KOLMOGOROV, *Interpolation and Extrapolation of stationary random sequences*, vol. 5, p. 3-14, Izv. Akad. Nauk SSSR Ser. Math., 1941. 15
- [135] T. KONDO, T. FUJIWARA, Y. OKUMURA et Y. NODE, « Picture conversion apparatus, picture conversion method, learning apparatus and learning method », Patent 6,323,905, Sony Corp., United States, nov. 2001. 165
- [136] J. KYBIC, T. BLU et M. UNSER, « Generalized sampling: A variational approach – Part I: Theory », *IEEE Trans. Signal Processing*, vol. 50, n° 8, p. 1965-1976, août 2002. 154
- [137] J. KYBIC, T. BLU et M. UNSER, « Generalized sampling: A variational approach – Part II: Applications », *IEEE Trans. Signal Processing*, vol. 50, n° 8, p. 1977-1985, août 2002. 154
- [138] M.-J. LAI, « Fortran subroutines for B-nets of box splines on three and four directional meshes », *Numerical Algorithms*, vol. 2, p. 33-38, 1992. 73, 79
- [139] C. LEE, M. EDEN et M. UNSER, « High-quality image resizing using oblique projection operators », *IEEE Trans. Image Processing*, vol. 7, n° 5, p. 679-692, mai 1998. 132
- [140] P. LEGRAND et J. L. VÉHEL, « Interpolation de signaux par conservation de la régularité Höldérienne », *Actes du GRETSI*, 2003. 160
- [141] T. M. LEHMANN, C. GÖNNER et K. SPITZER, « Survey: Interpolation methods in medical image processing », *IEEE Trans. Med. Imag.*, vol. 18, n° 11, p. 1049-1075, nov. 1999. 27
- [142] W. LI et A. FETTWEIS, « Interpolation filters for 2-D hexagonally sampled signals », *Int. J. Circuit Theory Applic.*, vol. 25, p. 259-277, 1997. 57
- [143] X. LI et M. T. ORCHARD, « New edge-directed interpolation », *IEEE Trans. Image Processing*, vol. 10, n° 10, p. 1521-1527, oct. 2001. 162, 164, 166
- [144] Y. LIU, « Irregular sampling for spline wavelet subspaces », *IEEE Trans. Inform. Theory*, vol. 42, n° 2, p. 623-627, 1996. 28, 95
- [145] A. K. LOUIS, « A unified approach to regularization methods for linear ill-posed problems », *Inverse Problems*, vol. 15, p. 489-498, 1999. 16
- [146] S. MALASSIOTIS et M. G. STRINTZIS, « Optimal biorthogonal wavelet decomposition of wire-frame meshes using box splines, and its application to the hierarchical coding of 3-D surfaces », *IEEE Trans. Image Processing*, vol. 8, n° 1, p. 41-57, jan. 1999. 71, 72

- [147] F. MALGOUYRES, « Total variation based oversampling of noisy images », *Proc. of Scale-Space*, p. 111-122, Lecture Notes in Computer Science 2106, juil. 2001. 155, 158
- [148] F. MALGOUYRES et F. GUICHARD, « Edge direction preserving image zooming: A mathematical and numerical analysis », *SIAM J. Numer. Anal.*, vol. 39, n° 1, p. 1-37, 2001. 150, 155, 156
- [149] S. MALLAT, *A wavelet tour of signal processing*, Academic Press, 2^{de} édition, 1999. 28, 31, 124, 129
- [150] S. MALLAT et S. ZHONG, « Characterization of signals from multiscale edges », *IEEE Trans. Pattern Anal. Machine Intell.*, vol. 14, n° 7, p. 2207-2232, juil. 1992. 160
- [151] D. MARR, *Vision*, Freeman, San Francisco, CA, 1982. 10
- [152] E. MEIJERING, « A chronology of interpolation », *Proc. IEEE*, vol. 90, n° 3, p. 319-342, mars 2002. 21, 23, 27
- [153] E. MEIJERING, W. NIESSEN, J. PLUIM et M. VIERGEVER, « Quantitative comparison of sinc-approximating kernels for medical image interpolation », *Proc. of MICCAI*, p. 210-217, 1999. 23
- [154] E. MEIJERING, W. NIESSEN et M. VIERGEVER, « Quantitative evaluation of convolution-based methods for medical image interpolation », *Medical Image analysis*, vol. 5, n° 2, p. 111-126, juin 2001. 23, 27
- [155] E. H. W. MEIJERING, K. J. ZUIDERVELD et M. A. VIERGEVER, « Image reconstruction by convolution with symmetrical piecewise n th-order polynomial kernels », *IEEE Trans. Image Processing*, vol. 8, n° 2, p. 192-201, Feb. 1999. 23
- [156] R. M. MERSEREAU, « The processing of hexagonally sampled two-dimensional signals », *Proc. IEEE*, vol. 67, n° 6, p. 930-949, juin 1979. 57, 70
- [157] C. A. MICHELLI et T. J. RIVLIN, éd., *Optimal Estimation in Approximation Theory*, chap. A Survey of Optimal Recovery, p. 1-54, Plenum, New York, 1976. 18, 19
- [158] L. MIDDLETON et J. SIVASWAMY, « Edge detection in a hexagonal-image processing framework », *Image and Vision Computing*, vol. 19, n° 14, p. 1071-1081, déc. 2001. 57
- [159] L. MOISAN, « Extrapolation de spectre et variation totale pondérée », *Actes du GRETSI*, 2001. 155
- [160] B. S. MORSE et D. SCHWARTZWALD, « Isophote-based interpolation », *Proc. of IEEE ICIP*, vol. 3, p. 227-231, oct. 1998. 154, 155, 156
- [161] B. S. MORSE et D. SCHWARTZWALD, « Level-set image reconstruction », *Proc. of CVPR*, p. 333-340, 2001. 154, 155

- [162] A. MUÑOZ BARRUTIA, T. BLU et M. UNSER, « Non-uniform to uniform grid conversion using least-squares splines », *Proc. of EUSIPCO*, vol. 4, p. 1997-2000, sept. 2000. 98
- [163] A. MUÑOZ BARRUTIA, T. BLU et M. UNSER, « ℓ_p -multiresolution analysis: How to reduce ringing and sparsify the error », *IEEE Trans. Image Processing*, vol. 11, n° 6, p. 656-669, juin 2002. 27, 176
- [164] D. MUMFORD et B. GIDAS, « Stochastic models for generic images », *Quarterly Appl. Math.*, vol. 59, p. 85-111, 2001. 10
- [165] D. D. MURESAN, « Review of optimal recovery », Technical report 14853, Electrical and Computer Engineering department at Cornell University, Ithaca, NY, 2002. 18, 19
- [166] D. D. MURESAN, « Fast edge directed polynomial interpolation », *Proc. of IEEE ICIP*, vol. 2, p. 990-993, 2005. 157, 190
- [167] D. D. MURESAN et T. W. PARKS, « Adaptively quadratic (AQua) image interpolation », *IEEE Trans. Image Processing*, vol. 13, n° 5, p. 690-698, mai 2004. 157, 159, 167, 170
- [168] F. NICOLIER et F. TRUCHETET, « Image magnification using decimated orthogonal wavelet transform », *Proc. of IEEE ICIP*, vol. 2, p. 355-358, sept. 2000. 160
- [169] B. NOUVEL et E. RÉMILA, « Configurations induced by discrete rotations: Periodicity and quasi-periodicity properties », *Discrete Applied Mathematics*, vol. 147, n° 2-3, p. 325-343, avr. 2005. 43
- [170] R. NOWAK, « Multiscale hidden Markov models for Bayesian image analysis », B. VIDAKOVIC et P. MÜLLER, édés., *Bayesian Inference in Wavelet Based Models*, p. 243-266, Springer Verlag, Berlin, 1999. 10
- [171] G. NÜRNBERGER et F. ZEILFELDER, « Developments in bivariate spline interpolation », *J. Comput. Appl. Math.*, vol. 121, n° 1, p. 1-29, 2000. 72
- [172] A. P. PENTLAND, « Fractal-based description of natural scenes », *IEEE Trans. Pattern Anal. Machine Intell.*, vol. 6, p. 661-674, 1984. 10
- [173] S. PERIASWAMY, « Detection of microcalcifications in mammograms using hexagonal wavelets », Master's thesis, University of South Carolina, 1996. 57
- [174] P. PERONA et J. MALIK, « Scale-space and edge detection using anisotropic diffusion », *IEEE Trans. Pattern Anal. Machine Intell.*, vol. 12, p. 629-639, 1990. 157
- [175] D. P. PETERSEN et D. MIDDLETON, « Sampling and reconstruction of wavenumber-limited functions in N -dimensional Euclidean spaces », *Inform. Control*, vol. 5, p. 279-323, 1962. 57
- [176] W. K. PRATT, *Digital Image Processing*, Wiley, New York, 1990. 8

- [177] H. PRAUTZSCH et W. BOEHM, « Box splines », *Handbook of Computer Aided Geometric Design*, Springer, Berlin, 2001. 71, 72
- [178] W. H. PRESS, S. A. TEUKOLSKY, W. T. VETTERLING et B. P. FLANNERY, *Numerical Recipes in C: The Art of Scientific Computing*, Cambridge University Press, New York, 1992. 103, 104, 107
- [179] J. PRICE et M. HAYES III, « Optimal prefiltering for improved image interpolation », *Proc. of 32nd Asilomar Conference on Signals, Systems, and Computers*, vol. 2, p. 959-963, nov. 1998. 137, 150
- [180] D. RAJAN et S. CHAUDHURI, « Generation of super-resolution images from blurred observations using Markov random fields », *Proc. of IEEE ICASSP*, vol. 3, p. 1837-1840, 2001. 150
- [181] R. RAMANATH, W. E. SNYDER, Y. YOO et M. S. DREW, « Color image processing pipeline », *IEEE Signal Processing Mag.*, vol. 22, n° 1, p. 34-43, jan. 2005. 8
- [182] S. RAMANI, D. VAN DE VILLE et M. UNSER, « Sampling in practice: Is the best reconstruction space bandlimited? », *Proc. of IEEE ICIP*, vol. 2, p. 153-156, sept. 2005. 17, 20, 22, 40, 54, 58
- [183] S. RAMANI, D. VAN DE VILLE et M. UNSER, « Non-ideal sampling and adapted reconstruction using the stochastic matérn model », *Proc. of IEEE ICASSP*, mai 2006. 10, 20, 40, 56
- [184] G. RAMPONI, « Warped distance for space-variant linear image interpolation », *IEEE Trans. Image Processing*, vol. 8, n° 5, p. 629-639, mai 1999. 162, 163
- [185] G. RAMPONI et S. CARRATO, « Interpolation of the DC component of coded images using a rational filter », *Proc. of IEEE ICIP*, vol. 1, p. 389-392, oct. 1997. 162, 185
- [186] X. RAN et N. FARVARDIN, « A perceptually motivated three-component image model. Part I: Description of the model — Part II: Application to image compression », *IEEE Trans. Image Processing*, vol. 4, n° 4, p. 401-415, 430-447, avr. 1995. 10
- [187] O. RIOUL et M. VETTERLI, « Wavelets and signal processing », *IEEE Signal Processing Mag.*, vol. 8, n° 4, p. 14-38, oct. 1991. 175
- [188] G. M. ROBBINS et T. S. HUANG, « Inverse filtering for linear shift-invariant imaging systems », *Proc. IEEE*, vol. 60, n° 7, p. 862-871, juil. 1972. 8
- [189] D. F. ROGERS et J. A. ADAMS, *Mathematical Elements for Computer Graphics*, McGraw-Hill, New York, 1990. 2
- [190] L. I. RUDIN, *Images, Numerical Analysis of Singularities and Shock Filters*, PhD thesis, California Institute of Technology, Pasadena, CA, 1987. 10
- [191] L. I. RUDIN, S. OSHER et E. FATEMI, « Nonlinear total variation based noise removal algorithms », *Physica D*, vol. 60, n° 1-4, p. 259-268, 1992. 18, 155
- [192] I. J. SCHOENBERG, « Spline functions and the problem of graduation », *Proc. Nat. Acad. Sci.*, vol. 52, p. 652-654, 1964. 97

- [193] W. F. SCHREIBER, *Fundamentals of Electronic Imaging Systems*, Springer-Verlag, Berlin, 1986. 128
- [194] R. R. SCHULTZ et R. L. STEVENSON, « A Bayesian approach to image expansion for improved definition », *IEEE Trans. Image Processing*, vol. 3, n° 3, p. 233-242, mai 1994. 150, 154, 157
- [195] D. SECREST, « Best integration formulas and best error bounds », *Math. Computat.*, vol. 19, p. 79-83, jan. 1965. 19
- [196] D. SECREST, « Error bounds for interpolation and differentiation by the use of spline functions », *SIAM J. Numer. Anal.*, vol. 2, n° 3, p. 440-447, 1965. 19
- [197] J. SERRA, *Image Analysis and Mathematical Morphology*, Academic Press, London, 1982. 91
- [198] C. E. SHANNON, « A mathematical theory of communication », *The Bell System Technical Journal*, vol. 27, p. 379-423 & 623-656, 1948. 21
- [199] C. E. SHANNON, « Communication in the presence of noise », *Proc. of the Inst. of Radio Eng.*, vol. 37, n° 1, p. 10-21, 1949. 21, 95
- [200] G. SHARMA et H. J. TRUSSEL, « Digital color imaging », *IEEE Trans. Image Processing*, vol. 6, n° 7, p. 901-932, 1997. 8
- [201] R. G. SHENOY et T. W. PARKS, « An optimal recovery approach to interpolation », *IEEE Trans. Signal Processing*, vol. 40, n° 8, p. 1987-1996, août 1992. 18, 19
- [202] E. SIMONCELLI, « Statistical models for images: Compression, restoration and synthesis », *Proc. of 31st ACSSC*, p. 673-678, IEEE Computer Society Press, 1997. 10
- [203] M. SONKA et J. M. FITZPATRICK, éd., *Handbook of Medical Imaging*, vol. 2, chap. 11: Echocardiography, SPIE Press, Washington, 2000. 93
- [204] R. C. STAUNTON, « The design of hexagonal sampling structures for image digitization and their use with local operators », *Image and Vision Computing*, vol. 7, n° 3, p. 162-166, 1989. 57
- [205] G. STRANG et G. FIX, « A Fourier analysis of the finite element variational method », *Constructive aspect of functional analysis*, p. 796-830, Edizioni Cremonese, Rome, Italy, 1971. 26, 60
- [206] T. STROHMER, « Computationally attractive reconstruction of bandlimited images from irregular samples », *IEEE Trans. Image Processing*, vol. 6, n° 4, p. 540-548, avr. 1997. 95
- [207] T. STROHMER, « Numerical analysis of the non-uniform sampling problem », *J. Comp. Appl. Math.*, vol. 112, n° 1-2, p. 297-316, 2000. 95
- [208] D. S. TAUBMAN et M. W. MARCELLIN, *JPEG2000: Image Compression Fundamentals, Standard and Practice*, Kluwer Academic Publishers, 2002. 124, 182

- [209] L. TENZE et S. CARRATO, « A rational $n \times$ image interpolator », *Proc. of EUSIPCO*, sept. 2000. 162
- [210] P. THÉVENAZ, U. E. RUTTIMANN et M. UNSER, « A pyramid approach to subpixel registration based on intensity », *IEEE Trans. Image Processing*, vol. 7, n° 1, p. 27-41, jan. 1998. 113
- [211] P. THÉVENAZ, T. BLU et M. UNSER, « Image interpolation and resampling », I.N. BANKMAN, éd., *Handbook of Medical Imaging, Processing and Analysis*, p. 393-420, Academic Press, San Diego, CA, 2000. 23, 26
- [212] P. THÉVENAZ, T. BLU et M. UNSER, « Interpolation revisited », *IEEE Trans. Med. Imag.*, vol. 19, n° 7, p. 739-758, juil. 2000. 23, 26, 27
- [213] A. N. TIKHONOV et V. Y. ARSENIN, *Solutions of Ill-posed Problems*, Wiley, New York, 1977. 16
- [214] S. TIROSH, D. VAN DE VILLE et M. UNSER, « Polyharmonic smoothing splines and the multi-dimensional Wiener filtering of fractal-like signals », *IEEE Trans. Image Processing*, 2006, sous presse. 10, 20, 22, 40, 56, 58, 175
- [215] B. TOM et A. KATSAGGELOS, « Resolution enhancement of monochrome and color video using motion compensation », *IEEE Trans. Image Processing*, vol. 10, p. 278-287, fév. 2001. 150
- [216] M. TREMBLAY, S. DALLAIRE et D. POUSSART, « Low level segmentation using CMOS smart hexagonal image sensor », *Proc. of Conf. Computer Arch. for Machine Perception*, p. 21-28, 1995. 57
- [217] M. UNSER, « Quasi-orthogonality and quasi-projections », *Applied and Computational Harmonic Analysis*, vol. 3, n° 3, p. 201-214, juil. 1996. 41
- [218] M. UNSER, « Splines: A perfect fit for signal and image processing », *IEEE Signal Processing Mag.*, vol. 16, n° 6, p. 22-38, nov. 1999. 17, 27, 33, 40, 83
- [219] M. UNSER, « Sampling—50 Years after Shannon », *Proc. IEEE*, vol. 88, n° 4, p. 569-587, avr. 2000. 28, 29
- [220] M. UNSER et A. ALDROUBI, « A general sampling theory for nonideal acquisition devices », *IEEE Trans. Signal Processing*, vol. 42, n° 11, p. 1915-2925, nov. 1994. 29
- [221] M. UNSER, A. ALDROUBI et M. EDEN, « B-spline signal processing: Part I: Theory & Part II: Efficient design and applications », *IEEE Trans. Signal Processing*, vol. 41, n° 2, p. 821-848, fév. 1993. 26, 27, 108, 173
- [222] M. UNSER, A. ALDROUBI et M. EDEN, « The L_2 polynomial spline pyramid », *IEEE Trans. Pattern Anal. Machine Intell.*, vol. 15, n° 4, p. 364-379, avr. 1993. 116, 132
- [223] M. UNSER, A. ALDROUBI et M. EDEN, « Enlargement or reduction of digital images with minimum loss of information », *IEEE Trans. Image Processing*, vol. 4, n° 3, p. 247-258, mars 1995. 116, 132

- [224] M. UNSER et T. BLU, « Cardinal exponential splines: Part I—Theory and filtering algorithms », *IEEE Trans. Signal Processing*, vol. 53, n° 4, p. 1425-1438, avr. 2005. 73
- [225] M. UNSER et T. BLU, « Generalized smoothing splines and the optimal discretization of the Wiener filter », *IEEE Trans. Signal Processing*, vol. 53, n° 6, p. 2146-2159, juin 2005. 17, 20, 73
- [226] M. UNSER et J. ZERUBIA, « A generalized sampling theory without band-limiting constraints », *IEEE Trans. Circuits Syst. II*, vol. 45, n° 8, p. 959-969, août 1998. 29
- [227] P. P. VAIDYANATHAN, *Multirate Systems and Filter Banks*, Prentice Hall, Englewood Cliffs, NJ, 1993. 118
- [228] P. P. VAIDYANATHAN et B. VRCELJ, « Biorthogonal partners and applications », *IEEE Trans. Signal Processing*, vol. 49, n° 5, p. 1013-1027, mai 2001. 31, 150
- [229] D. VAN DE VILLE, T. BLU, M. UNSER, W. PHILIPS, I. LEMAHIEU et R. VAN DE WALLE, « Hex-spline: A novel family for hexagonal lattices », *IEEE Trans. Image Processing*, vol. 13, n° 6, p. 758-772, June 2004. 59, 62, 63, 77, 79, 83, 90
- [230] D. VAN DE VILLE, L. CONDAT, T. BLU et M. UNSER, « Error estimates for generalized multidimensional sampling/interpolation », *IEEE Trans. Image Processing*, 2006, en préparation. 91
- [231] D. VAN DE VILLE, W. PHILIPS et I. LEMAHIEU, « Least-squares spline resampling to a hexagonal lattice », *Signal Processing: Image Communication*, vol. 17, n° 5, p. 393-408, mai 2002. 62
- [232] A. VAN DER SCHAAF et J. VAN HATEREN, « Modelling the power spectra of natural images », *Vision research*, vol. 36, n° 17, p. 2759-2770, 1996. 10
- [233] C. VÁZQUEZ, E. DUBOIS et J. KONRAD, « Reconstruction of nonuniformly sampled images in spline spaces », *IEEE Trans. Image Processing*, vol. 14, n° 6, p. 713-725, juin 2005. 108, 115
- [234] M. VETTERLI, « A theory of multirate filter banks », *IEEE Trans. Acoust., Speech, Signal Processing*, vol. 35, p. 356-372, avr. 1987. 175
- [235] D. Van De VILLE, T. BLU et M. UNSER, « Isotropic polyharmonic b-splines: Scaling functions and wavelets », *IEEE Trans. Image Processing*, vol. 14, n° 11, p. 1798-1813, nov. 2005. 175
- [236] M. VRHEL, E. SABER et H. J. TRUSSELL, « Color image generation and display technologies », *IEEE Signal Processing Mag.*, vol. 22, n° 1, p. 23-33, jan. 2005. 8
- [237] G. WAHBA, *Spline Models for Observational Data*, Society for Industrial and Applied Mathematics (SIAM), Philadelphia, PA, 1990. 96, 97
- [238] J. WALLIS, *Arithmetica Infinitorum*, Oxford, 1655. 21
- [239] Z. WANG, A. C. BOVIK et L. LU, « Why is image quality assessment so difficult? », *Proc. of IEEE ICASSP*, vol. 4, p. 3313-3316, mai 2002. 144

- [240] H. WENDLAND, *Scattered Data Approximation*, Cambridge University Press, 2005. 96
- [241] E. T. WHITTAKER, « On the functions which are represented by the expansions of interpolation-theory », *Proc. of the Royal Society of Edinburgh*, vol. 35, p. 181-194, 1915. 21
- [242] N. WIENER, *The Extrapolation, Interpolation and Smoothing of Stationary Time Series*, John Wiley & Sons, New York, 1949. 15
- [243] G. WOLBERG, *Digital Image Warping*, IEEE Computer Society Press, Los Alamitos, CA, 1990. 8, 114, 128
- [244] Q. WU, X. HE et T. HINTZ, « Bilateral filtering based edge detection on hexagonal architecture », *Proc. of IEEE ICASSP*, vol. 2, p. 713-716, 2005. 57
- [245] G. WYSZECKI et W. S. STILES, *Color Science: Concepts and Methods, Quantitative, Data and Formulae*, Wiley, New York, 2^{de} édition, 1982. 7
- [246] Y. YE et J. ZHU, « Geometric studies on variable radius spiral comebeam scanning », *Medical Physics*, vol. 31, p. 1473-1480, 2004. 93
- [247] D. C. YOULA et H. WEBB, « Image restoration by the method of convex projections: Part 1—Theory », *IEEE Trans. Med. Imag.*, vol. MI-1, n° 2, p. 81-94, oct. 1982. 179
- [248] X. YU, B. S. MORSE et T. W. SEDERBERG, « Image reconstruction using data-dependent triangulation », *IEEE Computer Graphics and Applications*, vol. 21, n° 3, p. 62-68, 2001. 161, 162
- [249] Y. ZHU, S. C. SCHWARTZ et M. T. ORCHARD, « Wavelet domain image interpolation via statistical estimation », *Proc. of IEEE ICIP*, 2001. 160
- [250] F. ZOTTA, P. CARRAI, G. FERRETTI et G. RAMPONI, « A directional edge-sensitive image interpolator », *Proc. of SPIE Int. Symp. Elec. Imaging*, jan. 2004. 162, 163

Méthodes d'approximation pour la reconstruction de signaux et le redimensionnement d'images

Résumé

Ce travail porte sur la reconstruction de signaux et le redimensionnement d'images. La reconstruction vise à estimer une fonction dont on ne connaît que des mesures linéaires éventuellement bruitées. Par exemple, le problème d'interpolation uniforme consiste à estimer une fonction $s(t)$ définie sur l'ensemble des réels, n'en connaissant que les valeurs $s(k)$ aux entiers k . L'approche proposée est originale et consiste à effectuer une quasi-projection dans un espace fonctionnel fixé, en minimisant l'erreur d'approximation asymptotique lorsque le pas d'échantillonnage tend vers zéro. Les cas 1D, 2D cartésien, et 2D hexagonal sont évoqués. Nous appliquons ensuite notre formalisme au problème d'agrandissement des images, pour lequel seules des méthodes non-linéaires s'avèrent à même de synthétiser correctement l'information géométrique à laquelle nous sommes le plus sensibles. Nous proposons une méthode appelée induction, à la fois simple, rapide et performante.

Mots clés : traitement d'images et de signaux, reconstruction, interpolation, approximation, quasi-projections, splines, problèmes linéaires inverses, réduction, agrandissement, redimensionnement.

Abstract

This work addresses the problems of signal reconstruction and image resizing. Reconstruction aims at estimating a function given only (eventually noisy) linear measurements performed on it at discrete locations. For instance, uniform interpolation is the problem of estimating $s(t)$ on the real line, given the values $s(k)$ at the integers k . The proposed approach is original, and consists in performing a quasi-projection in a given linear shift-invariant reconstruction space, such that the decay of the approximation error is maximal as the sampling step tends to zero. We illustrate our method in the 1D, 2D Cartesian and 2D hexagonal cases. We then apply our formalism to the problem of image resizing, for which only non-linear methods are able to correctly synthesize the geometric information to which the human visual system is sensitive. We propose a method called induction, which is simple, fast and efficient.