

N°d'ordre : 435

N°attribué par la bibliothèque : ENSL07435

THESE

en vue d'obtenir le grade de

Docteur de l'Université de Lyon - École Normale Supérieure de Lyon

Spécialité : Physique

Laboratoire de Physique

École doctorale de Physique



présentée et soutenue publiquement le 14 décembre 2007 par

Gabriel Rilling

Titre :

Décompositions Modales Empiriques

Contributions à la théorie, l'algorithmie et l'analyse de performances

Directeur de thèse : M. Patrick FLANDRIN

Après avis de : M. Jacques LEMOINE

Mme Valérie PERRIER

devant la commission d'examen formée de

M. Alain ARNEODO

Président

M. Patrick FLANDRIN

Directeur de thèse

M. Jacques LEMOINE

Membre/Rapporteur

Mme Valérie PERRIER

Membre/Rapporteur

Mme Béatrice PESQUET-POPESCU Membre

Remerciements

Les travaux rapportés dans ce manuscrit doivent beaucoup au soutien d'un certain nombre de personnes. Je ne les ai pour la plupart jamais vraiment remerciés et je souhaite donc le faire ici avec la plus grande sincérité.

Je tiens tout d'abord à remercier les membres du jury qui m'ont fait l'honneur d'évaluer ce travail. Je remercie Alain Arnéodo d'avoir apporté son regard personnel et de m'avoir fait le plaisir de présider le jury. Merci à Valérie Perrier et Jacques Lemoine d'avoir pris le temps de soigneusement étudier ce manuscrit et de le rapporter en me faisant part de leurs nombreuses questions et recommandations. Je remercie également Béatrice Pesquet-Popescu d'avoir accepté de participer à ce jury malgré les nombreuses sollicitations et d'avoir ainsi apporté son soutien à mes travaux.

Plus que tout autre, je veux remercier Patrick Flandrin sans qui cette thèse n'aurait jamais été. Scientifiquement, les travaux exposés ici doivent beaucoup à la justesse de ses conseils, à la pertinence de ses réflexions ainsi qu'au recul dont il est capable dans ses analyses. À cela s'ajoutent ses qualités de direction, notamment la grande liberté qu'il m'a laissée ainsi que sa constante disponibilité. Il a aussi par sa présence amicale et bienveillante contribué à faire de ces années une expérience toujours enrichissante et agréable. Enfin j'ai également beaucoup apprécié ses qualités humaines et en particulier sa patience dont j'ai sans doute éprouvé les limites. Pour toutes ces raisons, je lui suis infiniment reconnaissant.

Ce travail aurait aussi sûrement été très différent sans le cadre agréable du laboratoire de physique de l'ENS Lyon. En premier lieu, je remercie les membres de l'équipe Sisyphe qui formaient mon entourage direct et sont devenus mes amis : Pierre, dont la générosité et la disponibilité semblent inépuisables, Stéphane, bonne humeur et chaleur humaine personnifiées, Patrice, « le chef cool », Herwig, Jun et Bruno avec qui j'ai eu le plaisir de partager le bureau des thésards ainsi que Nico et Antoine, qui résidaient un peu plus loin mais répondaient toujours présent à l'appel du café. Au delà de l'équipe signal, je remercie également les autres partenaires de mes (très) nombreuses pauses, Yoann, Mathieu, Éric, Hervé ainsi que Mathieu et Sébastien, amis travailleurs du soir et de la nuit (qui m'ont plus souvent empêché de travailler qu'autre chose).

Plus généralement, je salue l'ensemble des membres du laboratoire, sans oublier les personnels administratifs et techniques, toujours prompts à résoudre les moindre problèmes.

Il me reste enfin à remercier mes parents et mon frère, à qui je dois la curiosité et la rigueur qui m'ont amené jusqu'ici, et dont le soutien même tacite m'a été précieux au long de ces années.

Table des matières

1	Décompositions modales empiriques	5
1	Motivations et contexte	5
1.1	Contexte	5
1.2	Méthodes dévolues à l'analyse des signaux non stationnaires	8
2	Présentation des décompositions modales empiriques	14
2.1	La décomposition modale empirique originale	14
2.2	Transformée de Hilbert–Huang (HHT)	18
2.3	Définition plus générale d'une EMD	18
2.4	L'EMD : une méthode « intuitive »	20
3	Propriétés élémentaires de la décomposition	21
3.1	Non-linéarité	21
3.2	Covariances vis-à-vis de certaines transformations	22
3.3	Quasi-orthogonalité	23
3.4	Localité	26
3.5	Multirésolution	30
4	Propriétés des IMFs	34
4.1	Une définition en toute rigueur trop restrictive	34
4.2	Les IMFs d'indice supérieur à 2 sont des splines	37
5	Questions liées à l'implantation	38
5.1	Conditions aux bords	38
5.2	Arrêt du processus de tamisage	45
5.3	Interpolation	47
5.4	Problèmes liés à l'échantillonnage	54
6	Variantes et extensions	69
6.1	EMD locale	69
6.2	EMD en ligne	71
6.3	« Ensemble Empirical Mode Decomposition »	74
6.4	EMD pour les images	76
6.5	Signal bivarié	76
2	Analyse	97
1	Cadre déterministe	97
1.1	Cas d'un mélange de 2 sinusoïdes	97
1.2	Généralisations	136
1.3	Cas d'un mélange à $n > 2$ ondes	150
1.4	Application à l'EMD bivariée	150
1.5	Conclusions sur l'EMD de sommes de composantes déterministes « simples »	161
2	Cadre stochastique : décomposition d'un bruit large bande	164
2.1	Bruits blancs de densités variées	164

2.2	Bruit gaussien fractionnaire — estimation de l'exposant d'une loi d'échelle . . .	186
2.3	EMD bivariée d'un bruit « blanc » complexe	217
2.4	Liens avec le modèle déterministe	250

Introduction

Dans tous les domaines de sciences, à l'exception peut-être de certaines branches des mathématiques, la connaissance progresse par la conjonction d'une approche théorique ou de modélisation et d'une approche d'observation du sujet d'étude. Dans le cadre de la seconde, un souci constant est de s'assurer que ce qui est observé est bien relatif à l'objet étudié. Pour étudier par exemple les phénomènes de marées, on peut s'intéresser à l'évolution du niveau de la mer en un point donné en fonction du temps. Cependant, cette évolution ne dépend pas que des marées mais peut être influencée par beaucoup d'autres phénomènes comme des ondes de gravité, des courants marins, des ondes de surface dues au vent, le passage de bateaux, la vie aquatique, etc. Pour pouvoir étudier les marées, il en découle que l'observateur doit nécessairement faire la part dans les observations de ce qui est relatif aux marées et de ce qui ne l'est pas. Plus généralement, dans toutes les situations d'observation se pose le problème d'extraire les informations « qui nous intéressent », c'est-à-dire celles qui sont pertinentes en vue d'une modélisation ou d'une application donnée. Dans certains cas, les données expérimentales sont suffisamment simples pour que le cerveau de l'observateur puisse extraire ces informations. Dans d'autres, il peut avoir besoin de l'aide d'outils d'analyse de données.

Les travaux rapportés dans ce manuscrit ont pour sujet central un nouvel outil d'analyse de données dévolu à des situations difficiles où les outils classiques sont mal adaptés. Plus qu'un outil, il s'agit en fait d'une approche novatrice qui, à partir d'un outil appelé « Empirical Mode Decomposition » [28], a donné lieu par la suite à une famille d'outils, les *décompositions modales empiriques*, variantes ou extensions de la méthode originale. Séduisante par son aspect particulièrement intuitif, la nouvelle approche souffre d'un défaut de jeunesse dans la mesure où elle n'est encore définie que par la sortie d'un algorithme complexe sans fondement théorique bien établi. En outre, la mise en œuvre de l'algorithme nécessite un certain nombre de choix pour lesquels les solutions proposées initialement sont fonctionnelles mais clairement pas optimales. Enfin, le principe de base de la méthode est certes très intuitif mais le fait qu'il soit inhabituel, que sa mise en œuvre dépende d'un certain nombre de degrés de liberté, et qu'il soit de plus utilisé de manière itérée fait que la sortie de l'algorithme n'est pas forcément si intuitive qu'on pourrait le penser a priori. Néanmoins, les propriétés uniques de cette nouvelle approche ont rapidement suscité l'intérêt dans des domaines d'application divers avant d'intéresser la communauté traitement du signal.

Lorsque les tous premiers travaux à l'origine de cette thèse ont été commencés en 2002 à l'occasion d'un stage, la méthode avait commencé à être diffusée en dehors du domaine de l'océanographie et de la climatologie où elle avait été introduite, mais rien ou presque n'avait encore été fait, ni pour proposer des choix plus performants pour les degrés de liberté de l'algorithme, ni pour en améliorer la compréhension.

Au cours de cette thèse, c'est précisément cet objectif de comprendre le fonctionnement de la décomposition modale empirique qui a été le moteur principal avec l'objectif annexe de lui apporter une base théorique. Il a fallu pour cela commencer par réaliser une implantation de l'algorithme. Cette étape a été l'occasion d'aborder les problématiques liées aux différents choix algorithmiques nécessaires mais on ne s'est pas attardé outre mesure sur ces questions, l'objectif étant plus d'avoir un algorithme fonctionnel rapidement pour commencer à analyser son fonctionnement. De plus, il paraît difficile de construire une implantation optimale à tout point de vue si on ne sait pas quel sens

donner à la notion d'optimalité, faute d'avoir une compréhension précise du fonctionnement de la méthode. L'implantation réalisée alors a été publiée sur internet en tant que logiciel libre et constitue la base d'une petite boîte à outils pour **MATLAB** consacrée aux décompositions modales empiriques qui a été développée au cours de cette thèse.

Une fois cette implantation réalisée, on a pu se pencher sur l'objectif premier qu'est l'étude des propriétés de la méthode. Dans ce but, on peut envisager schématiquement deux approches : une approche théorique consistant à partir de l'algorithme pour construire un cadre théorique adapté permettant par la suite d'en décrire les propriétés, ou une approche empirique consistant à déterminer les propriétés de l'algorithme à partir d'observations de son comportement dans des situations bien contrôlées. Pour de nombreuses raisons, la décomposition modale empirique se montre rêtive à l'approche théorique. Sa forme algorithmique aux nombreux degrés de liberté est clairement un obstacle important mais même l'élément central de l'algorithme, décrit par la succession de seulement quelques opérations, est suffisamment inhabituel pour constituer un obstacle majeur. Face à ces difficultés, l'approche empirique s'est naturellement imposée dans un premier temps. N'ayant a priori aucune connaissance précise des propriétés de l'algorithme, on a choisi d'étudier son comportement dans un petit nombre de situations simplifiées à l'extrême pour faciliter l'interprétation. On est ainsi parti de l'étude de la décomposition modale empirique de trois types de signaux : les signaux sinusoïdaux, les sommes de deux signaux sinusoïdaux et les bruits blancs gaussiens. La première situation a permis d'aborder la thématique de l'influence de l'échantillonnage. Les deux suivantes ont permis d'approcher le fonctionnement de la décomposition par des points de vue complémentaires : d'une part le cas de deux composantes déterministes simples et d'autre part le cas du bruit blanc qui peut-être vu comme la somme d'un grand nombre de composantes déterministes indépendantes. Dans tous les cas, des simulations extensives ont été réalisées pour observer l'influence des différents paramètres des modèles de signaux ainsi que celle des paramètres clés de l'algorithme de décomposition modale empirique. Les caractéristiques observées dans les deux premières situations ont alors motivé un certain nombre de simplifications de l'algorithme permettant d'aboutir dans chaque cas à une modélisation de son comportement justifiant en partie les observations et pouvant s'étendre à des cas plus généraux. La troisième situation s'est montrée plus difficile à approcher directement. Les résultats des simulations sont cependant riches d'observations et clairement pas déconnectés de ceux observés pour les sommes de deux signaux sinusoïdaux.

Indépendamment de l'objectif principal qu'est l'analyse du fonctionnement de la décomposition modale empirique, on s'est également intéressé à une extension permettant de traiter des signaux à valeur complexe, la méthode initiale étant limitée aux signaux réels. Ces travaux, initialement motivés par une collaboration en océanographie, ont abouti à la conception de nouveaux algorithmes étendant le principe de l'algorithme initial aux signaux complexes. Les propriétés de ces derniers sont étudiées au cours de cette thèse en parallèle de celles de la méthode originale.

Ce document est organisé de la façon suivante. Après une introduction rappelant le contexte dans lequel elles ont été introduites les décompositions modales empiriques, le chapitre 1 propose une présentation de la famille des décompositions modales empiriques ainsi que des considérations relatives aux questions liées à l'implantation ou aux conditions d'emploi de ces méthodes. Certaines de ces considérations reprennent et étendent les observations faites dans [55]. On trouvera également dans ce chapitre un compte rendu des travaux réalisés sur la question de l'influence de l'échantillonnage publiés dans [50, 51, 54]. Enfin, parmi les variantes et extensions de la décomposition modale empirique originale, une attention particulière est consacrée à celle permettant de traiter des signaux complexes présentée dans [57].

Le chapitre 2 est centré sur l'analyse du fonctionnement de la décomposition. Il se divise en deux grandes parties. La première part de l'étude de la décomposition d'une somme de deux sinusoïdes pour aboutir à un modèle qui est ensuite généralisé à des sommes de composantes périodiques faiblement non linéaires ainsi qu'aux sommes d'exponentielles complexes pour les algorithmes dédiés au cas

complexe. Cette partie reprend et étend les résultats proposés dans [53]. La seconde grande partie part de l'étude de la décomposition d'un bruit blanc gaussien et est consacrée plus généralement à l'étude de la décomposition de bruits large bande. On s'est intéressé notamment au cas de bruits présentant des caractéristiques d'invariance d'échelle ainsi qu'à l'estimation de l'exposant caractéristique de cette invariance à l'aide de la décomposition modale empirique. Les résultats rapportés dans cette partie reprennent en plus détaillé ceux proposés dans [21, 19, 20, 56]. On trouvera également dans cette partie une extension au cas de bruits à valeur complexe prolongeant les résultats rapportés dans [52].

Chapitre 1

Décompositions modales empiriques

1 Motivations et contexte

1.1 Contexte

Le contexte dans lequel s'inscrivent les méthodes de décompositions modales empiriques est très généralement celui de l'analyse de données. Plus précisément, ces méthodes ont été introduites pour faire face à une contradiction courante dans ce domaine, à savoir que la grande majorité des outils utilisés sont basés sur des hypothèses de stationnarité et de linéarité qui ne sont en réalité jamais strictement vérifiées. Elles le sont dans de nombreux cas de manière approximative, ce qui justifie l'emploi des outils, mais il existe aussi des situations où elles ne le sont pas, même approximativement, auquel cas les outils basés sur ces hypothèses sont alors inadaptés.

À ces difficultés s'ajoute aussi bien souvent le fait qu'il est impossible d'agir sur le système physique sous-jacent et qu'on ne peut reproduire le processus qui a généré les données. C'est en particulier le cas des données issues de systèmes comportant un très grand nombre de degrés de libertés comme les données climatiques ou économiques. Dans ces conditions, on ne dispose généralement que d'un seul enregistrement de données, ou une réalisation du processus aléatoire associé au processus physique comprenant la génération des données et leur mesure, et encore, une réalisation partielle dans la mesure où l'enregistrement est nécessairement de durée finie.

Dans ce contexte, les décompositions modales empiriques ont été introduites dans le but de proposer des méthodes souples d'emploi permettant de faciliter la lecture des données, c'est-à-dire l'extraction d'informations, généralement en vue d'une application donnée. Pour remplir au mieux l'objectif de souplesse, il faut idéalement relâcher les conditions de stationnarité et de linéarité et autant que possible limiter tout a priori sur le contenu des données. Cette dernière nécessité exclut de fait toutes les approches paramétriques qui consistent à décrire les données à partir des paramètres d'un modèle prédéterminé, calculés pour que la sortie du modèle ressemble au mieux aux données.

Pour préciser le cadre dans lequel les décompositions modales empiriques ont été introduites, on se propose maintenant de discuter les notions de stationnarité et de linéarité.

1.1.1 (Non-)stationnarité

1.1.1.1 Définitions La notion de stationnarité est une propriété relative aux processus aléatoires.

Définition 1. *Un processus aléatoire $\{x(t), t \in \mathbb{R}\}$ est dit strictement stationnaire si pour tout $\tau \in \mathbb{R}$, le processus translaté de τ $\{x_\tau(t), t \in \mathbb{R}\}$ défini par $\forall t \in \mathbb{R}, x_\tau(t) = x(t + \tau)$ est identique au processus $\{x(t), t \in \mathbb{R}\}$:*

$$\{x(t), t \in \mathbb{R}\} \stackrel{d}{=} \{x_\tau(t), t \in \mathbb{R}\}, \quad (1.1)$$

Pour une fonction déterministe $f(t)$, la propriété de stationnarité telle qu'énoncée ci-dessus implique que $f(t + \tau) = f(t)$ pour tout τ et donc que $f(t)$ est constante. En pratique cependant, on qualifie également de stationnaires les fonctions périodiques. Cette pratique se justifie en considérant l'astuce suivante : étant donnée une fonction $f(t)$ périodique de période T , on peut lui associer le processus aléatoire $\tilde{f}(t)$ défini par

$$\forall t \in \mathbb{R}, \tilde{f}(t) = f(t + \varphi), \quad (1.2)$$

où φ est une variable aléatoire uniformément distribuée dans $[0, T]$. $\tilde{f}(t)$ est alors strictement stationnaire.

L'hypothèse de stationnarité stricte étant une hypothèse forte, on utilise bien souvent la notion de stationnarité « au sens large ».

Définition 2. Un processus aléatoire $\{x(t), t \in \mathbb{R}\}$ est dit *faiblement stationnaire* (ou « au sens large », ou « à l'ordre 2 ») ssi $\forall t \in \mathbb{R}$,

- (i) $\mathbb{E}x(t) = m_x = \text{constante}$,
- (ii) $\mathbb{E}\{(x(t) - m_x)(x(s) - m_x)^*\} = r_x(t, s) = \gamma_x(t - s)$.

En pratique, la notion de stationnarité est pratiquement toujours considérée au sens large et c'est donc cette dernière définition qu'il faut avoir à l'esprit quand on parle d'outils consacrés au traitement des signaux stationnaires.

1.1.1.2 Lien avec la transformée de Fourier La transformée de Fourier est spécialement bien adaptée à la description des signaux stationnaires à l'ordre 2. En effet, on peut montrer que la stationnarité à l'ordre 2 garantit une représentation de Cramér de tout signal $x(t)$ [18]

$$x(t) = \int_{-\infty}^{\infty} e^{2i\pi ft} dX(f), \quad (1.3)$$

où l'intégrale est de type Fourier-Stieltjes et où l'égalité est à prendre au sens de la moyenne quadratique. Le grand intérêt de cette décomposition est que les incréments spectraux, qui mesurent les poids aléatoires $dX(f)$ associés aux fonctions de décomposition $t \mapsto e^{2i\pi ft}$ sont orthogonaux entre eux

$$\forall (f, f') \in \mathbb{R}^2, f \neq f' \implies \mathbb{E}\{dX(f)dX^*(f')\} = 0. \quad (1.4)$$

En d'autres termes les signaux stationnaires admettent une décomposition fréquentielle en variables décorréliées. Ce résultat peut se voir comme une conséquence de l'adéquation entre la localisation fréquentielle idéale induite par la décomposition de Fourier et la permanence au cours du temps du contenu spectral d'un signal stationnaire.

1.1.1.3 Non-stationnarité L'hypothèse de stationnarité, même faible, n'est pas vérifiée dans un grand nombre de situations pratiques [45, 28]. Il y a à cela plusieurs origines. La première est qu'un signal¹ expérimental a toujours une durée finie et n'est donc au mieux que la troncature d'un processus stationnaire. En pratique, on peut s'accomoder de cet aspect si l'échelle d'évolution du signal est courte devant sa durée, ce qui se traduit pour un processus stationnaire par le fait que le support de sa fonction d'autocorrélation $\gamma_x(\tau) = \mathbb{E}\{x(t)x^*(t + \tau)\} - \mathbb{E}\{x(t)\}\mathbb{E}\{x^*(t + \tau)\}$ est de longueur petite devant la durée du signal. En revanche, cet aspect peut représenter une cause de non-stationnarité si cette condition n'est pas vérifiée, bien que le processus physique sous-jacent puisse être effectivement considéré comme stationnaire sur une plus longue durée.

1. La notion de signal correspond ici essentiellement à la notion de processus aléatoire mais on l'utilisera aussi pour désigner des fonctions, éventuellement rendues aléatoires par l'ajout d'une phase aléatoire.

Au-delà des problèmes dus à l'acquisition nécessairement limitée dans le temps des signaux, de nombreux signaux expérimentaux sont aussi intrinsèquement non stationnaires. On peut citer parmi d'autres les signaux de radar et de sonar, les ondes sismiques, les signaux hydrodynamiques dans le régime turbulent, les signaux biomédicaux comme l'électrocardiogramme, les signaux de parole, les signaux musicaux ou encore les bruits de moteurs. Dans tous ces cas, les outils stationnaires sont mal adaptés et donnent lieu à des descriptions bien souvent plus difficiles à lire que les signaux eux-mêmes.

Pour faire face aux situations non stationnaires, un certain nombre d'outils spécifiques ont été développés. On fera une présentation générale des méthodes non paramétriques en 1.2.

1.1.2 (Non-)linéarité

1.1.2.1 Définition La notion de signal non linéaire se rattache à celle de système non linéaire, le premier étant un signal généré par le second [61, 17]. La définition d'un système non linéaire se rattache quant à elle au modèle mathématique décrivant le système : un système est linéaire si et seulement si sa dynamique est régie par une (des) équation(s) linéaire(s). Suivant le cas, ces équations peuvent être de plusieurs types : des équations différentielles, des équations aux différences, des équations fonctionnelles, etc.

1.1.2.2 Lien avec la transformée de Fourier La transformée de Fourier et plus généralement celle de Laplace sont particulièrement bien adaptées à l'analyse des systèmes linéaires. Ce résultat est très généralement lié au fait que les fonctions exponentielles $t \mapsto e^{pt}$, $p \in \mathbb{C}$, sont les vecteurs propres des opérateurs linéaires invariants dans le temps. Si on considère un système linéaire, la transformée de Fourier (ou de Laplace) constitue alors une décomposition en vecteurs propres des opérateurs linéaires définissant le système et est donc naturellement bien adaptée pour décrire les signaux qui en sont issus.

1.1.2.3 Non-linéarité De manière générale, l'hypothèse de linéarité d'un système physique est une approximation valable uniquement lorsque les grandeurs physiques en jeu ne sont pas trop grandes. On peut considérer à ce propos l'exemple canonique du pendule simple. L'équation différentielle régissant la dynamique du système est de la forme

$$\frac{d^2\theta}{dt^2} + \omega_0^2 \sin \theta = \text{constante}. \quad (1.5)$$

Cette équation est non linéaire sous cette forme, mais lorsque l'amplitude des oscillations est faible $\theta \ll 1$, on sait qu'on peut la remplacer par l'approximation

$$\frac{d^2\theta}{dt^2} + \omega_0^2 \theta = \text{constante}. \quad (1.6)$$

Plus précisément, cette approximation est valable quand la durée et la précision de l'acquisition du signal $\theta(t)$ ne sont pas suffisantes pour détecter d'une part le caractère non sinusoïdal des oscillations et d'autre part l'écart entre la période de l'oscillateur harmonique $2\pi/\omega_0$ et la période vraie de l'oscillateur non-linéaire (1.5). Plus généralement, cet exemple permet de voir que bon nombre de signaux issus de systèmes non linéaires peuvent être considérés comme linéaires si la précision de leur mesure n'est pas suffisante pour détecter les effets non linéaires.

En revanche, lorsqu'un système est étudié suffisamment finement pour percevoir les effets non linéaires, il devient nécessaire pour le décrire d'utiliser des outils adaptés aux situations non linéaires, ou plutôt qui ne reposent pas sur une hypothèse de linéarité. Si on reprend l'exemple du pendule simple non linéaire (1.5), les fonctions propres de l'opérateur différentiel non linéaire $d^2/dt^2 + \omega_0^2 \sin(\cdot)$ ne sont plus les fonctions exponentielles et il en résulte que l'oscillation propre $\theta(t)$ est non sinusoïdale.

On sait alors que la transformée de Fourier de $\theta(t)$ est constituée de plusieurs harmoniques qui n'ont pas de sens en elles-mêmes, ce qui permet de voir que la transformée de Fourier n'est pas adaptée à l'analyse d'une oscillation non linéaire.

L'analyse des signaux non linéaires commence généralement par une phase exploratoire non paramétrique. Cette phase consiste de manière générale à représenter des caractéristiques du signal dans un espace de plus grande dimension pour déceler à l'œil des structures particulières. La première étape est en général une simple représentation temporelle du signal. Si celui-ci est stationnaire, elle peut être suivie d'une analyse en corrélation, en densité spectrale de puissance, et plus généralement tous les outils de l'analyse des signaux stationnaires/linéaires peuvent être utilisés. Au-delà de ces méthodes, existent aussi un certain nombre d'outils spécialement adaptés à l'analyse de signaux non linéaires [61]. L'exemple canonique de tels outils est sans doute le diagramme de phase pour les systèmes à temps continu qui est la représentation de la trajectoire du signal dans le plan $x(t)$, dx/dt . Pour les systèmes à temps discret, ce dernier se décline en diagramme de dispersion (scatter plot). À ces deux méthodes s'ajoutent un certain nombre d'outils parmi lesquels, l'analyse du bispectre [7], les histogrammes bivariés (histogramme 2D du couple $(x(t), x(t+1))$), les diagrammes de récurrence (recurrence plot) [7] ou encore la régression retardée (lagged regression) [61]. De plus, tous ces outils peuvent être appliqués sur le signal directement ou sur un signal « délinéarisé », différence entre le signal et une modélisation linéaire de ce dernier, par exemple à l'aide d'un modèle autorégressif (linéaire). En pratique, on constate que la quasi-totalité de ces méthodes requièrent de fait une hypothèse de stationnarité ou d'ergodicité pour être efficaces.

En dehors de la phase exploratoire, les méthodes d'analyse non linéaires consistent généralement à adapter un modèle générique aux données. Il existe pour ce faire une grande variété de modèles non linéaires aux propriétés diverses, dont on pourra trouver une liste détaillée dans [23]. En dehors de ces approches paramétriques, existent aussi des approches non paramétriques d'« apprentissage » (machine learning), utilisant des outils tels que les réseaux de neurones ou plus récemment les machines à vecteurs de support [58]. Ces approches sont particulièrement utiles pour prédire l'évolution du signal en dehors de l'intervalle où il est connu mais, en revanche, elles ne facilitent généralement pas sa description dans la mesure où la représentation de ce dernier déterminée par ces méthodes est souvent dans un espace de grande dimension (voire infinie) qui n'a pas de sens en soi.

1.2 Méthodes dévolues à l'analyse des signaux non stationnaires

Une première solution pour décrire les signaux non stationnaires consiste à étendre l'outil stationnaire qu'est la transformée de Fourier au cadre non stationnaire. On aboutit alors aux approches temps-fréquence dont on pourra trouver des présentations modernes dans [18, 10, 6, 5, 24].

1.2.1 Représentations temps-fréquence / temps-échelle

Les approches temps-fréquence peuvent être classées en deux catégories suivant qu'elles sont linéaires par rapport au signal ou bilinéaires.

1.2.1.1 Représentations linéaires Une manière simple d'adapter la transformée de Fourier au cadre non stationnaire consiste à l'appliquer de manière locale à l'aide d'une fenêtre de pondération bien localisée en temps. On aboutit alors à la transformée de Fourier à court terme :

$$F(t, f) = \int_{-\infty}^{\infty} x(\tau) w(\tau - t) e^{-2i\pi f\tau} d\tau, \quad (1.7)$$

où $w(t)$ est non nulle uniquement sur un voisinage de 0. Cette relation est de plus inversible si $w(t)$ est d'énergie unité, ce qui permet de représenter le signal comme une combinaison linéaire d'atomes temps-fréquence de la forme $w(t - \tau) e^{2i\pi f\tau}$. En pratique, on peut voir qu'une telle représentation

impose le choix d'une échelle caractéristique donnée par le support temporel de la fenêtre $w(t)$. Ce choix a un certain nombre de conséquences. Sur la localisation des événements dans le plan temps-fréquence tout d'abord : ces derniers pourront d'autant mieux être localisés en temps que le support de $w(t)$ sera petit ; à l'inverse ils pourront d'autant mieux être localisés en fréquence que le support sera grand. L'autre conséquence importante est que la transformée de Fourier à court terme ne permet pas vraiment d'analyser les évolutions à des échelles plus grandes que le support de la fenêtre. L'information donnée par cette dernière aux grandes échelles se résume de fait à l'évolution temporelle du signal lissée par la fenêtre $w(t)$.

Face à ces limites, une autre solution populaire consiste d'une certaine manière à faire varier la taille de la fenêtre en fonction de la fréquence. De fait, la notion de fréquence est alors remplacée par la notion d'échelle et on arrive au cadre de la transformée en ondelettes continue qui s'exprime sous la forme

$$W(t, a) = \int_{-\infty}^{\infty} x(\tau) \frac{1}{\sqrt{a}} \psi^* \left(\frac{\tau - t}{a} \right) d\tau. \quad (1.8)$$

où $\psi(t)$ est l'ondelette « mère ». Il faut noter que cette seule définition ne définit pas vraiment une transformée en ondelettes continue : il faut y ajouter au minimum la condition d'admissibilité

$$\int_{-\infty}^{\infty} |\Psi(f)|^2 \frac{df}{|f|} = 1, \quad (1.9)$$

où $\Psi(f)$ est la transformée de Fourier de $\psi(t)$. Cette relation est en fait la condition nécessaire et suffisante à l'inversion de la transformée en ondelettes et donc à la décomposition du signal $x(t)$ sur la famille d'ondelettes $\psi((t - \tau)/a)$, de paramètres τ et a . Une conséquence de cette condition d'admissibilité est que $\psi(t)$ doit être de moyenne nulle et qu'ainsi elle présente au moins quelques oscillations, d'où le nom d'ondelette. De manière générale, la décomposition en échelles ainsi définie ne peut se ramener à une interprétation en fréquence. Cependant, pour l'analyse exploratoire des données, il est fréquent d'utiliser des ondelettes raisonnablement localisées en fréquence autour d'une fréquence f_0 , ce qui permet d'interpréter le coefficient $W(t, a)$ comme une contribution à la position $(t, f_0/a)$ du plan temps-fréquence. La transformée en ondelettes s'est très largement développée depuis son introduction et a donné naissance à une grande variété d'outils d'analyse des signaux non stationnaires dont on pourra trouver un éventail dans [42].

De manière générale, ces représentations linéaires sont à valeur complexe comme la transformée de Fourier. En pratique, tout comme pour la transformée de Fourier, on ne s'intéresse bien souvent qu'au module carré de ces transformations avec une interprétation en terme de distribution d'énergie dans le plan temps-fréquence ou temps-échelle. On parle alors de « spectrogramme » pour la transformée de Fourier à court terme et de « scalogramme » pour la transformée en ondelettes.

De manière plus générale, il n'est pas nécessaire de partir d'une décomposition linéaire pour construire une distribution d'énergie temps-fréquence ou temps-échelle mais il est au contraire avantageux de considérer directement la question à travers la construction de représentations bilinéaires.

1.2.1.2 Représentations bilinéaires Les représentations temps-fréquence et temps-échelle bilinéaires peuvent être vues comme toutes construites autour d'un élément central qu'est la distribution de Wigner-Ville

$$WV(t, f) = \int_{-\infty}^{\infty} x \left(t + \frac{\tau}{2} \right) x^* \left(t - \frac{\tau}{2} \right) e^{-2i\pi f\tau} d\tau. \quad (1.10)$$

Cette distribution présente par rapport au spectrogramme un certain nombre de bonnes propriétés dont une meilleure localisation dans le plan temps-fréquence et une localisation parfaite des modulations de fréquences linéaires. En revanche, ces bonnes propriétés viennent avec l'inconvénient qu'est la présence de termes d'interférences oscillants qui compliquent l'interprétation de la décomposition. De plus, la distribution de Wigner-Ville est contrairement au spectrogramme difficile à interpréter

en terme de distribution d'énergie parce qu'elle n'est pas nécessairement positive. Plus généralement, on pourra trouver une étude détaillée des propriétés de cette distribution dans [18, 10].

Pour limiter ou supprimer les difficultés liées à l'analyse de la représentation de Wigner-Ville, une solution consiste à lisser en quelque sorte la distribution par un filtre linéaire dans le plan temps-fréquence. On aboutit alors à ce qui est appelé la classe de Cohen dont les membres prennent la forme générale

$$C(t, f) = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} \Pi(t - \tau, f - \nu) WV(\tau, \nu) d\tau d\nu, \quad (1.11)$$

et qui regroupe en fait toutes les méthodes temps-fréquence bilinéaires covariantes par rapport aux translations en temps et en fréquence [18], ce qui inclut en particulier le spectrogramme. La liberté supplémentaire par rapport à la représentation de Wigner-Ville contenue dans le noyau de lissage permet d'obtenir un certain nombre de propriétés supplémentaires par rapport à cette dernière comme la positivité, la causalité et surtout une réduction des termes d'interférences. Cependant, on peut noter que toutes ces améliorations éventuelles viennent au prix d'une détérioration de la localisation temps-fréquence de la représentation de Wigner-Ville, ce qui est évident si on considère la forme (1.11) comme un filtrage linéaire passe-bas dans le plan temps-fréquence. On pourra trouver dans [18] une approche générale pour construire les noyaux de lissage à partir des propriétés souhaitées.

En dehors de la classe de Cohen, il existe d'autres classes de représentations bilinéaires dont une en particulier généralise la notion de scalogramme : la classe affine. Comme la classe de Cohen, elle peut être définie comme une forme de lissage de la distribution de Wigner-Ville, cette fois en temps-échelle

$$A(t, a) = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} \Pi\left(\frac{t - \tau}{a}, a\nu\right) WV(\tau, \nu) d\tau d\nu. \quad (1.12)$$

La richesse des représentations appartenant à cette classe est similaire à celle de la classe de Cohen et on pourra trouver également dans [18] une approche générale pour les construire en fonction des propriétés souhaitées.

1.2.1.3 Réallocation Pour améliorer la lisibilité des représentations temps-fréquence et temps-échelle, une technique particulièrement intéressante, proposée dans [3], est la technique de réallocation. Il s'agit d'un post-traitement qui peut être appliqué de manière générale sur toutes les distributions temps-fréquence et temps-échelle évoquées précédemment et qui permet d'améliorer la localisation de ces distributions pour en faciliter la lisibilité. La présentation proposée ici est inspirée de [24] où un chapitre est consacré à la méthode de réallocation, à sa mise en œuvre et à ses applications.

Si on prend l'exemple d'une distribution de la classe de Cohen

$$C(t, f) = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} \Pi(t - \tau, f - \nu) WV(\tau, \nu) d\tau d\nu, \quad (1.13)$$

le principe de la réallocation part du constat que le coefficient $C(t, f)$ peut être vu comme une moyenne pondérée de la distribution de Wigner-Ville sur une surface centrée en (t, f) correspondant au domaine où $\Pi(t - \tau, f - \nu)$ est non nul. Par analogie avec la mécanique, on peut voir que le fait d'affecter cette contribution en (t, f) revient à affecter la masse d'un solide à son *centre géométrique*. En mécanique, il est bien connu qu'il est plus naturel d'affecter la masse au *centre de gravité* et c'est justement le point de vue adopté par la réallocation. À chaque point du plan temps-fréquence, deux quantités complémentaires $\hat{t}$ et $\hat{f}$ sont calculées. Elles correspondent aux coordonnées du centre de gravité des valeurs de la distribution de Wigner-Ville, pondérées par la fenêtre temps-fréquence $\Pi(t - \tau, f - \nu)$, et c'est au point $(\hat{t}, \hat{f})$ que la valeur de la représentation temps-fréquence $C(t, f)$ calculée en (t, f) est ré-affectée. La distribution temps-fréquence réallouée se définit alors comme

$$\tilde{C}(t, f) = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} C(\tau, \nu) \delta(t - \hat{t}(\tau, \nu)) \delta(f - \hat{f}(\tau, \nu)) d\tau d\nu, \quad (1.14)$$

où $\hat{t}(t, f)$ et $\hat{f}(t, f)$ sont définis par

$$\hat{t}(t, f) = \frac{1}{C(t, f)} \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} \tau \Pi(t - \tau, f - \nu) WV(\tau, \nu) d\tau d\nu, \quad (1.15)$$

$$\hat{f}(t, f) = \frac{1}{C(t, f)} \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} \nu \Pi(t - \tau, f - \nu) WV(\tau, \nu) d\tau d\nu. \quad (1.16)$$

Il est bon de préciser que la méthode de réallocation permet certes une meilleure localisation dans le plan temps-fréquence ou temps-échelle mais elle ne permet pas une meilleure résolution. En effet, si la représentation temps-fréquence ou temps-échelle sous-jacente ne permet pas de distinguer deux contributions proches en temps ou en fréquence/échelle, la méthode de réallocation ne permettra pas non plus de les distinguer mais elle ré-affectera l'énergie correspondant aux deux contributions quelque part entre leurs deux localisations idéales.

1.2.2 Fréquence instantanée

Parallèlement aux approches temps-fréquence et temps-échelle présentées précédemment, une autre approche particulièrement séduisante est celle de la fréquence instantanée qui représente une tentative d'outrepasser la limitation intrinsèque à l'analyse conjointe en temps et en fréquence représentée par le principe d'incertitude d'Heisenberg-Gabor. Ce principe d'incertitude qui est l'expression de l'impossibilité d'une localisation parfaite en temps et en fréquence [18] peut s'énoncer de la manière suivante. Étant donnée une fonction $x(t)$, on définit $\Delta_t^2 = \left(\int t^2 |x(t)|^2 dt - \left(\int t |x(t)|^2 dt \right)^2 \right) / E_x$ (où $E_x = \int |x(t)|^2 dt$ est l'énergie de $x(t)$), qui mesure le support en temps de $x(t)$, et $\Delta_f^2 = \left(\int f^2 |\hat{x}(f)|^2 df - \left(\int f |\hat{x}(f)|^2 df \right)^2 \right) / E_x$ (où $\hat{x}(f)$ est la transformée de Fourier de $x(t)$), qui mesure son support en fréquence. Le principe d'incertitude d'Heisenberg-Gabor s'écrit alors

$$\Delta_t \Delta_f \geq \frac{1}{4\pi}. \quad (1.17)$$

Compte tenu de cette limitation, la définition d'une fréquence instantanée paraît a priori délicate. On pourra consulter à ce sujet l'excellent chapitre consacré à la fréquence instantanée dans [24]. Le problème général prend la forme suivante. Étant donné un signal de type sinusoïde modulée en amplitude et en fréquence

$$x(t) = a(t) \cos \varphi(t) \quad (1.18)$$

on cherche à remonter aux paramètres $a(t)$ et $\varphi(t)$ qui sont son amplitude et sa phase instantanée, dont la fréquence instantanée peut être définie comme la dérivée à un facteur 2π près

$$f(t) = \frac{1}{2\pi} \frac{d\varphi}{dt}. \quad (1.19)$$

On peut d'ores et déjà voir que le problème est particulièrement mal posé dans la mesure où, s'il est clair que la donnée de $a(t)$ et $\varphi(t)$ définit sans ambiguïté $x(t)$, il existe une infinité de couples $(a(t), \varphi(t))$ différents conduisant au même signal $x(t)$. De plus, toutes ces paramétrisations ne sont pas équivalentes du point de vue de l'interprétation et seules certaines peuvent effectivement être interprétées comme une amplitude et une phase instantanée. Par exemple, on peut tout à fait choisir comme paramétrisation $a(t) = x(t)$ et $\varphi(t) = 0$. Si l'on veut donc déterminer une paramétrisation pertinente du point de vue de l'interprétation, il faut introduire un certain nombre de conditions, la plus évidente étant que $a(t)$ doit être une fonction positive. L'autre condition nécessaire à une interprétation en amplitude et fréquence instantanée est plus délicate à formaliser et correspond à l'idée que $a(t)$ et $(d\varphi/dt)(t)$ doivent évoluer lentement par rapport $\cos \varphi(t)$.

1.2.2.1 Méthode du signal analytique En pratique, la méthode standard pour définir les paramètres $a(t)$ et $\varphi(t)$ consiste à construire à partir du signal réel $x(t) = a(t) \cos \varphi(t)$ un signal complexe $y(t) = a(t)e^{i\varphi(t)}$ où les paramètres $a(t)$ et $\varphi(t)$ peuvent alors être déterminés de manière univoque comme le module et l'argument de $y(t)$. Une solution à ce problème consiste à utiliser les propriétés de la transformée de Hilbert définie par

$$(\mathcal{H}x)(t) = \frac{1}{\pi} VP \int_{-\infty}^{\infty} \frac{x(t')}{t - t'} dt', \quad (1.20)$$

où VP indique que l'intégrale impropre est calculée selon la valeur principale au sens de Cauchy. De là, on définit le *signal analytique* $y(t)$ associé à $x(t)$ par

$$y(t) = x(t) + i(\mathcal{H}x)(t). \quad (1.21)$$

$y(t)$ est alors un signal complexe analytique, c'est-à-dire que sa transformée de Fourier est nulle pour les fréquences négatives. La transformée de Hilbert a en particulier la propriété de transformer un cosinus en sinus ($\mathcal{H}\{\cos t\} = \sin t$), ce qui fait que le signal analytique associé au cosinus $\cos t$ est l'exponentielle complexe e^{it} . De manière plus générale, cette propriété reste approximativement valide, sous certaines conditions, si le cosinus est modulé en amplitude et en fréquence

$$\mathcal{H}\{a(t) \cos \varphi(t)\} \simeq a(t) \sin \varphi(t). \quad (1.22)$$

Les conditions pour que cette égalité soit stricte sont particulièrement délicates. Une première condition est parfois présentée sous le nom de théorème de Bedrosian [4] et s'énonce sous la forme suivante. Si la transformée de Fourier de $a(t)$, $\hat{a}(f)$, est nulle pour $|f| > B$ et que celle de $\cos \varphi(t)$ est nulle pour $|f| < B$, alors

$$\mathcal{H}\{a(t) \cos \varphi(t)\} = a(t) \mathcal{H}\{\cos \varphi(t)\}. \quad (1.23)$$

Cette condition est une manière de formaliser la condition précédemment énoncée de manière floue selon laquelle $a(t)$ doit évoluer lentement par rapport à $\cos \varphi(t)$. Au-delà de cette condition, l'égalité dans (1.22) requiert de plus que

$$\mathcal{H}\{\cos \varphi(t)\} = \sin \varphi(t), \quad (1.24)$$

qui exige des propriétés particulièrement complexes sur $\varphi(t)$ [24]. En pratique, ces considérations sont d'un intérêt assez limité puisque la décomposition de $x(t)$ en $x(t) = a(t) \cos \varphi(t)$ n'est généralement pas connue a priori et qu'on cherche justement à la calculer. Une condition plus intéressante est en revanche celle sous laquelle la paramétrisation $(a(t), \varphi(t))$ déterminée par le module et l'argument du signal analytique peut effectivement donner lieu à une interprétation pertinente en termes d'amplitude et de phase instantanées. Celle-ci est généralement donnée sous la forme d'une condition de « bande étroite », c'est-à-dire que la transformée de Fourier du signal $\hat{x}(f)$ doit être non nulle uniquement pour $|f|$ appartenant à un intervalle $[f_0 - \Delta_f, f_0 + \Delta_f]$ dont la largeur est étroite par rapport à sa fréquence centrale $\Delta_f \ll f_0$. Cette formulation est malheureusement très restrictive dans la mesure où elle impose à la fréquence instantanée de rester dans les environs de f_0 alors que la méthode du signal analytique permet tout-à-fait de définir une fréquence et une amplitude instantanées pertinentes même si la fréquence varie beaucoup sur la durée du signal tant qu'elle ne varie pas trop vite par rapport à la période locale. De manière générale, ces conditions sont une émanation du principe d'incertitude d'Heisenberg-Gabor qui ne permet une interprétation en fréquence que si celle-ci « dure » suffisamment longtemps, ce qui est analogue à l'idée que la fréquence ne doit pas trop varier à l'échelle de la période locale.

Dans le contexte de l'analyse de signaux non linéaires, l'approche de la fréquence instantanée est en quelque sorte moins inadaptée que les représentations temps-fréquence et temps-échelle évoquées précédemment. En effet, si celles-ci ont tendance naturellement à décomposer une onde non linéaire

périodique en ses composantes de Fourier, l'approche de la fréquence instantanée lui associera au contraire toujours un unique couple amplitude/fréquence instantanées. Ce couple ne peut alors pas vraiment être interprété en termes d'amplitude et de fréquence au sens de Fourier mais l'évolution de ces grandeurs sur une période constitue une représentation de l'onde non linéaire qui n'est pas forcément dénuée d'intérêt [28]. En revanche, l'interprétation peut devenir délicate si, par exemple, la fréquence instantanée prend des valeurs négatives au cours de la période. Pour cette raison, cette description n'est a priori pertinente que pour des ondes faiblement non linéaires.

Enfin la méthode de la fréquence instantanée présente une limitation forte dans la mesure où elle n'est pertinente que quand le signal contient une seule composante analogue à une sinusoïde modulée en amplitude et en fréquence. En effet, si le signal contient deux composantes de cette forme, la définition du couple amplitude/fréquence instantanée à partir du signal analytique ne fournit qu'un seul couple amplitude/fréquence instantanées pour les deux composantes et l'interprétation de ce couple est alors loin des fréquences et amplitudes instantanées des deux composantes. On verra au paragraphe suivant une approche permettant de remonter à des couples amplitude/fréquence instantanées multiples à partir d'une transformée en ondelettes continue ou d'une transformée de Fourier à court terme. Les décompositions modales empiriques qui constituent le sujet central de cette thèse ont également été introduites dans cette optique [28].

1.2.2.2 Plusieurs fréquences instantanées : extraction de lignes de crête La méthode d'extraction de lignes de crêtes (ridges) est généralement utilisée dans le cadre de la transformée en ondelettes continue mais elle peut également s'appliquer sur des représentations temps-fréquence à interférences réduites comme le spectrogramme. On pourra en trouver une présentation dans [42] et de manière plus détaillée dans [6]. On ne présentera ici que le principe de base dans le cas du spectrogramme, les détails des méthodes évoluées se trouvant dans les deux ouvrages précités.

Étant donné un signal composé de plusieurs composantes sinusoïdales modulées en amplitude et en fréquence

$$x(t) = \sum_{k=1}^K a_k(t) \cos \varphi_k(t), \quad (1.25)$$

le problème consiste à remonter aux amplitudes et fréquences instantanées de chacune des composantes. L'idée de base des méthodes d'extraction de lignes de crêtes est alors que si à un instant donné t_0 , les fréquences instantanées des différentes composantes sont suffisamment éloignées, alors la section du spectrogramme en t_0 , $|F(t_0, f)|^2$, doit en fonction de f présenter K lobes centrés chacun sur la fréquence instantanée d'une des composantes. De plus, sous certaines conditions relatives à la fenêtre utilisée pour la transformée de Fourier à court terme et dans la mesure où les fréquences et amplitudes instantanées des composantes n'évoluent pas trop vite, ces lobes ne présentent qu'un seul maximum local dont la position permet de remonter à la fréquence instantanée de la composante correspondant au lobe. L'amplitude de $|F(t_0, f)|$ au niveau des maxima de chacun des lobes permet quant à elle de remonter aux amplitudes instantanées. Partant de là, on voit qu'on peut remonter aux couples amplitude/fréquence instantanées de chacune des composantes en analysant pour chaque valeur de t où le spectrogramme est calculé, les extrema de $|F(t, f)|$ en tant que fonction de f : les positions donnent les fréquences instantanées, et les valeurs de $|F(t, f)|$ à ces mêmes positions donnent les amplitudes instantanées. Un post-traitement simple permet alors de relier les différentes valeurs d'amplitude/fréquence instantanées obtenues à différents instants pour reconstituer les évolutions temporelles des amplitudes/fréquences instantanées de chacune des composantes constituant $x(t)$.

En pratique, les méthodes utilisées pour extraire les lignes de crêtes sont très différentes de l'approche simpliste proposée ici mais le principe reste le même. Ces méthodes sont en fait bien plus performantes et prennent notamment en compte la présence de bruit dans les données pour fournir une estimation robuste des amplitudes et fréquences instantanées. De plus, elles permettent

également d'estimer la phase instantanée des composantes $\varphi_k(t)$ en prenant en compte la phase du spectrogramme $F(t, f)$.

Les méthodes d'extraction de lignes de crêtes ont aussi un certain nombre de limites. Tout d'abord, elles héritent naturellement des limites de la représentation temps-fréquence ou temps-échelle sur laquelle elles sont basées. En particulier, elles ne peuvent pas en améliorer ni la résolution fréquentielle ni la résolution temporelle. De plus, les fréquences et amplitudes instantanées estimées sont en quelque sorte lissées par le noyau caractérisant la représentation. Enfin, si la représentation temps-fréquence ou temps-échelle décompose une onde non linéaire en ses composantes harmoniques, l'extraction de lignes de crêtes verra une ligne de crête par harmonique. Par conséquent, la représentation des ondes faiblement non linéaires en termes de fréquence et amplitude instantanées suggérée précédemment n'est pas possible par cette approche. Par rapport à toutes ces limites, le point de vue proposé par les décompositions modales empiriques constitue une alternative a priori attrayante.

2 Présentation des décompositions modales empiriques

Par rapport à des méthodes proposant une représentation globale du signal (approches temps-fréquence/temps-échelle) ou des méthodes permettant d'analyser finement une composante oscillante (fréquence instantanée), se pose aussi le problème étant donné un signal constitué de plusieurs composantes de séparer les différentes contributions, éventuellement dans le but de les analyser finement par la suite. Les méthodes de décompositions modales empiriques offrent pour ce problème une approche automatique originale présentant un certain nombre de propriétés remarquables.

2.1 La décomposition modale empirique originale

Cette section est consacrée à la présentation de la méthode de décomposition modale empirique sous sa forme originale telle qu'introduite dans [28]. La présentation qui est faite ici est assez différente de la présentation initiale et vise à mettre en valeur les concepts clés de la méthode.

2.1.1 Principe récursif

Au plus haut niveau hiérarchique, la décomposition modale empirique (EMD pour « Empirical Mode Decomposition ») peut être vue comme l'application récursive d'une opération de décomposition élémentaire $\mathcal{D}$ permettant d'extraire de tout signal *oscillant* x sa composante qui oscille localement le plus *rapidement* $\mathcal{D}x$. La différence $x - \mathcal{D}x$ est alors la partie de x qui oscille plus *lentement*, dont on peut à nouveau extraire la partie oscillant localement le plus rapidement. On obtient ainsi le principe général de l'EMD qu'on peut formuler par

$$EMD(x) = \{\mathcal{D}x\} \cup EMD(x - \mathcal{D}x), \quad (1.26)$$

où la décomposition s'arrête quand $x - \mathcal{D}x$ n'est plus un signal *oscillant*, auquel cas son EMD est par convention simplement lui-même.

À terme, le principe précédent aboutit à une décomposition de la forme

$$x(t) = a_K[x](t) + \sum_{k=1}^K d_k[x](t), \quad (1.27)$$

où on utilise les notations

$$d_{k+1}[x](t) = \mathcal{D}a_k[x](t), \quad (1.28)$$

$$a_{k+1}[x](t) = a_k[x](t) - \mathcal{D}a_k[x](t), \quad (1.29)$$

avec la convention $a_0[x](t) = x(t)$. La composante $a_K[x](t)$ étant ce qui reste du signal $x(t)$ après avoir extrait toutes les composantes oscillantes, elle est souvent appelée « résidu ». En pratique, la décomposition se limite toujours à un nombre fini de composantes $d_k[x](t)$ même si ce résultat n'est pas prouvé dans le cas général.

Pour préciser la présentation de l'EMD, il convient d'expliquer ce qu'on entend d'une part par « signal oscillant » et d'autre part par « composante oscillant le plus rapidement ». Le premier est simple et correspond bien à l'idée intuitive d'oscillation : on appelle signal oscillant tout signal qui est tantôt croissant, tantôt décroissant. Plus précisément, on pourra considérer qu'un signal est oscillant à partir du moment où il présente au moins deux extrema locaux, le cas d'un seul extremum étant plus éloigné de l'idée intuitive d'oscillation.

La notion de « composante oscillant le plus rapidement » est plus spécifique à l'EMD. Il s'agit tout d'abord d'un signal oscillant *autour de zéro*, ce qu'on peut définir par les deux caractéristiques

- tous les minima locaux sont strictement négatifs, tous les maxima locaux sont strictement positifs,
- le signal est raisonnablement symétrique par rapport à la ligne zéro.

La deuxième caractéristique est volontairement présentée de manière floue à ce stade. La définition plus précise étant en lien direct avec la manière de calculer $\mathcal{D}x$, elle sera présentée en même temps que cette dernière. De plus, il existe un certain nombre de variantes d'EMD qui utilisent en pratique des définitions légèrement différentes mais qui correspondent toutes au cadre de la définition floue proposée ici. De manière générale, la définition proposée ici de signal oscillant autour de zéro correspond à la définition de ce que les créateurs de l'EMD ont appelé « Intrinsic Mode Function » (IMF), qu'on peut traduire par « fonction modale intrinsèque », l'idée étant que les composantes $d_k[x](t)$ sont censées représenter des modes d'oscillations intrinsèques au signal $x(t)$. En pratique, la définition d'IMF englobe une grande variété de signaux : toute forme d'onde périodique avec un maximum positif et un minimum négatif par période peut être un IMF si les valeurs absolues du signal à ses minima et à ses maxima sont égales, et toute telle forme d'onde modulée en amplitude et en fréquence est aussi un IMF.

En plus d'être un signal oscillant autour de zéro, la « composante oscillant le plus rapidement » $\mathcal{D}x$ doit aussi, comme son nom l'indique, osciller plus rapidement que le reste du signal $x - \mathcal{D}x$. Là encore, il n'existe pas de définition précise, mais l'idée est en gros que la composante oscillant le plus rapidement doit avoir localement plus d'extrema que la composante oscillant plus lentement. Pour préciser cette notion, on pourrait proposer par exemple d'imposer que la composante oscillant plus lentement présente au maximum un extremum sur tout intervalle entre deux minima (ou maxima) locaux du signal. Cependant, une telle définition reste très floue et, de plus, l'EMD ne garantit pas une telle propriété, même si elle est souvent vérifiée.

Le cadre général de l'EMD qui vient d'être présenté n'est pas sans rappeler le cadre de la transformée en ondelettes discrète (dyadique) [42]. Cette dernière est en fait également réalisée en pratique de manière récursive : le signal est d'abord décomposé en « détail » et « approximation », correspondant respectivement à une partie haute fréquence et une partie basse fréquence, puis la même décomposition est appliquée à la partie « approximation ». La différence entre transformée en ondelettes discrète et EMD se trouve en fait dans la manière dont est calculée la décomposition élémentaire en « détail » plus « approximation » ou « oscillations rapides » plus « oscillations lentes ». Pour la transformée en ondelettes discrète, les deux composantes sont calculées par une opération prédéfinie de filtrage linéaire invariante dans le temps dont la réponse en fréquence est simplement contractée d'un facteur 2 lorsque la décomposition est ensuite appliquée à la partie « approximation ». Comme on le verra par la suite, les caractéristiques de la décomposition pour l'EMD sont à l'inverse déterminées par le signal et de manière locale. De plus, la transformée en ondelettes discrète est fondée sur une base mathématique bien établie alors que l'EMD n'est jusqu'ici définie que par la sortie d'un algorithme aux multiples degrés de liberté.

2.1.2 Extraction de la composante oscillant le plus rapidement : processus de tamisage

2.1.2.1 Les extrema locaux comme repères de la plus petite échelle du signal Si on considère dans un premier temps le cas simpliste d'un signal sinusoïdal, une manière efficace de mesurer ses caractéristiques, fréquence, amplitude et phase, consiste à repérer un certain nombre de points de contrôle. La fréquence peut par exemple être mesurée en repérant l'espacement entre les passages à zéro. Pour l'amplitude et la phase, ces derniers sont insuffisants mais on peut y accéder en considérant par exemple les extrema locaux, qui ne sont autres que les passages à zéro de la dérivée du signal. De manière plus générale, il est possible d'analyser le contenu d'un signal stationnaire composé d'une mixture de composantes sinusoïdales simplement en étudiant les caractéristiques de ses passages à zéro et des passages à zéro de ses dérivées successives [34]. L'idée de base est simple : lorsqu'on prend la dérivée du signal, les amplitudes relatives des composantes sont multipliées par leur fréquence. De là, lorsque l'ordre de la dérivée tend vers l'infini, les amplitudes des différentes composantes deviennent négligeables devant celle de la composante de plus haute fréquence, ce qui permet d'analyser en détail cette dernière. Par la suite, celle-ci est soustraite au signal pour analyser les composantes suivantes. En pratique cependant, cette technique est difficilement applicable, simplement parce que toute mesure est intrinsèquement bruitée et que les opérations de dérivées successives amplifient le bruit à haute fréquence. Par conséquent, l'analyse par les passages à zéro des dérivées successives est généralement limitée aux passages à zéro du signal et de sa dérivée, c'est-à-dire les extrema du signal.

Si on considère maintenant un signal non stationnaire composé d'une mixture de composantes sinusoïdales non stationnaires, c'est-à-dire modulées en amplitude et en fréquence, on peut légitimement supposer que ce qui est vrai pour un signal stationnaire devrait rester à peu près vrai pour un signal non stationnaire, surtout si l'évolution de la non-stationnarité est relativement lente par rapport aux périodes (locales) des composantes sinusoïdales non stationnaires.

C'est ainsi que les créateurs de l'EMD proposent de définir la « composante oscillant localement le plus rapidement » à l'aide des extrema locaux du signal [28]. Par ailleurs, ils considèrent également la possibilité d'utiliser les passages à zéro de dérivées d'ordre supérieur dans le cas où le signal ne présente pas d'extrema. Dans une contribution ultérieure [25], ils proposent même de remplacer les extrema locaux dans l'EMD par les « extrema de courbure » qui sont analogues aux passages à zéro de la dérivée troisième.

Dans le cas de l'EMD standard, la décomposition est basée sur les extrema locaux du signal. Plus précisément, si on considère par exemple l'évolution du signal entre deux minima locaux successifs t_0 et t_1 , l'idée est en gros de considérer que le signal est sur cet intervalle la somme d'une composante oscillante $d(t)$ et d'une tendance lente $a(t)$ à peu près constante sur l'intervalle $[t_0, t_1]$ et apparentée à la valeur moyenne de $x(t)$ sur cet intervalle. De fait, la composante $a(t)$ est souvent appelée « moyenne locale » du signal, parce que sur chaque intervalle entre deux minima elle est localement proche, ne serait-ce que philosophiquement, de la moyenne du signal sur cet intervalle. Si on suppose maintenant $a(t)$ et $d(t)$ définis proprement sur toute la durée du signal, alors $d(t)$ est a priori un bon candidat pour être la composante du signal oscillant localement le plus rapidement. En pratique, toute la difficulté se trouve en fait dans la manière de définir proprement les deux composantes $a(t)$ et $d(t)$.

2.1.2.2 Détermination de la moyenne locale à partir des extrema locaux Pour calculer la moyenne locale d'un signal, on pourrait penser intuitivement à une opération de la forme

$$a(t) = \frac{1}{\alpha(t)} \int x(u) w\left(\frac{u-t}{\alpha(t)}\right) du, \quad (1.30)$$

avec $w(u)$ une fonction de pondération (i.e. de masse unité) dont la masse serait concentrée autour de zéro et $\alpha(t)$ l'échelle locale à laquelle le signal est moyenné pour calculer $a(t)$. En pratique, proposer un choix de fonction de pondération n'est pas nécessairement difficile, mais le choix de l'échelle locale $\alpha(t)$ est un problème bien plus délicat pour lequel il n'existe pas de solution simple dans le cas général.

En outre, à supposer qu'on dispose d'une définition adéquate pour $\alpha(t)$, la définition (1.30) resterait insuffisante. Dans le cas d'une modulation de fréquence linéaire, par exemple, chaque demi-période est plus courte (ou plus longue) que la précédente. Ceci fait que si on se place par exemple au niveau d'un passage à zéro en t_0 , entouré à gauche par un minimum et à droite par un maximum, la demi-période négative à gauche de t_0 est plus longue que la demi-période positive à droite. De ce fait, la moyenne locale a toutes les chances d'être non nulle en t_0 et même probablement négative si la fréquence du signal est croissante. À l'inverse, si on considère un passage à zéro entre un maximum et un minimum, le problème est identique (à supposer que l'échelle locale est inversement proportionnelle à la fréquence instantanée) mais inversé et on en déduit que le signe de la moyenne locale est alors opposé au cas précédent. Finalement, on aboutit donc à une moyenne locale qui présente au moins autant d'extrema que le signal initial, ce qui est en désaccord avec l'idée selon laquelle la moyenne locale doit osciller plus lentement que le signal. Dans le cas d'une modulation d'amplitude linéaire, un raisonnement analogue aboutirait à la même conclusion. De fait, pour améliorer la définition (1.30), il faudrait que la forme de la fonction de pondération dépende des évolutions locales du signal, ce qui complique encore le problème.

Face à ces difficultés les créateurs de l'EMD proposent de définir la moyenne locale géométriquement à l'aide de la notion intuitive d'enveloppes d'un signal oscillant. Plus précisément, les enveloppes considérées s'appuient sur les extrema locaux du signal, l'enveloppe supérieure étant une courbe lisse interpolant les maxima et l'enveloppe inférieure une courbe lisse interpolant les minima. En pratique, elles sont le plus souvent calculées à l'aide d'une interpolation spline cubique mais la question du choix optimal de technique à utiliser pour calculer les enveloppes est ouverte. On discutera de ce problème plus en détail en 5.3.

Une fois les enveloppes déterminées, la moyenne locale est simplement définie comme la demi-somme de ces dernières. Par conséquent, la composante oscillant le plus rapidement peut a priori être obtenue à partir du signal en lui *soustrayant la moyenne de ses enveloppes*. Cependant, rien ne garantit que la composante ainsi définie vérifie la définition d'un IMF proposée précédemment. En effet, lorsqu'on soustrait la moyenne des enveloppes au signal, il est tout à fait possible de faire apparaître de nouveaux extrema qui ne sont pas nécessairement bien placés par rapport à la ligne zéro. Pour remédier à ce problème, les créateurs de l'EMD proposent d'itérer l'opération de *soustraction de la moyenne des enveloppes* jusqu'à ce que tous les maxima soient strictement positifs et tous les minima strictement négatifs. De plus, après un certain nombre d'itérations, on peut s'attendre à ce que la moyenne des enveloppes du signal soit proche de zéro, c'est-à-dire que les enveloppes sont presque symétriques par rapport à la ligne zéro. Il s'agit là en fait de la définition plus précise de la deuxième caractéristique de la définition d'un IMF qu'on avait laissée floue précédemment. En réalité, cette définition n'est pas très précise non plus et on la discutera plus en détail en 4.1. De plus, cette imprécision fait qu'on ne peut pas en pratique utiliser la définition d'un IMF comme critère pour décider quand arrêter l'itération et donc un autre critère doit être utilisé. Ce problème sera discuté en 5.2.

2.1.2.3 Opérateur de tamisage En résumé, la composante oscillant le plus rapidement est obtenue à partir du signal en itérant l'opération de *soustraction de la moyenne des enveloppes*. Cette procédure itérative est appelée dans la littérature « sifting process », ce qu'on peut traduire par « processus de tamisage ». Formellement, on définit ce processus comme l'itération d'un opérateur élémentaire de tamisage $\mathcal{S}$ défini par la procédure Algo. 1. La composante oscillant le plus rapidement est alors définie par l'itération de cet opérateur

$$(\mathcal{D}x)(t) = (\mathcal{S}^n x)(t), \quad (1.31)$$

où n est le nombre d'itérations déterminé selon un certain critère, par exemple l'un de ceux proposés en 5.2.

Algorithme 1 : Opérateur élémentaire de tamisage

-
- 1 Extraire les maxima et les minima de $x(t)$: $\{t_i^{max}, x_i^{max}\}, \{t_i^{min}, x_i^{min}\}$.
 - 2 Interpoler les ensembles de maxima $\{t_i^{max}, x_i^{max}\}$ et les ensembles de minima $\{t_i^{min}, x_i^{min}\}$ pour obtenir les enveloppes supérieure et inférieure : $e_{max}(t), e_{min}(t)$.
 - 3 Calculer la moyenne des enveloppes : $m(t) = (e_{max}(t) + e_{min}(t)) / 2$
 - 4 Soustraire la moyenne au signal : $\mathcal{S}[x](t) = x(t) - m(t)$
-

2.2 Transformée de Hilbert–Huang (HHT)

Dans la contribution originale, l'EMD n'est pas présentée individuellement mais comme la première étape d'un procédé d'analyse conçu pour les signaux non stationnaires et/ou non linéaires qui consiste à

1. appliquer l'EMD au signal : $x(t) \longrightarrow \{d_k(t), 1 \leq k \leq K\} \cup \{a_K(t)\}$
2. pour chaque IMF $d_k(t)$, calculer sa fréquence instantanée $f_k^i(t)$ (avec un exposant « i » pour instantané) et son amplitude instantanée $a_k^i(t)$ par la méthode du signal analytique (cf 1.2.2.1).

La combinaison de l'EMD et de l'estimation des fréquence et amplitude instantanées est connue sous le nom de transformée de Hilbert–Huang, (Hilbert–Huang Transform (HHT)).

Si on reprend la discussion, abordée en 1.2.2.1, des conditions nécessaires à l'interprétation des paramètres déterminés par la méthode du signal analytique en termes d'amplitude et de fréquence instantanées, on avait vu qu'une formulation standard de ces conditions était la condition de « bande étroite » selon laquelle la bande de fréquence occupée par le signal est étroite devant sa fréquence centrale. On avait également vu qu'un inconvénient de cette condition était qu'elle ne permettait pas de grandes variations de la fréquence instantanée, même lentes. Qualitativement, la définition d'IMF peut être vue comme proche d'une version locale de la condition de bande étroite et correspond donc à une tentative de formulation plus générale de cette dernière. En pratique, les signaux vérifiant cette définition ont en général effectivement des fréquences et amplitudes instantanées pertinentes mais il existe aussi des cas où les valeurs obtenues sont aberrantes [59].

Depuis la première présentation de la transformée de Hilbert–Huang [28], la méthode de calcul des paramètres instantanés par le signal analytique a montré un certain nombre de limites [26, 59], dues à la définition d'IMF qui n'est pas si bien adaptée à l'extraction des amplitude et fréquence instantanées par cette méthode. Ces limites ont conduit les créateurs de l'EMD et les membres de la communauté à envisager d'autres techniques [26, 8, 30]. À ce jour, un certain nombre de techniques sont étudiées et d'autres sont en cours d'élaboration. Le problème reste largement ouvert.

2.3 Définition plus générale d'une EMD

Réduite à ses principes les plus élémentaires, l'EMD se distingue essentiellement par deux caractéristiques :

- le choix de définir une échelle en s'appuyant sur les extrema locaux,
- l'extraction successive des échelles composant le signal de la plus fine à la plus grossière.

Plus généralement, on peut même proposer de ne pas considérer uniquement les extrema locaux mais un ensemble de points caractéristiques qui peuvent comprendre, par exemple, les extrema locaux, les points d'inflexion, les extrema de courbure, etc.

Partant de là, on peut proposer une formulation générale pour l'EMD ne retenant que ces éléments clés. Le cadre général reste donné par la formulation récursive proposée précédemment :

$$EMD(x) = \{\mathcal{D}x, EMD(x - \mathcal{D}x)\}. \quad (1.32)$$

Si $\mathcal{D}x$ correspond à la plus petite échelle du signal x , alors ce principe récursif rend compte de l'aspect « extraction successive des échelles de la plus fine à la plus grossière ».

Pour définir $\mathcal{D}$, on considère tout d'abord un ensemble de points caractéristiques du signal aux positions $\{t_i, i \in I\}$, correspondant par exemple aux extrema du signal. À cet ensemble de points, on associe le signal échantillonné

$$x_e(t) = \sum_{i \in I} x(t_i) \delta(t - t_i) \quad (1.33)$$

De là, on peut définir une moyenne locale du signal par une opération de filtrage généralisé

$$m(t) = \left(h \tilde{*} x_e \right) (t), \quad (1.34)$$

où h est un filtre passe-bas généralisé et $\tilde{*}$ représente une convolution généralisée. Dans cette formulation, il est important de noter que le filtre passe-bas généralisé h n'agit pas comme un filtre linéaire indépendant du signal échantillonné $x_e(t)$. En fait, l'opération de filtrage généralisé doit avoir un certain nombre de propriétés de covariance qui sont typiquement celles vérifiées par une opération d'interpolation :

multiplication par un scalaire : $\forall \lambda \in \mathbb{R}, x'_e(t) = \lambda x_e(t), \implies \left(h \tilde{*} x'_e \right) (t) = \lambda \left(h \tilde{*} x_e \right) (t)$

dilatation/inversion du temps : $\forall \lambda \in \mathbb{R}, x'_e(t) = x_e(\lambda t), \implies \left(h \tilde{*} x'_e \right) (t) = \left(h \tilde{*} x_e \right) (\lambda t)$

décalage temporel : $\forall \lambda \in \mathbb{R}, x'_e(t) = x_e(\lambda + t), \implies \left(h \tilde{*} x'_e \right) (t) = \left(h \tilde{*} x_e \right) (\lambda + t)$

ajout d'une constante : $\forall \lambda \in \mathbb{R}, x'_e(t) = \lambda + x_e(t), \implies \left(h \tilde{*} x'_e \right) (t) = \lambda + \left(h \tilde{*} x_e \right) (t)$

Étant donnée la moyenne locale $m(t)$ ainsi définie, on définit alors $\mathcal{D}x$

$$(\mathcal{D}x)(t) = x(t) - m(t) \quad (1.35)$$

$$= x(t) - \left(h \tilde{*} x_e \right) (t). \quad (1.36)$$

Dans cette formulation, on a choisi de définir $\mathcal{D}x$ de manière directe sans processus itératif, contrairement à ce qui est fait dans l'algorithme de l'EMD. Pour expliquer ce choix, on peut remarquer que l'itération du processus de tamisage dans l'EMD a essentiellement deux rôles : faire apparaître de nouveaux extrema qui étaient cachés du fait de la superposition d'échelles dans le signal et rendre les enveloppes plus symétriques. De fait, on peut voir que l'apparition d'extrema au cours du processus de tamisage peut être compensée dans la formulation générale par une meilleure sélection des points de contrôle $\{t_i, i \in I\}$. De même, l'amélioration de la symétrie des enveloppes peut être compensée par une meilleure définition de la moyenne locale.

On a également choisi de ne pas définir la moyenne locale à partir d'enveloppes mais directement à partir des points de contrôle par une opération analogue à un filtrage passe-bas. De fait, dans l'EMD standard, les enveloppes semblent jouer un rôle important qui pourrait justifier qu'on les inclue dans la formulation générale proposée ici. Cependant, il existe aussi des variantes de l'EMD standard aux propriétés très similaires qui n'utilisent pas la notion d'enveloppes et estiment la moyenne locale directement à partir des points de contrôle. Parmi ces variantes, on peut citer le cas où la moyenne locale est calculée en interpolant les points d'inflexion du signal [15], ou encore la méthode dite « B-spline EMD » [9]. Pour cette dernière méthode, la moyenne locale est définie directement à partir des extrema comme une somme de B-splines ayant pour ensemble de nœuds l'ensemble des extrema du signal. Les coefficients des B-splines dans la somme sont obtenus par des moyennes pondérées des extrema présents dans le support de chaque B-spline. La moyenne locale ainsi définie est comme dans le cas de l'EMD standard une spline ayant pour ensemble de nœuds les extrema mais elle n'est pas définie à partir d'enveloppes. Par rapport à l'EMD standard, cette méthode a l'avantage de mieux se prêter à une analyse théorique.

2.4 L'EMD : une méthode « intuitive »

L'EMD est souvent présentée comme une méthode intuitive. Pour éclairer cet aspect, on peut rappeler que l'EMD a été conçue pour imiter certains aspects de la vision de l'œil. Dans la contribution originale [28], les auteurs font en effet directement référence à la vision de l'œil en citant [17], selon qui « la première étape de l'analyse de données est d'examiner les données à l'œil ». Si on étudie en détail l'algorithme de l'EMD, on peut voir que deux étapes clés sont inspirées par la vision de l'œil : le choix des extrema comme points de contrôle et l'utilisation d'enveloppes.

Pour ce qui est des extrema, on sait que le cerveau humain lorsqu'il perçoit une image n'attribue pas la même attention à tous les endroits de l'image mais privilégie au contraire un certain nombre de points caractéristiques. Très grossièrement, ces points peuvent être vus comme les points où localement l'image varie le plus. La similarité avec l'EMD est ici patente puisque dans le cadre de l'EMD, on choisit également de s'appuyer sur un certain nombre de points caractéristiques et de plus les extrema sont assez proches de l'idée de points où le graphe du signal varie beaucoup. En fait, les extrema ne sont en général probablement pas exactement les points que l'œil aurait choisi comme points caractéristiques. En effet, pour repérer les extrema à partir du graphe du signal, il faut repérer des tangentes horizontales alors que la direction horizontale n'est pas nécessairement connue si on ne dispose que du graphe du signal. Si on considère par exemple un signal composé d'une onde sinusoïdale à laquelle est superposée une tendance linéaire, l'œil percevra sans difficulté la direction correspondant à la tendance linéaire mais pas nécessairement la direction horizontale. Pour se rapprocher de la vision humaine, il faudrait donc choisir des points de contrôle qui puissent être définis directement à partir du graphe du signal sans orientation privilégiée. Pour ce faire, on peut par exemple s'intéresser à la courbure du graphe du signal et repérer par exemple les points de courbure maximale qui correspondent bien à l'idée de points où le graphe varie le plus. C'est en fait à peu de chose près ce qu'ont proposé les créateurs de l'EMD dans [25] : ils proposent de remplacer les extrema du signal dans l'algorithme de l'EMD par les extrema de sa courbure définie par

$$c(t) = \frac{\frac{d^2x}{dt^2}}{\left(1 + \frac{dx}{dt}\right)^{3/2}}. \quad (1.37)$$

En pratique, un tel choix de points de contrôle peut paraître assez séduisant au premier abord mais il présente aussi un inconvénient important : le fait que les points de contrôle changent lorsque le signal est multiplié par un scalaire. Ce point est en particulier gênant si les grandeurs représentées en abscisse et en ordonnée sont de dimensions différentes (par exemple un temps et une position), auquel cas le graphe du signal n'a de sens qu'à des dilatations en abscisse et en ordonnée près. Ce problème explique peut-être pourquoi il n'y a pas eu à ce jour de suite à la proposition d'utiliser ces points pour l'EMD.

Plus encore que le choix des extrema comme points de contrôle, la notion d'enveloppes est typiquement visuelle. D'une certaine manière, la perception d'enveloppes par l'œil est un moyen de simplifier la représentation du graphe du signal : les enveloppes délimitent une forme globale et les évolutions du signal entre les enveloppes sont perçues à la manière d'une texture. L'utilisation des enveloppes dans l'EMD adopte une perspective très proche : les enveloppes contiennent l'information à grande échelle, la moyenne locale étant calculée uniquement à partir de ces dernières, alors que l'information à petite échelle est contenue dans les évolutions du signal entre les enveloppes. Toutefois, les enveloppes utilisées pour l'EMD sont clairement une version pauvre des enveloppes de la vision humaine. En effet, la vision humaine est capable de percevoir des enveloppes bien plus complexes à plusieurs échelles alors que les enveloppes de l'EMD sont limitées à l'échelle des oscillations du signal repérées par les extrema locaux. Si on considère par exemple le mouvement Brownien représenté Fig. 1.1, l'œil est capable de percevoir des enveloppes à plusieurs échelles, alors que les enveloppes de l'EMD sont limitées à la plus petite échelle. De fait, l'EMD est capable d'explorer les plus grandes

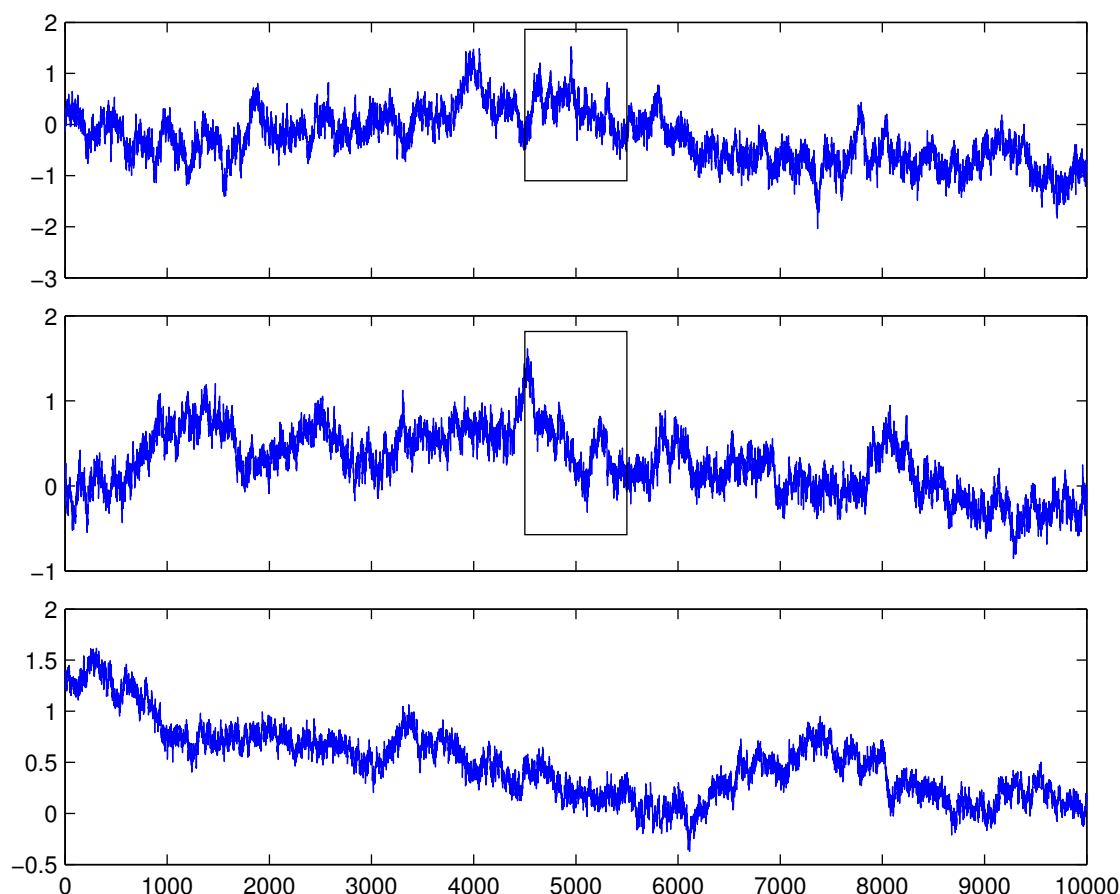


FIGURE 1.1 – Exemple d’un mouvement Brownien fractionnaire à plusieurs échelles. Les graphes du milieu et du bas représentent des agrandissements des parties encadrées des graphes du haut et du milieu respectivement. À chaque échelle, l’œil perçoit la courbe sous la forme d’un trait avec une certaine *épaisseur*, ce qui est très voisin de la perception d’enveloppes.

échelles après avoir extrait les plus petites alors que l’œil est capable de percevoir directement les grandes échelles. En revanche, l’œil n’est capable de percevoir effectivement les grandes échelles que quand celles-ci sont suffisamment importantes alors que l’EMD pourra théoriquement toujours les percevoir après avoir extrait les plus petites échelles.

3 Propriétés élémentaires de la décomposition

3.1 Non-linéarité

L’algorithme de l’EMD est globalement non linéaire. En effet, l’EMD d’une somme de deux signaux est en général différente de la somme des EMD des signaux séparés :

- le nombre d’IMF de la somme n’est pas contrôlé par les nombres d’IMFs des signaux séparés
- un IMF de la somme n’est pas généralement descriptible en termes de somme d’un ensemble quelconque d’IMFs des signaux séparés ni même une combinaison linéaire.

Dans l’algorithme de l’EMD, il y a au plus trois sources de non-linéarité. La première et la plus importante est dans le fait de s’appuyer sur les extrema. En effet, le nombre et la position des extrema dans une somme de signaux sont en général différents de ceux des signaux pris individuellement. Cette source de non-linéarité est fondamentale dans l’EMD dans la mesure où on choisit explicitement de définir les échelles de variations rapides d’un signal à l’aide de ses extrema.

Les deux autres sources de non-linéarité éventuelles sont l'interpolation et le critère d'arrêt du processus de tamisage qui selon l'implantation peuvent être non linéaires. En pratique, l'interpolation la plus utilisée (spline cubique) est linéaire mais la quasi-totalité des critères d'arrêt proposés sont susceptibles de créer des non linéarités (cf 5.2) en faisant varier les nombres d'itérations. En effet, si on suppose le reste de l'algorithme linéaire (en remplaçant par exemple une itération de tamisage par un filtrage linéaire prédéfini), la seule solution pour que le critère d'arrêt ne cause pas de non-linéarité est que les nombres d'itérations pour chaque IMF soient définis a priori², de manière à avoir

$$d_1[x + y] = \mathcal{S}^n[x + y] = \mathcal{S}^n[x] + \mathcal{S}^n[y] = d_1[x] + d_1[y], \quad (1.38)$$

et de même pour tous les IMFs extraits par la suite. En pratique, les critères d'arrêt sont justement faits pour adapter les nombres d'itérations au signal et ils créent donc par là même des non-linéarités supplémentaires. Cependant, utiliser des nombres d'itérations fixés a priori peut aussi être intéressant pour traiter des lots de données pour lesquels on voudrait autant que possible appliquer le même traitement à chaque échantillon. C'est en particulier le cas dans le cadre de l'une des variantes de l'EMD dite EMD d'ensemble (EEMD pour « Ensemble EMD ») qui sera décrite en 6.3. Enfin, on utilisera au cours de cette thèse essentiellement des nombres d'itérations fixés a priori, d'une part parce que tout critère d'arrêt aboutit in fine à un nombre d'itérations et d'autre part parce que les propriétés de l'EMD s'expriment plus simplement à partir des nombres d'itérations que des paramètres d'un critère d'arrêt.

3.2 Covariances vis-à-vis de certaines transformations

L'EMD est covariante par rapport aux transformations linéaires de l'espace et du temps, c'est-à-dire qu'elle commute avec les opérations suivantes

multiplication par un scalaire : L'opérateur élémentaire de tamisage commute avec les multiplications par un scalaire :

$$\mathcal{S}\{\lambda x\} = \lambda \mathcal{S}\{x\}. \quad (1.39)$$

Cette propriété provient du fait que la position des extrema d'un signal est inchangée par les multiplications par un scalaire positif. La sélection des extrema étant l'unique étape non linéaire de l'opérateur de tamisage, on en déduit que celui-ci commute avec les multiplications par des scalaires positifs. De plus, l'opérateur commute aussi avec l'opération de changement de signe puisque celle-ci a pour effet d'interchanger les deux enveloppes et de changer leur signe, ce qui résulte finalement en un simple changement de signe de leur moyenne. Au final, on obtient donc que l'opérateur de tamisage commute avec toutes les multiplications par un scalaire.

Si le critère d'arrêt utilisé pour arrêter le processus de tamisage est également indépendant de l'amplitude et du signe du signal, ce qui est le cas en général, l'EMD entière est alors covariante avec les multiplications par un scalaire.

inversion du sens du temps : Cette propriété est en fait dépendante du schéma d'interpolation. Si celui-ci est covariant par rapport à l'inversion du sens du temps

$$\forall t \in \mathbb{R}, \quad I(\{-t_i, x_i\})(-t) = I(\{t_i, x_i\})(t), \quad (1.40)$$

alors l'EMD aussi. En pratique, sauf cas particulier, il est généralement souhaitable d'avoir une telle propriété et tous les schémas d'interpolation envisagés pour l'EMD la vérifient.

2. Il n'est pas nécessaire qu'ils soient les mêmes pour tous les IMFs, tant qu'ils ne dépendent pas du signal.

dilatation du temps : L'algorithme de l'EMD est covariant par rapport aux dilatations du temps si et seulement si le schéma d'interpolation utilisé l'est

$$\forall t \in \mathbb{R}, \quad I(\{\lambda t_i, x_i\})(\lambda t) = I(\{t_i, x_i\})(t). \quad (1.41)$$

À nouveau, cette propriété est généralement souhaitable et tous les schémas d'interpolation envisagés pour l'EMD la vérifient.

décalage temporel : L'EMD est covariante par rapport aux décalages temporels si et seulement si l'interpolation l'est, ce qui est encore une fois systématique

$$\forall t \in \mathbb{R}, \quad I(\{t_i + \delta, x_i\})(t + \delta) = I(\{t_i, x_i\})(t). \quad (1.42)$$

ajout d'une constante : L'opérateur de tamisage est invariant par rapport à l'ajout d'une constante :

$$\mathcal{S}\{x + \alpha\} = \mathcal{S}\{x\}. \quad (1.43)$$

Il en découle que les IMFs sont invariants par rapport à l'ajout d'une constante et les approximations, dont le résidu, sont covariantes.

$$d_k[x + \alpha](t) = d_k[x](t), \quad (1.44)$$

$$a_k[x + \alpha](t) = a_k[x](t) + \alpha. \quad (1.45)$$

3.3 Quasi-orthogonalité

3.3.1 Le problème de l'orthogonalité

La quasi-orthogonalité de la décomposition est étroitement liée à son aspect multirésolution dans la mesure où le produit scalaire de deux ondes de fréquences différentes est habituellement faible et que les IMFs sont construits de telle manière qu'à tout instant ils oscillent avec des périodes différentes. Dans cet esprit, il est suggéré dans la contribution originale [28] que la décomposition en « oscillations rapides » et « oscillations lentes » serait (quasi-)orthogonale. Le raisonnement sous-jacent est d'assimiler localement la partie « oscillations lentes » à la composante continue (au sens de Fourier) du signal décomposé. La partie « oscillations rapides » n'ayant alors (localement) pas de composante continue, serait nécessairement (localement) orthogonale à la partie « oscillations lentes ».

Selon ce raisonnement, essentiellement deux aspects sont susceptibles de faire obstacle à l'orthogonalité de la décomposition :

1. L'orthogonalité est une propriété globale et non locale. La non-stationarité peut compromettre l'orthogonalité globale alors qu'elle semble localement assurée.
2. La moyenne des enveloppes d'un signal ne coïncide pas nécessairement avec la moyenne au sens intégrale, même localement.

S'il est clair que le premier point peut nuire à l'orthogonalité globale de la décomposition (cf Fig. 1.2), cela n'a pas forcément de conséquences fâcheuses dans la mesure où le problème est peut-être plutôt que la notion d'orthogonalité globale est inadaptée à une méthode qui veut analyser les signaux localement. La notion d'orthogonalité locale en revanche, tout comme les concepts de fréquence locale, pose des problèmes de définition, ce qui explique qu'on se contente de l'orthogonalité globale en pratique.

La différence entre moyenne intégrale et moyenne des enveloppes est plus gênante dans la mesure où celle-ci ne pose pas de problème de définition de l'orthogonalité. En effet, on peut aisément trouver des cas de signaux périodiques simples pour lesquels les deux moyennes diffèrent, par exemple

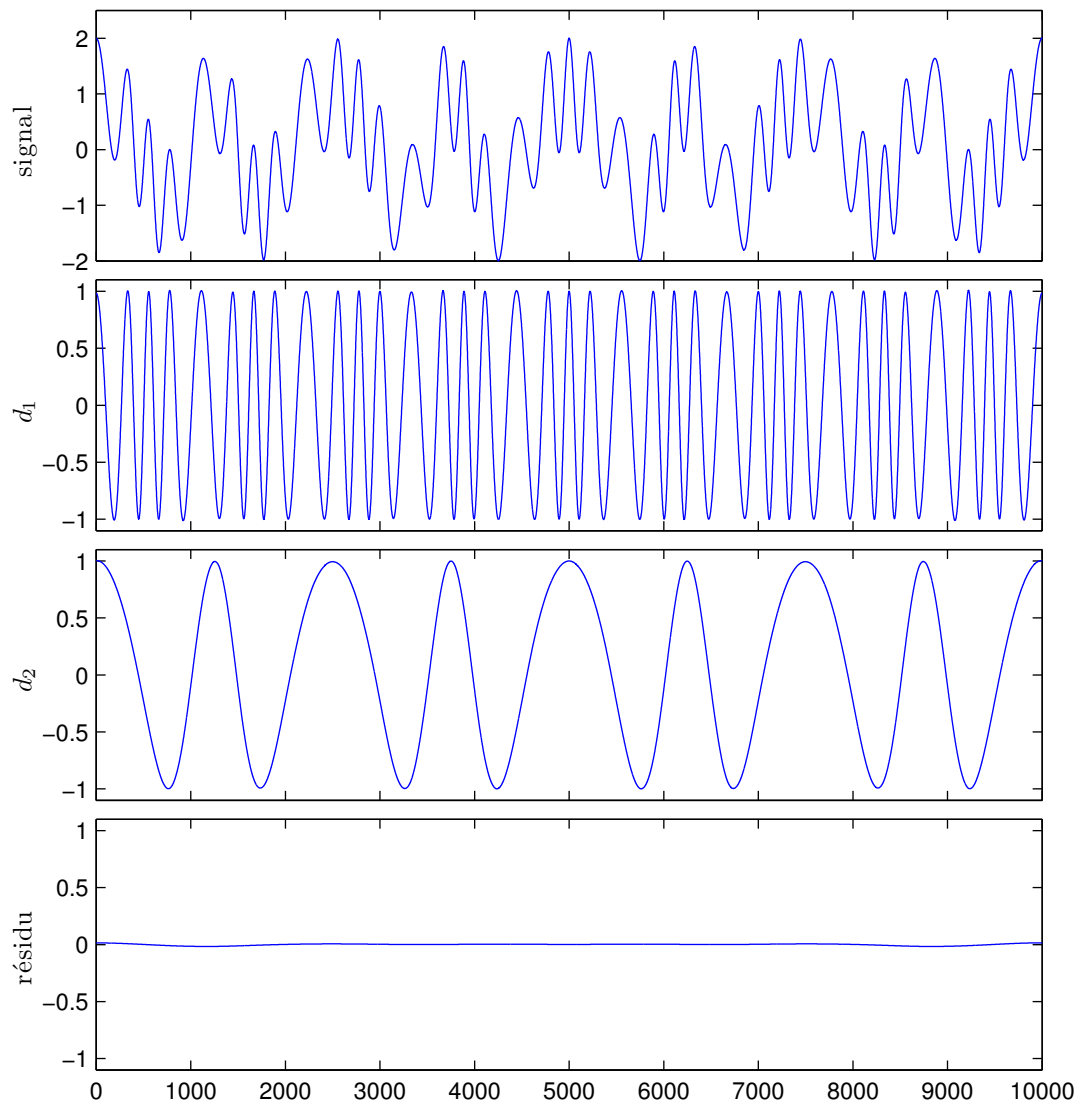


FIGURE 1.2 – Exemple où les non-stationnarités sont à l’origine d’une décomposition (très légèrement) non orthogonale. Du fait des non-stationnarités, les deux composantes ont une valeur moyenne non nulle et ne sont donc pas parfaitement orthogonales.

$x(t) = \cos t + 0.2 \cos(2t)$ a une moyenne intégrale nulle et une moyenne d'enveloppes de 0.2. Ainsi le signal

$$x(t) = \cos t + 0.2 \cos(2t) + \cos ft + 0.2 \cos(2ft), \quad (1.46)$$

avec $f \gg 1$ est décomposé par l'EMD en (cf Fig. 1.3) :

$$d_1(t) = \cos t + 0.2 \cos(2t) - 0.2 \quad (1.47)$$

$$d_2(t) = \cos ft + 0.2 \cos(2ft) - 0.2, \quad (1.48)$$

$$m_2(t) = 0.4 \quad (1.49)$$

où il est clair que les termes constants compromettent l'orthogonalité de la décomposition.

Enfin, toujours dans un contexte périodique, on remarque qu'au-delà de la différence entre moyenne intégrale et moyenne des enveloppes, il existe des situations où la non-orthogonalité peut avoir d'autres origines moins évidentes. Ainsi, le signal [31]

$$x(t) = \cos 2t + 0.25 \cos t + 0.25 \cos 3t \quad (1.50)$$

est décomposé par l'EMD en (cf Fig. 1.4) :

$$d_1(t) = \cos 2t + 0.25 \cos 3t - 0.25 \cos t + 0.0563 \quad (1.51)$$

$$d_2(t) = 0.5 \cos t - 0.0563. \quad (1.52)$$

Il s'agit d'un cas un peu particulier où l'EMD a un comportement contraire à l'intuition qui s'explique très bien par le modèle développé en 1.1.

3.3.2 Quantifier l'orthogonalité

L'orthogonalité entre deux IMFs peut être quantifiée directement par un indice d'orthogonalité égal à leur produit scalaire normalisé

$$IO_{ij} = \frac{\langle d_i(t), d_j(t) \rangle}{\sqrt{\langle d_i(t), d_i(t) \rangle \cdot \langle d_j(t), d_j(t) \rangle}}. \quad (1.53)$$

Globalement, l'orthogonalité de la décomposition peut alors être décrite par une matrice d'orthogonalité IO de terme général IO_{ij} . Cependant, une telle représentation présente l'inconvénient de donner la même importance aux IMFs de petite et de grande énergie alors qu'en pratique, il est clair que la non-orthogonalité est d'autant plus problématique que les énergies des IMFs correspondants sont grandes. Pour cette raison, il est généralement plus judicieux lorsqu'on s'intéresse à l'orthogonalité globale de la décomposition de ne pas normaliser les produits scalaires, ou de tous les normaliser par la même quantité, par exemple la norme L^2 du signal analysé. De là, on peut alors définir un indice d'orthogonalité global pour la décomposition de la forme

$$IO = \frac{\sum_{i < j} |\langle d_i(t), d_j(t) \rangle|}{\langle x(t), x(t) \rangle}. \quad (1.54)$$

Dans la contribution originale [28], un indice d'orthogonalité est également défini comme

$$IO = \sum_{i < j} \left\langle \frac{d_i(t)}{x(t)}, \frac{d_j(t)}{x(t)} \right\rangle. \quad (1.55)$$

Il est bon de souligner que cette définition provient en fait d'un raisonnement faux et qu'elle ne mesure donc en rien l'orthogonalité de la décomposition (sans parler de problèmes de définition si

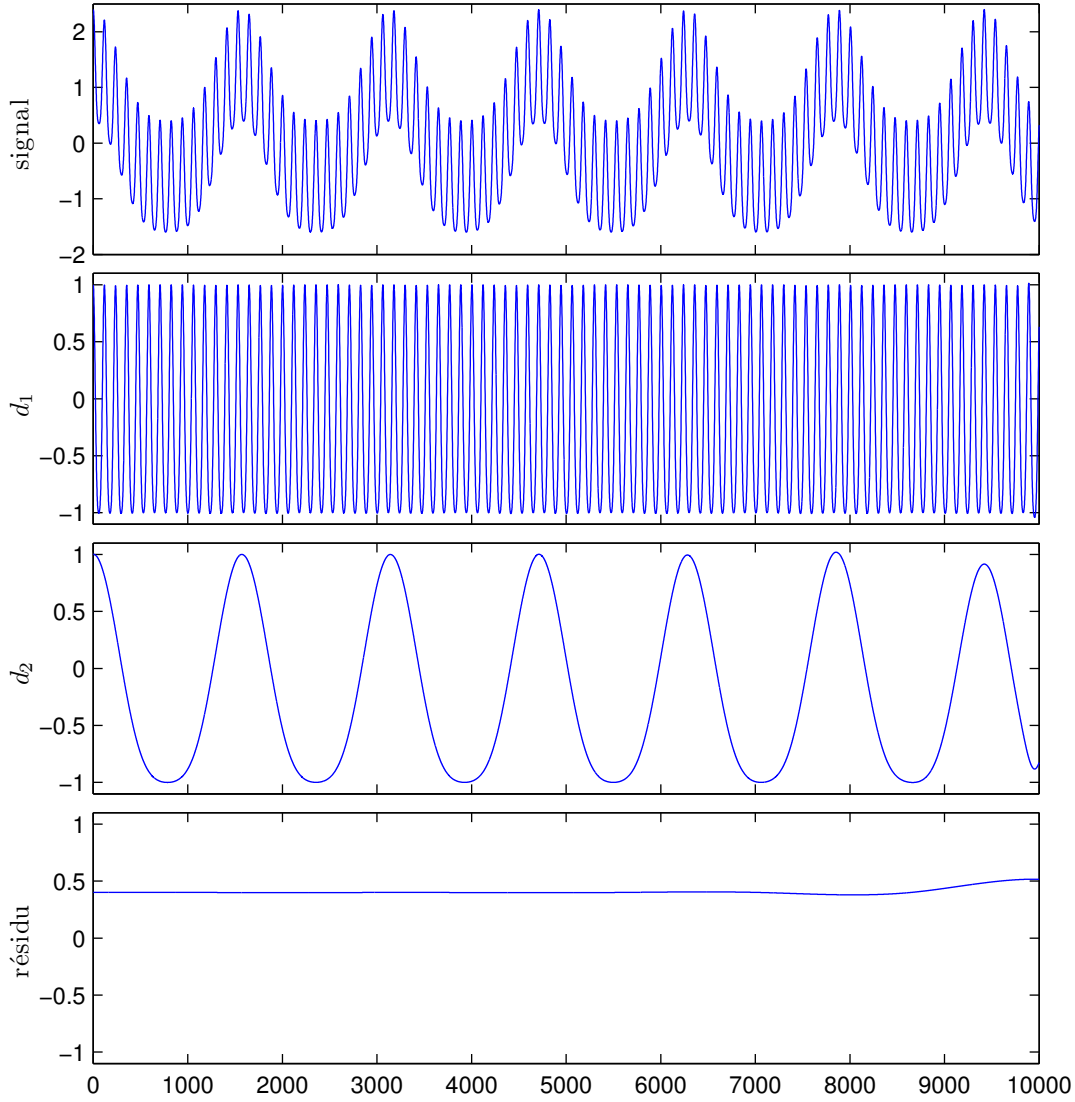


FIGURE 1.3 – Exemple de décomposition peu orthogonale due à la présence de composantes continues dans les IMFs. La décomposition a été limitée à 2 IMFs pour plus de clarté. Le signal analysé est (1.46)

$x(t)$ s'annule sur l'intervalle considéré). En corrigeant le raisonnement, on aboutit à la définition suivante :

$$IO = \sum_{i < j} \frac{\langle d_i(t), d_j(t) \rangle}{\langle x(t), x(t) \rangle}. \quad (1.56)$$

Il est clair avec cette dernière définition que si la décomposition est orthogonale, l'indice d'orthogonalité est nul. Malheureusement, la réciproque est fausse, raison pour laquelle on préférera par exemple la définition (1.54). Néanmoins, il semble que la définition originale (1.55), ou peut-être sa version corrigée (1.56), ait été employée dans un certain nombre de travaux utilisant l'EMD, certains s'appuyant dessus dans leurs analyses [29, 32].

3.4 Localité

Lors de l'étape centrale de l'EMD, on soustrait au signal une « moyenne locale », définie comme la demi-somme de son enveloppe supérieure et de son enveloppe inférieure. En pratique, le terme de

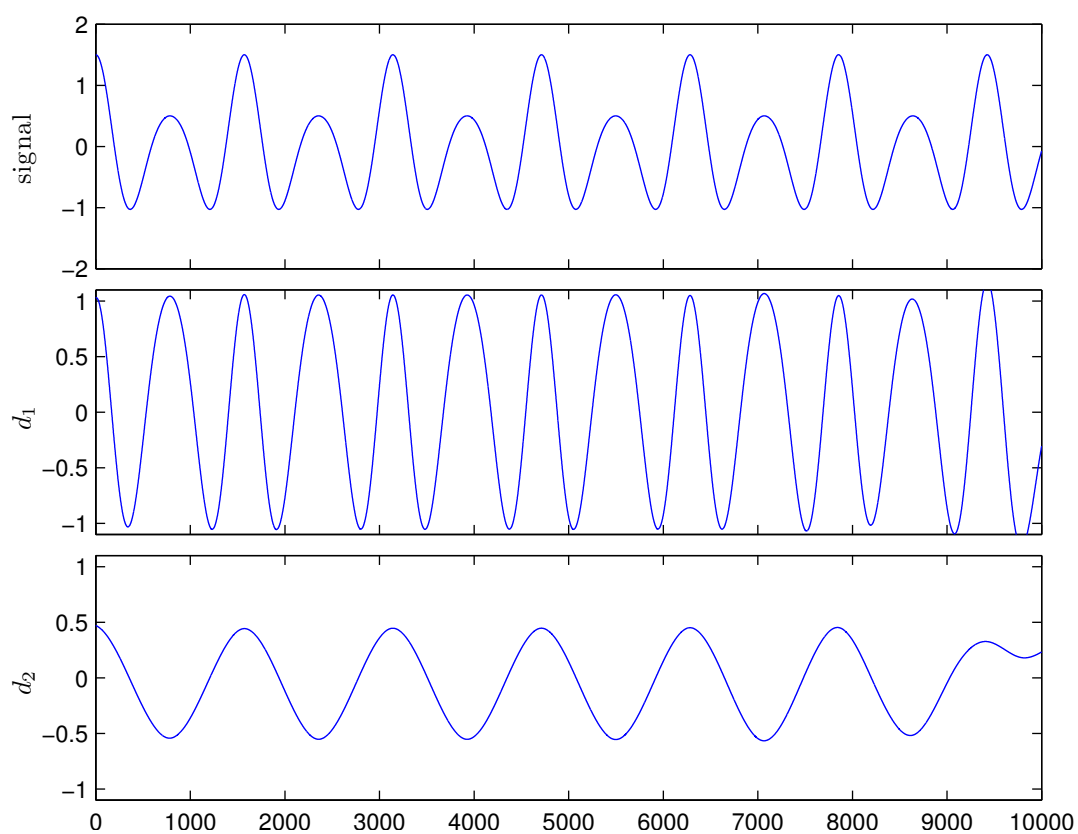


FIGURE 1.4 – Exemple de décomposition peu orthogonale où la non-orthogonalité n'est pas uniquement due aux composantes continues. Le signal analysé est (1.50)

moyenne locale est un oxymore, ce qui fait qu'il n'est pas évident de comprendre sa signification. Une terminologie plus appropriée serait sans doute de parler de « moyenne à l'échelle locale », qui permet de mettre en évidence le fait que la moyenne locale est intimement liée à une échelle qui dépend de la position. Dans le cadre de l'EMD, l'échelle locale est définie par les extrema. Pour s'en convaincre, il suffit de rappeler que la moyenne locale est définie à partir des enveloppes du signal qui interpolent les maxima et les minima. La notion d'enveloppe étant intrinsèquement locale, on en déduit que la valeur, par exemple de l'enveloppe supérieure, en un point donné dépend essentiellement des deux maxima qui entourent ce point et dans une moindre mesure des autres maxima plus éloignés. Par conséquent, on peut considérer que les enveloppes, et donc la moyenne locale, sont définies localement à l'échelle correspondant à l'espacement entre les maxima/minima, c'est-à-dire grossièrement deux fois l'échelle correspondant à l'espacement entre les extrema, maxima et minima confondus.

La moyenne locale d'un signal x étant définie de manière locale à l'échelle de l'espacement entre les maxima/minima, il en découle que le signal après une itération de tamisage $\mathcal{S}x$ dépend lui aussi de x de manière locale à l'échelle de l'espacement entre les maxima/minima. Si on itère maintenant le tamisage, le même raisonnement reste valable, du moins si on suppose qu'il n'y a pas apparition de nouveaux extrema et que les extrema présents initialement ne se déplacent pas trop d'une itération à la suivante. Si au contraire des extrema apparaissent au cours du processus de tamisage à l'itération k , alors $\mathcal{S}^{k+1}x$ dépend de $\mathcal{S}^k x$ à l'échelle correspondant à ses extrema et donc de x à l'échelle des extrema de $\mathcal{S}^k x$. À terme, si on suppose que les extrema initialement présents dans x ne se sont pas trop déplacés au cours du processus de tamisage, on peut considérer selon ce raisonnement que $\mathcal{S}^n x$, n étant le nombre d'itérations final, dépend de x localement à l'échelle correspondant aux extrema de $\mathcal{S}^n x$. En réalité, ce raisonnement ne donne qu'une borne inférieure de l'échelle à laquelle $\mathcal{S}^n x$ dépend de x , l'écart venant du fait que la valeur d'une enveloppe en un point ne dépend pas uniquement des deux maxima/minima voisins mais aussi de ceux qui sont plus éloignés, du moins dans le cas de l'interpolation spline cubique. Pour s'apercevoir de ce problème, on peut observer Fig. 1.5 que les enveloppes de $\mathcal{S}^n x$ où x est un bruit blanc gaussien sont de plus en plus lisses à mesure qu'on augmente le nombre d'itérations. Le fait que les enveloppes soient lisses implique que les échantillons de $\mathcal{S}^n x$ sont corrélés à une échelle correspondant au moins à l'échelle de variation des enveloppes. Les échantillons du bruit blanc étant par définition décorrélés, on en déduit que nécessairement $\mathcal{S}^n x$ dépend de x à une échelle qui augmente avec le nombre d'itérations. Malheureusement, déterminer l'échelle à laquelle $\mathcal{S}^n x$ dépend de x en fonction de n n'est pas simple dans le cas général. Néanmoins, le problème pourrait être approché dans des situations simples à l'aide du modèle développé en 1.1.

Sachant que $\mathcal{S}^n x$ dépend de x localement à une échelle au moins de l'ordre de l'espacement entre les maxima/minima de $\mathcal{S}^n x$, on a que le premier IMF $d_1[x]$ dépend de x localement à une échelle au moins de l'ordre de l'espacement entre les maxima/minima de $d_1[x]$. Par conséquent, la première approximation $a_1[x]$ dépend aussi de x à la même échelle. Si on poursuit le raisonnement, on obtient que le deuxième IMF $d_2[x]$ dépend de $a_1[x]$ à une échelle au moins de l'ordre de l'espacement des maxima/minima de $d_2[x]$. Par conséquent, l'échelle à laquelle le deuxième IMF dépend du signal est au moins de l'ordre de la somme de l'espacement des maxima/minima de $d_2[x]$ et de $d_1[x]$. En pratique, l'espacement des maxima/minima de $d_2[x]$ étant supérieur à celui de $d_1[x]$, on pourra souvent considérer que l'échelle à laquelle $d_2[x]$ dépend du signal est essentiellement contrôlée par les extrema de $d_2[x]$ mais le résultat peut-être faux en particulier si le nombre d'itérations de tamisage utilisé pour calculer $d_1[x]$ est très supérieur à celui utilisé pour calculer $d_2[x]$.

Finalement, si on applique le raisonnement aux IMFs et approximations suivants, on aboutit aux résultats :

- l'approximation $a_k[x]$ dépend localement de x à une échelle au moins de l'ordre de la somme des espacements des maxima/minima (autour du point considéré) dans les IMFs $d_{k'}[x]$, $k' < k$.
- l'IMF $d_k[x]$ dépend localement de x à une échelle au moins de l'ordre de la somme des espacements des maxima/minima (autour du point considéré) dans les IMFs $d_{k'}[x]$, $k' \leq k$.

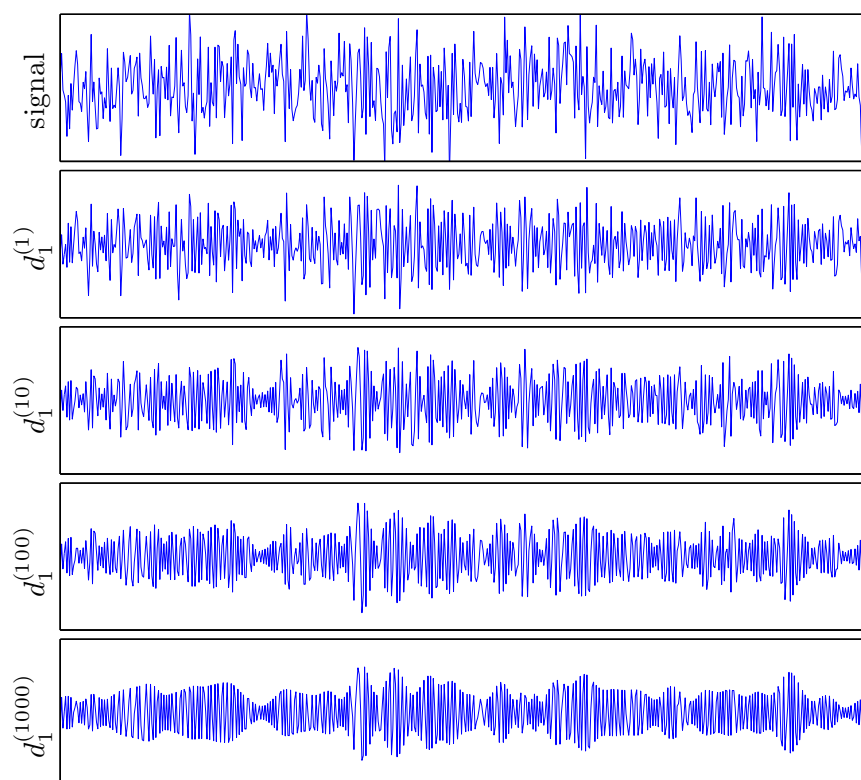


FIGURE 1.5 – Évolution du premier IMF issu d'un bruit blanc gaussien en fonction du nombre d'itérations. Au fur et à mesure que le nombre d'itérations augmente, les enveloppes sont de plus en plus lisses.

3.5 Multirésolution

L'EMD réalise une décomposition multi-échelles, ou multirésolution, dans la mesure où elle explore successivement les échelles du signal de la plus fine, représentée par le premier IMF, à la plus grossière, représentée par le dernier IMF ou le résidu.

Par rapport à d'autres méthodes d'analyse multirésolution, dont la référence est la transformée en ondelettes, l'EMD présente un certain nombre de particularités. Tout d'abord, à l'instar de la transformée en ondelettes discrète, elle propose une décomposition en échelles discrètes dans la mesure où la décomposition est constituée d'un nombre fini de composantes. En revanche, les échelles de l'EMD diffèrent significativement des échelles de la transformée en ondelettes discrète. Elles se distinguent par les caractéristiques suivantes :

définies par les extrema La notion d'échelle dans l'EMD est associée à l'espacement entre les extrema. Cette notion diffère fortement de la notion d'échelle dans le cadre de la transformée en ondelettes, par exemple, où l'échelle est définie de manière relative par comparaison avec une forme d'onde donnée, à savoir l'ondelette. De fait, le concept d'échelle dans l'EMD est tellement différent qu'il a été proposé d'utiliser à la place du terme « échelle », le terme « empquency » [39] de la contraction de « empirical mode frequency » qui correspond à l'inverse de l'échelle et qui est défini comme

$$f_e = \frac{1}{2d}, \quad (1.57)$$

où d est l'espacement entre les deux extrema qui entourent le point considéré. Le point de vue adopté par l'EMD permet d'associer une même échelle à des formes d'ondes très différentes. Cet aspect est mis en évidence par l'exemple proposé Fig. 1.7 où l'EMD réalise une décomposition en échelles avec des formes d'ondes variables. Ceci contraste avec le cas où les échelles sont définies par comparaison avec une forme d'onde donnée. Dans ce cas, une forme d'onde très différente de la forme d'onde de référence est généralement décrite par plusieurs échelles.

adaptativité Les échelles des IMFs sont déterminées par les échelles présentes dans le signal et non par une grille prédéterminée comme dans le cas des transformées en ondelettes discrètes. Pour ces dernières, les échelles sont le plus souvent de la forme $a2^k$, $k \in \mathbb{N}$ (transformée en ondelettes dyadique), où a est une échelle de référence.

localité En vertu de la propriété de localité de l'EMD étudiée en 3.4, l'échelle d'un IMF n'est pas définie de manière globale mais de manière locale, le caractère local étant relatif à l'espacement entre les extrema.

Les propriétés d'adaptativité et de localité des échelles sont bien mises en évidence par l'exemple proposé Fig. 1.6.

L'adaptativité et la localité sont aussi à la source d'un des défauts de l'EMD connu sous le nom de « mélange de modes » (ou « mode mixing »). Ce problème se rencontre notamment quand le signal est constitué de plusieurs composantes dont certaines ne sont pas présentes sur toute la durée du signal. Il arrive dans ce cas que des composantes qu'on aurait aimé voir dans un seul IMF soient réparties sur plusieurs comme dans le cas de l'exemple représenté Fig. 1.8. Pour corriger ce phénomène, diverses possibilités ont été envisagées qui reviennent toutes à imposer l'échelle d'un IMF [25] ou du moins à fortement la guider [65, 14].

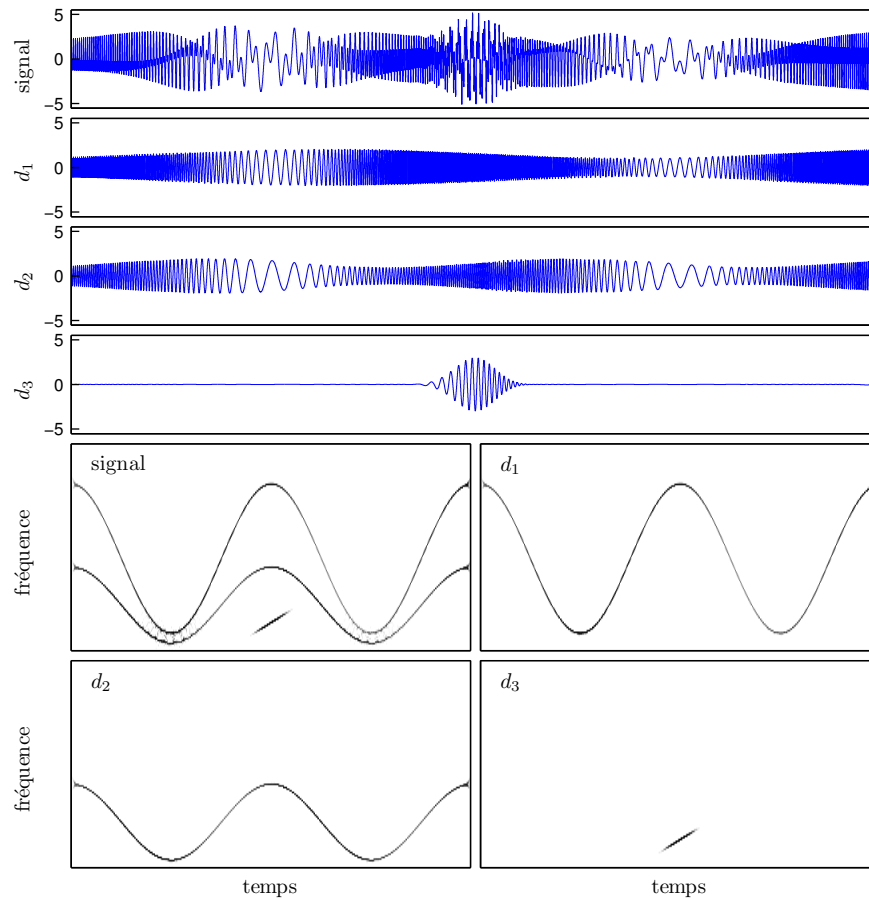


FIGURE 1.6 – Le signal de la rangée du haut est décomposé par l’EMD en les 3 IMF’s présentés et 6 autres qu’on n’a pas représentés parce qu’ils sont pratiquement nuls (ils ne contiennent que 0.3% de l’énergie du signal). L’analyse temps-fréquence (spectrogramme réalloué) du signal montre la présence de 3 composantes temps-fréquence dont les supports temporels et fréquentiels se recouvrent, ce qui exclut de les séparer par une méthode non adaptative. Les représentations temps-fréquence des 3 IMF’s principaux mettent en évidence le fait qu’ils capturent efficacement la structure en 3 composantes du signal.

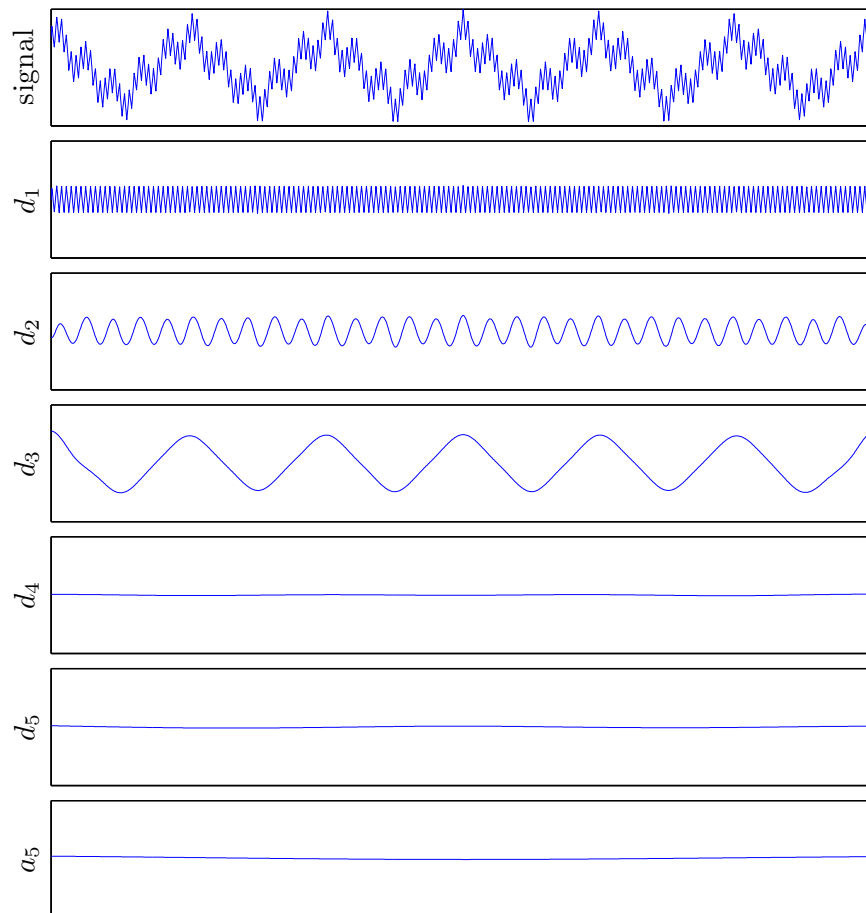


FIGURE 1.7 – Le signal de la rangée du haut est constitué de trois composantes : deux composantes oscillantes triangulaires, l’une à haute fréquence, l’autre à basse fréquence et une composante sinusoïdale entre les deux. La décomposition proposée par l’EMD reproduit presque exactement les 3 composantes, la seule différence étant que la composante triangulaire de basse fréquence voit ses extrema adoucis.

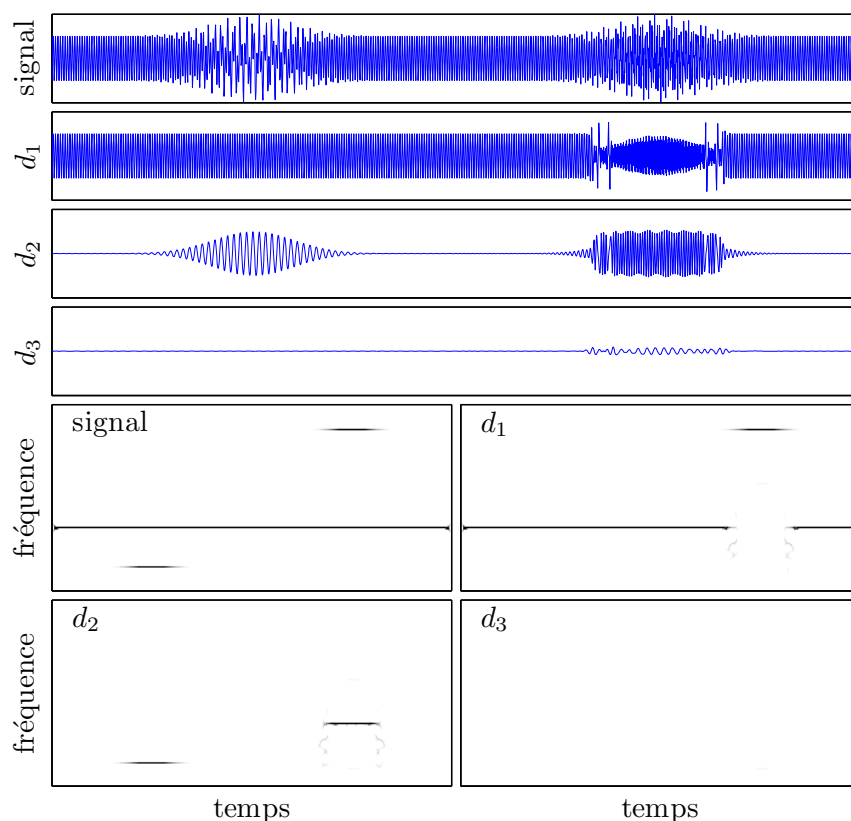


FIGURE 1.8 – Le signal de la rangée du haut est constitué de trois composantes : une composante sinusoïdale permanente et deux composantes sinusoïdales localisées en temps, une de fréquence plus haute que celle de la composante permanente et une de fréquence plus basse. Les représentations temps-fréquence (spectrogrammes réalloués) des IMFs montrent que le premier IMF capture à tout instant la composante de plus haute fréquence et qu'il contient donc la composante permanente sauf lorsque la composante localisée de plus haute fréquence est présente. La partie de la composante permanente située à l'emplacement de la composante localisée haute fréquence est alors décalée dans le deuxième IMF.

4 Propriétés des IMFs

4.1 Une définition en toute rigueur trop restrictive

4.1.1 Définition d'un IMF

Définition 3. Une fonction $f : [a, b] \rightarrow \mathbb{R}$ est une « fonction modale intrinsèque » ssi [28] :

- (i) tous les maxima locaux de f sont strictement positifs, tous les minima locaux sont strictement négatifs.
- (ii) la somme de l'enveloppe supérieure interpolant les maxima de f et de l'enveloppe inférieure interpolant les minima est nulle.

Quelques remarques sur cette définition. Tout d'abord, la première condition est plus souvent donnée sous la forme :

« le nombre d'extrema et de passages à zéro doivent différer d'au plus un »

Les deux formes sont équivalentes dès lors que les extrema et passages à zéro sont en nombre fini. De ce fait, la première forme, telle que donnée dans la définition, est un peu plus générale que la forme usuelle et peut donc être intéressante dans une perspective théorique. En revanche, les signaux étant en pratique discrets et de taille finie, les deux définitions sont totalement équivalentes dans la pratique. La différence est alors essentiellement une question de point de vue : la forme utilisée dans la définition prend un point de vue local alors que la forme alternative, plus courante dans la littérature, prend un point de vue global. L'EMD ayant été créée pour traiter de manière locale les signaux non stationnaires, il semble naturel que la définition d'un IMF soit formulée de manière locale.

Deuxièmement, il s'agit d'une définition incomplète. En effet, si la première clause est sans ambiguïté, ce n'est pas le cas de la seconde qui nécessite de préciser ce qu'on entend par « interpolation ». Le problème est que les enveloppes dépendent en fait non seulement du schéma d'interpolation utilisé mais également de la manière dont sont traitées les conditions aux bords. De ce fait, la définition d'un IMF est en fait indissociable d'une implantation de l'EMD.

4.1.2 Effet d'une interpolation polynomiale par morceaux

On s'intéresse ici à préciser les conséquences de la deuxième clause de la définition d'un IMF dans le cas où le schéma d'interpolation est polynomial par morceaux entre les maxima/minima. On montre que la prise en compte rigoureuse de la deuxième clause contraint de manière excessive les enveloppes en imposant qu'elles soient non plus polynomiales par morceaux mais globalement polynomiales, le degré du polynôme étant a priori le même que celui de l'interpolation.

Pour démontrer ce résultat, considérons une fonction $f : [a, b] \rightarrow \mathbb{R}$ et t_0 la position de l'un de ses minima locaux. Au voisinage de t_0 , l'enveloppe inférieure interpolant les minima est composée d'un polynôme $P_g(t)$ à gauche de t_0 et $P_d(t)$ à droite. En supposant que t_0 n'est pas également la position d'un maximum local de f — toujours vérifié en pratique —, l'enveloppe supérieure est localement autour de t_0 décrite par un polynôme $P_u(t)$. Dans ces conditions, si f est un IMF, la deuxième clause de la définition au voisinage de t_0 s'écrit alors :

$$P_u(t) = -P_g(t), \quad \text{si } t < t_0 \quad (1.58)$$

$$P_u(t) = -P_d(t), \quad \text{si } t > t_0. \quad (1.59)$$

Étant donné que P_u, P_g et P_d sont des polynômes, on en déduit que nécessairement $P_g = -P_u = P_d$. Le raisonnement étant valable pour tous les extrema, les enveloppes supérieure et inférieure d'un IMF

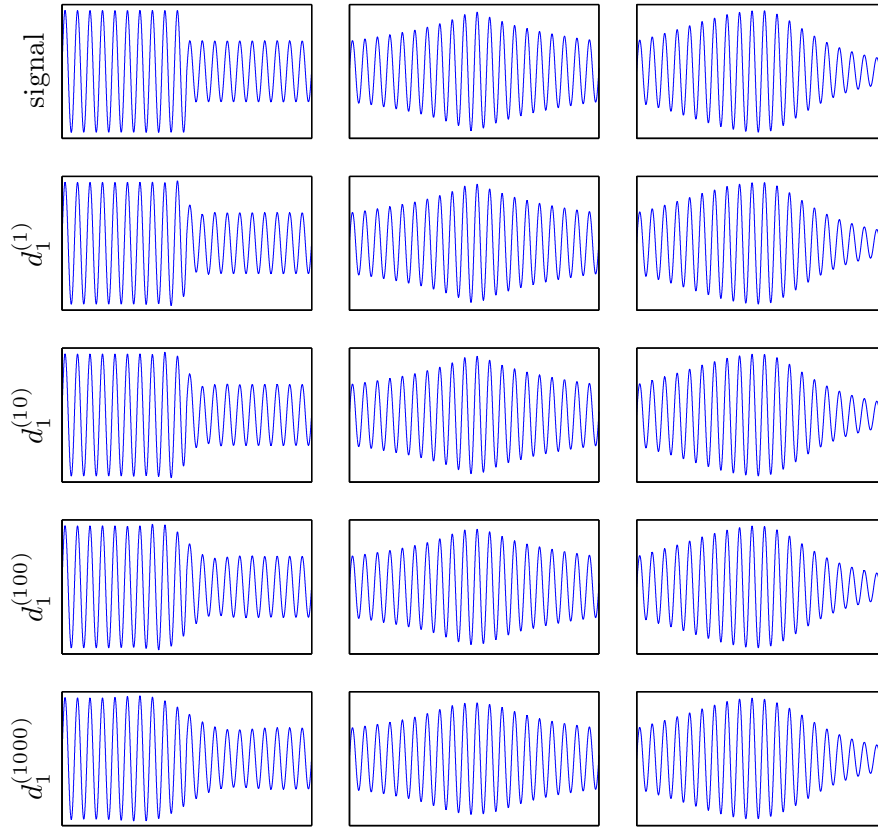


FIGURE 1.9 – Effets de lissage des discontinuités des enveloppes ainsi que de leurs dérivées première et seconde. Les trois colonnes correspondent de gauche à droite aux discontinuités des enveloppes, de leur dérivée première et de leur dérivée seconde. Le signal initial est représenté sur la première ligne, ses premiers IMFs obtenus avec 1, 10, 100 et 1000 itérations sur les lignes suivantes. Le gommage de la discontinuité de la dérivée seconde est difficile à voir sur cette figure mais on peut le voir en agrandissant l’enveloppe, voir Fig. 1.10.

sont alors globalement des polynômes. Dans le cas relativement courant d’interpolations cubiques par morceaux, les enveloppes sont donc contraintes à être des polynômes de degré 3³.

En pratique, le fait que les points fixes de l’opérateur de tamisage aient des enveloppes nécessairement cubiques ne signifie pas que le processus de tamisage itéré à l’infini converge nécessairement vers des IMFs aux enveloppes cubiques. En revanche, on peut dire que les enveloppes tendent vers une allure cubique localement dans la mesure où le processus de tamisage gomme plus ou moins rapidement les discontinuités des enveloppes ou de leurs dérivées 1 et 2 (cf Fig. 1.9).

4.1.3 Vers une définition pratique d’un IMF

Pour les raisons évoquées précédemment, la deuxième clause de la définition d’un IMF est peu utilisable en l’état car trop rigoureuse. Une possibilité est alors de considérer que la somme des deux enveloppes d’un IMF ne vaut pas strictement zéro mais est par exemple plus petite en valeur absolue

3. Les auteurs de [59] aboutissent quant à eux à la conclusion que les enveloppes sont nécessairement des polynômes de degré 2. La raison en est vraisemblablement qu’ils supposent que l’ensemble de définition de f peut être $\mathbb{R}$. Dans ce cas en effet, l’enveloppe supérieure d’un IMF étant typiquement positive, elle ne peut-être décrite que par un polynôme de degré pair.

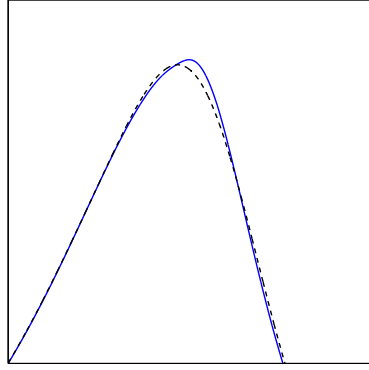


FIGURE 1.10 – Agrandissement de l’enveloppe supérieure du signal (trait plein) correspondant à la colonne de droite dans Fig. 1.9 et l’enveloppe supérieure de son premier IMF calculé à l’aide de 1000 itérations de tamisage (tirets). Le déplacement de l’enveloppe correspond à un gommage de la discontinuité de la dérivée seconde de l’enveloppe du signal.

qu’une certaine quantité $\epsilon > 0$ [59] :

$$\forall t \in [a, b], \quad |U(t) + L(t)| \leq \epsilon. \quad (1.60)$$

Une telle définition, si elle présente des intérêts théoriques, est en pratique limitée pour deux raisons :

- Elle dépend de la normalisation du signal. Pour que la définition d’IMF soit indépendante de toute normalisation, il conviendrait que ϵ soit par exemple proportionnel à une certaine norme du signal.
- Elle n’est pas locale. La valeur de ϵ est la même à tout instant même si l’amplitude du signal varie de manière importante au cours du temps. Ainsi, la fonction $x(t) = e^t(\alpha + \sin t)$ avec $\alpha < 1$ serait une IMF sur $] -\infty, \ln \epsilon/\alpha[$ et ne le serait pas sur $] \ln \epsilon/\alpha, +\infty[$ alors que l’allure locale de la fonction ne change pas au cours du temps si ce n’est par son amplitude $x(t + 2k\pi) = e^{2k\pi}x(t)$.

Pour pallier à ces deux défauts, on peut par exemple remplacer ϵ par une quantité proportionnelle à une « amplitude locale » du signal $A(t)$:

$$\forall t \in [a, b], \quad |U(t) + L(t)| \leq \epsilon A(t), \quad \text{avec } A(t) = |U(t) - L(t)|, \quad (1.61)$$

où on définit $A(t)$ comme l’écart entre l’enveloppe supérieure et l’enveloppe inférieure. On aurait bien sûr pu imaginer d’autres définitions pour l’« amplitude locale ». En pratique cependant, on peut remarquer que pour définir une amplitude locale on rencontre essentiellement le même problème que pour définir une moyenne locale, à savoir le choix d’une échelle d’observation locale adaptée aux données. Dans la mesure où la moyenne des enveloppes est introduite dans l’EMD justement comme un moyen d’estimer cette moyenne locale sans avoir à explicitement définir une échelle locale, définir l’amplitude locale comme — ou à partir de — l’écart entre les enveloppes semble nécessaire pour rester dans un cadre cohérent.

Si on remplace la deuxième clause de la définition d’un IMF par cette nouvelle formulation (1.61), on obtient une définition relativement satisfaisante dans la mesure où elle est a priori utilisable en pratique sans trop contraindre les enveloppes et qu’elle respecte les exigences de localité nécessaires à un traitement adéquat des signaux non stationnaires.

Enfin, le processus de tamisage étant censé être arrêté dès lors que le signal en cours de traitement est un IMF, une telle définition est a priori directement utilisable comme critère d’arrêt du processus de tamisage. Le critère d’arrêt « local » (cf 5.2) est basé sur ce principe. La forme utilisée est toutefois un peu plus souple encore pour éviter des situations de blocage où le critère $|U(t) + L(t)| \leq \epsilon |U(t) - L(t)|$ ne serait jamais vérifié pour tout t alors qu’il le serait sur quasiment tout l’intervalle de définition $[a, b]$.

4.2 Les IMFs d'indice supérieur à 2 sont des splines

On se propose de montrer ici que les approximations ainsi que les IMFs autres que le premier sont des sommes d'interpolations. Dans le cas particulier où l'interpolation utilisée est de la famille des splines, ce qu'on supposera par la suite, ce résultat a comme conséquence que les approximations et IMFs autres que le premier sont également des splines. Cette propriété provient du fait que la somme de deux splines de même degré est une spline du même degré. Plus précisément, si $s_1(t)$ et $s_2(t)$ sont des splines de degré n avec pour ensembles de nœuds respectifs N_1 et N_2 , alors $s_1(t) + s_2(t)$ est une spline de degré n sur $N_1 \cup N_2$. Pour se convaincre de ce résultat, il suffit de rappeler qu'une spline de degré n sur un ensemble ordonné de nœuds $\{t_i\}_{i \in I}$ n'est autre qu'une fonction polynomiale de degré n sur chacun des intervalles $[t_i, t_{i+1}]$ dont les dérivées successives sont continues en chaque t_i jusqu'à l'ordre $n - 1$ [13].

Dans la suite, nous utiliserons les notations suivantes :

$m[x](t)$: moyenne des enveloppes de $x(t)$.

$d_i[x](t)$: i^{e} IMF.

$a_i[x](t)$: i^{e} approximation. Par convention $a_0[x](t) = x(t)$.

$n_i[x]$: nombre d'itérations de tamisage utilisé pour calculer le i^{e} IMF. $d_i[x](t) = (\mathcal{S}^{n_i[x]} a_{i-1}[x])(t)$, où $\mathcal{S}$ est l'opérateur élémentaire de tamisage correspondant à l'opération de soustraire la moyenne des enveloppes : $(\mathcal{S}x)(t) = x(t) - m[x](t)$.

$d_i^{(j)}[x](t)$: i^{e} IMF en cours de calcul après j itérations de tamisage. $d_i^{(j)}[x](t) = (\mathcal{S}^j a_{i-1}[x])(t)$, avec $j \leq n_i[x]$. Par convention $d_i^{(0)}[x](t) = a_{i-1}[x](t)$.

$a_i^{(j)}[x](t)$: i^{e} approximation en cours de calcul après j itérations de tamisage. $a_i^{(j)}[x](t) = a_{i-1}[x](t) - d_i^{(j)}[x](t)$.

$E_i^j[x]$: ensemble des abscisses des extrema de $d_i^{(j)}[x](t)$.

S_E^n : espace des splines de degré n dont les nœuds sont dans E . Dans la suite, le degré des splines étant systématiquement celui des splines de l'interpolation, l'exposant n sera omis.

Proposition 1. *Les approximations successives par EMD, $a_k[x](t)$ d'un signal $x(t)$ sont toutes des splines*

$$a_k[x] \in S_{\mathcal{E}_k[x]}, \quad \text{où} \quad \mathcal{E}_i[x] = \bigcup_{j=0}^{n_i[x]-1} E_i^j[x]. \quad (1.62)$$

Démonstration. Commençons par remarquer que la moyenne des enveloppes d'un signal $x(t)$ est une spline sur les extrema de $x(t)$. En effet, l'enveloppe supérieure étant une spline sur les maxima et l'enveloppe inférieure une spline sur les minima leur demi-somme est alors une spline sur les extrema. Ainsi, $a_k^{(1)}[x]$ étant simplement la moyenne des enveloppes de $a_{k-1}[x]$, on peut écrire :

$$a_k^{(1)}[x] \in S_{E_k^0[x]}.$$

De là, on démontre aisément par récurrence que

$$a_k^{(i)}[x] \in S_{\bigcup_{j=0}^{i-1} E_k^j[x]},$$

sachant que

$$a_k^{(i+1)}[x](t) = a_k^{(i)}[x](t) + d_k^{(i)}[x](t) - d_k^{(i+1)}[x](t),$$

c'est-à-dire

$$a_k^{(i+1)}[x](t) = a_k^{(i)}[x](t) + m[d_k^{(i)}[x]](t). \quad \square$$

Corollaire 1. Les IMFs $d_k[x](t)$ d'ordre $k \geq 2$ d'un signal $x(t)$ sont des splines

$$d_k[x] \in S_{\mathcal{E}_k[x] \cup \mathcal{E}_{k-1}[x]}, \quad \text{où} \quad \mathcal{E}_i[x] = \bigcup_{j=0}^{n_i[x]} E_i^j[x]. \quad (1.63)$$

Démonstration.

$$d_k(t) = a_{k-1}(t) - a_k(t). \quad \square$$

5 Questions liées à l'implantation

5.1 Conditions aux bords

L'algorithme de l'EMD s'appuie intensivement sur la notion d'enveloppes d'un signal qui sont définies comme interpolant les maxima et minima locaux. S'il est clair qu'une telle définition est susceptible de fournir des enveloppes pertinentes pour tout point situé entre deux extrema, elle pose en revanche le problème de ne pas définir les enveloppes sur les bords du signal, au-delà du dernier extremum. De plus, la solution consistant à ne considérer que la portion du signal entre son second et son avant-dernier extremum (pour avoir les premiers et derniers maxima et minima) est vouée à l'échec parce que les positions des extrema peuvent changer d'une itération à la suivante, ce qui peut obliger à réduire à chaque itération la partie du signal réellement analysée. Enfin, si l'interpolation utilisée est suffisamment lisse, on peut toujours extrapoler la valeur des enveloppes au-delà des derniers extrema, mais il se trouve que les enveloppes ne sont dans ce cas pas suffisamment bien contrôlées aux bords et peuvent prendre des valeurs rapidement aberrantes. Ce problème ne serait pas très gênant si les variations aberrantes des enveloppes aux bords n'affectaient que la décomposition aux bords, mais il se trouve que les effets de bords ainsi créés se propagent bien au-delà des bords quand on itère le tamisage.

Pour mettre en œuvre l'algorithme de l'EMD, il faut donc une méthode pour définir les enveloppes jusqu'aux bords du signal qui doit

1. limiter l'erreur au niveau des bords
2. limiter la propagation de l'erreur vers la partie centrale du signal

En pratique, il est clair qu'on ne peut définir une telle erreur que si l'on connaît a priori la décomposition en IMFs, ce qui fait qu'on ne peut pas l'évaluer dans le cas général. En dehors des cas où l'erreur introduite est suffisamment importante pour être visible à l'œil, il est donc difficile de percevoir l'influence que peut avoir la méthode utilisée pour prolonger les enveloppes sur la décomposition finale. Pour évaluer différentes méthodes, on choisira donc de les appliquer sur des signaux de synthèse dont on connaît a priori la décomposition en IMFs.

À ce jour, plusieurs solutions ont été proposées, qu'on peut regrouper en trois familles

utilisation d'une fonction fenêtre l'idée est ici de multiplier le signal par une fonction fenêtre, valant 1 sur la partie centrale du signal et s'atténuant doucement jusqu'à zéro aux bords. Grâce à cette opération, on peut résoudre le problème de définition aux bords des enveloppes en leur imposant de valoir zéro aux bords. Proposée comme un pis-aller par [14], puis plus sérieusement par [46], cette méthode est en fait très limitée puisque le simple fait de rajouter une constante, ou pire une tendance non stationnaire, au signal peut aboutir à de grandes variations aux bords après avoir multiplié par la fenêtre.

prolonger le signal l'idée est ici de prolonger le signal au-delà des bords de manière à faire apparaître de nouveaux extrema permettant de définir correctement les enveloppes aux bords du signal. Le problème est vaste et constitue à lui seul tout le domaine de recherche de la « prédiction des séries temporelles ». Plusieurs techniques ont été envisagées, notamment des

méthodes à base de réseaux de neurones [16] et de machines à vecteurs de support [33, 63]. Les auteurs de [33] proposent ainsi d'utiliser un outil actuellement populaire dans le domaine : la régression par machine à vecteurs de support (SVM). Malheureusement, ils n'envisagent dans leur article que le noyau linéaire et n'appliquent leur méthode que sur des signaux parfaitement linéaires (sommées de trois modes de Fourier) ou avec une forte composante périodique (signal réel). De fait, il semble que la manière dont la méthode est employée la limite à une simple prédiction linéaire, ce qui explique qu'elle marche bien sur les exemples proposés, mais que pour des signaux plus compliqués comme ceux utilisés dans la suite pour évaluer un certain nombre de méthodes, elle échoue très fréquemment, simplement parce que l'apprentissage n'arrive pas à converger. Les auteurs de [63] proposent une méthode a priori plus performante à base de machine à vecteur de support utilisant un noyau apparemment gaussien. Le fait d'utiliser un noyau non linéaire constitue une amélioration très nette par rapport au noyau linéaire dans la mesure où cela permet de mieux prédire des données non linéaires. De plus, les auteurs proposent une méthode de choix automatique des paramètres de la machine à vecteurs de support (algorithme dit « particle swarm optimization »), nécessaire pour rendre la méthode utilisable dans le cas général.

prolonger les ensembles d'extrema plus simple que prolonger le signal on peut aussi se contenter de prolonger les ensembles d'extrema. D'un point de vue théorique, le problème est tout aussi compliqué que prévoir le signal mais le fait qu'en pratique il suffise de rajouter quelques points laisse la place à des méthodes empiriques avec disjonctions de cas qui auraient été trop compliquées à définir si l'on avait voulu prévoir le signal au complet. Le coût en temps de calcul de la prédiction d'extrema est aussi naturellement bien moindre que celui de la prédiction du signal. Une telle méthode de prédiction d'extrema consistant essentiellement à opérer une symétrie miroir aux bords du signal a été proposée dans [55], malheureusement sans la détailler, et est implantée dans les programmes que nous mettons à disposition sur internet. Deux autres méthodes ont été proposées dans [12]. La première s'apparente à une symétrie miroir et la seconde se rapproche d'une prédiction linéaire des extrema.

Dans la suite de cette section, on se propose d'évaluer les performances de plusieurs méthodes de prolongation du signal ou des extrema. Dans ce but, on se donne un modèle de signal constitué d'une tendance lente déterministe $x_t(t)$ et d'une composante oscillante aléatoire $x_o(t)$ dont l'amplitude varie lentement de manière déterministe et la période rapidement de manière aléatoire. Les échelles de temps des deux composantes sont choisies suffisamment différentes pour qu'on puisse considérer sans ambiguïté que la composante oscillante est le premier IMF de la décomposition et que la tendance est la première approximation (ou le résidu). Plus précisément, la composante oscillante est définie comme le produit d'une modulation d'amplitude lente et d'une fonction oscillante d'amplitude unité. Cette dernière est définie par sa fréquence instantanée qui est elle-même un bruit basse fréquence auquel on ajoute une constante pour le rendre positif. La fréquence de coupure du filtre utilisé pour définir le bruit et la constante qui lui est ajoutée sont choisies de telle manière que la fréquence instantanée varie à l'échelle d'une période. L'idée derrière ce choix est de mettre en difficulté toute méthode peu robuste qui s'appuierait fortement sur une quasi-périodicité du signal à proximité des bords. Finalement,

$$x(t) = x_t(t) + x_o(t), \quad (1.64)$$

avec $x_t(t)$ une tendance lente déterministe

$$x_o(t) = a(t)f(t), \quad (1.65)$$

avec $a(t) > 0$ une modulation d'amplitude lente déterministe et $f(t)$ une modulation de fréquence aléatoire variant rapidement à l'échelle de sa période.

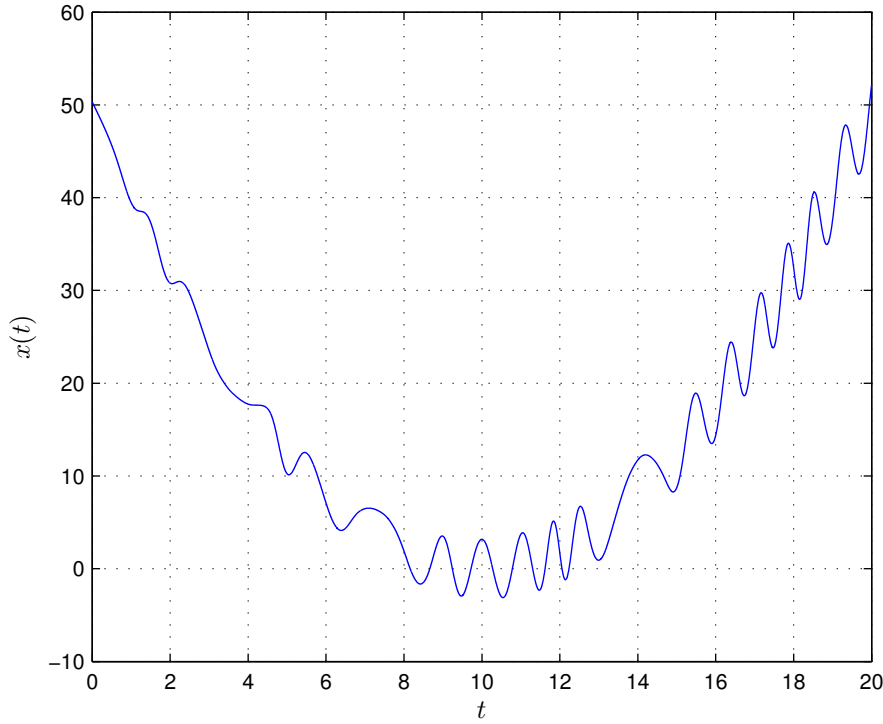


FIGURE 1.11 – Exemple de signal utilisé pour l'étude des performances des techniques de gestion des effets de bords. La composante oscillante a une période presque fixe mais qui fait des sauts relativement régulièrement.

Pour cette étude, on a choisi les paramètres suivants :

$$t \in [0, 20] \quad (1.66)$$

$$x_t(t) = \frac{(t - 10)^2}{2} \quad (1.67)$$

$$a(t) = 1 + e^{t/5}. \quad (1.68)$$

La fréquence instantanée de $f(t)$ est construite à partir d'un bruit blanc filtré à l'aide d'un filtre passe-bas de Butterworth d'ordre 2 et de fréquence de coupure 0.25. Le bruit basse fréquence ainsi obtenu est ensuite redimensionné pour avoir une valeur minimale de 0.1 et une valeur maximale de 1.6 avant d'être utilisé comme fréquence instantanée pour $f(t)$. Le signal $f(t)$ ainsi créé a une période qui varie en général relativement peu de proche en proche mais qui peut présenter des variations importantes avec une probabilité non négligeable. Un exemple d'une réalisation du signal $x(t)$ défini avec ces paramètres est proposé Fig. 1.11.

Pour présenter les méthodes testées, on explique comment prolonger les ensembles d'extrema (ou le signal) au niveau du bord gauche, le cas du bord droit étant symétrique. Le bord gauche sera associé à l'abscisse 0 pour simplifier. On utilisera les notations suivantes

$x(t)$: signal

$t_{\max/\min}(k), 1 \leq k \leq n_{\max/\min}$: abscisse du k^e maximum/minimum

$x_{\max/\min}(k), 1 \leq k \leq n_{\max/\min}$: ordonnée du k^e maximum/minimum

Les extrema rajoutés sur le bord gauche seront notés avec des indices $k \leq 0$. Ne connaissant pas le signal au-delà des bords, les points situés aux extrémités ne sont a priori pas considérés comme des

extrema. Certaines méthodes de prolongement proposeront de les rajouter comme extrema supplémentaires.

Les méthodes testées sont les suivantes :

p1 Cette méthode traite les enveloppes indépendamment et consiste simplement à rajouter un extremum par enveloppe en projetant le premier extremum à l'abscisse 0 en conservant son ordonnée

$$t_{max}(0) = 0, \quad x_{max}(0) = x_{max}(1).$$

Idem pour les minima.

p2 Idem **p1** si $x(0) < x_{max}(1)$, $x_{max}(0) = x(0)$ sinon. Idem pour les minima. Cette variante permet de mieux contrôler la valeur au bord.

s1 Deux maxima rajoutés par symétrie miroir par rapport au bord

$$x_{max}(0) = x_{max}(1), \quad t_{max}(0) = -t_{max}(1) \quad (1.69)$$

$$x_{max}(-1) = x_{max}(2), \quad t_{max}(-1) = -t_{max}(2). \quad (1.70)$$

Idem pour les minima.

s2 Deux maxima rajoutés par symétrie miroir avec contrôle de la valeur au bord. Pour commencer, le point du bord $(0, x(0))$ est rajouté à l'ensemble des maxima si $x(0) > x_{max}(1)$ ou à celui des minima si $x(0) < x_{min}(1)$. Ensuite, 2 maxima et 2 minima sont ajoutés par symétrie miroir, ou éventuellement 1 si le point du bord a déjà été ajouté à l'ensemble considéré. L'axe de symétrie est le même pour les maxima et les minima. Il est situé généralement à l'abscisse de l'extremum le plus proche du bord (celui-ci compte s'il a été rajouté précédemment) mais il est remplacé par le bord lui-même si pour l'une des deux enveloppes, les 2 extrema ainsi rajoutés ont des abscisses toutes deux positives. Cette méthode est celle proposée dans [55] et implantée dans les codes que nous mettons à disposition sur internet.

DS1 Deux maxima rajoutés selon

$$x_{max}(0) = x_{max}(1), \quad t_{max}(0) = t_{max}(1) - (t_{max}(2) - t_{max}(1)) \quad (1.71)$$

$$x_{max}(-1) = x_{max}(1), \quad t_{max}(-1) = t_{max}(1) - 2(t_{max}(2) - t_{max}(1)). \quad (1.72)$$

Idem pour les minima. Cette méthode est proche de la première méthode proposée dans [12], seul un extremum étant rajouté dans la version originale. Le fait d'ajouter un extremum supplémentaire rend la méthode plus robuste au cas où $t_{max}(2) - t_{max}(1) < t_{max}(1)$.

DS2 Deux extrema sont ajoutés à chaque enveloppe en itérant 2 fois la procédure suivante décrite dans le cas où le premier extremum est un maximum :

$$x_{max}(0) = x_{max}(1) + \left(\frac{x_{max}(2) - x_{min}(1)}{t_{max}(2) - t_{min}(1)} + \frac{x_{max}(1) - x_{min}(1)}{t_{max}(1) - t_{min}(1)} \right) (t_{max}(2) - t_{max}(1)), \quad (1.73)$$

$$t_{max}(0) = t_{max}(1) - (t_{max}(2) - t_{max}(1)) \quad (1.74)$$

$$x_{min}(0) = x_{min}(1) + \left(\frac{x_{max}(2) - x_{min}(1)}{t_{max}(2) - t_{min}(1)} + \frac{x_{max}(1) - x_{min}(1)}{t_{max}(1) - t_{min}(1)} \right) (t_{min}(2) - t_{min}(1)), \quad (1.75)$$

$$t_{min}(0) = t_{min}(1) - 2(t_{min}(2) - t_{min}(1)). \quad (1.76)$$

Cette procédure est pratiquement identique à la deuxième méthode proposée dans [12]. Comme précédemment, seul un extremum est rajouté dans la version originale. De plus, la méthode est normalement accompagnée de conditions aux limites pour l'interpolation spline cubique qui n'ont pas été utilisées ici faute de code performant permettant d'ajuster ces conditions aux limites. Le fait de rajouter un point de plus que dans la méthode originale permet dans une certaine mesure de compenser ce problème.

pl Prolongement linéaire des extrema avec gestion de cas particuliers. Pour commencer, deux extrema sont ajoutés selon (pour les maxima)

$$x_{max}(0) = x_{max}(1) - (x_{max}(2) - x_{max}(1)), \quad t_{max}(0) = t_{max}(1) - (t_{max}(2) - t_{max}(1)) \quad (1.77)$$

$$x_{max}(-1) = x_{max}(1) - 2(x_{max}(2) - x_{max}(1)), \quad t_{max}(-1) = t_{max}(1) - 2(t_{max}(2) - t_{max}(1)), \quad (1.78)$$

si $t_{max}(1) - 2(t_{max}(2) - t_{max}(1)) < 0$ et selon

$$x_{max}(0) = x(0), \quad t_{max}(0) = 0 \quad (1.79)$$

$$x_{max}(-1) = x(0) - (x_{max}(1) - x(0)), \quad t_{max}(-1) = -t_{max}(1), \quad (1.80)$$

sinon.

Dans le second cas, le point $(0, x(0))$ est ajouté à l'ensemble des maxima s'il se trouve au dessus de la ligne brisée reliant les maxima tels que définis à l'étape précédente. Les maxima ajoutés sont alors recalculés selon le premier cas.

SVM Prédiction du signal à l'aide d'une régression par machine à vecteurs de support. Le signal est prédit sur une durée variable de telle manière qu'au moins deux maxima et minima apparaissent sur cette durée. On utilise un noyau gaussien comme proposé dans [63] mais contrairement à ce qui est préconisé dans l'article, on utilise une méthode simple pour choisir les paramètres de la machine à vecteurs de support. La taille est choisie proportionnellement à la période moyenne (relative à l'écart entre les extrema) du signal au voisinage du bord, la valeur précise du paramètre ayant été réglée à la main pour avoir une prédiction satisfaisante sur le modèle de signal envisagé. La machine à vecteurs de support est entraînée sur les 4 périodes les plus proches du bord. Le paramètre de tolérance à l'erreur (parfois appelé « largeur du tube ») est réglé de manière à avoir une erreur faible devant l'amplitude d'oscillation (écart entre les valeurs des maxima et des minima) au voisinage du bord. Il en résulte que le nombre de vecteurs de support peut être important et le temps de calcul conséquent, de l'ordre de 5000 fois supérieur à celui des autres méthodes.

Considérant le modèle de signal décrit précédemment, on en génère 1000 réalisations et on calcule le premier IMF à l'aide des différentes méthodes proposées. Dans chaque cas, on définit la fonction erreur comme la différence entre la première approximation et la tendance lente du signal

$$err(t) = a_1(t) - x_t(t). \quad (1.81)$$

À partir de cette erreur, on calcule alors deux quantités dans le but d'évaluer d'une part l'erreur aux bords du signal et d'autre part la propagation de cette dernière au centre du signal

err_{moy} moyenne quadratique de l'erreur

$$err_{moy} = \sqrt{\frac{1}{20} \int_0^{20} err(t)^2 dt}. \quad (1.82)$$

err_{centre} moyenne quadratique de l'erreur sur la partie centrale du signal (entre $\max(t)/4$ et $3 \max(t)/4$)

$$err_{centre} = \sqrt{\frac{1}{10} \int_5^{15} err(t)^2 dt} \quad (1.83)$$

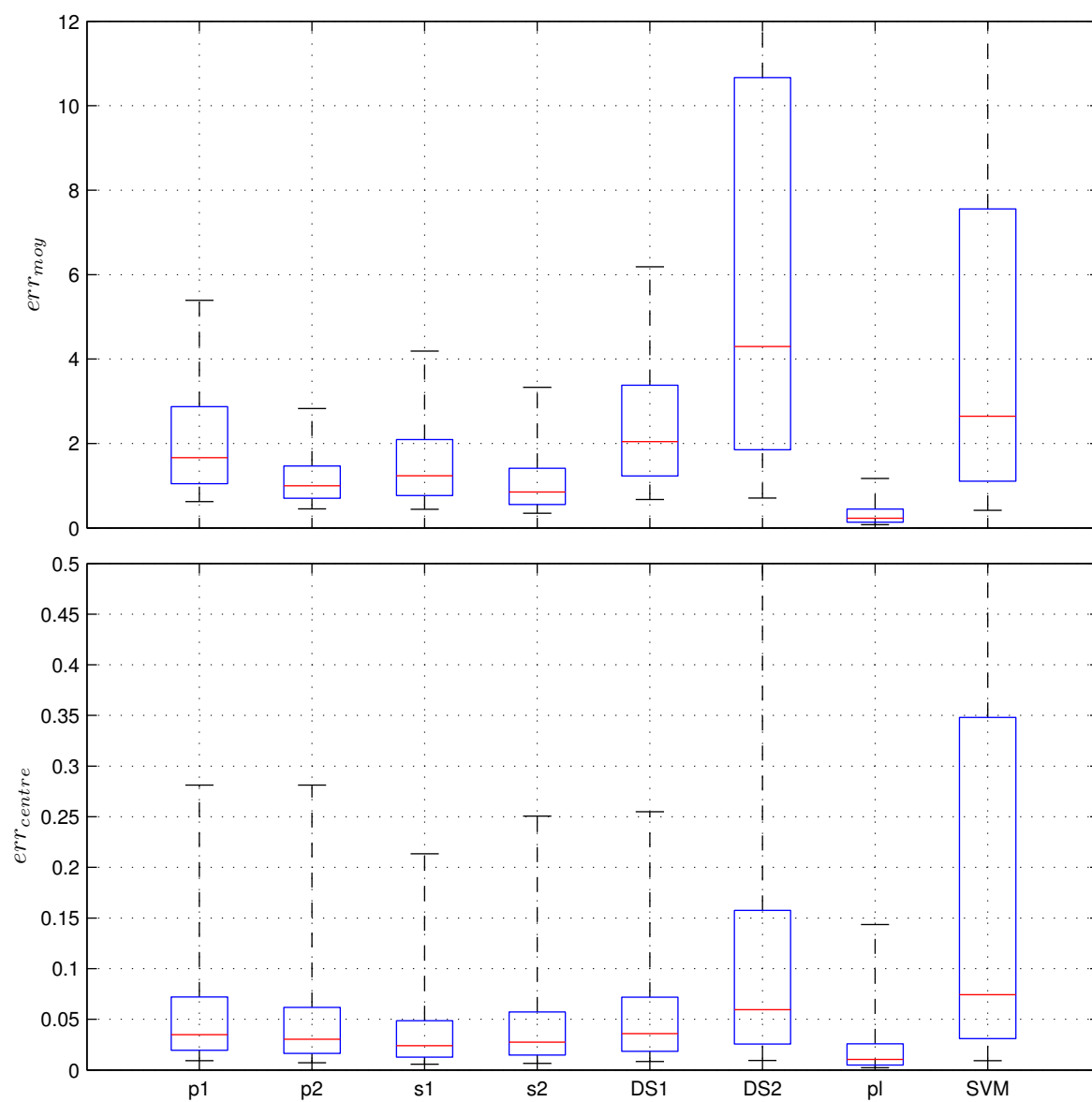


FIGURE 1.12 – Erreurs générées par les différentes méthodes sous formes de boîtes à moustaches. La boîte correspond aux premiers et derniers quartiles et le trait au milieu à la médiane. Les moustaches s'étendent jusqu'aux 5^e et 95^e centiles. En haut, la moyenne quadratique de l'erreur sur toute la durée du signal err_{moy} . En bas, la moyenne quadratique de l'erreur sur la partie centrale du signal err_{centre} . Du fait des différences d'ordres de grandeurs entre les deux, err_{moy} correspond en fait essentiellement à l'erreur sur les bords du signal.

L'erreur étant naturellement plus importante sur les bords, on peut considérer que err_{moy} rend essentiellement compte de l'erreur aux bords, ce qui est confirmé par les ordres de grandeur de err_{moy} et err_{centre} observés en pratique.

Les résultats des simulations sont présentés Fig. 1.12. Les différentes méthodes peuvent être réparties en trois grandes classes en fonction de leurs performances dans la situation de test proposée. Les méthodes **DS2** et **SVM** sont de loin les plus mauvaises, les méthodes **p1**, **p2**, **s1**, **s2** et **DS1** ont des performances moyennes et la méthode **p1** est nettement meilleure que toutes les autres.

Les mauvaises performances de **DS2** et **SVM** ont des origines différentes. **DS2** ne donne des résultats corrects que quand l'intervalle entre deux extrema successifs est stable près du bord. Le signal test ayant été créé spécifiquement pour avoir des intervalles entre extrema successifs variables, la méthode **DS2** est souvent mise en échec. Quant à la méthode **SVM**, ses mauvaises performances ont deux origines. La première est que de manière générale on observe que la méthode arrive bien à prédire un signal oscillant au-delà du bord mais souvent avec une amplitude légèrement différente et décalé verticalement par rapport aux oscillations de l'autre côté du bord. En fait, la prédiction par SVM prédit bien souvent la « bonne » allure de signal mais ne conserve pas forcément des caractéristiques comme l'amplitude d'oscillation qui peuvent avoir leur importance pour prédire les extrema. La deuxième origine des performances mauvaises de la méthode **SVM** est qu'il arrive aussi que le signal prédit au-delà du bord n'ait tout simplement pas d'extremum, ou alors très loin du bord. Dans ces cas, l'erreur finale est beaucoup plus importante, ce qui explique que 5% des cas aboutissent à une erreur finale err_{moy} de l'ordre de 10^5 ou plus.

Parmi les méthodes aux performances moyennes, on observe que la méthode **DS1** est généralement un peu moins bonne que les autres. Ce résultat est difficile à interpréter dans la mesure où **DS1** est conceptuellement assez proche de **p1** et que les cas où $2(t_{max}(2) - t_{max}(1)) < t_{max}(1)$ pour lesquels on pourrait penser que **DS1** va poser problème ne font même pas apparaître de différences de performances par rapport aux autres. La seule explication plausible semble être que le fait de ne pas contraindre explicitement les enveloppes sur le bord comme le fait **p1** permet de plus grandes fluctuations. En dehors de **DS1**, il est intéressant de comparer les résultats obtenus par les méthodes **p1**, **p2**, **s1** et **s2**. En particulier, il semble que le fait de prendre en compte la valeur au bord dans les versions 2 des deux méthodes ait une influence directe sur l'amplitude de l'erreur au niveau des bords. En revanche, la différence de performances entre les méthodes de projection **p1** et **p2** et de symétrie **s1** et **s2** semble tout aussi difficile à expliquer que précédemment pour les méthodes **DS1** et **p1**.

La méthode **p1** donne des résultats nettement meilleurs que les autres dans la situation de test proposée. Ce résultat est encourageant mais il doit être nuancé. En effet, il n'est pas difficile de voir que la situation de test proposée est clairement plus favorable à une méthode de prolongement linéaire qu'à une méthode de symétrie ou de projection. En pratique, dans des situations plus générales, il semble en revanche qu'une telle méthode est susceptible de produire de plus grandes déviations aux bords que par exemple la méthode **s2**. Par conséquent, on peut penser que la méthode **p1** est peut-être moins robuste que **s2** qui a fait ses preuves, puisqu'elle est implantée depuis 2003 dans le programme que nous mettons à disposition sur internet. Les performances observées ici sont cependant encourageantes et demandent à être confirmées plus largement.

Enfin, on s'est attaché pour cette étude à observer à la fois l'amplitude de l'erreur au niveau des bords du signal et la propagation de cette dernière vers le centre du signal. Un aspect qu'on n'a pas étudié et qui est sans doute important est aussi la propagation de l'erreur au fur et à mesure qu'on extrait des IMFs. En pratique, cet aspect peut se révéler assez délicat à étudier dans la mesure où plus on augmente le nombre d'IMFs extraits, plus on réduit la probabilité de retrouver comme résidu la tendance lente qu'on avait mise initialement dans le signal. Par conséquent l'erreur mesurée par (1.81), où $a_1(t)$ serait remplacé par le résidu, a plus de chances de ne pas être due uniquement aux effets de bords, ce qui complique la tâche.

5.2 Arrêt du processus de tamisage

Différentes approches ont été proposées pour déterminer quand arrêter le processus de tamisage. De manière générale, elles remplissent toutes deux objectifs

1. lors de l'arrêt du processus de tamisage, le signal en cours de traitement doit vérifier la définition d'un IMF Déf. 3.
2. le tamisage ne doit pas être itéré un trop grand nombre de fois au risque de dénaturer l'information contenue dans les IMFs. (cf 4.1.2)

5.2.1 Approche originale

Essentiellement focalisée sur le premier objectif, l'approche proposée dans la contribution d'origine [28] consiste à arrêter le processus de tamisage dès que

1. tous les maxima locaux sont strictement positifs et tous les minima locaux strictement négatifs
2. la différence entre l'IMF en cours et sa version à l'itération précédente mesurée par

$$SD = \int_0^T \left(\frac{d_1^{(n)}(t) - d_1^{(n-1)}(t)}{d_1^{(n-1)}(t)} \right)^2 dt \quad (1.84)$$

doit être inférieure à un seuil, par exemple de l'ordre de 0.2-0.3 pour un signal de 1024 points. Le problème de cette approche est essentiellement que le critère évalué pour mesurer la différence entre l'IMF en cours et sa version à l'itération précédente ne remplit pas la fonction requise. En effet, l'intégrande diverge généralement aux points où le dénominateur s'annule ce qui fait que la valeur du critère SD est plus contrôlée par le comportement de $d_1^{(n-1)}(t)$ au voisinage des ses passages à zéro que par la valeur du numérateur. Il n'est pas rare d'observer avec ce critère des comportements aberrants où SD est très grand ($10^5 \cdot \text{seuil}$) alors que la moyenne du signal peut très raisonnablement être considérée comme nulle. En corrigeant le raisonnement derrière la définition de SD , on aboutit à la définition

$$SD = \frac{\int_0^T (d_1^{(n)}(t) - d_1^{(n-1)}(t))^2 dt}{\int_0^T (d_1^{(n-1)}(t))^2 dt}, \quad (1.85)$$

également corrigée par les auteurs dans [26, 64]. Cette définition remplit mieux l'objectif recherché mais elle a le défaut d'être globale. De fait, pour une même valeur de SD , la moyenne peut s'écarter fortement de zéro très localement ou s'en écarter beaucoup moins mais partout à la fois. En pratique, les deux situations ne sont pas forcément identiques : on peut par exemple tolérer une forte erreur locale si elle s'accompagne d'une faible erreur partout ailleurs plus facilement qu'une erreur relativement moins forte mais répartie.

Du point de vue de l'interprétation, on constate également que cette définition compare de fait l'amplitude de la moyenne de $d_1^{(n-1)}(t)$ (au numérateur) à sa propre amplitude (au dénominateur), pour savoir si on arrête le processus de tamisage à $d_1^{(n)}(t)$. De fait, on ne vérifie donc pas vraiment que la moyenne de l'IMF final est faible mais plutôt que la moyenne de sa version à l'itération précédente l'est. On n'est donc pas vraiment dans le cadre de la vérification de la définition d'un IMF comme on aurait pu le penser mais plutôt dans celui d'un critère de convergence de type Cauchy.

5.2.2 Approche « locale »

Une deuxième approche d'inspiration voisine a été proposée dans [55] et est implantée dans les programmes que nous mettons à disposition sur internet. L'idée est ici de s'appuyer sur la définition pratique d'un IMF élaborée en 4.1.3. Le processus de tamisage est arrêté dès que

1. tous les maxima locaux sont strictement positifs et tous les minima locaux strictement négatifs
2. les enveloppes de l'IMF en cours vérifient

$$|U(t) + L(t)| \leq \epsilon |U(t) - L(t)|$$

sur un ensemble de mesure supérieure à $(1 - \alpha)T$, avec T la durée totale du signal et

$$|U(t) + L(t)| \leq \epsilon_2 |U(t) - L(t)|$$

pour tout t .

Contrairement à l'approche globale proposée précédemment, cette définition est conçue pour contrôler le plus localement possible la qualité de l'IMF en construction. Idéalement, on voudrait pouvoir imposer partout la même contrainte $|U(t) + L(t)| \leq \epsilon |U(t) - L(t)|$, mais on s'est aperçu en pratique que la quantité $|U(t) + L(t)|/|U(t) - L(t)|$ avait bien souvent tendance à présenter des pics très localisés alors que sa valeur ailleurs est relativement proche de 0. Sachant qu'imposer une condition $|U(t) + L(t)| \leq \epsilon |U(t) - L(t)|$ partout avec un paramètre ϵ faible de l'ordre de 0.1 avait tendance pour cette raison à aboutir à des nombres d'itérations importants, on a relâché un peu la contrainte en l'imposant par exemple sur seulement 95% de la durée du signal alors qu'ailleurs on impose une contrainte beaucoup moins forte avec par exemple $\epsilon_2 = 0.5$.⁴

Le simple fait d'accorder cette tolérance a permis de réduire fortement les nombres d'itérations en général. Cependant, il arrive pour certains signaux que cette tolérance ne soit pas suffisante et on observe alors de grands nombres d'itérations alors qu'une analyse plus fine montre que la moyenne peut être considérée comme raisonnablement nulle sur une très large partie du signal pour des nombres d'itérations moins élevés. Bien entendu, la notion d'IMF acceptable est très subjective et en pratique liée à une application donnée. Il est donc illusoire de vouloir trouver un critère d'arrêt qui convienne à toutes les situations pratiques. Néanmoins, puisqu'il n'est que rarement possible d'avoir un IMF dont la moyenne est suffisamment faible partout, la plupart des situations peuvent se ramener à un compromis entre une moyenne suffisamment nulle presque partout et un ensemble de points sur lequel une moyenne plus importante est tolérée. Malheureusement, la diversité des situations et des signaux fait qu'il est nécessaire d'ajuster les paramètres d'un tel compromis à chaque cas.

Par ailleurs, on a relâché pour ce critère la condition $|U(t) + L(t)| \leq \epsilon |U(t) - L(t)|$ évaluant la deuxième clause de la définition d'un IMF (Déf. 3) mais on a laissée intacte la première clause parce qu'elle est en général moins contraignante que la seconde. Ce résultat est en fait uniquement valable pour des signaux pas trop longs. Pour des signaux ayant plus de 100000 extrema, il est possible que la première clause devienne plus restrictive que la seconde et qu'elle cause à elle seule de très grands nombres d'itérations. En pratique, ce cas n'a pas encore été vraiment étudié pour la bonne raison qu'il n'y a pas à l'heure actuelle de programme permettant de réaliser une EMD sur des signaux aussi longs en des temps courts. Dans les rares cas approchés au cours de cette thèse, tolérer que 0.01% des extrema puissent ne pas satisfaire à la première clause a toujours permis de ramener les nombres d'itérations à des ordres de grandeur raisonnables.

5.2.3 Approche « robuste »

Plus axée sur l'objectif de ne pas itérer le processus de tamisage plus que nécessaire, les auteurs de la contribution originale ont aussi proposé une autre méthode pour déterminer les nombres d'itérations qui ne vérifie explicitement que la première clause de la définition d'un IMF [29]. Cette méthode consiste à arrêter le processus de tamisage quand pour un nombre donné S d'itérations successives,

- tous les maxima sont positifs et tous les minima négatifs

4. En pratique, on constate bien souvent que la valeur de ϵ_2 choisie n'influe pas vraiment sur les nombres d'itérations, du moins tant qu'elle n'est pas trop faible. Dans beaucoup de cas, la contrainte $|U(t) + L(t)| \leq \epsilon_2 |U(t) - L(t)|$, sert donc plus à assurer un certain niveau de nullité de la moyenne pour l'utilisateur qu'à déterminer les nombres d'itérations.

- le nombre d’extrema ne change pas

Étant donnée la difficulté qu’il peut y avoir à évaluer le degré de nullité de la moyenne d’un IMF, cette méthode propose tout simplement de ne pas le mesurer mais de supposer que S itérations de tamisage sont suffisantes pour que la moyenne soit suffisamment nulle. De fait, on observe de manière générale que lors d’une itération de tamisage, la moyenne après l’itération est presque toujours plus faible qu’avant sauf aux points où de nouveaux extrema sont apparus, ce qui justifie le point de vue selon lequel une itération de tamisage durant laquelle aucun extremum n’apparaît a pour effet de ramener la moyenne vers zéro sur toute la durée du signal.

En pratique, on observe que de faibles valeurs de S , entre 3 et 5 comme le suggèrent les auteurs, permettent très souvent d’arriver à des nombres d’itérations raisonnables. De plus, toujours dans l’objectif d’éviter de trop itérer le processus de tamisage, les auteurs recommandent aussi de fixer une limite aux nombres d’itérations pour faire face à toute éventualité. En revanche, contrairement aux méthodes évoquées précédemment, l’utilisateur n’a pas de contrôle explicite sur les caractéristiques des IMFs obtenus.

Enfin, on peut noter également que ce critère peut poser des problèmes s’il est utilisé pour traiter des signaux de durée variable. En effet, pour un signal stationnaire, si p est la probabilité qu’un extremum apparaisse à l’itération k du processus de tamisage pour un signal de durée T , elle sera entre p et $2p$ pour un signal de durée $2T$. Par conséquent, le nombre d’itérations déterminé par cette méthode pour extraire le premier IMF du signal de durée $2T$ peut être supérieur à celui utilisé pour le signal de durée T avec la même valeur de S dans les deux cas alors que les contenus des deux signaux sont identiques.

5.2.4 Approche « fixe »

Pour terminer, une dernière approche consiste simplement à déterminer les nombres d’itérations a priori sans se préoccuper du fait que les IMFs obtenus vérifient la définition. Cette approche est a priori à éviter de manière générale pour traiter des signaux isolés mais elle peut être intéressante pour traiter des lots de données puisqu’elle garantit une certaine forme d’égalité de traitement qui peut se révéler intéressante pour comparer par la suite les résultats. C’est en particulier le cas dans le cadre de l’une des variantes de l’EMD dite EMD d’ensemble (EEMD pour « Ensemble EMD ») qui sera décrite en 6.3.

5.3 Interpolation

Pour passer des ensembles de maxima et minima locaux du signal à ses enveloppes supérieure et inférieure, on utilise généralement une technique d’interpolation. Comme on l’a observé en 4.2, le choix de la technique d’interpolation utilisée n’est pas anodin dans la mesure où les IMFs, en dehors du premier, sont tous des sommes d’un plus ou moins grand nombre d’interpolations. De plus, comme on le montrera en 1.1, le comportement de l’algorithme en fonction du nombre d’itérations de tamisage dépend très fortement du schéma d’interpolation utilisé.

Pour approcher la problématique de l’interpolation, il convient pour commencer de préciser les objectifs recherchés. Dans le contexte de l’EMD, il y en a essentiellement deux :

localité : L’un des objectifs de l’EMD étant de proposer une méthode permettant de traiter des signaux non stationnaires, il faut qu’elle soit capable de s’adapter aux évolutions du signal. Dans ce but, il est nécessaire que la valeur d’une enveloppe en un instant donné dépende essentiellement des évolutions du signal à proximité de cet instant et moins, voire pas du tout, d’instant plus éloignés où le signal est susceptible d’évoluer très différemment. Ainsi, des interpolations particulièrement non locales telle l’interpolation de Lagrange, pour prendre un exemple extrême, sont totalement inadaptées. À l’inverse, il est a priori souhaitable que la valeur de l’interpolation en un instant donné dépende du signal à horizon fini, c’est-à-dire

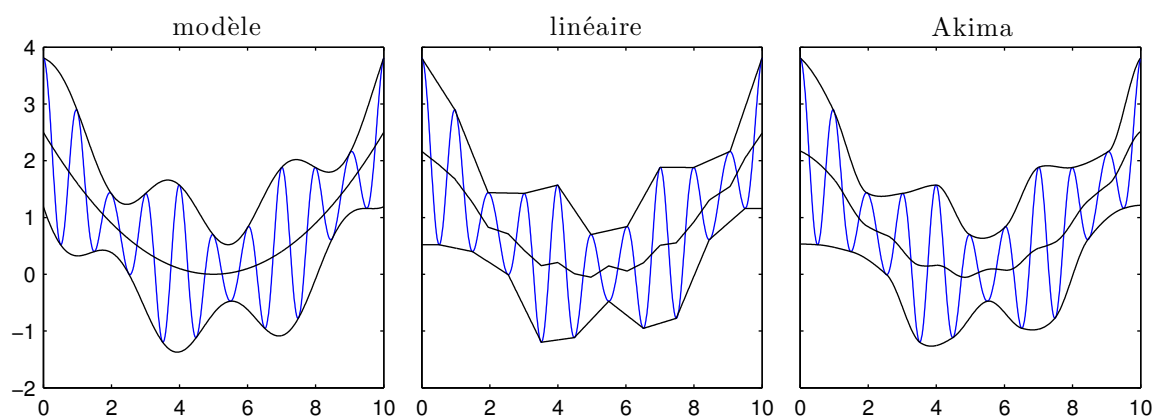


FIGURE 1.13 – Sur l'exemple de signal proposé, il est clair que la moyenne des enveloppes doit présenter au moins un minimum vers le centre du signal, mais les ondulations introduites par les interpolations linéaire par morceaux et Akima sont plus discutables.

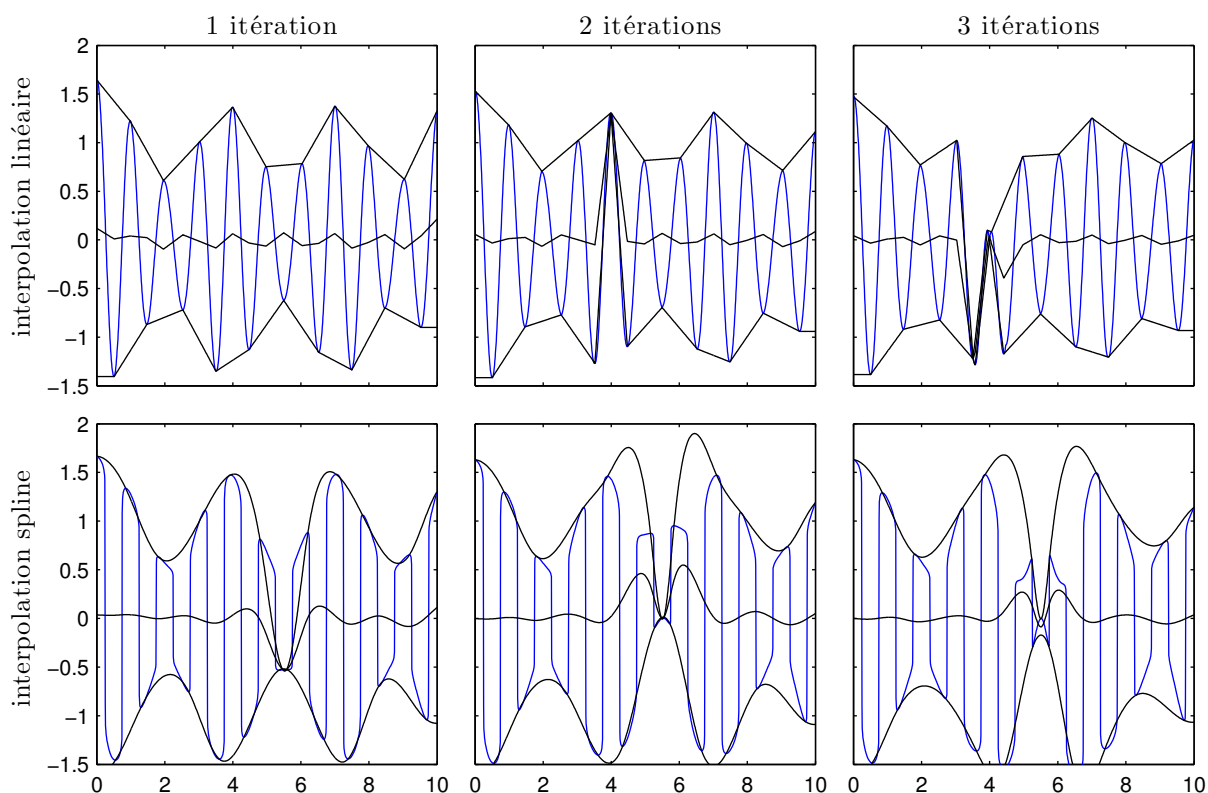


FIGURE 1.14 – Le manque de régularité de l'interpolation linéaire par morceaux peut être la cause de l'apparition d'extrema parasites causant de plus des problèmes de convergence. Le problème persiste pour des interpolations plus régulières comme la spline cubique, mais seulement lorsque le signal présente des extrema particulièrement « plats ».

qu'elle ne dépende que d'un nombre limité de points d'interpolation, ou nœuds, de part et d'autre de l'instant considéré. Au minimum, il faut que l'influence d'un nœud sur la valeur de l'interpolation en un instant donné décroisse en fonction de sa distance (exprimée, par exemple, en termes de nombre de nœuds) à l'instant considéré et tende vers zéro quand celle-ci tend vers l'infini. En pratique, la plupart des schémas d'interpolation usuels sont dans ce dernier cas mais la manière dont l'influence d'un nœud décroît en fonction de la distance peut varier selon le schéma.

régularité : L'origine de cet objectif est un peu plus subtile. Elle apparaît plus clairement si on prend en compte le fait que les approximations successives du signal par EMD sont toutes des sommes d'interpolations. Si on calcule par exemple le premier IMF avec différents schémas d'interpolation, on obtient en le soustrayant au signal différentes premières approximations qui en particulier ont des nombres d'extrema différents. Parmi ces extrema, certains sont clairement la trace d'échelles caractéristiques du signal analysé mais d'autres sont en fait introduits artificiellement par l'interpolation. Pour s'en convaincre, il suffit d'imaginer une interpolation particulièrement mauvaise qui produirait des ondulations entre chaque nœud : une somme de telles interpolations pourrait alors conserver un certain nombre de ces ondulations qui, si des extrema y sont associés, pourraient alors être perçues par l'EMD comme des échelles de temps caractéristiques du signal. Par conséquent, il est a priori souhaitable que l'interpolation n'ondule pas trop : au plus un extremum *ou* une inflexion entre deux nœuds semble un objectif raisonnable. De plus, si l'interpolation présente des extrema, il faut autant que possible que ces extrema ne soient pas trop piqués. Ainsi, sur l'exemple proposé Fig. 1.13, il est clair que la moyenne des enveloppes doit présenter au moins un minimum vers le centre du signal, mais la pertinence des extrema supplémentaires présents lorsque l'interpolation est linéaire par morceaux ou Akima (dérivée première continue seulement) est quant à elle plutôt discutable, puisqu'ils résultent directement des extrema des enveloppes. De plus, le manque de régularité de l'interpolation linéaire par morceaux pose même des problèmes de convergence pour l'EMD dans la mesure où il est fréquent de voir apparaître de nouveaux extrema très douteux lorsque la moyenne des enveloppes est soustraite au signal. Plus généralement, le problème persiste pour des schémas d'interpolation plus réguliers mais il ne se manifeste que quand les extrema du signal sont particulièrement « plats », c'est-à-dire que les dérivées successives au voisinage des extrema sont particulièrement faibles (cf Fig. 1.14).

À ces deux objectifs principaux, on peut ajouter un troisième qui est que les enveloppes respectent bien la notion intuitive d'enveloppe, c'est-à-dire que le signal doit toujours se trouver entre l'enveloppe inférieure et l'enveloppe supérieure. En pratique, on observe que les enveloppes calculées par l'interpolation spline cubique présentent souvent de petits dépassements où le signal se trouve localement du mauvais côté de l'enveloppe. Dans des cas plus extrêmes, il arrive même que l'enveloppe supérieure soit localement inférieure à l'enveloppe inférieure. Pour cette raison, les recherches effectuées sur la question de l'interpolation cherchent généralement à atteindre l'objectif de supprimer ces dépassements ou du moins de les limiter [15, 47, 36, 35]. Toutefois, comme précisé dans la contribution originale [28], l'effet de ces dépassements est indirect puisque c'est la moyenne des enveloppes qui intervient finalement dans le processus de tamisage et rien ne prouve que ces dépassements soient effectivement problématiques, si ce n'est d'un point de vue esthétique. De plus, ces dépassements ne sont pas nécessairement un défaut propre au schéma d'interpolation, mais à la procédure de calcul des enveloppes dans sa globalité. En particulier, le choix des extrema comme points d'interpolation n'est pas forcément idéal de ce point de vue. Si on considère par exemple le cas d'un signal de la forme $x(t) = \alpha t + \cos t$ avec $\alpha \neq 0$, il est clair a priori que ses enveloppes doivent être rectilignes mais il est aisé de voir que si celles-ci passent par les extrema, alors il y a nécessairement des dépassements au voisinage de ces derniers. Pour ces raisons, on n'insistera pas sur cet objectif pour comparer différents schémas d'interpolation.

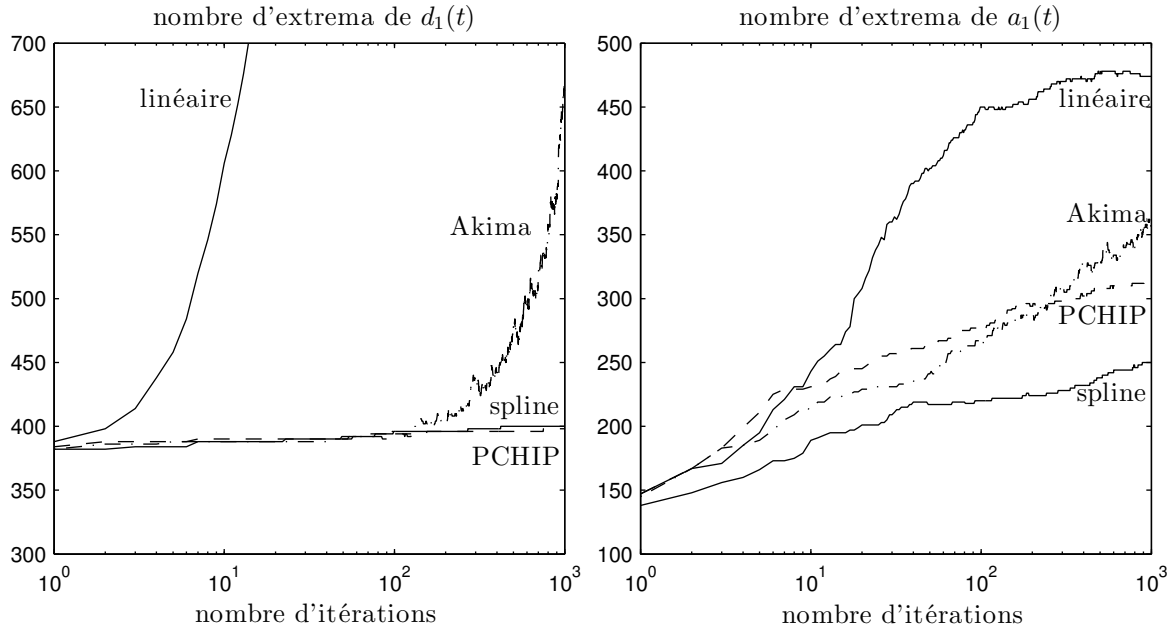


FIGURE 1.15 – Nombres d'extrema du premier IMF $d_1(t)$ et de la première approximation $a_1(t)$ en fonction du nombre d'itérations de tamisage pour différents schémas d'interpolation. Le signal utilisé pour cette étude est un bruit blanc gaussien suréchantillonné d'un facteur 8 à l'aide d'une interpolation spline cubique (on obtiendrait des résultats similaires avec l'interpolation PCHIP par exemple). Le fait d'utiliser un signal suréchantillonné est important pour pouvoir observer les instabilités générées par les interpolations linéaires par morceaux et Akima.

Si on cherche alors un schéma d'interpolation qui satisfasse les deux exigences de localité et de régularité, on se heurte au fait que ces deux exigences sont en fait contradictoires : pour obtenir une interpolation régulière, il est nécessaire que sa valeur sur un intervalle entre deux extrema dépende au moins de son évolution sur les intervalles voisins. Le plus souvent, cette dépendance aboutit de proche en proche au fait que la valeur de l'interpolation en un point dépend de tous les points d'interpolation mais, heureusement, la dépendance décroît avec la distance et tend vers zéro quand celle-ci tend vers l'infini. Cependant, la vitesse de décroissance de la dépendance en fonction de la distance peut varier sensiblement d'un schéma d'interpolation à l'autre. Ainsi, si on prend l'exemple de la famille des splines polynomiales, il n'y a pas de dépendances entre intervalles voisins pour les splines de degré 0 et 1 (interpolation constante et linéaire par morceaux respectivement) ; aux ordres supérieurs, la dépendance décroît exponentiellement, croît avec le degré et tend vers $1/n$, n étant la distance exprimée en termes de nombre de nœuds, quand le degré tend vers l'infini [62]. Parallèlement, la régularité des splines polynomiales augmente avec le degré, une spline polynomiale de degré n étant par définition un polynôme par morceaux de degré au plus n dont les dérivées successives sont continues jusqu'à l'ordre $n - 1$.

En pratique, l'interpolation la plus utilisée pour l'EMD est la spline cubique. Celle-ci n'est clairement pas optimale du point de vue de la localité, mais elle semble en revanche particulièrement bonne du point de vue de la régularité. Pour s'en apercevoir, on peut par exemple s'intéresser à l'évolution des nombres d'extrema du premier IMF et de la première approximation en fonction du nombre d'itérations de tamisage utilisées pour les calculer. On observe ainsi Fig. 1.15 que, comparée à d'autres interpolations cubiques populaires, l'interpolation Akima [2] et l'interpolation PCHIP [22], la première approximation calculée à l'aide de l'interpolation spline cubique contient nettement moins d'extrema que les autres. Par ailleurs, on observe aussi que les enveloppes calculées par l'interpola-

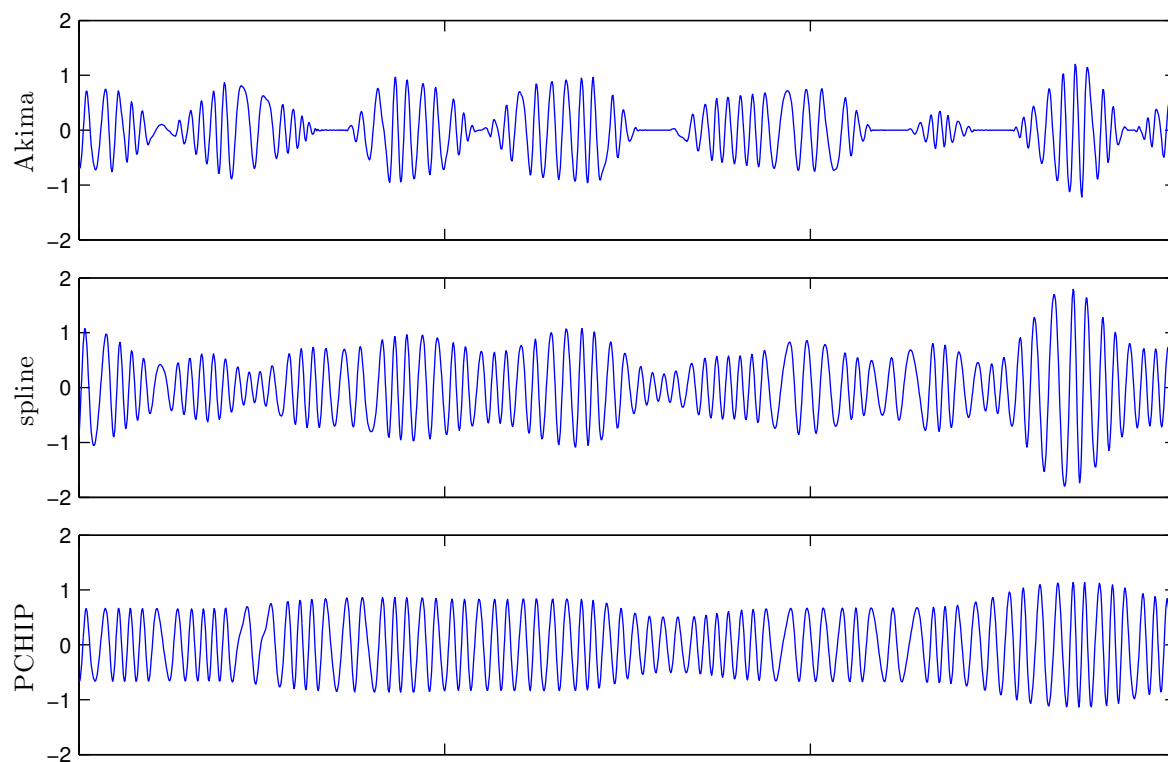


FIGURE 1.16 – Premier IMF après 1000 itérations de tamisage pour différents schémas d'interpolation. Le signal utilisé pour calculer ces IMFs est le même que celui utilisé pour la Fig. 1.15. Le nombre de 1000 itérations est très grand mais permet de bien voir les différences qualitatives entre les différentes interpolations. En particulier le fait que les enveloppes de l'IMF soient plus lisses avec l'interpolation PCHIP qu'avec l'interpolation spline est valable pour tous les nombres d'itérations.

tion spline cubique présentent en général plus d'extrema que celles calculées par les interpolations Akima et PCHIP. Sur l'exemple Fig. 1.15, les enveloppes supérieures/inférieures présentent pour l'interpolation spline cubique, pour l'interpolation Akima et pour l'interpolation PCHIP respectivement 128/134, 126/122 et 123/121 extrema et leurs moyennes respectivement 138, 145 et 147 extrema. Ceci montre que le nombre d'extrema des enveloppes n'est pas nécessairement un critère pertinent. En effet, l'interpolation spline cubique aboutit certes à des enveloppes avec plus d'extrema que les autres mais la moyenne des enveloppes en contient moins ! L'origine de ce résultat est apparemment le fait que les extrema dans les enveloppes splines cubiques sont généralement moins piqués (la dérivée seconde est localement plus faible) que pour les autres interpolations. Cette propriété est d'ailleurs sans doute en lien avec la propriété de « minimum de courbure » de l'interpolation spline cubique naturelle (cas où les dérivées secondes aux extrémités sont nulles), selon laquelle elle constitue l'interpolation dont la norme L_2 de la dérivée seconde est minimale.

Par ailleurs, l'évolution du nombre d'extrema du premier IMF permet d'observer que l'interpolation Akima n'est pas suffisamment régulière pour l'EMD. En effet, le fait que le nombre d'extrema augmente sans cesse indique une instabilité qu'on peut d'ailleurs observer très directement en regardant l'allure du premier IMF : celui-ci présente de nombreux intervalles où l'amplitude est beaucoup plus faible qu'ailleurs et où l'IMF ressemble à un bruit discret haute fréquence. En revanche, les évolutions des nombres d'extrema du premier IMF calculé avec l'interpolation spline cubique et PCHIP sont très proches, ce qui suggère que ces deux méthodes sont suffisamment régulières pour ne pas introduire (trop) d'extrema parasites. Cette conclusion est d'ailleurs cohérente avec le fait que les nombres d'extrema des approximations calculées par ces deux interpolations croissent relativement lentement par rapport aux autres et semblent presque se stabiliser à partir de 100-200 itérations. Enfin, si on s'intéresse à l'allure des IMFs produits, on observe que la méthode PCHIP, plus locale et moins régulière que la spline cubique, conduit paradoxalement à des IMFs aux enveloppes plus régulières. Si on analyse ce résultat en perspective avec la modélisation proposée en 1.1 du chapitre chapitre 2, tout semble indiquer que l'EMD utilisant l'interpolation PCHIP pourrait être modélisée par un modèle similaire à celui proposé pour l'interpolation spline cubique dans lequel le filtre équivalent pour un nombre d'itérations donné aurait une fréquence de coupure plus élevée. Malheureusement, l'interpolation PCHIP ne pouvant être mise sous la forme d'un filtrage linéaire quand les nœuds sont uniformément espacés, le modèle développé en 1.1 du chapitre 2 ne peut être adapté à cette interpolation.

En dehors des schémas d'interpolation classiques, on peut citer aussi un certain nombre de travaux réalisés sur la question de l'interpolation qui proposent des méthodes originales optimisées pour l'EMD :

- Dans [47], les auteurs proposent une méthode nouvelle adaptée de la méthode dite « parabolic parameter spline interpolation » communément utilisée pour interpoler des ensembles quelconques de points du plan. Selon l'article, la nouvelle interpolation permet de réduire les dépassements observés avec l'interpolation spline cubique tout en gardant un bon niveau de régularité contrairement à une interpolation comme Akima. Cette amélioration s'accompagne également d'améliorations sur le spectre de Hilbert–Huang qui présente moins d'artefacts, selon les auteurs. En contrepartie, la méthode n'est pas encore totalement opérationnelle dans la mesure où elle reste tributaire d'un paramètre de réglage dont la valeur n'est pas établie de manière définitive.
- Dans une perspective très différente, il a été proposé de calculer les enveloppes à l'aide d'une approche par équations aux dérivées partielles [15]. Contrairement au cas précédent, l'idée n'est pas tant de perfectionner l'EMD mais plutôt d'essayer de parvenir à une formulation qui se prête mieux à une analyse théorique que la formulation standard. En pratique, les enveloppes sont calculées à l'aide d'équations aux dérivées partielles de la forme

$$s_\theta = -g(t)s_{tttt} \quad (1.86)$$

où les indices correspondent aux dérivées partielles et θ est une variable supplémentaire telle que $s(t, \theta = 0)$ corresponde à une condition initiale pour le problème d'estimation de l'enveloppe et $s(t, \theta = \infty)$ à l'estimation finale. Du fait de la dérivée quatrième dans le second membre, si $g(t)$ est nulle pour un ensemble discret de points $\{t_i, i \in I\}$, alors la solution asymptotique $s(t, \theta = \infty)$ est un polynôme cubique par morceaux sur les intervalles entre les t_i . Cette caractéristique peut être un avantage dans la mesure où un polynôme par morceaux peut être décrit par un nombre restreint de caractéristiques. En revanche, la forme du second membre en dérivée quatrième peut aussi être pénalisante si la condition initiale, proposée comme égale au signal dans l'article, est malencontreusement un polynôme cubique sur un intervalle entre deux extrema. Par ailleurs, le choix de la valeur de $g(t)$ en dehors des points de contrôle (maxima ou minima suivant le cas) permet d'ajuster le comportement de l'interpolation. Les auteurs proposent un choix pour la fonction $g(t)$ qui prend en compte localement les signes des dérivées premières et secondes du signal. L'interpolation résultante présente apparemment de très bonnes caractéristiques. En outre, on peut observer que le choix de la fonction $g(t)$ proposée fait que l'interpolation utilise de fait plus d'informations que les simples positions des maxima ou minima, ce qui n'est peut-être pas étranger aux bonnes performances observées. De manière plus générale, ce constat amène la question de la quantité d'information nécessaire pour produire des enveloppes satisfaisantes. À ce propos, il est clair que s'il existe une méthode optimale (le problème d'optimisation restant à préciser) pour calculer les enveloppes, celle-ci utilise à coup sûr plus d'informations que les simples positions des extrema. De plus, les enveloppes optimales ne passent même très probablement pas par ces derniers. En revanche, le fait de se limiter aux extrema permet de réduire considérablement la complexité de la tâche à accomplir tout en fournissant des résultats acceptables en général.

- Une dernière approche très récente [36] prend le problème de manière différente : au lieu de proposer un schéma d'interpolation basé sur un certain nombre de critères préétablis (régularité, localité, non dépassement,...), les auteurs considèrent initialement un modèle de signal de type somme de composantes AM-FM

$$x(t) = \sum_{i=1}^N a(t) \cos(2\pi\varphi_i(t)), \quad (1.87)$$

et se fixent comme objectif d'extraire la composante dont la fréquence instantanée ($d\varphi_i/dt$) est la plus grande. Dans ce but, ils réalisent des optimisations successives, à l'aide d'algorithmes génétiques, pour dans un premier temps optimiser le choix des points d'interpolation et par la suite optimiser un polynôme par morceaux s'appuyant sur ces points. Les résultats de leur étude sur les points d'interpolation suggèrent que les extrema du signal ne sont pas les points optimaux et qu'ils sont avantageusement remplacés par les extrema de la composante qu'on souhaite extraire, c'est-à-dire celle dont la fréquence instantanée est localement la plus grande. En pratique, ce constat ne résoud évidemment rien puisque la composante de plus haute fréquence instantanée n'est généralement pas connue initialement. Face à ce problème, les auteurs ont suggéré plus récemment [35] une méthode permettant d'estimer ces points au sein même du processus de tamisage qui de plus reste philosophiquement proche de l'EMD originale. La méthode semble encore limitée dans la mesure où il faut un grand nombre d'itérations de tamisage pour que ses effets apparaissent mais elle ouvre une piste prometteuse. Enfin, les conclusions de la deuxième partie de l'étude sur l'interpolation optimale [36] sont moins claires : des interpolations de type Hermite d'ordre élevé (13) semblent être favorisées mais leurs performances chutent fortement dès que les points interpolés ne sont pas les points établis comme optimaux précédemment.

Comme le montrent ces études, la question de l'interpolation et plus généralement de la manière de définir les enveloppes est cruciale pour l'EMD, dans la mesure où ses propriétés dépendent fortement

de la technique employée. De plus, comme on le mettra en évidence dans la partie 1.1 du chapitre 2, les propriétés de résolution fréquentielle de l'EMD évoluent lorsqu'on itère le processus de tamisage, et la manière dont elles évoluent dépend également de cette technique. Au cours de cette thèse, on se restreindra à l'interpolation spline cubique parce que c'est la plus utilisée en pratique, l'objectif étant plus d'améliorer la compréhension de l'EMD que de perfectionner la technique.

5.4 Problèmes liés à l'échantillonnage

L'EMD est pratiquement toujours présentée dans un contexte de temps-continu, c'est-à-dire que les signaux sont généralement considérés comme des fonctions d'une variable de temps continue. En pratique cependant, l'EMD est presque exclusivement implantée en temps discret, c'est à dire que les signaux analysés sont échantillonnés ainsi que les IMFs obtenus en sortie. Il en découle que lorsqu'on veut réaliser l'EMD d'un signal défini à temps continu, on n'a d'autre choix que de le discrétiser avant de lui appliquer l'EMD. Or, si aucune précaution n'est prise lors de la discrétisation, celle-ci a un effet a priori non nul sur la sortie de l'algorithme. Dans un premier temps, on s'intéresse au cas très simple d'une sinusoïde et on étudie comment l'échantillonnage modifie son EMD. Par la suite, on propose une modélisation plus poussée des effets d'échantillonnage permettant d'aboutir dans un cadre plus général à une borne sur l'amplitude des effets d'échantillonnage pour une itération de tamisage.

5.4.1 Effets d'échantillonnage sur l'EMD d'une sinusoïde

Soit $x(t)$ le signal à temps continu⁵ :

$$x(t) = \cos 2\pi t, \quad t \in \mathbb{R} \quad (1.88)$$

$x(t)$ est clairement un IMF. Dans un contexte de temps continu, on s'attend donc à ce que l'EMD de $x(t)$ ne fournisse qu'un seul IMF égal à $x(t)$ et pas de résidu. En pratique, il se trouve que si on discrétise $x(t)$ et qu'on lui applique l'EMD à temps discret, il est rare qu'on obtienne effectivement un seul IMF égal au signal discrétisé. À l'origine, ceci provient du fait que lorsqu'on discrétise $x(t)$, on perd généralement la position et la valeur du signal à ses extrema et les positions et valeurs des extrema du signal discret sont donc généralement différentes de celles du signal à temps continu.

5.4.1.1 Simulations Pour évaluer l'influence de l'échantillonnage dans le cas d'une sinusoïde, on introduit les versions échantillonnées de $x(t)$:

$$x_{f_e, \varphi}[n] = x\left(\frac{n}{f_e} + \varphi\right) = \cos\left(\frac{2\pi}{f_e}n + 2\pi\varphi\right), \quad n \in \mathbb{Z}, \quad (1.89)$$

avec $f_e \geq 2$ la fréquence d'échantillonnage et φ la phase associée. Étant donné que $x(t)$ est un IMF, on peut évaluer l'influence de l'échantillonnage par l'écart entre $x_{f_e, \varphi}[n]$ et le premier IMF de sa décomposition $d_{f_e, \varphi}[n]$:

$$e(f_e, \varphi) = \lim_{N \rightarrow \infty} \frac{1}{2N+1} \sum_{n=-N}^N |x_{f_e, \varphi}[n] - d_{f_e, \varphi}[n]|. \quad (1.90)$$

Le nombre d'itérations de tamisage utilisé pour calculer l'EMD de $x_{f_e, \varphi}[n]$ n'est pas précisé explicitement ici dans la mesure où les modèles développés par la suite sont valables dès lors que celui-ci est non nul. On utilisera également une version de cet écart moyenné par rapport à la phase, assimilée à une variable aléatoire uniforme sur $[0, 1]$:

$$\bar{e}(f_e) = \mathbb{E}_{\varphi}\{e(f_e, \varphi)\} = \int_0^1 e(f_e, \varphi) d\varphi. \quad (1.91)$$

5. On considère le signal défini sur une durée infinie de manière à ne pas se soucier d'éventuels effets de bords.

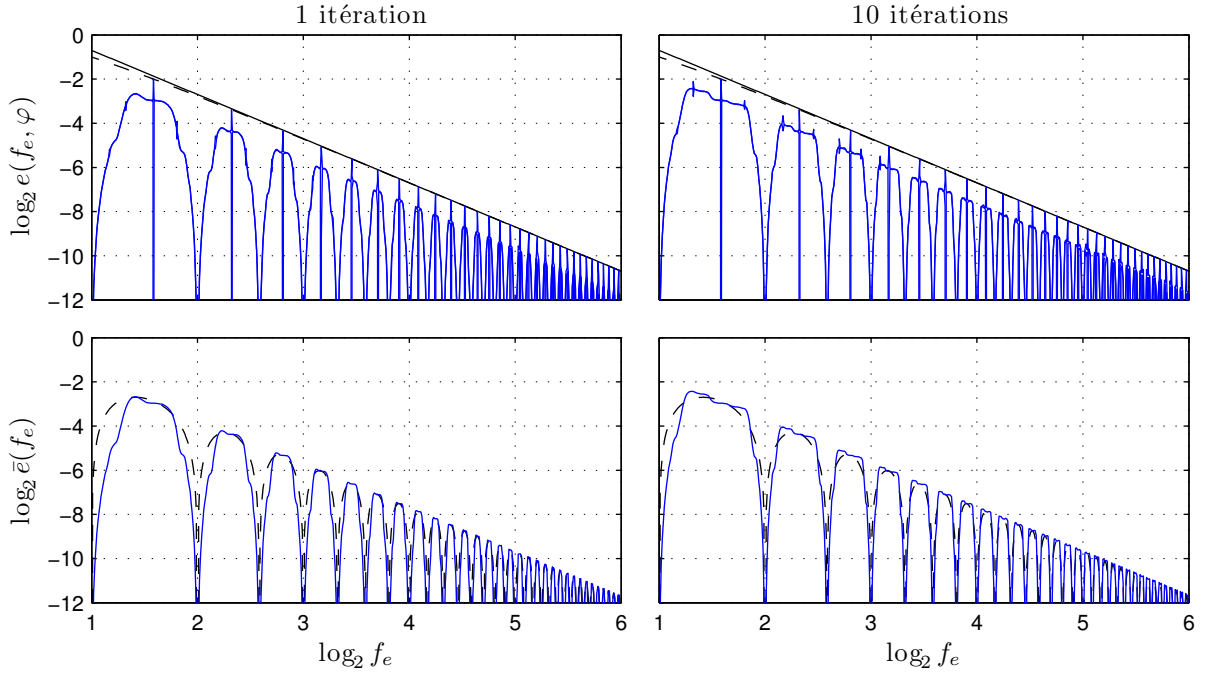


FIGURE 1.17 – Écart entre le premier IMF et le signal sinusoïdal pour 1 (à gauche) et 10 (à droite) itérations de tamisage. En haut, l'écart maximum et minimum (par rapport à la phase) en fonction de f_e . La courbe en trait plein au dessus correspond à la majoration $\pi^2/(4f_e^2)$; la courbe en tirets à la majoration $(1 - \cos(\pi/f_e))/2$. En bas, l'écart moyenné par rapport à la phase $\bar{e}(f_e)$. La courbe en tirets correspond à l'approximation (1.111).

Les évolutions de $e(f_e, \varphi)$ et $\bar{e}(f_e)$ en fonction de f_e sont représentées Fig. 1.17 pour des nombres d'itérations de tamisage fixés à l'avance et Fig. 1.18 pour des nombres d'itérations déterminés au cours du processus de tamisage par le critère « local » (cf 5.2). Ces évolutions présentent quatre caractéristiques remarquables :

1. $e(f_e, \varphi)$ et $\bar{e}(f_e)$ sont bornés par une fonction proportionnelle à f_e^{-2} .
2. $e(f_e, \varphi)$ et $\bar{e}(f_e)$ s'annulent pour $f_e = 2k$ avec $k \in \mathbb{N}$. $e(f_e, \varphi = 1/(2f_e))$ s'annule aussi pour $f_e = 2k + 1$.
3. dans les cas où le nombre d'itérations est fixé à l'avance, les comportements de $e(f_e, \varphi)$ et $\bar{e}(f_e)$ pour $2k \leq f_e \leq 2k + 2$ sont quasi-indépendants de k à renormalisation près.
4. dans les cas où le nombre d'itérations est déterminé au cours du processus de tamisage, l'écart s'annule pour f_e supérieure à une certaine limite.

5.4.1.2 Borne de l'écart On se propose de justifier le résultat suivant :

$$e(f_e, \varphi) \leq \frac{1}{2} \left(1 - \cos \frac{\pi}{f_e} \right) \leq \frac{\pi^2}{4f_e^2}, \quad (1.92)$$

la première inégalité étant une égalité pour $\varphi = 0$ et $f_e = 2k + 1$, $k \in \mathbb{N}$.

Pour ce faire, nous allons utiliser les deux approximations suivantes :

1. la valeur de l'enveloppe supérieure est comprise entre le plus petit maximum et 1. Similairement, celle de l'enveloppe inférieure est comprise entre -1 et le plus grand minimum.
2. le premier IMF est calculé en une seule itération de tamisage.

La première approximation nous permet d'encadrer les valeurs des enveloppes par les valeurs extrêmes des maxima et minima. Dans le cas d'une interpolation spline cubique, cette approximation est très

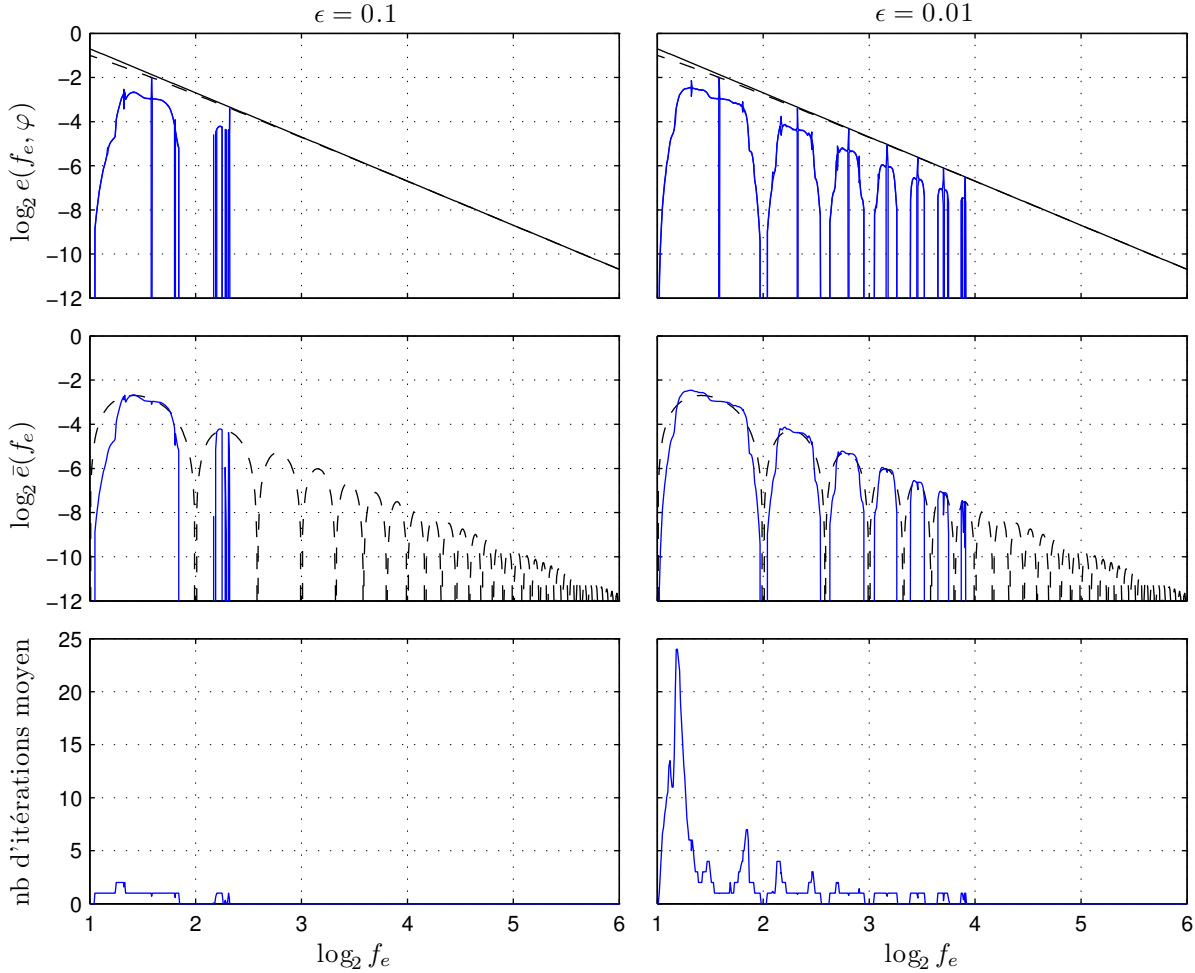


FIGURE 1.18 – Écart entre le premier IMF et le signal sinusoïdal lorsque le nombre d'itérations de tamisage est adapté selon le critère « local » (cf 5.2). Dans le cas de signaux sinusoïdaux, il n'y a pas d'inconvénient à utiliser un paramètre de tolérance $\alpha = 0$, ce qui permet de ramener le critère d'arrêt au seul paramètre ϵ . Les colonnes de gauche et de droite correspondent à deux paramètres ϵ différents, 0.1 et 0.01 respectivement. En haut, l'écart maximum et minimum (par rapport à la phase) en fonction de f_e . La courbe en trait plein au dessus correspond à la majoration $\pi^2/(4f_e^2)$; la courbe en tirets à la majoration $(1 - \cos(\pi/f_e))/2$. Au milieu, l'écart moyenné par rapport à la phase $\bar{e}(f_e)$. La courbe en tirets correspond à l'approximation (1.111). En bas les nombres moyens d'itérations de tamisage effectuées.

discutable dans le cas général mais il se trouve qu'elle est raisonnablement bien vérifiée dans le cas particulier des sinusoides. On peut par ailleurs signaler que s'il s'agit bien d'une approximation dans le cas d'une interpolation spline cubique, ce n'est pas nécessairement le cas pour tout schéma d'interpolation. Pour une interpolation linéaire par morceaux par exemple, elle serait exacte.

La deuxième approximation est un peu plus discutable dans la mesure où il faut généralement plus d'une itération de tamisage pour obtenir un IMF dont les enveloppes sont « bien symétriques ». Cependant, dans le cas d'une simple sinusoides, la première itération est toujours la plus importante et l'effet des itérations suivantes est moins important même si elles sont nombreuses. Ce résultat est valable pour une simple sinusoides parce qu'aucun extremum n'apparaît au cours du processus de tamisage. Dans des cas plus compliqués, il est fréquent que des extrema apparaissent au cours du processus de tamisage auquel cas les itérations suivant ces apparitions peuvent avoir des effets aussi importants que la toute première itération.

Partant de ces deux approximations, le résultat (1.92) s'obtient de la manière suivante. En utilisant la deuxième hypothèse, on obtient que la différence entre $x_{f_e, \varphi}[n]$ et $d_{f_e, \varphi}[n]$ n'est autre que la moyenne des enveloppes de $x_{f_e, \varphi}[n]$

$$m_{f_e, \varphi}[n] = \frac{e_{\max}[n] + e_{\min}[n]}{2}. \quad (1.93)$$

Si on considère l'une de ces enveloppes, par exemple l'enveloppe supérieure, la première hypothèse permet d'encadrer sa valeur par

$$\alpha(f_e) \leq e_{\max}[n] \leq 1, \quad (1.94)$$

où $\alpha(f_e)$ est la borne inférieure des valeurs des maxima de $x_{f_e, \varphi}[n]$ pour φ quelconque. Symétriquement, on a les inégalités inverses pour l'enveloppe inférieure

$$-1 \leq e_{\min}[n] \leq -\alpha(f_e). \quad (1.95)$$

La valeur de $\alpha(f_e)$ s'obtient en considérant le cas particulier où 2 points d'échantillonnage consécutifs sont placés symétriquement par rapport à un extremum de $x(t)$: l'écart entre ces deux points étant la période d'échantillonnage, quelle que soit la valeur de φ , il y a toujours un point d'échantillonnage entre ces deux points. Par conséquent, la valeur de $x(t)$ en ces deux points fournit la borne inférieure recherchée : $\alpha(f_e) = \cos(\pi/f_e)$, qui est positive si $f_e \geq 2$. De là, on obtient

$$|m_{f_e, \varphi}[n]| = \frac{||e_{\max}[n]| - |e_{\min}[n]||}{2} \quad (1.96)$$

$$\leq \frac{1 - \alpha}{2}, \quad (1.97)$$

$$\leq \frac{1}{2} \left(1 - \cos \frac{\pi}{f_e} \right). \quad (1.98)$$

La deuxième partie du résultat (1.92) découle alors de la formule classique « $\cos u \geq 1 - u^2/2$ ».

Dans le cas où $f_e = 2k + 1$, $k \in \mathbb{N}$ et $\varphi = 0$, tous les maxima de $x_{f_e, \varphi}[n]$ valent 1 et tous les minima $-\cos(\pi/f_e)$. Il en découle que la première inégalité de (1.92) est dans ce cas une égalité et la borne est donc atteinte.

Enfin, dans le cas où le nombre d'itérations est déterminé au cours du processus de tamisage par le critère « local », cette borne permet d'évaluer à partir de quelle fréquence d'échantillonnage, le signal sinusoidal est considéré comme étant déjà un IMF par l'algorithme. Dans le cas de signaux sinusoidaux, on peut considérer que l'écart entre les enveloppes $|U(t) - L(t)|$ vaut approximativement 2, ce qui permet d'exprimer le critère sous la forme

$$|m_{f_e, \varphi}| < \epsilon. \quad (1.99)$$

On en déduit donc la condition suffisante : si

$$\frac{\pi^2}{4f_e^2} < \epsilon, \quad (1.100)$$

i.e.

$$f_e > \frac{\pi}{2\sqrt{\epsilon}}, \quad (1.101)$$

alors le signal sinusoïdal est considéré comme étant déjà un IMF par l'algorithme. Dans le cas de la figure Fig. 1.18, on obtient ainsi que le nombre d'itérations est nul dès que $f_e > 5.0$ pour $\epsilon = 0.1$ et $f_e > 15.7$ pour $\epsilon = 0.01$.

5.4.1.3 Estimation de l'erreur moyenne À partir des deux hypothèses utilisées pour obtenir la borne précédente, on propose un modèle plus perfectionné permettant de retrouver l'ordre de grandeur de $\bar{e}(f_e)$, notamment le fait que $\bar{e}(f_e)$ s'annule pour les fréquences $f_e = 2k$, $k \in \mathbb{N}$.

Sachant que φ est uniformément distribuée sur $[0, 1]$, le processus $x_{f_e, \varphi}[n]$ est stationnaire à tout ordre. En supposant les signaux de durée infinie, de manière à ignorer les effets de bords, et en utilisant le fait que l'EMD est covariante par rapport aux décalages temporels, on peut montrer que les IMFs et approximations issus de $x_{f_e, \varphi}[n]$ sont également stationnaires à tout ordre.

De manière générale, si $x(t)$, $t \in \mathbb{R}$ est un processus stationnaire à tout ordre, $x(t)$ et le processus translaté défini par $\forall t \in \mathbb{R}, x_\tau(t) = x(t - \tau)$ sont identiques en distribution :

$$\{x(t), t \in \mathbb{R}\} \stackrel{d}{=} \{x_\tau(t), t \in \mathbb{R}\}, \quad (1.102)$$

et donc le k^e IMF $d_k[x_\tau](t)$ vérifie

$$\{d_k[x](t), t \in \mathbb{R}\} \stackrel{d}{=} \{d_k[x_\tau](t), t \in \mathbb{R}\} \stackrel{d}{=} \{(d_k[x])_\tau(t), t \in \mathbb{R}\} \stackrel{d}{=} \{d_k[x](t - \tau), t \in \mathbb{R}\}, \quad (1.103)$$

où on a utilisé le fait que l'EMD commute avec les translations temporelles. En revanche, l'EMD étant non linéaire, si $x(t)$ n'est pas stationnaire à tout ordre, on ne peut rien conclure sur la stationnarité de ses IMFs à aucun ordre.

$x_{f_e, \varphi}[n]$ et le premier IMF $d_{f_e, \varphi}[n]$ étant stationnaires, on en déduit que pour tout n ,

$$\bar{e}(f_e) = \mathbb{E}_\varphi\{|x_{f_e, \varphi}[n] - d_{f_e, \varphi}[n]|\}, \quad (1.104)$$

et donc en utilisant la deuxième hypothèse

$$\bar{e}(f_e) = \mathbb{E}_\varphi\{|m_{f_e, \varphi}[n]|\} = \mathbb{E}_\varphi\left\{\frac{|e_{\max}[n] + e_{\min}[n]|}{2}\right\}. \quad (1.105)$$

L'étape suivante est d'estimer la valeur de la moyenne des enveloppes. Dans ce but, on fait l'approximation que la valeur d'une enveloppe est égale à celle de l'extremum le plus proche :

$$e_{\max/\min}[n] \approx \cos\left(\frac{2\pi}{f_e}n_{\max/\min} + 2\pi\varphi\right), \quad (1.106)$$

où $n_{\max/\min}$ correspond à l'extremum (maximum pour l'enveloppe supérieure, minimum pour l'enveloppe inférieure) le plus proche de n . De là, sachant que tout point n se trouve toujours entre un maximum et un minimum, le problème se ramène à l'étude des relations entre par exemple un maximum et le minimum suivant. Si on considère un maximum en $n = n_{\max}$, la phase vérifie

$$\exists k \in \mathbb{Z}, \quad \varphi_{\max} = \frac{n_{\max}}{f_e} + \varphi - k \in \left[-\frac{1}{2f_e}, \frac{1}{2f_e}\right]. \quad (1.107)$$

Par conséquent, si on suppose que n_{max} est la position d'un maximum, la phase φ_{max} en n_{max} est uniformément distribuée sur $[-1/(2f_e), 1/(2f_e)]$. Soit alors $\tilde{n}$ tel que la position du minimum suivant directement soit $n_{max} + \tilde{n}$. D'après le raisonnement qu'on vient de proposer, on a alors

$$\bar{e}(f_e) \approx \mathbb{E}_\varphi\{|m_{f_e, \varphi}[n]|\} \approx \frac{1}{2} \mathbb{E}_\psi \left\{ \left| \cos \psi + \cos \left(\frac{2\pi \tilde{n}(\psi)}{f_e} + \psi \right) \right| \right\}, \quad (1.108)$$

où ψ est distribuée uniformément dans $[-\pi/f_e, \pi/f_e]$. Si on note K la partie entière de $f_e/2$, $K = \lfloor f_e/2 \rfloor$ et $\tilde{\psi} = \pi - (2K + 1)\pi/f_e$, on peut calculer $\tilde{n}(\psi)$:

$$\tilde{n}(\psi) = \begin{cases} K & \text{si } \psi > \tilde{\psi}, \\ K + 1 & \text{si } \psi < \tilde{\psi}. \end{cases} \quad (1.109)$$

Si on remplace maintenant $\tilde{n}(\psi)$ dans (1.108), on obtient :

$$\begin{aligned} \mathbb{E}_\varphi\{|m_{f_e, \varphi}|\} &\approx \frac{f_e}{4\pi} \int_{-\pi/f_e}^{\pi/f_e} \left| \cos \psi + \cos \left(\frac{2\pi \tilde{n}(\psi)}{f_e} + \psi \right) \right| d\psi, \\ &\approx \frac{f_e}{4\pi} \left[\int_{\tilde{\psi}}^{\pi/f_e} \left| \cos \psi + \cos \left(\frac{2\pi K}{f_e} + \psi \right) \right| d\psi \right. \\ &\quad \left. + \int_{-\pi/f_e}^{\tilde{\psi}} \left| \cos \psi + \cos \left(\frac{2\pi(K+1)}{f_e} + \psi \right) \right| d\psi \right]. \end{aligned}$$

Le calcul mène finalement à

$$\bar{e}(f_e) \approx \frac{f_e}{\pi} \left[\cos \left(\frac{K\pi}{f_e} \right) \left(1 - \sin \left(\frac{(K+1)\pi}{f_e} \right) \right) + \cos \left(\frac{(K+1)\pi}{f_e} \right) \left(\sin \left(\frac{K\pi}{f_e} \right) - 1 \right) \right]. \quad (1.110)$$

Enfin, cette formule peut encore être simplifiée si l'on utilise les approximations classiques « $\sin u \approx u$ » et « $1 - \cos u \approx u^2/2$ » valables quand $u \rightarrow 0$. Si on suppose $\pi/f_e \ll 1$, il s'en suit que $|\pi/2 - K\pi/f_e| \ll 1$ également et donc, (1.110) se simplifie en

$$\bar{e}(f_e) \approx \frac{\pi^2}{8f_e^2} (2(K+1) - f_e)(f_e - 2K). \quad (1.111)$$

Finalement, on obtient donc un modèle parabolique sur chaque intervalle $[2k, 2(k+1)]$, $k \in \mathbb{N}$ qui permet de retrouver à la fois le fait que l'écart s'annule pour les fréquences d'échantillonnage égales à des entiers pairs et la décroissance globale en f_e^{-2} mise en évidence dans le préfacteur.

De plus, on remarque que le modèle (1.110) est exact pour les fréquences d'échantillonnage entières puisque pour ces fréquences, tous les maxima et tous les minima sont égaux et donc les enveloppes sont effectivement égales à la valeur du maximum ou minimum le plus proche.

5.4.2 Généralisation à une itération de tamisage pour un signal quelconque

On se propose d'élargir les modèles précédents à un cadre beaucoup plus général englobant une très grande partie des signaux oscillants. Dans ce but, il nous faut d'abord déterminer des conditions minimales sur la fréquence d'échantillonnage d'un signal à temps continu pour pouvoir lui appliquer l'EMD avec un minimum de garanties. Dans la mesure où l'EMD s'appuie sur les extrema locaux du signal pour analyser les différentes échelles de temps qui le composent, il semble important que tous les extrema locaux du signal continu aient une contrepartie dans sa version discrétisée. En effet, si un extremum est perdu lors de l'échantillonnage, c'est en fait généralement une paire d'extrema qui est perdue et donc une échelle de temps à un certain endroit. Dans le cas général, pour une fréquence

d'échantillonnage f_e donnée, il est difficile d'assurer que, pour toute phase d'échantillonnage, tous les extrema du signal continu ont une contrepartie à temps discret. Par exemple dans un cas extrême comme celui du signal

$$x(t) = \begin{cases} 1 & \text{si } t \in \mathbb{Z}, \\ -1 & \text{si } t + \frac{1}{2} \in \mathbb{Z}, \\ 0 & \text{sinon,} \end{cases} \quad (1.112)$$

il est aisé de voir que les seuls couples (f_e, φ) tels que $x(n/f_e + \varphi)$ conserve tous les extrema de $x(t)$ sont ceux de la forme $(f_e = 2k, \varphi = k'/(2k))$, $k, k' \in \mathbb{N}$. En revanche, si l'on impose que

1. les intervalles où $x(t)$ est constant sont au plus de longueur Δ
2. l'écart entre deux extrema de $x(t)$ est supérieur à 2Δ

alors $x(n/f_e + \varphi)$ conserve tous les extrema de $x(t)$ dès lors que $f_e > 1/\Delta$. En effet, dans ces conditions, on est sûr qu'entre deux extrema successifs de $x(t)$, il y a au moins deux points d'échantillonnage et que de plus, ces deux points n'ont pas la même ordonnée, ce qui permet de repérer le « signe de la dérivée » ou plutôt le sens croissant ou décroissant. Ainsi, si on considère un maximum de $x(t)$ en t_0 entouré de deux minima en t_1 et t_{-1} , le signal échantillonné $x(n/f_e + \varphi)$ a une dérivée (discrete) strictement positive en au moins un point dans l'intervalle $[t_{-1}, t_0]$ et strictement négative en au moins un point de $[t_0, t_1]$. On en déduit que le signal discret présente nécessairement un maximum sur $[t_{-1}, t_1]$.

Remarque : Les conditions sur l'écart entre deux extrema et sur la longueur des intervalles constants considérées ici pour déterminer la fréquence d'échantillonnage minimale sont totalement indépendantes du classique critère de Shannon selon lequel un signal est correctement échantillonné s'il est à bande limitée et que la fréquence d'échantillonnage est supérieure à la largeur de la bande. En effet un signal à bande limitée peut avoir des extrema arbitrairement proches quelle que soit la largeur de la bande. Par exemple, $t \mapsto (1 - \epsilon) \sin \epsilon f t - \epsilon \sin f t$ présente deux extrema aux environs de $\pm \sqrt{2\epsilon}/f$ lorsque $\epsilon \ll 1$.

Étant donné un signal $x(t)$ dont l'écart entre les extrema est supérieur à 2Δ et dont la longueur des intervalles constants est inférieure à Δ , on définit ses versions échantillonnées

$$x_{f_e, \varphi}[n] = x\left(\frac{n}{f_e} + \varphi\right), \quad n \in \mathbb{Z}, \quad (1.113)$$

avec $f_e > 1/\Delta$ pour assurer qu'aucun extremum ne disparaît au cours de l'échantillonnage. Si on reprend le raisonnement qui a été développé pour modéliser les effets d'échantillonnages sur les signaux sinusoïdaux, on s'est basé sur deux hypothèses. Parmi celles-ci, la deuxième selon laquelle on peut considérer que le premier IMF est obtenu à l'aide d'une seule itération de tamisage s'est révélée une bonne approximation pour le cas des signaux sinusoïdaux mais elle devient très mauvaise dans le cas général. En effet, des extrema peuvent apparaître au cours du processus de tamisage et les itérations suivant ces apparitions peuvent avoir des effets aussi importants que la première, ce qui oblige à prendre en compte toutes les itérations de tamisage. Malheureusement, ceci présente deux inconvénients majeurs pour l'analyse des effets d'échantillonnage :

1. De manière générale, le nombre d'itérations nécessaire à l'extraction d'un IMF n'est pas connu a priori mais déterminé au cours du processus de tamisage. Il en résulte que ce nombre d'itérations dépend en général des paramètres de l'échantillonnage f_e et φ .
2. Si pour un couple (f_e, φ) donné, une paire d'extrema apparaît à une certaine itération de tamisage, il n'y a aucune garantie qu'une paire d'extrema similaire apparaisse pour un autre couple (f_e, φ) à une quelconque itération du processus de tamisage.

Face à ces difficultés, on voit que le cadre du modèle développé précédemment est insuffisant pour appréhender pleinement les effets d'échantillonnage. En revanche, il permet a priori d'étudier les effets d'échantillonnage sur une seule itération de tamisage, ce à quoi on se restreindra par la suite.

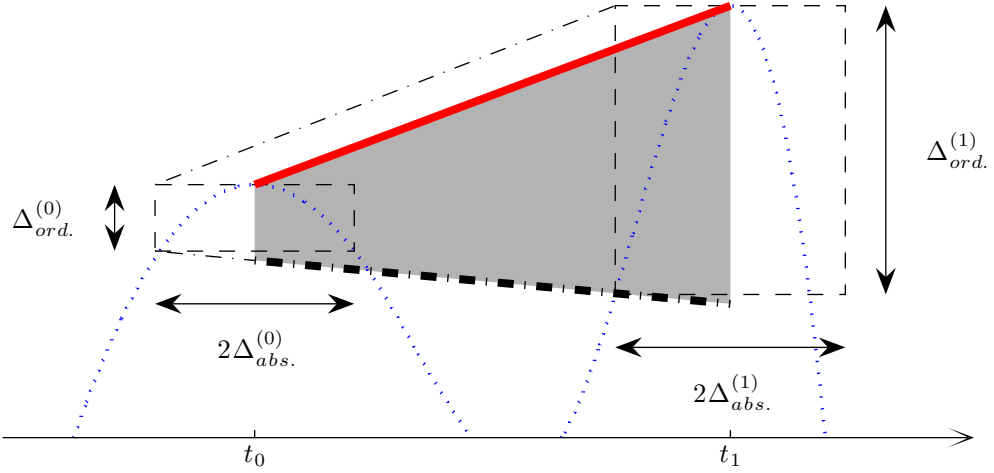


FIGURE 1.19 – Pour chacun des maxima, les rectangles en tirets délimitent les zones où sont situés les maxima correspondants dans le signal échantillonné. La ligne épaisse rouge reliant les deux maxima est l'enveloppe calculée à partir des extrema du signal continu dans l'approximation linéaire par morceaux. Les lignes en tirets-points délimitent la zone où peut se trouver l'interpolation calculée à partir des extrema du signal échantillonné. Enfin, la ligne épaisse en tirets-points correspond au scénario aboutissant à l'écart le plus important.

5.4.2.1 Détection d'extrema Considérons un maximum de $x(t)$ situé en t_0 . Les conditions sur $x(t)$ permettent d'assurer que si $f_e > 1/\Delta$, alors $x_{f_e, \varphi}[n]$ présente un maximum en n_0 tel que $|n_0/f_e + \varphi - t_0| < 1/f_e$, c'est-à-dire que soit le point d'échantillonnage immédiatement après t_0 soit celui immédiatement avant correspond à un maximum du signal échantillonné. Plus précisément, on peut définir l'indice de l'extremum dans le signal échantillonné comme l'unique indice (ou à la rigueur l'un des deux indices) n tel que $n/f_e + \varphi \in [t_0 - a, t_0 + b]$ avec $a, b > 0$ tels que $a + b = 1/f_e$ et $x(t_0 - a) = x(t_0 + b)$. On peut donc définir une incertitude en abscisse pour la position du maximum dans le signal échantillonné

$$\Delta_{abs.}^{(0)} = \max\{a, b\}. \quad (1.114)$$

L'incertitude en ordonnée correspondante vaut alors

$$\Delta_{ord.}^{(0)} = |x(t_0) - x(t_0 - a)| = |x(t_0) - x(t_0 + b)| \quad (1.115)$$

$$= \min\{|x(t_0) - x(t_0 + \Delta_{abs.}^{(0)})|, |x(t_0) - x(t_0 - \Delta_{abs.}^{(0)})|\}. \quad (1.116)$$

5.4.2.2 Répercussions sur une itération de tamisage Connaissant les incertitudes sur les positions des extrema, on s'intéresse maintenant à leurs répercussions sur les enveloppes. Dans ce but, il faudrait théoriquement calculer l'impact de ces incertitudes sur la valeur de l'interpolation en tout point du signal. En pratique, le problème est bien souvent inextricable dans la mesure où la valeur d'une interpolation en un point dépend bien souvent de tous les points d'interpolation. C'est le cas notamment de la plupart des schémas d'interpolation qui garantissent un certain niveau de régularité comme les splines d'ordre supérieur ou égal à 2. Devant cette difficulté, on propose d'approcher les effets d'échantillonnage en considérant le cas d'une interpolation linéaire par morceaux, qui présente l'avantage de ne dépendre localement que de deux points d'échantillonnage. En pratique, il semble que cette approximation soit en fait très raisonnable, les incertitudes obtenues étant en bon accord avec les observations.

Pour modéliser l'incertitude sur les enveloppes, on considère deux maxima successifs de $x(t)$ en $(t_0, x(t_0))$ et $(t_1, x(t_1))$. Les maxima correspondants dans le signal échantillonné sont situés dans les

boîtes rectangulaires $\{[t_i - \Delta_{abs.}^{(i)}, t_i + \Delta_{abs.}^{(i)}] \times [x(t_i) - \Delta_{ord.}^{(i)}, x(t_i)]; i = 0, 1\}$ (cf Fig. 1.19). Grâce à la figure, on peut voir que l'écart le plus important correspond dans le cas général au cas de la ligne tirets-points épaisse sur la figure. Si on intègre cet écart entre t_0 et t_1 , on obtient alors une borne de l'écart dû à l'échantillonnage sur $[t_0, t_1]$ pour l'enveloppe supérieure

$$\int_{t_0}^{t_1} |\delta e_{max}(t)| dt \leq \frac{(t_1 - t_0)(\Delta_{ord.}^{(0)} + \Delta_{ord.}^{(1)})}{2} + (t_1 - t_0) \frac{|x(t_1) - x(t_0) - \Delta_{ord.}^{(1)} + \Delta_{ord.}^{(0)}|(\Delta_{abs.}^{(1)} + \Delta_{abs.}^{(0)})}{2(t_1 - t_0 - |\Delta_{abs.}^{(1)} - \Delta_{abs.}^{(0)}|)}, \quad (1.117)$$

qu'on peut découper grossièrement pour aboutir à la borne moins stricte :

$$\begin{aligned} \int_{t_0}^{t_1} |\delta e_{max}(t)| dt &\leq (t_1 - t_0) \cdot \frac{\Delta_{ord.}^{(0)} + \Delta_{ord.}^{(1)}}{2} + |x(t_1) - x(t_0)| \cdot \frac{\Delta_{abs.}^{(1)} + \Delta_{abs.}^{(0)}}{2} \\ &+ \frac{t_1 - t_0}{t_1 - t_0 - |\Delta_{abs.}^{(1)} - \Delta_{abs.}^{(0)}|} \cdot \frac{\Delta_{abs.}^{(1)} + \Delta_{abs.}^{(0)}}{2} \cdot |\Delta_{ord.}^{(1)} - \Delta_{ord.}^{(0)}| + \frac{|\Delta_{abs.}^{(1)} - \Delta_{abs.}^{(0)}|}{t_1 - t_0 - |\Delta_{abs.}^{(1)} - \Delta_{abs.}^{(0)}|} \cdot |x(t_1) - x(t_0)|, \end{aligned} \quad (1.118)$$

Dans cette dernière formule, les deux premiers termes sont généralement les plus importants. Le premier rend compte de l'incertitude en ordonnée de l'interpolation directement due à l'incertitude en ordonnée des extrema. Le second dépend, lui, des positions relatives des extrema successifs : son importance dépend donc de celle des variations en ordonnée des extrema du signal continu. Le troisième terme correspondant à l'écart généré par la différence des incertitudes en ordonnée des deux extrema est généralement plus faible que les deux premiers, sauf éventuellement pour des fréquences d'échantillonnage à peine au dessus de la limite $1/\Delta$ (voire en dessous comme dans le cas de la figure Fig. 1.19). Enfin, le dernier terme est généralement beaucoup plus faible, d'une part parce qu'il est proportionnel à la différence entre les incertitudes en abscisse et d'autre part parce que celles-ci sont inférieures au quart de $t_1 - t_0$ pour $f_e = 1/\Delta$ et encore plus petites sinon.

Finalement, si on calcule la borne (1.117) ou (1.118) sur tous les intervalles entre deux maxima on obtient une borne des effets d'échantillonnages moyens pour l'enveloppe supérieure. La même opération sur l'enveloppe inférieure permet d'aboutir finalement à une borne sur les effets d'échantillonnage moyens pour la moyenne des enveloppes qui est simplement la moyenne des deux bornes obtenues pour les deux enveloppes. Enfin, cette borne sur les effets d'échantillonnage pour la moyenne des enveloppes constitue également une borne des effets d'échantillonnage sur une itération de tamisage.

5.4.2.3 Estimation des paramètres à partir du signal Étant donné un signal continu $x(t)$ le calcul des bornes (1.117) ou (1.118) requiert un petit nombre de paramètres faciles à calculer à partir du signal : les positions des extrema $(t_i, x(t_i))$ et les incertitudes $\Delta_{abs.}^{(i)}$ et $\Delta_{ord.}^{(i)}$ correspondantes. Cependant, si on ne dispose que d'une version échantillonnée du signal, on manque d'informations pour déterminer tous ces paramètres. Les positions des extrema peuvent être estimées par leurs positions dans le signal discret : $\hat{t}_i = n_i/f_e + \varphi$ et $\hat{x}_i = x_{f_e, \varphi}(n_i)$, où n_i est l'indice du i^e extremum du signal discret. Pour ces estimations, on remarque que les positions des extrema en abscisse sont évaluées avec une précision de l'ordre de $1/f_e$. Ceci a malheureusement comme conséquence qu'il est impossible d'estimer correctement les incertitudes en abscisse $\Delta_{abs.}^{(i)} \in [1/(2f_e), 1/f_e]$ à partir du signal discret, ce qui impose d'utiliser la valeur maximale $1/f_e$ pour calculer la borne, à moins d'avoir plus d'informations sur le comportement du signal continu autour de ses extrema. Les incertitudes en ordonnée posent un problème similaire mais on peut les estimer si on peut faire l'hypothèse que le signal est au moins deux fois continuellement dérivable autour de ses extrema. Dans ce cas, un développement de Taylor au voisinage des extrema suggère un comportement localement proche de $x(t) - x(t_i) \approx 0.5(t - t_i)^2 x''(t_i)$ qui mène à une évaluation de l'incertitude en ordonnée $\Delta_{ord.}^{(i)} \approx$

$0.5|x''(t_i)|/f_e^2$. Enfin, les dérivées secondes peuvent être estimées par des méthodes de différences finies au niveau des extrema du signal échantillonné $x_i'' \approx x''(t_i)$. Si on utilise ces valeurs dans la borne (1.118), on aboutit alors à une estimation de l'écart maximal engendré par l'échantillonnage sur une itération de tamisage sur l'intervalle $[t_i, t_{i+1}]$:

$$\int_{t_i}^{t_{i+1}} |\delta e_{max}(t)| dt \leq \frac{(\hat{t}_{i+1} - \hat{t}_i) |\hat{x}_i'' + \hat{x}_{i+1}''|}{4f_e^2} + \frac{|\hat{x}_{i+1} - \hat{x}_i|}{f_e} + \frac{|\hat{x}_{i+1}'' - \hat{x}_i''|}{2f_e^3}. \quad (1.119)$$

Finalement, ceci mène à la borne sur la norme L_1 de l'écart du à l'échantillonnage sur une itération de tamisage :

$$\|\delta \mathcal{S}x\|_1 \leq \frac{\lambda}{f_e} + \frac{\mu}{f_e^2} + \frac{\nu}{f_e^3}, \quad (1.120)$$

avec

$$\lambda = \frac{1}{2} \left(\sum_i |\hat{x}_{i+1}^m - \hat{x}_i^m| + \sum_i |\hat{x}_{i+1}^M - \hat{x}_i^M| \right), \quad (1.121)$$

$$\mu = \frac{1}{8} \left(\sum_i (\hat{t}_{i+1}^m - \hat{t}_i^m) |\hat{x}_i^{m''} + \hat{x}_{i+1}^{m''}| + \sum_i (\hat{t}_{i+1}^M - \hat{t}_i^M) |\hat{x}_i^{M''} + \hat{x}_{i+1}^{M''}| \right), \quad (1.122)$$

$$\nu = \frac{1}{4} \sum_i |\hat{x}_{i+1}^{m''} - \hat{x}_i^{m''}| + |\hat{x}_{i+1}^{M''} - \hat{x}_i^{M''}|, \quad (1.123)$$

où les exposants m et M se rapportent aux minima et aux maxima respectivement. L'intégration pour la norme L_1 dans (1.120) est faite sur la durée entre le second et l'avant dernier extremum de manière à être toujours entre deux maxima et entre deux minima. Pour des signaux comportant de nombreux extrema jusqu'à proximité des bords, on approximera cette durée par la durée totale du signal.

Remarque : Dans le cas sinusoïdal, on avait abouti à une borne proportionnelle à f_e^{-2} . La nouvelle borne obtenue ici (1.120) généralise ce résultat, puisque dans le cas sinusoïdal les paramètres λ et ν sont nuls. En revanche, le paramètre μ calculé selon (1.122) pour un signal sinusoïdal est 8 fois plus grand que ce qu'on avait calculé précédemment (1.92). Cette différence a deux origines. D'une part, le signal sinusoïdal est parfaitement symétrique autour de ses extrema : ceci permet de réduire $\Delta_{abs.}$ à $1/(2f_e)$ et les paramètres λ, μ et ν à $\lambda' = \lambda/2$, $\mu' = \mu/4$ et $\nu' = \nu/8$. D'autre part, les écarts générés par l'enveloppe supérieure et l'enveloppe inférieure se compensent partiellement en général, ce qui permet de réduire la borne encore de moitié dans le cas sinusoïdal.

5.4.2.4 Confrontation avec l'expérience Pour évaluer la validité de la borne (1.120) sur une itération de tamisage, il nous faut pouvoir comparer le résultat d'une itération de tamisage réalisée à temps continu sur un signal à temps continu avec celui d'une itération de tamisage réalisée à temps discret sur une version échantillonnée du signal. Dans ce but, on a réalisé une série de simulations basées sur un modèle de signal polynomial par morceaux. Comme on peut aisément calculer les extrema d'un polynôme de degré inférieur ou égal à 4, il est tout à fait possible de calculer une EMD à temps continu pour un signal polynomial par morceaux dont le degré est inférieur ou égal à 4. De là, la procédure d'évaluation est la suivante :

1. synthétiser un signal polynomial par morceaux $x(t)$, $t \in [0, L]$ tel que la distance minimale entre deux extrema soit supérieure à 2. T est choisi suffisamment grand pour pouvoir négliger les effets de bords. Le modèle de signal utilisé est constitué d'extrema dont les positions t_i sont générées aléatoirement avec des temps d'attente aléatoires selon une variable exponentielle de moyenne 1 à laquelle on ajoute 2 pour être sûr d'avoir un écart supérieur à 2. Pour les valeurs aux extrema, on génère un bruit discret essentiellement haute fréquence $b[i]$ (le modèle utilisé

est quelque peu alambiqué mais il semble que les résultats ne dépendent en pratique que peu de ce dernier.) et on définit $x(t_i) = b[i]$. Enfin, le signal $x(t)$ est défini entre les extrema par interpolation selon le schéma d'interpolation cubique `pchip` de `Matlab`⁶. Par construction, le signal n'a presque sûrement pas d'intervalles où il est constant.

2. appliquer l'opérateur de tamisage continu à $x(t) : d_\infty(t) = (\mathcal{S}x)(t)$
3. définir les versions échantillonnées $x_{f_e, \varphi}[n] = x(n/f_e + \varphi)$ avec $0 \leq n \leq N(f_e) = \lfloor T f_e \rfloor - 1$ et $0 \leq \varphi \leq 1/f_e$.
4. pour chaque couple (f_e, φ) ,
 - (a) estimer la borne (1.120) : $b(f_e, \varphi)$
 - (b) appliquer une itération de tamisage discret à chacune des versions échantillonnées : $d_{f_e, \varphi}[n] = (\mathcal{S}x_{f_e, \varphi})[n]$.
 - (c) calculer l'écart au continu :

$$e(f_e, \varphi) \equiv \frac{1}{N(f_e) + 1} \sum_{n=0}^{N(f_e)} \left| d_{f_e, \varphi}[n] - d_\infty\left(\frac{n}{f_e} + \varphi\right) \right|. \quad (1.124)$$

Des résultats de ces simulations sont proposés Fig. 1.20 pour deux exemples représentatifs. De manière générale, on observe que le comportement de la borne en fonction de la fréquence d'échantillonnage présente deux régimes : pour les faibles valeurs de f_e , elle se comporte comme f_e^{-2} alors que pour les grandes valeurs, elle se comporte comme f_e^{-1} . Cependant, étant donné que le paramètre λ peut être nul si tous les maxima et tous les minima ont la même ordonnée, c'est à dire si le signal s'apparente à une modulation de fréquence, la borne peut se comporter comme f_e^{-2} même pour les grandes fréquences d'échantillonnage (cf Fig. 1.20 (a)). L'écart au continu présente également les deux régimes en général mais il arrive que le régime en f_e^{-2} n'apparaisse pas pour les fréquences d'échantillonnage $f_e > 1/\Delta$, contrairement au cas de la borne (1.120). De plus, lorsque les deux régimes sont présents pour la borne et pour l'écart au continu mesuré, il est rare que les fréquences de coupures séparant les deux régimes soient les mêmes. Enfin, le comportement en f_e^{-3} auquel on pourrait s'attendre pour les petites valeurs de f_e n'a jamais été observé.

Quantitativement, on observe que la borne est toujours de très loin supérieure à l'écart au continu mesuré en pratique : aux environs de 12 dB au dessus dans la partie f_e^{-1} et 25 dB dans la partie f_e^{-2} . Ceci est en fait normal a priori puisque la borne est établie en considérant le scénario menant à l'écart le plus important partout alors qu'en pratique, on observe plutôt un écart moyen. Pour évaluer plus précisément l'adéquation entre le modèle et la réalité, on peut au lieu de calculer une borne calculer plutôt l'écart engendré en moyenne par les effets d'échantillonnage. Dans ce but, on définit les versions moyennes des écarts en abscisse et ordonnée :

$$\begin{aligned} \forall i, \quad \bar{\Delta}_{abs.}^{(i)} &= \mathbb{E}_\varphi\{|t_i - (n_i f_e + \varphi)|\}, \\ \bar{\Delta}_{ord.}^{(i)} &= \mathbb{E}_\varphi\{|x(t_i) - x_{f_e, \varphi}[n_i]|\}, \end{aligned}$$

où pour estimer l'écart moyen, on moyenne simplement par rapport à la phase de l'échantillonnage φ sans tenir compte des dépendances entre les différents extrema. Pour estimer ces quantités, on peut à nouveau faire appel à l'approximation parabolique autour des extrema, ce qui fournit

$$\begin{aligned} \bar{\Delta}_{abs.}^{(i)} &= \frac{\Delta_{abs.}^{(i)}}{2}, \\ \bar{\Delta}_{ord.}^{(i)} &= \frac{\Delta_{ord.}^{(i)}}{3}. \end{aligned}$$

6. Le signal obtenu n'est pas nécessairement deux fois dérivable à ses extrema mais il est suffisamment régulier pour que ça ne pose pas de problème.

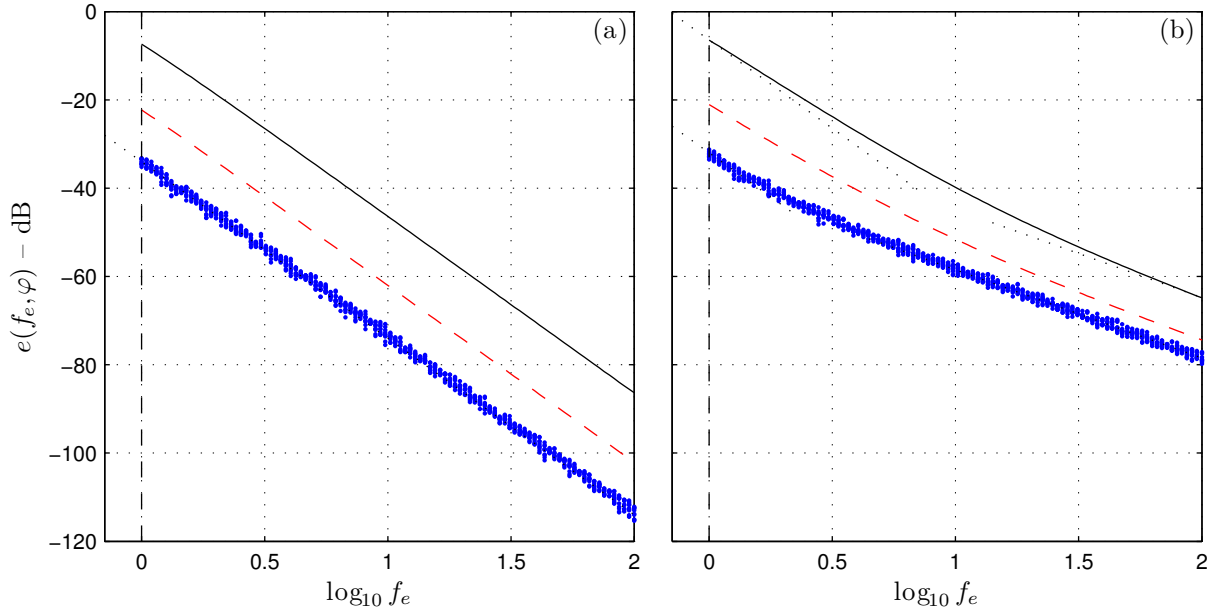


FIGURE 1.20 – Pour les deux graphes, chaque point représente une mesure d'écart au continu (1.124). Pour chaque couple (f_e, φ) , on calcule la borne (1.120) et l'écart moyen (1.126), la moyenne de ces quantités par rapport à la phase est représentée en trait plein noir et tirets rouges respectivement. (a) le signal continu a des enveloppes constantes ; on peut le voir comme une modulation de fréquence au sens large. (b) les maxima/minima du signal continu ont en moyenne une valeur de $1/-1$ avec une variance de 0.1. Les comportements asymptotiques en f_e^{-1} ou f_e^{-2} sont mis en évidence par des pointillés.

Enfin, on a également surestimé l'incertitude en abscisse $\Delta_{abs.}^{(i)} = 1/f_e$ faute de pouvoir l'estimer précisément. Si l'on a par ailleurs des informations sur le signal permettant de réduire cette incertitude à α/f_e , alors les coefficients λ , μ et ν deviennent :

$$\lambda' = \alpha\lambda, \quad \mu' = \alpha^2\mu \quad \text{et} \quad \nu' = \alpha^3\nu. \quad (1.125)$$

Dans le cas du modèle de signal utilisé lors des simulations, l'incertitude en abscisse est pratiquement toujours inférieure à $0.7/f_e$.

La combinaison des incertitudes en abscisse réduites α/f_e et des incertitudes moyennes permet d'aboutir finalement à l'estimation de l'écart moyen

$$\bar{e}(f_e) = \frac{\alpha\lambda}{2f_e} + \frac{\alpha^2\mu}{3f_e^2} + \frac{\alpha^3\nu}{6f_e^3}. \quad (1.126)$$

Cette estimation de l'écart moyen est nettement plus proche de l'écart mesuré en pratique mais reste quelques dB au dessus, typiquement de l'ordre de 10 dB pour la partie f_e^{-2} et 4 dB pour la partie f_e^{-1} . Cette inadéquation s'explique essentiellement par le fait que lorsqu'on fait la moyenne des enveloppes les écarts au continu pour chacune des enveloppes se compensent généralement en partie et parfois même totalement alors qu'ils sont simplement sommés dans le modèle. De plus, cette compensation est plus importante dans la partie où l'écart se comporte comme f_e^{-2} parce que l'écart correspondant au premier terme de (1.117) va toujours dans le sens de la sous-estimation pour l'enveloppe supérieure et de la surestimation pour l'enveloppe inférieure. Par conséquent, ces écarts ne s'ajoutent jamais mais se compensent toujours au moins partiellement.

5.4.3 Cas de plusieurs itérations de tamisage

Pour étudier comment l'échantillonnage influence l'EMD sur plusieurs itérations de tamisage, on peut s'intéresser comme précédemment à l'écart entre les premiers IMFs issus d'un signal à temps continu et ceux issus de versions discrétisées de ce signal. Dans ce but, on peut par exemple reprendre la procédure d'évaluation précédente à la seule différence que les IMFs $d_\infty[n]$ et $d_{f_e, \varphi}[n]$ ne sont plus restreints au premier IMF et ne sont plus calculés en une unique itération de tamisage mais en un nombre variable dépendant du signal, donc des paramètres d'échantillonnage (f_e, φ) . Les résultats de ces simulations reprenant le signal utilisé Fig. 1.20 (b) sont présentés Fig. 1.21 pour un paramètre d'arrêt $\epsilon = 0.05$ et Fig. 1.22 pour $\epsilon = 0.01$.

Si on s'intéresse tout d'abord au premier IMF, on observe que pour toutes les fréquences d'échantillonnage, il existe le plus souvent des valeurs de phase pour lesquelles le nombre d'itérations de tamisage est identique à celui utilisé pour le signal continu (3 pour Fig. 1.21, 13 pour Fig. 1.22). De plus, dans ces cas, l'écart au continu pour plusieurs itérations est du même ordre que ce qui était observé pour une itération. Tout se passe donc comme si l'essentiel de l'écart provenait de la première itération. Dans la plupart des situations étudiées, on constate de plus que le nombre d'itérations est égal à celui utilisé dans le cas continu lorsque la fréquence d'échantillonnage est suffisamment élevée. En revanche, pour des fréquences d'échantillonnage plus basses, le nombre d'itérations peut différer et aboutir à un écart significativement plus important. Il reste toutefois généralement inférieur à ϵ fois l'amplitude du signal (écart entre les enveloppes divisé par 2, qui vaut approximativement 1 ici), lorsque le nombre d'itérations ne diffère pas trop. Cet écart correspond typiquement à l'amplitude maximale autorisée pour la moyenne du premier IMF et donc à peu près à l'écart entre deux itérations de tamisage lorsque l'itération du processus de tamisage arrive à son terme.

Concernant les IMFs suivants, on observe des propriétés similaires. Comme pour le premier IMF, lorsque les nombres d'itérations utilisés pour calculer l'IMF concerné et les précédents sont identiques au cas continu, l'écart se comporte pratiquement comme pour le premier IMF obtenu avec 1 seule itération. On observe même une diminution de l'écart entre le premier et le deuxième IMF et à nouveau entre le deuxième et le troisième. En revanche, l'écart augmente significativement du troisième au quatrième IMF dans le domaine des grandes fréquences d'échantillonnage. Ces variations d'un IMF à l'autre sont encore inexplicables. Le fait que les écarts restent généralement du même ordre de grandeur que celui observé sur le premier IMF est probablement dû au fait que la première approximation dont sont extraits les IMFs suivants a des extrema naturellement plus écartés que le signal, ce qui fait que les effets d'échantillonnage sont moins prononcés. Cependant, comme la première approximation est la différence entre le signal et le premier IMF, l'écart observé sur le premier IMF se transmet aux IMFs suivants, ce qui justifie qu'ils présentent un écart du même ordre. Si on s'intéresse maintenant aux cas où les nombres d'itérations diffèrent, on observe le même type d'écart que ce qui était observé pour le premier IMF en éventuellement plus prononcé, ce qui peut s'expliquer par l'accumulation d'écarts entre le premier IMF et l'IMF concerné. De fait, la caractéristique la plus importante est peut-être que les situations où les nombres d'itérations diffèrent sont de plus en plus nombreuses à mesure que l'indice de l'IMF augmente et donc que l'écart est de plus en plus imprévisible.

5.4.4 Conclusions sur l'étude de l'influence de l'échantillonnage

Dans les paragraphes précédents, on s'est intéressé à quantifier l'influence de l'échantillonnage sur la décomposition. L'étude a tout d'abord été menée avec succès dans le cas simple d'un signal sinusoïdal, avant de considérer un modèle plus général pour lequel les résultats sont malheureusement plus mitigés.

Dans le cas sinusoïdal, une modélisation précise a permis de retrouver les principales caractéristiques des résultats expérimentaux. Au terme de cette étude, la conclusion la plus importante est

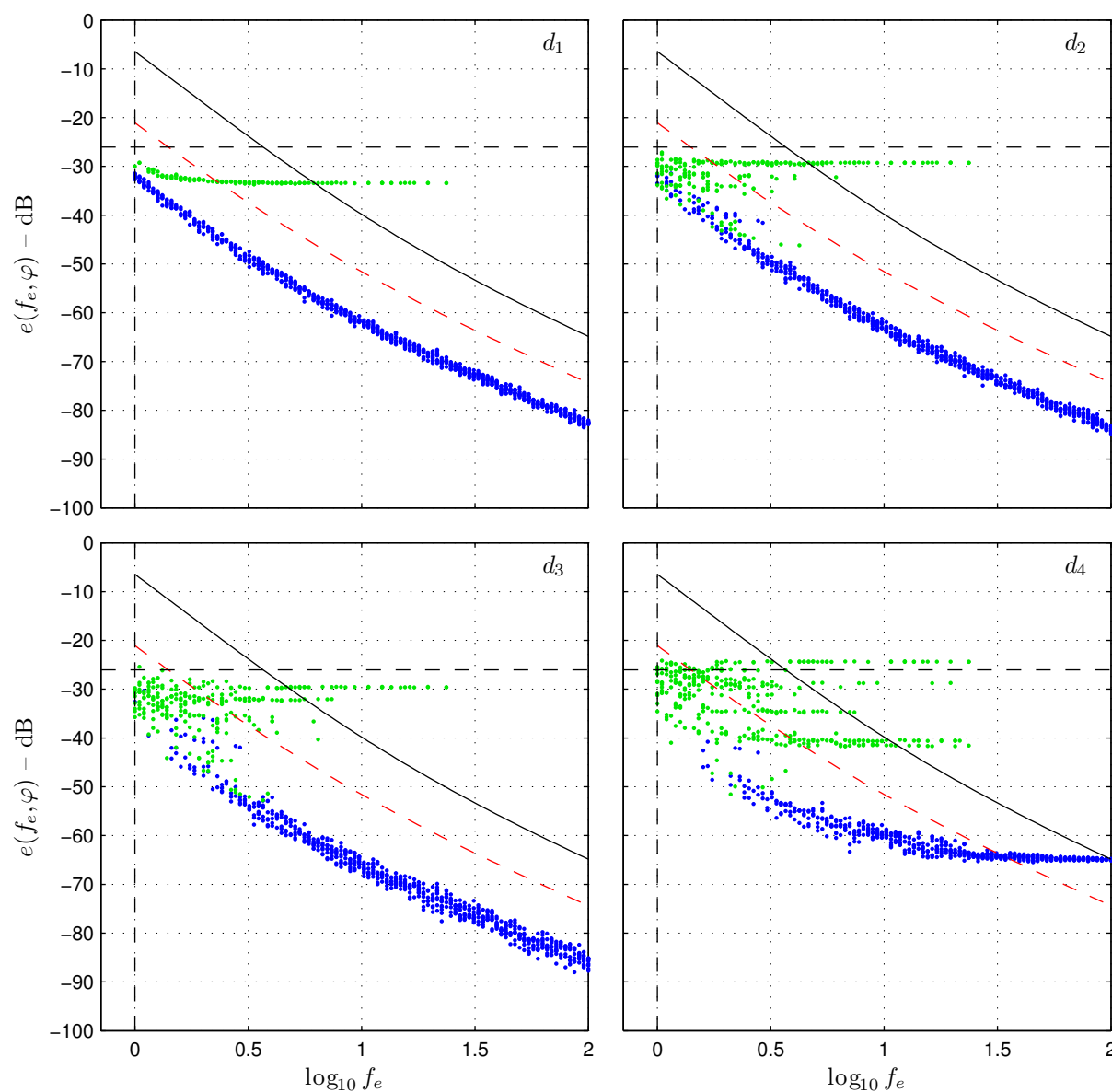


FIGURE 1.21 – Écart au continu pour les 4 premiers IMFs obtenus avec un nombre variable d'itérations déterminé par le critère d'arrêt « local » (cf 5.2) avec $\epsilon = 0.05$. Les cas pour lesquels les nombres d'itérations pour l'IMF concerné et les précédents sont identiques à ceux utilisés pour le signal continu sont représentés par un point bleu foncé ; les autres par un point vert clair. La ligne horizontale en tirets correspond à un écart moyen de l'ordre de ϵ fois l'amplitude du signal (écart entre les enveloppes divisé par 2) qui vaut approximativement 1. Le signal utilisé est le même que pour Fig. 1.20 (b).

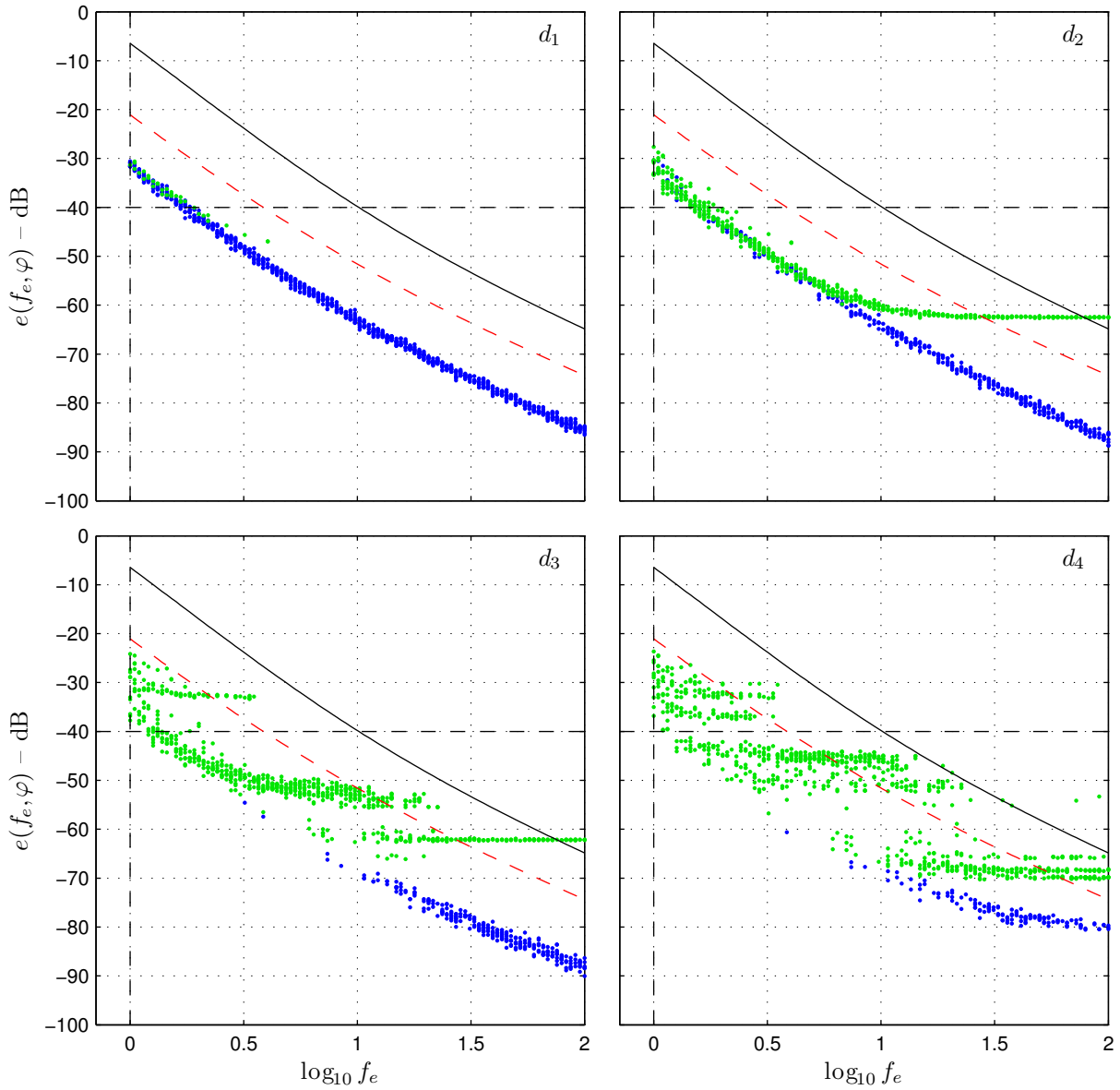


FIGURE 1.22 – Écart au continu pour les 4 premiers IMFs obtenus avec un nombre variable d'itérations déterminé par le critère d'arrêt « local » (cf 5.2) avec $\epsilon = 0.01$. Les cas pour lesquels les nombres d'itérations pour l'IMF concerné et les précédents sont identiques à ceux utilisés pour le signal continu sont représentés par un point bleu foncé ; les autres par un point vert clair. La ligne horizontale en tirets correspond à un écart moyen de l'ordre de ϵ fois l'amplitude du signal (écart entre les enveloppes divisé par 2) qui vaut approximativement 1. Le signal utilisé est le même que pour Fig. 1.20 (b).

sans doute la mise en évidence d'une borne sur l'écart généré par l'échantillonnage qui se comporte comme l'inverse du carré de la fréquence d'échantillonnage. C'est en effet la seule caractéristique observée à être susceptible d'être encore valable si l'on considère une sinusoïde légèrement déformée, par exemple par une modulation de fréquence et/ou d'amplitude.

Dans cet esprit, on s'est par la suite intéressé à un modèle plus polyvalent de signal oscillant dans l'espoir d'étendre le résultat précédent à une situation nettement plus proche du cas général. La première conclusion de cette étude est que l'EMD s'appuyant fortement sur les extrema locaux du signal, l'influence de l'échantillonnage se manifeste de manière très différente si le signal échantillonné comporte autant d'extrema locaux que le signal continu ou si des extrema disparaissent lors de l'échantillonnage. Dans le premier cas, il est a priori possible d'estimer un écart dû à l'échantillonnage dans la mesure où les enveloppes dans les cas continu et échantillonné interpolent des points relativement proches. Dans le second, les ensembles d'extrema du signal continu et du signal échantillonné sont suffisamment différents pour que l'écart soit nettement plus important et bien plus difficile à prévoir. En particulier, il est impossible à prévoir depuis le signal échantillonné uniquement, contrairement au premier cas. En se restreignant à ce premier cas, l'étude a permis d'aboutir à un résultat généralisant celui obtenu dans le cas sinusoïdal, tendant à montrer que dans le cas général, l'écart (borne ou écart moyen) généré par l'échantillonnage sur une itération de tamisage se comporte comme $\lambda f_e^{-1} + \mu f_e^{-2} (+\nu f_e^{-3})$, le dernier terme étant mis entre parenthèse car il semble généralement négligeable. Malheureusement, l'intérêt de ce résultat est limité dans la mesure où quand on se place dans les conditions d'utilisation réelles de l'algorithme, le nombre d'itérations de tamisage est généralement variable et dépend des paramètres de l'échantillonnage. Dès lors les écarts entre les IMFs issus des signaux échantillonnés et de la référence deviennent incontrôlables, même s'il semble qu'ils aient souvent tendance à être limités en ordre de grandeur par le paramètre du critère d'arrêt ϵ lorsqu'on utilise le critère « local ». Ce contrôle très limité est de plus essentiellement restreint aux quelques premiers IMFs, les IMFs d'ordre supérieur étant de plus en plus incontrôlables du fait de l'accumulation d'écarts dus aux IMFs précédents. Si au contraire, on réalise un nombre prédéterminé d'itérations de tamisage dans tous les cas, l'étude n'a été qu'ébauchée ici mais il semble qu'on observe dans ce cas entre les premiers IMFs continus et échantillonnés, un écart proche de celui observé sur une seule itération de tamisage pour le premier IMF. Plus précisément, on observe même un écart légèrement inférieur qui semble décroître avec le nombre d'itérations de tamisage par IMF et avec l'ordre de l'IMF, du moins pour les quelques premiers IMFs. Cette évolution pour le moment inexplicée est peut être due au fait que le processus de tamisage, en rendant les enveloppes des IMFs plus symétriques, rapproche de manière inattendue les IMFs issus du signal continu et du signal échantillonné. Un peu à la manière dont les images de deux points par une projection orthogonale sont nécessairement plus proches que les points eux-mêmes, on peut imaginer que lorsque les IMFs continus et échantillonnés, par l'itération du processus de tamisage, se rapprochent de l'ensemble des signaux vérifiant les propriétés caractéristique des d'un IMF, ils se rapprochent également les uns des autres.

6 Variantes et extensions

6.1 EMD locale

Une caractéristique importante de l'EMD est qu'elle agit à une échelle locale, ce qui lui permet d'être a priori bien adaptée pour traiter des signaux non stationnaires. En particulier, la définition d'un IMF est constituée de deux caractéristiques locales, la seconde dépendant d'une *échelle locale* de l'ordre de l'espacement entre les extrema. En revanche, lorsque le nombre d'itérations de tamisage est déterminé au cours du processus de tamisage, le calcul des IMFs présente une caractéristique peu locale dans la mesure où le tamisage est appliqué sur *toute la durée* du signal tant qu'il existe une

zone où la définition d'IMF⁷ n'est pas vérifiée. Bien sûr, les effets du tamisage sont plus importants là où la moyenne des enveloppes est plus importante, ce qui fait que le tamisage agit essentiellement là où il y en a besoin. Cependant, appliquer un grand nombre d'itérations de tamisage là où la moyenne est faible n'est pas non plus sans effet. L'effet le plus visible est que les enveloppes deviennent de plus en plus lisses au fur et à mesure qu'on itère le processus de tamisage, au risque de perdre l'information contenue dans l'évolution temporelle de l'amplitude de l'enveloppe [28]. De plus, cet effet s'accompagne d'une augmentation du nombre d'IMFs nécessaire pour reconstituer les évolutions rapides des enveloppes, un peu à la manière dont une modulation d'amplitude sinusoïdale est décomposée en trois composantes de Fourier aux amplitudes constantes.

D'un point de vue général, il est difficile d'établir si une décomposition avec plus d'IMFs aux enveloppes plus lisses est plus ou moins bonne qu'une autre avec moins d'IMFs et des enveloppes plus variables. De fait, il est généralement déconseillé d'itérer le tamisage jusqu'à obtenir des enveloppes trop lisses mais les raisons ne sont pas toujours très claires et semblent s'appuyer notamment sur l'hypothèse selon laquelle l'EMD permettrait d'extraire des composantes pertinentes d'un point de vue physique, ou sur l'idée voisine que l'EMD doit expliquer le contenu du signal avec aussi peu de composantes que possible. En pratique, avec l'expérience acquise au cours de cette thèse, il semble que la seule raison valable — c'est-à-dire non basée sur un a priori discutable — pour justifier qu'on limite l'itération du processus de tamisage de manière à ne pas trop lisser les enveloppes, soit que celles-ci sont calculées de manière approchée et que sommer un trop grand nombre de ces estimations approchées finit par polluer le contenu du signal analysé. Cependant, il est très difficile de justifier rigoureusement un tel argument puisque dans tous les cas la somme des IMFs et du résidu reconstitue le signal initial et qu'on ne trouve jamais un sous-ensemble d'IMFs dont la somme est nulle, cas où on aurait pu indiscutablement dire que ces IMFs n'étaient pas pertinents.

En revanche, ce dont on peut être sûr, c'est qu'il n'est pas toujours souhaitable d'obtenir une décomposition avec des enveloppes très lisses sur une partie du signal, juste parce que la méthode utilisée pour déterminer le nombre d'itérations de tamisage a décidé qu'il fallait en faire un grand nombre pour qu'une *autre partie* du signal satisfasse correctement à la définition d'IMF retenue. Pour remédier à ce problème, on a proposé [55] une variante de l'algorithme de l'EMD où le tamisage est appliqué *seulement* là où il y en a besoin. Dans sa version actuelle, disponible sur internet, cette variante est liée à la méthode de détermination des nombres d'itérations « locale » (cf 5.2) mais on pourrait étendre le principe à d'autres méthodes dès lors qu'elles sont en lien avec une mesure *locale* de la qualité d'un IMF, c'est-à-dire une fonction binaire qui en chaque instant détermine si le signal en cours de traitement est localement un IMF ou non. Dans la version actuelle, cette fonction est de la forme

$$s(t) = \begin{cases} 1 & \text{si } \left| \frac{e_{\max}(t) + e_{\min}(t)}{e_{\max}(t) - e_{\min}(t)} \right| < \epsilon \\ 0 & \text{si } \left| \frac{e_{\max}(t) + e_{\min}(t)}{e_{\max}(t) - e_{\min}(t)} \right| \geq \epsilon \end{cases} \quad (1.127)$$

où $e_{\max}(t)$ et $e_{\min}(t)$ désignent respectivement les enveloppes supérieure et inférieure du signal et ϵ est un paramètre. De là, l'algorithme de l'« EMD locale » est alors le même que celui de l'EMD originale avec un opérateur de tamisage modifié (cf Algo. 2).

Dans cette définition, on a pris soin d'étendre les zones où la fonction d'évaluation $s(t)$ vaut 1 par une fonction « douce ». Il y a deux raisons à cette précaution. Premièrement, le fait d'adoucir les bords des zones où $s(t) = 1$ permet d'éviter de soustraire au signal une fonction discontinue, ce qui amènerait des discontinuités dans l'IMF et dans l'approximation qui pourraient finir par créer des extrema parasites. Le fait d'adoucir les bords permet a priori de limiter l'apparition d'extrema parasites mais il peut rester des effets indésirables. En particulier, la manière dont $f(t)$ décroît vers

7. Cette définition peut varier suivant la méthode retenue pour déterminer le nombre d'itérations (cf 5.2) mais elle implique toujours au moins la condition sur le signe des extrema.

Algorithme 2 : Opérateur de tamisage local

-
- 1 Extraire les maxima et les minima de $x(t)$: $\{t_i^{max}, x_i^{max}\}, \{t_i^{min}, x_i^{min}\}$.
 - 2 Interpoler les ensembles de maxima $\{t_i^{max}, x_i^{max}\}$ et les ensembles de minima $\{t_i^{min}, x_i^{min}\}$ pour obtenir les enveloppes supérieure et inférieure : $e_{max}(t), e_{min}(t)$.
 - 3 Calculer la moyenne des enveloppes : $m(t) = (e_{max}(t) + e_{min}(t)) / 2$
 - 4 **Calculer la fonction d'évaluation $s(t)$ (1.127).**
 - 5 **Créer une fonction positive $f(t)$ qui vaut 1 là où $s(t)$ vaut 1 et qui descend doucement vers zéro autour.**
 - 6 Soustraire la moyenne **multipliée par $f(t)$** au signal : $S^{local}[x](t) = x(t) - m(t)f(t)$
-

zéro peut influencer la forme des IMFs suivants surtout s'ils ont une amplitude plus faible ou de l'ordre de ϵ fois l'amplitude de l'IMF en cours. Deuxièmement, après les quelques premières itérations de tamisage la moyenne des enveloppes oscille souvent autour de zéro, ce qui fait que dans les zones où elle n'est pas assez proche de zéro, la fonction d'évaluation $s(t)$ s'annule en fait localement autour de chaque passage à zéro et prend la valeur 1 autour des extrema de la moyenne. Dans ces conditions, il semble préférable que $f(t)$ vaille 1 uniformément sur tout l'intervalle pour limiter les risques de contamination des IMFs suivants.

En pratique, dans la version actuelle du programme, l'extension de $s(t)$ à $f(t)$ est faite par une fonction linéaire par morceaux dont les paramètres sont déterminés en fonction de la « période locale » du signal, évaluée par l'écart entre les extrema. Le résultat semble fonctionner très correctement. Toutefois, il faut noter qu'aucune étude n'a jamais été réalisée sur les performances de cet algorithme modifié ce qui fait qu'il est difficile d'évaluer la confiance qu'on peut attribuer à ses résultats. De plus, la manière de définir $f(t)$ à partir de $s(t)$ mériterait sans doute une étude plus poussée. La méthode utilisée actuellement a été testée sur un certain nombre d'exemples et ne produit apparemment pas d'artefacts visibles mais il ne fait aucun doute qu'il y aurait des bénéfices à tirer, par exemple du remplacement de la fonction linéaire par morceaux $f(t)$ par une fonction plus régulière.

6.2 EMD en ligne

Partant de l'algorithme de l'EMD locale décrit précédemment, il vient l'idée que si on peut appliquer le tamisage localement, on peut aussi le faire progressivement en partant d'un bout du signal pour arriver jusqu'à l'autre. Bien entendu, l'intérêt d'un tel point de vue est limité pour un signal mesuré préalablement, c'est-à-dire dont on connaît déjà tous les échantillons (en admettant que le signal ait un début et une fin). En revanche, cela peut s'avérer particulièrement intéressant pour un signal en cours de mesure où les échantillons arrivent au fur et à mesure.

En pratique, l'algorithme d'EMD en ligne est plus compliqué que la simple idée d'appliquer le tamisage localement de façon progressive. Il est en fait composé de deux parties dont la deuxième correspond à l'idée de tamisage local progressif et la première peut être vue comme un prétraitement permettant d'appliquer ensuite la seconde. Ces deux parties sont appelées par la suite *tamisage par blocs* et *tamisage local en ligne*.

6.2.1 Tamisage par blocs

Comme son nom l'indique, le tamisage par blocs consiste essentiellement à découper le signal en sous-intervalles, appliquer le tamisage sur chacun de ces sous-intervalles puis recoller ensuite les morceaux.

Plus précisément, le tamisage par blocs est basé sur le fait que la valeur d'une enveloppe en un point dépend essentiellement des positions des extrema proches du point considéré. Dans le cas d'une interpolation spline cubique, l'influence d'un extremum décroît exponentiellement avec la distance

au point considéré, exprimée par exemple en termes de nombre de nœuds. Par conséquent, si au lieu de tenir compte de tous les extrema pour calculer la valeur de l'enveloppe en un point donné, on n'en considère que n de part et d'autre de ce point, on peut aboutir à une valeur très proche de celle qu'on obtiendrait avec tous les extrema. De plus, il est possible d'ajuster la précision d'un tel calcul en prenant n plus ou moins grand. Typiquement, pour le calcul des enveloppes d'un bruit blanc gaussien de variance 1, on observe que l'écart entre l'enveloppe calculée avec tous les extrema et celle calculée sur un intervalle entre deux extrema en ne tenant compte que de n extrema de part et d'autre est de l'ordre de $10^{-n/2}$, ce qui permet d'obtenir de très bonnes estimations des enveloppes en ne considérant qu'un nombre restreint d'extrema de part et d'autre. En pratique, il peut s'avérer nécessaire d'adapter la valeur de n au signal analysé mais le fait que la décroissance soit a priori toujours exponentielle permet d'assurer qu'on trouvera toujours une valeur de n satisfaisante qui ne soit pas trop grande.

Une fois choisie une valeur pour n , le tamisage par morceaux est réalisé de la manière suivante. Soit $d_k^{(i)}(t)$ le k^e IMF en cours d'élaboration après i itérations de tamisage. Si on connaît $d_k^{(i)}(t)$ sur un intervalle $[a, b]$ et qu'il présente suffisamment d'extrema sur $[a, b]$, on peut calculer $d_k^{(i+1)}(t)$ sur $[a + \epsilon_1, b - \epsilon_2]$, avec ϵ_1 et ϵ_2 choisis de telle manière que $d_k^{(i)}(t)$ ait au moins n extrema dans $[a, a + \epsilon_1]$ et $[b - \epsilon_2, b]$. Si par la suite, on connaît $d_k^{(i)}(t)$ sur une durée plus longue $[a, c]$, $c > b$, on pourra calculer $d_k^{(i+1)}(t)$ sur $[b - \epsilon_2, c - \epsilon_3]$ et recoller presque parfaitement avec la partie déjà connue sur $[a + \epsilon_1, b - \epsilon_2]$. Entre-temps, ce qui est déjà calculé de $d_k^{(i+1)}(t)$ sur $[a + \epsilon_1, b - \epsilon_2]$ peut être utilisé pour calculer $d_k^{(i+2)}(t)$ jusque $b - \epsilon_2 - \epsilon_4$.

Avec cette méthode, on peut appliquer un nombre prédéterminé d'itérations de tamisage de manière progressive et donc extraire des IMFs progressivement. De plus, une fois qu'une partie de l'IMF en cours d'élaboration a subi toutes les itérations prévues, celui-ci n'est plus modifié et peut donc être soustrait au signal (ou à l'approximation précédente) pour commencer à extraire l'IMF suivant. On peut donc ainsi extraire plusieurs IMFs simultanément avec juste pour chaque IMF un retard sur le précédent.

En revanche, si on souhaite appliquer un nombre variable d'itérations de tamisage déterminé au cours du processus de tamisage, on ne peut le faire que lorsque le signal est connu sur toute sa durée et, pour les IMFs autres que le premier, que lorsque tous les IMFs précédents ont été calculés complètement. La solution est alors d'adapter le principe de l'EMD locale pour avoir de fait des nombres d'itérations de tamisage qui varient en fonction de la position et qui du coup ne dépendent plus de toute la durée du signal.

6.2.2 Tamisage local en ligne

Le tamisage local en ligne (cf Algo. 3) est obtenu en ajoutant une étape à l'opérateur de tamisage local Algo. 2. L'opérateur de tamisage ainsi défini réalise de fait un tamisage local sur une fenêtre glissante $w(t)$ non nulle sur un intervalle $[a, b]$. L'idée est ensuite de faire progresser la fenêtre sur le signal de telle manière qu'après son passage, l'IMF en cours d'élaboration soit considéré comme localement achevé : $s(t) = 0, t < a$. Pour atteindre cet objectif, on fait avancer l'avant b et l'arrière a de la fenêtre séparément. L'avant de la fenêtre avance en fonction des données disponibles : si le signal est connu jusqu'à un instant t , on positionne l'avant de la fenêtre en $t - \epsilon$ de telle manière que le signal ait n extrema dans $[t - \epsilon, t]$ pour pouvoir calculer les enveloppes de manière satisfaisante. Le cas de l'arrière de la fenêtre est plus compliqué. L'idée est de le faire avancer quand le tamisage local est terminé ($s(t) = 0$) sur un intervalle $[a, c]$, $c < b$ suffisamment grand devant la période locale mesurée par l'écart entre les extrema. On fait alors avancer l'arrière de la fenêtre a jusqu'à c ou, par précaution, seulement jusqu'à $(a + c)/2$. Le problème de cette méthode est que sachant que la fonction $f(t)$ est non nulle sur un ensemble plus grand que celui où $s(t)$ est non nulle, il est tout à fait possible qu'en appliquant l'opérateur de tamisage local en ligne, une zone où $s(t)$ vaut 1 s'étende de part et

d'autre, par exemple du fait de l'apparition d'une paire d'extrema. Par conséquent, si $s(t)$ est nulle sur un intervalle $[a, c]$ à une itération donnée, rien ne dit qu'elle le sera encore à l'itération suivante si elle est non nulle à proximité des frontières de $[a, c]$. Il faut donc faire très attention lorsqu'on fait avancer l'arrière de la fenêtre d'observation, raison pour laquelle on suggère la précaution de ne le faire avancer que jusqu'à $(a + c)/2$ si $s(t) = 0$ sur $[a, c]$.

Algorithme 3 : Opérateur de tamisage local en ligne

- 1 Extraire les maxima et les minima de $x(t)$: $\{t_i^{max}, x_i^{max}\}, \{t_i^{min}, x_i^{min}\}$.
 - 2 Interpoler les ensembles de maxima $\{t_i^{max}, x_i^{max}\}$ et les ensembles de minima $\{t_i^{min}, x_i^{min}\}$ pour obtenir les enveloppes supérieure et inférieure : $e_{max}(t), e_{min}(t)$.
 - 3 Calculer la moyenne des enveloppes : $m(t) = (e_{max}(t) + e_{min}(t)) / 2$
 - 4 Calculer la fonction d'évaluation $s(t)$ (1.127).
 - 5 Créer une fonction positive $f(t)$ qui vaut 1 là où $s(t)$ vaut 1 et qui descend doucement vers zéro autour.
 - 6 **Multiplier par une fenêtre glissante $w(t)$ qui parcourt la durée du signal :**
 $f_2(t) = f(t)w(t)$.
 - 7 Soustraire la moyenne **multipliée par $f_2(t)$** au signal : $\mathcal{S}^{online}[x](t) = x(t) - m(t)f_2(t)$
-

Le tamisage local en ligne ainsi défini permet d'extraire un IMF en parcourant le signal. Sachant de plus que l'IMF n'est plus modifié après le passage de la fenêtre, il est possible d'extraire plusieurs IMFs simultanément avec un retard entre chaque.

Le problème de ce procédé est qu'il ne fonctionne en fait bien que quand le signal d'origine a déjà une moyenne proche de zéro. En effet, s'il présente au contraire une moyenne grande par rapport à son amplitude, le fait de lui soustraire la moyenne de ses enveloppes sur la partie avant de la fenêtre $w(t)$ peut poser de sérieux problèmes. Si on suppose que $w(t)$ peut être divisée en trois parties $[a, b]$, $[b, c]$ et $[c, d]$, où $w(t) = 1$ sur $[b, c]$ et décroît doucement vers zéro sur $[a, b]$ et $[c, d]$, le fait de soustraire la moyenne des enveloppes multipliée par $w(t)$ fait qu'après une itération de tamisage local en ligne, l'IMF en cours d'élaboration a une moyenne grande en d et proche de zéro en c . Dans ce cas, on perd généralement une grande partie des extrema initialement présents sur $[c, d]$ et le résultat final de la décomposition est alors très mauvais.

Pour éviter ce problème, il suffit de combiner les deux procédés du tamisage par blocs et du tamisage local en ligne. L'idée est de faire un petit nombre prédéterminé d'itérations de tamisage par blocs et seulement après d'appliquer le tamisage local en ligne. Le tamisage par blocs permet de ramener la moyenne des enveloppes pas trop loin de zéro pour qu'il soit ensuite possible d'appliquer le tamisage local en ligne. En pratique, il suffit généralement de très peu d'itérations de tamisage par blocs (typiquement 1 ou 2) pour ramener la moyenne suffisamment près de zéro.

6.2.3 Commentaires

Tout comme l'EMD locale, la version actuelle de l'EMD en ligne, disponible sur internet, est encore au stade de prototype. En particulier, le programme actuel n'est pas conçu pour être appliqué à un flot de données comme il le devrait : il s'applique en fait à des signaux complets comme l'EMD standard et se contente de simuler en interne une arrivée progressive des données. De plus, un certain nombre de choix ont été réalisés pour construire le programme mais ils n'ont pas été étudiés de manière approfondie. On s'est contenté en général de vérifier sur quelques exemples que ces choix ne créaient pas d'artefacts suffisamment importants pour être visibles à l'œil. Parmi ces choix, on peut citer

- le nombre d'extrema n de part et d'autres d'un point donné considéré comme suffisant pour estimer la valeur des enveloppes. Sa valeur est actuellement de 5 par défaut dans le programme.

Pour les enveloppes d'un bruit blanc gaussien de variance 1, ce nombre conduit à un écart par rapport à l'enveloppe calculée avec tous les extrema de l'ordre de 10^{-3} , ce qui n'est pas forcément négligeable.

- la forme de la fenêtre $w(t)$, actuellement linéaire par morceaux.
- les règles utilisées pour décider comment on fait progresser l'avant et l'arrière de la fenêtre.

Finalement, le programme actuel permet de faire la démonstration qu'il est possible d'appliquer le tamisage de manière progressive sur un flot de données. Cependant, si on compare ses résultats avec l'EMD classique ou l'EMD locale, on observe généralement que les quelques premiers IMFs sont proches mais que les autres peuvent être assez différents. Étant donnée la complexité de l'algorithme en ligne par rapport à l'EMD classique, ou même locale, ainsi que le fait que ses paramètres ont été déterminés empiriquement sans étude rigoureuse, il est clair que les IMFs calculés avec cet algorithme doivent être considérés avec beaucoup de précautions.

6.3 « Ensemble Empirical Mode Decomposition »

6.3.1 Principe

Autre variante de l'EMD, très différente des deux proposées précédemment, l'EMD d'ensemble, ou EEMD pour « Ensemble Empirical Mode Decomposition » a été introduite [65] initialement pour limiter les effets de mélanges de modes. Étant donné un signal $x(t)$, le principe est le suivant :

1. on génère N réalisations $n_i(t)$, $1 \leq i \leq N$ de bruit blanc gaussien de variance σ^2
2. on calcule N jeux d'IMFs $d_k^{(i)}(t)$, $1 \leq i \leq N$ à partir des N signaux $x(t) + n_i(t)$
3. les IMFs EEMD sont alors les moyennes d'ensemble des IMFs précédents

$$d_k^{EEMD}(t) = \frac{1}{N} \sum_{i=1}^N d_k^{(i)}(t). \quad (1.128)$$

Si $\sigma/\sqrt{N}$ est petit devant l'échelle des variations du signal $x(t)$, l'idée est alors de considérer que la dépendance des IMFs EEMD par rapport aux réalisations de bruit mises en jeu devient négligeable et que les IMFs EEMD ainsi définis ne dépendent donc que du signal $x(t)$ et de σ .

Cette méthode a en fait été initialement proposée dans [20] pour permettre de calculer l'EMD d'une impulsion, le problème étant qu'un tel signal ne présente pas assez d'extrema pour pouvoir être traité par l'EMD classique. De fait, le résultat obtenu dépend de l'amplitude du bruit utilisé et ne converge pas quand celle-ci tend vers zéro ce qui fait qu'on ne peut pas vraiment parler de l'EMD d'une impulsion, mais qu'on a au contraire affaire à une variante de l'EMD.

Dans une perspective très différente, l'auteur de [65] a vu dans cette méthode un moyen de supprimer certains effets de mélanges de modes apparaissant lorsque le signal contient une composante intermittente, c'est-à-dire dont l'amplitude est parfois nulle, parfois non nulle. Dans ces conditions, comme on l'a vu en 3.5, pour peu qu'il y ait une composante plus basse fréquence dans le signal, l'EMD produit typiquement un IMF mélangeant la composante intermittente avec la composante plus basse fréquence alors que les évolutions de cette dernière là où la composante intermittente est non nulle sont reléguées dans l'IMF suivant. Pour supprimer ce phénomène, une solution simple consiste à séparer les deux composantes par un filtrage linéaire. Le principe qui fait que l'EEMD peut se révéler efficace dans ces conditions n'est en fait pas très éloigné de l'idée d'un filtrage linéaire. Il est basé sur le fait que l'action de l'EMD sur un bruit large bande s'apparente à celle d'un banc de filtres quasi-dyadique (cf 2 au chapitre 2). De là, l'idée est que si on ajoute un bruit de grande amplitude (par rapport aux amplitudes des composantes) au signal et qu'on calcule l'EMD on peut supposer a priori que c'est essentiellement le bruit qui pilote la décomposition et que les composantes vont donc se retrouver, noyées dans le bruit, dans l'IMF dont la bande de fréquence leur correspond le mieux. Si ces bandes sont différentes pour les deux composantes et que celles-ci se retrouvent toujours dans

les mêmes IMFs quelle que soit la réalisation de bruit, les IMFs EEMD contiennent alors les deux composantes isolées dans deux IMFs différents.

Dans cette description, il faut noter que les notions de fréquence d'une composante et de bande de fréquence d'un IMF sont ici très différentes de leur significations habituelles liées à la transformée de Fourier. Pour l'EMD, la notion de fréquence est en fait locale et se rapporte plus ou moins à l'écartement entre les extrema. Une conséquence importante de cette différence de point de vue est que l'EMD ne fait en pratique pas vraiment de différence entre un signal stationnaire et un signal non stationnaire, contrairement aux méthodes basées sur la transformée de Fourier.

6.3.2 Questions relatives à l'implantation

Lors de l'étape finale de l'EEMD, on définit les IMFs comme les moyennes d'ensemble des IMFs bruités. En pratique, cependant, on obtient généralement des nombres d'IMFs qui dépendent de la réalisation. De plus, si le premier IMF correspond environ à la bande de fréquences $[0.25, 0.5]$, le dernier contient en général 2 à 3 extrema et correspond donc lui aussi à une certaine bande de fréquence. Par conséquent, si on considère deux réalisations donnant des nombres d'IMFs différents, il n'est pas du tout sûr qu'on somme des quantités cohérentes lorsqu'on somme deux IMFs de même indice, surtout si l'indice est grand. Pour contourner ce problème, on peut envisager différentes solutions dont aucune n'est totalement satisfaisante :

- ne définir l' IMF EEMD d'indice k que quand toutes les réalisations ont fourni au moins k IMFs, résidu compris ou non.
- étendre artificiellement tous les jeux d'IMFs en rajoutant des IMFs nuls après le dernier de manière à avoir les mêmes nombres d'IMFs pour toutes les réalisations.
- réduire artificiellement tous les jeux d'IMFs à un nombre d'IMFs inférieur ou égal au nombre d'IMFs minimal en sommant les derniers IMFs surnuméraires.

Une autre question est celle du choix de la méthode utilisée pour déterminer les nombres d'itérations. Comme on le verra en 2 du chapitre 2, les nombres d'itérations sont importants parce que les caractéristiques du banc de filtres équivalent à l'EMD dépendent essentiellement de ce paramètre. Dans le cas de bruits (sans signal supplémentaire), on constate cela dit qu'on peut obtenir avec des nombres d'itérations variables, un banc de filtres équivalent quasiment identique à celui obtenu avec des nombres d'itérations fixes. Toutefois, rien ne garantit que les nombres d'itérations ne seront pas très différents si on ajoute un signal au bruit. Pour cette raison, dans le cadre de l'EEMD, on utilise en général des nombres d'itérations fixés à l'avance par précaution. De plus, les IMFs EEMD ne vérifiant pas forcément la définition d'un IMF, l'intérêt que l'EMD de chaque réalisation aboutisse à des IMFs bien formés paraît plutôt limité.

6.3.3 Quelques exemples d'applications

Le premier exemple (cf Fig. 1.23) correspond à la situation décrite précédemment pour expliquer l'intérêt de la méthode. Le signal est constitué de deux composantes : une sinusoïde d'amplitude constante et une de fréquence supérieure modulée en amplitude par une gaussienne. Dans le cas particulier considéré, on observe que les deux composantes sont séparées en plusieurs IMFs EEMD mais la composante haute fréquence est étrangement répartie sur deux IMFs. L'origine de ce fait est simplement que la composante haute fréquence ne se trouve pas toujours dans le même IMF. Dans le cas particulier présenté, il semble que la composante haute fréquence soit grossièrement dans le 4^e IMF pour la moitié des réalisations et dans le 5^e pour l'autre moitié. De plus, si on observe les IMFs d'un peu plus près, on peut aussi voir que le 7^e IMF contient aussi une fraction de la composante basse fréquence. De manière générale, on voit sur cet exemple que le fait d'ajouter du bruit au signal impose d'une certaine manière une grille quasi-dyadique en fréquence qui n'est pas nécessairement adaptée au signal. De plus, il semble que l'appartenance d'une composante à une case de la grille soit

toujours un minimum aléatoire, ce qui aboutit au fait qu'on retrouve souvent au moins une faible partie de chaque composante dans les IMFs adjacents à ceux où elles se trouvent pour la majorité des réalisations.

Le second exemple est identique au premier à ceci près que la composante basse fréquence n'est plus sinusoïdale mais triangulaire. Étant donnée la nature a priori très différente d'un filtrage linéaire de l'EEMD, on aurait pu penser qu'elle serait capable comme l'EMD d'extraire des composantes non linéaires. En pratique, force est de constater que ses capacités non linéaires sont en fait limitées puisqu'elle décompose le signal triangulaire en plusieurs composantes quasi-sinusoïdales.

Enfin, les deux derniers exemples mettent en évidence une limitation de l'EEMD par rapport aux signaux non stationnaires. Ces deux exemples sont d'une part une modulation de fréquence linéaire et d'autre part une modulation d'amplitude quadratique. Dans les deux cas, l'EEMD répartit le signal sur plusieurs composantes.

6.4 EMD pour les images

Cette extension de l'EMD n'a pas du tout été étudiée au cours de cette thèse mais elle suscite un certain enthousiasme de la part de la communauté. Les problématiques relatives à l'EMD pour les images étant très différentes du cas de l'EMD classique, on ne les présentera pas ici. Le lecteur intéressé est renvoyé aux travaux de Linderhed [38, 39], Nunes [44, 43], Liu [41, 40], Damerval [11] et Xu [66].

6.5 Signal bivarié

L'EMD a également été étendue au cas de signaux bivariés, ie à valeur dans $\mathbb{R}^2$, ou dans $\mathbb{C}$. Il faut préciser cependant qu'une telle EMD ne peut avoir de sens que si les deux composantes $x_1(t)$ et $x_2(t)$ du signal analysé peuvent être assimilées aux coordonnées cartésiennes d'un point mobile au cours du temps. Ainsi, un signal bivarié donnant par exemple la direction et la vitesse du vent au cours du temps pourra a priori être adéquatement traité par une EMD bivariée — en passant des coordonnées polaires aux coordonnées cartésiennes — mais la décomposition d'un signal dont les composantes correspondent par exemple aux potentiels de deux électrodes placées à des endroits différents (même proches) d'un cerveau aurait peu de sens. Cette contrainte est simplement liée au fait que toutes les approches proposées remplacent de fait la notion d'« oscillation » par la notion de « rotation » dans le principe de l'EMD et que la notion de rotation sous-entend l'existence d'un certain objet tournant dans un espace à au moins deux dimensions. Ainsi, là où l'EMD classique décompose un signal (scalaire) en une « composante oscillant rapidement » et une « composante oscillant lentement », le principe des extensions bivariées est essentiellement de décomposer un signal bivarié en une « composante tournant rapidement » et une « composante tournant lentement ». De fait, le point de vue est très similaire à celui adopté par Ptolémée dans sa description du mouvement des planètes : on décompose un mouvement compliqué en une superposition de rotations de plus ou moins longue période. Deux approches ont été proposées :

1. analyser séparément les composantes analytiques et antianalytiques du signal à l'aide de l'EMD classique appliquée à leurs parties réelles. Les IMFs bivariés sont alors les parties analytiques des IMFs réels obtenus à partir de la composante analytique et les parties antianalytiques des autres.
2. construire un algorithme bivarié en transposant directement le processus de tamisage au cadre bivarié.

Dans la suite, on présentera tout d'abord succinctement la première approche en précisant les avantages et limites puis on s'intéressera de plus près à la deuxième approche. Les signaux bivariés seront souvent assimilés à des signaux à valeur complexe pour simplifier les écritures.

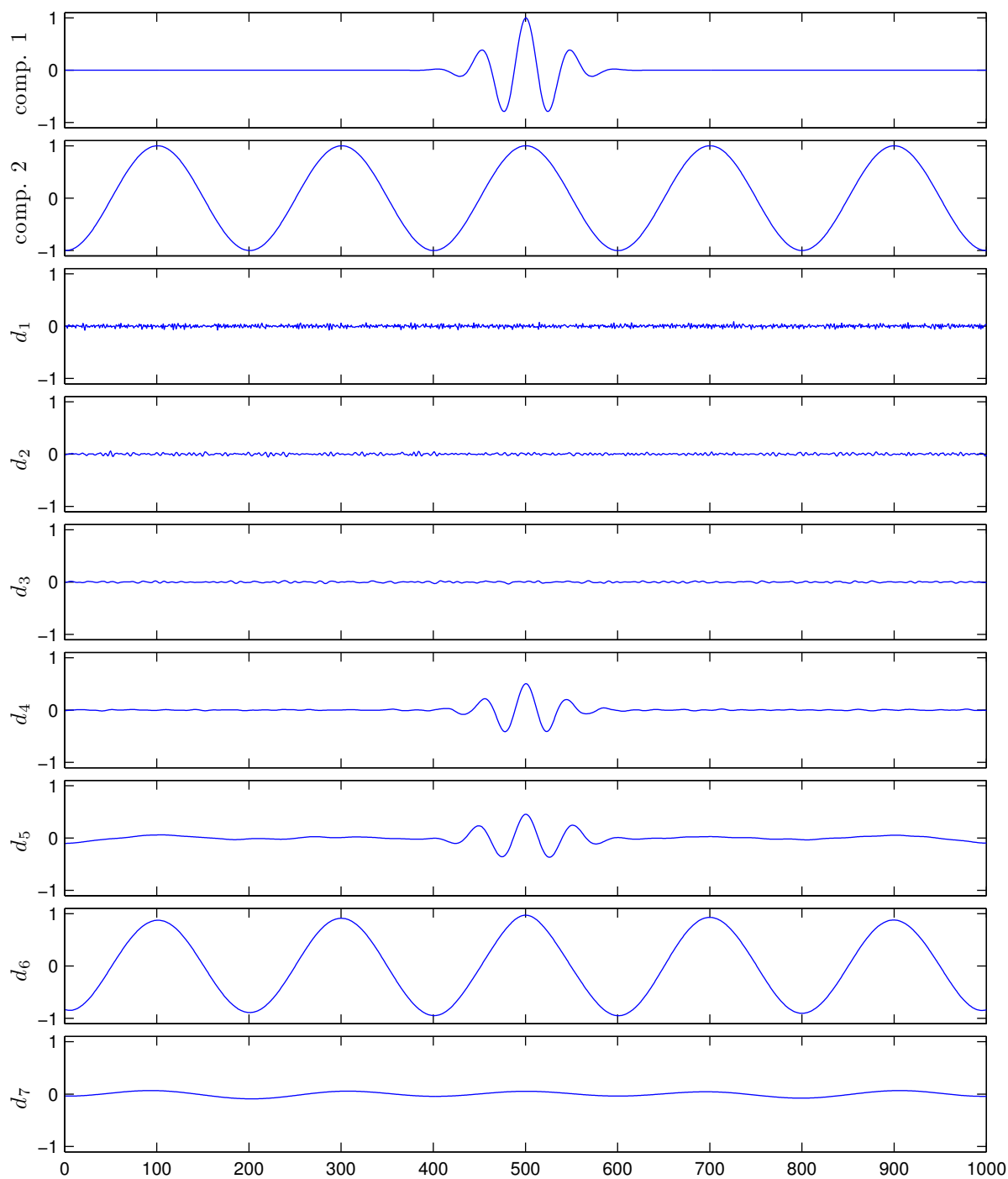


FIGURE 1.23 – Exemple d'application de l'EEMD. Cas d'une composante sinusoïdale et d'une composante intermittente. Les IMFs EEMD ont été calculés à l'aide de 10000 réalisations de bruit de variance 1.

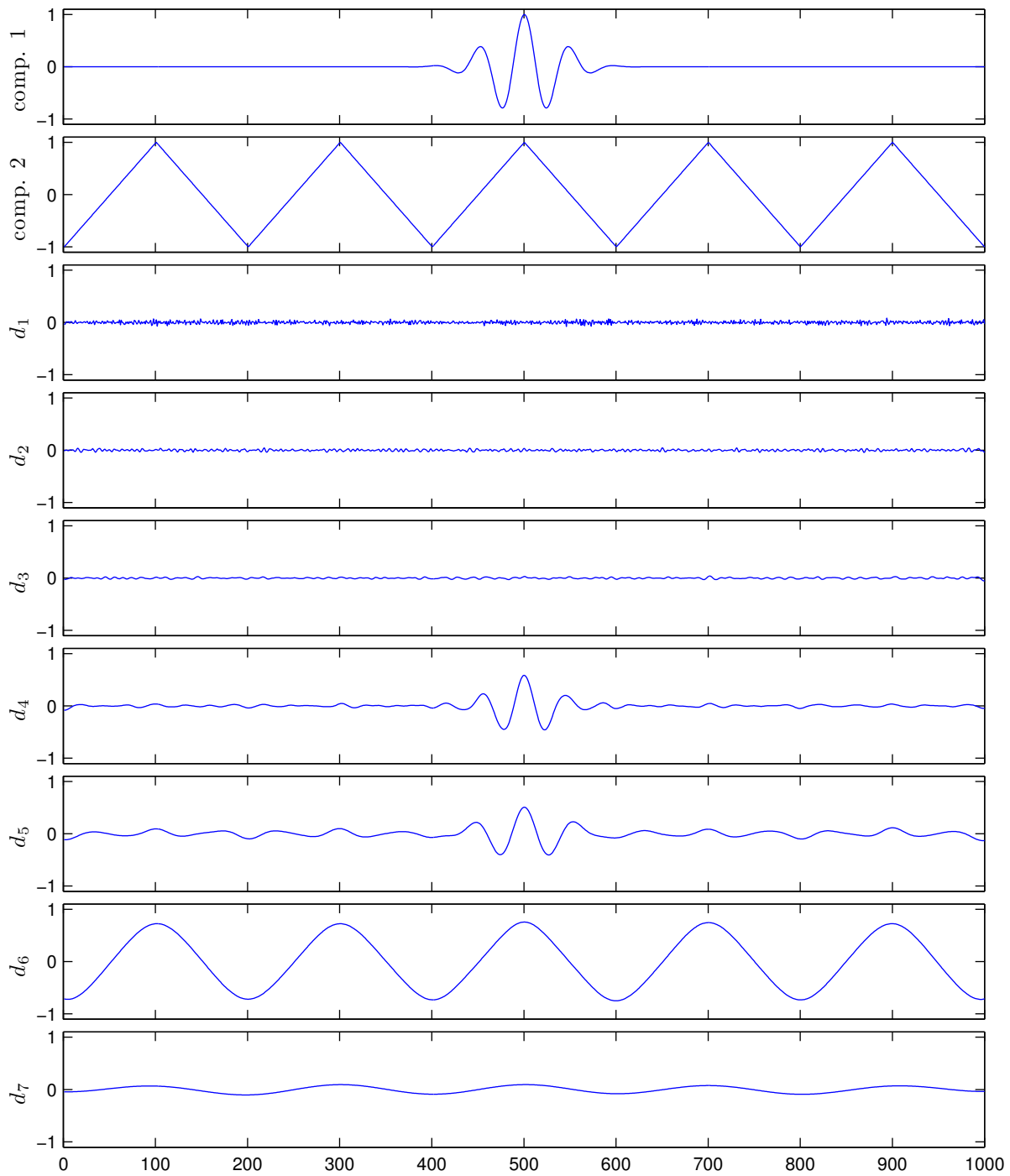


FIGURE 1.24 – Exemple d’application de l’EEMD. Cas d’une composante sinusoïdale et d’une composante triangulaire. Les IMFs EEMD ont été calculés à l’aide de 10000 réalisations de bruit de variance 1.

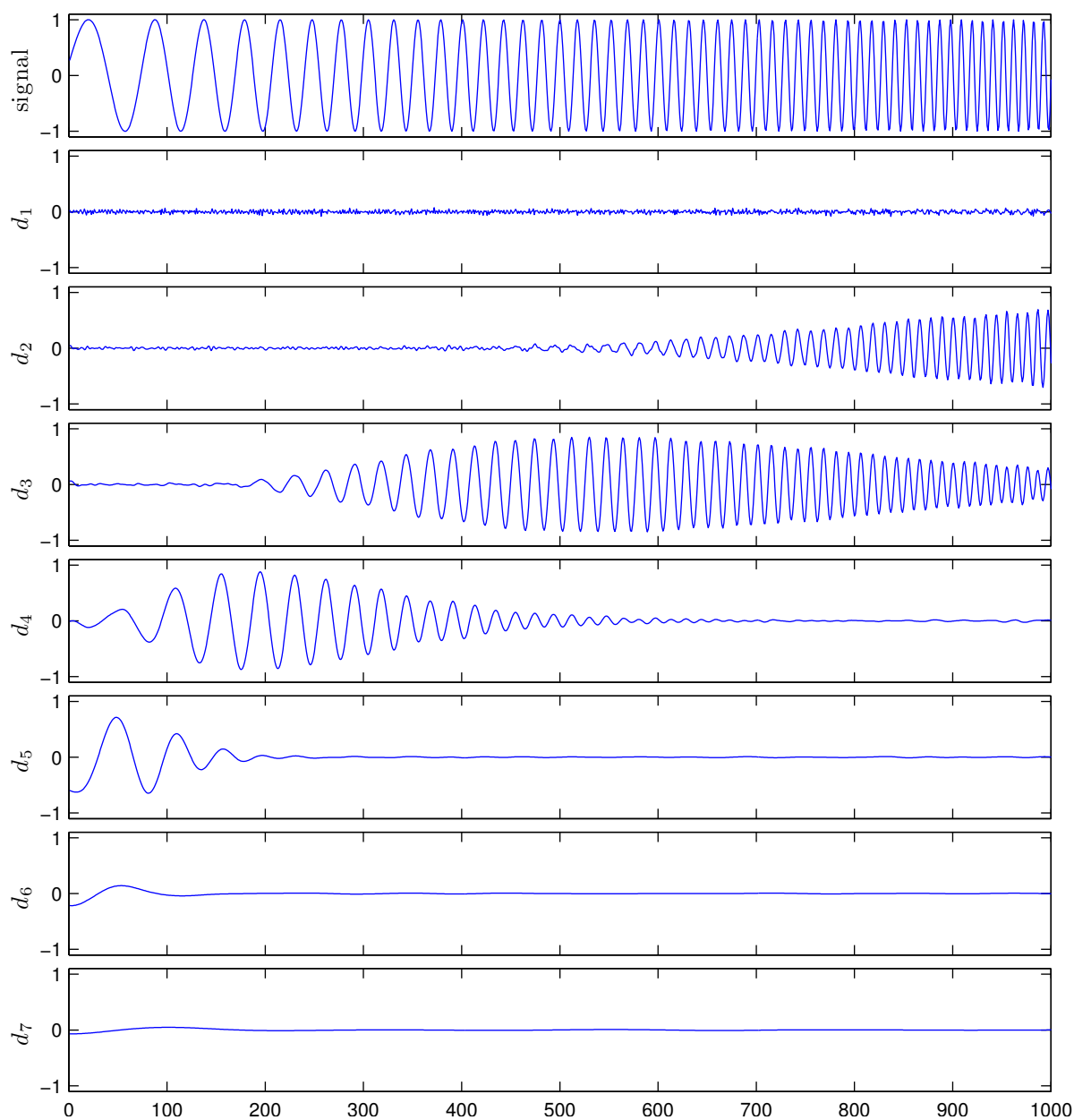


FIGURE 1.25 – Exemple d'application de l'EEMD. Cas d'une composante sinusoïdale modulée en fréquence linéairement. Les IMFs EEMD ont été calculés à l'aide de 10000 réalisations de bruit de variance 1.

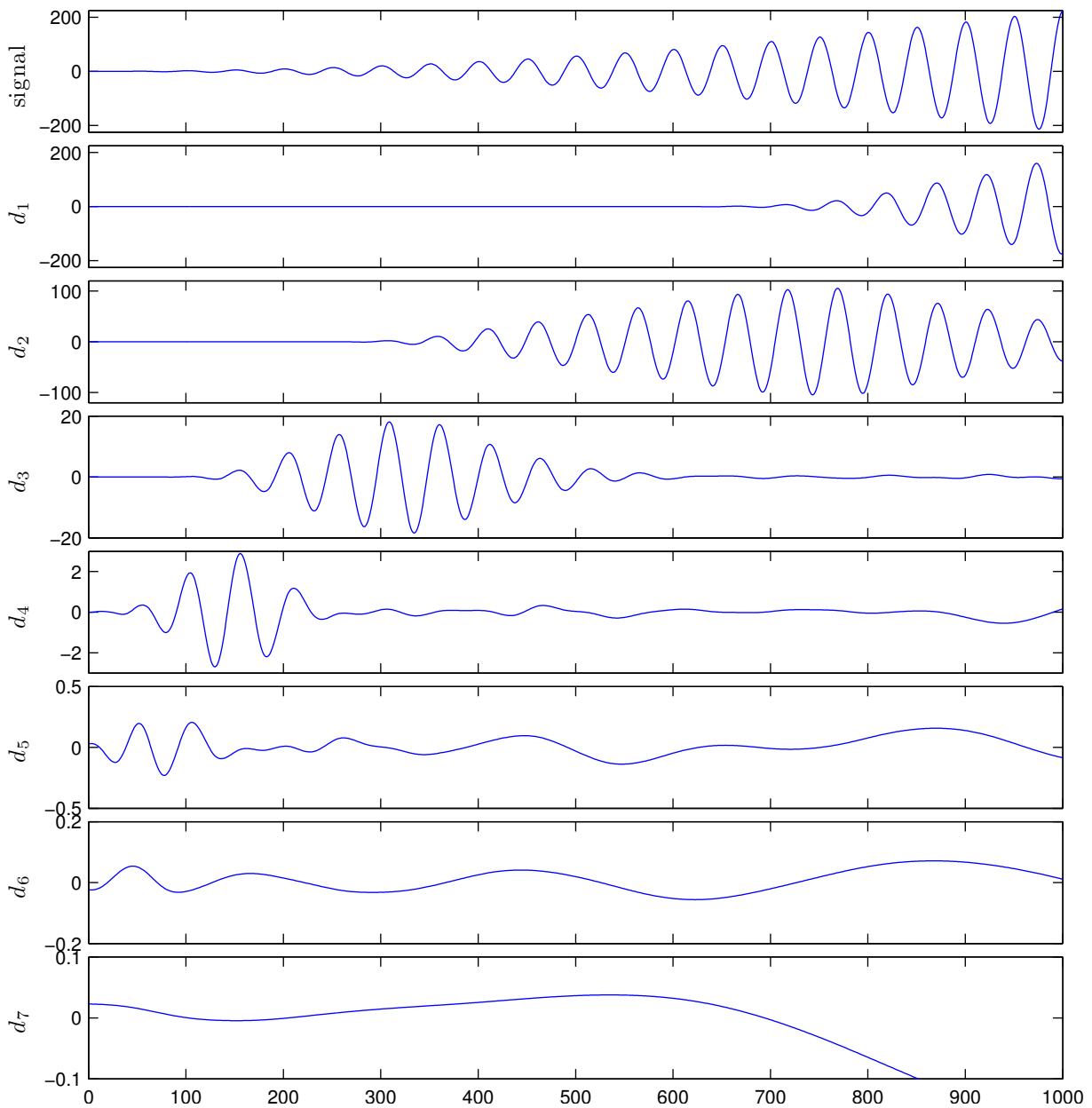


FIGURE 1.26 – Exemple d’application de l’EEMD. Cas d’une composante sinusoïdale modulée en amplitude quadratiquement. Les IMFs EEMD ont été calculés à l’aide de 10000 réalisations de bruit de variance 1. Dans les 4 derniers IMFs, on observe une ondulation à droite due à la grande amplitude du signal dans cette région. Il ne s’agit pas d’un effet de bord : le signal utilisé en entrée de l’EEMD est en fait plus long que ce qui est représenté.

6.5.1 « Complex Empirical Mode Decomposition »

6.5.1.1 Description Cette première approche, introduite sous le nom de « Complex Empirical Mode Decomposition » [60], est basée sur le fait que prendre la partie réelle d'un signal (anti-)analytique est une opération inversible. Ainsi, le signal bivarié à analyser $x(t)$ est tout d'abord séparé en ses 2 composantes analytique et antianalytique :

$$x_+(t) = h_+ * x(t) \quad (1.129)$$

$$x_-(t) = h_- * x(t), \quad (1.130)$$

où h_+ et h_- sont les filtres analytique et antianalytique :

$$H_+[x](\omega) = \mathbf{1}_{\mathbb{R}_+^*}(\omega)X(\omega) + \frac{1}{2}\mathbf{1}_{\{0\}}(\omega)X(0) \quad (1.131)$$

$$H_-[x](\omega) = \mathbf{1}_{\mathbb{R}_-^*}(\omega)X(\omega) + \frac{1}{2}\mathbf{1}_{\{0\}}(\omega)X(0). \quad (1.132)$$

Ensuite, l'EMD est appliquée séparément à la partie réelle de chacune des composantes :

$$\text{Re}(x_+(t)) = \sum_{i=1}^{n_+} d_i^+(t) + a_{n_+}^+(t) \quad (1.133)$$

$$\text{Re}(x_-(t)) = \sum_{i=1}^{n_-} d_i^-(t) + a_{n_-}^-(t). \quad (1.134)$$

Dans ces décompositions, les IMFs obtenus (réels) $d_i^+(t)$ et $d_i^-(t)$ peuvent alors être vus comme les parties réelles d'IMFs bivariés décomposant $x_+(t)$ et $x_-(t)$. Ces deux signaux étant (anti-)analytiques, il est légitime de supposer que les composantes de leurs décompositions respectives partagent la même propriété. On peut ainsi définir les EMD complexes de $x_+(t)$ et $x_-(t)$ en prenant respectivement les parties analytiques et antianalytiques des décompositions de leurs parties réelles :

$$x_+(t) = \sum_{i=1}^{n_+} h_+ * d_i^+(t) + h_+ * a_{n_+}^+(t) \quad (1.135)$$

$$x_-(t) = \sum_{i=1}^{n_-} h_- * d_i^-(t) + h_- * a_{n_-}^-(t). \quad (1.136)$$

On a ainsi obtenu séparément des EMD des parties positives et négatives du spectre, l'EMD globale est définie simplement comme la juxtaposition de ces deux décompositions :

$$x(t) = \sum_{-n_- \leq i \leq n_+, i \neq 0} d_i(t) + a(t), \quad (1.137)$$

avec

$$d_i(t) = \begin{cases} h_+ * d_i^+(t), & 1 \leq i \leq n_+ \\ h_- * d_{-i}^-(t), & -n_- \leq i \leq -1, \end{cases} \quad (1.138)$$

et

$$a(t) = x(t) - \sum_{-n_- \leq i \leq n_+, i \neq 0} d_i(t) \quad (1.139)$$

$$a(t) = h_+ * a_{n_+}^+(t) + h_- * a_{n_-}^-(t). \quad (1.140)$$

Remarque : On aurait également pu garder deux composantes séparées pour le résidu $a(t)$. Le choix des auteurs de les regrouper est probablement dû au fait que le résidu est généralement vu comme une tendance globale du signal et n'est pas interprété en termes d'oscillations — ou de rotations dans le cas présent — comme les IMFs.

6.5.1.2 Commentaires Le principal attrait de cette approche est qu'elle garantit dans une certaine mesure que les IMFs correspondent bien à l'idée qu'on peut se faire d'un signal tournant autour de l'origine du plan complexe. En effet les IMFs étant les parties (anti-)analytiques d'IMFs réels, ils devraient dans une assez large mesure avoir une fréquence instantanée soit positive soit négative (voir [59] pour des exemples d'exceptions) et donc bien correspondre à la notion de signal tournant autour de l'origine du plan complexe. On verra que l'autre EMD bivariée proposée ci-après offre moins de garanties quant au caractère tournant des IMFs.

Le défaut majeur de cette première approche est probablement qu'elle s'appuie fortement sur les filtres h_+ et h_- qui en pratique posent les mêmes problèmes que la transformée de Hilbert, à savoir des effets de bords très importants dus d'une part au caractère singulier de la réponse impulsionnelle et d'autre part à sa décroissance lente en $\mathcal{O}(1/t)$ en $\pm\infty$.

Enfin un deuxième défaut important de cette méthode est que la décomposition n'est pas covariante vis-à-vis des rotations du plan complexe. On peut remarquer en effet que le fait de prendre la partie réelle des composantes (anti-)analytiques privilégie de fait une direction donnée du plan. Malheureusement, s'il est évidemment possible de choisir une autre direction (en prenant par exemple la partie imaginaire), il n'y a aucune raison pour que le résultat final soit identique. Cet aspect peut cependant être amélioré si, comme dans les algorithmes proposés par la suite, au lieu de considérer seulement la partie réelle, on considère une série de projections sur des directions bien réparties dans $[0, 2\pi]$. On se rapprocherait alors très fortement de l'algorithme Algo. 5 présenté en 6.5.3.

6.5.2 Principe de la deuxième approche

Contrairement à l'« EMD complexe » qu'on vient de présenter qui définit une EMD bivariée à l'aide d'une utilisation astucieuse de l'EMD classique, la deuxième approche propose une EMD bivariée à tous niveaux dans la mesure où le principe même du processus de tamisage est construit dans un cadre bivarié. Ainsi, le principe de base sur lequel sont fondés les deux algorithmes proposés par la suite, est de donner corps au principe précédemment énoncé selon lequel un signal bivarié peut être décomposé en la somme d'une composante « tournant rapidement » et d'une composante « tournant lentement ».

6.5.2.1 Une enveloppe en trois dimensions Pour séparer ces deux composantes, l'idée est comme pour l'EMD classique de définir la composante tournant lentement de manière géométrique comme la moyenne d'une « enveloppe du signal ». Il faut pour cela représenter le signal bivarié en 3 dimensions : le temps et ses deux composantes. Dans une telle représentation, il est clair qu'on ne peut se contenter comme dans l'EMD classique de deux enveloppes (dessus / dessous) mais qu'on a besoin au contraire de considérer toutes les directions possibles. De fait, la notion d'enveloppe qui s'impose est une surface tubulaire qui enserre le signal (cf Fig. 1.27). De là, il reste à définir la moyenne d'une telle enveloppe, c'est-à-dire à chaque instant le « centre » de la section (courbe fermée) de l'enveloppe tridimensionnelle. En pratique, la définition adoptée par les algorithmes bivariés n'est pas intuitive et sera présentée ultérieurement. On peut remarquer que la question se pose en fait également dans le cas univarié où on choisit de prendre la demi-somme des deux enveloppes à défaut d'une définition plus performante [28]. On a vu par exemple en 3.3 que cette définition était peu adaptée aux signaux fortement asymétriques. La définition du centre de l'enveloppe utilisée par les algorithmes bivariés est certainement tout aussi discutable.

6.5.2.2 Les « armatures latérales » du tube enveloppe Pour calculer le centre de l'enveloppe, l'idée mise en œuvre par les deux algorithmes proposés est de s'appuyer sur des « armatures latérales » du tube enveloppe (cf Fig. 1.27 (b)). Si l'on suppose que la section du tube entoure un domaine convexe à tout instant, chacune de ces armatures correspond alors à une ligne de crête de

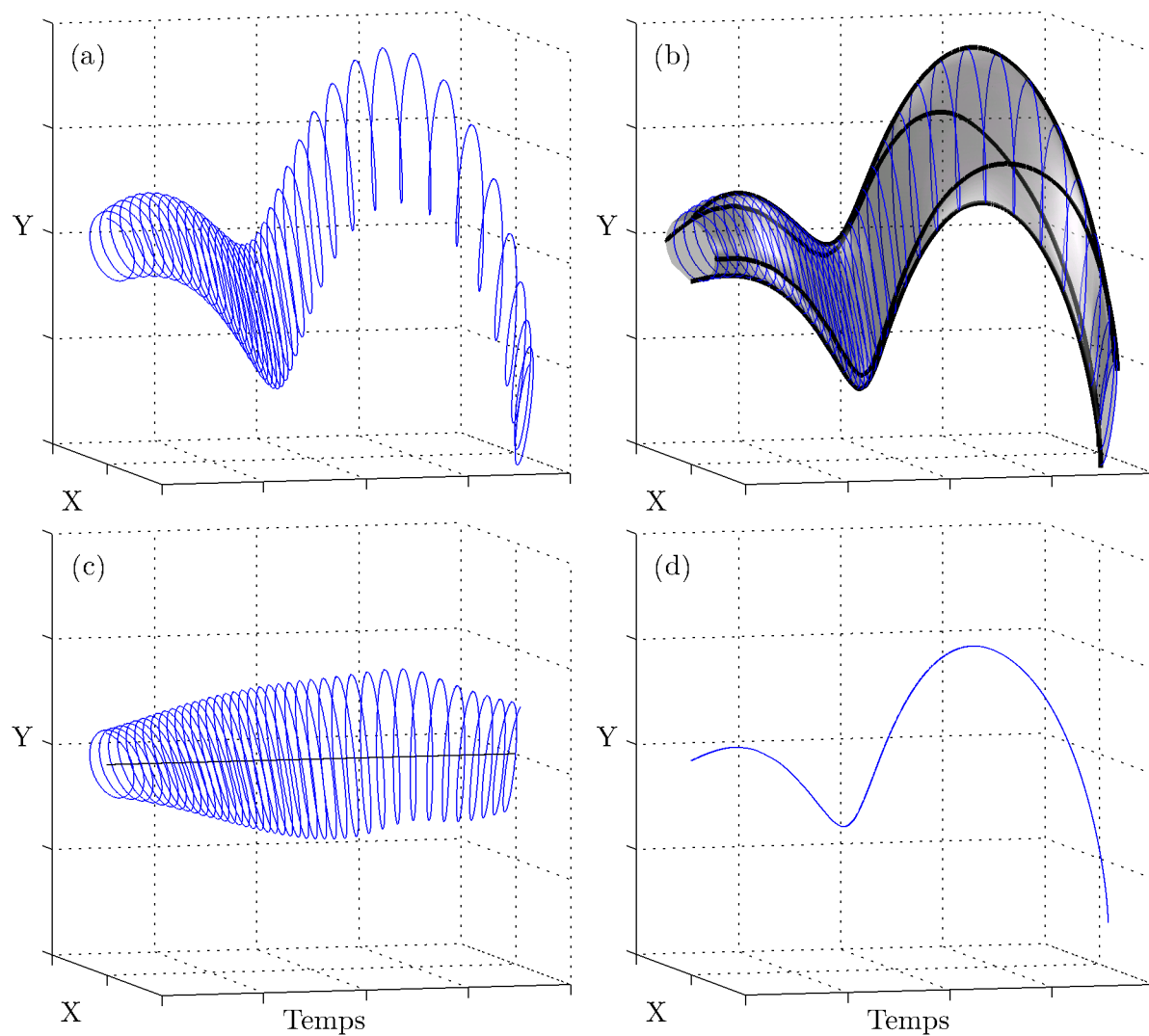


FIGURE 1.27 – Principe des extensions bivariées. (a) Un signal tournant composite. (b) Le signal entouré de son enveloppe-tube 3D. Les lignes noires épaisses sont les « armatures latérales » utilisées pour calculer l'axe central. (c) Composante tournant rapidement. (d) Composante tournant plus lentement définie comme l'axe central du tube en (b).

l'enveloppe dans une direction donnée. Ainsi, l'armature associée à la direction « vers le haut » pourrait être définie comme à tout instant le point de la section du tube dont l'altitude est maximale. Si la section est suffisamment régulière la tangente en ce point à la section est horizontale et la courbure dirigée vers le bas. En pratique cependant il n'y a pas besoin de définir le tube enveloppe pour définir de telles armatures latérales. En effet, ces armatures étant incluses dans le tube enveloppe, elles sont en contact avec la trajectoire (3D) du signal en un certain nombre de points et on peut alors les définir en tout point par interpolation. Par rapport à l'EMD classique, ces points particuliers sont en fait les équivalents des maxima locaux (ou des minima locaux) et peuvent être vus comme des maxima locaux *dans une direction particulière*. Si l'on considère une direction φ , les maxima locaux dans cette direction sont les points t_i^φ où la dérivée du signal est orthogonale à cette direction et la courbure tournée dans la direction opposée :

$$\left\langle \frac{dx}{dt}(t_i^\varphi), u_\varphi \right\rangle = 0 \quad \text{et} \quad \left\langle \frac{d^2x}{dt^2}(t_i^\varphi), u_\varphi \right\rangle < 0, \quad (1.141)$$

où u_φ est un vecteur unitaire dirigé selon la direction φ . On constate par ailleurs que dans la mesure où les dérivées peuvent être sorties des produits scalaires ces maxima locaux dans la direction φ coïncident avec les maxima locaux de la projection du signal sur cette même direction.

À partir de ces maxima dans la direction φ , on peut alors définir l'armature latérale $e_\varphi(t)$ comme la spline cubique interpolante. Ce choix d'interpolation provient simplement du fait que la spline cubique est à l'heure actuelle encore considérée comme le meilleur schéma d'interpolation pour l'EMD classique malgré un certain nombre de défauts connus [28, 26, 47] et des travaux prometteurs sur des interpolations potentiellement plus performantes (cf 5.3).

6.5.2.3 Le « centre » des armatures latérales À partir d'un jeu de N armatures latérales, il y a plusieurs moyens de définir leur centre. Le centre étant défini à chaque instant par rapport à la section du tube à cet instant, la question est en fait de définir le centre de N points dans le plan. Si l'on prend l'exemple de 4 armatures latérales — et donc 4 points sur la section du tube — correspondant aux 4 directions haut, bas, gauche et droite, on peut définir leur centre d'au moins deux manières :

1. l'isobarycentre des 4 points (cf Fig. 1.28 (a)).
2. l'intersection de deux droites, l'une étant à mi-chemin entre les deux tangentes horizontales à la section du tube passant par les points « haut » et « bas », l'autre étant à mi-chemin entre les deux tangentes verticales passant par les points « gauche » et « droite » (cf Fig. 1.28 (b)).

Plus généralement, pour N directions $\varphi_k, 1 \leq k \leq N$, on a N armatures latérales $e_{\varphi_k}(t)$ et leur « centre » $m(t)$ peut être défini comme l'isobarycentre :

$$m(t) = \frac{1}{N} \sum_{k=1}^N e_{\varphi_k}(t), \quad (1.142)$$

ou comme

$$m(t) = \frac{2}{N} \sum_{k=1}^N \langle e_{\varphi_k}(t), u_{\varphi_k} \rangle u_{\varphi_k}, \quad (1.143)$$

qui généralise la méthode utilisant les tangentes.

Remarque : La différence de normalisation entre les deux versions provient simplement du fait qu'une tangente $\langle e_{\varphi_k}(t), u_{\varphi_k} \rangle u_{\varphi_k}$ contient seulement une des deux composantes de l'armature $e_{\varphi_k}(t)$ et qu'il faut donc par exemple deux tangentes associées à des directions orthogonales pour obtenir l'équivalent d'une armature.

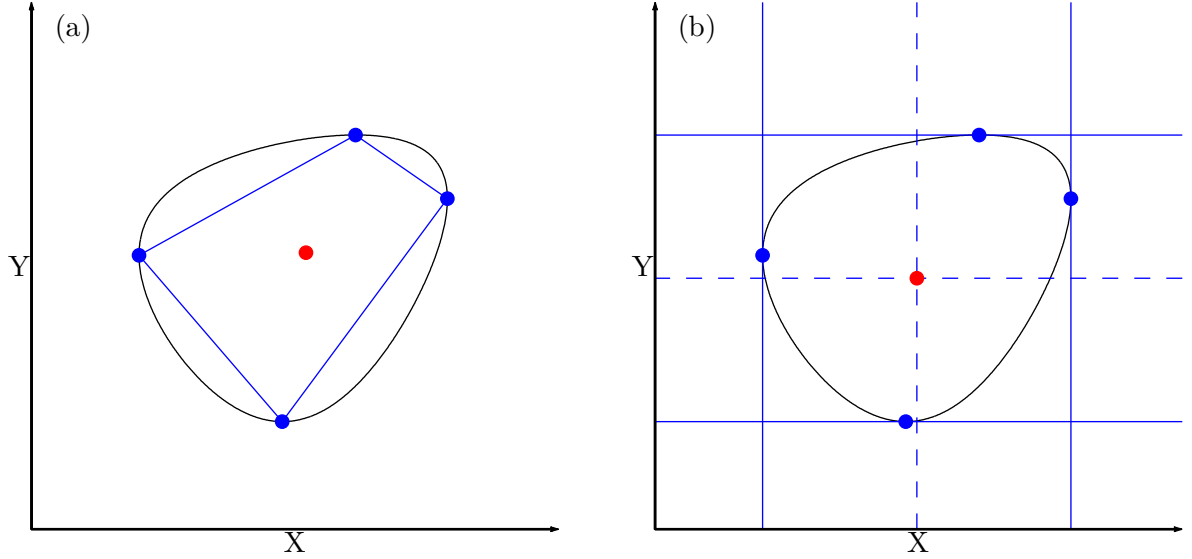


FIGURE 1.28 – Deux possibilités pour définir l'axe central de l'enveloppe.

6.5.2.4 Des méthodes inégales vis-à-vis de l'échantillonnage A priori plus compliquée, l'intérêt de la deuxième méthode est essentiellement qu'elle s'avère plus robuste par rapport aux erreurs dues à l'échantillonnage. En effet, tout comme pour l'EMD classique, la détection des extrema comporte une certaine incertitude du fait de l'échantillonnage. En revanche, une différence majeure entre les deux est que dans le cas univarié, la dérivée du signal est nulle aux extrema alors que dans le cas bivarié, seule sa composante dans la direction considérée est nulle. Pour estimer cette incertitude dans le cas bivarié, on peut s'intéresser au développement de Taylor du signal autour d'un de ses extrema associé à une direction φ :

$$x(t) - x(t_i) = (t - t_i^\varphi) \frac{dx}{dt}(t_i^\varphi) + \frac{(t - t_i^\varphi)^2}{2} \frac{d^2x}{dt^2}(t_i^\varphi) + \mathcal{O}((t - t_i^\varphi)^3). \quad (1.144)$$

Si l'on projette maintenant cette expression sur u_φ et la direction orthogonale, on obtient :

$$\langle x(t) - x(t_i), u_\varphi \rangle = \frac{(t - t_i^\varphi)^2}{2} \left\langle \frac{d^2x}{dt^2}(t_i^\varphi), u_\varphi \right\rangle + \mathcal{O}((t - t_i^\varphi)^3) \quad (1.145)$$

$$x(t) - x(t_i) - \langle x(t) - x(t_i), u_\varphi \rangle u_\varphi = (t - t_i^\varphi) \frac{dx}{dt}(t_i^\varphi) + \mathcal{O}((t - t_i^\varphi)^2), \quad (1.146)$$

où l'on voit que l'incertitude est beaucoup plus importante dans la direction orthogonale à u_φ que dans la direction colinéaire, où elle est comparable à celle observée pour l'EMD classique. En pratique, cette incertitude dans la direction orthogonale fait que la définition du centre (1.142) génère des erreurs relativement importantes même pour des taux de suréchantillonnage habituellement considérés comme grands. Par contre, l'erreur générée par la définition (1.143) n'utilisant que la projection sur la direction u_φ est beaucoup plus faible, de l'ordre de celle observée dans l'EMD classique. De ce fait, cette dernière définition est a priori préférable si le signal analysé est assez peu suréchantillonné ou si la précision sur la décomposition est particulièrement importante. En revanche la définition (1.142) a sans doute son intérêt dans les cas très bien échantillonnés.

6.5.3 Algorithmes

À partir des deux définitions proposées pour le « centre » des armatures latérales, on peut écrire deux opérateurs de tamisage différents définis par les deux algorithmes Algo. 4 et Algo. 5. Pour

obtenir les algorithmes d'EMD bivariées associés à ces opérateurs il suffit de remplacer l'opérateur de tamisage dans l'algorithme de l'EMD classique par l'un de ces nouveaux opérateurs.

Algorithme 4 : EMD bivariée : opérateur de tamisage associé à la définition (1.142)

- 1 **pour** $1 \leq k \leq N$ **faire**
- 2 Projeter le signal $x(t)$ sur la direction $\varphi_k : p_{\varphi_k}(t) = \langle x(t), u_{\varphi_k} \rangle$
- 3 Extraire les positions $\{t_j^k\}$ des maxima de $p_{\varphi_k}(t)$.
- 4 Interpoler les points $\{(t_j^k, x(t_j^k))\}$ pour obtenir l'armature latérale dans la direction φ_k : $e_{\varphi_k}(t)$.

Calculer la moyenne des armatures latérales : $m(t) = \frac{1}{N} \sum_{k=1}^N e_{\varphi_k}(t)$

- 5
 - 6 Soustraire la moyenne au signal : $\mathcal{S}^{B1}[x](t) = x(t) - m(t)$
-

Algorithme 5 : EMD bivariée : opérateur de tamisage associé à la définition (1.143)

- 1 **pour** $1 \leq k \leq N$ **faire**
- 2 Projeter le signal $x(t)$ sur la direction $\varphi_k : p_{\varphi_k}(t) = \langle x(t), u_{\varphi_k} \rangle$
- 3 Extraire les maxima de $p_{\varphi_k}(t) : \{t_j^k, p_j^k\}$.
- 4 Interpoler les points $\{(t_j^k, p_j^k)\}$ pour obtenir la composante selon u_{φ_k} de l'armature latérale dans la direction $\varphi_k : e'_{\varphi_k}(t)$.
- 5 Calculer la moyenne de ces composantes :

$$m(t) = \frac{2}{N} \sum_{k=1}^N e'_{\varphi_k}(t) u_{\varphi_k}$$

- 6 Soustraire la moyenne au signal : $\mathcal{S}^{B2}[x](t) = x(t) - m(t)$
-

Concernant l'algorithme Algo. 5, on remarque que l'interpolation est identique à celle réalisée pour obtenir l'enveloppe supérieure dans l'opérateur de tamisage original. Il en découle que si le nombre N de directions considéré est pair (et que chaque direction a son opposé), on peut exprimer le deuxième algorithme à partir de l'opérateur de tamisage de l'EMD classique (Algo. 6). Cette dernière formulation est essentiellement intéressante dans la mesure où elle permet d'étudier le comportement de cette EMD bivariée aux lumières des connaissances acquises — ou à venir — sur l'EMD classique.

Algorithme 6 : Reformulation de l'algorithme Algo. 5

- 1 **pour** $1 \leq k \leq N/2$ **faire**
- 2 Projeter le signal $x(t)$ sur la direction $\varphi_k : p_{\varphi_k}(t) = \langle x(t), u_{\varphi_k} \rangle$
- 3 Calculer l'estimation partielle de $\mathcal{S}^{B2}[x](t)$ selon la direction $\varphi_k : s_{\varphi_k}(t) = \mathcal{S}[p_{\varphi_k}](t)$

4 Finalement : $\mathcal{S}^{B2}[x](t) = \frac{4}{N} \sum_{k=1}^{N/2} s_{\varphi_k}(t) u_{\varphi_k}$

6.5.4 Symétries

Les algorithmes bivariés préservent les propriétés de symétrie de l'EMD classique (cf 3.2) et en ajoutent deux sous certaines hypothèses concernant le jeu de directions utilisé : la covariance par

rapport à certaines réflexions dont l'axe passe par l'origine et la covariance par rapport à certaines rotations centrées à l'origine. Si on assimile le plan $\mathbb{R}^2$ au plan complexe, la première de ces deux symétries est liée à la notion de covariance par rapport à la conjugaison qui est la forme sous laquelle elle est présentée ici.

conjugaison : Si le jeu de directions utilisé est composé de paires de directions conjuguées, avec éventuellement en plus les directions correspondant aux arguments 0 et π , les algorithmes bivariés commutent avec la conjugaison. En effet, si $\{\varphi_k, 1 \leq k \leq N\} = \{-\varphi_k, 1 \leq k \leq N\}$ (en supposant les φ_k différents les uns des autres), alors on peut trouver une permutation p telle que $\forall k, \varphi_{p(k)} = -\varphi_k$ et il en découle que la projection de $x(t)$ sur u_{φ_k} est identique à la projection de x^* sur $u_{\varphi_{p(k)}}$:

$$\langle x(t), u_{\varphi_k} \rangle = \operatorname{Re} (x(t)e^{-i\varphi_k}) \quad (1.147)$$

$$= \frac{x(t)e^{-i\varphi_k} + x^*(t)e^{i\varphi_k}}{2} \quad (1.148)$$

$$= \frac{x(t)e^{i\varphi_{p(k)}} + x^*(t)e^{-i\varphi_{p(k)}}}{2} \quad (1.149)$$

$$= \langle x^*(t), u_{\varphi_{p(k)}} \rangle \quad (1.150)$$

Par conséquent, les positions des extrema $t_j^k[x]$ de la projection $p_{\varphi_k}[x](t)$ de $x(t)$ sur u_{φ_k} sont les mêmes que celles des extrema $t_j^{p(k)}[x^*]$ de la projection de $x^*(t)$ sur $u_{\varphi_{p(k)}}$. Sachant que l'opération d'interpolation commute avec la conjugaison, on obtient la relation entre les armatures latérales pour l'algorithme Algo. 4 : $e_{\varphi_k}[x](t) = (e_{\varphi_{p(k)}}[x^*](t))^*$. Pour l'algorithme Algo. 5, la relation est encore plus simple : $e'_{\varphi_k}[x](t) = e'_{\varphi_{p(k)}}[x](t)$. Au final, on aboutit dans les deux cas au fait que la moyenne des armatures latérales pour $x^*(t)$ est conjuguée de celle pour $x(t)$ et donc que

$$\mathcal{S}^{B1}[x^*](t) = (\mathcal{S}^{B1}[x](t))^* \quad \text{et} \quad \mathcal{S}^{B2}[x^*](t) = (\mathcal{S}^{B2}[x](t))^*. \quad (1.151)$$

Dans la limite où le nombre de directions tend vers l'infini, le jeu de directions peut être remplacé par une distribution continue de directions. Si celle-ci est symétrique par rapport à 0 (en supposant les directions dans $[-\pi, \pi]$), alors les algorithmes bivariés commutent asymptotiquement avec la conjugaison.

rotations : Si le jeu de directions possède des symétries de rotation, alors les algorithmes d'EMD bivariés les conservent. Soit $\psi \in [0, 2\pi]$ tel qu'il existe une permutation p telle que $\forall k, \varphi_k = \varphi_{p(k)} + \psi$. La projection de $x(t)e^{i\psi}$ sur la direction φ_k vérifie alors

$$\langle x(t)e^{i\psi}, u_{\varphi_k} \rangle = \operatorname{Re} (x(t)e^{i\psi}e^{-i\varphi_k}) \quad (1.152)$$

$$= \langle x(t), u_{\varphi_{p(k)}} \rangle. \quad (1.153)$$

Sachant que l'interpolation commute avec les rotations du plan complexe, il en résulte que les armatures latérales pour l'algorithme Algo. 4 vérifient $e_{\varphi_k}[xe^{i\psi}](t) = e^{i\psi}e_{\varphi_{p(k)}}[x](t)$. Pour l'algorithme Algo. 5, la relation est $e'_{\varphi_k}[x](t) = e'_{\varphi_{p(k)}}[x](t)$. Au final, on obtient dans les deux cas que la moyenne des armatures $m(t)$ vérifie $m[xe^{i\psi}](t) = e^{i\psi}m[x](t)$ et donc que

$$\mathcal{S}^{B1}[xe^{i\psi}](t) = e^{i\psi}\mathcal{S}^{B1}[x](t) \quad \text{et} \quad \mathcal{S}^{B2}[xe^{i\psi}](t) = e^{i\psi}\mathcal{S}^{B2}[x](t). \quad (1.154)$$

Dans la limite où le nombre de directions tend vers l'infini, les algorithmes bivariés commutent avec les rotations qui laissent invariante la distribution de directions. En particulier, si cette distribution est uniforme, alors les algorithmes bivariés commutent asymptotiquement avec toute rotation du plan.

6.5.5 IMFs bivariés

Les algorithmes bivariés Algo. 4 et Algo. 5 sont conçus pour que des signaux tournant autour de l'origine du plan soient a priori des IMFs admissibles. Sachant que la notion de signal tournant autour d'un point est plutôt vague et que de plus il n'est a priori pas impossible d'obtenir des IMFs qui ne tournent pas autour de l'origine, l'objectif de cette section est de préciser dans une certaine mesure les caractéristiques des IMFs produits par ces algorithmes, c'est-à-dire des signaux qui sont des points fixes — éventuellement approximatifs — des opérateurs de tamisage :

$$\mathcal{S}^{B1}[x](t) \approx x(t) \quad \text{ou} \quad \mathcal{S}^{B2}[x](t) \approx x(t). \quad (1.155)$$

Étant donnée la complexité des opérateurs de tamisage bivariés, il n'existe pas actuellement de caractérisation théorique de leurs points fixes. Par la suite, nous tenterons de cerner quelques propriétés de ces points fixes à l'aide d'exemples simples. Dans la mesure où, à l'instar de l'opérateur de tamisage classique, les opérateurs de tamisage bivariés sont construits pour agir de manière locale on peut se retreindre dans un premier temps à l'étude de leur comportement sur des signaux périodiques. On observe en effet que le comportement de ces opérateurs sur des signaux qui sont localement proches de périodiques est très semblable au cas de signaux rigoureusement périodiques (cf Fig. 1.29).

6.5.5.1 Cas de signaux dont le tube enveloppe est constant On s'intéresse ici au cas de signaux périodiques tournant autour de l'origine du plan. Pour simplifier l'étude, on peut s'intéresser tout d'abord au cas simple où le signal n'effectue qu'un tour par période. Parmi ceux-ci, il existe un cas très simple dans lequel on peut facilement calculer la position du centre de l'enveloppe : le cas où les armatures sont constantes. C'est en particulier le cas si chacune de ces armatures n'interpole qu'un point par période, c'est à dire que quelle que soit la direction sur laquelle on projette le signal, celui-ci n'a que deux extrema par période (un maximum et un minimum). Géométriquement, ceci est encore équivalent à dire que le *sens de rotation instantané* du signal est constant, ce dernier étant décrit par exemple par le sens du produit vectoriel des vecteurs vitesse et accélération

$$\left\langle \frac{dx}{dt} \wedge \frac{d^2x}{dt^2}, e_z \right\rangle, \quad (1.156)$$

où e_z est un vecteur orthogonal au plan dans lequel vit le signal $x(t)$. Ceci est aussi équivalent à dire que la trajectoire du signal dans le plan entoure un domaine convexe. Enfin, si le signal tourne par exemple dans le sens trigonométrique, cette propriété est encore équivalente à dire que la phase de la dérivée de $x(t)$, pris comme un signal à valeur complexe, est strictement croissante

$$\frac{dx}{dt} = r(t)e^{i\psi(t)}, \quad r(t) > 0, \quad \frac{d\psi}{dt} > 0. \quad (1.157)$$

Pour un tel signal, on peut alors expliciter les armatures latérales (pour les deux méthodes) et donc leur centre. En identifiant le vecteur unitaire u_φ avec le nombre complexe $e^{i\varphi}$, on peut alors écrire que, pour chaque période, le maximum dans la direction u_φ correspond au point où

$$\psi(t) = \varphi \pm \frac{\pi}{2}, \quad (1.158)$$

avec un signe $+$ si le signal tourne dans le sens trigonométrique et $-$ sinon. Les armatures latérales étant des constantes égales à la valeur du maximum dans la direction considérée, on en déduit les armatures latérales pour les deux algorithmes

$$e_\varphi(t) = x\left(\psi^{-1}\left(\varphi \pm \frac{\pi}{2}\right)\right), \quad (1.159)$$

$$e'_\varphi(t) = \left\langle x\left(\psi^{-1}\left(\varphi \pm \frac{\pi}{2}\right)\right), u_\varphi \right\rangle, \quad (1.160)$$

et les centres des armatures correspondants

$$m_N^{B1}(t) = \frac{1}{N} \sum_{k=1}^N x \left(\psi^{-1} \left(\varphi_k \pm \frac{\pi}{2} \right) \right), \quad (1.161)$$

$$m_N^{B2}(t) = \frac{2}{N} \sum_{k=1}^N \left\langle x \left(\psi^{-1} \left(\varphi_k \pm \frac{\pi}{2} \right) \right), u_{\varphi_k} \right\rangle u_{\varphi_k}. \quad (1.162)$$

Si on prend maintenant la limite de ces expressions quand N tend vers l'infini en supposant que la distribution des φ_k est uniforme dans $[0, 2\pi]$, on obtient

$$m_\infty^{B1}(t) = \frac{1}{2\pi} \int_0^{2\pi} x \left(\psi^{-1}(\varphi) \right) d\varphi, \quad (1.163)$$

$$= \frac{1}{2\pi} \int_0^T x(t) \frac{dx}{dt} dt. \quad (1.164)$$

et de même

$$m_\infty^{B2}(t) = \frac{1}{\pi} \int_0^T \left\langle x(t), u_{\psi(t) \mp \frac{\pi}{2}} \right\rangle u_{\psi(t) \mp \frac{\pi}{2}} dt, \quad (1.165)$$

qui donne en notation complexe

$$m_\infty^{B2}(t) = \frac{1}{\pi} \int_0^T e^{i\psi(t) \mp \frac{i\pi}{2}} \operatorname{Re} \left(x(t) e^{-i\psi(t) \pm \frac{i\pi}{2}} \right) dt, \quad (1.166)$$

$$= m_\infty^{B1}(t) - \frac{1}{2\pi} \int_0^T e^{2i\psi(t)} x^*(t) dt \quad (1.167)$$

$$= m_\infty^{B1}(t), \quad (1.168)$$

puisque

$$\int_0^T e^{2i\psi(t)} x^*(t) \frac{d\psi}{dt} dt = \frac{i}{2} \left(\left[-e^{2i\psi(t)} x^*(t) \right]_0^T + \int_0^T \frac{dx^*}{dt} e^{2i\psi(t)} dt \right) \quad (1.169)$$

$$= \frac{i}{2} \int_0^T r(t) e^{i\psi(t)} dt = \frac{i}{2} \int_0^T \frac{dx}{dt} dt = 0. \quad (1.170)$$

Ainsi, dans ce cas très simple, on voit que le centre des armatures est en fait le même pour les deux algorithmes lorsque $N \rightarrow \infty$ et correspond à la formule

$$m = m_\infty^{B1}(t) = m_\infty^{B2}(t) = \frac{1}{2\pi} \int_0^T x(t) \frac{d\psi}{dt} dt, \quad (1.171)$$

moyenne du signal pondérée par la dérivée de la phase de sa dérivée. Si on se place du point de vue de la trajectoire, il s'agit en fait de la *moyenne de la trajectoire pondérée par sa courbure*. En effet, la courbure de la trajectoire vaut

$$\gamma(t) = \frac{\left| \frac{dx}{dt} \frac{d^2x}{dt^2} \right|}{\left\| \frac{dx}{dt} \right\|^3} \quad (1.172)$$

qui donne en notation complexe

$$\gamma(t) = \frac{\operatorname{Im} \left(\frac{dx^*}{dt} \frac{d^2x}{dt^2} \right)}{\left| \frac{dx}{dt} \right|^3} \quad (1.173)$$

$$= \frac{1}{r(t)} \frac{d\psi}{dt}, \quad (1.174)$$

et l'abscisse curviligne s vérifie

$$ds = \left| \frac{dx}{dt} \right| dt = r(t) dt \quad (1.175)$$

d'où le résultat

$$m = \frac{1}{2\pi} \int_0^T x(t) \gamma(t) \underbrace{r(t) dt}_{ds}. \quad (1.176)$$

Remarque : La formule donnant le centre de l'enveloppe (1.176) est valable dès lors que les armatures sont des constantes. Elle reste donc valable pour de tels signaux périodiques modulés en fréquence, c'est-à-dire qu'elle est valable de manière générale dès lors que le signal suit (sans changer de sens) une trajectoire qui est une courbe fermée à un seul tour englobant une surface convexe.

6.5.6 Exemples d'application

6.5.6.1 Filtrage adaptatif de composantes quasi-monochromatiques modulées AM-FM

Tout comme avec l'EMD classique, l'EMD bivariée permet de décomposer efficacement des sommes de composantes monochromatiques modulées en amplitude et en fréquence dès lors que les modulations sont suffisamment lentes par rapport à la période locale et que les composantes sont suffisamment éloignées en fréquence à tout instant. Un exemple d'une telle décomposition est présenté Fig. 1.29.

6.5.6.2 Exemple d'un signal expérimental

Les algorithmes bivariés sont particulièrement adaptés pour traiter des signaux de la forme de celui proposé Fig. 1.30 qui contiennent très clairement au moins une composante tournante non stationnaire. Le signal correspond à la position au cours du temps d'un flotteur océanographique sous-surfacique suivi acoustiquement. Ce flotteur a été déployé dans le nord de l'océan atlantique pour suivre le mouvement d'eaux salées provenant de la méditerranée au cours de l'expérience « Eastern basin » [49]. Le signal est disponible en ligne à <http://wfdac.who.edu>. Les boucles de la trajectoire indiquent que le flotteur est piégé dans un « vortex cohérent », c'est-à-dire un tourbillon en langage courant.

Appliquées à de tels signaux où des composantes tournantes sont clairement présentes, les algorithmes bivariés donnent typiquement la décomposition proposée Fig. 1.31 où les composantes tournantes ont été isolées dans des IMFs séparés. La décomposition est constituée de composantes tournantes et de composantes non tournantes. Les composantes tournantes sont principalement les deux premiers IMFs qui correspondent au vortex cohérent présumé. Ces deux composantes peuvent a priori être analysées plus avant pour extraire des informations concernant le vortex cohérent, comme sa vitesse de rotation, son ellipticité, etc. Une telle étude des composantes tournantes a déjà été réalisée à l'aide de méthodes de crêtes d'ondelettes pour extraire la partie tournante du signal [37]. Les autres composantes ne semblent pas vraiment tournantes et sont supposées décrire les courants à grande échelle dans lesquels se déplace le vortex cohérent. Il est encore impossible de dire quelles informations peuvent être extraites de ces composantes à grande échelle, mais notre collaborateur océanographe J. M. Lilly est particulièrement intéressé par ces composantes dans la mesure où peu de méthodes permettent d'extraire efficacement les informations à grande échelle de signaux comme celui étudié ici.

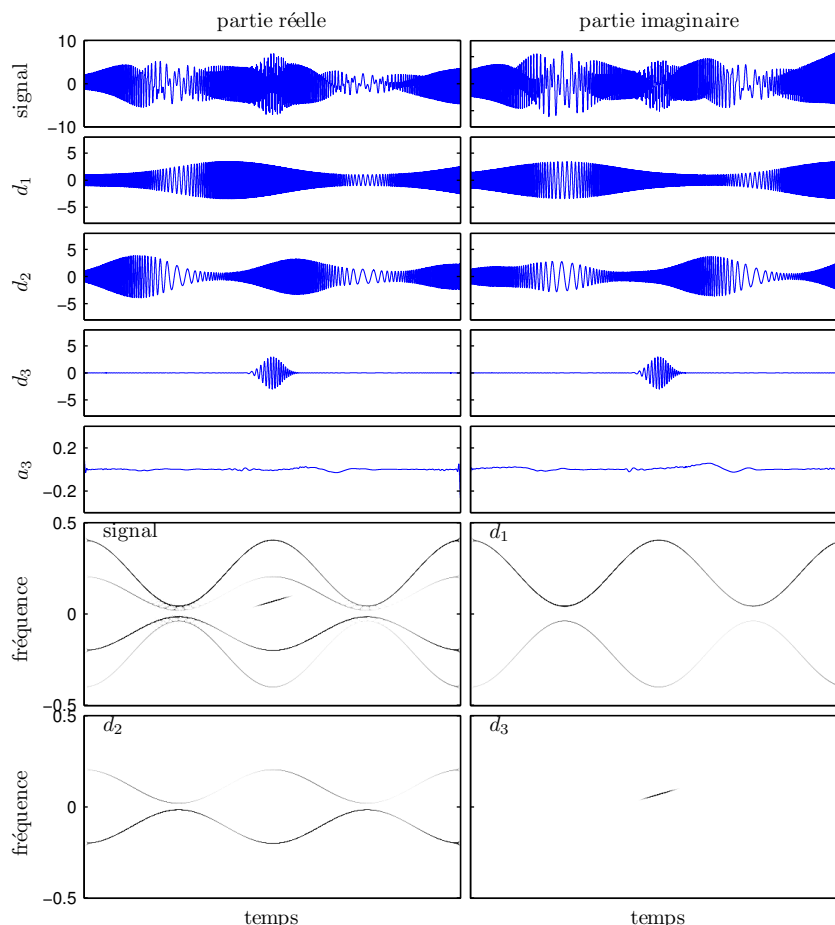


FIGURE 1.29 – Filtrage adaptatif de composantes quasi-monochromatiques modulées AM-FM. Le signal (rangée du haut) est décomposé par l’algorithme Algo. 5 avec 32 directions itéré 10 fois pour extraire chaque IMF. La décomposition est stoppée après avoir extrait 3 IMFs représentés en dessous du signal accompagnés du résidu a_3 où on voit qu’il ne reste pratiquement rien des 3 composantes constituant le signal. Les spectrogrammes réalloués du signal et des trois IMFs occupent les 4 quadrants du bas de la figure. Le signal initial contient 3 composantes modulées en amplitude et en fréquence et étirées (avec une amplitude variable au cours du temps) dans une direction (également variable). Le fait que les deux premières composantes soient étirées dans une direction fait que celles-ci tournent localement de manière elliptique. Sachant qu’une rotation elliptique se décompose en Fourier en deux rotations circulaires de fréquences opposées, ceci explique que leurs représentations temps-fréquence contiennent à la fois des fréquences positives et négatives.

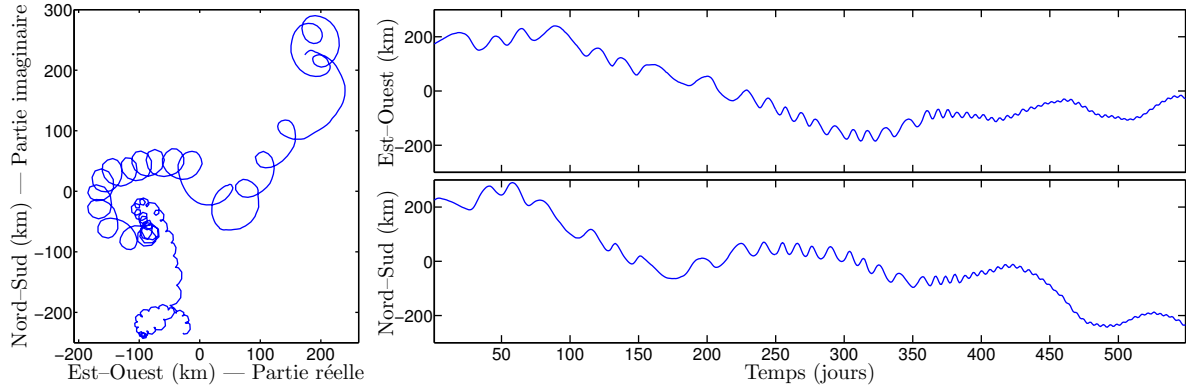


FIGURE 1.30 – Signal correspondant à la position au cours du temps d’un flotteur océanographique sous-surface suivi acoustiquement. Ce flotteur a été déployé dans le nord de l’océan atlantique pour suivre le mouvement d’eaux salées provenant de la méditerranée au cours de l’expérience « Eastern basin » [49]. Le signal est disponible en ligne à <http://wfdac.whoi.edu>.

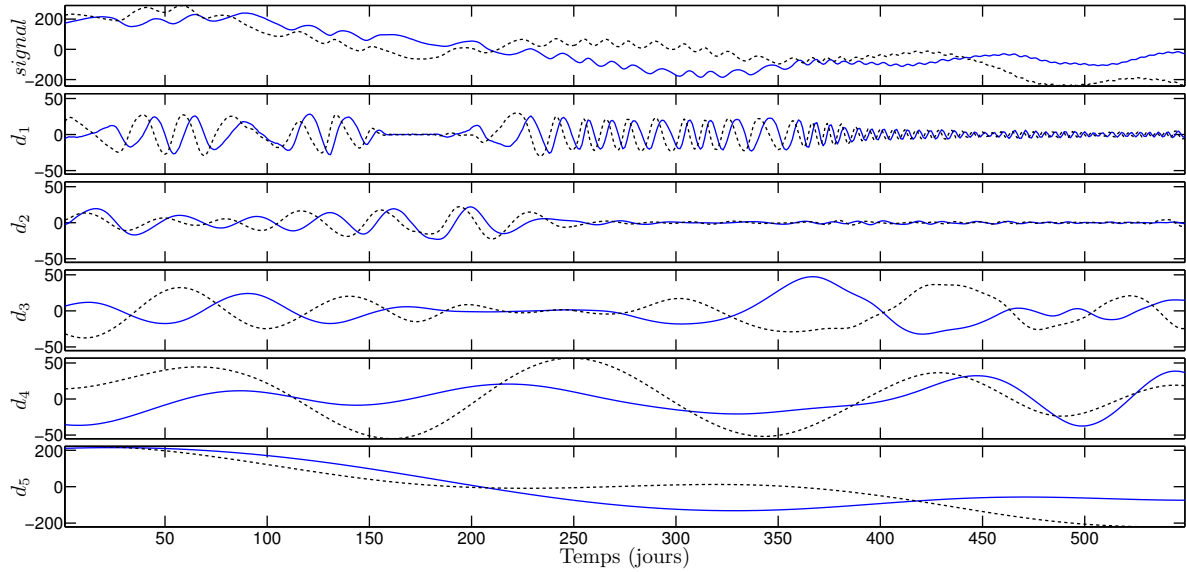


FIGURE 1.31 – EMD bivariée du signal Fig. 1.30. Les parties réelles sont représentées par des traits pleins bleus et les parties imaginaires par des tirets noirs. Cette décomposition a été obtenue avec l’algorithme Algo. 5 avec 64 directions et 10 itérations pour chaque IMF. Les endroits où les composantes sont tournantes peuvent être identifiées par un décalage de phase à peu près constant entre les deux composantes. Si le décalage correspond à un quart de période (composantes en quadrature), la rotation est circulaire. D’autres décalages correspondent à des rotations elliptiques.

6.5.7 Limites de l'approche proposée

L'EMD bivariée étant encore très récente, ses performances et ses limites restent à évaluer de manière plus exhaustive. On propose cependant quelques observations concernant des limites éventuelles de la méthode.

L'une de ces limites est sans doute que dans sa formulation l'EMD bivariée est conçue pour extraire des composantes tournantes alors qu'en pratique, il n'est pas rare d'obtenir des IMFs qui ne sont pas vraiment tournants. Si la décomposition n'est pas pour autant dénuée d'intérêt, elle est néanmoins discutable dans la mesure où l'« enveloppe » de tels signaux ne correspond plus au cadre dans lequel les algorithmes ont été définis. Plus précisément, les IMFs non tournants peuvent être vus comme localement tournants mais leur sens de rotation change fréquemment, parfois sans que le signal puisse faire un tour complet entre deux changements. Ces changements de sens semblent être de deux sortes :

- des changements progressifs au cours desquels le cycle de rotation du signal s'aplatit progressivement, passe par un stade totalement aplati, où le signal oscille plus qu'il ne tourne, puis se remet à tourner mais dans l'autre sens. Par exemple

$$x(t) = t \cos t + i \sin t \quad (1.177)$$

est un IMF bivarié tout à fait acceptable pour les deux algorithmes et qui change de sens de rotation en $t = 0$.

- des changements où le signal tournant dans un sens donné rebrousse chemin pour tourner dans l'autre sens. Un exemple est proposé Fig. 1.32.

En pratique, le premier type de changement de sens n'est pas problématique dans la mesure où le tube enveloppe tel qu'il est défini dans les algorithmes correspond bien à l'idée qu'on peut se faire de l'enveloppe d'un tel signal.

En revanche, le changement de sens par rebroussement pose plus de problèmes dans la mesure où la forme de l'enveloppe d'un signal présentant de tels changements de sens est nettement moins évidente à définir. Si on considère la définition utilisée dans les algorithmes sur l'exemple Fig. 1.32, on peut voir que toutes les armatures latérales « oscillent » à des fréquences voisines de celle du signal et que donc l'enveloppe du signal est ondulée à l'échelle de la période du signal ce qui contredit l'intuition de la notion d'enveloppe qui implicitement est censée être plutôt lisse à l'échelle de variation du signal. Malgré cela, il n'est pas a priori impossible que le centre des armatures latérales soit nul. En pratique cependant, le signal proposé en exemple Fig. 1.32 n'est pas un IMF parfait dans la mesure où le centre de ses armatures latérales n'est pas vraiment nul mais juste de faible amplitude. En itérant davantage le tamisage, on pourrait d'ailleurs le décomposer en séparant sa partie imaginaire et sa partie réelle. Il faudrait cependant pour cela un très grand nombre d'itérations de tamisage. Pour donner un ordre de grandeur, au bout de 100 itérations, l'algorithme Algo. 5 le décompose à peu près en $d_1(t) = \cos(4t) + 0.35i \cos(3t)$ et $d_2(t) = 0.65 \cos(3t)$. L'algorithme Algo. 4 arrive quant à lui à quelque chose qui s'apparente à des versions assez déformées de $d_1(t) \cos(4t) + 0.25i \cos(3t)$ et $d_2(t) = 0.75 \cos(3t)$. Ainsi, même si on n'a pas d'exemple d'IMF parfait — c'est à dire avec un centre des armatures parfaitement nul —, il est tout à fait possible de rencontrer de tels rebroussements dans les IMFs puisqu'ils sont tout à fait compatibles avec un centre des armatures de faible amplitude par rapport au signal.

Une autre limite des algorithmes proposés peut être qu'ils sont incapables de dissocier deux fréquences de signes opposés lorsqu'elles sont proches en valeur absolue. Cette propriété a pour origine le fait que les algorithmes bivariés s'appuient sur des extrema de projections du signal où les fréquences négatives et positives sont naturellement mélangées. Elle sera mise en évidence en 1.4. Pour pouvoir séparer les fréquences positives et négatives, une solution pourrait être de faire comme dans l'algorithme « Complex Empirical Mode Decomposition » (cf 6.5.1) une séparation préalable du signal en partie analytique et antianalytique. Si une telle solution est sans doute intéressante dans

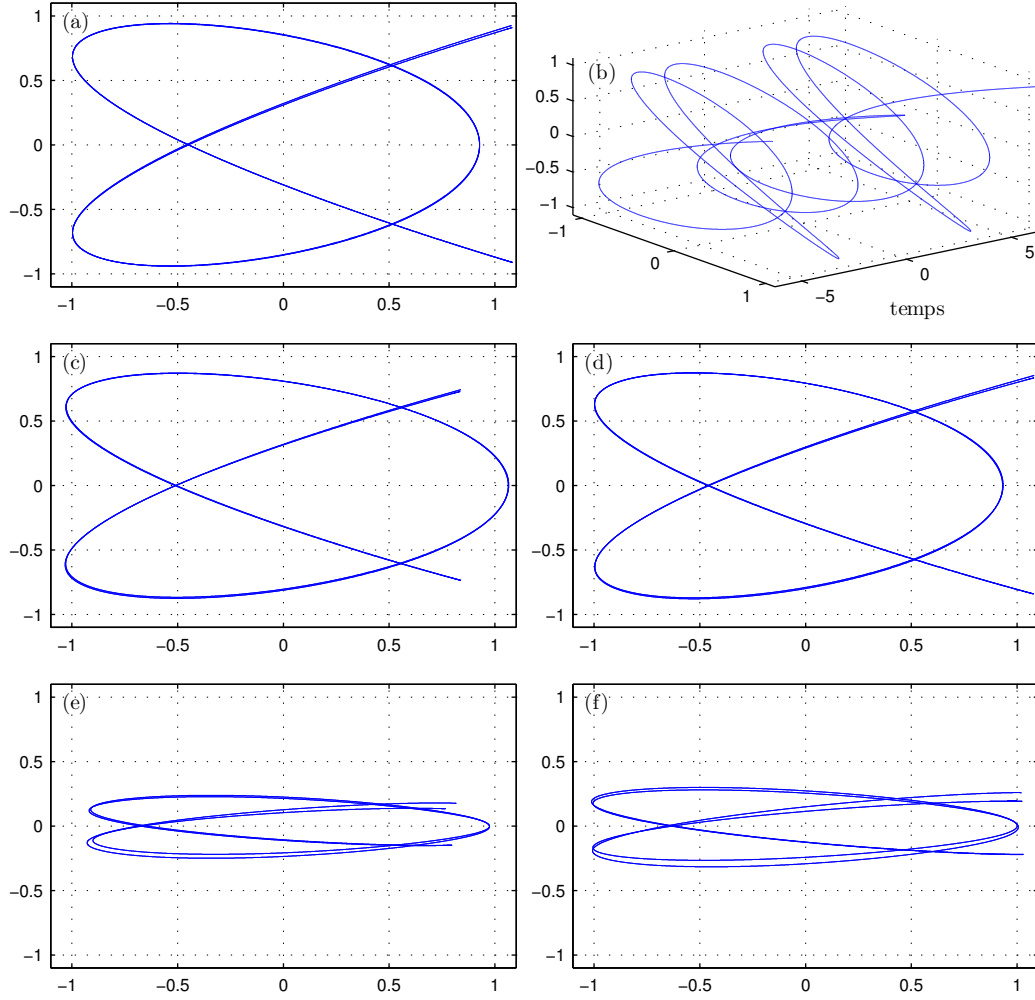


FIGURE 1.32 – Exemple d’IMF tournant qui change de sens fréquemment en rebroussant chemin. (a) la trajectoire du signal. (b) une vue tridimensionnelle de son évolution temporelle. Cet IMF est le premier IMF issu de la décomposition de $x(t) = \cos(4t) + i \cos(3t)$ avec 10 itérations de tamisage (algorithme Algo. 5, 256 directions). On peut également obtenir un IMF très similaire en prenant le premier IMF de la décomposition de $x(t) = \exp(i\pi \cos t)$ si le nombre d’itérations de tamisage est de l’ordre de 10. (c) et (e) les résultats de 10 et 100 itérations de tamisage sur ce signal à l’aide de l’algorithme Algo. 4 avec 256 directions. (d) et (f) idem avec l’algorithme Algo. 5.

certain cas, elle peut aussi être un inconvénient si le signal contient des composantes qui ne tournent pas circulairement. Par exemple, si on considère une rotation elliptique

$$x(t) = e^{i\theta} (a \cos(\omega t + \varphi) + ib \sin(\omega t + \varphi)), \quad (1.178)$$

celle-ci se décompose en deux composantes de Fourier aux fréquences ω et $-\omega$

$$x(t) = \frac{a+b}{2} e^{i(\theta+\varphi)} e^{i\omega t} + \frac{a-b}{2} e^{i(\theta-\varphi)} e^{-i\omega t}. \quad (1.179)$$

Il en découle que si on sépare préalablement les composantes analytiques et antianalytiques de $x(t)$, cette rotation elliptique est décomposée en deux IMFs distincts avec des fréquences opposées alors que les algorithmes bivariés sont capables de la considérer comme une unique composante.

Chapitre 2

Analyse

1 Cadre déterministe

L'objet de ce chapitre est l'étude des propriétés de la décomposition dans des situations déterministes. En vertu de la propriété de localité de l'EMD, l'approche proposée consiste à étudier comment se fait la décomposition lorsque le signal est constitué de composantes périodiques. Pour ce faire, on s'intéresse dans un premier temps au cas simplifié au maximum qui est celui d'un signal constitué d'une somme de deux sinusoides sur lequel on construira un modèle du comportement de l'EMD qui permet d'expliquer une grande partie des comportements observés et qui reste en partie valable lorsque les composantes sinusoidales sont modulées en amplitude et en fréquence. Par la suite, on verra comment ce modèle permet également d'expliquer la décomposition de signaux constitués de deux composantes faiblement non linéaires et éventuellement de plus de deux composantes. Enfin, on proposera une adaptation de ce modèle pour décrire le comportement des extensions bivariées de l'EMD introduites au chapitre 1 en 6.5.2.

1.1 Cas d'un mélange de 2 sinusoides

1.1.1 Introduction

La question de la séparation de deux composantes sinusoidales est reliée à la notion de résolution fréquentielle. À première vue, cette question peut paraître bien posée dans l'idée qu'il est toujours préférable de séparer deux composantes sinusoidales de fréquences différentes mais cet a priori n'est pas toujours justifié. Si on considère l'exemple simple d'un signal composé de deux composantes sinusoidales de mêmes amplitudes, on peut l'écrire sous la forme $x(t) = \cos 2\pi f_1 t + \cos 2\pi f_2 t$ avec une interprétation sous-jacente de deux composantes indépendantes, mais on peut tout aussi bien l'écrire sous la forme $x(t) = 2 \cos(\pi(f_1 - f_2)t) \cos(\pi(f_1 + f_2)t)$ avec une interprétation cette fois-ci en termes d'une composante sinusoidale modulée en amplitude. Cette interprétation est particulièrement pertinente lorsque $f_1 \approx f_2$ où elle est reliée à la notion de battement. Dans ce contexte, il n'est pas a priori évident que l'interprétation de type Fourier en termes de deux composantes soit plus pertinente que celle en termes d'une sinusoidale modulée en amplitude. Dans le cas de l'EMD, on voit intuitivement que la méthode va privilégier l'une ou l'autre des interprétations suivant que les fréquences sont très éloignées ou au contraire très rapprochées. L'objectif de cette section est de passer des considérations intuitives à une compréhension objective et quantitative de la manière dont l'EMD choisit l'une ou l'autre interprétation.

1.1.2 Objectifs

Dans l'esprit de l'exemple schématique du battement proposé en introduction, l'objectif de cette section est de caractériser, et si possible de modéliser, le comportement de l'EMD sur des mélanges de 2 sinusoïdes. Pour ce faire, on peut dans un premier temps chercher à répondre aux deux questions suivantes :

1. *Dans quelles conditions l'EMD sépare-t-elle les deux composantes sinusoïdales ?*
2. *Dans quelles conditions l'EMD considère-t-elle le signal comme une seule composante modulée en amplitude et en fréquence ?*

Sachant que l'EMD peut également avoir un comportement intermédiaire, on s'intéressera en fait à quantifier un « degré de séparation » mesurant la proximité d'une situation donnée avec les cas limites « 1 composante » et « 2 composantes ».

Au-delà de ces 2 cas limites, il n'est pas non plus a priori impossible d'observer des comportements ne correspondant à aucun de ces 2 cas et n'étant pas non plus intermédiaire. En conséquence, on tentera donc de répondre à une troisième question :

3. *Dans quelles conditions l'EMD décompose-t-elle le signal en plusieurs composantes qui ne sont pas simplement des combinaisons linéaires des deux composantes initiales ?*

1.1.3 Modèle de signal

L'algorithme standard de l'EMD étant généralement implanté pour traiter des signaux à temps discret, on s'intéressera dans un premier temps aux sommes de sinusoïdes à temps discret :

$$x[n] = a_1 \cos(2\pi f_1 n + \varphi_1) + a_2 \cos(2\pi f_2 n + \varphi_2), \quad 1 \leq n \leq N,$$

Sachant que l'EMD est covariante par rapport à la multiplication du signal par un scalaire non nul, on s'intéressera en fait à un modèle où seul intervient le rapport des amplitudes des deux composantes $a = a_2/a_1$:

$$x[n; a, f_1, f_2, \varphi_1, \varphi_2] = \cos(2\pi f_1 n + \varphi_1) + a \cos(2\pi f_2 n + \varphi_2), \quad 1 \leq n \leq N. \quad (2.1)$$

Le modèle étant symétrique, on se restreindra de plus au cas où $f_2 < f_1$ pour simplifier les discussions. La composante

$$x_1[n; f_1, \varphi_1] = \cos(2\pi f_1 n + \varphi_1) \quad (2.2)$$

sera par la suite appelée composante haute fréquence (HF) et

$$x_2[n; a, f_2, \varphi_2] = a \cos(2\pi f_2 n + \varphi_2) \quad (2.3)$$

composante basse fréquence (BF).

En pratique, ce modèle peut encore être simplifié si l'on s'intéresse à la limite dans laquelle les fréquences des sinusoïdes tendent vers zéro. Compte tenu de l'étude qui a été menée sur l'influence de l'échantillonnage dans l'EMD (cf 5.4), on sait a priori que le comportement de l'EMD dans cette limite doit se rapprocher de celui d'une EMD à temps continu sur le signal à temps continu

$$x(t) = \cos(2\pi f_1 t + \varphi_1) + a \cos(2\pi f_2 t + \varphi_2), \quad t \in \mathbb{R}.$$

Comme dans la limite continue l'EMD est covariante vis-à-vis des dilatations et des translations de l'axe temporel, on pourra en fait se restreindre à un modèle ne faisant intervenir que des quantités relatives a , $f = f_2/f_1$ et $\varphi = \varphi_2 - \varphi_1$:

$$x(t; a, f, \varphi) = \cos 2\pi t + a \cos(2\pi f t + \varphi), \quad t \in \mathbb{R}, \quad (2.4)$$

où comme dans le cas discret, on se limitera également au cas $f < 1$ pour simplifier. Enfin, ce dernier modèle étant considérablement plus simple que son homologue à temps discret, il permettra également de proposer une approche théorique modélisant le comportement de l'EMD dans la limite continue pour certaines valeurs des paramètres a , f et φ .

On définit comme dans le cas discret les deux composantes haute fréquence (HF)

$$x_1(t) = \cos 2\pi t, \quad (2.5)$$

et basse fréquence (BF)

$$x_2(t; a, f, \varphi) = a \cos(2\pi f t + \varphi). \quad (2.6)$$

1.1.4 Critères examinés

Pour apporter une réponse aux questions formulées précédemment, on définit 2 critères c_1 et c_2 qui quantifient la distance d'une situation donnée avec le cas limite « 1 composante » pour c_1 et « 2 composantes » pour c_2 . Ces 2 critères permettent de plus de répondre dans une certaine mesure à la troisième question qui correspond aux situations qui sont loin des 2 cas limites précédents. Pour y répondre plus précisément, on introduit un troisième critère c_3 qui quantifie la proximité d'une situation donnée avec l'ensemble des situations intermédiaires entre les deux cas limites.

Pour répondre à ces questions, on peut remarquer qu'il n'y a pas besoin de réaliser une EMD complète du signal mais seulement d'examiner le premier stade de la décomposition correspondant à la définition du premier IMF. En effet, celui-ci est nécessairement égal au signal dans le cas « 1 composante », à l'une des composantes dans le cas « 2 composantes » et doit être différent d'une combinaison linéaire des deux composantes dans le troisième cas. De plus, comme les IMFs sont ordonnés en fonction de leur échelle, le premier IMF dans le cas « 2 composantes » correspond nécessairement à la composante HF, ce qui permet de quantifier la proximité d'une situation avec ce cas (c_1) en examinant simplement la distance entre le premier IMF et la composante HF. Similairement, la proximité avec le cas « 1 composante » (c_2) est quantifiée par la distance entre le premier IMF et le signal initial et enfin la proximité avec toute situation intermédiaire (c_3) est quantifiée par la distance entre le premier IMF et l'espace des combinaisons linéaires des deux composantes. On utilisera pour les deux derniers critères deux normalisations différentes pour mieux faire apparaître certaines caractéristiques.

Les critères sont définis de la manière suivante :

$$c_1^{(n)}(a, f_1, f_2, \varphi_1, \varphi_2) = \frac{\left\| d_1^{(n)}[\cdot; a, f_1, f_2, \varphi_1, \varphi_2] - x_1[\cdot; f_1, \varphi_1] \right\|_{\ell_2}}{\left\| x_2[\cdot; a, f_2, \varphi_2] \right\|_{\ell_2}}, \quad (2.7)$$

$$c_{2,1}^{(n)}(a, f_1, f_2, \varphi_1, \varphi_2) = \frac{\left\| d_1^{(n)}[\cdot; a, f_1, f_2, \varphi_1, \varphi_2] - x[\cdot; a, f_1, f_2, \varphi_1, \varphi_2] \right\|_{\ell_2}}{\left\| x_1[\cdot; f_1, \varphi_1] \right\|_{\ell_2}}, \quad (2.8)$$

$$c_{2,2}^{(n)}(a, f_1, f_2, \varphi_1, \varphi_2) = \frac{\left\| d_1^{(n)}[\cdot; a, f_1, f_2, \varphi_1, \varphi_2] - x[\cdot; a, f_1, f_2, \varphi_1, \varphi_2] \right\|_{\ell_2}}{\left\| x_2[\cdot; a, f_2, \varphi_2] \right\|_{\ell_2}}, \quad (2.9)$$

$$c_{3,1}^{(n)}(a, f_1, f_2, \varphi_1, \varphi_2) = \frac{\left\| d_1^{(n)}[\cdot; a, f_1, f_2, \varphi_1, \varphi_2] - P_{x_1, x_2} \left(d_1^{(n)}[\cdot; a, f_1, f_2, \varphi_1, \varphi_2] \right) \right\|_{\ell_2}}{\left\| x_1[\cdot; f_1, \varphi_1] \right\|_{\ell_2}}, \quad (2.10)$$

$$c_{3,2}^{(n)}(a, f_1, f_2, \varphi_1, \varphi_2) = \frac{\left\| d_1^{(n)}[\cdot; a, f_1, f_2, \varphi_1, \varphi_2] - P_{x_1, x_2} \left(d_1^{(n)}[\cdot; a, f_1, f_2, \varphi_1, \varphi_2] \right) \right\|_{\ell_2}}{\left\| x_2[\cdot; a, f_2, \varphi_2] \right\|_{\ell_2}}, \quad (2.11)$$

avec les notations :

$d_1^{(n)}[\cdot; a, f_1, f_2, \varphi_1, \varphi_2] = \mathcal{S}^n x[\cdot; a, f_1, f_2, \varphi_1, \varphi_2]$, : premier IMF de la décomposition de $x[\cdot; a, f_1, f_2, \varphi_1, \varphi_2]$ obtenu à l'aide de n itérations de tamisage

$$\|u\|_{\ell_2} = \sqrt{\frac{1}{N} \sum_{n=1}^N u^2[n]}, : \text{norme } \ell_2 \text{ de } u,$$

$$P_{x_1, x_2}(u) = \frac{\langle u, x_1 \rangle}{\langle x_1, x_1 \rangle} x_1 + \frac{\langle u, x_2 \rangle}{\langle x_2, x_2 \rangle} x_2 : \text{projection de } u \text{ sur l'espace engendré par } x_1 \text{ et } x_2. \text{ (les paramètres ont été omis pour alléger les notations)}$$

Remarque : Le choix des normes ℓ_2 pour mesurer les distances se justifie dans la mesure où il permet de simplifier certains résultats théoriques. Les résultats expérimentaux sont quant à eux relativement peu sensibles à ce choix.

1.1.5 Simulations

1.1.5.1 Cas discret Une première série de simulations a été réalisée pour observer le comportement des différents critères sur le modèle discret (2.1) en fonction des paramètres $a \in \mathbb{R}_+^*$, $f_1 \in]0, 0.5]$, $f_2 \in]0, f_1]$, et dans une moindre mesure $\varphi_1 \in [0, 1/f_1]$ et $\varphi_2 \in [0, 1/f_2]$. L'objectif de ces premières simulations étant principalement qualitatif, le nombre d'itérations de tamisage a été fixé arbitrairement à 10. L'influence de ce paramètre important sera étudiée plus avant lors de la deuxième série de simulations dont l'objectif est plus quantitatif.

Méthodologie Pour les 9 valeurs de a : 10^{-2} , $10^{-1.5}$, 10^{-1} , $10^{-0.5}$, 1 , $10^{0.5}$, 10 , $10^{1.5}$ et 10^2 , on calcule les premiers IMFs $d_1^{(10)}[\cdot; a, f_1, f_2, \varphi_1, \varphi_2]$ issus des signaux $x[\cdot; a, f_1, f_2, \varphi_1, \varphi_2]$, avec dans un premier temps $f_1 = k_1/480$, $0 \leq k_1 \leq 240$, $f_2 = k_2/480$, $0 \leq k_2 < k_1$, $\varphi_i = p/(4f_i)$, $0 \leq p \leq 3$. Dans un deuxième temps, on réalise les mêmes simulations en divisant les fréquences f_1 et f_2 par 10 pour s'approcher du cas continu. Pour chaque cas, on calcule les quantités c_1 , $c_{2,1}$, $c_{2,2}$, $c_{3,1}$ et $c_{3,2}$. Les signaux ont une taille de 3840 points dont seuls les 1920 points centraux sont conservés pour évaluer les différents critères. Le fait de ne considérer que la partie centrale de l'IMF permet de limiter les effets de bords liés à la nécessité de prolonger artificiellement les séquences d'extrema pour définir les interpolations sur les bords du signal. Le choix de $1920 = 8 \times 240$ points, lié aux valeurs des fréquences f_1 et f_2 utilisées, permet de réaliser toutes les opérations d'intégration sur un nombre entier de périodes pour chacune des deux composantes évitant ainsi des effets parasites. Enfin, on considère des fréquences multiples de $1/240$ parce que 240 a beaucoup de diviseurs ce qui permet de faire mieux apparaître les singularités aux fréquences $f = p/q$ avec p et q relativement « petits ».

Résultats Les valeurs des critères moyennées par rapport aux paramètres de phase φ_1 et φ_2 sont représentées Fig. 2.1 à Fig. 2.5, pour f_1 et f_2 entre 0 et 0.5, et Fig. 2.6 à Fig. 2.10, pour f_1 et f_2 entre 0 et 0.05. Toutes les images ont été artificiellement coupées à 2 pour avoir la même dynamique partout.

De manière générale, la première série de figures est difficile à lire dans la mesure où d'importants effets d'échantillonnage cachent souvent les variations ayant trait à la notion de séparation. Ces effets d'échantillonnage se manifestent sur les images par des bandes verticales lorsque $a < 1$ et horizontales sinon. Le caractère vertical ou horizontal de ces bandes montre qu'elles sont liées à une des deux composantes seulement, qui se trouve assez logiquement être celle qui a la plus grande amplitude. De plus, on observe que ces bandes sont centrées sur les fréquences de la forme $1/(2k+1)$ et disparaissent pour les fréquences $1/(2k)$, tout comme les effets d'échantillonnage sur une sinusoïde (cf 5.4.1). Au-delà de ces effets d'échantillonnage, on observe que dès lors que les fréquences ne sont

pas trop élevées $f_i \lesssim 0.25$, bon nombre de caractéristiques semblent dépendre essentiellement du rapport f_2/f_1 .

Pour y voir plus clair, on peut s'intéresser plus particulièrement à la deuxième série de simulations où les fréquences ne vont que jusqu'à 0.05, ce qui permet de limiter les effets d'échantillonnage. On observe de plus beaucoup plus nettement sur cette deuxième série le fait que la dépendance des critères vis-à-vis des deux fréquences passe essentiellement par leur rapport. De manière plus détaillée, les différents critères apportent chacun un certain nombre d'informations :

c_1 indique que les deux composantes peuvent être séparées ssi leur rapport est inférieur à une certaine limite. Cette limite dépend du rapport d'amplitudes, mais semble-t-il seulement quand celui-ci est supérieur à 1.

$c_{2,1}$ et $c_{2,2}$ permettent de voir quand l'EMD perçoit le signal comme une seule composante. On observe ainsi grâce à $c_{2,2}$ que cela semble être le cas lorsque $a < 1$ et que le rapport de fréquences est supérieur au rapport limite observé sur c_1 . Lorsque $a > 1$, $c_{2,2}$ n'apporte pas beaucoup d'information mais $c_{2,1}$ permet de voir que le comportement en fonction du rapport de fréquences est plus compliqué et alterne les cas où le signal est considéré comme une unique composante et d'autres où on ne sait pas bien ce qui se passe puisque, d'après c_1 , les deux composantes ne sont pas non plus séparées.

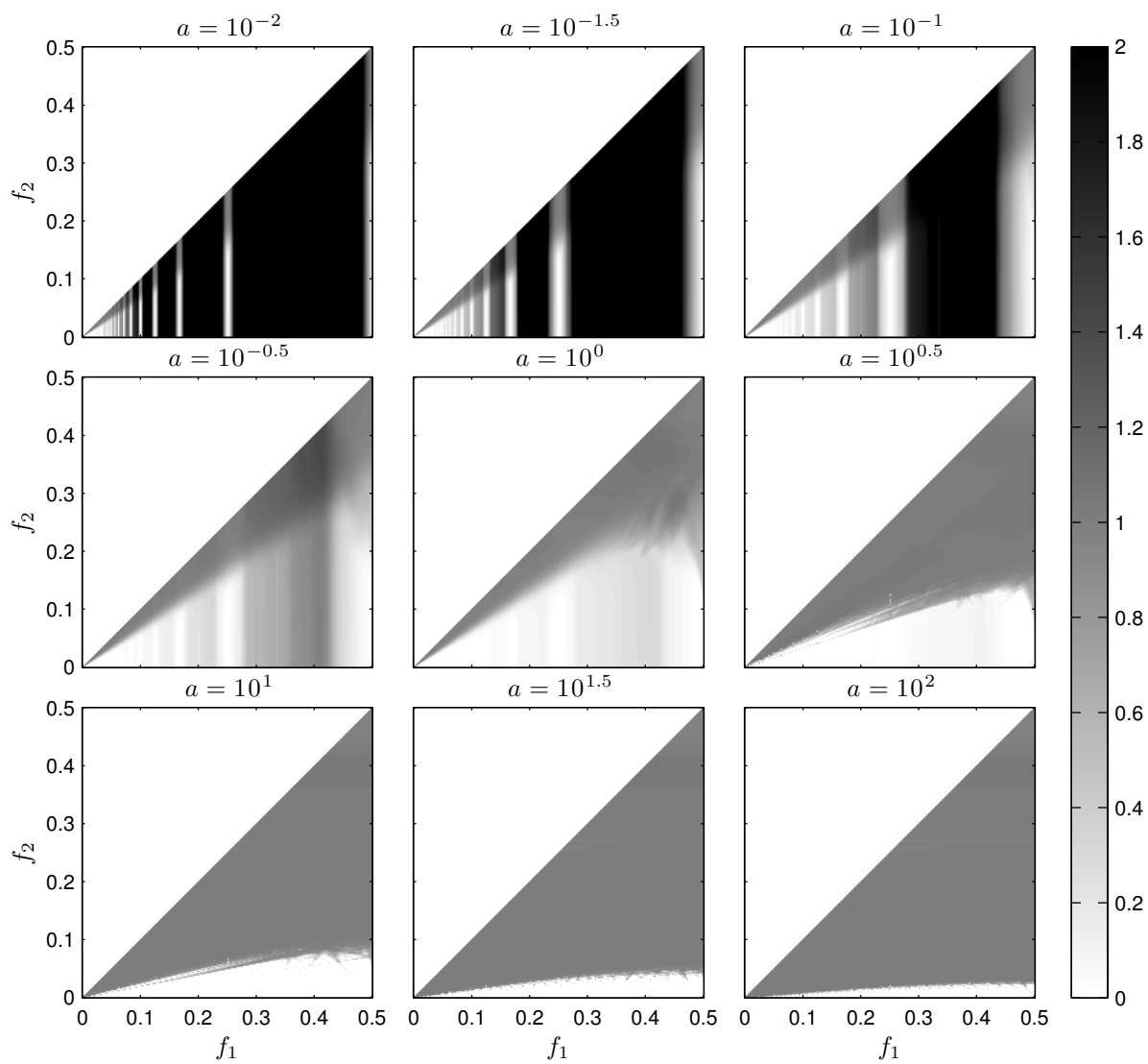
$c_{3,1}$ et $c_{3,2}$ apportent un autre éclairage sur les cas où les deux composantes ne sont pas séparées et où le signal n'est pas non plus considéré comme une unique composante. $c_{3,1}$ permet en particulier d'observer que les comportements pour $a > 1$ mal identifiés par les autres critères ne sont pas intermédiaires entre les deux situations « 2 composantes » et « 1 composante » mais font intervenir d'autres composantes que les deux initiales. $c_{3,2}$ enfin, permet d'observer (sur les images $a = 10^{-0.5}$ et $a = 1$) que la transition entre les deux états « 2 composantes » et « 1 composante » lorsque $a < 1$ ne se fait pas uniquement en passant par des combinaisons linéaires des 2 composantes mais qu'il y a dans les situations intermédiaires une (très) faible proportion d'autres composantes.

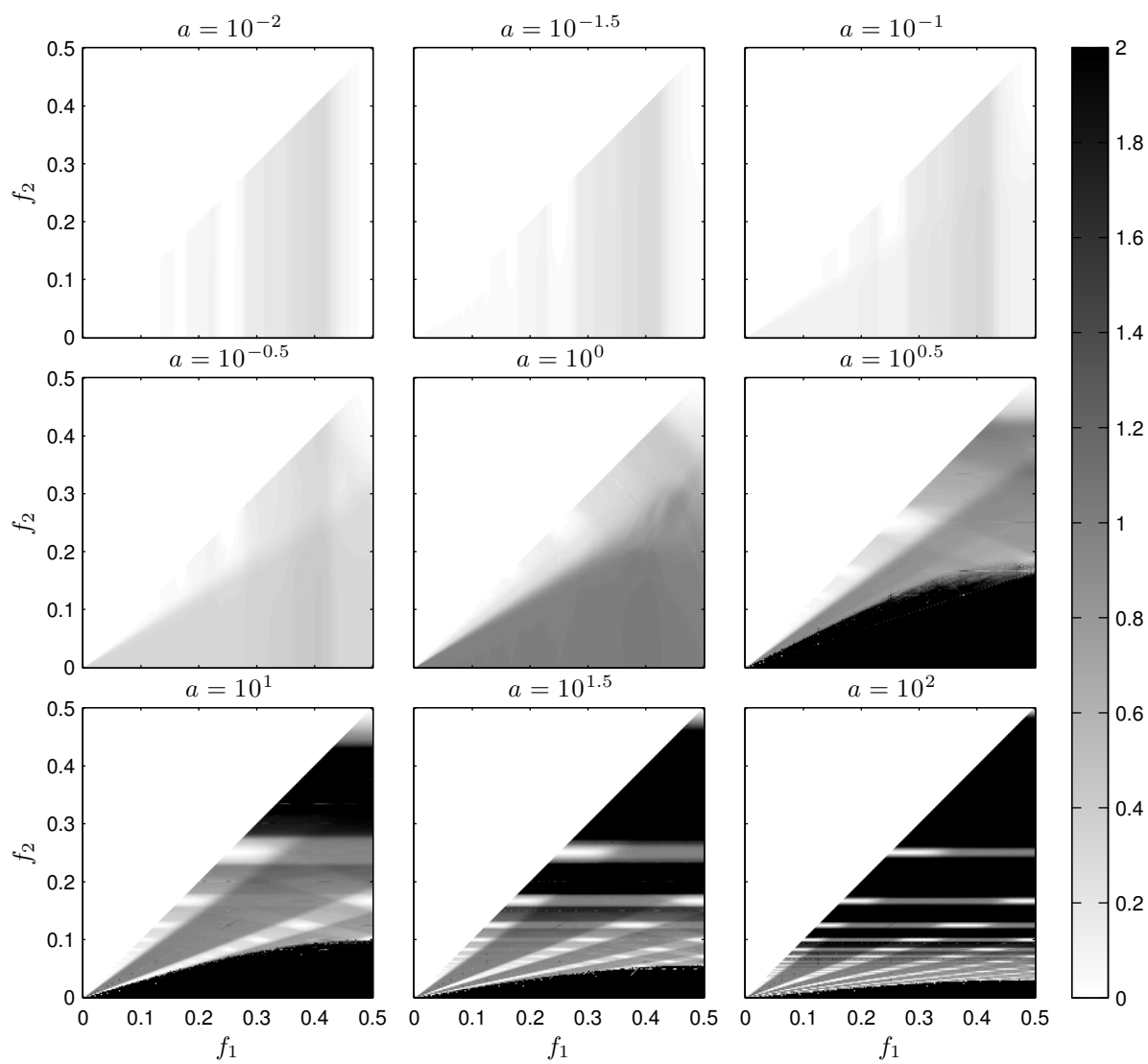
1.1.5.2 Cas continu Suite à la première série de simulations, on réalise une deuxième série dans l'objectif de caractériser le comportement de l'EMD dans la limite continue.

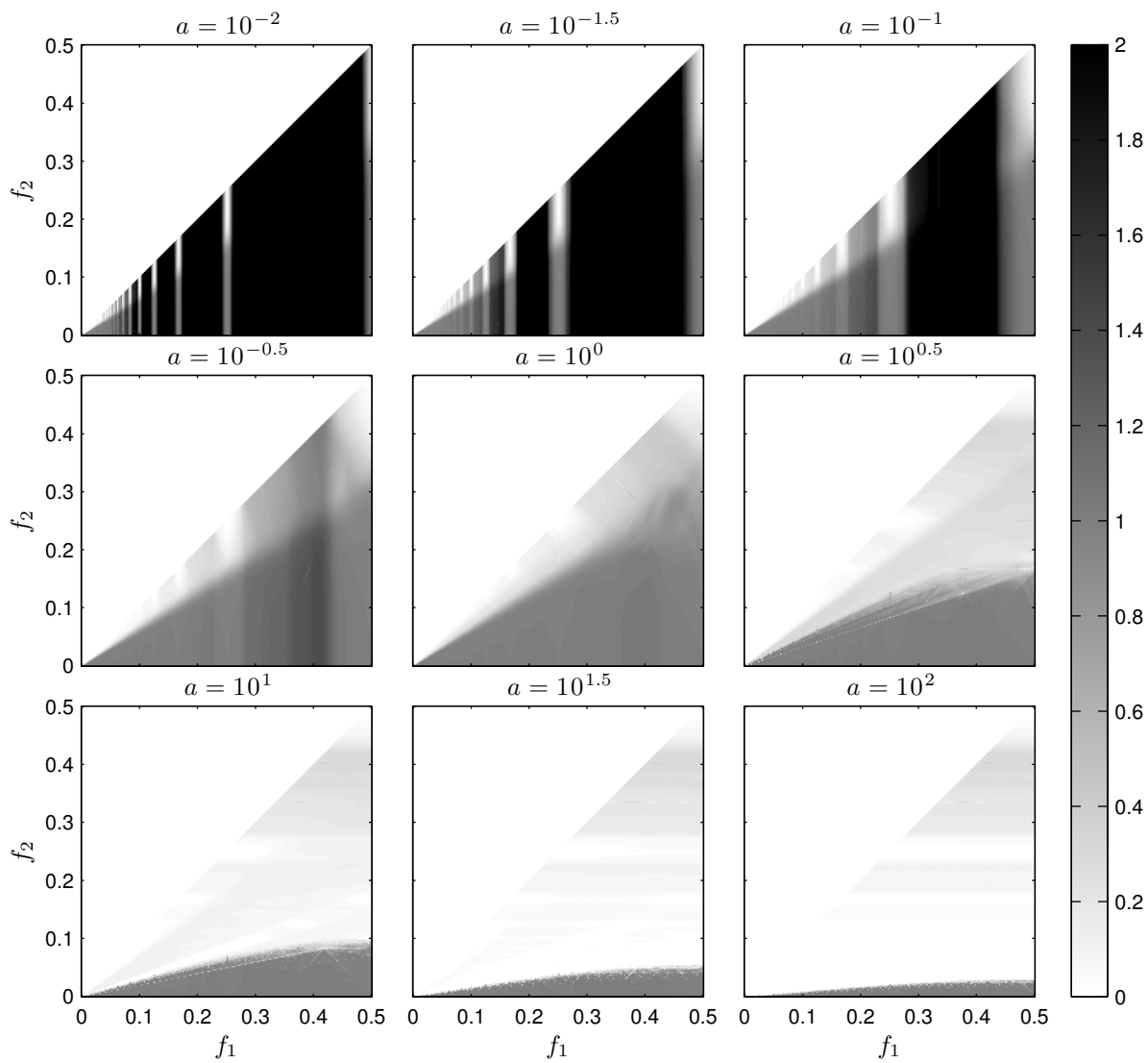
Méthodologie Outre les précautions prises pour la première série de simulations, on s'attache ici à également limiter tout effet lié à l'échantillonnage. Pour ce faire, on utilise les résultats précédemment établis quant aux effets de l'échantillonnage sur l'EMD d'une sinusoïde qui montrent que ceux-ci s'annulent pour les fréquences $f = 1/(2k)$, $k \in \mathbb{N}$. Dans le cas d'une somme de deux sinusoïdes, on observe que les effets d'échantillonnage disparaissent (ou presque) de manière similaire lorsque la composante dominante, c'est-à-dire celle dont les extrema sont les plus proches de ceux du signal somme (cf discussion sur les extrema 1.1.6.1), a une fréquence de la forme $1/(2k)$. Suivant le domaine d'intérêt par rapport aux paramètres a et $f = f_2/f_1$ du modèle continu (2.4), on présentera donc plutôt des simulations, nécessairement réalisées sur le modèle discret, où $f_1 = 1/32$ et $f_2 = kf_1/240$ ou alors $f = f_2/f_1 = k/240$ avec $f_2 = 1/(2p)$ où $2p$ est l'entier pair le plus proche de $32/f$.

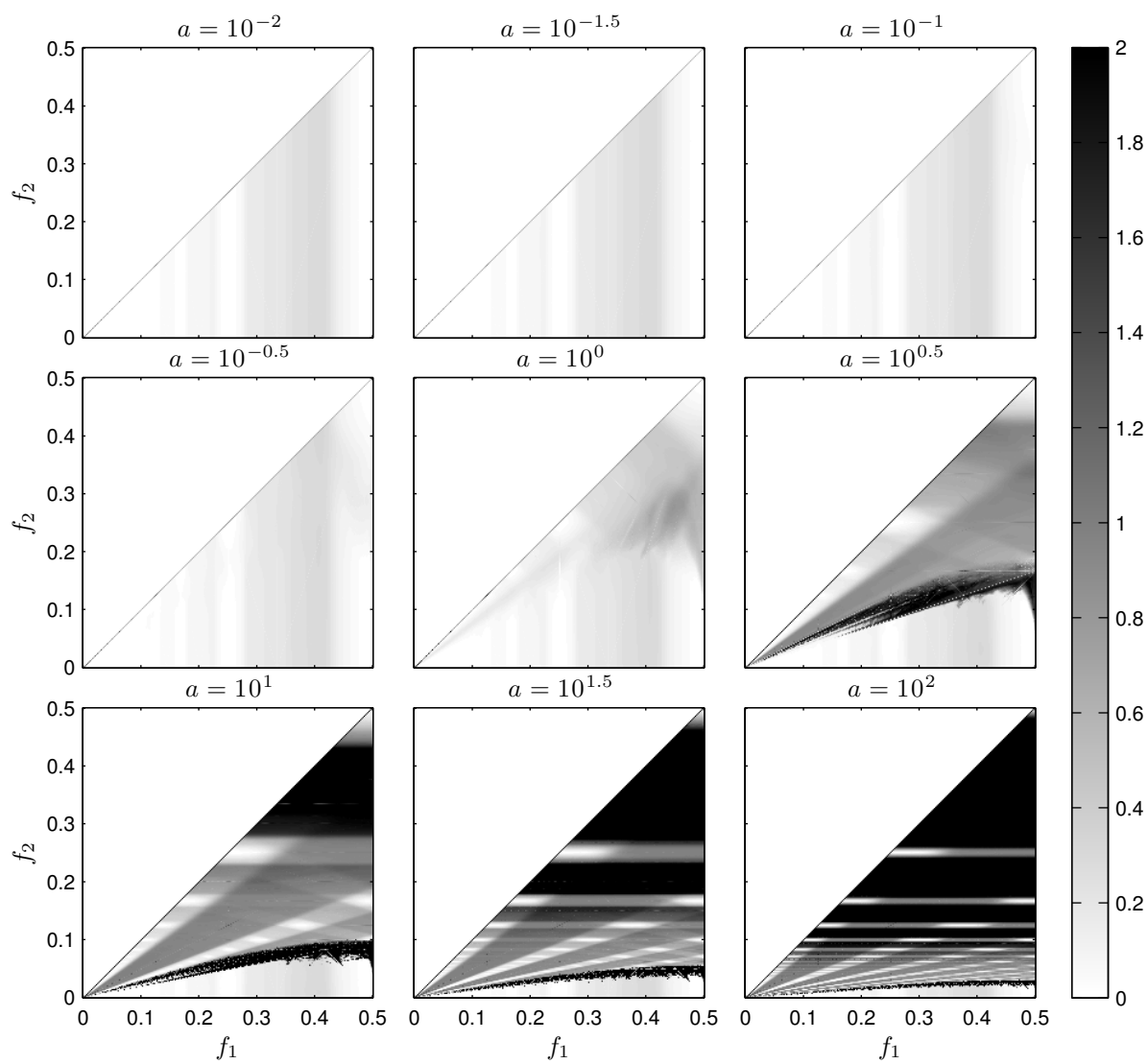
Pour amorcer la théorie développée dans la suite, on s'intéressera également à l'évolution de la densité d'extrema $d_e(a, f, \varphi)$ (ou nombre moyen d'extrema par unité de longueur) du signal somme $x(t; a, f, \varphi)$. On observera en particulier la quantité

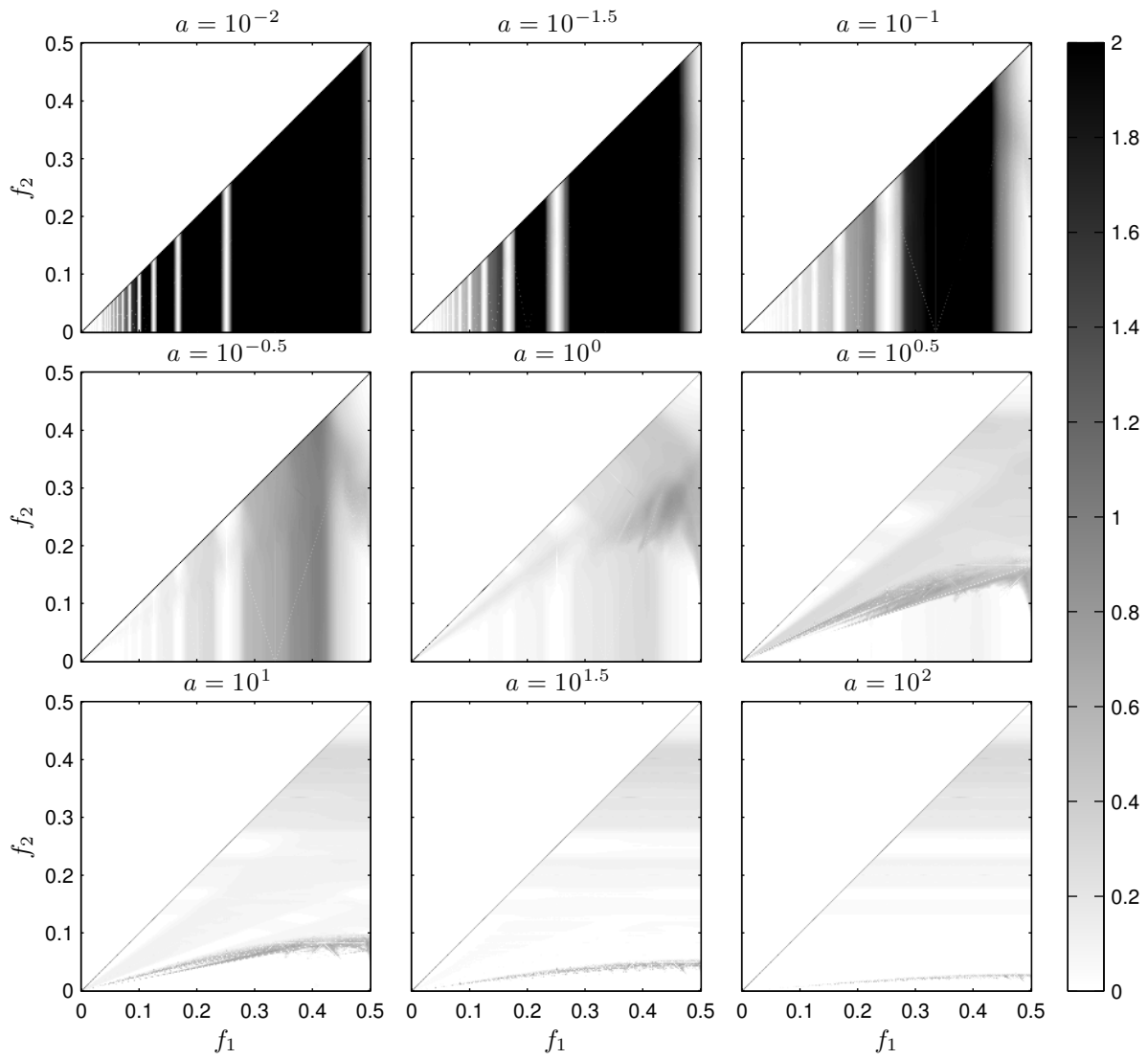
$$\tilde{d}_e(a, f, \varphi) = \frac{d_e(a, f, \varphi) - 2f}{2 - 2f} \quad (2.12)$$

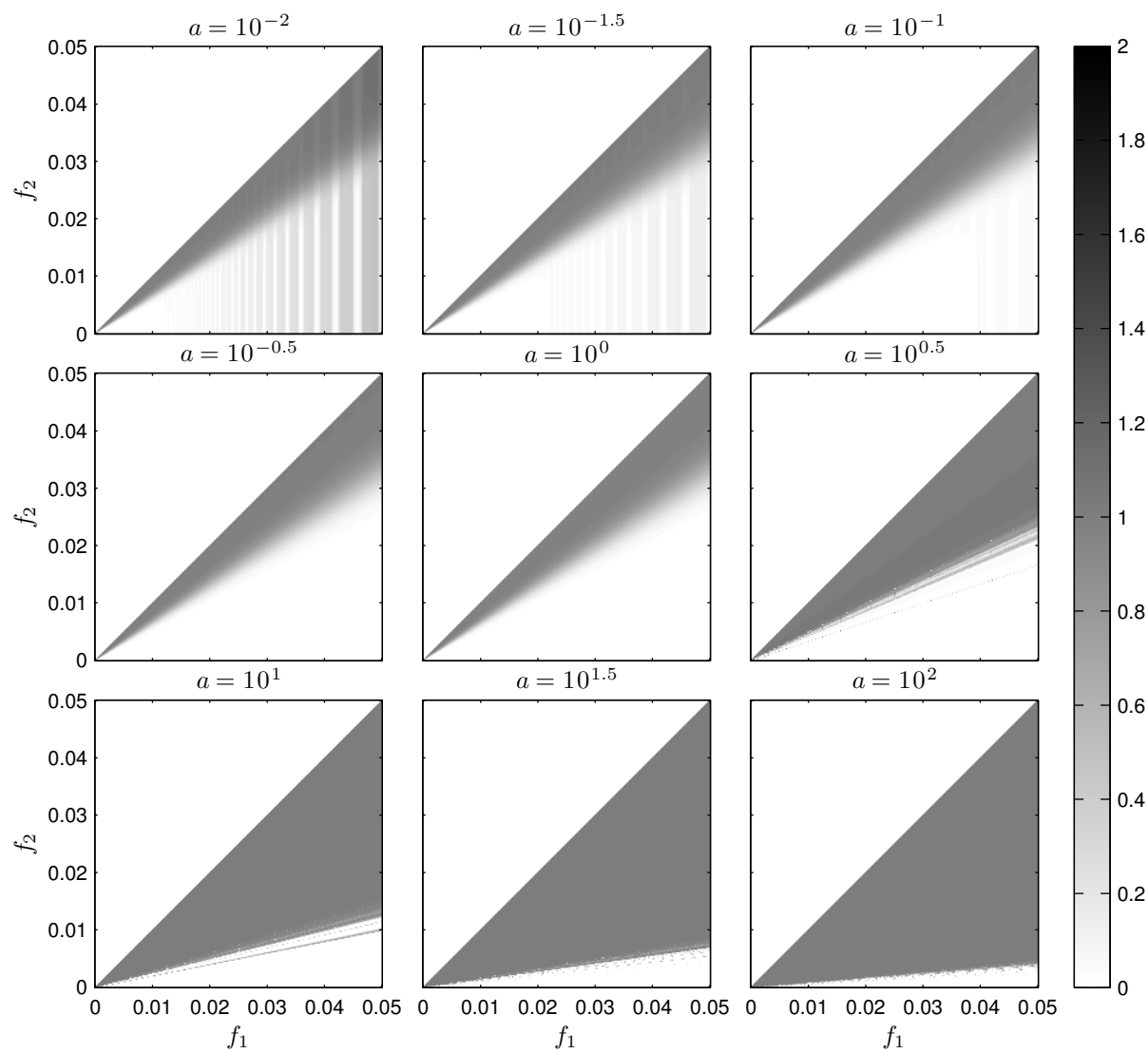
FIGURE 2.1 – c_1

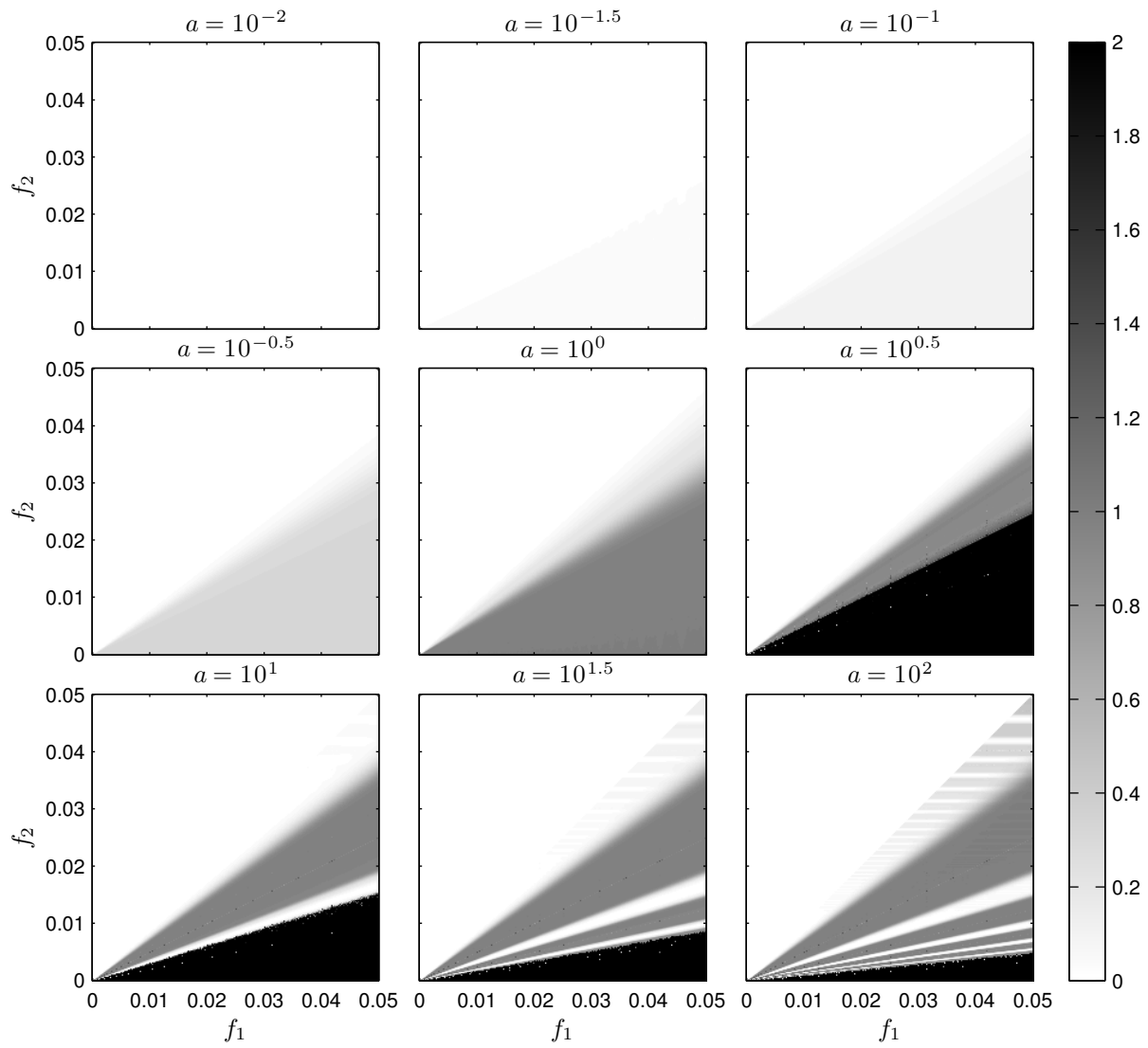
FIGURE 2.2 – $c_{2,1}$

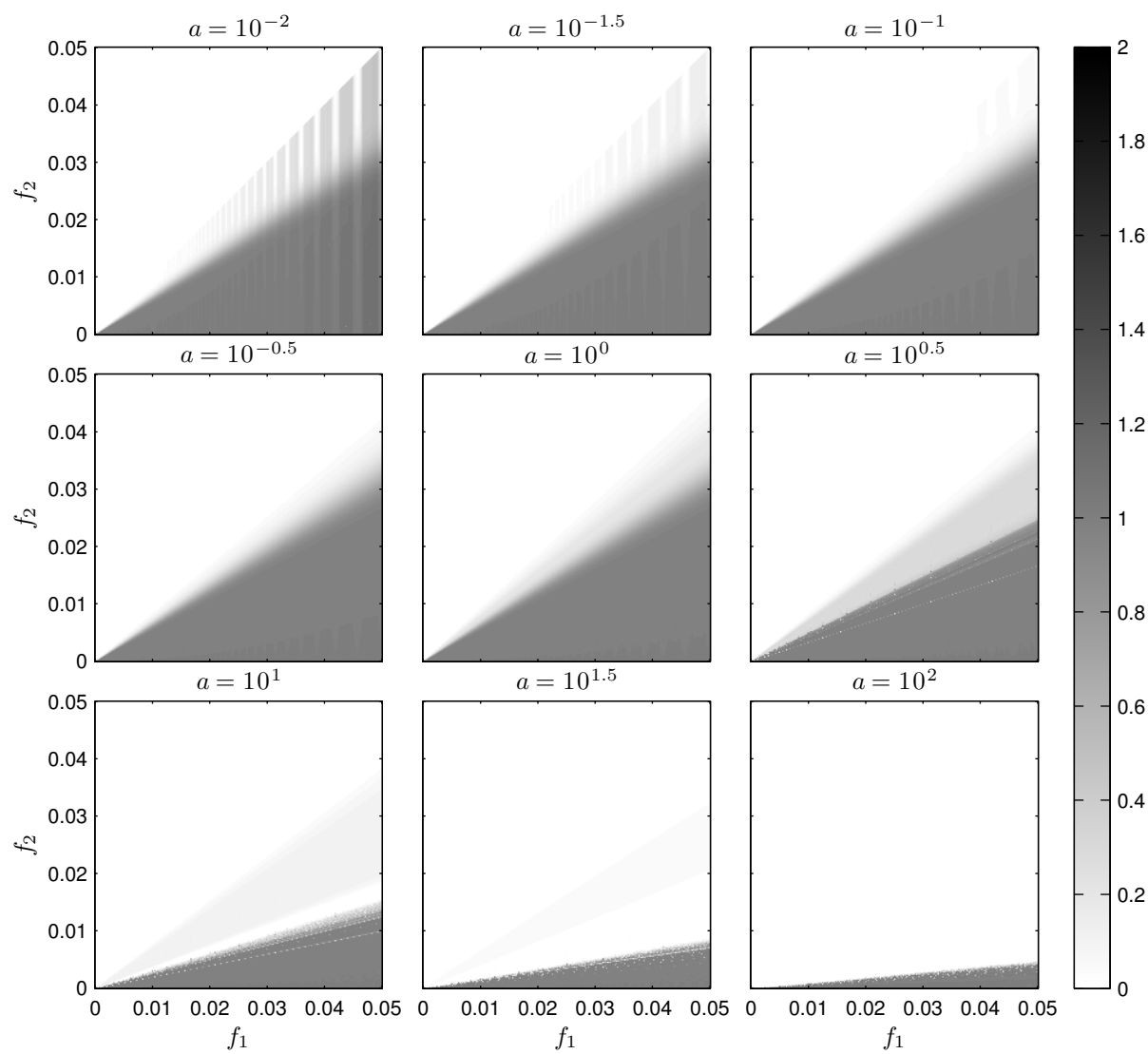
FIGURE 2.3 – $c_{2,2}$

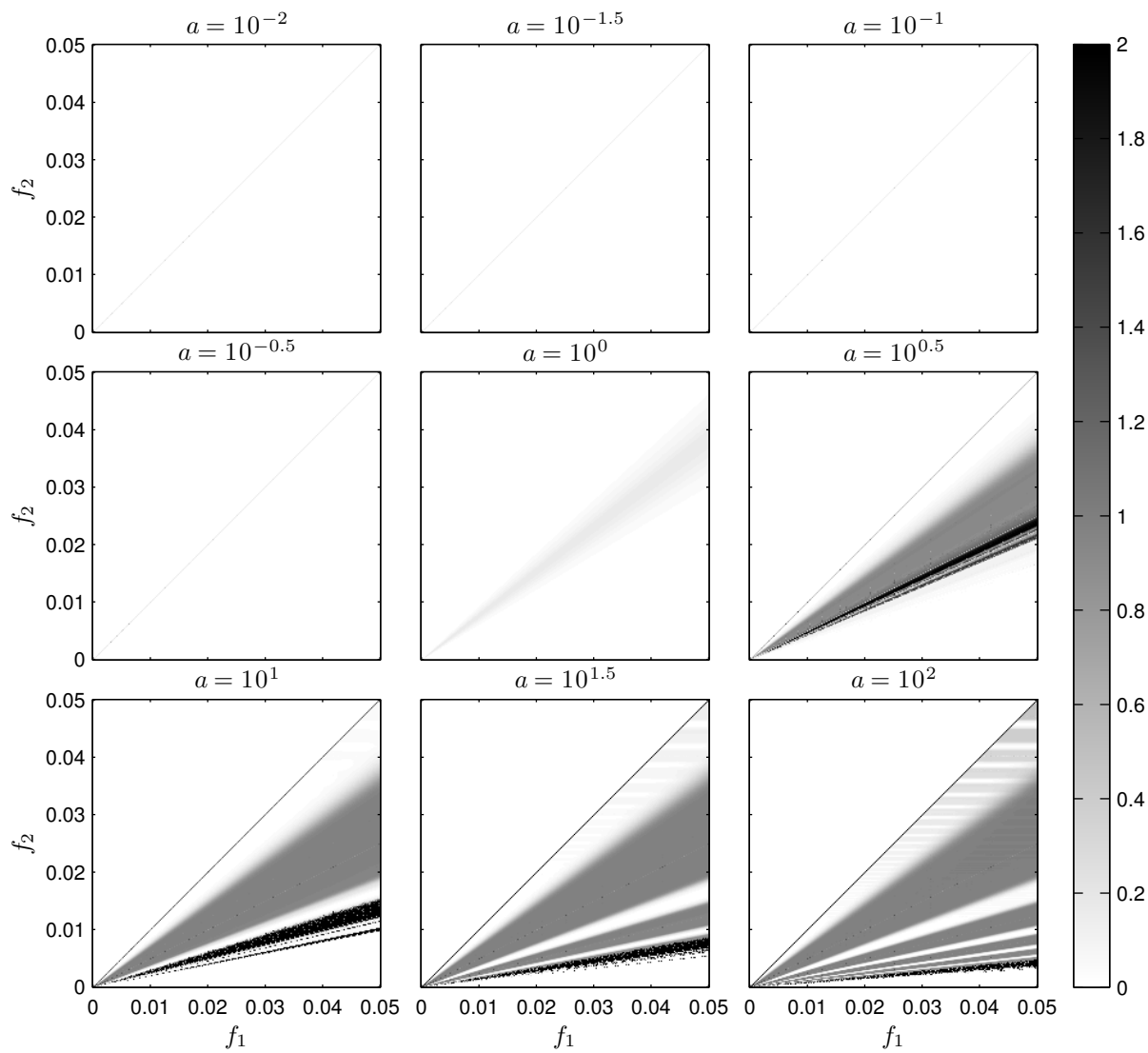
FIGURE 2.4 – $c_{3,1}$

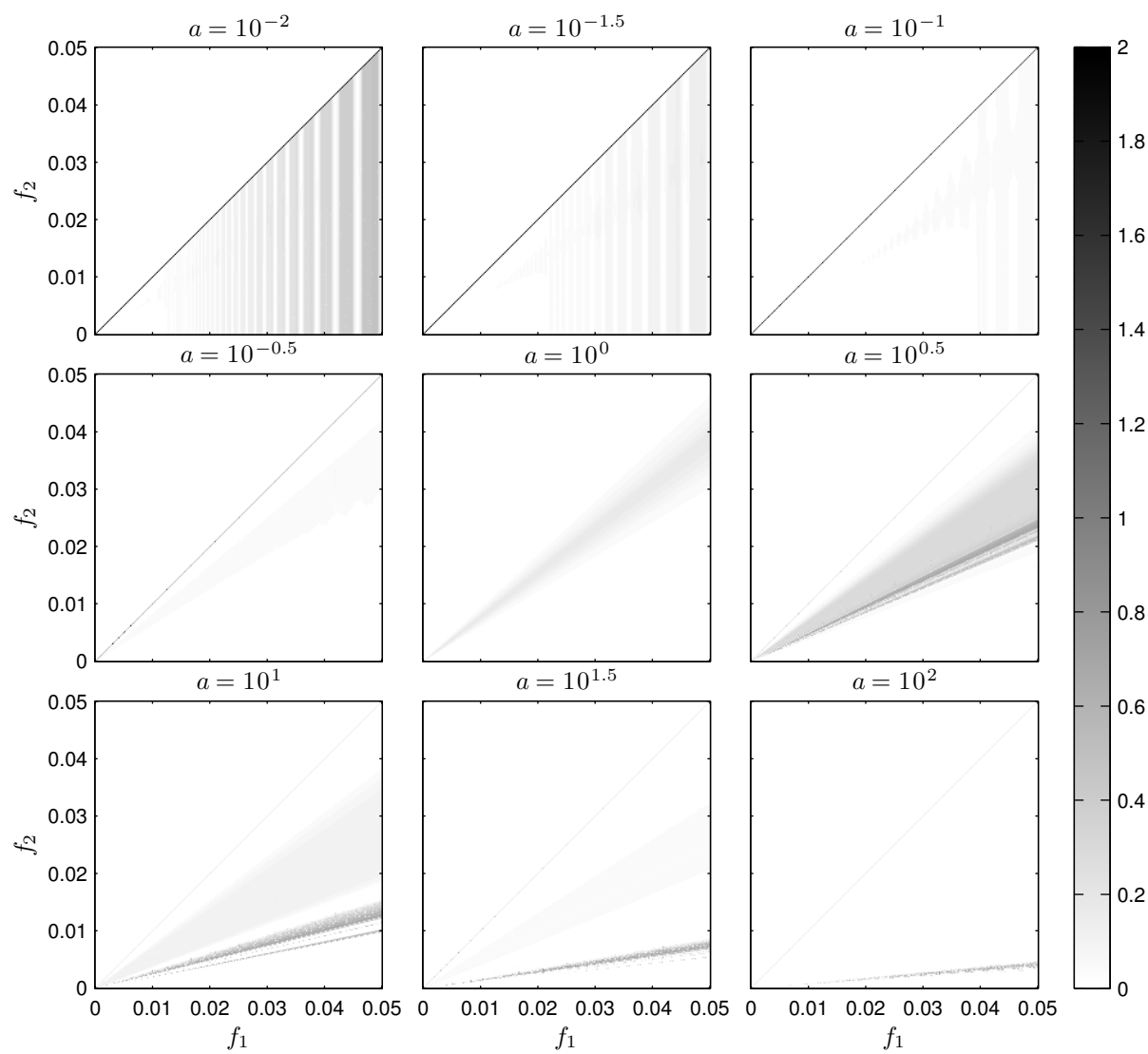
FIGURE 2.5 – $c_{3,2}$

FIGURE 2.6 – c_1

FIGURE 2.7 – $c_{2,1}$

FIGURE 2.8 – $c_{2,2}$

FIGURE 2.9 – $c_{3,1}$

FIGURE 2.10 – $c_{3,2}$

conçue pour valoir 0 lorsque la densité d'extrema du signal somme est identique à celle de la composante BF (qui vaut $2f$ dans le modèle continu(2.4)) et 1 lorsqu'elle est identique à celle de la composante HF (qui vaut 2).

Résultats Les valeurs des critères moyennées par rapport à la différence de phases φ sont représentées Fig. 2.11 (c_1), Fig. 2.12 ($c_{2,1}$ et $c_{2,2}$) et Fig. 2.13 ($c_{3,1}$ et $c_{3,2}$). Les écart types correspondant à certaines de ces figures sont représentés Fig. 2.14 ainsi que la densité d'extrema réduite $\tilde{d}_e(a, f, \varphi)$ (2.12) moyennée par rapport à φ .

De manière générale, les différentes quantités font apparaître deux grandes zones dans le plan (a, f) avec des comportements homogènes à l'intérieur de ces zones et une zone de transition entre les deux. Celles-ci sont délimitées sur les différentes images par les courbes tracées en tirets et tiret-point, qui correspondent respectivement aux équations $af = 1$ et $af^2 = 1$, et qui délimitent en fait très précisément les zones où la densité d'extrema est exactement celle de la composante HF ($af < 1$) ou BF ($af^2 > 1$), résultat mis en évidence Fig. 2.14, et démontré en 1.1.6.1. Plus précisément, la courbe $af < 1$ délimite bien la zone de transition entre les 2 grandes zones mais seulement pour $f > 1/3$. Pour $f < 1/3$, on expliquera en 1.1.6.1 pourquoi la frontière de la zone de transition est au dessus de la courbe $af = 1$ et on en déduira l'équation, correspondant à la courbe en pointillés sur les images.

Intéressons-nous maintenant plus précisément au comportement de l'EMD au sein de chaque zone.

zone $af < 1$, jusqu'à la courbe en pointillés : Dans cette zone, on observe que le comportement des critères normalisés par la norme de la composante BF (c_1 , $c_{2,2}$ et $c_{3,2}$) est pratiquement indépendant du rapport d'amplitude a , et ce d'autant plus qu'on s'écarte de la zone de transition. Pour un nombre d'itérations donné, le comportement de l'EMD évolue progressivement entre le cas où les deux composantes sont parfaitement séparées $f \rightarrow 0$ et celui où elles sont considérées comme une unique composante $f \rightarrow 1$. Cette évolution en fonction de f , se fait en fait avec un saut assez doux (allure sigmoïde) dont la position dépend du nombre d'itérations : proche de 0.5 pour 1 itération, celle-ci se rapproche de 1 quand le nombre d'itérations augmente. Au cours de cette évolution, de plus, on constate que le premier IMF est en très bonne approximation tout le temps une combinaison linéaire des deux composantes initiales. En résumé, tout semble en fait indiquer un comportement analogue à celui d'un filtre linéaire où le premier IMF serait obtenu en filtrant le signal à l'aide d'un filtre passe-haut dont la réponse en fréquence est justement la sigmoïde observée pour l'évolution de c_1 en fonction de f quand $a \rightarrow 0$.

zone $af^2 > 1$: Dans cette zone, on peut déjà observer grâce au critère c_1 que quelles que soient les valeurs de a et f , il n'est jamais possible de séparer les deux composantes. Au-delà de ça, on observe que contrairement au cas précédent, ce sont les critères normalisés par la norme de la composante HF ($c_{2,1}$ et $c_{3,1}$) qui semblent indépendants du rapport d'amplitude a ¹. L'étude de ces critères montre en fait que le comportement de l'EMD alterne entre considérer que le signal est une seule composante si f est plutôt proche de l'ensemble des $\{1/(2k+1), k \in \mathbb{N}\}$, ou alors proposer un premier IMF qui n'est pas combinaison linéaire des deux composantes HF et BF si f est plutôt proche de $\{1/(2k+1), k \in \mathbb{N}\}$. Plus précisément, on observe que dans ce second cas, la distance entre le premier IMF et le signal est proche de la norme de la composante HF et que de plus, le premier IMF est à la même distance de l'espace des combinaisons linéaires des composantes HF et BF. En en déduit donc que dans ce second cas, l'EMD produit le premier IMF en ajoutant au signal une nouvelle composante de même norme que la composante HF et qui est orthogonale aux deux composantes initiales.

1. c_1 semble également indépendant du rapport d'amplitude mais on pourra voir grâce à la modélisation qu'il en dépend à la manière de $(a+1)/a$ et que cette dépendance est donc peu visible dans cette zone où a est plutôt grand. La

zone de transition : Le comportement de l'EMD dans cette dernière zone est nettement plus compliqué que dans les deux précédentes. En pratique, on observe souvent un comportement intermédiaire entre les deux précédents avec des intervalles où le comportement est proche de celui de la zone $af < 1$ intercalés avec d'autres où il est proche de celui de la zone $af^2 > 1$ (cf Fig. 2.15). La largeur des différents intervalles n'évolue pas de manière simple en fonction des paramètres a et f et est a priori liée à des détails comme la proximité du rapport de fréquences avec certains nombres rationnels. Une trace de ces effets dans le comportement de c_1 ou de $c_{2,2}$ est le caractère irrégulier de la frontière $c_1 = 0.5$ (ou encore mieux de $c_{2,2} = 0.5$), où l'on peut voir en particulier des excroissances pour tous les rapports $f = 1/k$, $k \in \mathbb{N}^*$. De plus, on peut dans une certaine mesure expliquer l'évolution de ces frontières en fonction du nombre d'itérations : les intervalles où le comportement est proche de celui de la zone $af < 1$ ont tendance à s'allonger lorsque le nombre d'itération augmente jusqu'à faire disparaître ceux où le comportement est proche de la zone $af^2 > 1$. Intuitivement, ce comportement se justifie par le fait que des extrema sont susceptibles d'apparaître au cours du processus de tamisage, mais qu'il est en revanche beaucoup plus rare de les voir disparaître. On observe ainsi un déplacement de la frontière vers la zone $af^2 > 1$, limité par certaines valeurs de f comme les rapports $f \rightarrow 1/k$ cités précédemment pour lesquels les longueurs des intervalles tendent vers l'infini. Lorsque $f = 1/k$, le comportement dépend de la différence de phases φ et est globalement soit celui de la zone $af < 1$ soit celui de la zone $af^2 > 1$. Par conséquent, il n'y a pas d'évolution des longueurs des intervalles avec le nombre d'itérations. Pour des rapports de fréquences moins particuliers, la différence de phase a essentiellement une influence sur les positions des intervalles mais pas sur leurs longueurs.

1.1.6 Modélisation dans le cas continu

1.1.6.1 A propos des extrema

Quelques résultats sur les positions des extrema Les extrema jouent un rôle central dans l'algorithme de l'EMD et il n'est donc pas surprenant de voir que leurs propriétés, notamment leur densité examinée lors des simulations, semblent fortement liées au comportement de l'algorithme. En effet, il paraît intuitivement clair que l'EMD ne peut extraire une composante que si elle peut détecter des extrema liés à cette dernière. Dans le cas de la somme de deux sinusoïdes (2.4), il n'y a que deux cas asymptotiques pour lesquels on peut aisément caractériser complètement les positions des extrema : si l'une des deux composantes a une amplitude infiniment plus grande que l'autre ($a \rightarrow 0$ ou $a \rightarrow \infty$). En pratique cependant, les simulations semblent indiquer que le comportement de l'EMD est très semblable aux cas asymptotiques dès lors que la densité d'extrema est identique à celle de l'une des deux composantes. On a observé à ce propos que la densité d'extrema était égale à celle de la composante HF si $af < 1$ et à celle de la composante BF si $af^2 > 1$ avec une transition relativement douce entre les deux. Commençons par démontrer ces résultats.

Proposition 2. *La densité d'extrema*

$$d_e(a, f, \varphi) = \lim_{T \rightarrow \infty} \frac{\text{nombre d'extrema de } x(\cdot; a, f, \varphi) \text{ dans } \left[-\frac{T}{2}, \frac{T}{2}\right]}{T} \quad (2.13)$$

seule trace discernable (surtout sur les vues 3D pour 3 et 10 itérations) de cette évolution est en fait la zone où $c_1 > 1$ pour $f \gtrsim 0.5$ et $\log_{10} a \approx 0.5$.

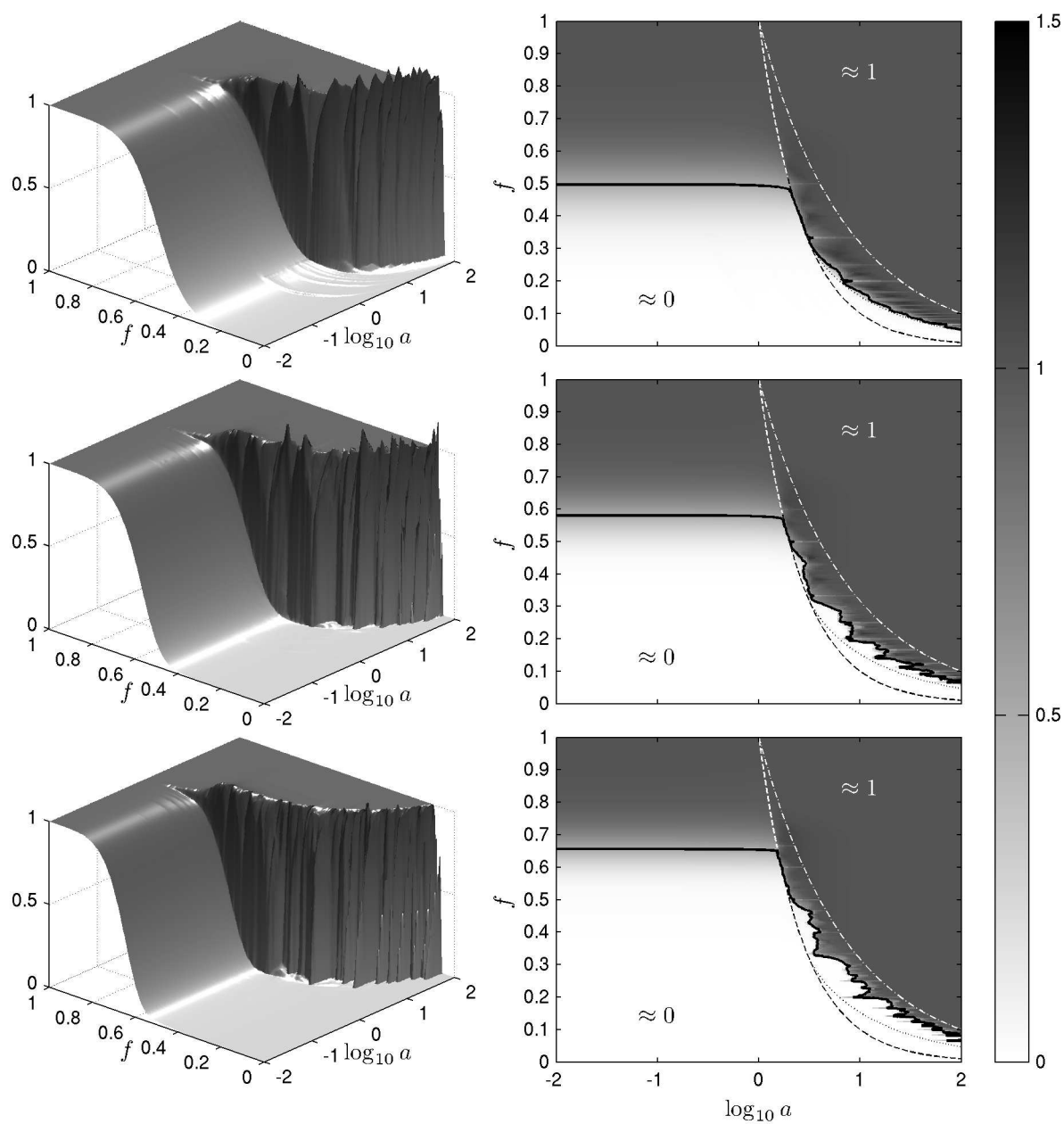


FIGURE 2.11 – Colonne de gauche : vue 3D de c_1 moyenné par rapport à φ pour (de haut en bas) 1, 3 et 10 itérations de tamisage. Colonne de droite : idem en images. Les lignes épaisses correspondent aux courbes de niveau = 0.5.

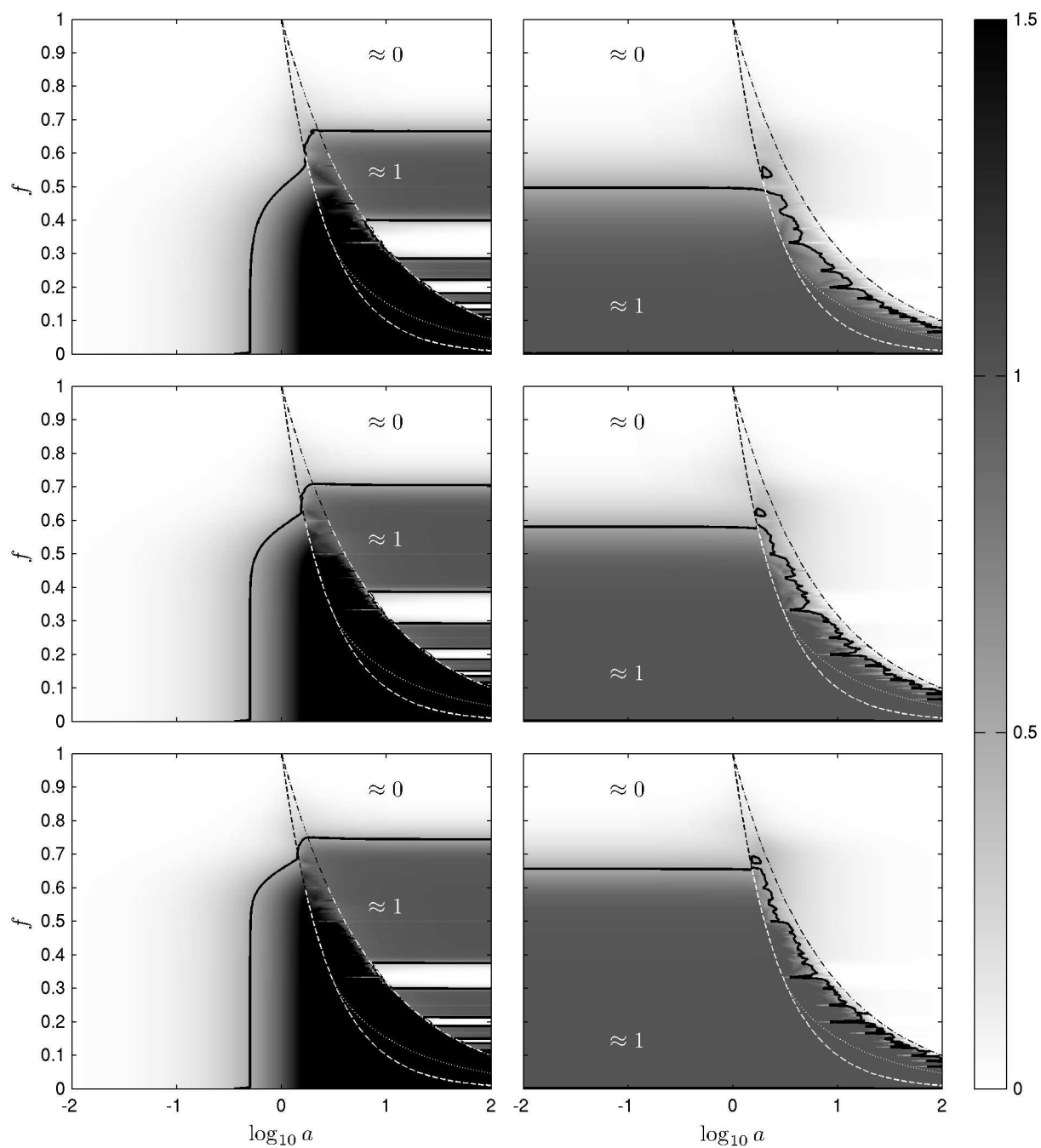


FIGURE 2.12 – Colonne de gauche : $c_{2,1}$ moyenné par rapport à φ pour (de haut en bas) 1, 3 et 10 itérations de tamisage. Colonne de droite : idem pour $c_{2,2}$. Les lignes épaisses correspondent aux courbes de niveau = 0.5.

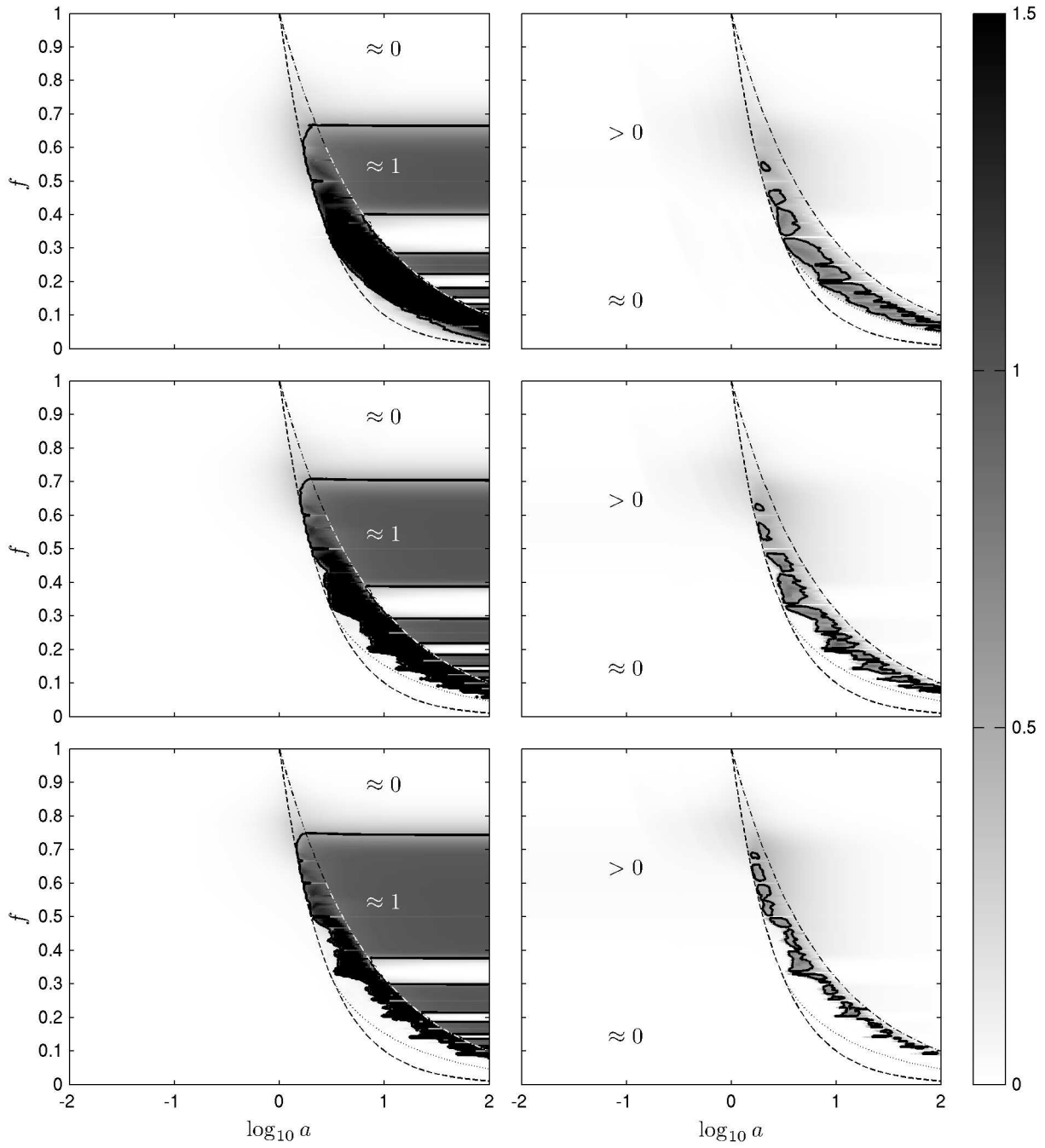


FIGURE 2.13 – Colonne de gauche : $c_{3,1}$ moyenné par rapport à φ pour (de haut en bas) 1, 3 et 10 itérations de tamisage. Colonne de droite : idem pour $c_{3,2}$. $c_{3,2}$ s'écarte très légèrement de 0 (jusque 0.03 pour 10 itérations) pour f aux alentours de 0.6, et ce même quand $a \rightarrow 0$. Les lignes épaisses correspondent aux courbes de niveau = 0.5.

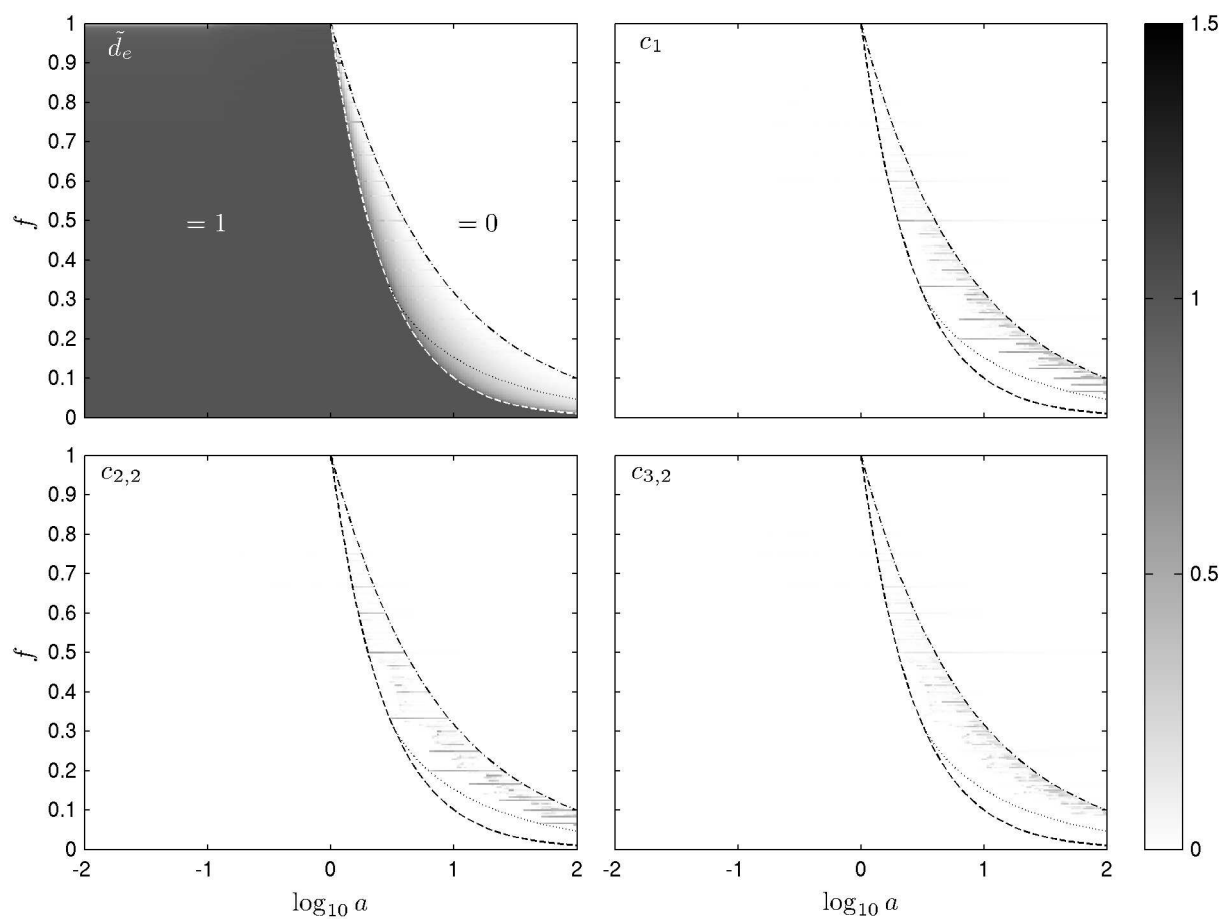


FIGURE 2.14 – En haut à gauche : densité d'extrema réduite (2.12) moyennée par rapport à φ . Autres : écarts type des quantités c_1 , $c_{2,2}$ et $c_{3,2}$ par rapport à φ pour 10 itérations de tamisage (rares sont les valeurs dépassant 0.5).

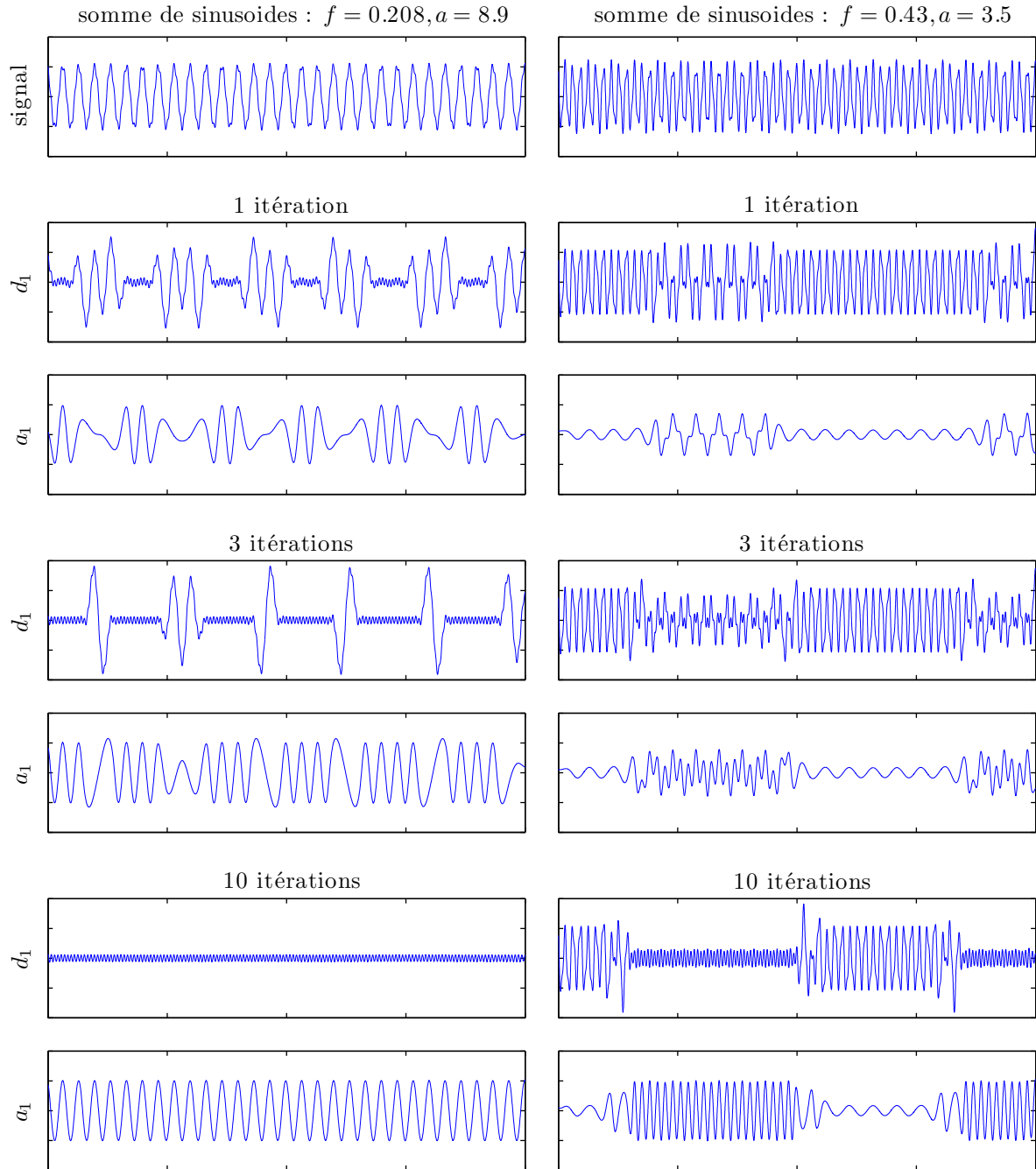


FIGURE 2.15 – Exemples de comportements dans la zone de transition. Initialement, seule une partie des extrema de la composante HF donne lieu à des extrema dans le signal composite. Au fil des itérations de tamisage, il en apparaît de plus en plus, mais le nombre d'itérations nécessaire à tous les faire ressortir peut être très grand. A priori, celui-ci diverge quand $f \rightarrow 1/k, k \in \mathbb{N}^*$.

vaut en moyenne par rapport à $\varphi \in [0, 2\pi]$

$$\langle d_e(a, f, \varphi) \rangle_\varphi = \begin{cases} 2 & \text{si } af < 1 \\ 2f + \frac{4}{\pi} \left[\sin^{-1} \left(\frac{1}{af} \sqrt{\frac{1 - a^2 f^4}{1 - f^2}} \right) - f \sin^{-1} \sqrt{\frac{1 - a^2 f^4}{1 - f^2}} \right] & \text{si } af^2 < 1 < af \\ 2f & \text{si } af^2 > 1 \end{cases} \quad (2.14)$$

Démonstration. Commençons par prouver les 2 résultats simples :

$$\langle d_e(a, f, \varphi) \rangle_\varphi = \begin{cases} 2 & \text{si } af < 1 \\ 2f & \text{si } af^2 > 1 \end{cases} \quad (2.15)$$

Pour ce faire, on montre d'abord que le signe de la dérivée seconde du signal somme (2.4) au niveau de ses extrema est le même que celui de la dérivée seconde de la composante HF si $af < 1$ et BF si $af^2 > 1$. Supposons que $x(t; a, f, \varphi)$ a un extremum en $t = t_0$:

$$\partial_t x(t; a, f, \varphi)|_{t=t_0} \propto \sin 2\pi t_0 + af \sin(2\pi f t_0 + \varphi) = 0. \quad (2.16)$$

La dérivée seconde de $x(t; a, f, \varphi)$ vaut

$$\partial_t^2 x(t; a, f, \varphi) \propto \cos 2\pi t + af^2 \cos(2\pi f t + \varphi), \quad (2.17)$$

et on veut montrer que

$$|af^2 \cos(2\pi f t_0 + \varphi)| < |\cos 2\pi t_0| \text{ si } af < 1, \quad (2.18)$$

$$|af^2 \cos(2\pi f t_0 + \varphi)| > |\cos 2\pi t_0| \text{ si } af^2 > 1. \quad (2.19)$$

En mettant au carré et en utilisant (2.16), on obtient les équations équivalentes

$$a^2 f^4 \cos^2(2\pi f t_0 + \varphi) + a^2 f^2 \sin^2(2\pi f t_0 + \varphi) < 1 \text{ si } af < 1, \quad (2.20)$$

$$a^2 f^4 \cos^2(2\pi f t_0 + \varphi) + a^2 f^2 \sin^2(2\pi f t_0 + \varphi) > 1 \text{ si } af^2 > 1, \quad (2.21)$$

qui sont vraies puisque $a^2 f^4 < a^2 f^2 < 1$ si $af < 1$ et $1 < a^2 f^4 < a^2 f^2$ si $af^2 > 1$.

Étant donné ce premier résultat, on en déduit qu'il ne peut y avoir qu'un extremum dans le signal somme entre deux passages à zéro de la composante HF si $af < 1$ (ou BF si $af^2 > 1$) parce que son type (maximum ou minimum) est en fait imposé par le signe de la dérivée seconde de HF (ou BF si $af^2 > 1$). Il ne peut donc y avoir qu'un seul extremum dans le signal somme par demi-période de HF si $af < 1$ (ou BF si $af^2 > 1$), d'où les densités d'extrema (indépendantes de φ) dans ces deux cas.

Montrons maintenant que si $af^2 < 1 < af$, la densité d'extrema moyenne est donnée par

$$\langle d_e(a, f, \varphi) \rangle_\varphi = 2f + \frac{4}{\pi} \left[\sin^{-1} \left(\frac{1}{af} \sqrt{\frac{1 - a^2 f^4}{1 - f^2}} \right) - f \sin^{-1} \sqrt{\frac{1 - a^2 f^4}{1 - f^2}} \right] \quad (2.22)$$

Si $x(t; a, f, \varphi)$ est extrémal en t_e , la dérivée du signal doit être nulle en ce point

$$\partial_t x(t; a, f, \varphi)|_{t=t_e} \propto \sin 2\pi t_e + af \sin(2\pi f t_e + \varphi) = 0. \quad (2.23)$$

Comme $af > 1$, on en déduit que nécessairement

$$\sin(2\pi f t_e + \varphi) \in \left[-\frac{1}{af}, \frac{1}{af} \right], \quad (2.24)$$

ce qui implique que

$$2\pi f t_e + \varphi \in \left[-\sin^{-1}\left(\frac{1}{af}\right) + k\pi, \sin^{-1}\left(\frac{1}{af}\right) + k\pi \right], \quad k \in \mathbb{Z} \quad (2.25)$$

et donc

$$t_e \in \left[-\frac{1}{2\pi f} \left(\sin^{-1}\left(\frac{1}{af}\right) + k\pi - \varphi \right), \frac{1}{2\pi f} \left(\sin^{-1}\left(\frac{1}{af}\right) + k\pi - \varphi \right) \right], \quad k \in \mathbb{Z}. \quad (2.26)$$

On vient de montrer que les extrema du signal somme sont nécessairement localisés dans des intervalles situés au voisinage des extrema de la composante BF (en $(k\pi - \varphi)/(2\pi f)$, $k \in \mathbb{Z}$). De là, si l'on calcule le nombre moyen (par rapport à φ) d'extrema dans l'un de ces intervalles, on voit que ce nombre est également le nombre moyen d'extrema par demi-période (entre deux passages à zéro) de la composante BF. La densité moyenne d'extrema dans le signal somme sera alors simplement ce nombre divisé par $1/(2f)$, longueur d'une demi-période de la composante BF. Intéressons-nous donc à l'un de ces intervalles, par exemple

$$I = \left[-\frac{1}{2\pi f} \sin^{-1}\left(\frac{1}{af}\right) - \varphi, \frac{1}{2\pi f} \sin^{-1}\left(\frac{1}{af}\right) - \varphi \right] \quad (2.27)$$

et faisons un changement de variables de manière à le rendre fixe par rapport à φ : on pose $t' = t + \frac{\varphi}{2\pi f}$, ce qui correspond à

$$x(t'; a, f, \varphi) = \cos(2\pi t' - \varphi) + \cos(2\pi f t'). \quad (2.28)$$

L'intervalle I précédent devient alors

$$I = \left[-\frac{1}{2\pi f} \sin^{-1}\left(\frac{1}{af}\right), \frac{1}{2\pi f} \sin^{-1}\left(\frac{1}{af}\right) \right] \quad (2.29)$$

Dans la suite, on omettra les « primes » pour alléger les notations.

On introduit le point $t_0 \in I \cap \mathbb{R}_+$ tel que $x(t; a, f, \varphi)$ puisse présenter une inflexion horizontale (point où la dérivée première s'annule sans changer de signe) en t_0 , ce qui correspond au cas où les deux courbes de la figure Fig. 2.16 sont tangentes. t_0 vérifie

$$\exists \varphi_0 \text{ tel que, } \begin{cases} \sin(2\pi t_0 - \varphi_0) + af \sin(2\pi f t_0) = 0 \\ \cos(2\pi t_0 - \varphi_0) + af^2 \cos(2\pi f t_0) = 0. \end{cases} \quad (2.30)$$

Si l'on suppose l'existence d'un tel couple (t_0, φ_0) , il doit vérifier

$$\begin{cases} \sin^2(2\pi t_0 - \varphi_0) = a^2 f^2 \sin^2(2\pi f t_0) \\ \cos^2(2\pi t_0 - \varphi_0) = a^2 f^4 \cos^2(2\pi f t_0), \end{cases} \quad (2.31)$$

ce qui implique que

$$\sin^2(2\pi f t_0) = \frac{1 - a^2 f^4}{a^2 f^2 - a^2 f^4}, \quad (2.32)$$

qui est possible car $af^2 < 1 < af$ et $f < 1$. Par conséquent,

$$t_0 = \frac{1}{2\pi f} \sin^{-1}\left(\frac{1}{af} \sqrt{\frac{1 - a^2 f^4}{1 - f^2}}\right), \quad (2.33)$$

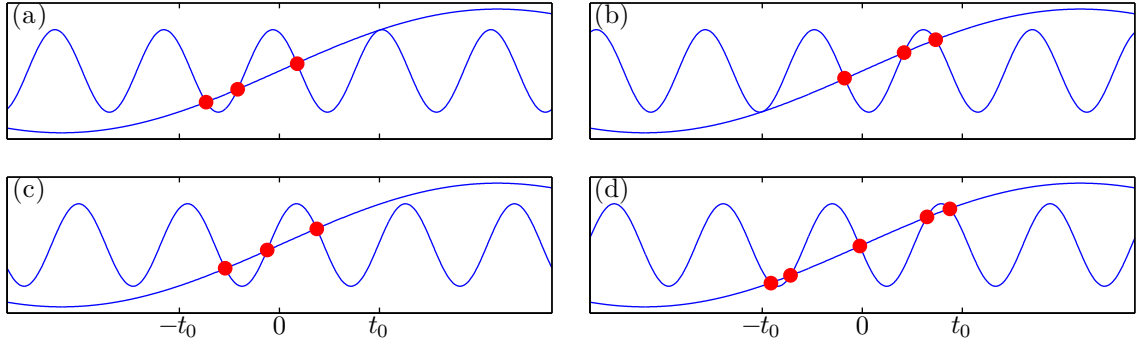


FIGURE 2.16 – Dans chaque cadre sont représentées la dérivée de la composante BF et l'opposée de la dérivée de la composante HF pour différentes valeurs de φ . Chaque croisement (points rouges) correspond à un extremum dans le signal. Là où les courbes sont tangentes, le signal présente une inflexion horizontale. (a) $\varphi = \varphi_0$. (b) $\varphi = \varphi_1$. (c) $\varphi \in]\varphi_0, \varphi_1[$. (d) $\varphi \in]\varphi_1, \varphi_0 + 2\pi[$.

puisque $t_0 \in I \subset \left[-\frac{1}{4f}, \frac{1}{4f}\right]$. Au passage, on a montré, sous réserve d'existence, que t_0 est le seul point de $I \cap \mathbb{R}_+$, où il peut y avoir une inflexion horizontale. Par ailleurs, φ_0 vérifie

$$\sin(2\pi t_0 - \varphi_0) = -af \sin(2\pi f t_0) = -\underbrace{\frac{1 - a^2 f^4}{1 - f^2}}_A. \quad (2.34)$$

où $A < 1$ car $af > 1$. On en déduit que l'une des deux valeurs de φ_0 suivantes convient :

$$\varphi_0 = 2\pi t_0 - \sin^{-1} A - \pi \quad \text{ou} \quad \varphi_0 = \sin^{-1} A - 2\pi t_0. \quad (2.35)$$

La deuxième équation de (2.30) permet alors de choisir

$$\varphi_0 = 2\pi t_0 - \sin^{-1} A - \pi. \quad (2.36)$$

Si l'on remplace t_0 et φ_0 par ces valeurs dans (2.30), on voit que le point t_0 correspond bien à une inflexion horizontale de $x(t; a, f, \varphi_0)$.

Dans la suite, on prendra φ dans $[\varphi_0, \varphi_0 + 2\pi]$ au lieu de $[0, 2\pi]$. Soit K le nombre d'extrema de $x(t; a, f, \varphi)$ dans l'intervalle ouvert $] -t_0, t_0[$ pour $\varphi = \varphi_0$ (l'extremum double n'est pas compté) et soit $\varphi_1 \in [\varphi_0, \varphi_0 + 2\pi[$ tel que $\varphi_1 \equiv -\varphi_0 [2\pi]$. Pour $\varphi = \varphi_1$, $x(t; a, f, \varphi)$ présente un extremum double en $-t_0$. Le calcul fournit

$$\varphi_1 = \pi + \sin^{-1} A - 2\pi t_0 + 2\pi \left\lceil \frac{2\pi t_0 - \sin^{-1} A - \pi}{\pi} \right\rceil \quad (2.37)$$

Si $\varphi_1 \equiv \varphi_0 \equiv 0 [2\pi]$, il est aisé de voir que pour toute valeur de $\varphi \in [\varphi_0, \varphi_0 + 2\pi]$, il y a exactement $K + 2$ extrema dans I . Si au contraire $\varphi_0 \not\equiv 0 [2\pi]$, le nombre d'extrema dépend de φ . En s'aidant de la Fig. 2.16, on peut voir que si $\varphi \in]\varphi_0, \varphi_1[$, il y a exactement K extrema et $K + 2$ si $\varphi \in]\varphi_1, \varphi_0 + 2\pi[$. Par conséquent, si φ est uniformément distribué dans $[\varphi_0, \varphi_0 + 2\pi]$, le nombre moyen d'extrema dans I est

$$K + 2 - \frac{\varphi_1 - \varphi_0}{\pi}. \quad (2.38)$$

Cette formule est également valable pour le cas $\varphi_1 \equiv \varphi_0 \equiv 0 [2\pi]$.

Il reste à calculer K . Pour ce faire, introduisons δ , l'écart entre t_0 et le maximum de $-\sin(2\pi t - \varphi_0)$ situé juste à côté. Étant donné que $\sin(2\pi t_0 - \varphi_0) = A$ et que $\sin(2\pi(t_0 + \delta) - \varphi_0) = 1$, on a

$$\delta = \frac{1}{4} - \frac{\sin^{-1} A}{2\pi}. \quad (2.39)$$

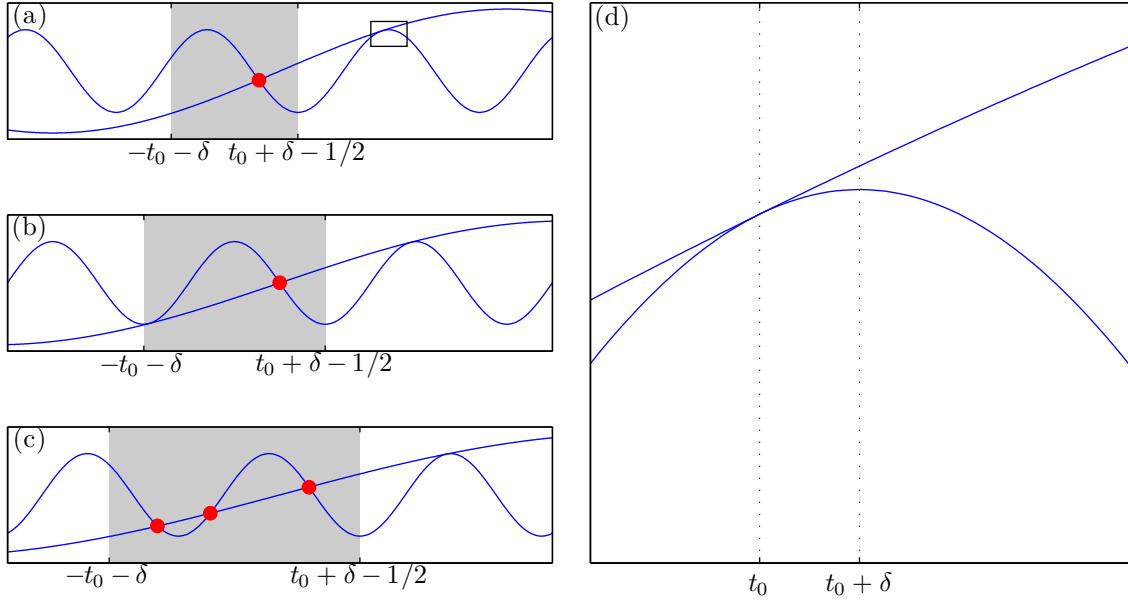


FIGURE 2.17 – Dans chaque cadre sont représentés la dérivée de la composante BF et l’opposé de la dérivée de la composante HF pour différentes valeurs de φ . Chaque croisement (points rouges) correspond à un extremum dans le signal. (a),(b),(c) nombre d’extrema dans $] -t_0, t_0[$ pour $\varphi = \varphi_0$. (d) définition de t_0 et δ (agrandissement de la figure (a)).

En s’aidant de la figure Fig. 2.17, on peut voir que le nombre K d’extrema dans I pour $\varphi = \varphi_0$ dépend du nombre K' de périodes entières de $\sin(2\pi t - \varphi_0)$ dans l’intervalle $] -t_0 - \delta, t_0 + \delta - 1/2[$, c’est-à-dire

$$K' = \lceil 2t_0 + 2\delta - 1/2 \rceil - 1 = \left\lceil 2t_0 - \frac{\sin^{-1} A}{\pi} \right\rceil - 1. \quad (2.40)$$

Dans le cas de Fig. 2.17 (a), il y a une période incomplète et un extremum (croisement). Dans le cas limite de Fig. 2.17 (b) où $\varphi_0 \equiv 0 [2\pi]$, il y a une période quasi-complète dans $] -t_0 - \delta, t_0 + \delta - 1/2[$ et toujours 1 seul extremum dans $] -t_0, t_0[$. Enfin, dans le cas de Fig. 2.17 (c), il y a 1 période complète et 3 extrema. Plus généralement, il en découle que K' périodes complètes donnent lieu à $K = 2K' + 1$ extrema dans $] -t_0, t_0[$. Finalement, si l’on reprend l’équation (2.38), on obtient que le nombre moyen d’extrema dans I vaut

$$K + 2 - \frac{\varphi_1 - \varphi_0}{\pi} = 1 + 4t_0 - \frac{2 \sin^{-1} A}{\pi}. \quad (2.41)$$

Ce nombre étant également le nombre moyen d’extrema dans une demi-période de la composante BF, si on remplace t_0 et A par leurs valeurs et qu’on divise l’expression par $1/(2f)$, on obtient finalement le résultat (2.22). $\square$

Au-delà de ces résultats sur la densité d’extrema, on sait également que lorsque l’une des deux composantes a une amplitude infiniment plus grande que l’autre ($a \rightarrow 0$ ou $a \rightarrow \infty$), les extrema du signal somme sont situés aux mêmes endroits que ceux de la composante dominante. Pour des valeurs finies du rapport d’amplitude, on peut montrer que, dès lors que la densité d’extrema est la même que celle de l’une des deux composantes ($af < 1$ ou $af^2 > 1$), les extrema du signal somme sont situés dans un certain rayon autour de ceux de la composante dominante. La proposition 3 apporte un contrôle quantitatif sur la proximité entre les extrema de la composante dominante et ceux du signal somme supportant un peu plus l’idée selon laquelle les extrema du signal somme sont d’autant plus proches de ceux de la composante dominante que af est petit devant 1 ou af^2 est grand devant 1.

Proposition 3. Si $af < 1$, les extrema de $x(t; a, f, \varphi)$ sont dans un rayon de $\frac{1}{2\pi} \sin^{-1}(af)$ autour des extrema de $\cos(2\pi t)$, c'est-à-dire des points de l'ensemble $\frac{1}{2}\mathbb{Z}$.

Si $af^2 > 1$, les extrema de $x(t; a, f, \varphi)$ sont dans un rayon de $\frac{1}{2\pi f} \sin^{-1}\left(\frac{1}{af}\right)$ autour des extrema de $\cos(2\pi ft + \varphi)$, c'est-à-dire des points de l'ensemble $-\frac{\varphi}{2\pi f} + \frac{1}{2f}\mathbb{Z}$.

Démonstration.

* Cas $af < 1$:

Considérons un extremum de $x(t; a, f, \varphi)$ situé en $t = t_0$. La dérivée du signal est nulle en t_0

$$\partial_t x(t; a, f, \varphi)|_{t=t_0} \propto \sin 2\pi t_0 + af \sin(2\pi f t_0 + \varphi) = 0. \quad (2.42)$$

Ceci implique que

$$|\sin 2\pi t_0| \leq af, \quad (2.43)$$

et donc

$$t_0 \in \left[\frac{k}{2} - \frac{1}{2\pi} \sin^{-1} af, \frac{k}{2} + \frac{1}{2\pi} \sin^{-1} af \right], \quad k \in \mathbb{Z}. \quad (2.44)$$

De plus, il y a nécessairement un extremum de $x(t; a, f, \varphi)$ dans chacun de ces intervalles puisque $\partial_t x(t; a, f, \varphi)$ est continue et que $\forall k \in \mathbb{Z}$,

$$\partial_t x\left(\frac{k}{2} - \frac{1}{4}; a, f, \varphi\right) = 2\pi \left(-1 + af \sin \left(2\pi f \left(\frac{k}{2} - \frac{1}{4} \right) + \varphi \right) \right) < 0, \quad (2.45)$$

et

$$\partial_t x\left(\frac{k}{2} + \frac{1}{4}; a, f, \varphi\right) = 2\pi \left(1 + af \sin \left(2\pi f \left(\frac{k}{2} + \frac{1}{4} \right) + \varphi \right) \right) > 0. \quad (2.46)$$

On pourrait de plus montrer qu'il est unique en reprenant l'argument utilisé dans la démonstration de la proposition 2 selon lequel le type de l'extremum est imposé par le signe de la dérivée seconde de la composante HF sur l'intervalle considéré.

* Cas $af^2 > 1$:

Similairement, si la dérivée de $x(t; a, f, \varphi)$ s'annule en t_0 , on doit avoir

$$|af \sin(2\pi f t_0 + \varphi)| \leq 1, \quad (2.47)$$

et donc

$$t_0 \in \left[\frac{k}{2f} - \frac{\varphi}{2\pi f} - \frac{1}{2\pi f} \sin^{-1} \left(\frac{1}{af} \right), \frac{k}{2f} - \frac{\varphi}{2\pi f} + \frac{1}{2\pi f} \sin^{-1} \left(\frac{1}{af} \right) \right], \quad k \in \mathbb{Z}. \quad (2.48)$$

De plus, tous ces intervalles contiennent un extremum de $x(t; a, f, \varphi)$, puisque $\forall k \in \mathbb{Z}$,

$$\partial_t x\left(\frac{k}{2f} - \frac{\varphi}{2\pi f}; a, f, \varphi\right) = 2\pi \left(\sin \left(2\pi t_k - \frac{1}{4f} \right) - af \right) < 0, \quad (2.49)$$

et

$$\partial_t x\left(\frac{k}{2f} - \frac{\varphi}{2\pi f} + \frac{1}{2f}; a, f, \varphi\right) = 2\pi \left(\sin \left(2\pi t_k + \frac{1}{4f} \right) + af \right) > 0, \quad (2.50)$$

où $t_k = \frac{k}{2f} - \frac{\varphi}{2\pi f}$. Celui-ci est de plus unique pour la même raison que dans le cas $af < 1$. $\square$

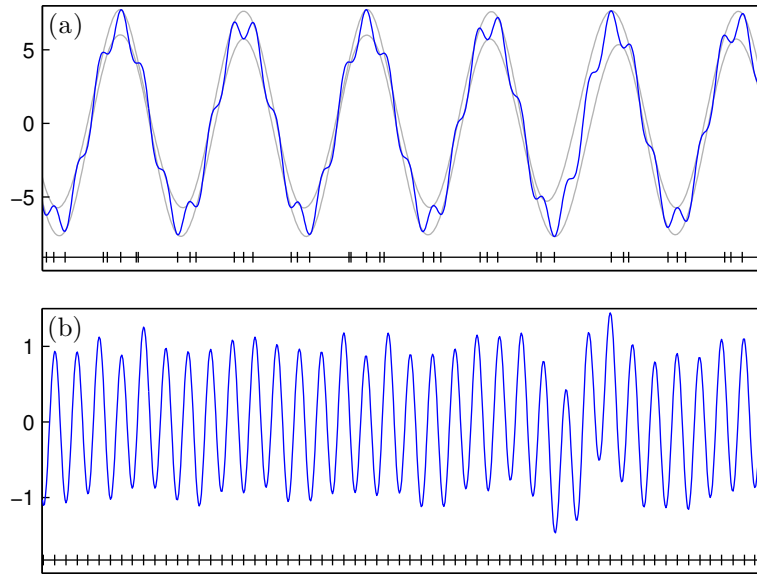


FIGURE 2.18 – Exemple de somme de sinusoïdes dans le cas où $af \gtrsim 1$ et $f < 1/3$. (a) le signal $x(t; a, f, \varphi)$. Ses enveloppes sont tracées en gris clair. (b) le signal après une itération de tamisage. Dans les deux cas, les positions des extrema sont identifiées par des marques sur la ligne horizontale en dessous de chaque signal.

Pourquoi la courbe $af = 1$ délimite moins bien la zone de transition quand $f < 1/3$ L'explication de ce phénomène est simplement que dans certains cas, des extrema peuvent apparaître après quelques itérations de tamisage. Dans le cas des sommes de sinusoïdes, ceci se produit essentiellement lors de la première itération. En pratique, on observe que pour $af \gtrsim 1$ et $f < 1/3$, l'EMD semble se comporter comme si la densité d'extrema était identique à celle de la composante HF alors qu'elle est légèrement inférieure. Si l'on regarde l'allure des signaux dans cette zone du plan (a, f) (cf Fig. 2.18), on s'aperçoit que les extrema sont essentiellement localisés au voisinage des extrema de la composante BF et qu'ailleurs, il y a des épaulements qui clairement donneraient lieu à des extrema si le rapport d'amplitude était plus faible. Étant donné que le signal présente des minima à proximité des maxima de la composante BF (et inversement), on devine que la moyenne des enveloppes suit grossièrement la composante BF. De ce fait, en la soustrayant lors de la première itération de tamisage, on voit bien que les extrema cachés sous la forme d'épaulements vont devenir de véritables extrema et que par conséquent la densité d'extrema après une itération augmente pour égaler celle de la composante HF.

Pour trouver une équation de la véritable frontière observée pour $f < 1/3$, il suffit en fait de trouver pour quelles valeurs de (a, f) il y a systématiquement au moins un minimum du signal somme entre deux passages à zéro de la composante BF encadrant un maximum de cette dernière. Plus généralement, il suffit de trouver dans quelles conditions il y a systématiquement au moins deux extrema du signal somme entre deux passages à zéro de la composante BF, l'un des deux étant nécessairement du type voulu. Il se trouve que si $f < 1/3$, ceci est possible sans que pour autant la densité d'extrema soit la même que celle de la composante HF. Si l'on considère le cas limite de Fig. 2.19 (a) où le rapport d'amplitude a la valeur limite $a = \tilde{a}(f)$, il est clair que quelle que soit la valeur de la phase φ , si $a < \tilde{a}(f)$, il y a systématiquement au moins trois extrema entre deux passages à zéro de la composante BF (qui sont des extrema dans la figure puisqu'on représente la dérivée des composantes). Inversement, il peut n'y en avoir qu'un seul si $a > \tilde{a}(f)$ en prenant la phase telle que dans le cas de la figure. La figure Fig. 2.19 (b) représente le même cas limite avec $f \lesssim 1/3$ et permet de voir pourquoi le phénomène d'apparition d'extrema après la première itération

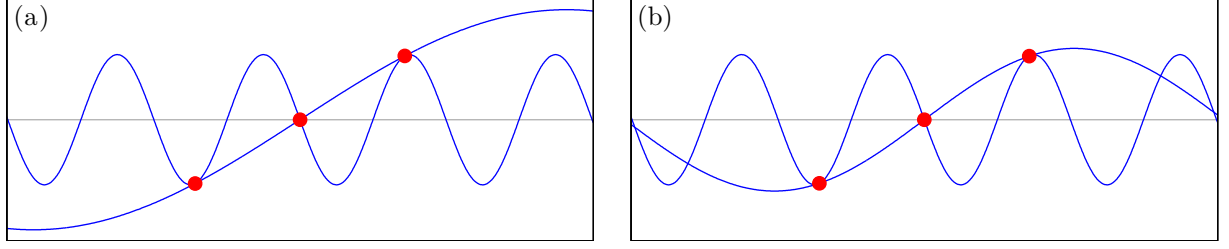


FIGURE 2.19 – Cas limite $a = \tilde{a}(f)$. Dans chaque graphe, on représente la dérivée de la composante HF et l'opposé de la dérivée de la composante BF de sorte que les extrema du signal $x(t; a, f, \varphi)$ apparaissent sous la forme de croisements entre les deux courbes, marqués par des gros points. Les cas (a) et (b) diffèrent par les valeurs de f , celle du cas (b) étant assez proche de $f = 1/3$ de manière à pouvoir visualiser le changement de comportement pour $f = 1/3$ où le cas limite correspondrait à $af = 1$ (même amplitude pour les deux dérivées).

de tamisage n'existe que dans le cas $f < 1/3$.

Pour trouver une expression de la frontière $\tilde{a}(f)$, l'idée est d'écrire les conditions de tangence de Fig. 2.19. Malheureusement, celles-ci ne permettent pas de donner une expression analytique de $\tilde{a}(f)$. On peut toutefois en trouver une bonne approximation si l'on suppose que les points de tangence coïncident avec les extrema de la dérivée de la composante HF situés à proximité, ce qui fournit :

$$\tilde{a}(f) = \frac{1}{f \sin\left(\frac{3\pi f}{2}\right)} \quad (2.51)$$

Enfin, on peut également remarquer que cette frontière à partir de laquelle il y a systématiquement au moins trois extrema du signal somme entre deux passages à zéro de la composante BF correspond exactement à la frontière sur laquelle la densité d'extrema vaut exactement $6f$ ce qui correspond à l'équation

$$\sin^{-1}\left(\frac{1}{af} \sqrt{\frac{1-a^2f^4}{1-f^2}}\right) - f \sin^{-1} \sqrt{\frac{1-a^2f^4}{1-f^2}} - \pi f = 0. \quad (2.52)$$

1.1.6.2 Modèle dans le cas d'extrema uniformément espacés Dans le cas où les extrema du signal sont uniformément espacés, maxima et minima indépendamment, il est possible de décrire le comportement de l'EMD par un modèle analytique. Si ce dernier n'est exact que pour les cas asymptotiques $a \rightarrow 0$ ou $a \rightarrow \infty$, seuls cas où les extrema sont effectivement uniformément espacés, on verra que le comportement de l'EMD est en fait très proche de ce modèle dès lors que la densité d'extrema est identique à celle de l'une des deux composantes ($af < 1$ ou $af^2 > 1$).

Présentation du modèle Soit un signal $x(t)$ défini sur $\mathbb{R}$ tel que pour tout $k \in \mathbb{Z}$, $x(t)$ admette un maximum local en $t = k$ et un minimum local en $t = k + \alpha$, avec $\alpha \in [0, 1]$.

Les extrema d'un tel signal étant uniformément espacés, les ensembles des maxima et des minima de ce signal peuvent être vus comme des échantillonnages particuliers qu'on peut caractériser par les peignes de Dirac dans le domaine de Fourier :

$$\text{III}_{\max}(\nu) = \sum_{k \in \mathbb{Z}} \delta(\nu - k), \quad (2.53)$$

$$\text{III}_{\min}(\nu) = \sum_{k \in \mathbb{Z}} e^{-2ik\pi\alpha} \delta(\nu - k), \quad (2.54)$$

où le facteur $e^{-2ik\pi\alpha}$ provient du décalage entre les deux échantillonnages.

En faisant la convolution de ces peignes de Dirac avec le signal dans le domaine de Fourier, on obtient une représentation de Fourier des ensembles d'extrema :

$$\mathcal{E}_{max} = \text{III}_{max} * \hat{x}(\nu), \quad (2.55)$$

$$\mathcal{E}_{min} = \text{III}_{min} * \hat{x}(\nu). \quad (2.56)$$

Remarquons au passage que comme tout échantillonnage, ces échantillonnages par les maxima et minima sont susceptibles d'introduire du repliement, ce qui a été initialement rapporté dans [31]. On verra par la suite que dans de nombreux cas, il y a effectivement des effets de repliement dans la décomposition proposée par l'EMD.

Si on utilise alors un schéma d'interpolation qui, dans le cas de nœuds uniformément espacés, peut être exprimé en termes de filtrage linéaire avec une réponse en fréquence $I(\nu)$ — ce qui est le cas notamment des interpolations splines de tout ordre [62] —, on peut alors écrire les représentations de Fourier des enveloppes :

$$\hat{e}_{max} = I(\nu) (\text{III}_{max} * \hat{x})(\nu), \quad (2.57)$$

$$\hat{e}_{min} = I(\nu) (\text{III}_{min} * \hat{x})(\nu), \quad (2.58)$$

où $I(\nu)$ dans le cas de l'interpolation spline cubique qui nous intéresse plus particulièrement est donné par [62]

$$I^{c.s.}(\nu) = \left(\frac{\sin \pi\nu}{\pi\nu} \right)^4 \frac{3}{2 + \cos 2\pi\nu}. \quad (2.59)$$

La moyenne des enveloppes peut alors être exprimée comme

$$\frac{\hat{e}_{min}(\nu) + \hat{e}_{max}(\nu)}{2} = I(\nu) (\text{III}_{mean}^\alpha * \hat{x})(\nu), \quad (2.60)$$

avec III_{mean}^α la demi-somme des deux échantillonnages

$$\text{III}_{mean}^\alpha(\nu) = \sum_{k \in \mathbb{Z}} e^{-ik\pi\alpha} \cos(k\pi\alpha) \delta(\nu - k). \quad (2.61)$$

Un cas intéressant est celui où maxima et minima sont intercalés symétriquement pour lequel III_{mean}^α devient

$$\text{III}_{mean}^{1/2} = \sum_{k \in \mathbb{Z}} \delta(\nu - 2k), \quad (2.62)$$

qui correspond donc à un échantillonnage à la fréquence 2. On voit que dans ce cas particulier, les deux échantillonnages par les maxima et les minima se combinent en un seul échantillonnage par tous les extrema. Une conséquence importante est que le repliement éventuellement introduit par chacun des échantillonnages est ici en partie compensé lorsqu'on fait la moyenne des deux enveloppes. En particulier, on peut considérer qu'il n'y a pas de repliement si le spectre de $x(t)$ ne contient pas de fréquences en dehors de l'intervalle $[-1, 1]$, même s'il peut y avoir du repliement dans chacune des enveloppes si le spectre de $x(t)$ contient des fréquences en dehors de $[-1/2, 1/2]$.

En soustrayant la moyenne des enveloppes au signal, on obtient finalement une représentation de Fourier de l'opérateur de tamisage :

$$(\widehat{\mathcal{S}x})(\nu) = \hat{x}(\nu) - I(\nu) (\text{III}_{mean}^\alpha * \hat{x})(\nu). \quad (2.63)$$

Une particularité importante de cette représentation est qu'elle est *linéaire* si on considère III_{mean}^α fixé a priori. Ceci est en fait dû au fait que la seule étape non linéaire dans l'opérateur de tamisage est la sélection des extrema du signal comme points caractéristiques (cf 3.1).

Formulation dans le cas général Si les maxima locaux de $x(t)$ sont situés en $t = k\lambda + \alpha_1$, $k \in \mathbb{Z}$ et les minima locaux en $t = k\lambda + \alpha_2$, $k \in \mathbb{Z}$, les peignes de Dirac correspondant deviennent

$$\text{III}_{\max}(\nu) = \frac{1}{\lambda} \sum_{k \in \mathbb{Z}} e^{-\frac{2ik\pi\alpha_1}{\lambda}} \delta\left(\nu - \frac{k}{\lambda}\right), \quad (2.64)$$

$$\text{III}_{\min}(\nu) = \frac{1}{\lambda} \sum_{k \in \mathbb{Z}} e^{-\frac{2ik\pi\alpha_2}{\lambda}} \delta\left(\nu - \frac{k}{\lambda}\right). \quad (2.65)$$

La réponse impulsionnelle du filtre d'interpolation est dilatée d'un facteur λ , ce qui donne la réponse en fréquence $\lambda I(\lambda\nu)$. Finalement, le modèle (2.63) devient dans le cas général

$$(\widehat{\mathcal{S}x})(\nu) = \hat{x}(\nu) - I(\lambda\nu) \left(\text{III}_{\text{mean}}^{\lambda, \alpha_1, \alpha_2} * \hat{x} \right)(\nu), \quad (2.66)$$

avec

$$\text{III}_{\text{mean}}^{\lambda, \alpha_1, \alpha_2} = \sum_{k \in \mathbb{Z}} e^{-\frac{ik\pi(\alpha_1 + \alpha_2)}{\lambda}} \cos\left(\frac{k\pi(\alpha_1 - \alpha_2)}{\lambda}\right) \delta\left(\nu - \frac{k}{\lambda}\right) \quad (2.67)$$

1.1.6.3 Application au cas de la somme de deux sinusoïdes

Cas $a \rightarrow 0$ et plus généralement $af < 1$ Dans le cas du signal $x(t; a, f, \varphi)$ pour $a \rightarrow 0$, le modèle précédent prend la forme de (2.63) avec $\alpha = 1/2$, c'est-à-dire

$$\text{III}_{\text{mean}}^{1/2} = \sum_{k \in \mathbb{Z}} \delta(\nu - 2k). \quad (2.68)$$

Initialement, le signal $x(t; a, f, \varphi)$ contient 4 composantes de Fourier aux fréquences $-1, 1, -f, f$ avec les coefficients $c_{\pm 1} = 1/2$ et $c_{\pm f} = ae^{\pm i\varphi}$. Les signaux étant toujours réels, on pourra se contenter dans la suite de ne considérer que les composantes aux fréquences positives. De plus, l'opérateur décrit par (2.63) étant linéaire si on fixe l'opérateur d'échantillonnage (2.68), on pourra l'appliquer indépendamment sur chacune des composantes de Fourier. Commençons par nous intéresser à $\mathcal{S}e_1$, où de manière générale on note $e_\nu(t) = e^{2i\pi\nu t}$ la fonction de la base de Fourier à la fréquence ν :

$$(\widehat{\mathcal{S}e_1})(\nu) = \hat{e}_1(\nu) - I(\nu) \left(\text{III}_{\text{mean}}^{1/2} * \hat{e}_1 \right)(\nu), \quad (2.69)$$

$$= \delta(\nu - 1) - \sum_{k \in \mathbb{Z}} I(2k + 1) \delta(\nu - 2k - 1) \quad (2.70)$$

Or dès lors que $I(\nu)$ correspond à un filtre d'interpolation qui reproduit l'unité — ce qui est la moindre des exigences pour une interpolation — on sait que $I(\nu)$ doit s'annuler pour $\nu \in \mathbb{Z}^*$ et valoir 1 en 0^2 . Par conséquent l'expression de $\mathcal{S}e_1$ se simplifie en

$$(\widehat{\mathcal{S}e_1})(\nu) = \delta(\nu - 1) = \hat{e}_1(\nu), \quad (2.71)$$

c'est-à-dire que la composante HF est laissée inchangée par l'opérateur de tamisage.

Pour la composante BF, on s'intéresse à $\mathcal{S}e_f$:

$$(\widehat{\mathcal{S}e_f})(\nu) = \delta(\nu - f) - \sum_{k \in \mathbb{Z}} I(2k + f) \delta(\nu - 2k - f). \quad (2.72)$$

2. Ce résultat se démontre aisément à l'aide de la formule de sommation de Poisson : si $i(t)$ est la réponse impulsionnelle de l'interpolation (et $I(\nu)$ sa réponse en fréquence), elle reproduit l'unité ssi $\forall x, \sum_{k \in \mathbb{Z}} i(x + k) = 1$ ce qui est équivalent à $\forall x, \sum_{k \in \mathbb{Z}} I(k) e^{2ik\pi x} = 1$ qui implique $\forall k \in \mathbb{Z}^*, I(k) = 0$ et $I(0) = 1$.

Contrairement au cas précédent, on ne peut pas simplifier plus l'expression. Toutefois, si on peut supposer que $I(\nu)$ est négligeable pour $\nu \notin [-1, 1]$, on pourra arriver au même niveau de simplification que pour la composante HF. Cette approximation est en particulier très raisonnable dans le cas de l'interpolation spline cubique dans la mesure où $I^{c.s.}(\nu)$ décroît rapidement (en $\mathcal{O}(\nu^{-4})$) en dehors de $[-1, 1]$ son maximum étant de 0.0064 pour $\nu \simeq 1.46$. Avec cette approximation, on obtient

$$(\widehat{\mathcal{S}e_f})(\nu) \simeq (1 - I(f))\hat{e}_f, \quad (2.73)$$

qui correspond à un filtrage passe-haut de la composante BF par le filtre $1 - I(\nu)$.

Si on combine maintenant ces résultats, on obtient que

$$d_1^{(1)}(t; a, f, \varphi) = (\mathcal{S}x(\cdot; a, f, \varphi))(t) \simeq x(t; a(1 - I(f)), f, \varphi), \quad (2.74)$$

c'est-à-dire que la somme de deux sinusôides après une itération de tamisage est très proche d'une somme de sinusôides avec un rapport d'amplitude inférieur, $I(\nu)$ étant généralement à valeur dans $[0, 1]$. Étant donné que le modèle utilisé est a priori valide seulement dans le cas où $a \rightarrow 0$, il est donc d'autant plus valide après une itération de tamisage, ce qui nous permet finalement d'exprimer les IMFs après un nombre n d'itérations de tamisage :

$$d_1^{(n)}(t; a, f, \varphi) = \cos 2\pi t + (1 - I(f))^n a \cos(2\pi f t + \varphi), \quad (2.75)$$

$$d_2^{(n)}(t; a, f, \varphi) = a_1^{(n)}(t; a, f, \varphi) = (1 - (1 - I(f))^n) a \cos(2\pi f t + \varphi), \quad (2.76)$$

où on s'est permis d'exprimer le deuxième IMF $d_2^{(n)}(t; a, f, \varphi)$ dans la mesure où $a_1^{(n)}(t; a, f, \varphi)$ est clairement déjà un IMF.

On voit ainsi que dans le cas particulier d'une somme de sinusôides, l'EMD se comporte comme un filtre linéaire de réponse en fréquence $(1 - I(\nu/f_1))^n$ pour n itérations de tamisage. Sachant que $I(\nu)$ est un filtre passe-bas, on déduit que $(1 - I(\nu/f_1))^n$ est la réponse en fréquence d'un filtre pass-haut dont la fréquence de coupure croît avec n . Si on veut définir plus concrètement cette fréquence de coupure, il semble naturel de choisir la fréquence à laquelle la réponse du filtre vaut exactement $1/2$. On peut donc définir $f_c^{(n)}$, la fréquence de coupure du filtre équivalent à l'EMD lorsque la composante HF à une fréquence de 1 par :

$$(1 - I(f_c^{(n)}))^n = \frac{1}{2}, \quad (2.77)$$

c'est-à-dire

$$I(f_c^{(n)}) = 1 - \frac{1}{2^{\frac{1}{n}}}. \quad (2.78)$$

Malheureusement cette équation n'a pas de solution analytique dans le cas d'une interpolation spline cubique. Cependant, la formule asymptotique

$$f_c^{(n)} = \left(1 + \left(\frac{\ln 2}{k}\right)^{\frac{1}{4}}\right)^{-1} + \mathcal{O}\left(n^{-\frac{5}{4}}\right) \quad (2.79)$$

en fournit une excellente approximation³, valable même pour les petites valeurs de n (cf Fig. 2.20). Cette approximation s'obtient en écrivant en partie le développement limité de $I(\nu)$ au voisinage de 1 :

$$I(\nu) = \frac{(1 - \nu)^4}{\nu^4} + \mathcal{O}((\nu - 1)^6), \quad (2.80)$$

3. On pourrait encore simplifier le développement asymptotique mais ce serait au détriment de la qualité de l'approximation pour les petites valeurs de n .

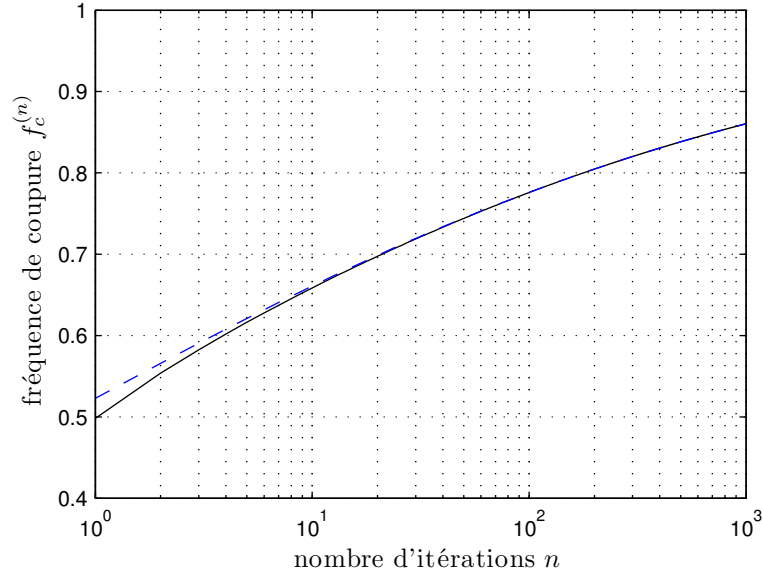


FIGURE 2.20 – Fréquence de coupure du filtre $(1 - I(\nu))^n$ en fonction du nombre d'itérations n . Trait plein noir : valeur exacte. Tirets : approximation (2.79).

qui implique pour $\nu < 1$ que

$$\nu = \frac{1}{1 + I(\nu)^{\frac{1}{4}}} + \mathcal{O}\left(I(\nu)^{\frac{3}{2}}\right). \quad (2.81)$$

Finalement, le résultat (2.79) est obtenu en insérant dans cette dernière formule le développement :

$$I(\nu) = 1 - \frac{1}{2^{\frac{1}{n}}} = \frac{\ln 2}{n} + \mathcal{O}\left(\frac{1}{n^2}\right). \quad (2.82)$$

Cas $a \rightarrow \infty$ et plus généralement $af^2 > 1$ Dans le cas où $a \rightarrow \infty$, les maxima du signal sont situés en $t = -\varphi/(2\pi f) + k/f$, $k \in \mathbb{Z}$ et les minima en $t = -\varphi/(2\pi f) + k/f + 1/(2f)$, $k \in \mathbb{Z}$. Le modèle (2.66) prend alors la forme

$$(\widehat{\mathcal{S}x})(\nu) = \hat{x}(\nu) - I(\nu/f) (\text{III}_{mean} * \hat{x})(\nu), \quad (2.83)$$

avec

$$\text{III}_{mean}(\nu) = \sum_{k \in \mathbb{Z}} \delta(\nu - 2kf) e^{2ik\varphi}. \quad (2.84)$$

Comme dans le cas précédent, on peut se permettre de regarder séparément le comportement du modèle sur les modes de Fourier e_1 et e_f . On montre ainsi que la composante BF est inchangée, comme précédemment la composante HF :

$$(\widehat{\mathcal{S}e_f})(\nu) = \hat{e}_f(\nu) - \sum_{k \in \mathbb{Z}} I(2k+1) \delta(\nu - (2k+1)f) e^{2ik\varphi} \quad (2.85)$$

$$= \hat{e}_f(\nu). \quad (2.86)$$

Le comportement est assez différent en revanche pour la composante HF :

$$(\widehat{\mathcal{S}e_1})(\nu) = \hat{e}_1(\nu) - \sum_{k \in \mathbb{Z}} I\left(\frac{2kf+1}{f}\right) \delta(\nu - 2kf - 1) e^{2ik\varphi} \quad (2.87)$$

qui si on utilise à nouveau l'approximation $I(\nu) \simeq 0$ pour $\nu \notin [-1, -1]$ donne

$$(\widehat{\mathcal{S}e_1})(\nu) \simeq \hat{e}_1(\nu) - I\left(\frac{1 - 2k_f f}{f}\right) \delta(\nu - 1 + 2k_f f) e^{-2ik_f \varphi}, \quad (2.88)$$

avec $k_f \in \mathbb{Z}$ tel que $2k_f - 1 < 1/f < 2k_f + 1$.

En combinant ces deux résultats, on obtient que le premier IMF après une itération vaut dans le cadre du modèle

$$d_1^{(1)}(t; a, f, \varphi) \simeq x(t; a, f, \varphi) - I\left(\frac{2k_f f - 1}{f}\right) \cos(2\pi(2k_f f - 1)t + 2k_f \varphi). \quad (2.89)$$

Contrairement au cas $a \rightarrow 0$, on constate que le signal après une itération sort clairement du modèle à deux composantes sinusoïdales. Cependant, la nouvelle composante ayant une amplitude inférieure à celle de la composante HF, on peut tout de même supposer qu'elle perturbera peu les positions des extrema du signal, et ce d'autant plus que af^2 est grand devant 1. Si l'on applique alors le même modèle pour les itérations de tamisage suivantes, on obtient finalement les IMFs pour n itérations :

$$d_1^{(n)}(t; a, f, \varphi) \simeq x(t; a, f, \varphi) - \lambda_n \cos(2\pi(2k_f f - 1)t + 2k_f \varphi), \quad (2.90)$$

$$d_2^{(n)}(t; a, f, \varphi) = a_1^{(n)}(t; a, f, \varphi) \simeq \lambda_n \cos(2\pi(2k_f f - 1)t + 2k_f \varphi), \quad (2.91)$$

avec

$$\lambda_n = 1 - \left(1 - I\left(\frac{2k_f f - 1}{f}\right)\right)^n. \quad (2.92)$$

Comme on pouvait l'attendre au vu des résultats des simulations, le comportement de l'EMD ne peut être décrit en termes de filtrage linéaire. Celui-ci est en effet assez surprenant dans la mesure où au lieu de séparer plus ou moins les deux composantes, elle en rajoute une avec une fréquence encore plus basse que celle de la composante BF. En pratique, il faut noter que ceci peut fortement compromettre l'analyse par HHT de certains signaux puisque, dans le cas de la somme de sinusoïdes, on observerait des fréquences instantanées oscillant autour de f pour le premier IMF et égale à $|2k_f - 1|$ pour le second. Si on utilise plutôt une représentation temps-fréquence pour analyser les IMFs, on verra apparaître les trois fréquences $1, f$ et $|2k_f - 1|$, mais il sera délicat de distinguer celles qui étaient dans le signal initialement et celle qui apparaît par repliement. Dans tous les cas, il est très difficile de défendre le point de vue offert par l'EMD. La seule circonstance atténuante qu'on puisse lui trouver est finalement que la nouvelle fréquence $|2k_f - 1|$ se « voit » dans les enveloppes lorsqu'on observe le signal, mais ce n'est même pas toujours si évident (cf Fig. 2.21).

1.1.7 Confrontation avec l'expérience

Pour attester de la validité des modèles proposés dans les deux cas $af < 1$ ou $af^2 > 1$, on représente Fig. 2.22 l'écart entre le premier IMF issu de la décomposition par EMD de la somme de deux sinusoïdes et les modèles précédents en fonction des rapports de fréquences et d'amplitudes. Plus précisément, les quantités représentées sont :

$$\begin{cases} \frac{\|d_1^{(n)} - d_{1,\text{modèle } a \rightarrow 0}^{(n)}\|_{\ell_2}}{\|x_2\|_{\ell_2}} & \text{pour la colonne de gauche} \\ \frac{\|d_1^{(n)} - d_{1,\text{modèle } a \rightarrow \infty}^{(n)}\|_{\ell_2}}{\|x_1\|_{\ell_2}} & \text{pour la colonne de droite.} \end{cases} \quad (2.93)$$

On visualise ainsi l'adéquation entre le modèle et le comportement de l'EMD dans la zone $af < 1$, colonne de gauche et dans la zone $af^2 > 1$, colonne de droite. Dans chaque cas, l'écart au modèle est normalisé par la norme de la composante dont l'amplitude est la plus faible.

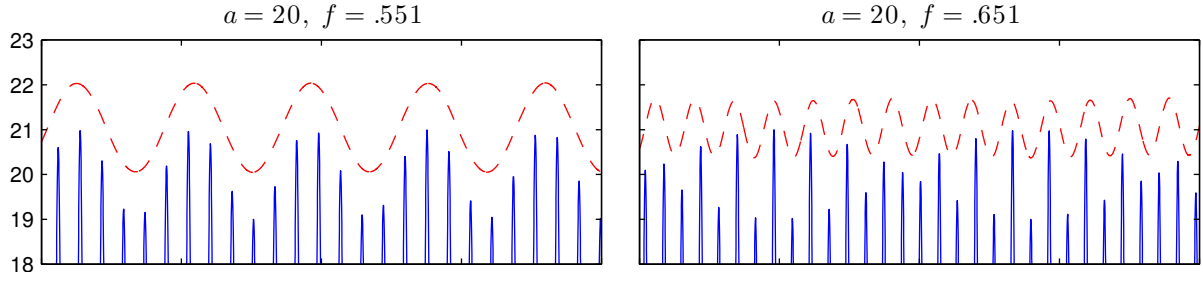


FIGURE 2.21 – Ressemblances entre moyenne des enveloppes et enveloppe supérieure. tirets rouges : moyenne des enveloppes (surélevée). trait plein bleu : signal $x(t; a, f, \varphi)$. A gauche, la moyenne des enveloppes est très similaire à l'enveloppe supérieure. A droite, la ressemblance est plus discutable. De manière générale, les enveloppes sont similaires à leur moyenne quand f est proche de $1/(2k)$, $k \in \mathbb{N}^*$ et pas du tout quand f est proche de $1/(2k + 1)$.

Les résultats montrent que le modèle décrit bien le comportement de l'EMD lorsque $a \rightarrow 0$ ou $a \rightarrow \infty$, mais surtout on observe qu'il reste très acceptable même pour des rapports d'amplitudes finis, et ce pratiquement jusqu'aux frontières $af < 1$ et $af^2 > 1$. Pour appuyer ce résultat, on représente Fig. 2.23 le critère c_1 tel que mesuré lors des simulations et la valeur prédite par le modèle pour deux valeurs de rapport d'amplitudes ($a = 0.01$ et $a = 1$) et trois nombres d'itérations différents. On observe sur cette figure que, en dehors d'un certain écart aux alentours de $f \approx 0.5$ sur les courbes correspondant à 10 itérations, les simulations sont très bien décrites par le modèle pour $a = 0.01$ et restent très proches même pour $a = 1$. L'écart aux alentours de $f \approx 0.5$ serait en fait lié à l'approximation faite dans le modèle selon laquelle $I(\nu) = 0$ dès que $|\nu| > 1$. En effet $f \approx 0.5$, donne $1 + f \approx 1.5$, et si $I(1.5) \approx 0.006$ reste petit, un modèle plus précis ferait intervenir $(1 - I(1.5))^{10} \approx 0.94$ qui s'écarte suffisamment de 1 pour créer des effets visibles. Cet écart est également presque visible dans la colonne de gauche de la figure Fig. 2.22 et correspond à la faible composante non combinaison linéaire des deux sinusoïdes initiales observée dans la colonne de droite de la figure Fig. 2.13.

On observe de plus que le modèle dans la zone où la composante HF domine ($af < 1$) semble rester valable jusqu'à un certain point dans la zone de transition. En fait, comme on l'a vu précédemment, même si la densité d'extrema du signal somme est strictement inférieure à celle de la composante HF, celle-ci peut augmenter au cours du processus de tamisage jusqu'à atteindre celle de la composante HF. Dès lors il n'est pas très surprenant de retrouver un résultat similaire au cas où la densité d'extrema est initialement égale à celle de la composante HF. Ceci fonctionne d'autant mieux que le filtre équivalent à une itération de tamisage $(1 - I(\nu))$ est quasi nul dès lors que $f \lesssim 0.3$ et que donc les effets de une ou plusieurs itérations de ce filtre sont assez proches dans cette zone.

1.1.8 Cas de composantes modulées en amplitude et en fréquence

L'EMD opérant à une échelle locale, on peut espérer que les résultats obtenus précédemment sur les sommes de sinusoïdes restent valables si les deux sinusoïdes sont modulées en amplitude et en fréquence, du moins tant que les modulations ne sont pas trop rapides et/ou trop prononcées.

1.1.8.1 Cas d'une composante modulée en amplitude et/ou en fréquence Dans ce but, on peut dans un premier temps s'intéresser à la question de savoir dans quelles conditions une sinusoïde modulée en amplitude et en fréquence peut être considérée comme un IMF. Dans le cas d'une simple modulation de fréquence, il est clair qu'on a toujours un IMF parfait dans la mesure où tous les maxima valent 1 et tous les minima -1, du moins si la fréquence instantanée reste toujours positive. Le problème est déjà moins évident dans le cas d'une modulation d'amplitude simple. En

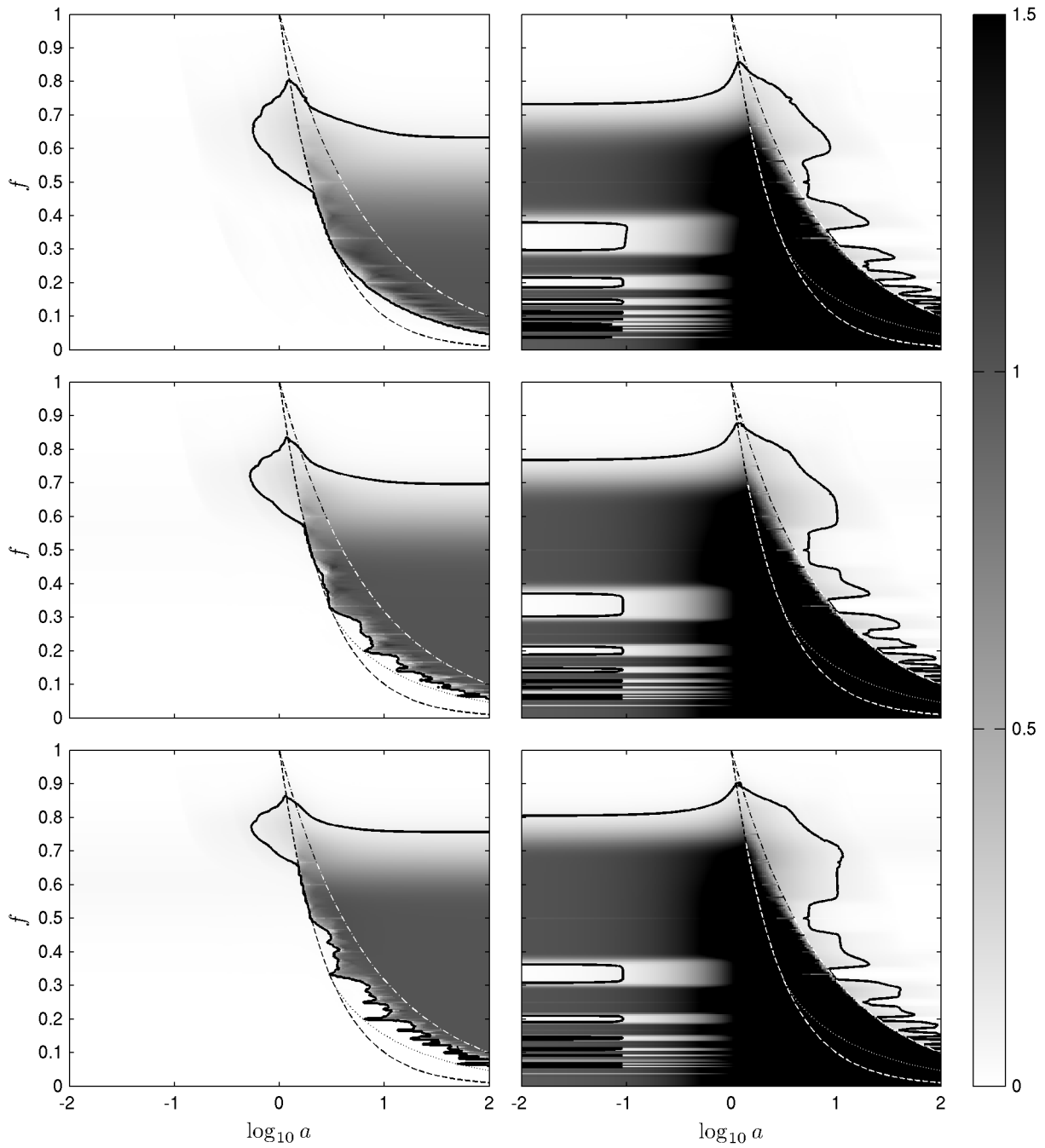


FIGURE 2.22 – Écart entre le modèle et les simulations pour 10 itérations de tamisage. Colonne de gauche : modèle $a \rightarrow 0$. Colonne de droite : modèle $a \rightarrow \infty$. Dans chaque cas la courbe épaisse noire correspond à la ligne de niveau = 0.1, c'est à dire que sur cette ligne la norme de l'écart entre le modèle et la simulation est un dixième de celle de la composante dont l'amplitude est la plus faible (x_2 quand $af < 1$ et x_1 quand $af^2 > 1$).

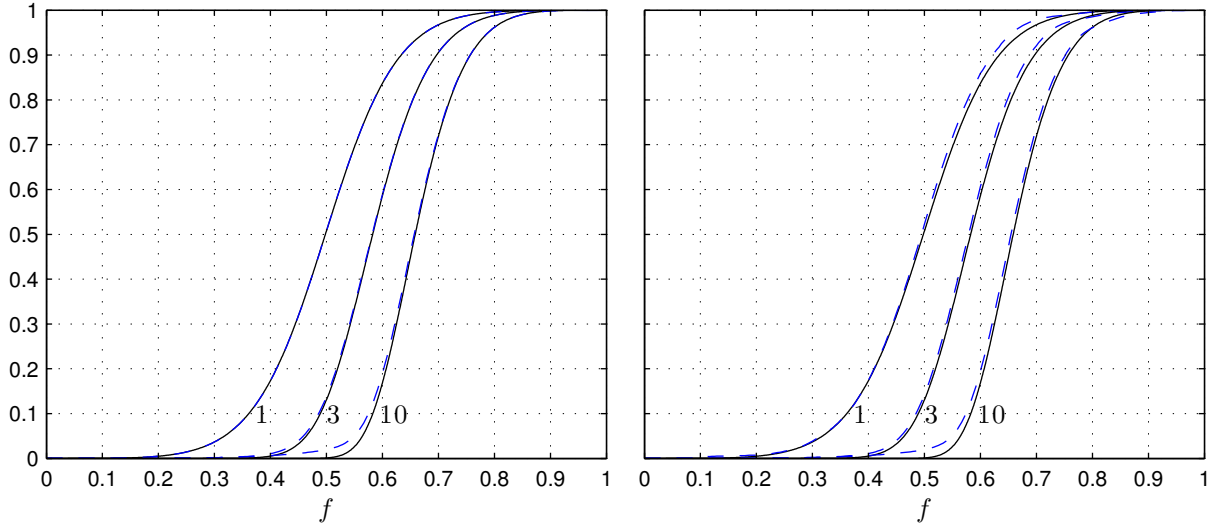


FIGURE 2.23 – Comparaison entre le critère c_1 mesuré lors des simulations (trait plein noir) et le c_1 prédit par le modèle (tirets bleus) pour 1, 3 et 10 itérations de tamisage. À gauche : $a = 0.01$. À droite : $a = 1$.

effet, il est clair qu'une modulation d'amplitude très lente par rapport à la fréquence de la sinusoïde donne a priori un IMF mais dans le cas inverse où la modulation varie rapidement par rapport à cette fréquence, c'est nettement plus délicat. Dans le cas d'une modulation sinusoïdale, on peut en fait montrer à l'aide du modèle que la sinusoïde modulée en amplitude ne peut être considérée comme un IMF que si le rapport entre la fréquence de modulation et la fréquence de la porteuse est inférieur à un certain seuil qui dépend du nombre d'itérations de tamisage. Considérons le signal modulé en amplitude suivant :

$$x(t) = (1 + m \cos(2\pi f_{AM}t)) \cos(2\pi t) \quad (2.94)$$

$$= \frac{m}{2} \cos(2\pi(1 - f_{AM})t) + \cos(2\pi t) + \frac{m}{2} \cos(2\pi(1 + f_{AM})t) \quad (2.95)$$

Si l'on suppose que ses extrema sont proches de ceux du signal non modulé, on peut lui appliquer le modèle (2.63) qui permet d'aboutir, avec les mêmes approximations que précédemment, à

$$d_1^{(n)} \simeq \cos(2\pi t) + \frac{m}{2} \cos(2\pi(1 + f_{AM})t) + m \left((1 - I(1 - f_{AM}))^n - \frac{1}{2} \right) \cos(2\pi(1 - f_{AM})t) \quad (2.96)$$

$$d_2^{(n)} \simeq m (1 - (1 - I(1 - f_{AM}))^n) \cos(2\pi(1 - f_{AM})t). \quad (2.97)$$

On voit ainsi que selon le modèle, on peut grossièrement considérer que le signal modulé en amplitude est un IMF uniquement si $(1 - I(1 - f_{AM}))^n$ est proche de 1, c'est-à-dire $f_{AM} < 1 - f_c^{(n)}$. Ce résultat est vérifié empiriquement Fig. 2.24. Dans le cas où, à l'inverse, $f_{AM} > 1 - f_c^{(n)}$, le comportement de l'EMD est un peu curieux. Grossièrement, si l'on considère que $(1 - I(1 - f_{AM}))^n = 1$ si $f_{AM} > 1 - f_c^{(n)}$ et 0 sinon, on voit que le signal modulé en amplitude (2.95) est converti en

$$d_1^{(n)} \approx \cos(2\pi t) + \frac{m}{2} \cos(2\pi(1 + f_{AM})t) - \frac{m}{2} \cos(2\pi(1 - f_{AM})t), \quad (2.98)$$

qui peut être vu comme une sinusoïde modulée en fréquence, à laquelle il faut ajouter $m \cos(2\pi(1 - f_{AM})t)$ pour retrouver la modulation d'amplitude. Ce résultat a été initialement observé dans [31] dans le cas particulier $m = 0.5$ et $f_{AM} = 0.5$ et constitue la première observation de la présence de repliement spectral dans l'EMD.

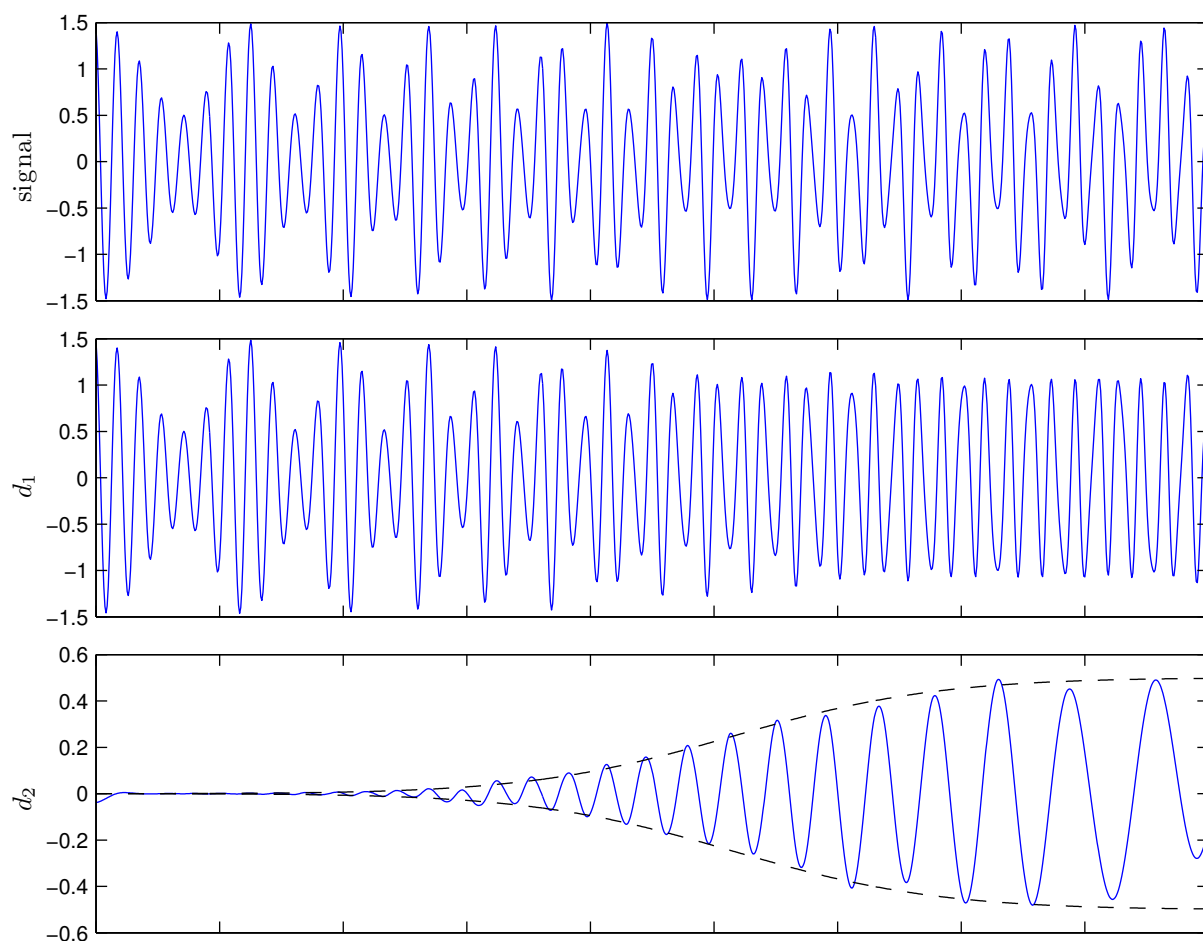


FIGURE 2.24 – EMD (10 itérations de tamisage) d’une sinusoïde modulée en amplitude par une modulation de fréquence linéaire dont la fréquence va, de gauche à droite, de 0.1 à 0.8 fois celle de la sinusoïde. Les « enveloppes » superposées au deuxième IMF correspondent à l’amplitude prédite par le modèle (2.97) en supposant la modulation localement sinusoïdale.

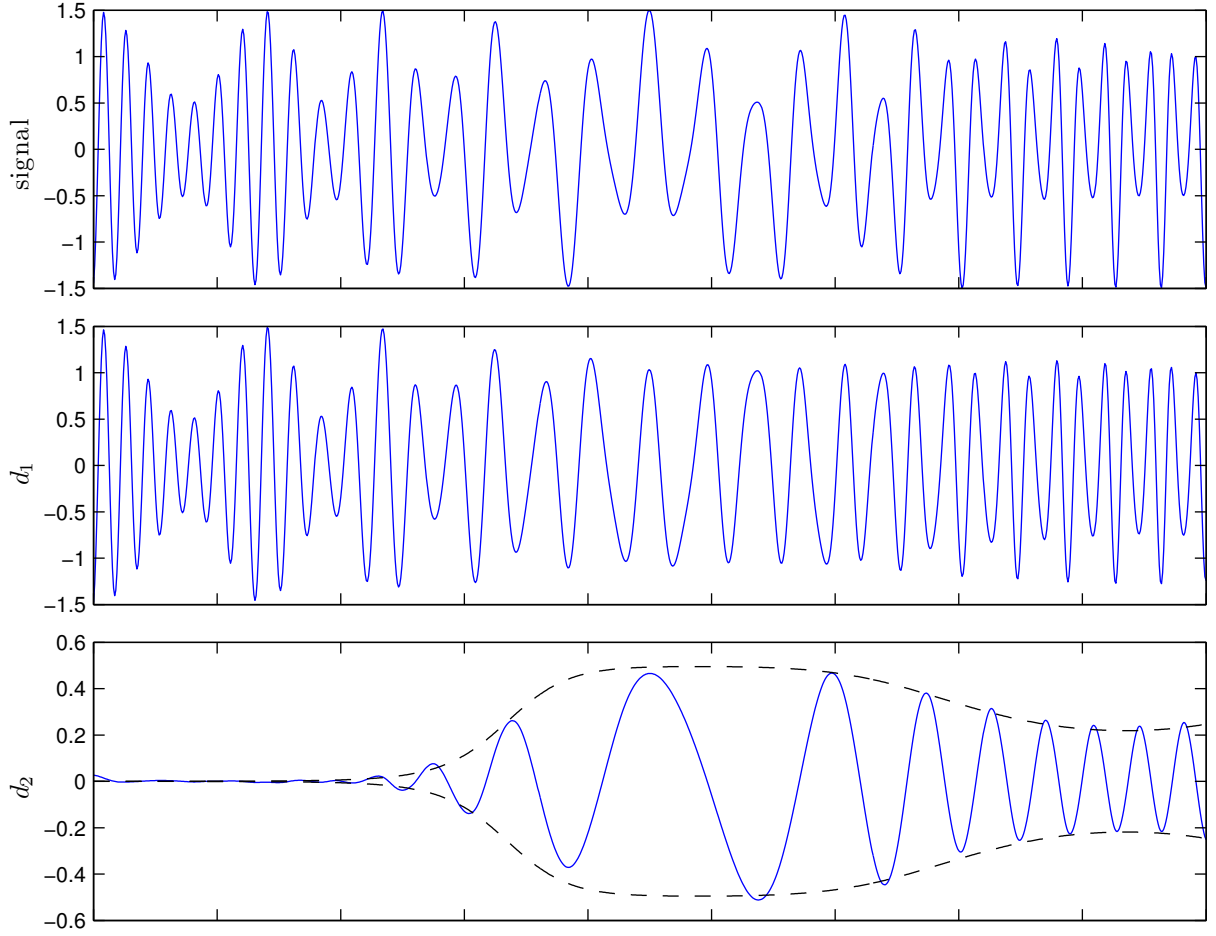


FIGURE 2.25 – EMD (10 itérations de tamisage) d'une sinusoïde modulée en amplitude et en fréquence. Les « enveloppes » superposées au deuxième IMF correspondent à l'amplitude prédite par le modèle (2.97) en supposant la modulation localement sinusoïdale et la fréquence localement constante.

Plus généralement, pour une modulation d'amplitude quelconque

$$x(t) = a(t) \cos(2\pi t), \quad (2.99)$$

on pourra considérer que $x(t)$ est un IMF si

- $a(t) > 0$, nécessaire pour avoir des extrema uniformément répartis ou presque,
- la modulation ne crée pas d'extrema supplémentaires,
- le spectre de $a(t)$ est nul au-delà de $1 - f_c^{(n)}$ environ.

Si on s'intéresse enfin au cas d'une sinusoïde modulée en amplitude et en fréquence

$$x(t) = a(t) \cos(2\pi \int^t f(t') dt'), \quad (2.100)$$

l'étude précédente reste en grande partie valable si on peut considérer les extrema comme localement uniformément espacés (cf Fig. 2.25). Ainsi, on pourra donc considérer que cette sinusoïde est un IMF dès lors que le spectre de $a(t)$ est localement nul au-delà de $f(t)(1 - f_c^{(n)})$ et que les variations de $f(t)$ sont relativement lentes par rapport à celles de $a(t)$. Malheureusement, le modèle ne permet pas de donner un critère plus quantitatif.

1.1.8.2 Cas de deux composantes modulées On considère maintenant un signal constitué de 2 composantes sinusoïdales modulées en amplitude et en fréquence qui sont séparément des IMFs :

$$x(t) = a_1(t) \cos(\varphi_1(t)) + a_2(t) \cos(\varphi_2(t)), \quad (2.101)$$

et on définit les fréquences instantanées :

$$f_i(t) = \frac{1}{2\pi} \frac{d\varphi_i}{dt}. \quad (2.102)$$

Pour adapter les résultats obtenus sur les sommes de sinusoïdes, on applique simplement ces derniers à chaque instant t_0 indépendamment en supposant que les amplitudes et les fréquences des deux composantes sont constantes égales à $a_i(t_0)$ et $f_i(t_0)$. Ainsi, si on suppose par exemple $f_1(t_0) > f_2(t_0)$ le modèle devient :

$$d_1^{(n)}(t_0) = \begin{cases} a_1(t_0) \cos(\varphi_1(t_0)) + \left(1 - I\left(\frac{f_2(t_0)}{f_1(t_0)}\right)\right)^n a_2(t_0) \cos(\varphi_2(t_0)) & \text{si } \left[\frac{a_2(t_0)f_2(t_0)}{a_1(t_0)f_1(t_0)} < 1\right] \text{ ou } \left[\frac{f_2(t_0)}{f_1(t_0)} < \frac{1}{3} \text{ et } \frac{a_2(t_0)f_2(t_0)}{a_1(t_0)f_1(t_0)} \sin\left(\frac{3\pi f_2(t_0)}{2f_1(t_0)}\right) < 1\right] \\ x(t) - \left(1 - \left(1 - I\left(2k_f - \frac{f_1(t_0)}{f_2(t_0)}\right)\right)^n\right) \cos(2k_f \varphi_2(t_0) - \varphi_1(t_0)) & \text{si } \frac{a_2(t_0)f_2^2(t_0)}{a_1(t_0)f_1^2(t_0)} > 1 \\ \text{on ne sait pas} & \text{sinon,} \end{cases} \quad (2.103)$$

avec k_f l'entier tel que $2k_f - 1 < f_1(t_0)/f_2(t_0) < 2k_f + 1$.

Un exemple d'application de ce modèle est proposé Fig. 2.26. Les résultats sont clairement encourageants et nécessiteraient sans doute une étude plus approfondie.

1.2 Généralisations

On s'intéresse ici à la généralisation du modèle développé pour les sommes de sinusoïdes au cas plus général des mélanges de signaux non linéaires dont les extrema sont uniformément espacés. Dans un premier temps, on généralisera les résultats sur les extrema du paragraphe 1.1.6.1 à une classe de signaux englobant une très grande majorité des situations pratiques puis dans un deuxième temps, on appliquera le modèle (2.66) à des sommes de signaux non linéaires dans l'objectif de pointer quelques comportements généraux.

1.2.1 À propos des extrema

Tout comme pour les sommes de deux composantes sinusoïdales, on peut montrer dans un cadre très général que, d'une part les extrema d'une somme de signaux tendent vers ceux de l'une des deux composantes si l'amplitude de l'autre tend vers zéro, et que d'autre part, la densité d'extrema de la somme est la même que celle de la composante dominante pour peu que son amplitude soit supérieure à un certain seuil.

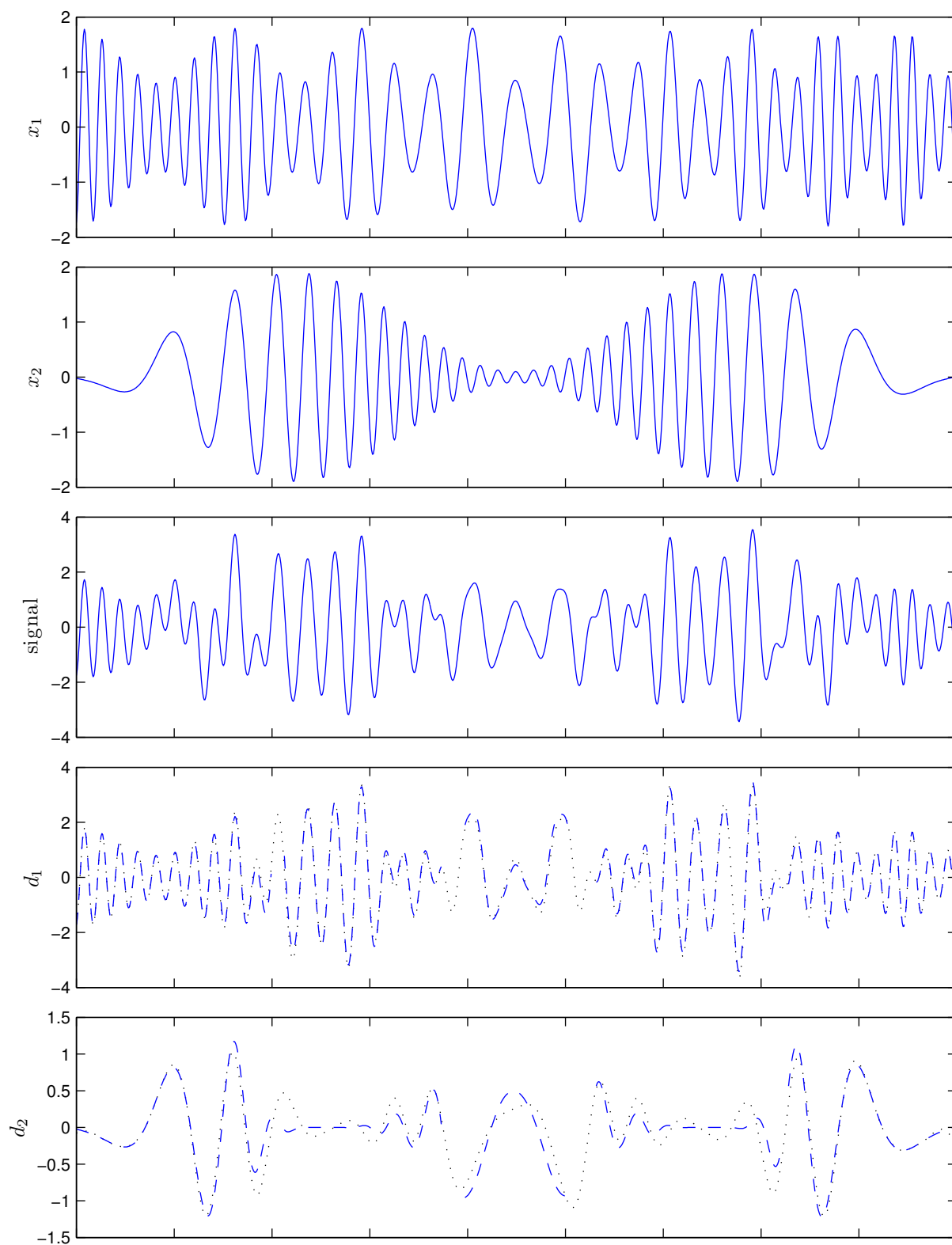


FIGURE 2.26 – EMD (10 itérations de tamisage) d'une somme de deux composantes sinusoïdales modulées en amplitude et en fréquence. Dans les deux cadres du bas de la figure, on a représenté en pointillés noirs les IMFs obtenus et en tirets bleus les prédictions du modèle (2.103).

Proposition 4. Soient $x_1 : t \mapsto x_1(t)$ et $x_2 : t \mapsto x_2(t)$ deux fonctions de $\mathbb{R} \rightarrow \mathbb{R}$ vérifiant

- (i) les fonctions sont au minimum deux fois continuellement dérivables : $x_i(t) \in \mathcal{C}^2(\mathbb{R})$.
- (ii) les dérivées premières et secondes sont bornées : $\forall t \in \mathbb{R}, |\partial_t x_i(t)| < M$ et $|\partial_t^2 x_i(t)| < M'$.
- (iii) les ensembles d'extrema sont non bornés : $\inf \mathcal{E}_{x_i} = -\infty$ et $\sup \mathcal{E}_{x_i} = +\infty$.
- (iv) $\exists \eta, \epsilon \in \mathbb{R}_+^*, \forall t \in \mathbb{R}, |\partial_t x_i(t)| < \eta \implies |\partial_t^2 x_i(t)| > \epsilon$.

La fonction $x : t \mapsto x(t; a, f, \varphi) = x_1(t) + ax_2(ft + \varphi)$ avec $f \in [0, 1]$, $a \in \mathbb{R}$ et $\varphi \in \mathbb{R}$ vérifie

- Les positions des extrema de $x(t; a, f, \varphi)$ sont les mêmes que celles de $x_1(t)$ quand $a \rightarrow 0$ et que celles de $x_2(ft + \varphi)$ quand $a \rightarrow \pm\infty$
- Il existe deux réels positifs A et B tels que pour tout intervalle $I \subseteq \mathbb{R}$, le nombre d'extrema de x dans I ($\# [\mathcal{E}_x \cap I]$) vérifie

$$\# [\mathcal{E}_x \cap I] \in \begin{cases} [\# [\mathcal{E}_{x_1} \cap I] - 2, \# [\mathcal{E}_{x_1} \cap I] + 2] & \text{si } |a|f < A \\ [\# [\mathcal{E}_{x_2(ft+\varphi)} \cap I] - 2, \# [\mathcal{E}_{x_2(ft+\varphi)} \cap I] + 2] & \text{si } |a|f^2 > B \end{cases} \quad (2.104)$$

Quelques remarques par rapport aux hypothèses :

- L'hypothèse ((iv)) signifie généralement que les dérivées premières et secondes ne peuvent s'annuler simultanément. Cette condition a dû être renforcée pour faire face à des cas délicats avec des ensembles d'extrema $\mathcal{E}_{x_i}$ infinis tels que $\forall t_j \in \mathcal{E}_{x_i}, \partial_t^2 x_i|_{t=t_j} \neq 0$, et $\inf \{|\partial_t^2 x_i|_{t=t_j}|, t_j \in \mathcal{E}_{x_i}\} = 0$.
- L'hypothèse ((iii)) et le fait de considérer des fonctions de $\mathbb{R} \rightarrow \mathbb{R}$ permettent de s'assurer que chaque extremum de $x(t)$ a un homologue parmi les extrema de $x_1(t)$ ou de $x_2(ft + \varphi)$ suivant le cas sans avoir à se soucier d'effets de bords.

Démonstration. La preuve n'est donnée que dans le cas où les extrema de $x(t; a, f, \varphi)$ sont proches de ceux de la composante $x_1(t)$ ($|a|f < A$) dans la mesure où elle est quasiment identique dans l'autre cas.

Soient A et B définis par

$$A = \min \left\{ \frac{\eta}{M}, \frac{\epsilon}{M'} \right\}, \quad B = \max \left\{ \frac{M}{\eta}, \frac{M'}{\epsilon} \right\}. \quad (2.105)$$

La preuve s'appuie sur la définition d'une fonction bijective $\mathcal{F} : \mathcal{E}_x \rightarrow \mathcal{E}_{x_1}$ existant dès lors que $|a|f < A$ et telle que $\forall t \in \mathcal{E}_x, \mathcal{F}(t) - t \xrightarrow{a \rightarrow 0} 0$.

* Définition de $\mathcal{F}$:

Soit $t_0 \in \mathcal{E}_x$ un extremum de $x(t; a, f, \varphi)$ avec $|a|f < A$. On supposera par la suite que t_0 est un maximum, la correspondance pour l'autre cas étant évidente. Remarquons tout d'abord que $|\partial_t x_1(t_0)| < \eta$:

$$t_0 \in \mathcal{E}_x \implies \partial_t x(t_0; a, f, \varphi) = 0 \implies |\partial_t x_1(t_0)| = |-af \partial_t x_2(ft + \varphi)| < |a|fM < \eta \quad (2.106)$$

Soit alors I_0 le plus grand intervalle ouvert $I_0 =]b, c[$ contenant t_0 tel que $\forall t \in I_0, |\partial_t x_1(t)| < \eta$. I_0 vérifie :

- I_0 est non vide : $\partial_t x_1(t)$ est continue et t_0 est dans l'image réciproque par $\partial_t x_1(t)$ de $] -\eta, \eta[$ qui est ouvert.
- I_0 est borné :

D'après l'hypothèse ((iii)), il existe nécessairement deux minima t_- et t_+ de $x_1(t)$ tels que $t_- < t_0 < t_+$. En outre, l'hypothèse ((iv)) et le type des extrema impliquent que $\partial_t^2 x_1(t_-) > \epsilon$, $\partial_t^2 x_1(t_0) < -\epsilon$ et $\partial_t^2 x_1(t_+) > \epsilon$.

Les dérivées secondes étant continues, il existe alors deux points t'_- et t'_+ tels que

$$t_- < t'_- < t_0 < t'_+ < t_+ \quad \text{et} \quad \partial_t^2(t'_-) = \partial_t^2(t'_+) = 0, \quad (2.107)$$

ce qui implique finalement $|\partial_t(t'_-)| \geq \eta$ et $|\partial_t(t'_+)| \geq \eta$ et donc $I_0 \subset]t'_-, t'_+[$.

Montrons maintenant que I_0 contient un unique extremum de $x_1(t)$ et qu'il s'agit d'un maximum. $\partial_t x_1(t)$ étant continue et I_0 étant le plus grand intervalle contenant t_0 et tel que $\forall t \in I_0, |\partial_t x_1(t)| < \eta$ on a nécessairement $|\partial_t x_1(b)| = |\partial_t x_1(c)| = \eta$.

D'après l'hypothèse ((iv)),

$$\forall t \in I_0, |\partial_t^2 x_1(t)| > \epsilon. \quad (2.108)$$

De plus, t_0 étant un maximum, on a nécessairement $\partial_t^2 x_1(t_0) < -\epsilon$. $\partial_t^2 x_1(t)$ étant continue, on en déduit que

$$\forall t \in I_0, \partial_t^2 x_1(t) < -\epsilon. \quad (2.109)$$

ce qui implique en particulier que $\partial_t x_1(t)$ est strictement décroissante sur I_0 . Par conséquent, on a nécessairement $\partial_t x_1(b) = \eta$ et $\partial_t x_1(c) = -\eta$ et on en déduit qu'il existe un unique $t'_0 \in I_0$ tel que $\forall t \in I_0, t < t'_0 \implies \partial_t x_1(t) > 0$ et $t > t'_0 \implies \partial_t x_1(t) < 0$. t'_0 est alors l'unique maximum de $x_1(t)$ dans I_0 ce qui permet de définir $\mathcal{F}(t_0) = t'_0$.

On définit également la fonction $\mathcal{G}$ qui à tout élément de l'image réciproque par $\partial_t x_1$ de $] -\eta, \eta[$ associe un intervalle de $\mathbb{R}$ ($I(\mathbb{R})$ désigne l'ensemble des intervalles de $\mathbb{R}$) selon le principe précédemment utilisé pour définir I_0 :

$$\begin{aligned} \mathcal{G} : (\partial_t x_1)^{-1} (] -\eta, \eta[) &\longrightarrow I(\mathbb{R}) \\ t &\longmapsto I = \bigcup \{I' \in I(\mathbb{R}), t \in I' \text{ et } \forall t' \in I', |\partial_t x_1(t')| < \eta\} \end{aligned} \quad (2.110)$$

En reprenant les arguments développés précédemment, on peut montrer que

- $\forall t \in (\partial_t x_1)^{-1} (] -\eta, \eta[), \mathcal{G}(t)$ est borné, non vide
- $\forall t \in (\partial_t x_1)^{-1} (] -\eta, \eta[), \partial_t x_1(t)$ est strictement monotone sur $\mathcal{G}(t) =]b, c[$ et ses bornes vérifient $\partial_t x_1(b) = \pm\eta$ et, avec le signe opposé, $\partial_t x_1(c) = \mp\eta$ (ce qui implique au passage que $\mathcal{G}(t)$ est bien un intervalle ouvert)
- $\mathcal{E}_x \subseteq (\partial_t x_1)^{-1} (] -\eta, \eta[)$ si $|a|f < A$

On peut également ajouter que

- $\mathcal{E}_{x_1} \subseteq (\partial_t x_1)^{-1} (] -\eta, \eta[)$
- $\forall t \in \mathcal{E}_x, \mathcal{F}(t) \in \mathcal{G}(t)$, par construction
- $\forall t, t' \in (\partial_t x_1)^{-1} (] -\eta, \eta[), \mathcal{G}(t) \cap \mathcal{G}(t') \neq \emptyset \implies \mathcal{G}(t) = \mathcal{G}(t')$, qui découle directement de la définition de $\mathcal{G}$

Enfin on démontrera par la suite que $\forall t \in (\partial_t x_1)^{-1} (] -\eta, \eta[), \partial_t x(t)$ est également strictement monotone sur $\mathcal{G}(t)$.

* $\mathcal{F}$ est injective :

Soient $t_0, t_1 \in \mathcal{E}_x$ tels que $\mathcal{F}(t_0) = \mathcal{F}(t_1) = t'_0$. Soient $I_0 = \mathcal{G}(t_0)$ et $I_1 = \mathcal{G}(t_1)$. Comme $t'_0 \in I_0, t'_0 \in I_1$ et que $|\partial_t x_1(t)| < \eta$ sur I_0 et I_1 on a nécessairement $I_1 = I_0$. De plus, $|\partial_t x_1(t)| < \eta \implies |\partial_t^2 x_1(t)| > \epsilon$. En supposant que t'_0 est un maximum de $x_1(t)$, on en déduit que $\forall t \in I_0, \partial_t^2 x_1(t) < -\epsilon$ ce qui implique $\forall t \in I_0, \partial^2 x(t) < -\epsilon + af^2 M' < 0$ et donc que $\partial x(t)$ est strictement décroissante sur I_0 , ce qui implique qu'il ne peut y avoir qu'un extremum de $x(t)$ dans I_0 . D'où $t_0 = t_1$ et l'injectivité de $\mathcal{F}$.

* $\mathcal{F}$ est surjective :

Soit $t'_0 \in \mathcal{E}_{x_1}$. On a $|\partial_t x_1(t'_0)| < \eta$ et on peut donc définir l'intervalle $I'_0 =]b', c'[= \mathcal{G}(t'_0)$. En supposant par exemple que t'_0 est un maximum de $x_1(t)$, on a de plus que $\partial_t x_1(b') = \eta$ et $\partial_t x_1(c') = -\eta$. Si $|a|f < A$, ceci implique que $\partial_t x(b') > \eta - AM \geq 0$ et $\partial_t x(c') < -\eta + AM \leq 0$ et donc que $x(t)$ passe par un maximum local t_0 dans l'intervalle I'_0 . Il est alors clair que $\mathcal{G}(t_0) = I'_0$ et donc que $\mathcal{F}(t_0) = t'_0$. De plus, on a montré au passage que $\mathcal{F}^{-1}(t'_0) \in \mathcal{G}(t'_0)$.

* Les extrema de $x(t; a, f, \varphi)$ tendent vers ceux de $x_1(t)$ quand $a \rightarrow 0$:

Soit $t'_0 \in \mathcal{E}_{x_1}$ et $I'_0 = \mathcal{G}(t'_0)$. Sachant que $\partial_t x_1(t)$ est strictement monotone sur I'_0 , elle est bijective et on peut donc définir g_0 , la fonction définie sur $\partial_t x_1(I'_0)$, inverse de la restriction de $\partial_t x_1(t)$ à I'_0 .

De là,

$$\mathcal{F}^{-1}(t'_0) - t'_0 = g_0(\partial_t x_1(\mathcal{F}^{-1}(t'_0))) - g_0(\partial_t x_1(t'_0)) \quad (2.111)$$

$$= g_0(-af\partial_t x_2(f\mathcal{F}^{-1}(t'_0) + \varphi)) - g_0(0), \quad (2.112)$$

et donc $|\mathcal{F}^{-1}(t'_0) - t'_0| \xrightarrow{a \rightarrow 0} 0$ parce que g_0 est continue et $|af\partial_t x_2(f\mathcal{F}^{-1}(t'_0) + \varphi)| \leq |afM| \xrightarrow{a \rightarrow 0} 0$

* Pour tout intervalle I , les nombres d'extrema de $x_1(t)$ et de $x(t)$ dans I diffèrent d'au plus 2

Soit $I =]b, c[$ un intervalle éventuellement non borné, ouvert ou fermé étant sans grande importance. Si on considère sa borne inférieure, on peut distinguer deux cas suivant que la proposition « $\exists t_0 \in \mathcal{E}_x, \inf \mathcal{G}(t_0) < b$ et $\sup \mathcal{G}(t_0) > b$ » est vraie ou fausse :

– si elle est vraie, alors

$$\forall t \in \mathcal{E}_x, t \neq t_0 \text{ et } t > b \implies \mathcal{F}(t) > b. \quad (2.113)$$

En effet, $\mathcal{F}(t) \leq b \implies \inf \mathcal{G}(t) < b \implies b \in \mathcal{G}(t) \cap \mathcal{G}(t_0) \implies \mathcal{G}(t) = \mathcal{G}(t_0)$ et $\partial_t x(t)$ étant strictement monotone sur $\mathcal{G}(t_0)$, il ne peut y avoir qu'un extremum de $x(t)$ dans l'intervalle ce qui implique que $t = t_0$. On peut reformuler (2.113) de la manière suivante

$$\mathcal{F}(\mathcal{E}_x \cap]b, +\infty[\setminus \{t_0\}) \subseteq]b, +\infty[\quad (2.114)$$

qui implique

$$\mathcal{F}(\mathcal{E}_x \cap]b, +\infty[\setminus \{t_0\}) \subseteq \mathcal{E}_{x_1} \cap]b, +\infty[\setminus \{\mathcal{F}(t_0)\}, \quad (2.115)$$

car $\mathcal{F}(\mathcal{E}_x \setminus \{t_0\}) = \mathcal{E}_{x_1} \setminus \{\mathcal{F}(t_0)\}$.

– si elle est fausse,

$$\forall t \in \mathcal{E}_x, t > b \implies \inf \mathcal{G}(t) \geq b \implies \mathcal{F}(t) > b, \quad (2.116)$$

ce qui implique que

$$\mathcal{F}(\mathcal{E}_x \cap]b, +\infty[) \subseteq \mathcal{E}_{x_1} \cap]b, +\infty[. \quad (2.117)$$

Avec le même raisonnement sur la borne supérieure de I , on peut montrer que

– soit $\exists t_1 \in \mathcal{E}_x$ tel que

$$\mathcal{F}(\mathcal{E}_x \cap]-\infty, c[\setminus \{t_1\}) \subseteq \mathcal{E}_{x_1} \cap]-\infty, c[\setminus \{\mathcal{F}(t_1)\}, \quad (2.118)$$

– soit

$$\mathcal{F}(\mathcal{E}_x \cap]-\infty, c[) \subseteq \mathcal{E}_{x_1} \cap]-\infty, c[. \quad (2.119)$$

En combinant les résultats aux deux bornes de I , on est dans l'un des 4 cas :

- $\mathcal{F}(\mathcal{E}_x \cap I) \subseteq \mathcal{E}_{x_1} \cap I$,
- $\mathcal{F}(\mathcal{E}_x \setminus \{t_0\} \cap I) \subseteq \mathcal{E}_{x_1} \cap I \setminus \{\mathcal{F}(t_0)\}$,
- $\mathcal{F}(\mathcal{E}_x \setminus \{t_1\} \cap I) \subseteq \mathcal{E}_{x_1} \cap I \setminus \{\mathcal{F}(t_1)\}$,
- $\mathcal{F}(\mathcal{E}_x \setminus \{t_0, t_1\} \cap I) \subseteq \mathcal{E}_{x_1} \cap I \setminus \{\mathcal{F}(t_0, t_1)\}$,

ce qui implique que

$$\#[\mathcal{E}_x \cap I] \leq \#[\mathcal{E}_{x_1} \cap I] + 2. \quad (2.120)$$

Enfin le même raisonnement en renversant les rôles de $x(t)$ et $x_1(t)$ permet de conclure. $\square$

1.2.2 Application du modèle à des mélanges de deux ondes faiblement non linéaires

L'objectif de cette section est d'étudier en quoi le modèle (2.66) peut nous éclairer sur les capacités de l'EMD à extraire des composantes non linéaires. Dans cette optique, on s'intéresse ici à la manière dont l'EMD décompose des sommes de deux composantes non linéaires vérifiant quelques propriétés permettant d'appliquer le modèle, au moins dans les cas asymptotiques. Ainsi, on s'intéressera à des signaux qui d'une part ont des ensembles d'extrema périodiques, maxima et minima indépendamment,

et d'autre part vérifient les hypothèses de la proposition 4 pour s'assurer que le modèle puisse être valide au moins asymptotiquement. En outre, on considérera uniquement des signaux périodiques avec seulement deux extrema par période de manière à pouvoir donner un sens aux mesures globales de comportement effectuées grâce aux critères définis à la section 1.1.4. L'autre intérêt de cette dernière simplification est de pouvoir aisément construire de tels signaux qui soient également des IMFs « parfaits », c'est-à-dire des signaux dont la moyenne des enveloppes est parfaitement nulle.

Le modèle de signal bi-composantes dans la suite sera

$$x(t; a, f, \varphi) = s_1(t) + as_2(ft + \varphi), \quad (2.121)$$

où $s_1(t)$ et $s_2(t)$ sont deux fonctions périodiques de période 1 avec deux extrema par période qu'on caractérise par leurs développements en séries de Fourier :

$$s_i(t) = \sum_{p \in \mathbb{Z}} c_{i,p} e^{i2p\pi t}. \quad (2.122)$$

On supposera en outre que ces fonctions ont leurs maxima à des instants entiers ($t_{max} \in \mathbb{Z}$), leurs minima décalés de α_{s_i} par rapport aux maxima ($t_{min} \in \alpha_{s_i} + \mathbb{Z}$) et qu'elles sont des IMFs, c'est-à-dire que $s_i(0) = -s_i(\alpha_{s_i})$.

Comme dans le cas sinusoïdal, on se référera à la composante hautes fréquences (HF) :

$$x_1(t) = s_1(t), \quad (2.123)$$

et basses fréquences (BF) :

$$x_2(t) = as_2(ft + \varphi). \quad (2.124)$$

1.2.2.1 Premier IMF dans le cas $a \rightarrow 0$ On détaille ici l'application du modèle pour des signaux non linéaires en supposant que les extrema du signal composite sont proches de ceux de la composante HF. L'objectif de cette section étant d'étudier les capacités de l'EMD à extraire des composantes non linéaires, on ne s'intéressera pas au cas où les extrema sont à l'inverse proches de ceux de la composante BF parce qu'on sait qu'il n'y a dans ce cas aucune chance de récupérer les composantes non linéaires. On pourrait montrer que le comportement de l'EMD dans ce cas est similaire à celui observé sur les sommes de sinusoïdes lorsque les extrema sont proches de ceux de la composante BF ($af^2 > 1$) : elle introduit des composantes basses fréquences (une à deux par harmonique de la composante HF), absentes du signal, qui se compensent entre les différents IMFs.

Par la suite, on supposera donc que les extrema du signal composite coïncident avec ceux de la composante HF situés aux instants entiers pour les maxima et décalés de α_{s_1} pour les minima. De plus, on supposera que les extrema restent à ces emplacements après un nombre quelconque d'itérations de tamisage. Le modèle prend alors la forme

$$(\widehat{\mathcal{S}x})(\nu) = \hat{x}(\nu) - I(\nu) (\text{III}_{mean} * \hat{x})(\nu), \quad (2.125)$$

avec

$$\text{III}_{mean} \sum_{k \in \mathbb{Z}} e^{-ik\pi\alpha} \cos(k\pi\alpha) \delta(\nu - k), \quad (2.126)$$

où $\alpha = \alpha_{s_1}$. Par la suite, on notera $\alpha_k = e^{-ik\pi\alpha} \cos(k\pi\alpha)$ pour alléger les notations.

Dans la mesure où le modèle est linéaire dès lors que l'opérateur d'échantillonnage est fixé, on traitera chaque composante indépendamment comme on l'a fait dans le cas sinusoïdal.

La composante HF est inaltérée par le processus de tamisage L'application du modèle sur la composante HF donne

$$(\widehat{\mathcal{S}x_1})(\nu) = \hat{x}_1(\nu) - \sum_p \sum_k \alpha_k c_{1,p} I(k+p) \delta(\nu - k - p), \quad (2.127)$$

qui en utilisant le fait que $I(k) = \delta_{k,0}$ (symbole de Kronecker) donne

$$(\widehat{\mathcal{S}x_1})(\nu) = \hat{x}_1(\nu) - \sum_p \alpha_p c_{1,p} \delta(\nu) = \hat{x}_1(\nu). \quad (2.128)$$

En effet, $x_1(t)$ étant un IMF, on a $x_1(0) + x_1(\alpha) = 2 \sum_p \alpha_p c_{1,p} = 0$. On en déduit que si l'opérateur d'échantillonnage reste identique, l'effet d'un nombre quelconque d'itérations de tamisage sur la composante HF est nul.

Cas de la composante BF Étant donnée la complexité des calculs à suivre, on applique le modèle sur chaque harmonique séparément. Considérons l'harmonique d'ordre p à la fréquence pf . Le comportement du modèle sur cet harmonique est décrit par $\mathcal{S}e_{pf}$:

$$(\widehat{\mathcal{S}e_{pf}})(\nu) = \delta(\nu - pf) - \sum_k \alpha_k I(pf + k) \delta(\nu - k - pf), \quad (2.129)$$

qu'on peut simplifier en faisant à nouveau l'approximation que $I(\nu) \simeq 0$ si $\nu \notin [-1, 1]$:

$$(\widehat{\mathcal{S}e_{pf}})(\nu) \simeq \delta(\nu - pf) - \alpha_{k_1} I(pf + k_1) \delta(\nu - k_1 - pf) - \alpha_{k_2} I(pf + k_2) \delta(\nu - k_2 - pf), \quad (2.130)$$

où k_1 et k_2 sont les deux entiers consécutifs tels que $-1 \leq pf + k_i \leq 1$.

On observe ainsi que suivant que $pf \in [-1, 1]$ ou non, l'opérateur de tamisage introduit une ou deux nouvelles fréquences. Lors des itérations suivantes, toutes ces fréquences interagissent entre elles ce qui oblige à les considérer toutes simultanément, le coefficient de chacune de ces composantes à l'itération $n+1$ étant une combinaison linéaire des coefficients de toutes ces composantes à l'itération n .

Dans le cas où $pf \in [-1, 1]$, les coefficients $a_{pf}^{(n)}$ et $a_{pf+k}^{(n)}$ des deux composantes à pf et $pf+k$ dans $\mathcal{S}e_{pf}$, où $k = \pm 1$ est tel que $pf+k \in [-1, 1]$, évoluent selon l'équation

$$\begin{pmatrix} a_{pf}^{(n+1)} \\ a_{pf+k}^{(n+1)} \end{pmatrix} \simeq \begin{pmatrix} 1 - I(pf) & -\alpha_{-k} I(pf) \\ -\alpha_k I(pf+k) & 1 - I(pf+k) \end{pmatrix} \begin{pmatrix} a_{pf}^{(n)} \\ a_{pf+k}^{(n)} \end{pmatrix}, \quad (2.131)$$

et donc $(\mathcal{S}^n e_{pf})(t) = a_{pf}^{(n)} e_{pf}(t) + a_{pf+k}^{(n)} e_{pf+k}(t)$ avec

$$\begin{pmatrix} a_{pf}^{(n)} \\ a_{pf+k}^{(n)} \end{pmatrix} \simeq \begin{pmatrix} 1 - I(pf) & -\alpha_{-k} I(pf) \\ -\alpha_k I(pf+k) & 1 - I(pf+k) \end{pmatrix}^n \begin{pmatrix} 1 \\ 0 \end{pmatrix} \quad (2.132)$$

ce qui donne finalement

$$\begin{cases} a_{pf}^{(n)} \simeq \frac{1 - \lambda_- - I(pf)}{\lambda_+ - \lambda_-} \lambda_+^n + \frac{I(pf) + \lambda_+ - 1}{\lambda_+ - \lambda_-} \lambda_-^n, \\ a_{pf+k}^{(n)} \simeq -\frac{\alpha_k I(pf+k)}{\lambda_+ - \lambda_-} (\lambda_+^n - \lambda_-^n), \end{cases} \quad (2.133)$$

où $\lambda_{\pm}$ sont les deux valeurs propres de la matrice de transition :

$$\lambda_{\pm} = 1 - \frac{I(pf) + I(pf + k)}{2} \pm \frac{1}{2} \sqrt{(I(pf) - I(pf + k))^2 + 4I(pf)I(pf + k) \cos^2(\pi\alpha)} \quad (2.134)$$

De manière analogue, dans le cas où $pf \notin [-1, 1]$, les coefficients $a_{pf}^{(n)}$, $a_{pf+k_1}^{(n)}$ et $a_{pf+k_2}^{(n)}$ évoluent selon

$$\begin{pmatrix} a_{pf}^{(n+1)} \\ a_{pf+k_1}^{(n+1)} \\ a_{pf+k_2}^{(n+1)} \end{pmatrix} \simeq \begin{pmatrix} 1 & 0 & 0 \\ -\alpha_{k_1}I(pf + k_1) & 1 - I(pf + k_1) & -\alpha_{k_1-k_2}I(pf + k_1) \\ -\alpha_{k_2}I(pf + k_2) & -\alpha_{k_2-k_1}I(pf + k_2) & 1 - I(pf + k_2) \end{pmatrix} \begin{pmatrix} a_{pf}^{(n)} \\ a_{pf+k_1}^{(n)} \\ a_{pf+k_2}^{(n)} \end{pmatrix}, \quad (2.135)$$

ce qui donne au final

$$\left\{ \begin{array}{l} a_{pf}^{(n)} \simeq 1, \\ a_{pf+k_1}^{(n)} \simeq I(pf + k_1) \left[\frac{(\alpha_{k_1}(\lambda_+ + \lambda_- + I(pf + k_1) - 2) + \alpha_{k_1-k_2}\alpha_{k_2}I(pf + k_2))}{(1 - \lambda_+)(1 - \lambda_-)} \right. \\ \quad \left. - \frac{\alpha_{k_1}(\lambda_- + I(pf + k_1) - 1) + \alpha_{k_1-k_2}\alpha_{k_2}I(pf + k_2)}{(\lambda_+ - \lambda_-)(1 - \lambda_+)} \lambda_+^n \right. \\ \quad \left. + \frac{\alpha_{k_1}(\lambda_+ + I(pf + k_1) - 1) + \alpha_{k_1-k_2}\alpha_{k_2}I(pf + k_2)}{(\lambda_+ - \lambda_-)(1 - \lambda_-)} \lambda_-^n \right], \\ a_{pf+k_2}^{(n)} \simeq I(pf + k_2) \left[\frac{(\alpha_{k_2}(\lambda_+ + \lambda_- + I(pf + k_2) - 2) + \alpha_{k_2-k_1}\alpha_{k_1}I(pf + k_1))}{(1 - \lambda_+)(1 - \lambda_-)} \right. \\ \quad \left. - \frac{\alpha_{k_2}(\lambda_- + I(pf + k_2) - 1) + \alpha_{k_2-k_1}\alpha_{k_1}I(pf + k_1)}{(\lambda_+ - \lambda_-)(1 - \lambda_+)} \lambda_+^n \right. \\ \quad \left. + \frac{\alpha_{k_2}(\lambda_+ + I(pf + k_2) - 1) + \alpha_{k_2-k_1}\alpha_{k_1}I(pf + k_1)}{(\lambda_+ - \lambda_-)(1 - \lambda_-)} \lambda_-^n \right], \end{array} \right. \quad (2.136)$$

avec $1, \lambda_+, \lambda_-$ les trois valeurs propres de la matrice de transition :

$$\lambda_{\pm} = 1 - \frac{I(pf + k_1) + I(pf + k_2)}{2} \pm \frac{1}{2} \sqrt{(I(pf + k_1) - I(pf + k_2))^2 + 4I(pf + k_1)I(pf + k_2) \cos^2(\pi\alpha)}. \quad (2.137)$$

Selon le modèle, on obtient donc que $\mathcal{S}e_{pf}$ est la somme de deux à trois composantes avec les coefficients calculés précédemment. Pour obtenir $\mathcal{S}x_2$, il suffit alors de faire la somme de ces composantes pour chaque harmonique :

$$\mathcal{S}x_2(t) = \sum_p c_{2,p} \mathcal{S}e_{pf}(t), \quad (2.138)$$

avec

$$\mathcal{S}e_{pf} = \begin{cases} a_{pf}e_{pf} + a_{pf+k}e_{pf+k} & \text{si } pf \in [-1, 1], \\ a_{pf}e_{pf} + a_{pf+k_1}e_{pf+k_1} + a_{pf+k_2}e_{pf+k_2} & \text{si } pf \notin [-1, 1]. \end{cases} \quad (2.139)$$

Dans le cas particulier où les extrema de la composante HF sont symétriquement intercalés ($\alpha = 0.5$), le modèle se simplifie considérablement. Les coefficients deviennent :

$$\begin{cases} a_{pf}^{(n)} \simeq (1 - I(pf))^n, \\ a_{pf+k}^{(n)} = 0, \end{cases} \quad \text{si } pf \in [-1, 1] \quad (2.140)$$

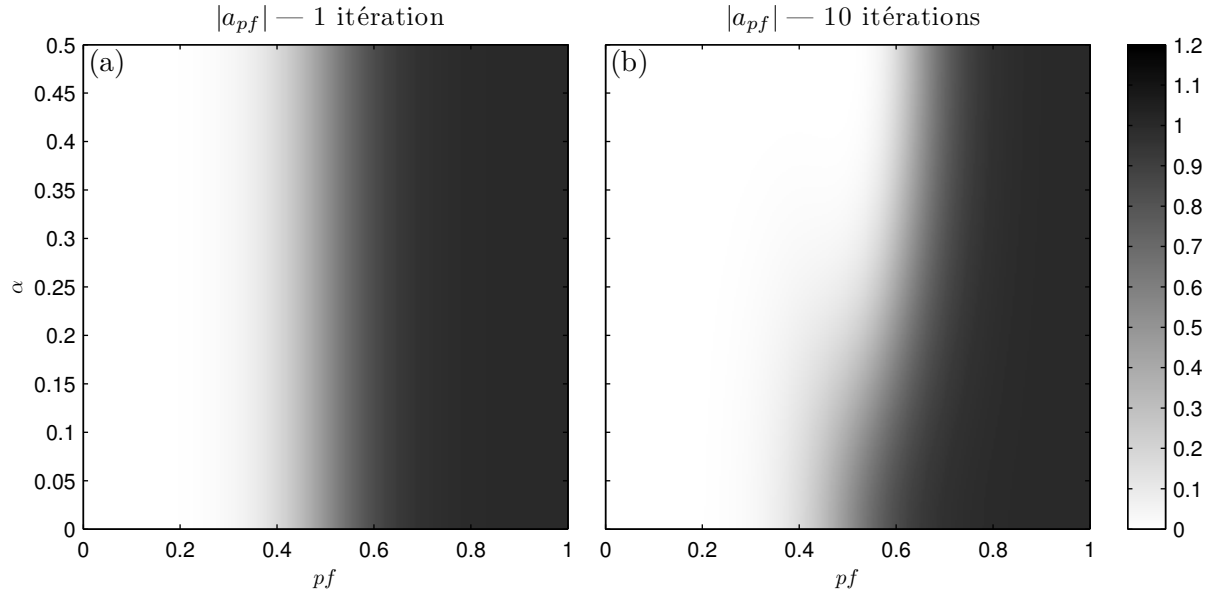


FIGURE 2.27 – Module du coefficient de l'harmonique p (normalisé par son coefficient dans la composante BF, $c_{2,p}$) dans le premier IMF en fonction de α et f pour 1 et 10 itérations de tamisage. Le coefficient étant égal à 1 dès que $pf > 1$, on n'a représenté que la région $pf < 1$. À α fixé, le module du coefficient en fonction de la fréquence pf présente une évolution d'allure sigmoïde. La fréquence de coupure (saut de la sigmoïde) varie de $f_c^{(n)}$ (2.77), (2.79) quand $\alpha = 0.5$ (extrema de HF symétriquement intercalés) à ≈ 0.5 pour $\alpha = 0$ (maxima et minima de HF aux mêmes points).

et si $pf \notin [-1, 1]$:

$$\begin{cases} a_{pf}^{(n)} \simeq 1, \\ a_{pf+k_1}^{(n)} = 0 \\ a_{pf+k_2}^{(n)} \simeq (1 - I(pf + k_2))^n - 1 \end{cases} \quad \begin{array}{l} \text{avec } k_1 \text{ impair tel que } pf + k_1 \in [-1, 1], \\ \text{avec } k_2 \text{ pair tel que } pf + k_2 \in [-1, 1]. \end{array} \quad (2.141)$$

On peut alors exprimer le premier IMF dans les deux cas comme

$$d_1^{(n)}(t) = x(t) - \sum_p (1 - (1 - I(pf + 2k))^n) c_{2,p} e^{2i\pi(pf+2k)t}, \quad (2.142)$$

où $2k$ est l'entier pair tel que $pf + 2k \in [-1, 1]$.

Pour y voir un peu plus clair dans le cas général où les extrema ne sont pas nécessairement symétriquement intercalés, on a représenté graphiquement les modules des différents coefficients en fonction de pf et α pour 1 et 10 itérations de tamisage : Fig. 2.27 pour a_{pf} et Fig. 2.28 pour a_{pf+k_1} et a_{pf+k_2} . On observe ainsi que l'harmonique d'ordre p de la composante BF est capturé par le premier IMF dès lors que la fréquence pf dépasse un certain seuil qui vaut $f_c^{(n)}$ (2.77) lorsque les extrema de la composante HF sont symétriquement intercalés ($\alpha = 0.5$) et ≈ 0.5 dans le cas opposé où les maxima et les minima de HF sont situés aux mêmes endroits ($\alpha = 0$). Ce dernier résultat est simplement dû au fait que,

$$a_{pf}^{(n)} \approx \frac{I(pf - 1)}{I(pf) + I(pf - 1)} \quad (2.143)$$

pour $pf \in [0, 1]$, $\alpha \rightarrow 0$. Dans le cas où $\alpha \neq 0.5$, on constate de plus que dès lors que pf est supérieur à cette fréquence de coupure, les coefficients des nouvelles fréquences apparaissant par repliement

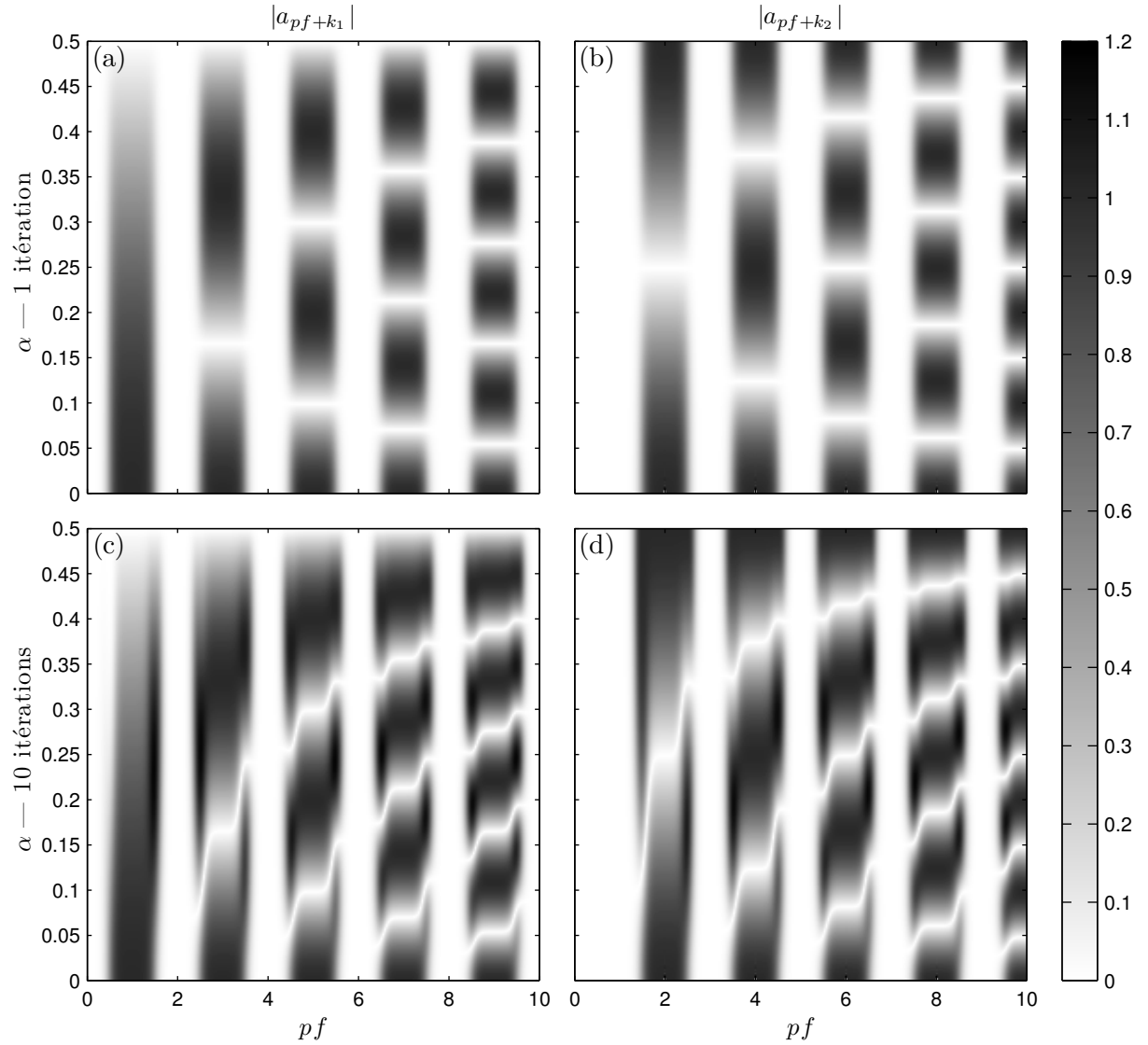


FIGURE 2.28 – Modules des coefficients des nouvelles fréquences dues à l’harmonique p (normalisés par le coefficient de l’harmonique dans la composante BF $c_{2,p}$) dans le premier IMF (ainsi que dans la première approximation $a_1(t) = x(t) - d_1(t)$) en fonction de α et f pour 1 et 10 itérations de tamisage. Colonne de gauche : fréquences $pf + k_1 \in [-1, 1]$ avec k_1 impair. Colonne de droite : fréquences $pf + k_2 \in [-1, 1]$ avec k_2 pair (rien n’est représenté pour $pf < 1$). A de rares exceptions près, dès que $\alpha < 0.5$, il y a toujours au moins un coefficient non nul sur les deux. En revanche si $\alpha = 0.5$ (extrema symétriquement intercalés), a_{pf+k_1} avec k_1 impair est toujours nul et il n’y a donc pas de fréquences dues au repliement si pf est proche d’un entier impair. Cependant, la largeur des bandes autour des entiers impairs dans lesquelles $a_{pf+k_2} \approx 0$ diminue avec le nombre d’itérations.

s'écartent de zéro et il est alors exceptionnel que les coefficients des deux nouvelles fréquences s'annulent simultanément. En revanche, si les extrema sont symétriquement intercalés, le coefficient de la nouvelle fréquence $pf + k_1$ où k_1 est impair est toujours nul et il y a régulièrement des bandes de fréquences où le coefficient de l'autre nouvelle fréquence s'annule également. Toutefois, on peut observer que la largeur de ces bandes diminue avec le nombre d'itérations.

Si l'on prend en compte tous les harmoniques de la composante BF, on voit finalement que les deux composantes $x_1(t)$ et $x_2(t)$ sont correctement séparées si et seulement si tous les harmoniques de BF ont une fréquence inférieure à une certaine fréquence de coupure qui varie de ≈ 0.5 quand $\alpha = 0$ à $f_c^{(n)}$ quand $\alpha = 0.5$. Dans le cas contraire, tout harmonique dont la fréquence est supérieure à cette fréquence coupure est capturé par le premier IMF. De plus, chaque harmonique dans ce cas génère une à deux fréquences supplémentaires avec des coefficients opposés dans le premier IMF et dans la première approximation. De ce fait, il est rare qu'il y ait des valeurs de $f \in [0, 1]$ pour lesquelles l'EMD considère le signal composite comme une unique composante modulée en amplitude et en fréquence.

1.2.2.2 Illustration Pour illustrer le comportement de l'EMD sur les sommes de signaux périodiques non linéaires, on a réalisé des simulations avec les deux formes d'onde :

$$y_1(t) = \cos(2\pi t) + 0.25 \cos(6\pi t) - 0.25, \quad (2.144)$$

$$y_2(t) = \cos(2\pi t) - 0.25 \cos\left(4\pi t + \frac{3\pi}{4}\right) - 0.1404, \quad (2.145)$$

qui sont des IMFs de période 1 avec deux extrema par période. $y_2(t)$ a ses extrema symétriquement intercalés et $y_1(t)$ présente un décalage de ≈ 0.39 entre un maximum et le minimum suivant. On a choisi ici des signaux avec un seul harmonique pour simplifier mais les résultats seraient tout aussi valables avec une structure harmonique plus complexe. Les résultats des simulations pour 10 itérations de tamisage avec $x(t; a, f, \varphi) = y_1(t) + ay_2(ft + \varphi)$ sont présentés figure Fig. 2.29 et ceux avec $x(t; a, f, \varphi) = y_2(t) + ay_1(ft + \varphi)$ Fig. 2.30. Dans les deux cas on ne s'est intéressé qu'à la zone du plan (a, f) où les extrema du signal composite sont proches de ceux de la composante HF.

Tout comme dans le cas de sommes de sinusoïdes, on observe que le comportement de l'EMD semble ne pas dépendre du rapport d'amplitudes a dès lors que celui-ci est soit suffisamment grand, soit suffisamment petit. On peut ainsi diviser le plan (a, f) en trois zones : les deux zones où le comportement de l'EMD est semblable soit au cas $a \rightarrow 0$ soit au cas $a \rightarrow \infty$, et la zone intermédiaire. Tout comme dans le cas sinusoïdal, on constate également que ces trois zones coïncident à peu de chose près avec les 3 zones observées pour la densité d'extrema. Une différence notable en revanche est la largeur de la zone intermédiaire qui est ici bien plus grande que dans le cas sinusoïdal. De plus, la densité d'extrema dans la zone intermédiaire n'est plus intermédiaire entre les densités d'extrema des deux composantes mais peut éventuellement être supérieure.

Si on s'intéresse plus précisément au comportement de l'EMD dans la zone où la densité d'extrema est identique à celle de la composante HF, on peut voir que globalement l'évolution des critères c_1 et $c_{2,2}$ est similaire au cas sinusoïdal mais avec un certain nombre de détails supplémentaires. En revanche le critère $c_{3,2}$ est assez différent du cas sinusoïdal où il était proche de zéro dans la quasi-totalité de la zone $af < 1$. Le fait que sa valeur s'écarte significativement de zéro dès lors que le rapport de fréquence est supérieur à une certaine valeur indique essentiellement que les harmoniques et le fondamental de la composante BF ne sont pas traités de la même manière par l'EMD. Comme on peut le vérifier à l'aide de la figure Fig. 2.31, la marche située à $f \approx 0.3$ dans le comportement de c_1 Fig. 2.29 et celle à $f \approx 0.2$ Fig. 2.30 correspondent respectivement au passage de l'harmonique 2 de $y_2(t)$ et de l'harmonique 3 de $y_1(t)$ dans le premier IMF. Le fait que $c_{2,2}$ soit différent de zéro pour $f \rightarrow 1$ s'explique par la présence de nouvelles composantes à $2f - 2$ dans le cas de la figure Fig. 2.29 et à $f - 1$, $3f - 2$ et $3f - 3$ dans le cas de la figure Fig. 2.30. Enfin, la bosse mystérieuse aux alentours

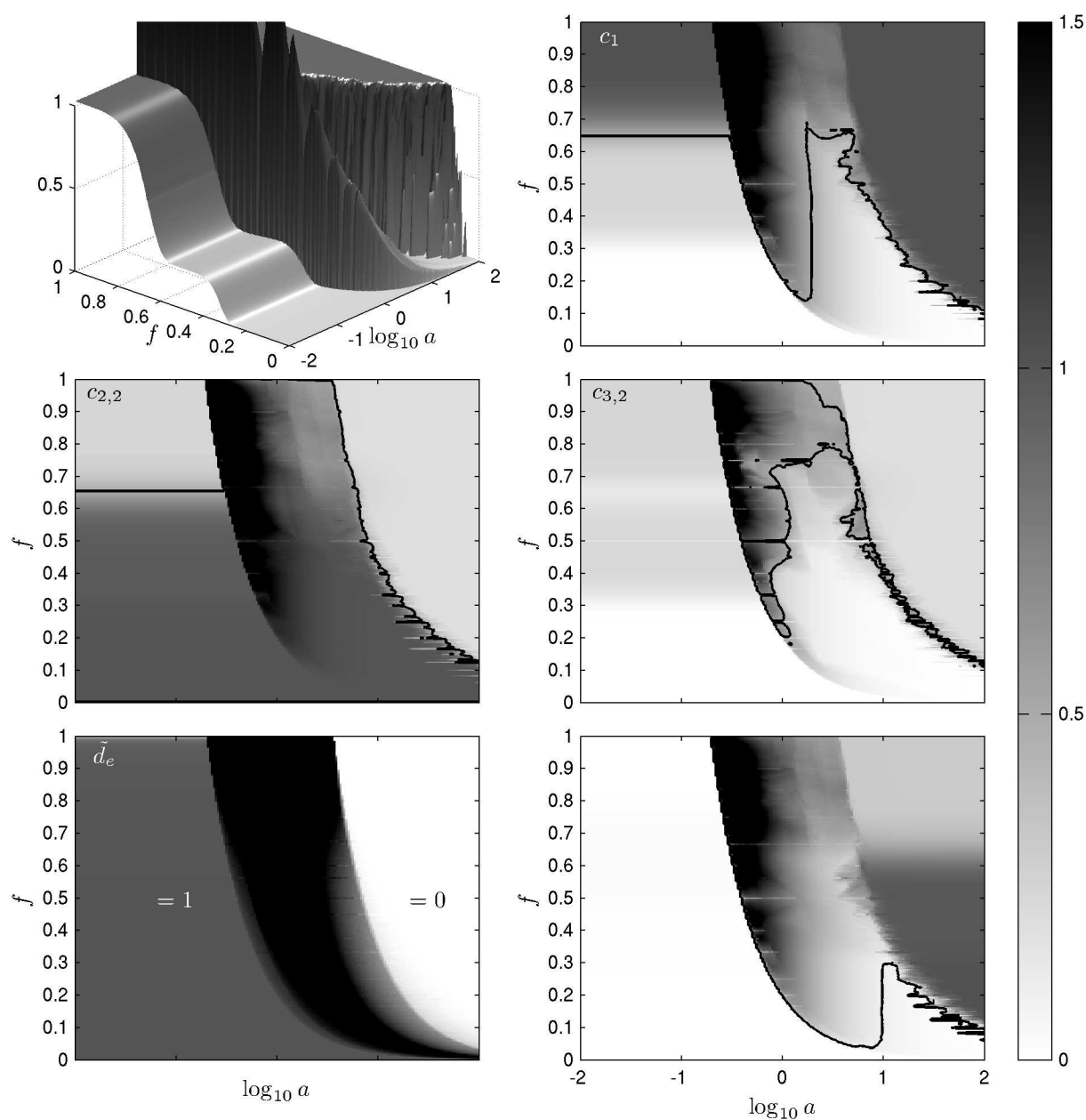


FIGURE 2.29 – Caractérisation du comportement de l'EMD sur $x(t; a, f, \varphi) = y_1(t) + ay_2(ft + \varphi)$ pour 10 itérations de tamisage. 1^{re} ligne : c_1 en 3D et en image. 2^e ligne : $c_{2,2}$ et $c_{3,2}$. 3^e ligne : densité d'extrema réduite $\tilde{d}_e$ (2.12) et écart au modèle normalisé par la composante BF. Les lignes épaisses noires représentent les lignes de niveau = 0.5 sauf pour l'écart au modèle, où le niveau est 0.1.

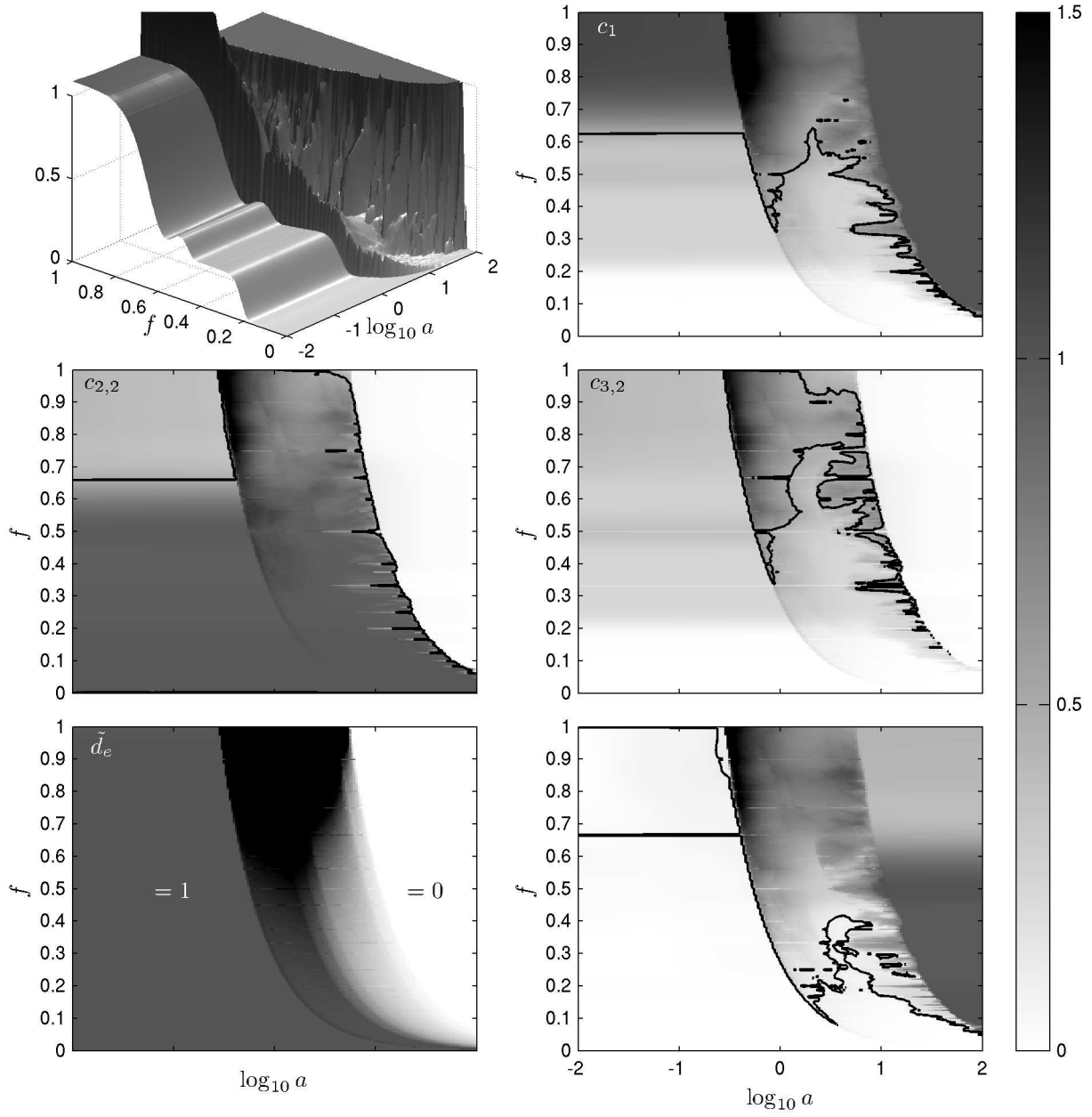


FIGURE 2.30 – Caractérisation du comportement de l'EMD sur $x(t; a, f, \varphi) = y_2(t) + ay_1(ft + \varphi)$ pour 10 itérations de tamisage. 1^{re} ligne : c_1 en 3D et en image. 2^e ligne : $c_{2,2}$ et $c_{3,2}$. 3^e ligne : densité d'extrema réduite $\tilde{d}_e$ (2.12) et écart au modèle normalisé par la composante BF. Les lignes épaisses noires représentent les lignes de niveau = 0.5 sauf pour l'écart au modèle, où le niveau est 0.1.

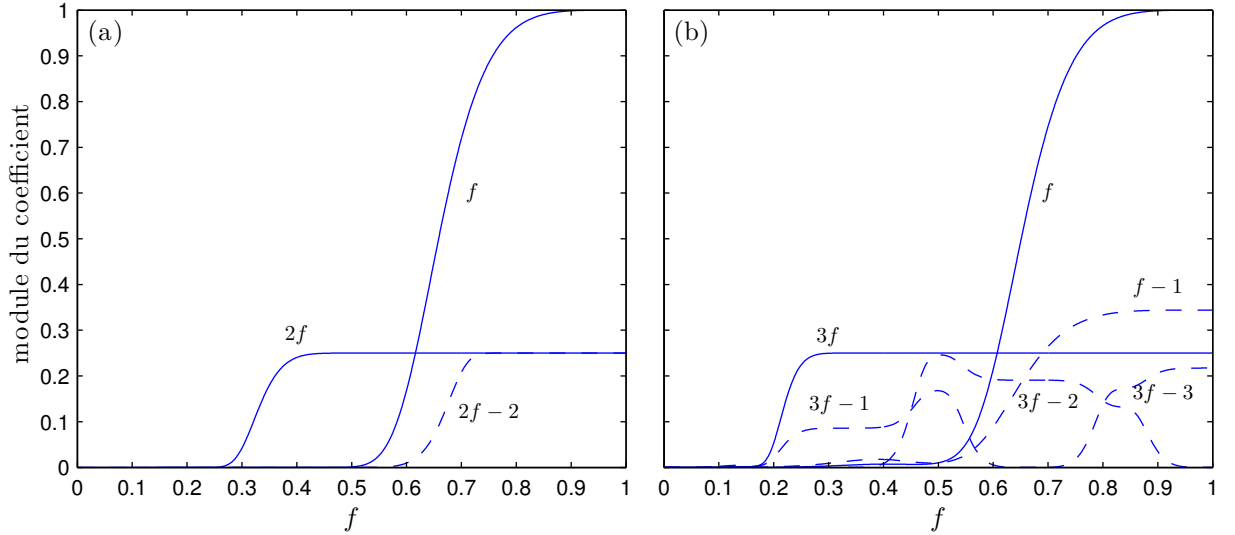


FIGURE 2.31 – Modules des coefficients dans le premier IMF du fondamental ou des harmoniques de la composante BF (trait plein) et des nouvelles composantes apparaissant du fait du repliement spectral (tirets). À gauche : $x(t; a, f, \varphi) = y_1(t) + ay_2(ft + \varphi)$. À droite : $x(t; a, f, \varphi) = y_2(t) + ay_1(ft + \varphi)$.

de $f \approx 0.5$ dans l'évolution de c_1 figure Fig. 2.30 s'explique par un sursaut d'amplitude de deux des nouvelles composantes apparaissant par repliement.

1.2.2.3 Conclusions L'étude précédente constitue une première approche de l'évaluation des capacités de l'EMD à décomposer des sommes de signaux non linéaires. Dans le cas particulier des sommes de signaux périodiques faiblement non linéaires étudiées, les résultats obtenus permettent de déterminer les conditions nécessaires sur le rapport de fréquence pour séparer correctement deux ondes : celui-ci doit être tel que les fréquences de tous les harmoniques de la composante BF soient inférieures à environ $f_c^{(n)}$ fois la fréquence de la composante HF si les extrema de cette dernière sont symétriquement intercalés et la moitié de la fréquence de la composante HF, si ce n'est pas le cas.

Dans le meilleur des cas, on voit donc que les performances de l'EMD ne sont pas meilleures que celles d'un filtre linéaire dont la fréquence de coupure aurait été ajustée, soit autour de $f_c^{(n)}$ fois la fréquence de la composante HF, soit autour de la moitié de cette fréquence. En outre, il n'est a priori pas évident de trouver un avantage à la décomposition complexe proposée par l'EMD lorsque les deux composantes ne sont pas correctement séparées par rapport au résultat d'un simple filtre linéaire. Par ailleurs, un filtre linéaire ajusté pourrait, contrairement à l'EMD, séparer les deux composantes quel que soit le rapport de leurs amplitudes.

Finalement, pour des signaux périodiques, le seul avantage de l'EMD par rapport à un filtre linéaire est que l'EMD ajuste automatiquement sa « fréquence de coupure ». Cet avantage est toutefois très relatif puisque la fréquence de coupure choisie par l'EMD n'est pas non plus totalement libre et pour des nombres d'itérations qui varient par exemple de 1 à 100, elle ne varie que de 0.5 à 0.78 fois la fréquence de la composante HF.

Au final, la grande force de l'EMD reste son aspect adaptatif puisque, si les composantes initiales ne sont plus périodiques mais par exemple modulées en amplitude et en fréquence, elle garde des performances similaires à celles observées sur les signaux périodiques, alors qu'un filtre linéaire est loin d'avoir une telle souplesse.

1.3 Cas d'un mélange à $n > 2$ ondes

Le modèle utilisé jusqu'à présent étant valable dès lors que les extrema sont localement approximativement uniformément espacés, il est tout à fait possible de l'utiliser dans des cas où le signal contient plus de deux composantes. De fait, les composantes non linéaires étudiées dans la section précédente peuvent être vues comme plusieurs composantes sinusoïdales. Ainsi, il est a priori tout à fait possible de prédire le premier IMF d'un signal constitué de plusieurs composantes, qu'on pourra choisir périodiques pour simplifier, à la condition que les extrema du signal soient approximativement uniformément espacés. De là, si les extrema de la première approximation $a_1(t)$ prédite par le modèle sont à nouveau approximativement uniformément espacés, on peut réitérer le processus. Au final, on voit que si le signal a des extrema quasi-uniformément espacés et que toutes les approximations $a_k(t)$ successivement prédites par le modèle ont cette même propriété, on peut a priori prédire toute la décomposition. En pratique cependant, il n'est pas rare de rencontrer des situations où les extrema ne sont pas uniformément espacés même localement. On peut penser par exemple à tous les cas correspondant aux zones de transition des figures Fig. 2.29 et Fig. 2.30. De plus, même dans le cas où les extrema sont uniformément espacés tout au long de la décomposition, il faut compter avec le fait que la prédiction de chaque IMF comporte une certaine imprécision, d'autant plus grande que l'on s'écarte de la situation idéale où les extrema sont exactement uniformément espacés. De ce fait, les imprécisions peuvent s'accumuler au fur et à mesure que des IMFs sont extraits jusqu'à fausser fortement le résultat à partir d'un certain rang.

1.4 Application à l'EMD bivariée

Les algorithmes d'EMD bivariée 6.5.2 étant proches de l'algorithme de l'EMD, on s'est également intéressé à leur comportement sur des sommes de deux exponentielles complexes, pendant à 2 composantes des sinusoïdes. On verra que le modèle proposé précédemment pour l'EMD dans le cas où les extrema du signal sont uniformément espacés peut se généraliser au cas bivarié où il garde globalement les mêmes propriétés que dans le cas classique.

1.4.1 EMD bivariée d'une somme de deux exponentielles complexes

Le modèle de signal qui nous intéresse dans cette section est le suivant

$$x(t; a, f, \varphi) = e^{2i\pi t} + ae^{2i\pi ft + \psi}, \quad (2.146)$$

avec $a \in \mathbb{R}$, $f \in [-1, 1]$ et $\psi \in [0, 2\pi]$. Comme précédemment, on appellera composante haute fréquence (HF) le premier terme de (2.146) et composante basse fréquence (BF) le second.

1.4.2 Modèle dans le cas bivarié

1.4.2.1 Adaptation du modèle Dans la continuité de ce qui a été fait pour l'EMD classique, le modèle dans le cas bivarié est basé sur l'hypothèse que quelle que soit la direction, les maxima du signal projeté sur cette direction sont uniformément espacés. Si on considère ainsi un signal complexe $x(t)$ tel que les maxima de sa projection sur toute direction soient toujours espacés de 1^4 , son armature latérale dans la direction φ_k est décrite par :

$$\hat{e}_{\varphi_k}(\nu) = I(\nu)(\text{III}_k * \hat{x})(\nu), \quad (2.147)$$

4. En toute rigueur, cette condition est très restrictive dans la mesure où seuls des signaux périodiques tels que ceux étudiés en 6.5.5.1 la vérifient. En pratique cependant, on s'intéressera à l'application de ce modèle dans des cas où cette condition n'est vérifiée qu'approximativement et la restriction est donc moins gênante.

où $\mathbb{I}\mathbb{I}_k$ est l'opérateur d'échantillonnage aux points $t = t_k + k'$, $k' \in \mathbb{Z}$ où la projection de $x(t)$ sur la direction φ_k est maximale :

$$\mathbb{I}\mathbb{I}_k(\nu) = \sum_{m \in \mathbb{Z}} e^{-2mi\pi t_k} \delta(\nu - m). \quad (2.148)$$

De là, découle un modèle pour chacun des deux algorithmes Algo. 4 et Algo. 5. Celui-ci est identique au cas classique pour l'algorithme Algo. 4 :

$$(\widehat{S^{B1}x})(\nu) = \hat{x}(\nu) - I(\nu) \mathbb{I}\mathbb{I}_{mean} * \hat{x}(\nu), \quad (2.149)$$

à l'exception de l'opérateur d'échantillonnage :

$$\mathbb{I}\mathbb{I}_{mean}(\nu) = \frac{1}{N} \sum_{k=1}^N \mathbb{I}\mathbb{I}_k(\nu) \quad (2.150)$$

$$= \sum_{m \in \mathbb{Z}} \alpha_m \delta(\nu - m), \quad (2.151)$$

où $\alpha_m = \frac{1}{N} \sum_{k=1}^N e^{-2mi\pi t_k}$.

Pour l'algorithme Algo. 5, le modèle se complique un peu. Sachant que la transformée de Fourier de la composante dans la direction φ (associée au nombre complexe $e^{i\varphi}$) d'un signal complexe $u(t)$ s'exprime sous la forme :

$$\mathfrak{F} [\text{Re} (u(t)e^{-i\varphi})e^{i\varphi}] (\nu) = \frac{1}{2} (\mathfrak{F}u(\nu) + (\mathfrak{F}u)^*(-\nu)e^{2i\varphi}), \quad (2.152)$$

on peut montrer que le modèle pour l'algorithme Algo. 5 est le même que le précédent avec un terme supplémentaire

$$(\widehat{S^{B2}x})(\nu) = \hat{x}(\nu) - I(\nu) ((\mathbb{I}\mathbb{I}_{mean} * \hat{x})(\nu) + (\mathbb{I}\mathbb{I}'_{mean} * (\hat{x})^*)(-\nu)), \quad (2.153)$$

avec $\mathbb{I}\mathbb{I}_{mean}$ tel que défini par (2.151) et

$$\mathbb{I}\mathbb{I}'_{mean}(\nu) = \sum_{m \in \mathbb{Z}} \alpha'_m \delta(\nu - m), \quad \text{où} \quad \alpha'_m = \frac{1}{N} \sum_{k=1}^N e^{i(2m\pi t_k + 2\varphi_k)}. \quad (2.154)$$

Enfin, dans le cas où le nombre de directions N tend vers l'infini, les modèles restent identiques à l'exception des opérateurs d'échantillonnage $\mathbb{I}\mathbb{I}_{mean}$ et $\mathbb{I}\mathbb{I}'_{mean}$ dont les paramètres α_m et α'_m deviennent

$$\alpha_m = \frac{1}{2\pi} \int_0^1 e^{-2i\pi m t} \frac{d\phi}{dt} dt, \quad (2.155)$$

$$\alpha'_m = \frac{1}{2\pi} \int_0^1 e^{i(2\pi m t + 2\phi(t))} \frac{d\phi}{dt} dt, \quad (2.156)$$

où ϕ est la fonction bijective de $[0, 1] \rightarrow [0, 2\pi]$ qui à $t_0 \in [0, 1]$ associe la direction $\phi(t_0)$ telle que la projection de $x(t)$ sur cette direction soit maximale en t_0 .

1.4.2.2 Application au cas d'une somme de deux exponentielles complexes Considérons une somme d'exponentielles complexes (2.146). Étant donné que la projection d'un tel signal sur une direction quelconque est un signal à deux composantes sinusoïdales tel que ceux étudiés en 1.1, on peut affirmer que les extrema de cette projection tendent vers ceux de la projection de la composante HF si $a \rightarrow 0$ et de la composante BF si $a \rightarrow \infty$.

Cas $a \rightarrow 0$: Les extrema de $e^{2i\pi t}$ projeté sur la direction $e^{i\varphi}$ sont situés en $t = \varphi/(2\pi) + k$, $k \in \mathbb{Z}$ ce qui implique $\phi(t) = 2\pi t$. De là, on obtient

$$\alpha_m = \delta_{m,0}, \quad (2.157)$$

$$\alpha'_m = \delta_{m,-2}, \quad (2.158)$$

ce qui implique que les peignes de Dirac traduisant l'échantillonnage dans les modèles sont considérablement simplifiés en $\text{III}_{mean} = \delta(\nu)$ et $\text{III}'_{mean} = \delta(\nu + 2)$. On observe en fait le même type de simplification que dans le cas classique lorsque maxima et minima sont symétriquement intercalés où tout se passe comme si on échantillonnait le signal à une fréquence double de celle à laquelle chaque enveloppe est échantillonnée. Comme on a ici une infinité d'enveloppes, tout se passe comme si le signal était échantillonné à un taux infini et donc les effets d'échantillonnage disparaissent. Malheureusement, il faut noter que ce raisonnement ne marche que quand on peut supposer que $\phi(t) = 2\pi t$, c'est-à-dire uniquement dans le cas où les extrema généralisés sont situés aux mêmes endroits que ceux d'une exponentielle complexe, ce qui en toute rigueur n'est le cas que quand le signal est exactement une exponentielle complexe à des constantes additives et multiplicatives près. Il ne faut donc pas espérer que les effets de sous-échantillonnage disparaissent dans le cas général.

Les modèles deviennent alors :

$$(\widehat{SB^1x})(\nu) = (1 - I(\nu))\hat{x}(\nu), \quad (2.159)$$

$$(\widehat{SB^2x})(\nu) = \hat{x}(\nu) - I(\nu)(\hat{x}(\nu) + (\delta(\nu + 2) * (\hat{x})^*)(-\nu)). \quad (2.160)$$

Si l'on tient compte de plus du fait que le spectre du signal $x(t)$ est nul en dehors de $[-1, 1]$, on peut alors voir que le terme supplémentaire dans le modèle correspondant au deuxième algorithme est en fait annulé ou presque par le filtre d'interpolation. Si on fait cette dernière approximation, on obtient finalement que les deux algorithmes suivent le même modèle :

$$(\widehat{SB^1x})(\nu) = (1 - I(\nu))\hat{x}(\nu), \quad (2.161)$$

$$(\widehat{SB^2x})(\nu) \simeq (1 - I(\nu))\hat{x}(\nu). \quad (2.162)$$

On voit ainsi que les deux algorithmes seraient équivalents dans le cas $a \rightarrow 0$ à un filtre linéaire de réponse en fréquence $1 - I(\nu)$ tout comme dans le cas de l'EMD classique d'une somme de deux sinusoïdes lorsque $a \rightarrow 0$. On en déduit l'expression du premier IMF pour les deux algorithmes :

$$d_1^{B1,(n)}(t) = e^{2i\pi t} + a(1 - I(f))^n e^{2i\pi ft + \psi}, \quad (2.163)$$

$$d_1^{B2,(n)}(t) \simeq e^{2i\pi t} + a(1 - I(f))^n e^{2i\pi ft + \psi}, \quad (2.164)$$

Cas $a \rightarrow \infty$: Les extrema de $e^{i(2\pi ft + \psi)}$ projeté sur la direction $e^{i\varphi}$ sont situés en $t = (\varphi - \psi)/(2\pi f) + k$, $k \in \mathbb{Z}$. Les modèles deviennent alors :

$$(\widehat{SB^1x})(\nu) = \hat{x}(\nu) - I\left(\frac{\nu}{f}\right)(\text{III}_{mean} * \hat{x})(\nu), \quad (2.165)$$

$$(\widehat{SB^2x})(\nu) = \hat{x}(\nu) - I\left(\frac{\nu}{f}\right)((\text{III}_{mean} * \hat{x})(\nu) + (\text{III}'_{mean} * (\hat{x})^*)(-\nu)), \quad (2.166)$$

avec

$$\text{III}_{mean}(\nu) = \sum_{m \in \mathbb{Z}} \alpha_m \delta(\nu - mf) \quad \text{où } \alpha_m = \frac{e^{2im\psi}}{2\pi} \int_0^{\frac{1}{f}} e^{-2i\pi mft} \frac{d\phi}{dt} dt \quad (2.167)$$

$$\text{III}'_{mean}(\nu) = \sum_{m \in \mathbb{Z}} \alpha'_m \delta(\nu - mf) \quad \text{où } \alpha'_m = \frac{e^{-2im\psi}}{2\pi} \int_0^{\frac{1}{f}} e^{i(2\pi mft + 2\phi(t))} \frac{d\phi}{dt} dt, \quad (2.168)$$

Sachant que de plus, on peut utiliser $\phi(t) = 2\pi ft + \psi$, on a $\alpha_m = \delta_{m,0}$ et $\alpha'_m = e^{2i\psi}\delta_{m,-2}$ et les modèles deviennent :

$$(\widehat{S^{B1}x})(\nu) = \left(1 - I\left(\frac{\nu}{f}\right)\right) \hat{x}(\nu), \quad (2.169)$$

$$(\widehat{S^{B2}x})(\nu) = \hat{x}(\nu) - I\left(\frac{\nu}{f}\right) (\hat{x}(\nu) + e^{2i\psi}(\delta(\nu + 2f) * (\hat{x})^*)(-\nu)). \quad (2.170)$$

Si l'on ajoute maintenant l'approximation $I(\nu) = 0$ si $\nu \notin [-1, 1]$, ces modèles se simplifient en

$$(\widehat{S^{B1}x})(\nu) \simeq \hat{x}(\nu), \quad (2.171)$$

$$(\widehat{S^{B2}x})(\nu) \simeq \hat{x}(\nu) - I\left(\frac{\nu}{f}\right) e^{2i\psi}(\delta(\nu + 2f) * (\hat{x})^*)(-\nu). \quad (2.172)$$

Selon ces dernières expressions, l'algorithme Algo. 4 devrait considérer le signal comme une unique composante. Pour ce qui est de l'algorithme Algo. 5, on peut voir dans un premier temps que celui-ci laisse la composante BF intacte :

$$(\widehat{S^{B2}x_2})(\nu) \simeq e^{i\psi}\delta(\nu - f) - I\left(\frac{\nu}{f}\right) (e^{i\psi}\delta(\nu - f)) \quad (2.173)$$

$$\simeq 0, \quad (2.174)$$

puisque $I(1) = 0$. Concernant la composante HF, on obtient

$$(\widehat{S^{B2}x_1})(\nu) \simeq \delta(\nu - 1) - I\left(\frac{\nu}{f}\right) (e^{2i\psi}\delta(\nu - 2f + 1)), \quad (2.175)$$

qui s'apparente à l'expression obtenue pour les sommes de sinusoides lorsque les extrema sont proches de ceux de la composante BF. De là, on peut exprimer le premier IMF après un nombre quelconque d'itérations de tamisage :

$$d_1^{B2,(n)}(t) \simeq e^{2i\pi t} + ae^{2i\pi ft + \psi} - \left(1 - \left(1 - I\left(2 - \frac{1}{f}\right)\right)^n\right) e^{2i\pi(2f-1)t + 2\psi}. \quad (2.176)$$

Selon cette expression, on voit que l'algorithme Algo. 5 ajoute une nouvelle composante basse fréquence au mélange dès que $|2 - 1/f| < f_c^{(n)}$, c'est-à-dire $(2 + f_c^{(n)})^{-1} < f < (2 - f_c^{(n)})^{-1}$, et le considère comme une unique composante sinon.

1.4.2.3 Confrontation avec l'expérience Comme on l'a fait pour les sommes de sinusoides, on peut analyser le comportement de l'EMD sur les sommes d'exponentielles complexes à l'aide des critères c_1 , $c_{2,1}$, $c_{2,2}$, $c_{3,1}$ et $c_{3,2}$. Les résultats sont présentés Fig. 2.32, Fig. 2.34 et Fig. 2.36 pour l'algorithme Algo. 4 et Fig. 2.33, Fig. 2.35 et Fig. 2.37 pour l'algorithme Algo. 5.

Résultats On peut d'emblée constater que les diagrammes correspondant aux deux algorithmes sont très similaires et qu'ils présentent de fortes ressemblances avec le cas des sommes de sinusoides. Ces ressemblances avec le cas sinusoidal s'expliquent simplement par le fait que les algorithmes bivariés utilisent tous les deux les extrema des projections du signal et que la projection d'une somme d'exponentielles complexes n'est autre qu'une somme de sinusoides de mêmes rapports d'amplitudes et de fréquences. Ainsi, les diagrammes se découpent en un certain nombre de grandes zones, délimitées par les quatre courbes $a|f| = 1$ (remplacé par $af \sin(3\pi f/2) = 1$ pour $f < 1/3$) et $af^2 = 1$, qui sont l'écho des 3 grandes zones observées pour les sommes de sinusoides. La deuxième conséquence importante du fait de s'appuyer sur des projections du signal complexe est que les algorithmes se

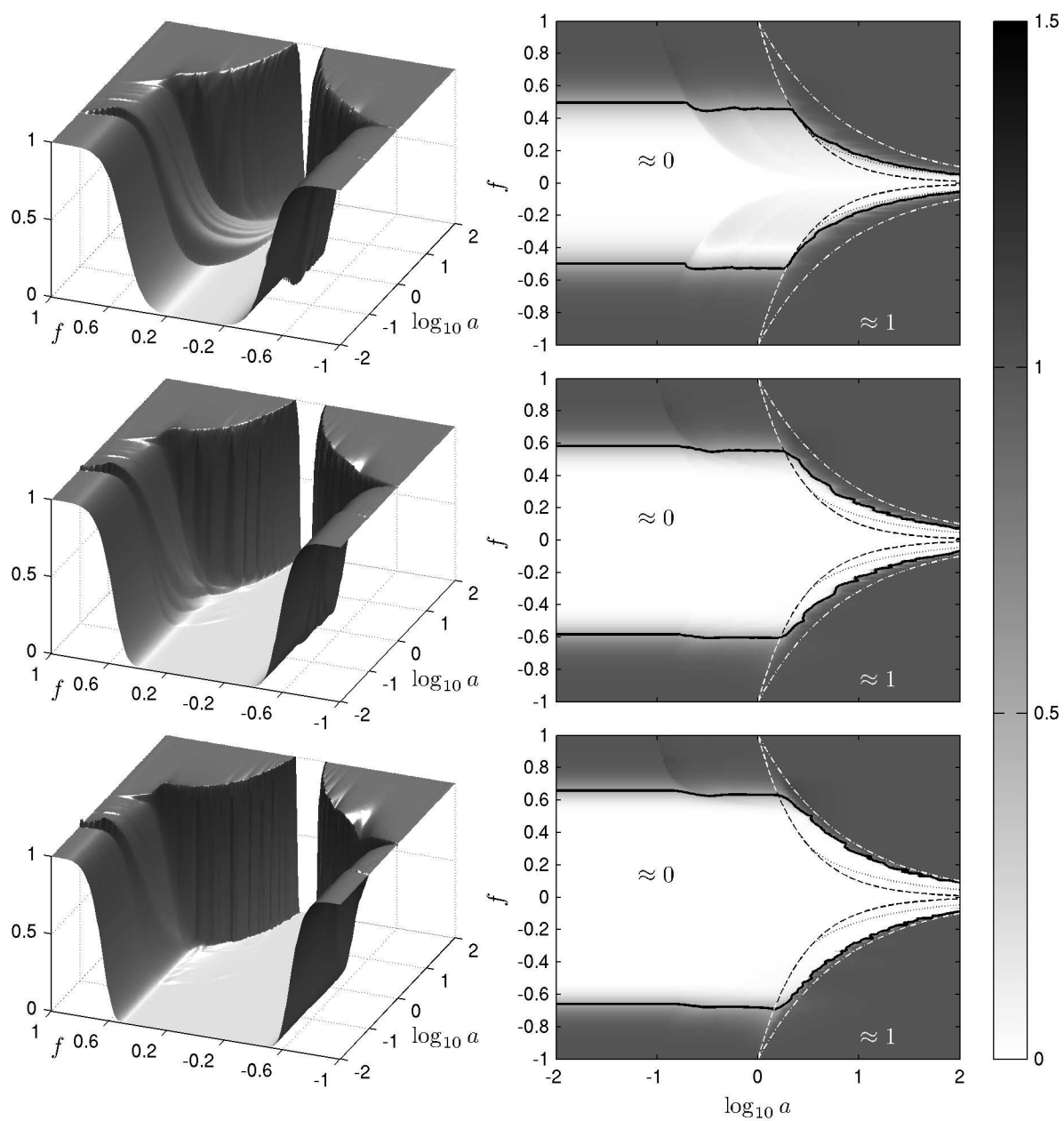


FIGURE 2.32 – Algo. 4. Colonne de gauche : vue 3D de c_1 moyenné par rapport à φ pour (de haut en bas) 1, 3 et 10 itérations de tamisage. Colonne de droite : idem en images.

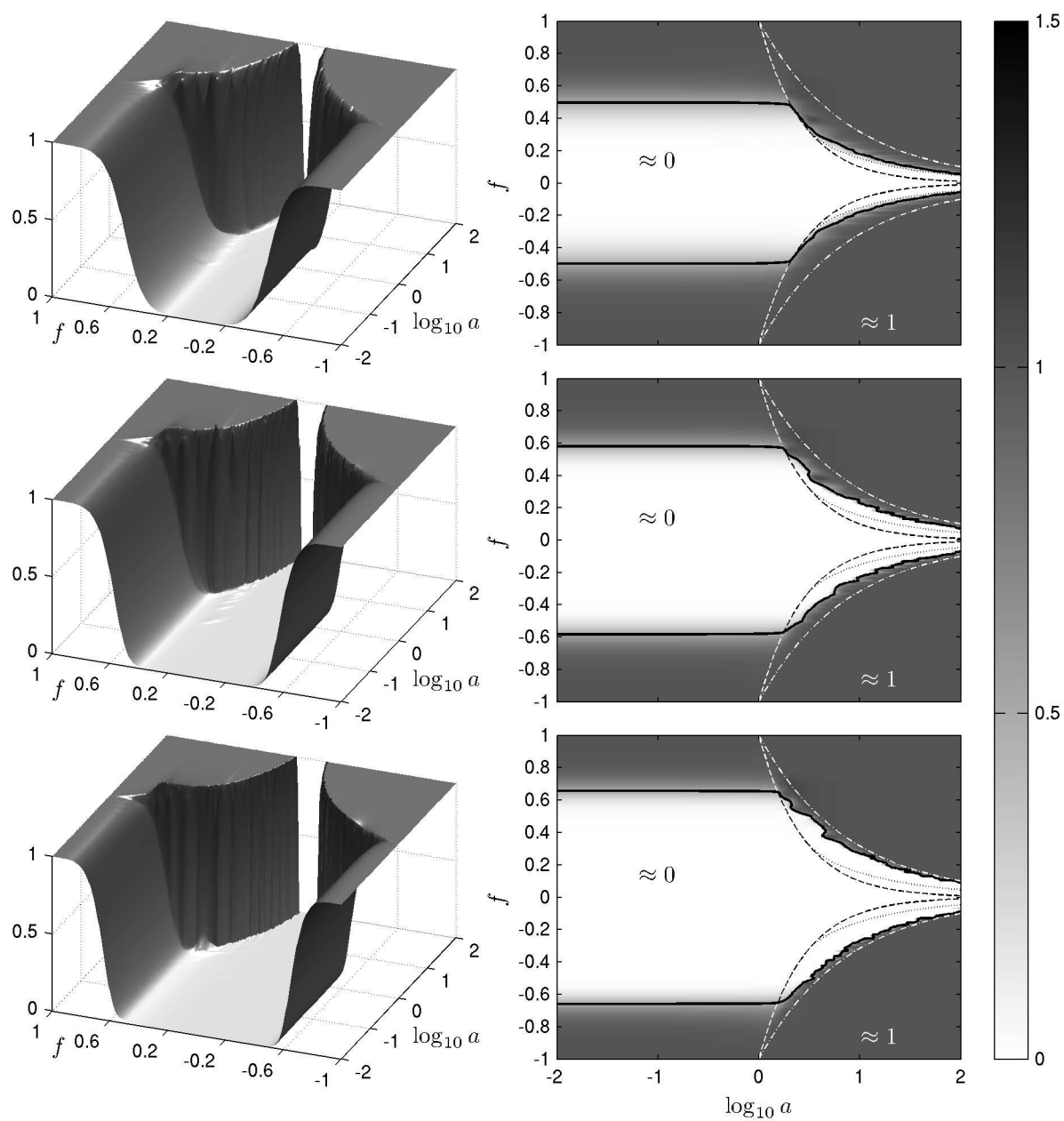


FIGURE 2.33 – Algo. 5. Colonne de gauche : vue 3D de c_1 moyenné par rapport à φ pour (de haut en bas) 1, 3 et 10 itérations de tamisage. Colonne de droite : idem en images.

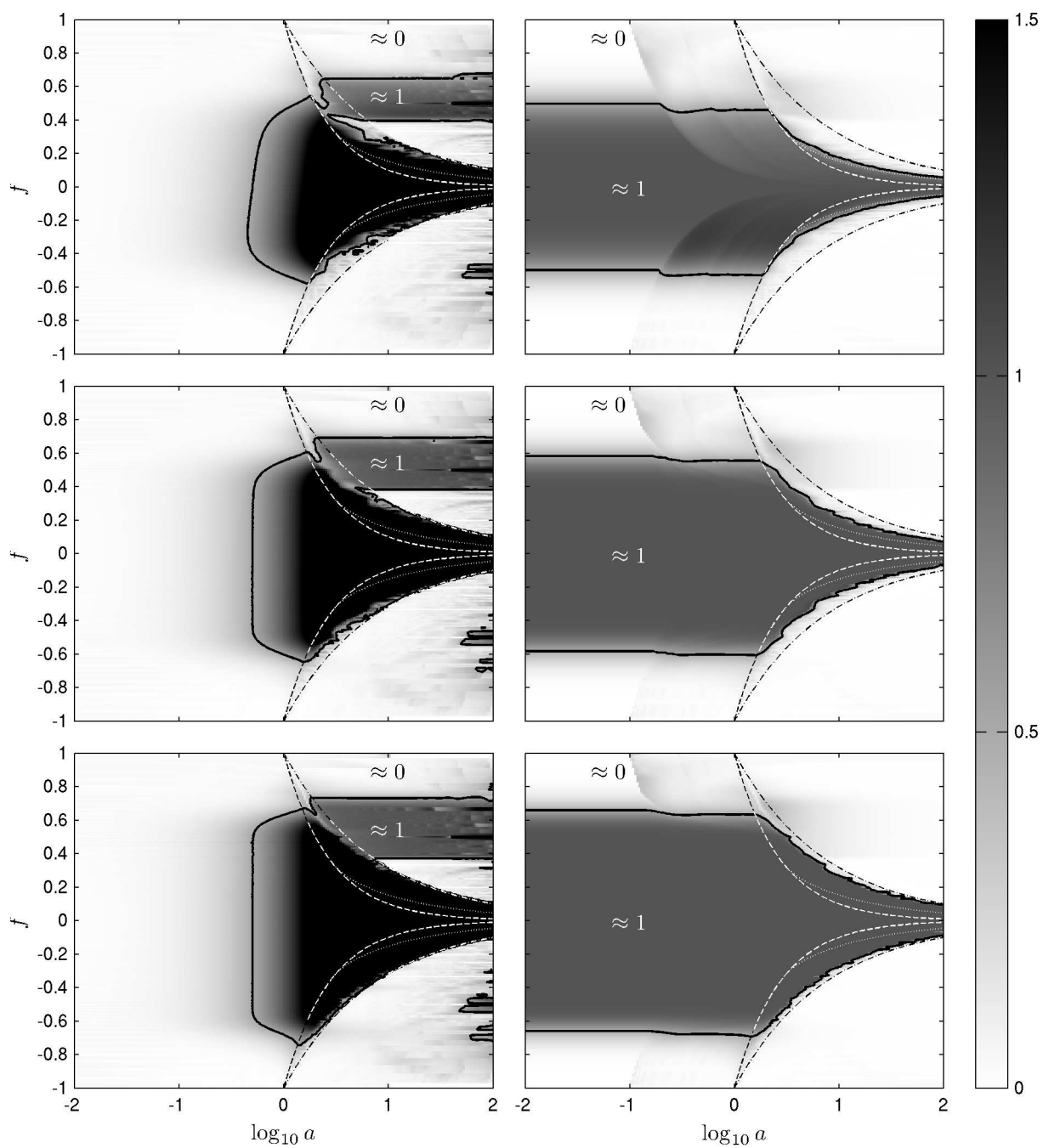


FIGURE 2.34 – Algo. 4. Colonne de gauche : $c_{2,1}$ moyenné par rapport à φ pour (de haut en bas) 1, 3 et 10 itérations de tamisage. Colonne de droite : idem pour $c_{2,2}$.

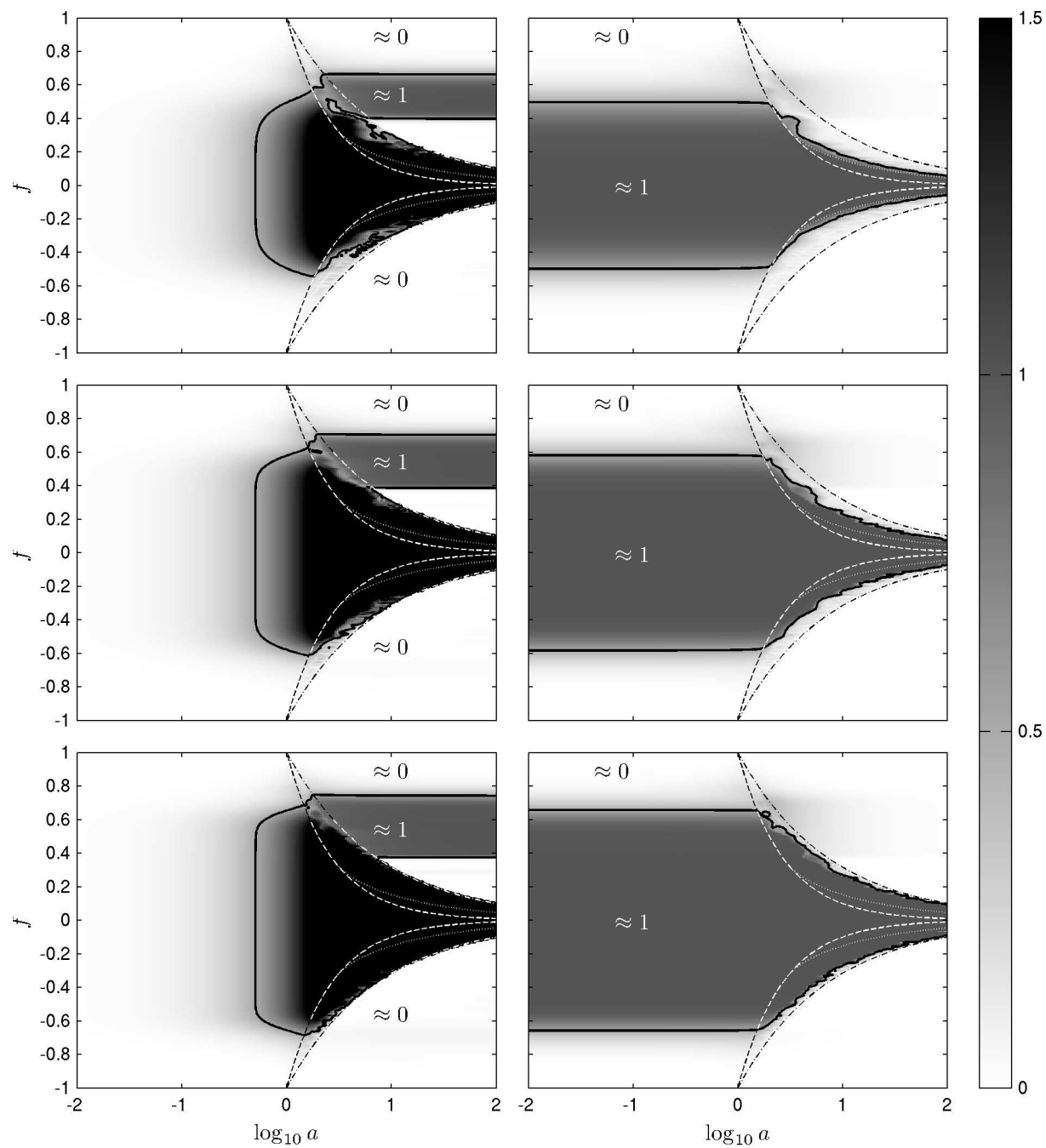


FIGURE 2.35 – Algo. 5. Colonne de gauche : $c_{2,1}$ moyenné par rapport à φ pour (de haut en bas) 1, 3 et 10 itérations de tamisage. Colonne de droite : idem pour $c_{2,2}$.

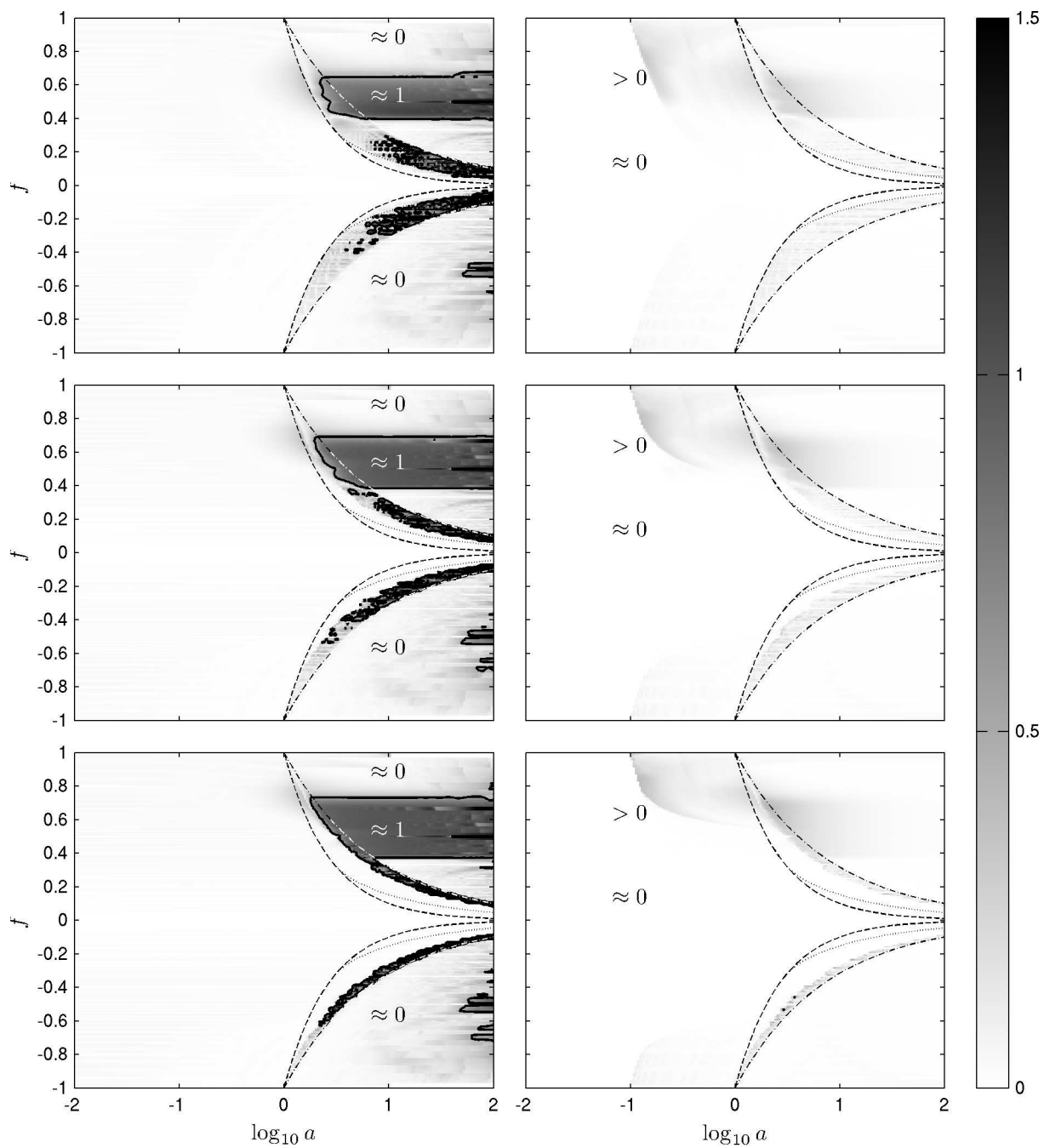


FIGURE 2.36 – Algo. 4. Colonne de gauche : $c_{3,1}$ moyenné par rapport à φ pour (de haut en bas) 1, 3 et 10 itérations de tamisage. Colonne de droite : idem pour $c_{3,2}$.

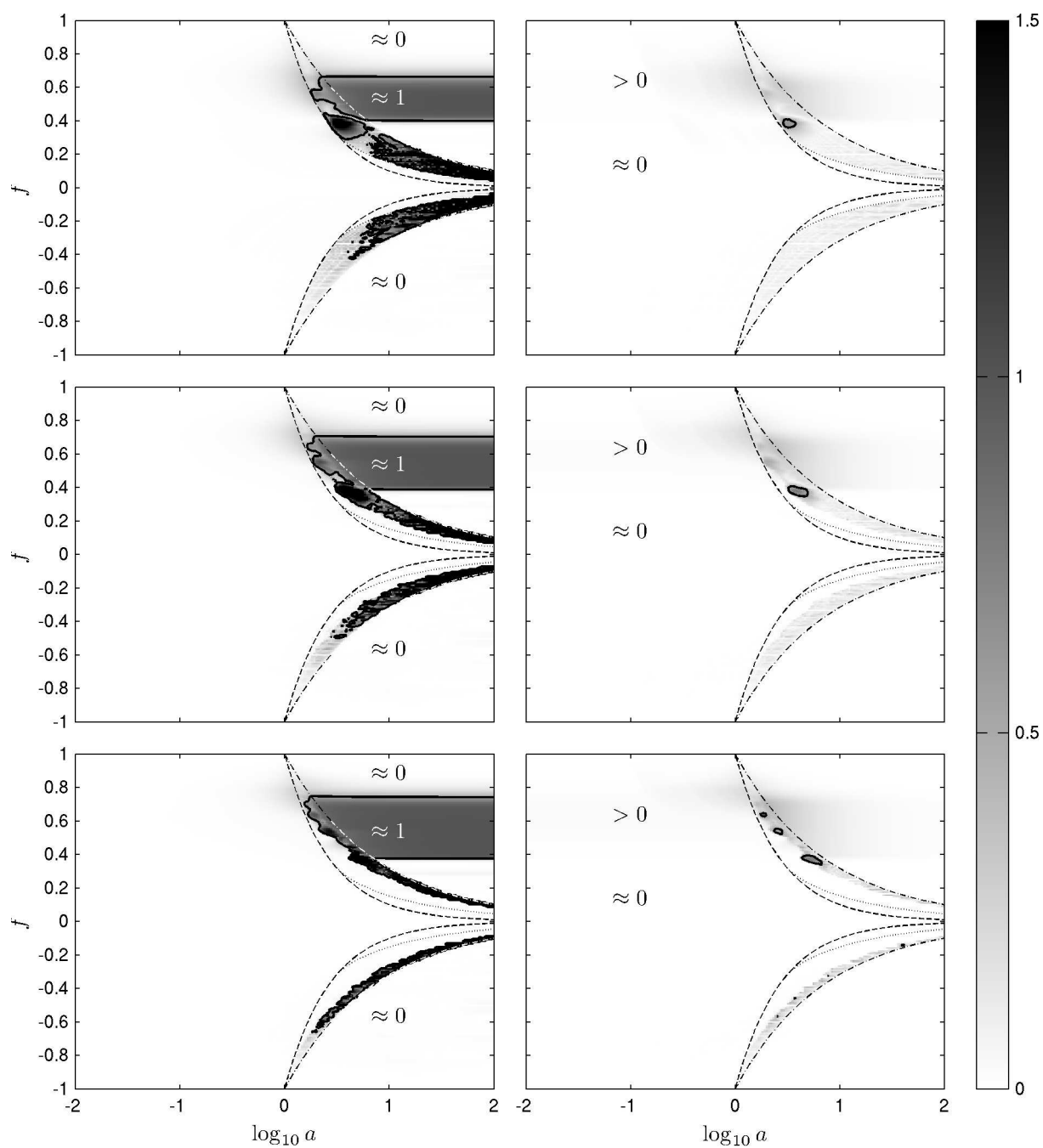


FIGURE 2.37 – Algo. 5. Colonne de gauche : $c_{3,1}$ moyenné par rapport à φ pour (de haut en bas) 1, 3 et 10 itérations de tamisage. Colonne de droite : idem pour $c_{3,2}$. $c_{3,2}$ s'écarte très légèrement de 0 (jusque 0.03 pour 10 itérations) pour f aux alentours de 0.6, et ce même quand $a \rightarrow 0$. C'est en fait la trace de l'approximation utilisée dans le modèle qui nous a permis de négliger la composante à la fréquence $2 - f$.

révèlent incapables de séparer des fréquences de signes opposés mais proches en valeur absolue, simplement parce que la différence de signe disparaît lorsque le signal est projeté.

Par la suite, on s'intéresse plus particulièrement au comportement de l'EMD au sein de chacune des grandes zones.

zone $a|f| < 1$, jusqu'aux courbes en pointillés : Le comportement des deux algorithmes dans cette zone est très proche de celui observé pour l'EMD des sommes de sinusoides : l'EMD agit comme un filtrage linéaire de réponse en fréquence $(1 - I(f))^n$, comme le prévoit le modèle. Le seul écart à ce modèle qu'on peut observer sur les diagrammes de la figure Fig. 2.32 est la petite vague qui suit approximativement les courbes $a|f| = 0.1$. Cet effet est en fait dû à l'échantillonnage (la vague se déplace si on change la fréquence d'échantillonnage) auquel l'algorithme Algo. 4 est malheureusement très sensible.

zone $af^2 > 1$: Le comportement des deux algorithmes diffère assez nettement de celui de l'EMD classique puisque, à l'exception d'une bande de rapports $(0.35 \lesssim f \lesssim 0.75$ pour 10 itérations de tamisage), les algorithmes bivariés considèrent tous deux que le signal est un unique IMF. De plus, il se trouve que la nouvelle composante apparaissant dans cette bande de rapports présente les mêmes caractéristiques pour les deux algorithmes. Si ceci était prévu par le modèle pour l'algorithme Algo. 5, ce n'est pas le cas de l'algorithme Algo. 4, pour lequel le modèle prévoit que le signal est systématiquement un unique IMF quel que soit le rapport de fréquences. L'explication de ce phénomène est en fait en lien avec la sensibilité à l'échantillonnage de cet algorithme. Le problème est que l'approximation des extrema uniformément espacés est bien moins pertinente pour cet algorithme puisqu'il est beaucoup plus sensible à la position précise de chaque extremum.

Pour percevoir où se situe l'inadéquation entre le modèle et la réalité, on peut par exemple s'intéresser à la partie imaginaire du signal

$$x(t) = e^{2i\pi t} + ae^{2i\pi ft}, \quad (2.177)$$

avec $f > 1$ et $af^2 < 1$, aux emplacements des maxima t_k de sa partie réelle. Le modèle fournit $t_k = k \in \mathbb{Z}$ et donc

$$x(t_k) = 1 + ae^{2i\pi fk}, \quad (2.178)$$

ce qui implique en particulier

$$\text{Im}(x(t_k)) = a \sin(2\pi fk). \quad (2.179)$$

Or en réalité, les t_k doivent vérifier l'équation

$$\sin(2\pi t_k) + af \sin(2\pi f t_k) = 0, \quad (2.180)$$

ce qui implique entre autres

$$\text{Im}(x(t_k)) = a(1 - f) \sin(2\pi f t_k), \quad (2.181)$$

qui n'est pas compatible avec l'équation (2.179) à supposer que les extrema soient situés en $t_k = k \in \mathbb{Z}$. En poussant le raisonnement un peu plus loin, on pourrait montrer que le comportement de l'algorithme Algo. 4 est en fait très similaire à l'autre algorithme.

zone de transition : Le comportement dans cette zone s'écarte aussi assez nettement de celui de l'EMD classique. En effet, là où la frontière était très irrégulière pour les sommes de sinusoides, on observe ici une frontière beaucoup plus régulière. L'explication est en fait assez simple sachant que l'irrégularité de la frontière dans le cas classique est liée à des situations particulières de rapports de fréquences de la forme $1/k$ (ou éventuellement certains rapports p/q avec p

et q pas trop grands) pour lesquels la densité d'extrema du signal dépend de la différence de phase entre les deux composantes. Dans le cas bivarié, ces situations n'existent pas vraiment puisqu'un décalage de phase entre les deux composantes peut être assimilé à une rotation du signal accompagnée d'une translation temporelle, et que les algorithmes bivariés sont asymptotiquement covariants par rapport à ces deux transformations quand le nombre de directions utilisé et la durée du signal tendent vers l'infini.

On observe de plus un déplacement de la frontière quand le nombre d'itérations de tamisage augmente. La frontière, initialement collée aux courbes $a|f| = 1$ et $af \sin(3\pi f/2) = 1$, se déplace avec les itérations en direction de la courbe $af^2 = 1$. Cette évolution a été également observée dans le cas des sommes de sinusoides mais la frontière se déplaçait moins rapidement parce qu'elle était « retenue » par les rapports de fréquences particuliers évoqués précédemment.

Conclusions De ces résultats, on peut finalement tirer deux conclusions :

- les deux algorithmes semblent avoir des comportements très similaires sur les sommes d'exponentielles complexes, la différence principale entre les deux étant la précision supérieure de l'algorithme Algo. 5. On verra en 2.3 que leurs comportements peuvent être plus éloignés dans d'autres situations. Dans le cas de sommes d'exponentielles complexes, leur comportement est aussi très proche du comportement de l'EMD classique sur les sommes de sinusoides, sauf dans la région $af^2 > 1$ où, à l'exception d'une bande de rapports de fréquences, le signal est toujours considéré comme un seul IMF bivarié, là où l'EMD classique rajoutait dans beaucoup de cas une composante dérivée de la composante HF par repliement.
- le modèle développé ne modélise pas correctement les résultats de l'algorithme Algo. 4 dans la partie $af^2 > 1$. Même si la modélisation semble correcte dans la partie $af < 1$, aux imprécisions dues à l'échantillonnage près, on peut légitimement douter de la capacité du modèle à s'adapter à des situations autres que les sommes d'exponentielles complexes.

1.5 Conclusions sur l'EMD de sommes de composantes déterministes « simples »

Dans cette partie concernant l'analyse des performances de l'EMD sur les sommes de sinusoides, et plus généralement de signaux faiblement non linéaires périodiques, on a d'une part caractérisé le comportement de l'EMD à l'aide de simulations numériques, et d'autre part développé un modèle permettant de retrouver les résultats observés dans un grand nombre de situations, du moins dans la limite continue où les effets liés à l'échantillonnage peuvent être négligés.

Concernant l'aspect caractérisation, l'étude a été guidée par les trois questions proposées en 1.1.2 et rappelées ici :

1. *Dans quelles conditions l'EMD sépare-t-elle les deux composantes sinusoidales ?*
2. *Dans quelles conditions l'EMD considère-t-elle le signal comme une seule composante modulée en amplitude et en fréquence ?*
3. *Dans quelles conditions l'EMD décompose-t-elle le signal en plusieurs composantes qui ne sont pas simplement des combinaisons linéaires des deux composantes initiales ?*

Les réponses de l'EMD à ces trois questions sont représentées de manière schématique dans la colonne de gauche Fig. ?? pour les trois situations étudiées.

Pour résumer les résultats de l'analyse des performances de l'EMD on a tout d'abord observé que les positions et la densité des extrema avaient un rôle déterminant. Le premier IMF a en effet dans une très grande majorité des cas autant d'extrema que le signal, l'exception correspondant aux cas similaires à ceux abordés en 1.1.6.1 où des inflexions importantes peuvent se transformer en extrema après quelques itérations de tamisage. Dans les trois situations étudiées, on a ainsi pu montrer que le plan (a, f) des rapports d'amplitude et de fréquence se divisait en trois grandes régions (cinq dans le cas complexe) : dans les deux régions extrêmes où le rapport d'amplitudes a tend vers 0

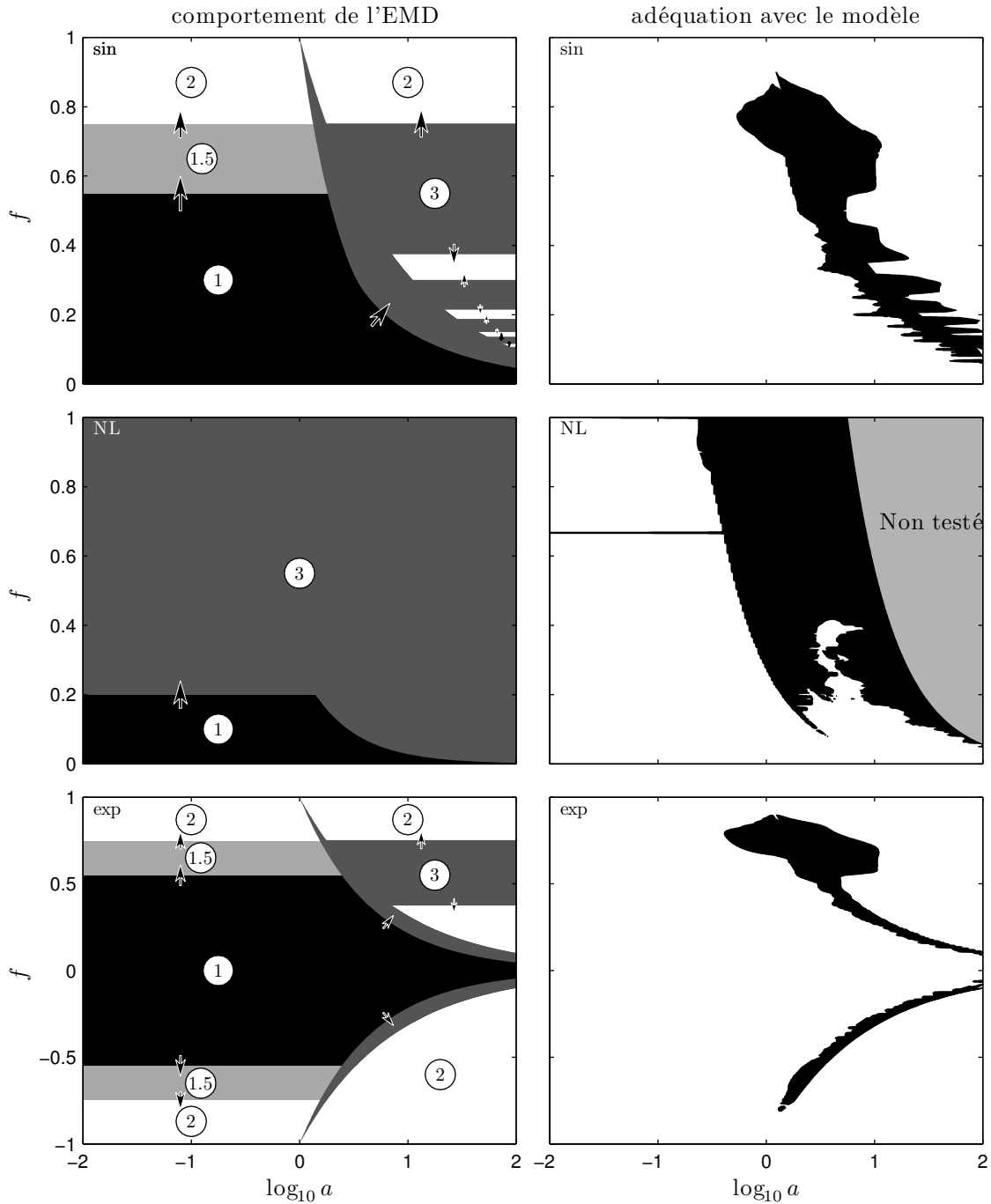


FIGURE 2.38 – Récapitulatif schématique des performances de l'EMD (colonne de gauche) et du modèle (colonne de droite) dans les trois situations : (sin) sommes de sinusoides, (NL) sommes de composantes faiblement non linéaires, (exp) sommes d'exponentielles complexes (algorithme Algo. 5). Les couleurs et étiquettes de la colonne de gauche correspondent aux trois questions posées initialement en 1.1.2 : (1) correspond aux cas où les deux composantes sont correctement séparées; (2) aux cas où l'EMD considère le signal comme une seule composante; (1.5) aux situations intermédiaires entre (1) et (2); (3) aux cas où l'EMD produit des composantes qui ne sont pas simplement des combinaisons linéaires des composantes de départ. Les flèches indiquent schématiquement les déplacements des frontières quand le nombre d'itérations augmente (les schémas correspondent à 10 itérations). Les zones noires de la colonne de droite correspondent aux cas où les prévisions du modèle s'écartent de plus de 0.1 fois la norme de la composante de plus faible amplitude.

ou l'infini, les extrema du signal correspondent approximativement à ceux de la composante de plus grande amplitude ; dans la région intermédiaire, les extrema peuvent provenir de manière générale des deux composantes voire de simples inflexions de ces dernières. Pour cette raison, le comportement de l'EMD dans cette zone transitoire devient nettement plus compliqué à décrire, même si on observe qu'il se ramène parfois à celui observé dans l'une ou l'autre région extrême après quelques itérations de tamisage. Cet aspect est notamment important dans le cas de sommes de composantes faiblement non linéaires dans la mesure où la zone transitoire peut alors être bien plus large et où les composantes sont susceptibles de contenir des inflexions pouvant donner lieu à des extrema dans la somme qui s'ajoutent à ceux provenant directement des extrema des deux composantes.

Concernant maintenant les zones extrêmes où les extrema sont approximativement aux mêmes positions que ceux de la composante de plus grande amplitude, on observe des comportements très différents suivant que ce soit la composante BF ou HF qui domine.

Lorsque la composante HF domine (côté a tend vers 0), on a observé que l'effet d'une itération de tamisage sur des sommes de sinusoides s'apparentait fortement à celui d'un filtre linéaire passe-haut ne dépendant que de la fréquence de la composante HF. De là, le comportement de l'EMD complète est dans ce cas également très proche d'un filtre linéaire correspondant simplement à l'itération du filtre correspondant à une itération. Par conséquent, la séparation ou non des deux composantes sinusoidales par l'EMD dépend uniquement du rapport de leurs fréquences et du nombre d'itérations de tamisage avec une transition progressive entre les deux régimes. Cette modélisation simpliste s'avère très bonne lorsque a tend vers 0 mais se dégrade lorsque a se rapproche de 1. Elle reste cependant une approximation correcte du comportement de l'EMD même pour $a = 1$, avec une marge d'erreur typiquement de 15% dans le pire des cas.

Lorsque l'on considère des sommes de composantes faiblement non linéaires en lieu et place des sinusoides, la description se complique. En pratique, on observe que si la composante BF des harmoniques jusqu'à l'ordre p , l'EMD est capable de séparer totalement les deux composantes uniquement lorsque le rapport de fréquences est inférieur à p fois le rapport de fréquences nécessaire à séparer deux composantes sinusoidales de mêmes fréquences. De plus, lorsque le rapport de fréquences est supérieur à $1/p$, on observe que les IMFs contiennent des fréquences supplémentaires par rapport à celles présentes initialement dans les deux composantes. De fait, il n'est plus possible dans ce cas de décrire le comportement de l'EMD uniquement en termes de filtrage linéaire. Il faut alors faire intervenir une autre caractéristique importante du processus de tamisage qui est l'échantillonnage. En effet, le fait de contruire les enveloppes à partir des extrema du signal fait implicitement intervenir une opération d'échantillonnage (par les extrema) qui s'accompagne parfois d'effets de repliement. Ces effets sont négligeables pour les sommes de sinusoides mais ne le sont plus dès lors que la composante BF n'est pas simplement sinusoidale du fait que ses harmoniques peuvent subir des effets de repliement. En ajoutant ces effets de repliement au modèle de filtrage linéaire précédent, on aboutit à un modèle plus complet de l'EMD qui permet de rendre compte aussi efficacement son comportement sur les sommes de signaux faiblement non linéaires que le simple modèle de filtrage linéaire le permettait sur les sommes de sinusoides.

Une autre situation où les effets de repliement dus à l'échantillonnage sont importants est le cas où la composante BF domine (côté a tend vers l'infini). On observe en effet pour les sommes de sinusoides que le comportement de l'EMD se divise essentiellement en deux régimes suivant le rapport de fréquences : soit elle considère les deux sinusoides comme une seule composante modulée en amplitude et en fréquence, soit elle introduit une nouvelle composante sinusoidale qui ne peut être expliquée autrement que par des effets de repliement. Là encore, le modèle comprenant échantillonnage et filtrage linéaire permet de rendre compte efficacement du comportement de l'EMD.

Le cas correspondant des sommes de composantes faiblement non linéaires n'a pas été abordé dans cette étude mais il est probable que le modèle permette là aussi de rendre compte du comportement de l'EMD. Il n'y aurait alors pas ou que très peu de rapports de fréquences pour lesquels l'EMD

considérerait la somme de signaux comme une seule composante. En revanche, il y aurait davantage de nouvelles fréquences susceptibles d'être introduites dans la décomposition : une par harmonique de la composante HF.

En conclusion, le modèle proposé permet dans de nombreux cas de prévoir le comportement de l'EMD dès lors que le signal est composé de deux composantes périodiques faiblement non linéaires. De plus, l'EMD agissant de manière locale, on a également observé que le modèle pouvait dans une certaine mesure fournir une bonne indication du comportement de l'EMD lorsque les composantes sont modulées en amplitude et en fréquence. Par extension, on peut supposer que le modèle peut s'avérer utile dans des situations plus générales dès lors que localement les extrema du signal sont approximativement uniformément espacés. En particulier, il est sans doute possible de l'appliquer localement dans certains cas se présentant dans les zones de transition du modèle à deux composantes étudié dans ces pages. En revanche, l'adéquation entre les prévisions du modèle et le comportement réel de l'EMD n'est bonne que si les extrema sont effectivement pratiquement uniformément espacés. S'ils ne le sont qu'approximativement, les prévisions du modèle sont généralement également approximatives. Enfin, on s'est borné à étudier les performances du modèle pour des nombres d'itérations de tamisage inférieurs à 10. Il semble probable que pour des nombres d'itérations nettement supérieurs, la qualité de la modélisation soit également réduite.

2 Cadre stochastique : décomposition d'un bruit large bande

En principe, l'EMD a été construite dans la perspective de décomposer des signaux qu'on peut mettre sous la forme de sommes de composantes oscillantes. Cependant, même si l'évolution d'une grandeur physique peut être décrite ainsi, sa mesure comporte inévitablement une certaine incertitude qu'on modélise bien souvent par un bruit large bande. Dès lors, le comportement de l'EMD peut être très différent de celui attendu intuitivement.

Dans le but de clarifier ce comportement, on s'intéresse dans ce chapitre aux caractéristiques de la décomposition proposée par l'EMD quand on lui soumet en entrée un bruit large bande. On étudiera dans un premier temps les cas de bruits blancs de densités différentes, puis on s'intéressera de plus près au cas de bruits gaussiens fractionnaires qui constituent un bon modèle de processus présentant un comportement de loi d'échelle. Une troisième partie est ensuite consacrée à une étude des mêmes types de propriétés pour les EMD bivariées introduites au chapitre 1 en 6.5.2. Enfin, on termine cette étude du comportement de l'EMD sur les bruits large bande par un rapprochement avec le modèle développé dans la partie précédente qui permet de retrouver un certain nombre des caractéristiques observées.

2.1 Bruits blancs de densités variées

On considère dans cette partie quatre modèles de bruits blancs ayant chacun une densité de probabilité marginale différente. L'objectif de cette étude comparative est de tenter de cerner en quoi la densité de probabilité marginale peut influencer la décomposition. On choisit pour ce faire quatre densités de formes différentes qu'on normalise toutes pour qu'elles soient de moyenne nulle et de variance unité. Les quatre modèles de bruits sont les suivants

gaussien : bruit de densité marginale gaussienne

uniforme : bruit de densité uniforme

asymétrique : bruit de densité quadratique asymétrique $p(x) \propto (x - a)^2$ pour $x \in [a, b]$.

bimodal : bruit de densité quadratique bimodale symétrique $p(x) \propto x^2$ pour $x \in [-a, a]$.

2.1.1 Stationnarité

Les IMFs issus de la décomposition d'un signal stationnaire à *tout ordre* sont stationnaires à des effets de bords près. En effet, l'EMD étant covariante par rapport aux décalages temporels, les IMFs issus d'un signal stationnaire à tout ordre de durée infinie (ce qui permet d'ignorer d'éventuels effets de bords) sont nécessairement stationnaires (cf chapitre 1 5.4.1.3 pour une démonstration). Si l'on ajoute à cela que l'EMD agit à une échelle locale, ce qui signifie qu'un IMF donné en un instant donné ne dépend du signal que dans un certain intervalle autour de cet instant, on peut alors en déduire que les IMFs issus d'un signal stationnaire de durée finie ne dépendent des bords du signal qu'à un horizon fini. Par conséquent la durée finie du signal n'a aucune incidence sur la partie centrale des IMFs qui est donc stationnaire. Cette affirmation doit cependant être nuancée en tenant compte du fait que la « localité » de l'EMD n'a pas le même sens pour chaque IMF. Ainsi, si le premier IMF à un instant donné peut ne dépendre du signal que dans un intervalle d'une dizaine de points autour de cet instant, il n'est pas rare que le dernier IMF ait une « période » du même ordre de grandeur que la durée du signal, auquel cas sa valeur à un instant donné dépend alors de toute la durée du signal et en particulier de la manière dont sont traités les effets de bords. En pratique, on pourra toutefois tolérer un peu d'effets de bords dans les simulations dans la mesure où on regarde systématiquement des quantités globales moyennées sur la durée du signal. Ainsi, dans le cas des bruits large bande qui nous intéressent ici, on pourra considérer que les IMFs sont tous à peu près stationnaires à l'exception des 2 derniers. Pour limiter les effets de bords restants, on retranchera généralement une durée fixe sur les bords des IMFs avant de leur appliquer les analyses.

2.1.2 Statistiques d'ordre 1

Les densités de probabilité marginales des IMFs et des approximations ont été estimées pour les différents modèles de bruits blancs à l'aide d'histogrammes. Les résultats sont représentés Fig. 2.39 et Fig. 2.40 pour les IMFs et Fig. 2.41 et Fig. 2.42 pour les approximations.

Pour ce qui concerne les IMFs, on observe qu'à la notable exception du premier IMF, et parfois des second et troisième, les densités marginales sont toutes pratiquement gaussiennes, avec seulement une petite singularité en 0. Cette caractéristique est en particulier bien vérifiée pour les IMFs issus des bruits gaussien et uniforme. Dans le cas du bruit asymétrique et surtout du bruit de densité bimodale, on observe que les premiers IMFs sont nettement moins gaussiens que dans les autres cas, sans doute parce que ces densités marginales sont plus éloignées de la gaussienne que la densité uniforme. De manière générale, on observe donc qu'en dehors du premier IMF les IMFs ont tendance à être de plus en plus gaussiens lorsque leur indice augmente, la convergence vers la gaussienne étant plus ou moins rapide suivant la proximité de la densité marginale du bruit avec la gaussienne. Si on s'intéresse maintenant aux moyennes et variances de ces gaussiennes, on voit très clairement que les moyennes sont toutes pratiquement nulles — elles s'écartent très légèrement de zéro pour la densité asymétrique avec au maximum ≈ 0.002 pour le premier IMF — et que les variances décroissent quand l'indice de l'IMF augmente. Comme on peut le voir Fig. 2.43, la décroissance est très proche d'exponentielle dans tous les cas.

Le cas du premier IMF est particulier puisque c'est le seul à avoir une densité marginale bimodale. On peut montrer en fait très simplement que cet effet est dû au fait que l'écart moyen entre deux extrema est dans le cas du premier IMF proche du pas de discrétisation. Pour s'en convaincre, il suffit d'observer que la densité marginale d'une version « suréchantillonnée » à l'aide d'une interpolation linéaire par morceaux du premier IMF issu du bruit blanc gaussien est tout aussi gaussienne que les densités marginales des IMFs suivants (cf Fig. 2.44). On peut également proposer une justification à l'aide d'une analyse simple des statistiques des extrema dans le premier IMF. Sachant que les bruits considérés sont blancs, on peut montrer très simplement que dans ces bruits, en moyenne 2 points sur 3 sont des extrema locaux. Dans la mesure où des extrema peuvent apparaître, mais rarement

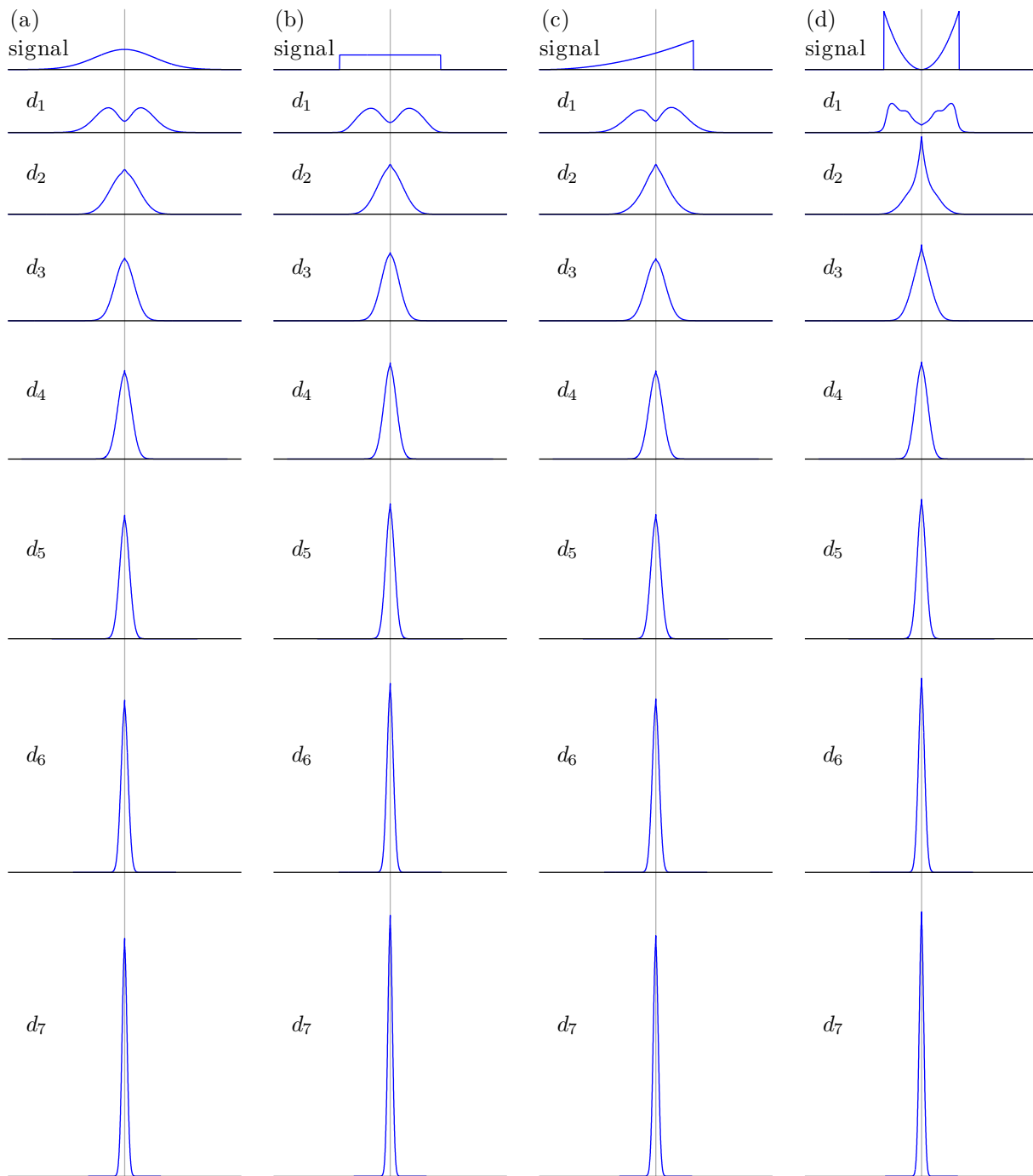


FIGURE 2.39 – Densités marginales des IMFs. (a) Gaussien. (b) uniforme. (c) asymétrique. (d) bi-modal.

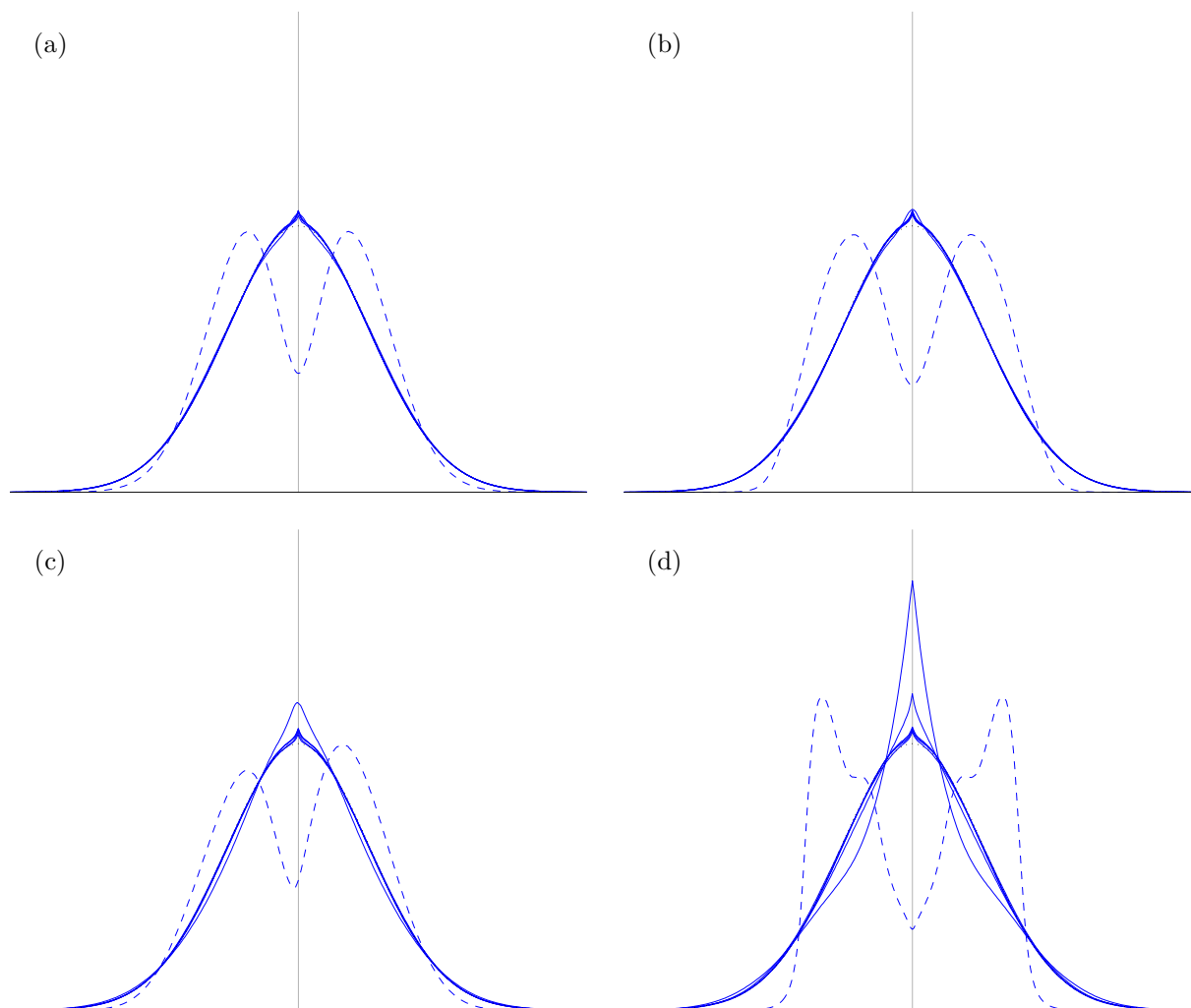


FIGURE 2.40 – Densités marginales des IMFs renormalisées. La courbe en pointillés est la gaussienne de référence. La courbe en tirets correspond au premier IMF. (a) Gaussien. (b) uniforme. (c) asymétrique. (d) bimodal.

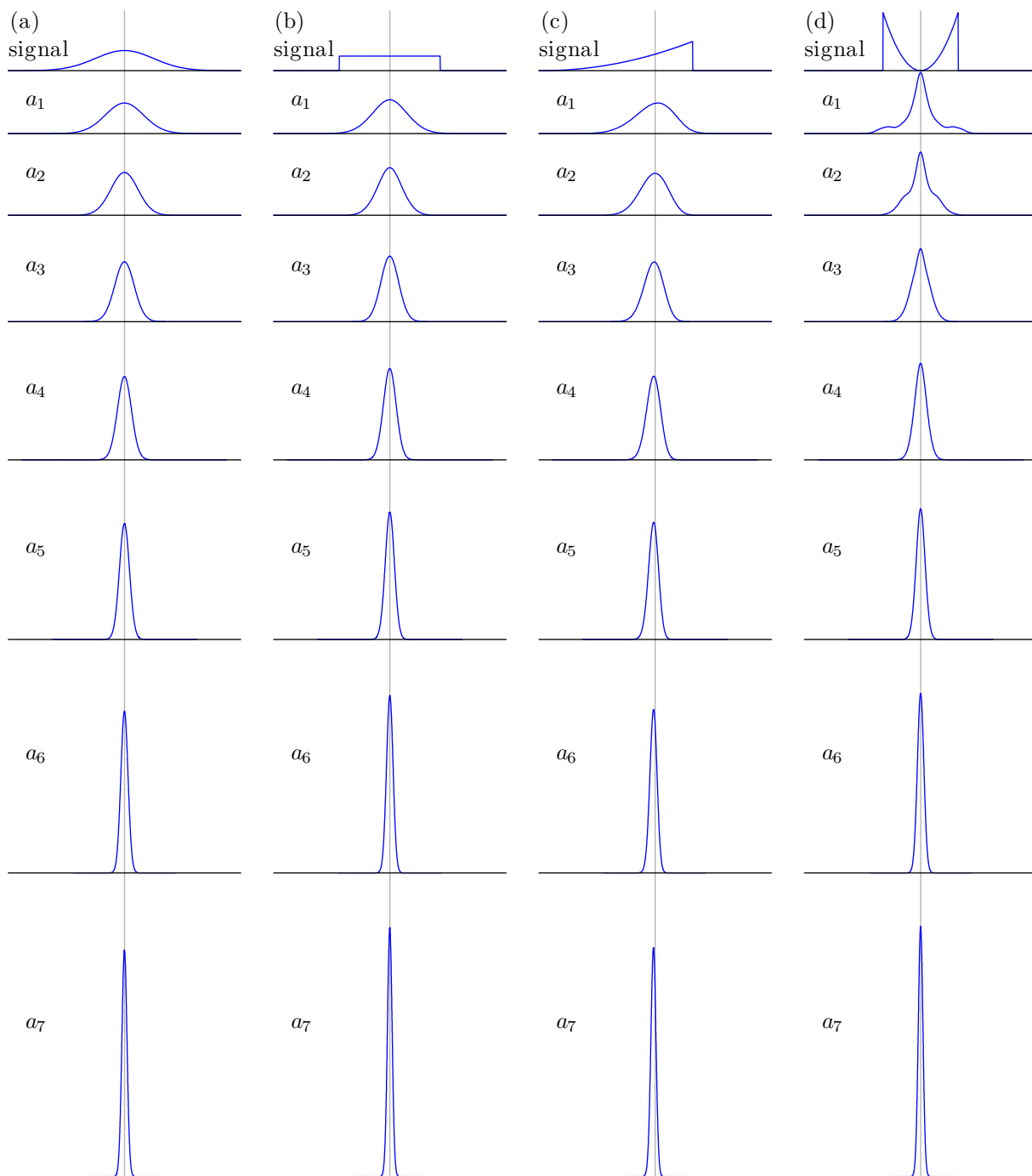


FIGURE 2.41 – Densités marginales des approximations. (a) Gaussien. (b) uniforme. (c) asymétrique. (d) bimodal.

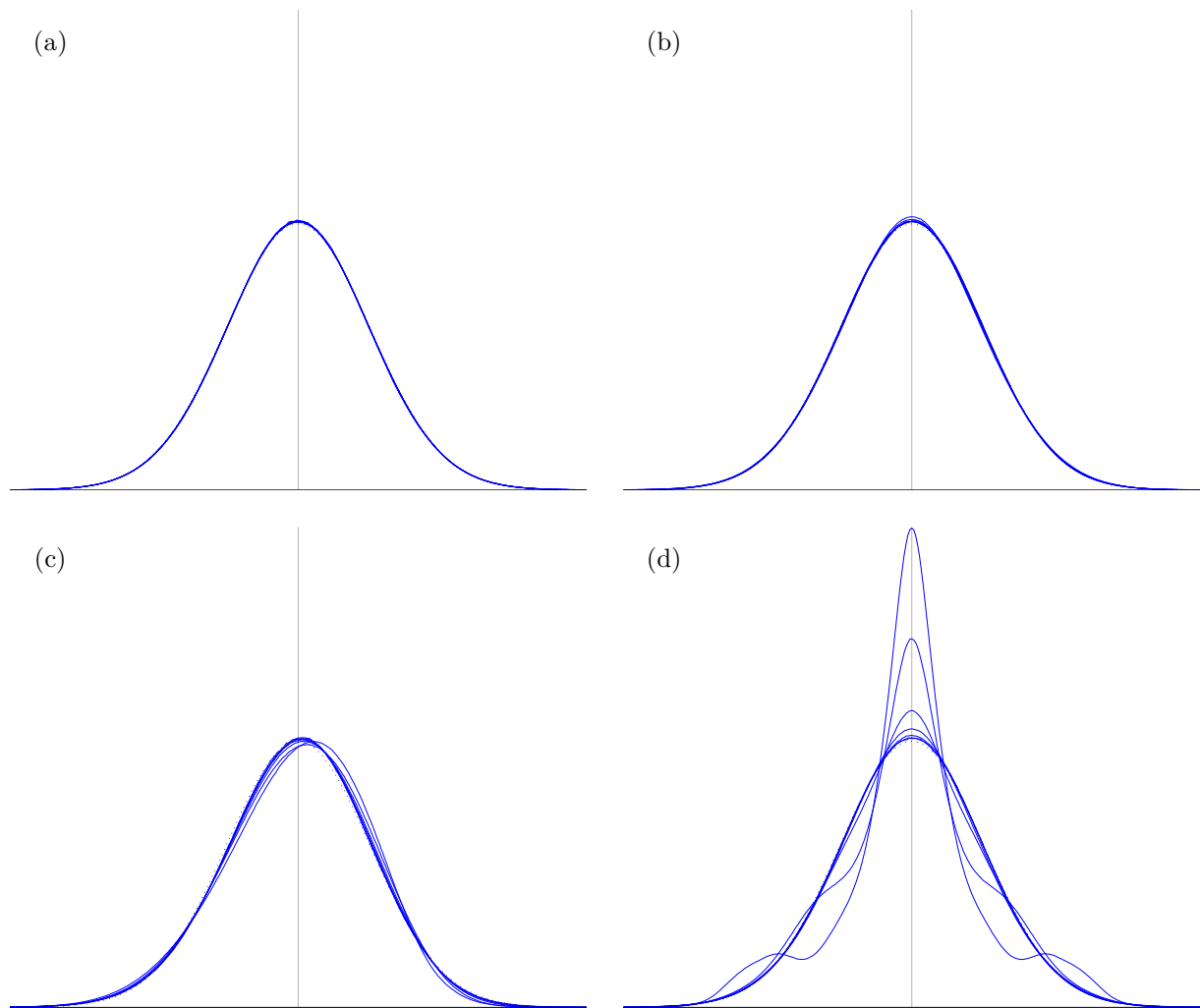


FIGURE 2.42 – Densités marginales des approximations renormalisées. La courbe en pointillés est la gaussienne de référence. La courbe en tirets correspond au premier IMF. (a) Gaussien. (b) uniforme. (c) asymétrique. (d) bimodal.

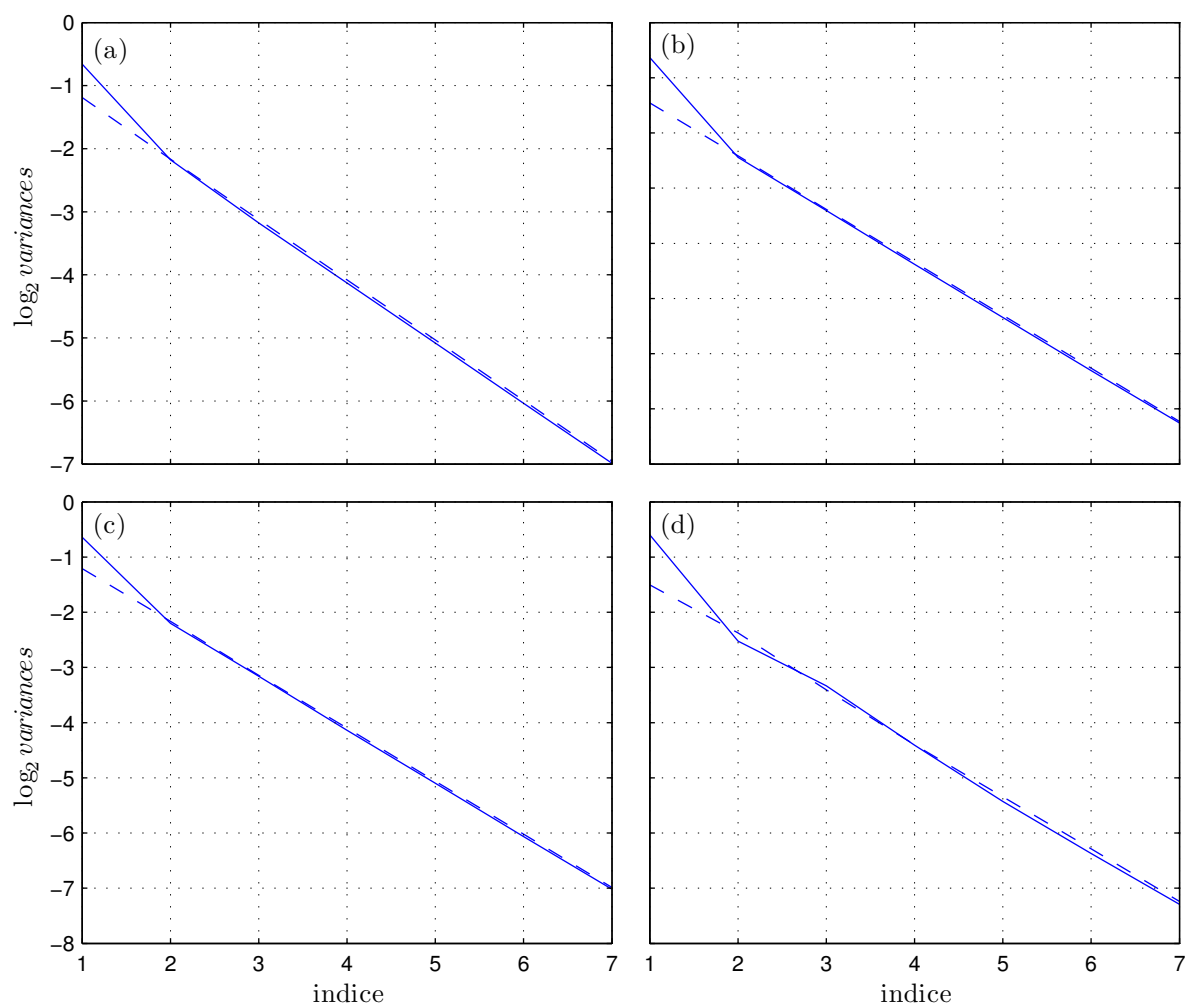


FIGURE 2.43 – Variances des IMFs (trait plein) et des approximations (tirets) en fonction de leur indice.

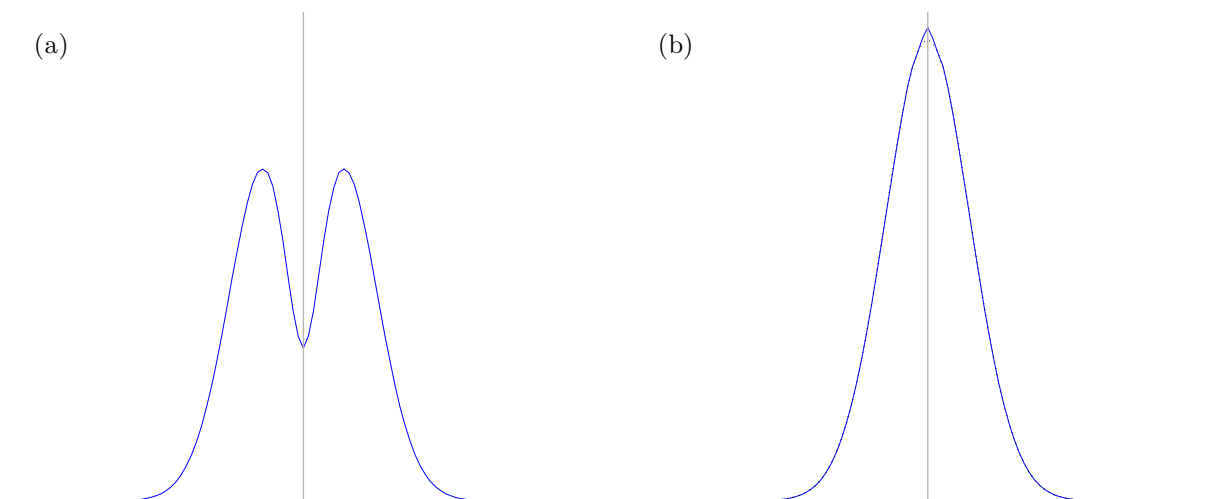


FIGURE 2.44 – Densité marginale du premier IMF issu du bruit blanc gaussien. (a) Densité marginale du premier IMF. (b) Densité marginale du premier IMF suréchantillonné d'un facteur 10 à l'aide d'une interpolation linéaire par morceaux. La courbe en pointillés est la gaussienne de même variance.

disparaître, au cours du processus de tamisage, le premier IMF a au moins autant d'extrema. Si l'on tient compte du fait qu'il s'agit d'un IMF on peut voir intuitivement que la symétrie entre les enveloppes fait que les extrema ont tendance à ne pas être trop proches de zéro. Au final, on obtient donc qu'au moins deux tiers des points auraient tendance à ne pas être trop proches de zéro, ce qui pourrait justifier une densité marginale bimodale.

Concernant les approximations, on observe que celles-ci sont elles aussi généralement gaussiennes, et même bien plus gaussiennes que les IMFs, puisqu'il n'y a pas de singularité en zéro. Les seuls cas où la densité marginale s'écarte notablement de la gaussienne sont ceux des premières approximations pour les densités asymétrique et bimodale. Enfin, comme dans le cas des IMFs on observe que les densités marginales des approximations sont d'autant plus gaussiennes que leur indice est grand. On observe de plus que leurs moyennes, contrairement à celles des IMFs, peuvent dépendre de la nature du bruit. Ainsi, les moyennes des approximations pour les bruits symétriques sont nulles comme on pouvait s'y attendre mais celles correspondant au bruit de densité asymétrique ne le sont pas, même pour les grands indices. Du point de vue des variances, on observe comme dans le cas des IMFs une décroissance en fonction de l'indice très proche d'exponentielle dans tous les cas.

2.1.3 Statistiques d'ordre 2

On a vu précédemment que les densités marginales des IMFs et des approximations étaient gaussiennes à l'exception du premier IMF et éventuellement des quelques suivants si la densité marginale du bruit est trop loin de la gaussienne. La situation est très différente pour ce qui est des densités à l'ordre 2, c'est-à-dire les densités de probabilité jointes des couples $(x(t), x(t+k))$ en fonction de k , avec x un IMF ou une approximation. En effet les IMFs quel que soit leur indice, ne sont pas du tout gaussiens à l'ordre 2. Les densités jointes du troisième IMF pour les différents types de bruits sont représentées Fig. 2.45, Fig. 2.46, Fig. 2.47 et Fig. 2.48. La même étude sur d'autres IMFs montre que l'allure de ces densités ne dépend que peu de l'indice de l'IMF, tant qu'il ne s'agit pas du premier. On constate de plus certaines similarités entre les densités des IMFs issus des différents types de bruits qui laissent penser qu'il y a sans doute un lien à trouver avec l'algorithme de l'EMD.

Concernant les approximations en revanche (Fig. 2.49, Fig. 2.50, Fig. 2.51 et Fig. 2.52), on constate qu'elles sont très raisonnablement gaussiennes lorsque les densités marginales sont gaussiennes comme dans le cas des bruits gaussien et uniforme. Dans les cas des bruits asymétrique et bimodal, celles-ci

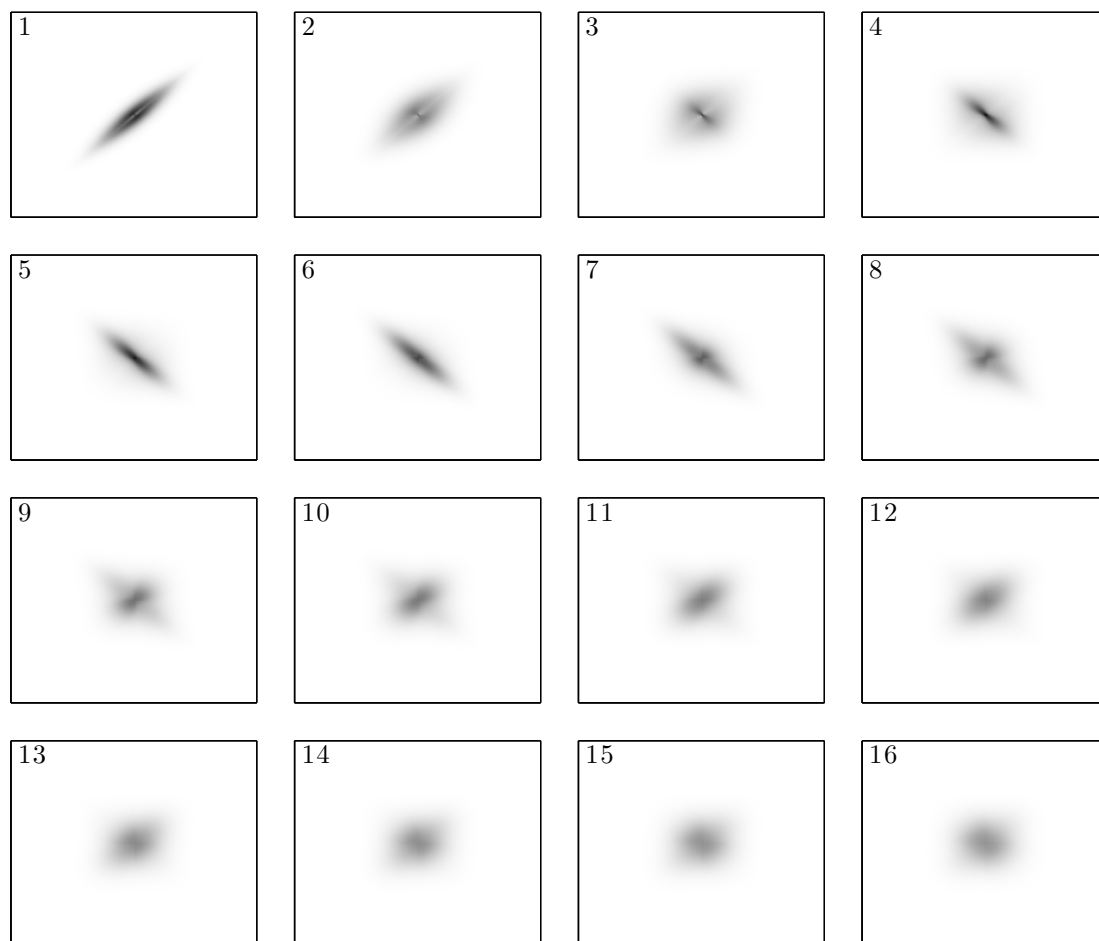


FIGURE 2.45 – Densités de probabilité à l'ordre 2 des IMFs issus du bruit gaussien. Chaque image correspond à une valeur de décalage indiquée dans le coin en haut à gauche.

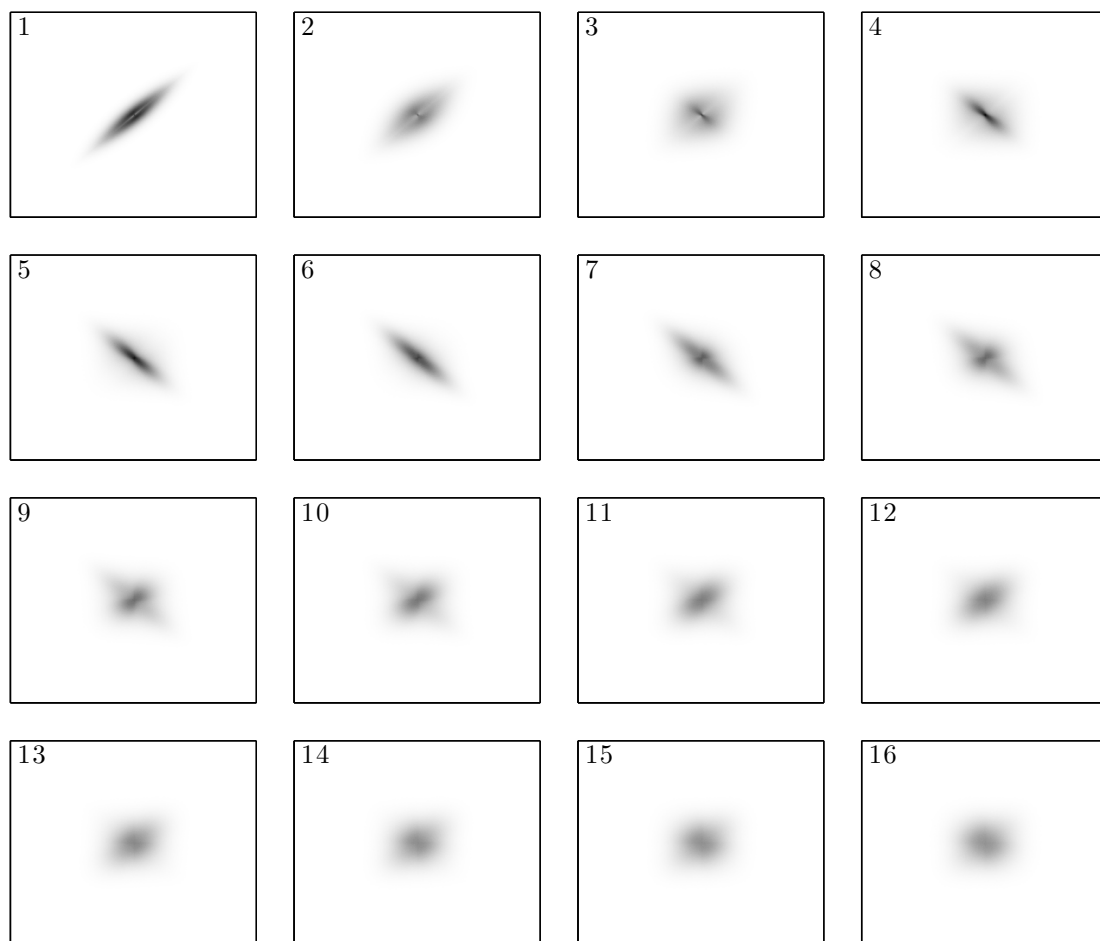


FIGURE 2.46 – Densités de probabilité à l'ordre 2 des IMFs issus du bruit uniforme. Chaque image correspond à une valeur de décalage indiquée dans le coin en haut à gauche.

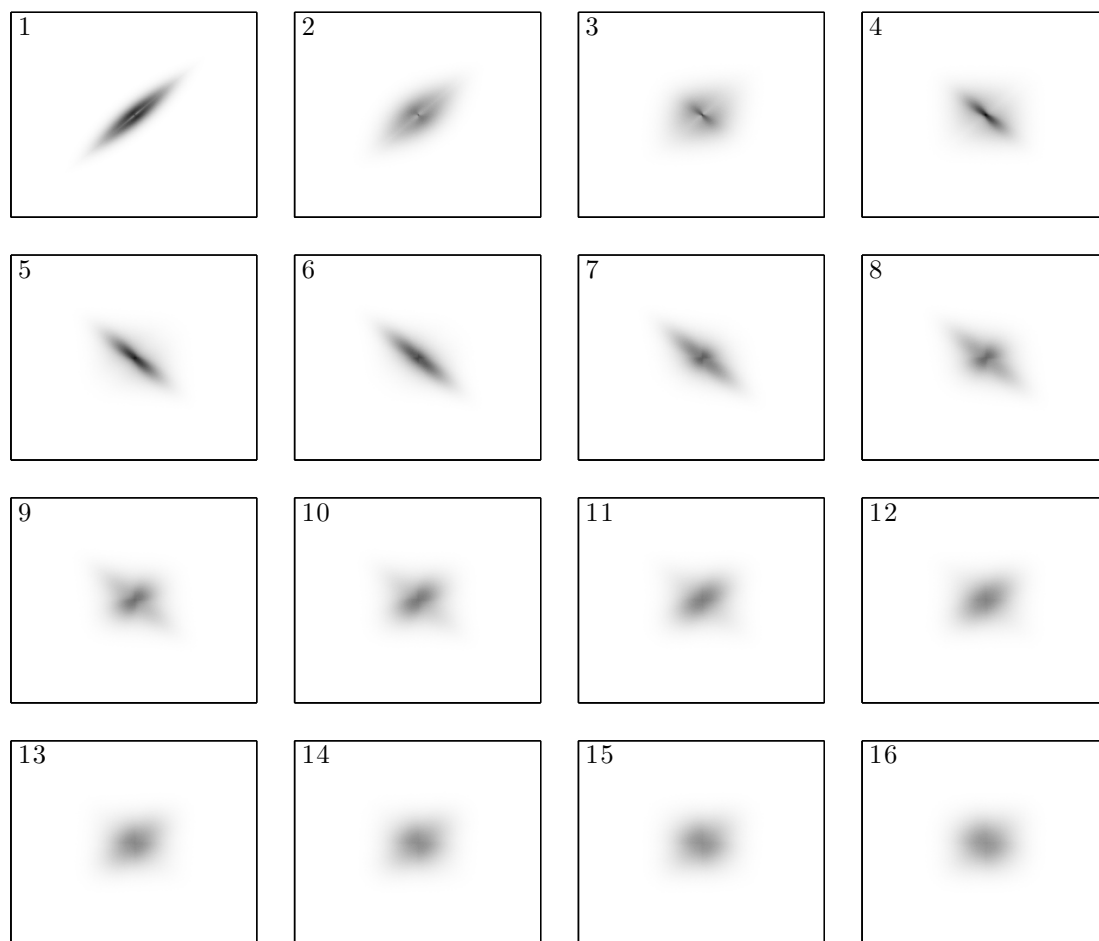


FIGURE 2.47 – Densités de probabilité à l'ordre 2 des IMFs issus du bruit asymétrique. Chaque image correspond à une valeur de décalage indiquée dans le coin en haut à gauche.

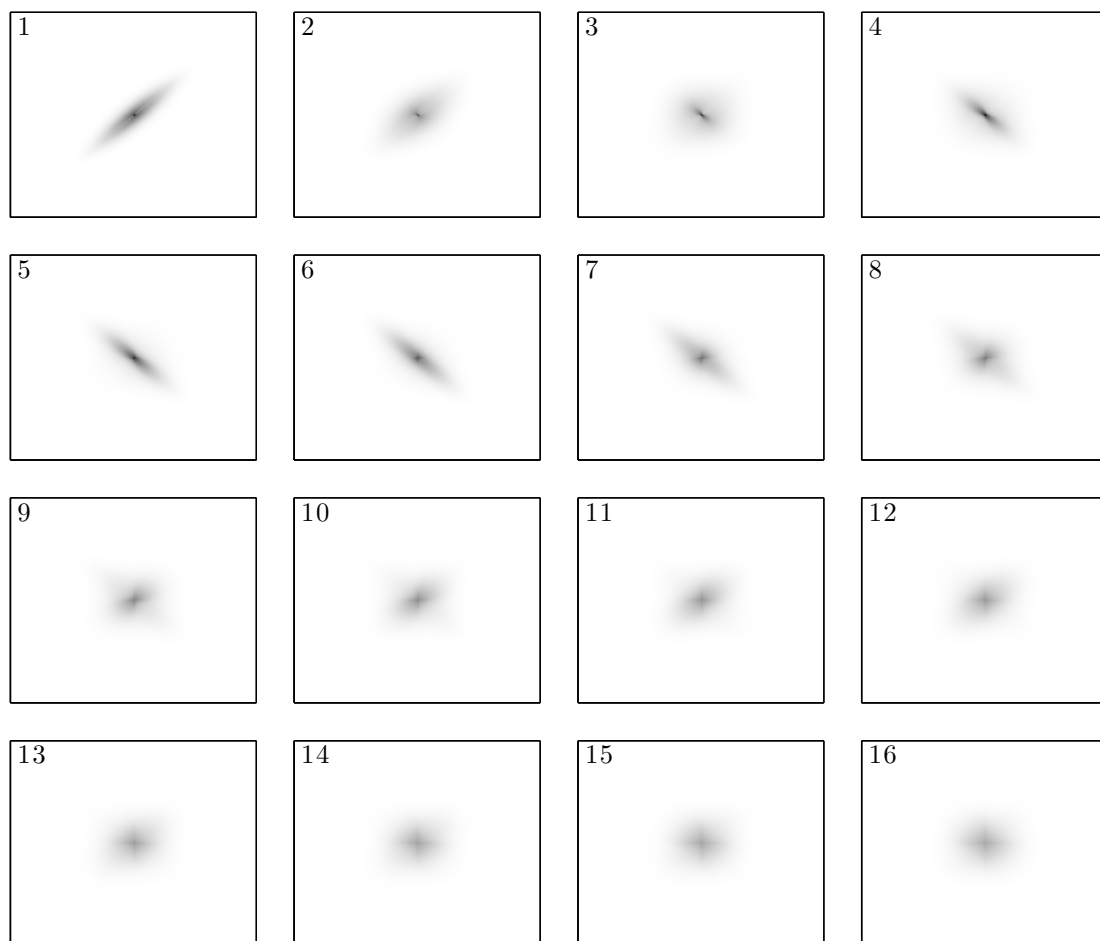


FIGURE 2.48 – Densités de probabilité à l'ordre 2 des IMFs issus du bruit bimodal. Chaque image correspond à une valeur de décalage indiquée dans le coin en haut à gauche.

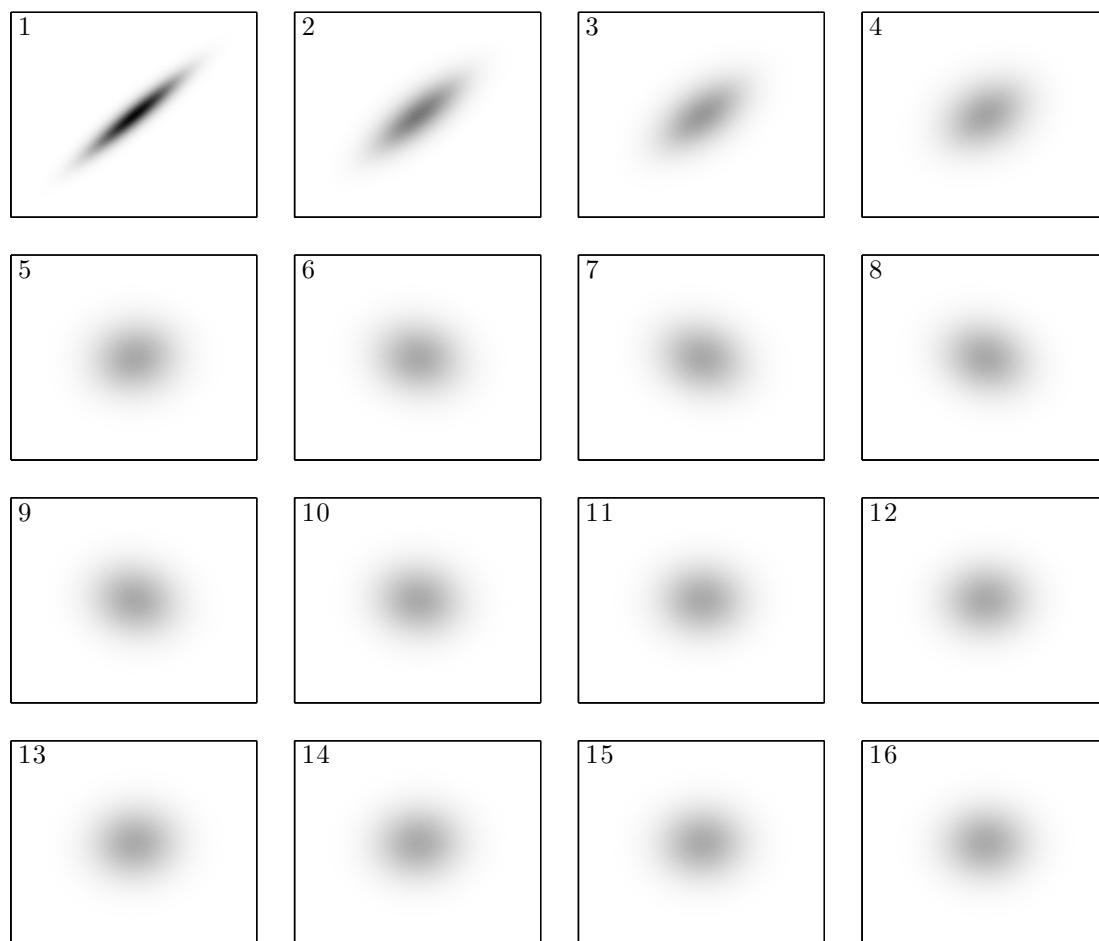


FIGURE 2.49 – Densités de probabilité à l'ordre 2 des approximations issues du bruit gaussien. Chaque image correspond à une valeur de décalage indiquée dans le coin en haut à gauche.

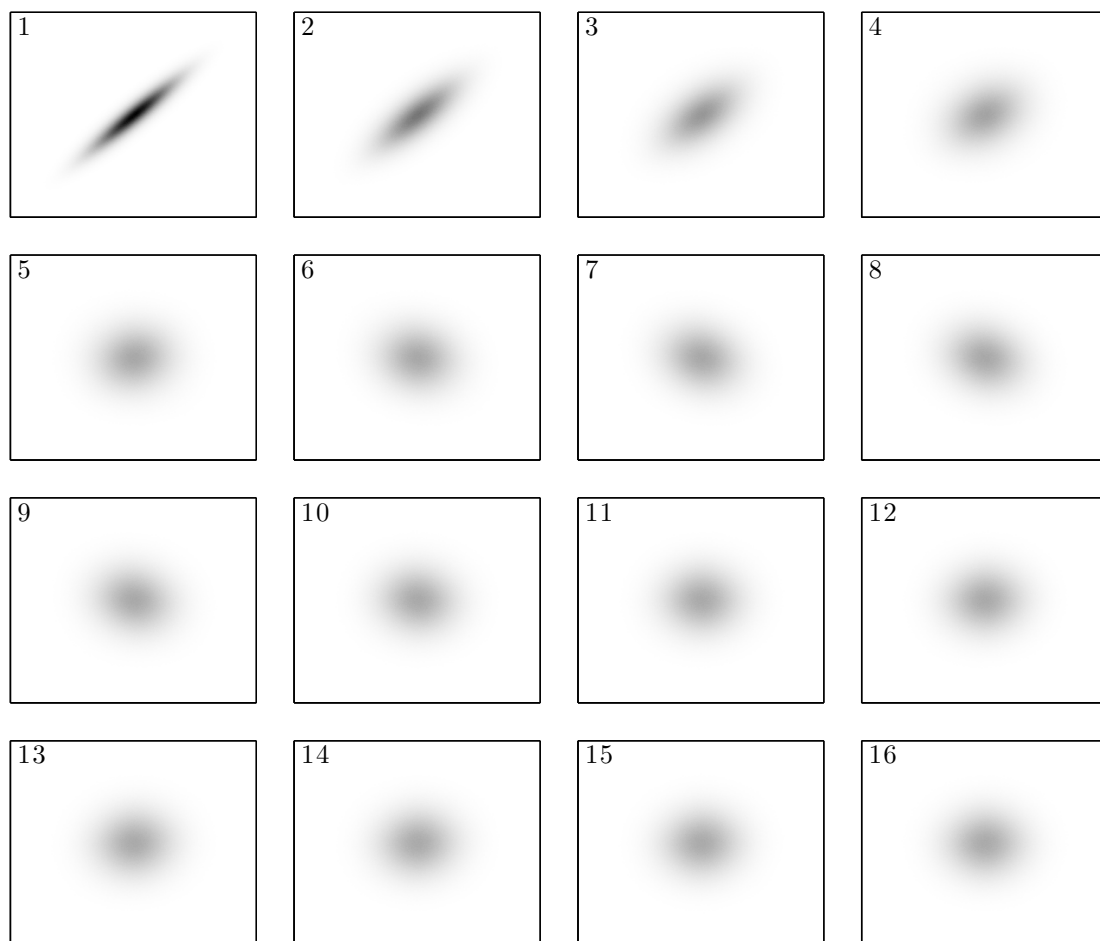


FIGURE 2.50 – Densités de probabilité à l'ordre 2 des approximations issues du bruit uniforme. Chaque image correspond à une valeur de décalage indiquée dans le coin en haut à gauche.

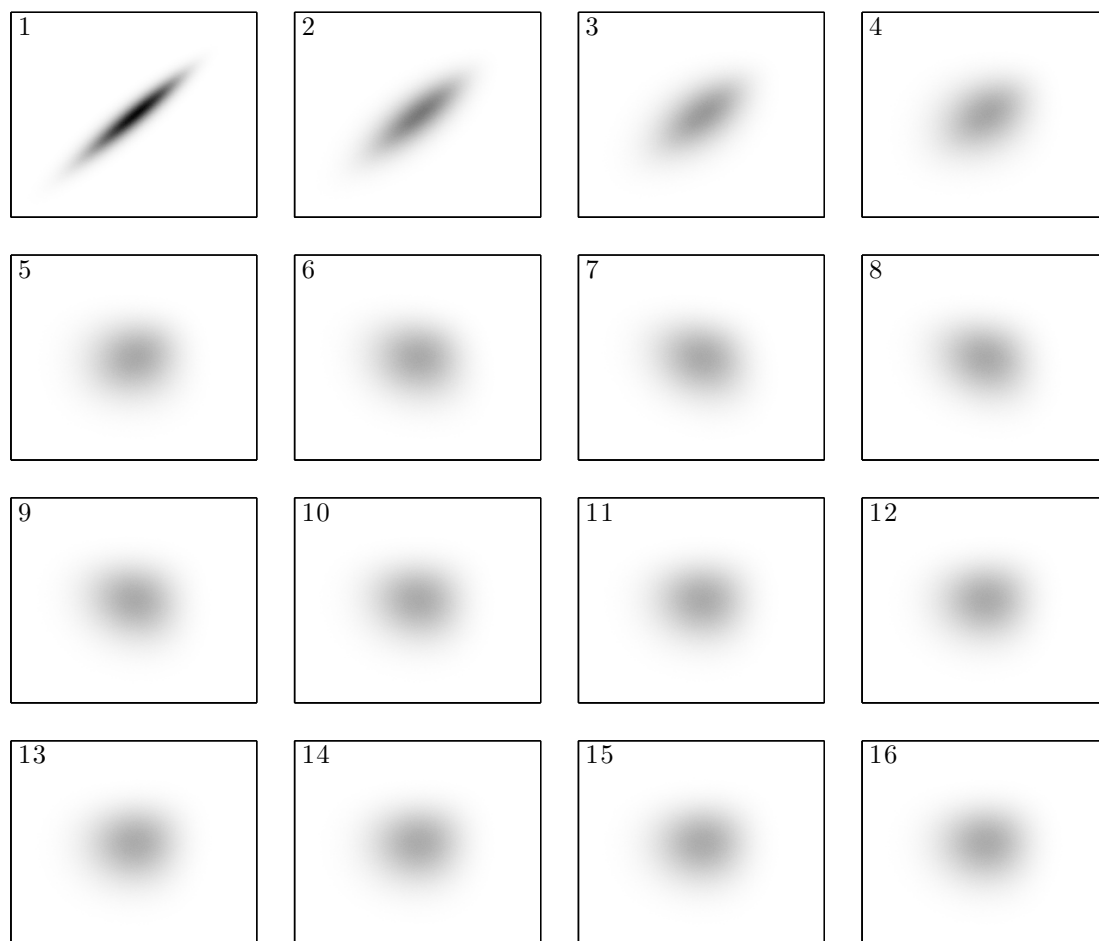


FIGURE 2.51 – Densités de probabilité à l'ordre 2 des approximations issues du bruit asymétrique. Chaque image correspond à une valeur de décalage indiquée dans le coin en haut à gauche.

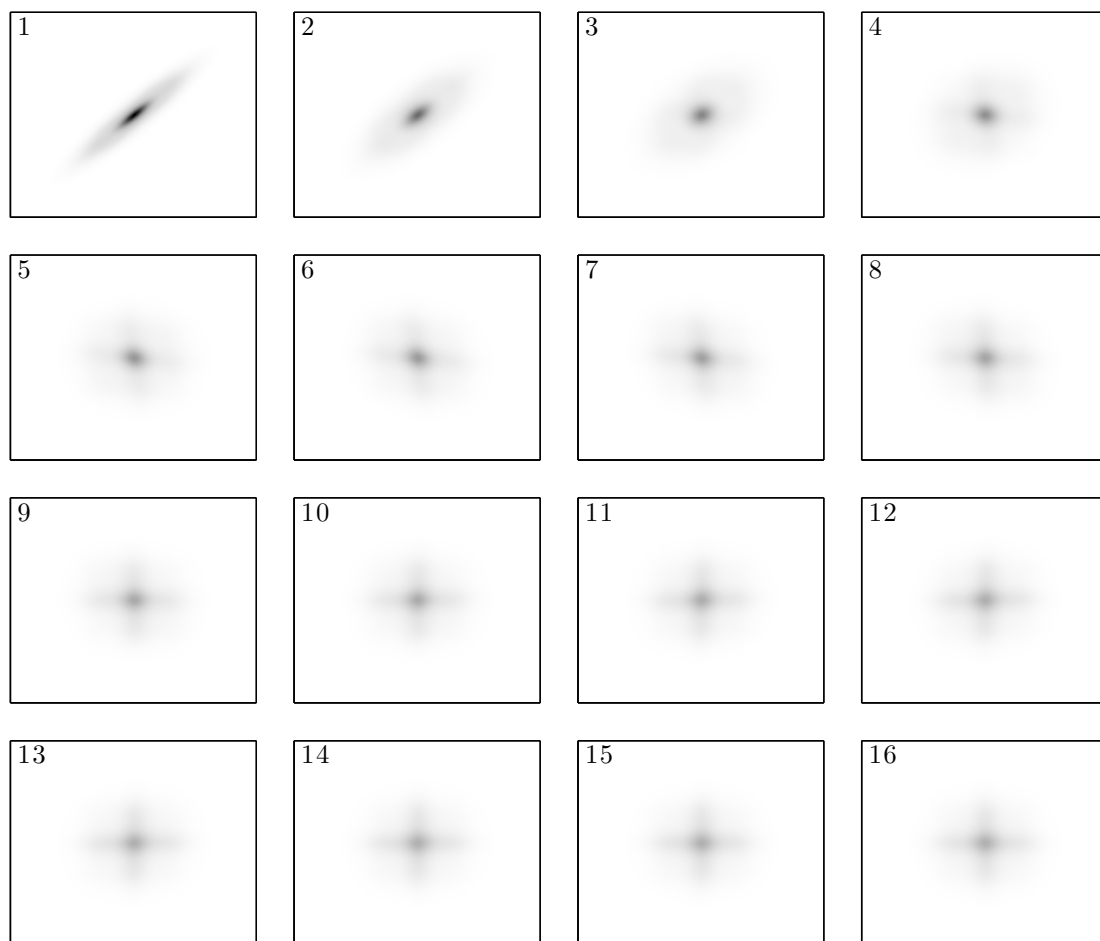


FIGURE 2.52 – Densités de probabilité à l'ordre 2 des approximations issues du bruit bimodal. Chaque image correspond à une valeur de décalage indiquée dans le coin en haut à gauche.

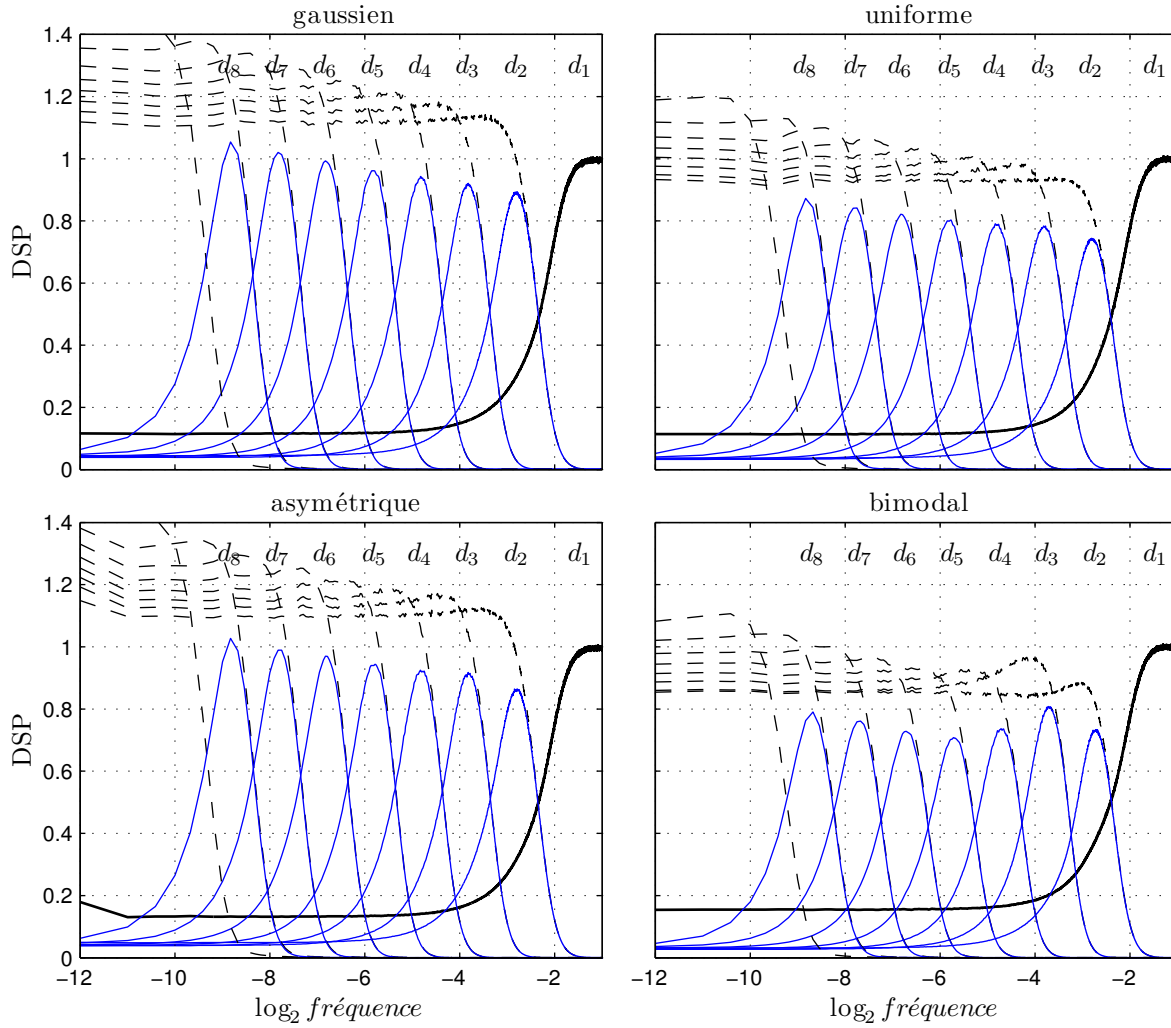


FIGURE 2.53 – Spectres des IMFs (traits pleins) et des approximations (tirets) pour les 4 types de bruit blanc. Le trait plus épais correspond au premier IMF. Les autres IMFs occupent des bandes de fréquences dont la fréquence centrale décroît avec l'indice de l'IMF.

semblent s'écarter de la gaussienne, mais seulement parce que le troisième IMF sur lequel sont basées les figures a une densité marginale qui s'écarte de la gaussienne dans ces deux cas.

Au-delà de ces densités de probabilité, on s'intéresse également aux spectres des IMFs et approximations Fig. 2.53. Ceux-ci sont estimés simplement pour un processus x par :

$$\hat{S}_x(f) = \sum_{m=-N+1}^{N-1} \hat{r}_x[m] w[m] e^{-2i\pi f m}, \quad |f| \leq \frac{1}{2}, \quad (2.182)$$

où w est une fenêtre de Hanning et $\hat{r}_x[m]$ l'estimée de l'autocorrélation du processus x à partir d'une série de réalisations $x^{(j)}, 1 \leq j \leq J = 100000$:

$$\hat{r}_x[m] = \frac{1}{J} \sum_{j=1}^J \left(\frac{1}{N-|m|} \sum_{n=1}^{N-|m|} x^{(j)}[n] x^{(j)}[n+|m|] \right) - \left(\frac{1}{NJ} \sum_{n=1}^N x^{(j)}[n] \right)^2, \quad |m| \leq N-1. \quad (2.183)$$

Les résultats montrent une organisation spontanée des spectres des IMFs qui rappelle fortement la sortie d'un banc de filtres autosimilaire, tel que ceux habituellement utilisés pour calculer les

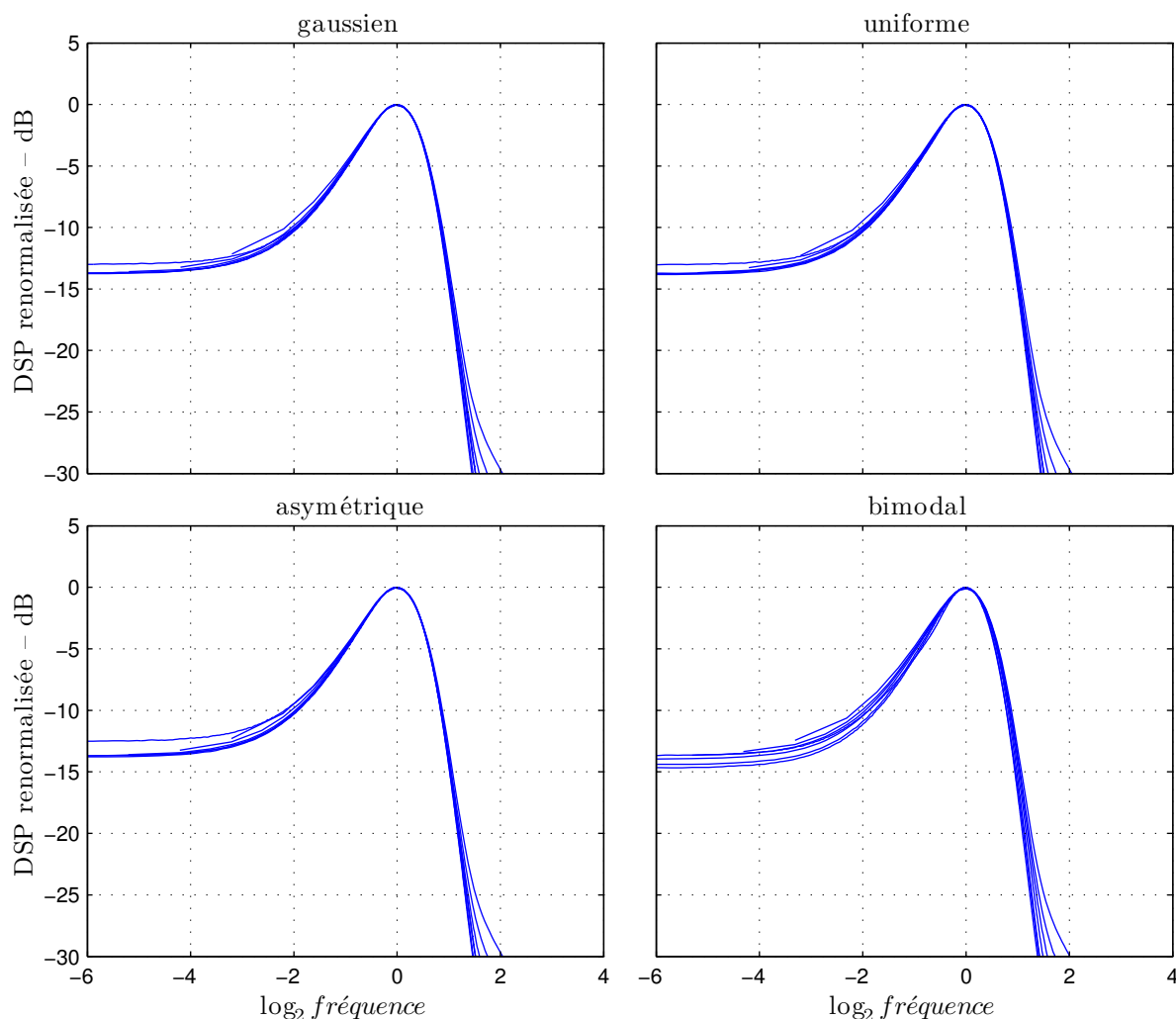


FIGURE 2.54 – Spectres des IMFs d'indice strictement supérieur à 1 renormalisés empiriquement.

transformées en ondelettes discrètes. En effet, on observe qu'à l'exception du premier IMF, les spectres des IMFs sont tous passe-bandes et quasi-identiques à une dilatation en fréquence — et un peu en amplitude — près, ce qu'on peut visualiser Fig. 2.54. De plus, le facteur de dilatation en fréquence d'un IMF au suivant semble être toujours le même. On verra par la suite que ce facteur, ici de l'ordre de 2, dépend en fait essentiellement du nombre d'itérations de tamisage qui a été fixé arbitrairement à 10 jusqu'ici. De fait, il semble que la principale différence entre les différents types de bruit soit l'amplitude des spectres des IMFs autres que le premier. On observe ainsi que les spectres des IMFs issus des bruits gaussien et asymétrique ont des amplitudes quasi-identiques, entre 0.9 et 1, et que ceux issus des bruits uniforme et bimodal ont des amplitudes plus faibles, de l'ordre de 0.8 pour le bruit uniforme et de 0.75 pour le bruit bimodal. En outre, on peut observer que les amplitudes des spectres des IMFs autres que le premier sont généralement croissantes avec l'indice, la seule exception étant le cas des premiers IMFs du bruit bimodal. Pour ce qui est du premier IMF, son comportement se démarque à nouveau des autres dans la mesure où il n'est pas passe-bande mais plutôt pass-haut et que, de plus, il a une contribution non négligeable à basses fréquences qui est en particulier beaucoup plus importante que celles des autres IMFs.

Concernant les approximations, on constate tout d'abord que leurs spectres sont tous passe-bas et qu'à l'instar des IMFs d'ordre supérieur à 1, leurs spectres ont tous pratiquement la même forme.

En outre, ils peuvent se déduire les uns des autres par une dilatation en fréquence et en amplitude qui est apparemment toujours la même d'une approximation à la suivante. L'autre caractéristique importante des approximations est le fait que leur densité de puissance à basses fréquences augmente d'une approximation à la suivante et que celle-ci dépasse systématiquement la densité spectrale de puissance du bruit à partir d'un certain indice, parfois dès la première approximation.

Pour clore la description à l'ordre 2 des IMFs, on s'est également intéressé aux corrélations entre eux. Pour ce faire, on a estimé leurs interspectres à l'aide de la même formule que celle utilisée pour l'estimation des spectres (2.182), en changeant juste l'estimateur de l'autocorrélation (2.183) en estimateur de l'intercorrélacion. Les normes des interspectres divisées par la moyenne géométrique des normes des spectres correspondants sont représentées Fig. 2.55 et les interspectres entre les 4 premiers IMFs pour le bruit blanc gaussien Fig. 2.56. Les interspectres entre les IMFs issus des autres types de bruit n'ont pas été inclus car ils sont qualitativement très proches du cas gaussien. En revanche, on propose figure Fig. 2.57 les interspectres entre premier IMF et première approximation qui font bien apparaître les différences de corrélations entre les différents types de bruit.

Les résultats montrent qu'à l'exception du premier IMF, les IMFs sont peu corrélés entre eux. De plus, il semble que la corrélation entre deux IMFs dépende essentiellement de l'écart d'indice entre les deux. Ainsi la norme de leurs interspectres renormalisés est au maximum de l'ordre de 0.1 entre deux IMFs successifs, de l'ordre de 0.03 à deux d'écart, et elle tombe à 0.01 pour de quatre à six d'écart suivant le type de bruit. Au-delà des normes, on peut analyser la nature de la corrélation à l'aide des interspectres entre les IMFs. Remarquons tout d'abord que ces derniers sont systématiquement réels. Ceci s'explique par le fait que les intercorrélations sont toutes réelles et paires : réelles parce que les IMFs sont réels, et paires parce que le bruit et l'EMD ne privilégient aucune direction temporelle. Plus en détail, on observe comme pour les normes que l'allure de l'interspectre entre deux IMFs dépend essentiellement de l'écart entre leurs indices. On peut ainsi voir que la nature de la corrélation diffère fortement entre deux IMFs successifs et entre deux IMFs plus éloignés. Dans le premier cas, on observe que les deux IMFs sont corrélés positivement dans la zone où leurs spectres se recouvrent, puis négativement dans la partie plus basses fréquences de la bande où l'IMF de plus grand indice domine. Dans le second cas, il n'y a pas de zone de recouvrement et l'interspectre est négatif et concentré dans la bande de fréquence de l'IMF de plus grand indice. Du point de vue de l'interprétation des résultats, ces corrélations peuvent apparaître plus ou moins gênantes. En effet, les corrélations positives ont comme conséquence que la somme des puissances des IMFs successifs dans la zone où leurs spectres se recouvrent est inférieure à la puissance de la somme, ce qui peut conduire à la sous-estimation de certaines quantités. Mais les corrélations négatives sont a priori bien plus gênantes puisqu'elles indiquent que les IMFs contiennent des informations qui se compensent entre elles quand on les somme et qu'il est donc possible de voir dans un IMF particulier des informations absentes du signal.

Le cas du premier IMF est différent puisqu'il est soit plus (bruits gaussien et asymétrique) soit moins (bruits uniforme et bimodal) corrélé avec les IMFs suivants que ces derniers entre eux. On peut d'ailleurs constater que ce dernier résultat se voit également sur les interspectres entre premier IMF et première approximation (cf Fig. 2.57) : ceux des bruits uniforme et bimodal sont très faibles à basses fréquences, ce qui implique que le premier IMF est peu corrélé avec les IMFs basses fréquences ; ceux des bruits gaussien et asymétrique sont en revanche assez fortement négatifs à basses fréquences et le premier IMF est donc assez fortement anticorrélé avec les IMFs basses fréquences. La raison pour laquelle les premiers IMFs issus des bruits uniforme et bimodal sont nettement moins corrélés avec les autres que pour les autres types de bruit est loin d'être claire a priori mais une piste séduisante semble être le lien entre l'aspect bimodal de ces densités⁵ et celui de la densité marginale du premier IMF.

5. La densité uniforme est en fait entre unimodale et bimodale, mais elle est nettement plus proche de bimodale que les densités gaussienne et asymétrique.

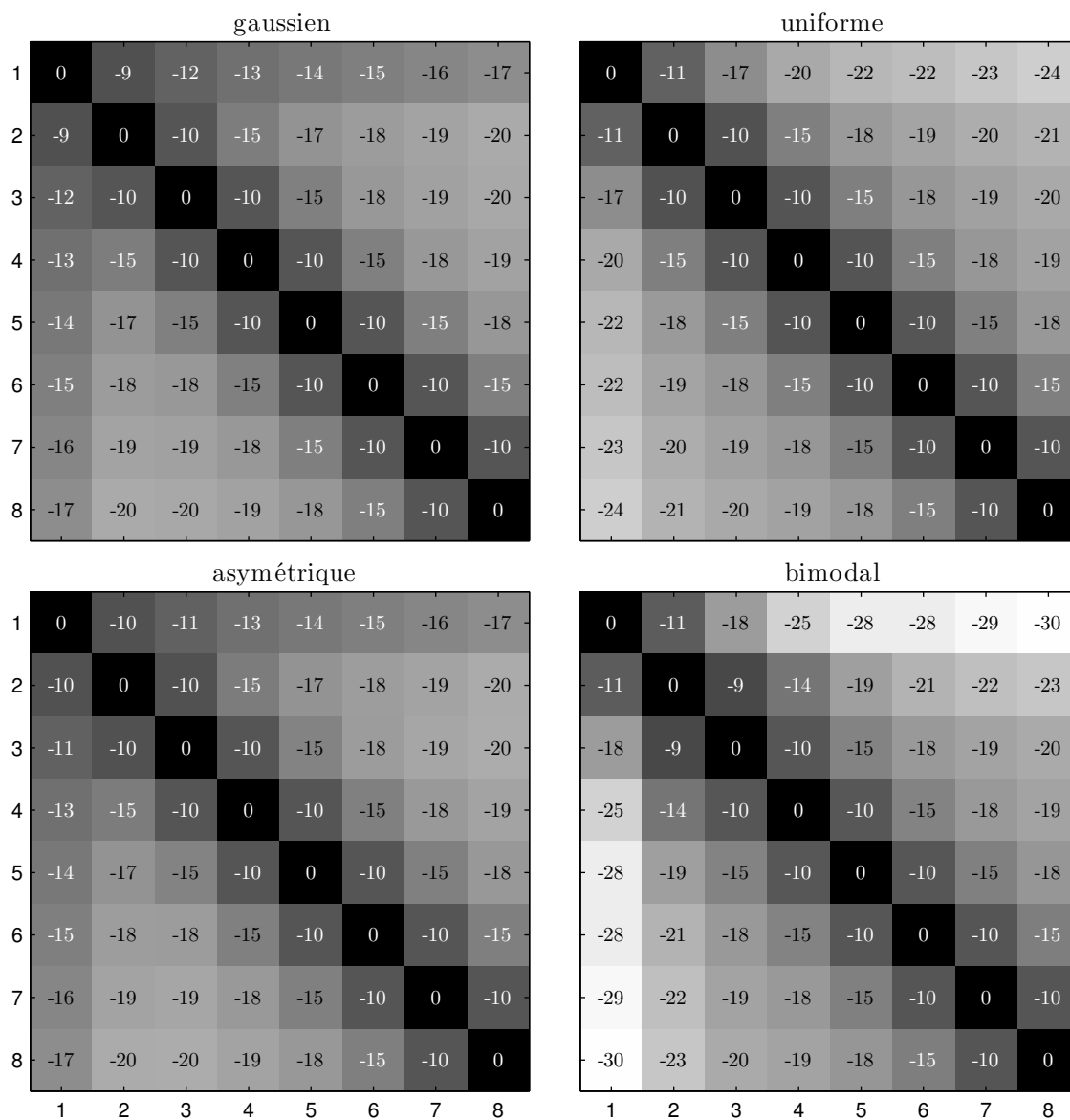


FIGURE 2.55 – Normes des interspectres entre les IMFs pour les 4 types de bruit blanc. Les normes des interspectres sont toutes normalisées par la moyenne géométrique des normes des spectres des deux IMFs correspondants. (dynamique en dB)

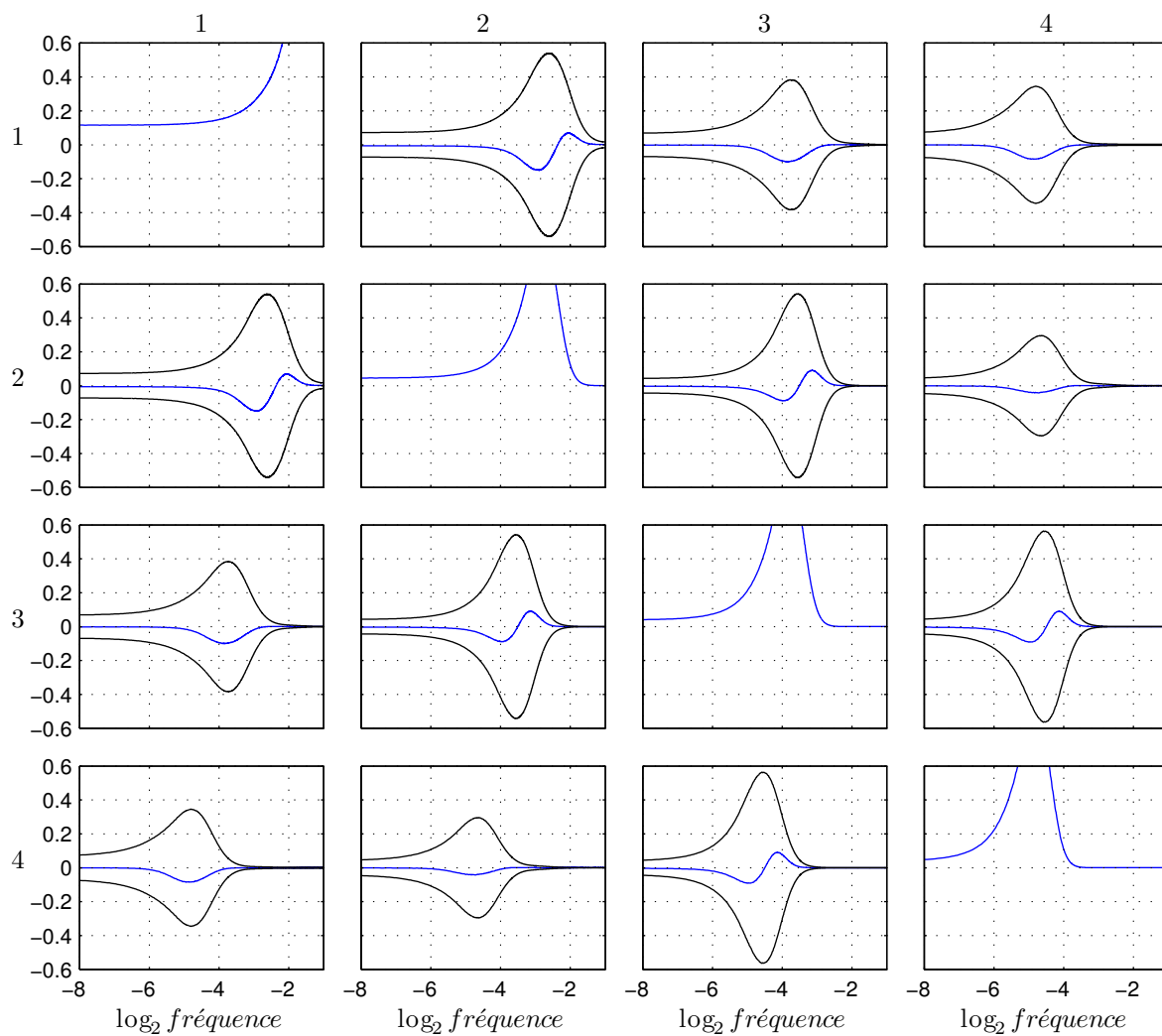


FIGURE 2.56 – Interspectres entre les 4 premiers IMFs pour le bruit blanc gaussien. Les courbes noires enveloppant les interspectres sont les moyennes géométriques des spectres des IMFs correspondants. Les spectres des IMFs sur la diagonale sont écrêtés pour améliorer la lisibilité des interspectres.

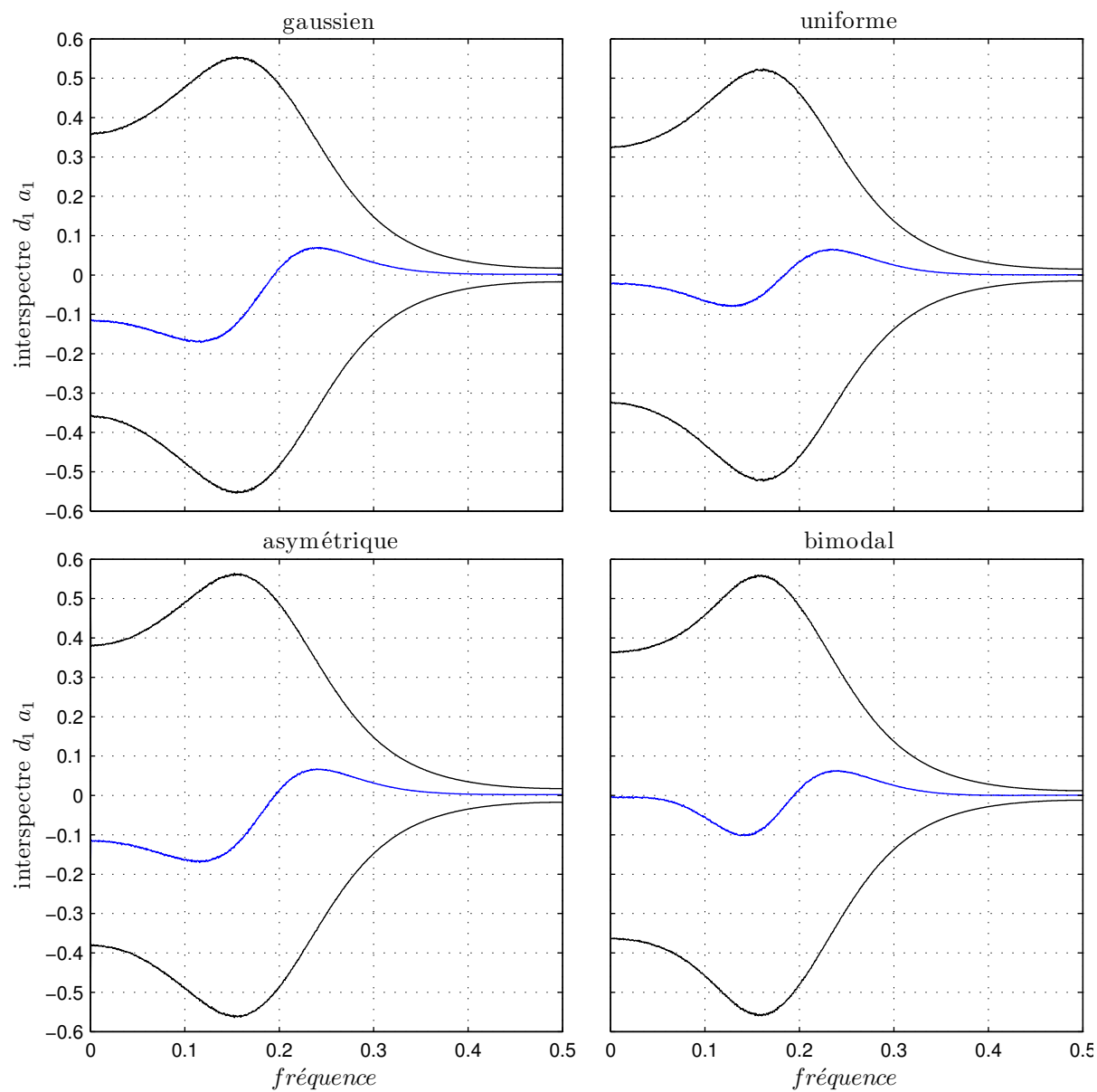


FIGURE 2.57 – Interspectre entre le premier IMF et la première approximation. Les courbes noires enveloppant les interspectres sont les moyennes géométriques des spectres du premier IMF et de la première approximation.

2.2 Bruit gaussien fractionnaire — estimation de l'exposant d'une loi d'échelle

2.2.1 Bruit Gaussien fractionnaire

Les bruits gaussiens fractionnaires (fGn pour “fractional Gaussian noise”) sont des processus intrinsèquement discrets définis comme les processus d'accroissements des mouvements browniens fractionnaires, seuls processus gaussiens auto-similaires à accroissements stationnaires. Les fGns sont ainsi des processus gaussiens stationnaires de moyenne nulle et ils sont donc entièrement caractérisés par leur structure au second ordre. Cette dernière présente l'intérêt de ne dépendre que d'un seul paramètre : l'exposant de Hurst ⁶ $H \in]0, 1[$ qui intervient dans sa fonction d'auto-corrélation $r_H[k] := \mathbb{E} \{x_H[n]x_H[n+k]\}$:

$$r_H[k] = \frac{\sigma^2}{2} (|k-1|^{2H} - 2|k|^{2H} + |k+1|^{2H}). \quad (2.184)$$

Le cas $H = 1/2$ se ramène au bruit blanc alors que les autres cas correspondent à des bruits corrélés, soit anticorrélés pour $H \in]0, 1/2[$, soit positivement corrélés pour $H \in]1/2, 1[$.

La transformée de Fourier de (2.184) fournit la densité spectrale de puissance associée :

$$S_H(f) = C\sigma^2 \left| e^{i2\pi f} - 1 \right|^2 \sum_{k=-\infty}^{k=\infty} \frac{1}{|f+k|^{2H+1}}, \quad (2.185)$$

avec $|f| < \frac{1}{2}$. Pour $H \neq \frac{1}{2}$, on peut approximer cette dernière par

$$S_H(f) \simeq C(2\pi\sigma)^2 |f|^{1-2H} \quad (2.186)$$

quand $|f| \rightarrow 0$. On notera que cette approximation est en pratique relativement bien vérifiée pour toute la bande $[0, 0.5]$. Ceci fait du fGn un bon modèle de signaux dont le spectre prend la forme d'une loi de puissance, ce qui est souvent associé à la notion d'invariance d'échelle.

D'un point de vue spectral, les fGns peuvent être séparés en deux classes par la valeur $H = \frac{1}{2}$. Quand $H \in]0, \frac{1}{2}[$, $S_H(0) = 0$ et le spectre est essentiellement passe-haut. À l'inverse, quand $H \in]\frac{1}{2}, 1[$, $S_H(0) = \infty$ avec une divergence de type « $1/f$ » et le spectre est essentiellement passe-bas.

2.2.2 Densités de probabilité aux ordres 1 et 2

Les densités de probabilité des IMFs et approximations issus de fGns présentent de manière générale les mêmes caractéristiques que pour le bruit blanc gaussien (cf Fig. 2.58, Fig. 2.59, Fig. 2.60 et Fig. 2.61) :

- les densités marginales des IMFs sont pratiquement gaussiennes (avec la même singularité en zéro) à l'exception de celle du premier IMF qui est toujours bimodale.
- les densités marginales des approximations sont gaussiennes.
- les densités à l'ordre 2 des IMFs sont non gaussiennes et ont toutes la même allure. Celles des approximations sont gaussiennes.

2.2.3 Spectres et filtres équivalents des IMFs

Les densités spectrales de puissance pour chacun des IMFs extraits à l'aide de dix itérations de tamisage ont été représentées Fig. 2.62. On propose également Fig. 2.64 les mêmes densités spectrales mais calculées à partir d'IMFs extraits à l'aide d'un nombre variable d'itérations, contrôlé par le critère d'arrêt « local » (cf 5.2) avec $\epsilon = 0.1$. Les paramètres utilisés pour le critère ont été choisis pour garder le même ordre de grandeur de 10 itérations de tamisage au moins pour la première

6. On parlera par la suite abusivement de « l'exposant de Hurst du fGn » alors qu'il s'agit en fait de l'exposant de Hurst du mouvement Brownien fractionnaire associé.

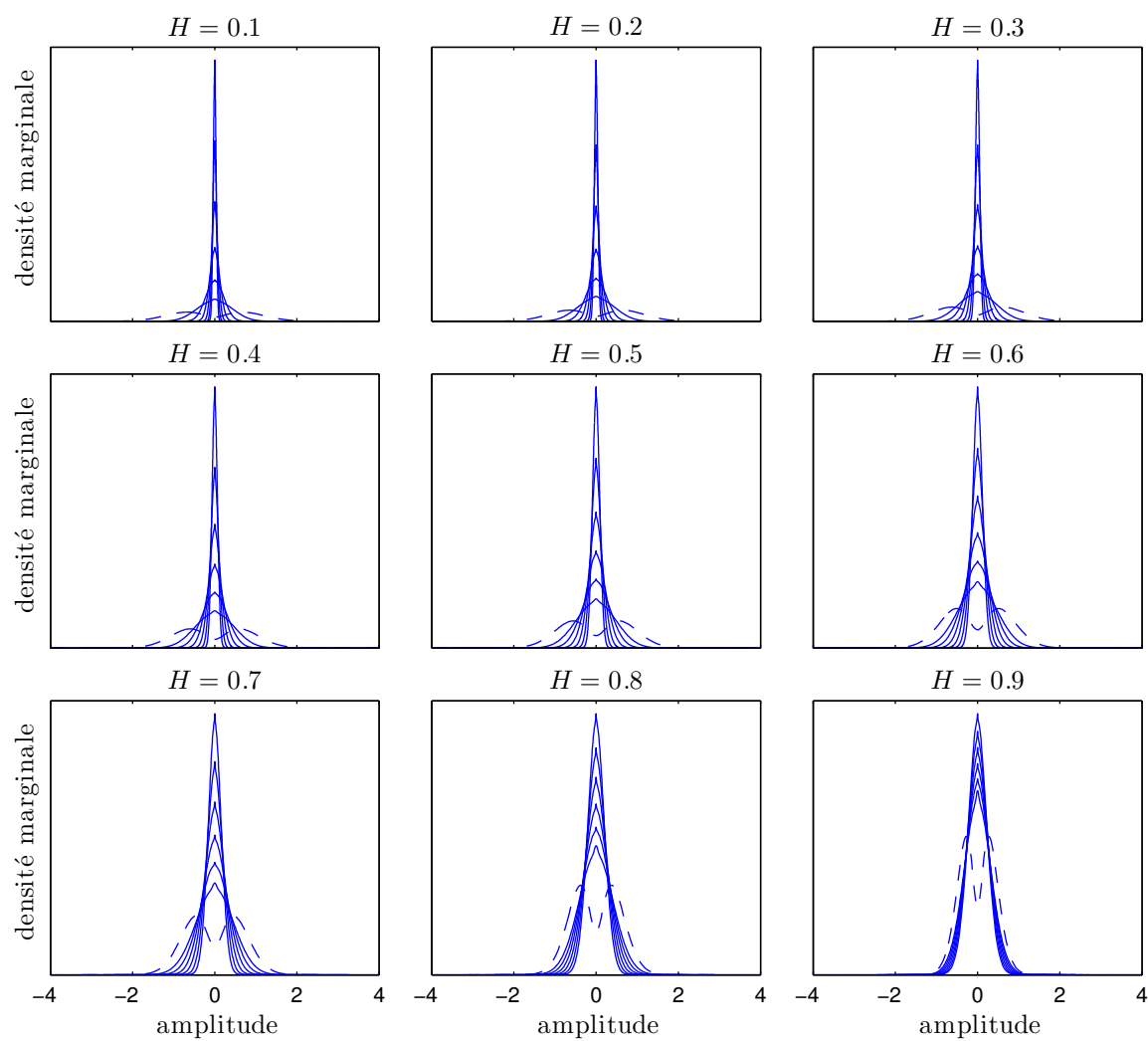


FIGURE 2.58 – Densités marginales des IMFs. La courbe en tirets correspond au premier IMF.

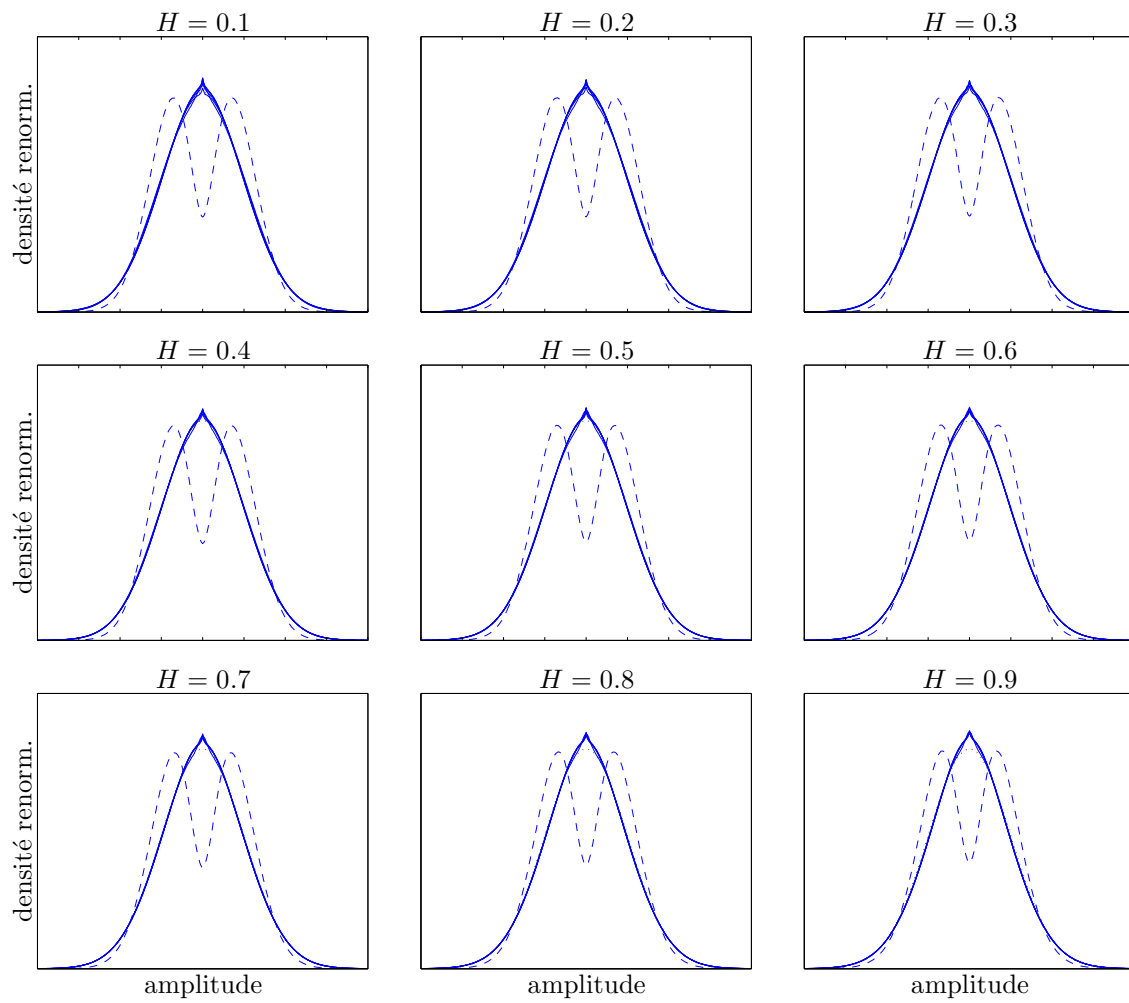


FIGURE 2.59 – Densités marginales des IMFs renormalisées. La courbe en pointillés est la gaussienne de référence. La courbe en tirets correspond au premier IMF.

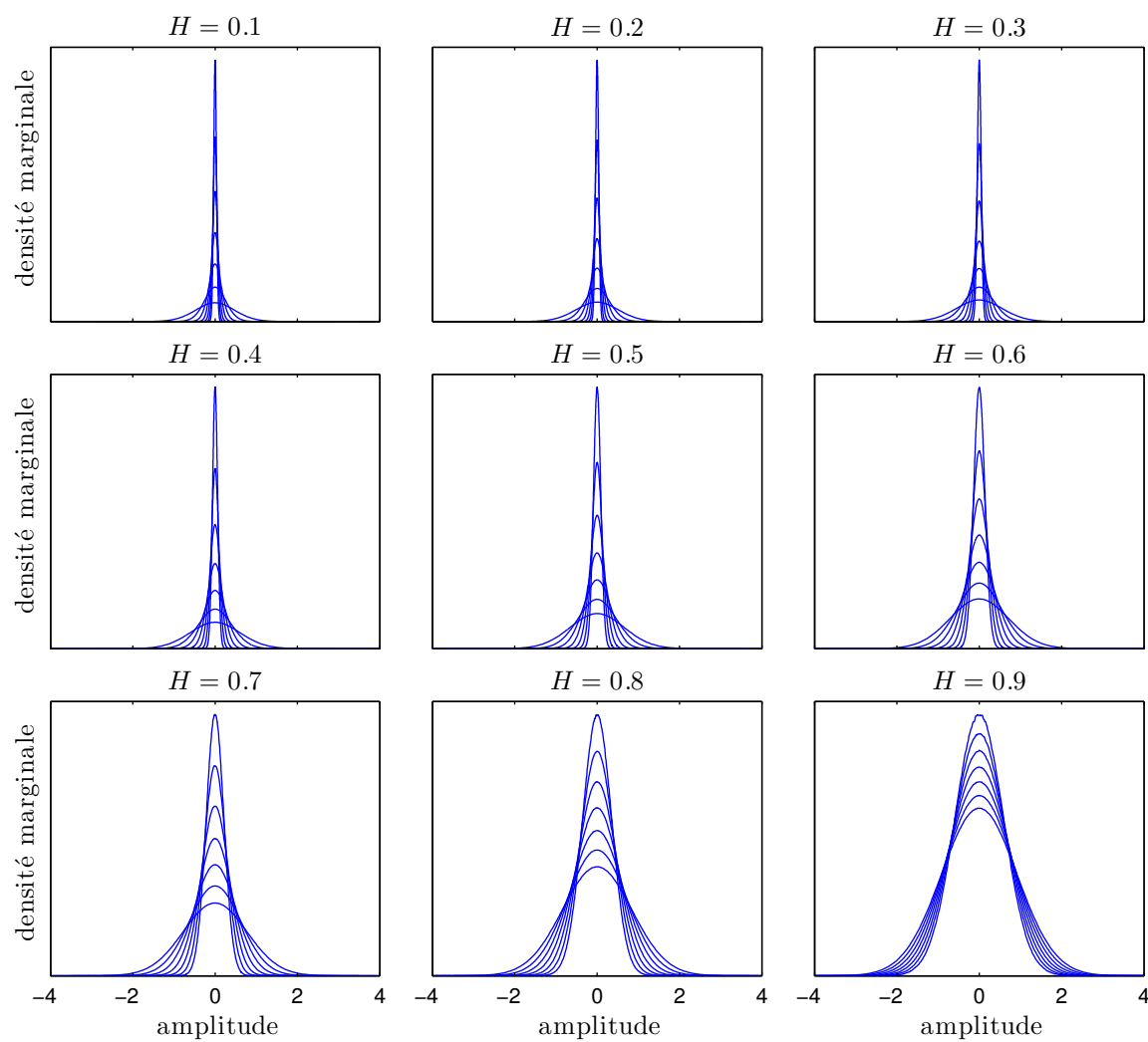


FIGURE 2.60 – Densités marginales des approximations.

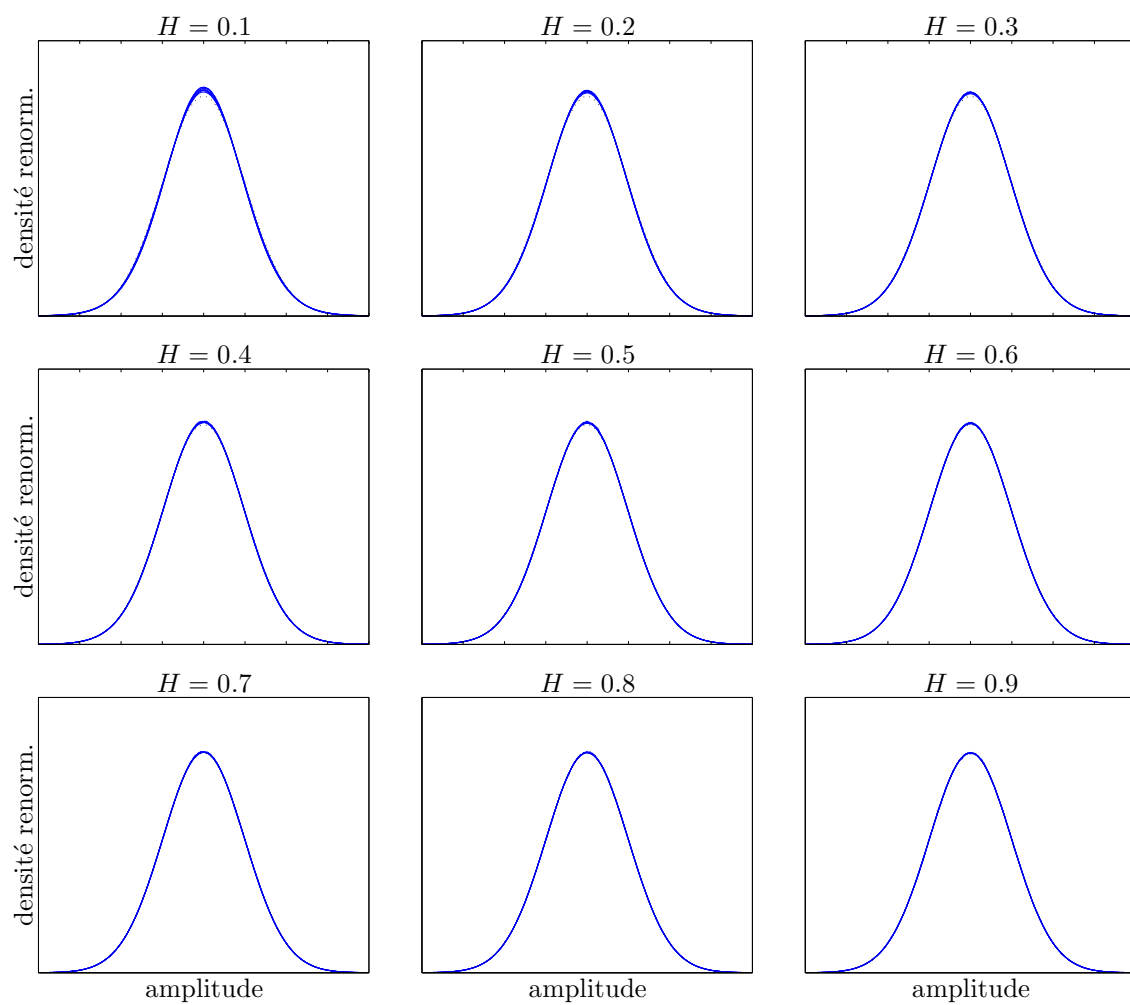


FIGURE 2.61 – Densités marginales des approximations renormalisées. La courbe en pointillés est la gaussienne de référence.

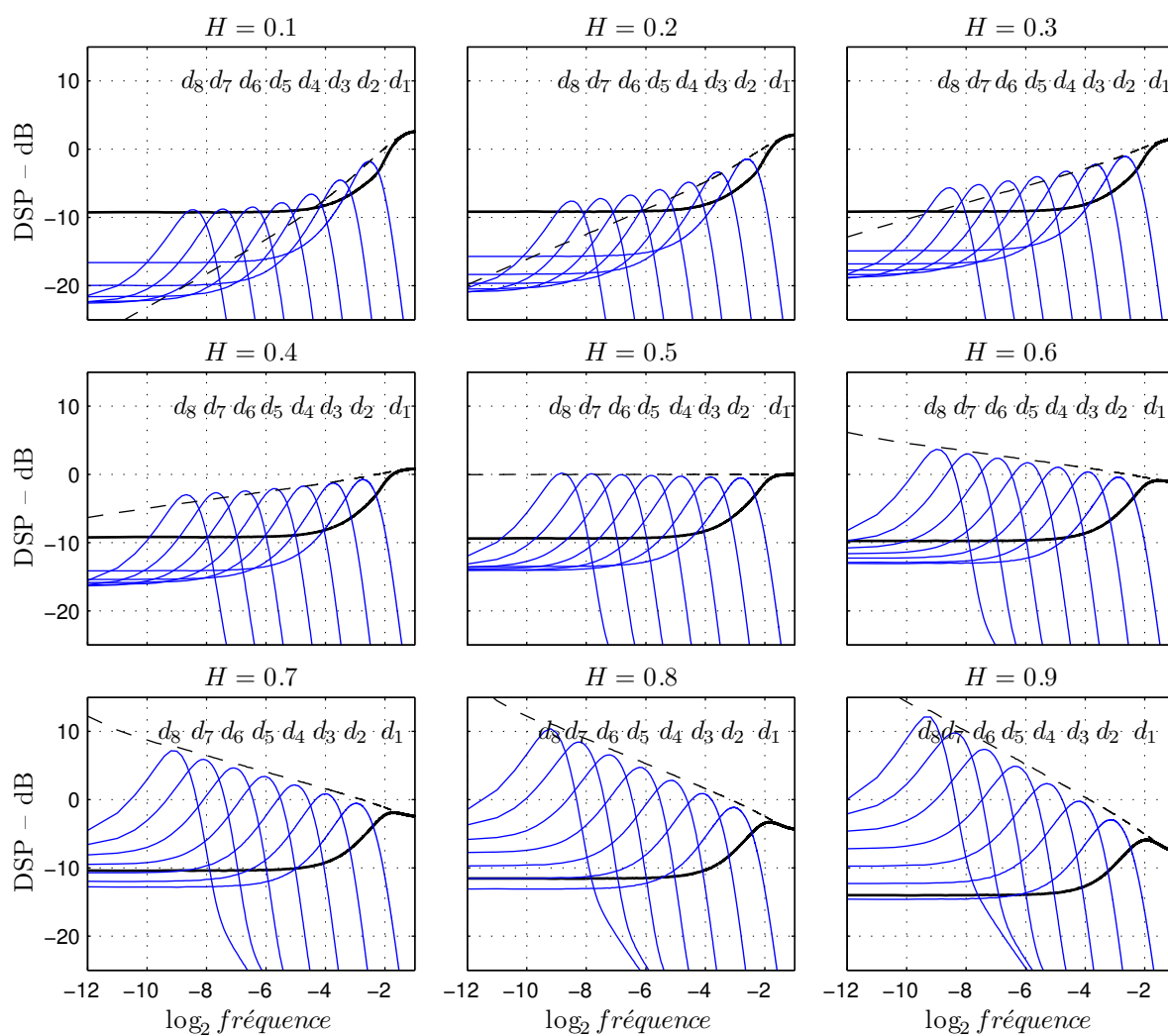


FIGURE 2.62 – Spectres des IMFs obtenus à l'aide de 10 itérations de tamisage par IMF.

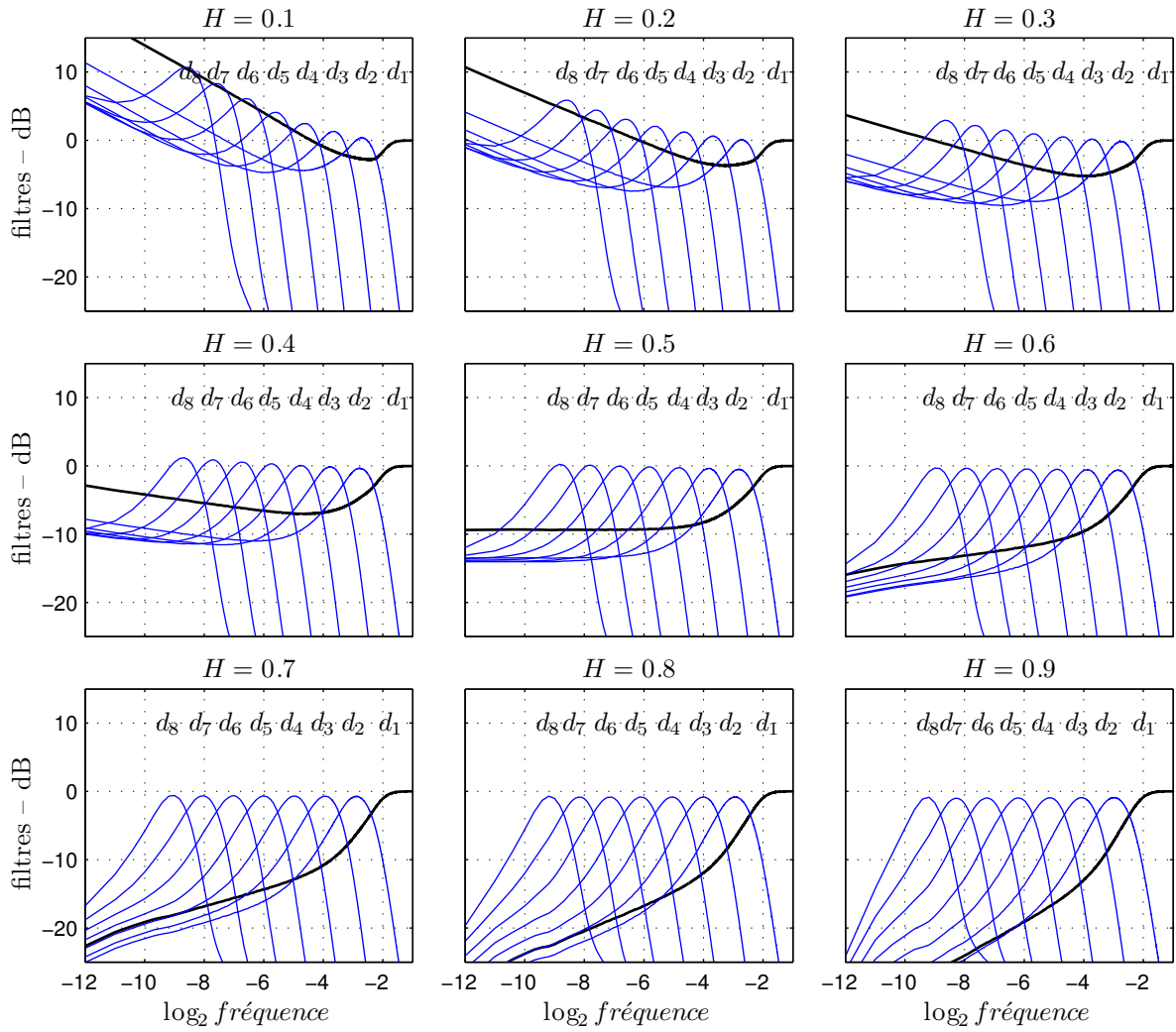


FIGURE 2.63 – Spectres des IMFs obtenus à l’aide de 10 itérations de tamisage par IMF divisés par le spectre du fGn.

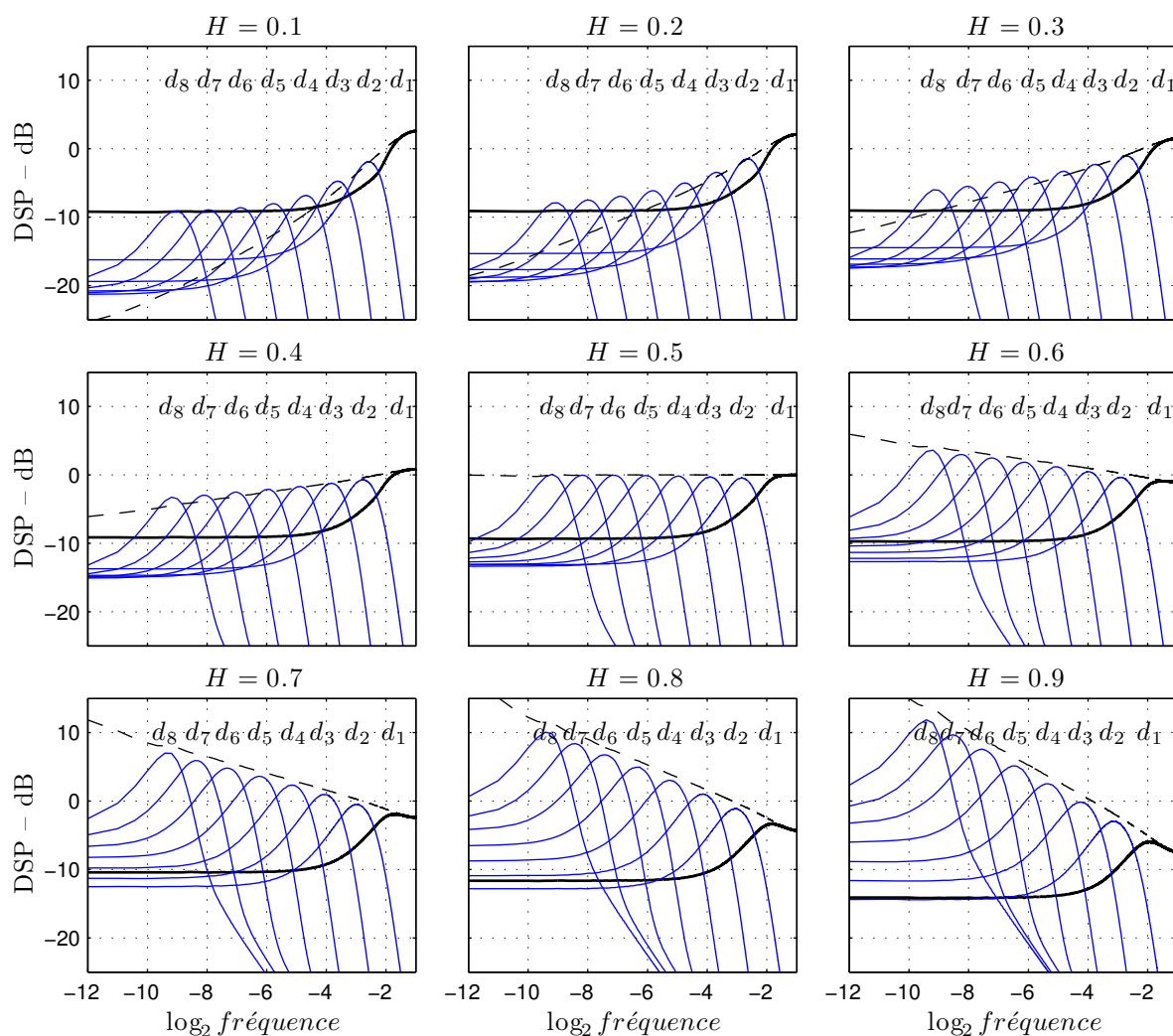


FIGURE 2.64 – Spectres des IMFs obtenus avec des nombres d'itérations contrôlés par le critère d'arrêt « local » (cf 5.2) avec $\epsilon = 0.1$.

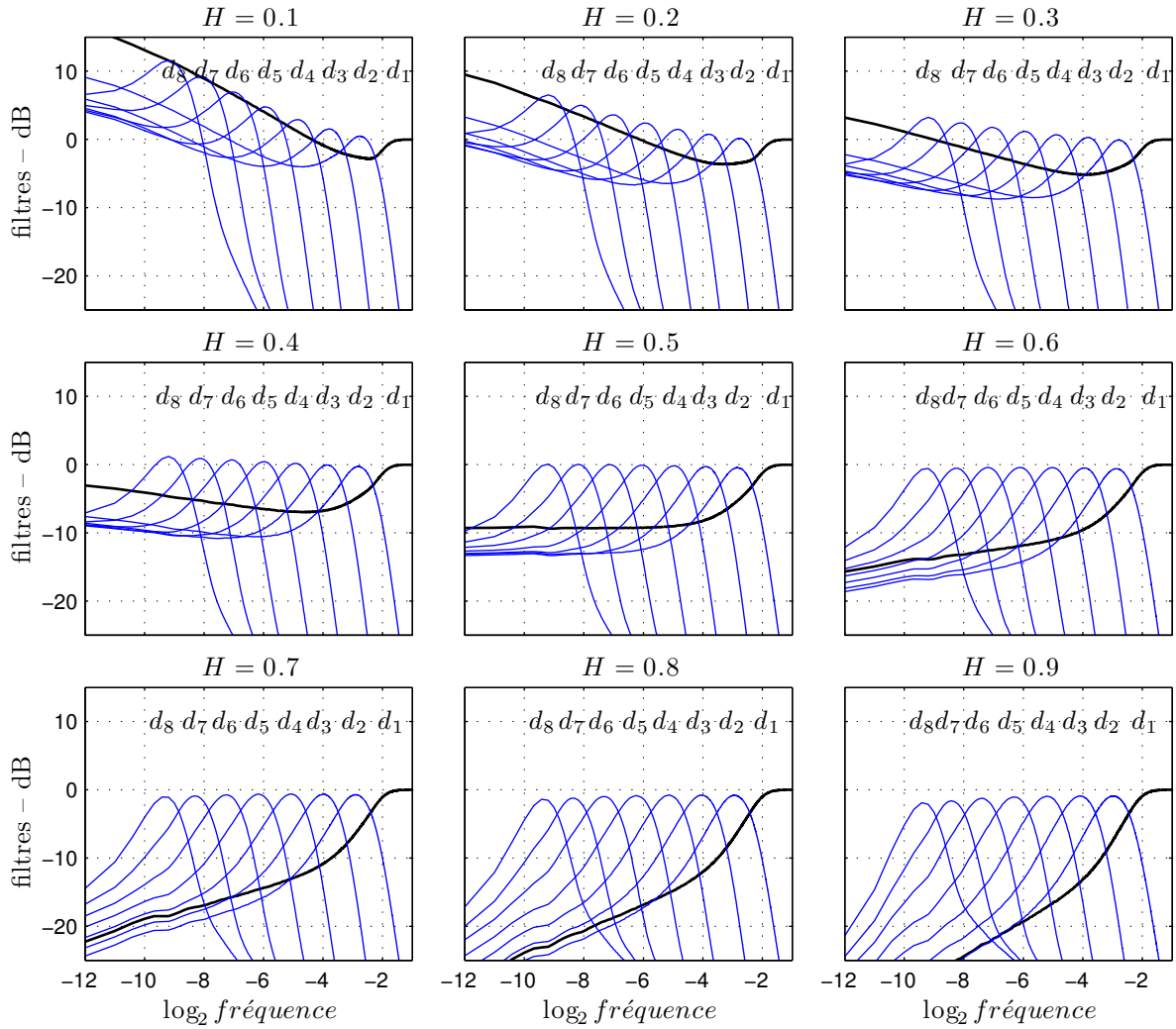


FIGURE 2.65 – Spectres des IMFs obtenus avec des nombres d'itérations contrôlés par le critère d'arrêt « local » (cf 5.2) avec $\epsilon = 0.1$, divisés par le spectre du fGn.

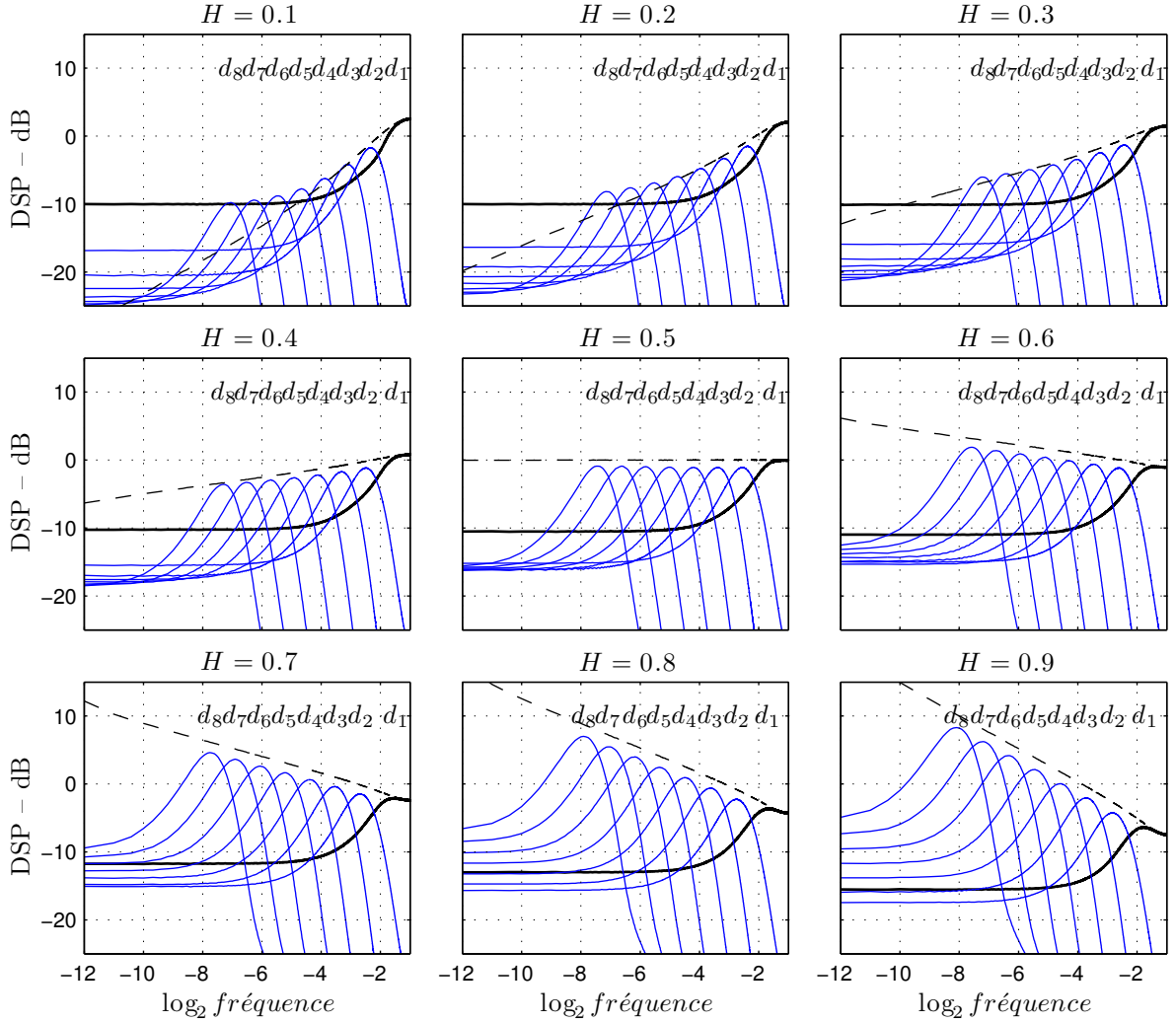


FIGURE 2.66 – Spectres des IMFs obtenus à l'aide de 50 itérations de tamisage par IMF.

composante, de manière à pouvoir comparer les deux résultats. En pratique, le nombre d'itérations pour le premier IMF avec ce critère d'arrêt est en moyenne de 10 avec une médiane à 7 et il décroît pour les IMFs suivants. Il vaut en moyenne typiquement 8, 7, 6, et 5 itérations pour les IMFs 2 à 5 et décroît de moins en moins vite par la suite. Dans les deux cas, qui sont extrêmement proches, on observe que comme dans le cas des bruits blancs étudiés précédemment, les spectres des IMFs suggèrent que le comportement de l'EMD s'apparente à celui d'un banc de filtres. Pour approfondir cette analogie, on s'intéresse également aux filtres équivalents à chacun des IMFs obtenus en divisant les densités spectrales de puissance des IMFs par la densité spectrale de puissance du fGn (cf Fig. 2.63 et Fig. 2.65). Enfin, pour observer l'influence du nombre d'itérations, on propose également les spectres et filtres équivalents pour 50 itérations de tamisage par IMF Fig. 2.66 et Fig. 2.67.

On observe ainsi que la modélisation de l'EMD par un banc de filtres semble pertinente pour $H > 0.5$ puisque les filtres équivalents dépendent relativement peu de H , ou plutôt les dépendances par rapport à H sont restreintes à des domaines où le gain est faible, ce qui limite fortement leur influence. Le banc de filtres équivalent est ainsi constitué d'un filtre passe-haut, correspondant au premier IMF, et d'une série de filtres passe-bandes, correspondant aux IMFs suivants. Le caractère passe-haut du filtre du premier IMF est cependant très relatif puisqu'il comporte également une contribution

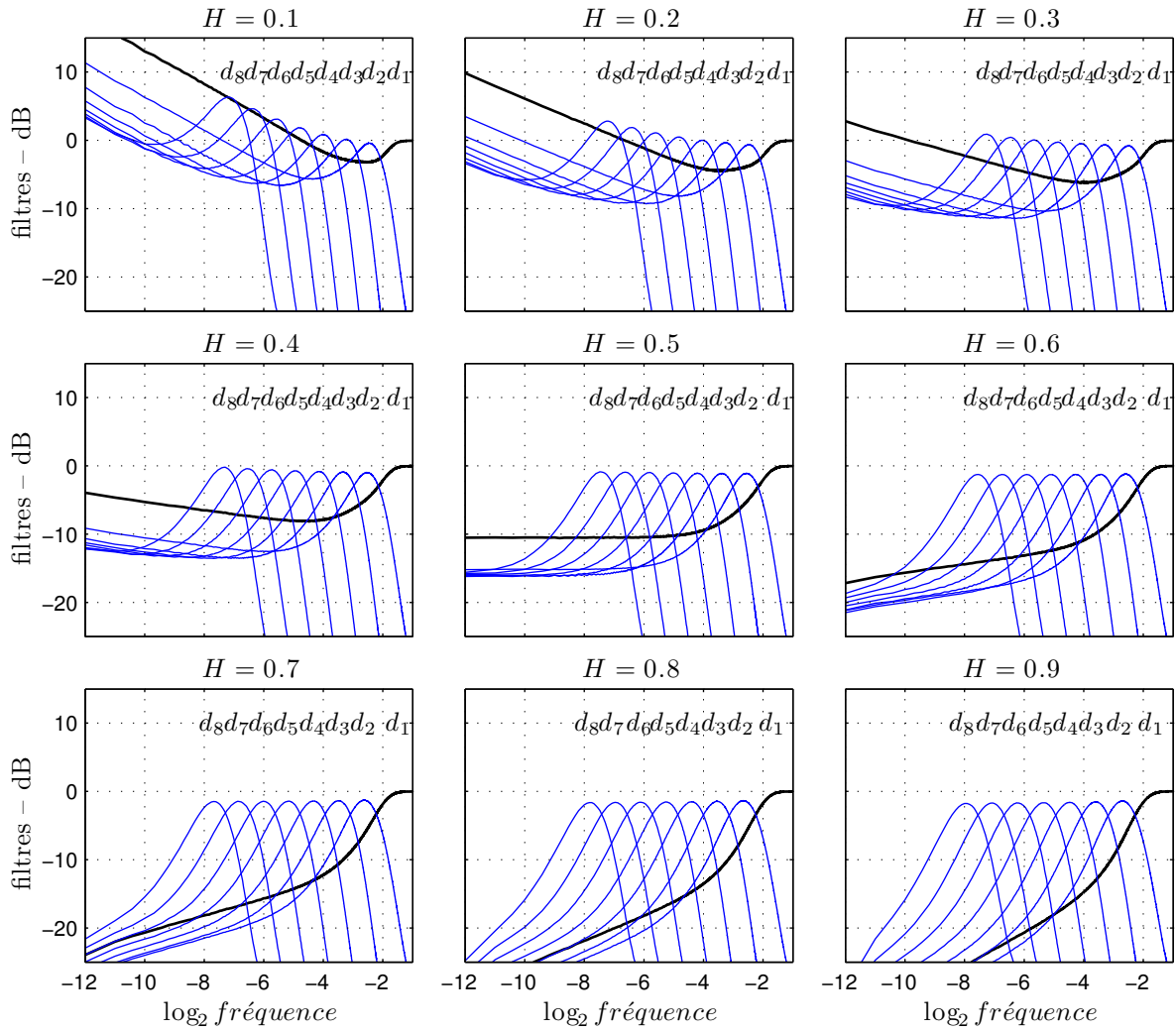


FIGURE 2.67 – Spectres des IMFs obtenus à l’aide de 50 itérations de tamisage par IMF divisés par le spectre du fGn.

relativement importante à basse fréquence, inférieure de seulement 10 dB à sa contribution à haute fréquence. Quant aux filtres équivalents aux autres IMFs, ils présentent la particularité d'être tous pratiquement identiques à une dilatation en fréquence près, qui est de plus toujours la même entre deux IMFs consécutifs. En première approximation, pour 10 itérations de tamisage par IMF, le filtre du premier IMF occupe la bande de fréquences $[0.25, 0.5]$ et les filtres passe-bandes des IMFs suivants les bandes $[1/2^{k+1}, 1/2^k]$, c'est-à-dire que le banc de filtres équivalent à l'EMD pour dix itérations de tamisage est quasi-dyadique. Pour 50 itérations, le banc de filtres est qualitativement très proche de celui observé pour 10 itérations, la principale différence étant que les IMFs occupent des bandes de fréquences plus étroites et plus rapprochées les unes des autres.

L'analogie avec un banc de filtres quasi-dyadique est beaucoup moins pertinente dans le cas où $H < 0.5$. Les filtres équivalents pour $H < 0.5$ dépendent plus fortement de H et sont très différents du cas $H > 0.5$: ils ne sont en particulier plus passe-bandes et leur gain à basse fréquence est même supérieur à un. Pour ce qui est de l'interprétation, l'origine de ce phénomène semble être la forme particulière du spectre du premier IMF. On observe en effet que celui-ci présente systématiquement une composante basses fréquences assez importante. L'amplitude de celle-ci pour $f \rightarrow 0$ est approximativement 10 dB en dessous de l'amplitude maximale du spectre et elle semble dépendre de l'amplitude du spectre du signal analysé dans la bande $[0.25, 0.5]$ mais pas dans la partie plus basses fréquences. Dans le cas des fGns avec $H < 0.5$, le spectre du signal tendant vers zéro en zéro contrairement à celui du premier IMF, il y a toujours une fréquence limite en dessous de laquelle le spectre de ce dernier dépasse celui du signal qui devient négligeable à fréquence nulle. De là, la première approximation étant la différence du signal et du premier IMF, elle contient nécessairement une composante basses fréquences plus importante que celle du signal qui est à son tour répartie dans les IMFs suivants. Le phénomène « d'amplification » à basses fréquences observé entre le signal et la première approximation se reproduit en fait par la suite pour les approximations suivantes mais l'ordre de grandeur d'amplification est nettement moins important, de l'ordre de celui observé pour les bruits blancs. Pour compléter l'analyse de ce phénomène d'amplification, on peut également s'intéresser à l'interspectre entre le premier IMF et la première approximation (cf Fig. 2.68). On constate en effet que celui-ci est systématiquement négatif à basses fréquences, ce qui signifie qu'une partie de l'information basses fréquences contenue dans le premier IMF et dans la première approximation se compense quand on en fait la somme pour retrouver le signal. Lorsque H s'approche de zéro, on observe que la cohérence entre le premier IMF et la première approximation (définie comme l'interspectre divisé par la moyenne géométrique des spectres) s'approche de -1 lorsque $f \rightarrow 0$ ce qui signifie que les informations à très basses fréquences contenues dans le premier IMF et dans la première approximation tendent à se compenser totalement.

2.2.4 Relations entre les spectres/filtres

Si on note $S_{k,H,n}^d(f)$ la densité spectrale de puissance de l'IMF $d_k(t)$ issu de la décomposition du fGn d'exposant H à l'aide de n itérations de tamisage, et $G_{k,H,n}^d(f)$ cette même densité spectrale de puissance divisée par la densité spectrale de puissance du fGn les relations entre les filtres équivalents des IMFs d'indice supérieur à 2 peuvent être formalisées par :

$$S_{k',H,n}^d(f) = \rho_{H,n}^{\alpha_{H,n}(k'-k)} S_{k,H,n}^d(\rho_{H,n}^{k'-k} f), \quad (2.187)$$

$$G_{k',H,n}^d(f) = G_{k,H,n}^d(\rho_{H,n}^{k'-k} f), \quad (2.188)$$

avec $k, k' \geq 2$ et $\rho_{H,n}$ et $\alpha_{H,n}$ des paramètres à ajuster. En utilisant l'approximation de la densité spectrale de puissance du fGn (2.186), on obtient immédiatement que nécessairement $\alpha_{H,n} = 2H - 1$. La valeur de $\rho_{H,n}$ en revanche dépend du signal et des paramètres de l'EMD. Plus précisément, on verra que cette valeur dépend essentiellement du nombre d'itérations de tamisage utilisé. Les relations (2.187) et (2.188) sont vérifiées empiriquement Fig. 2.69 et Fig. 2.70.

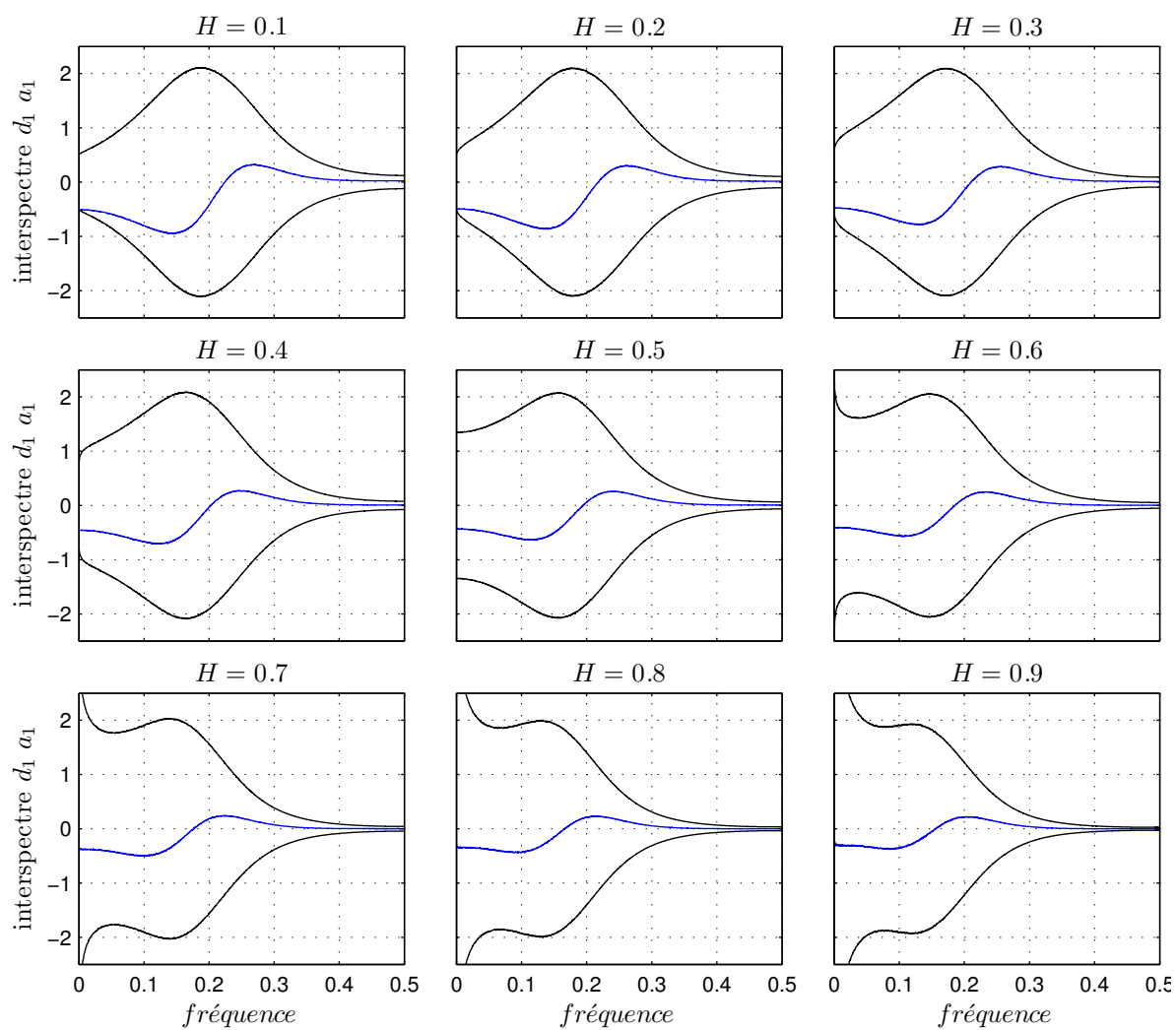


FIGURE 2.68 – Interspectre entre le premier IMF et la première approximation. Les courbes noires enveloppant les interspectres sont les moyennes géométriques des spectres du premier IMF et de la première approximation.

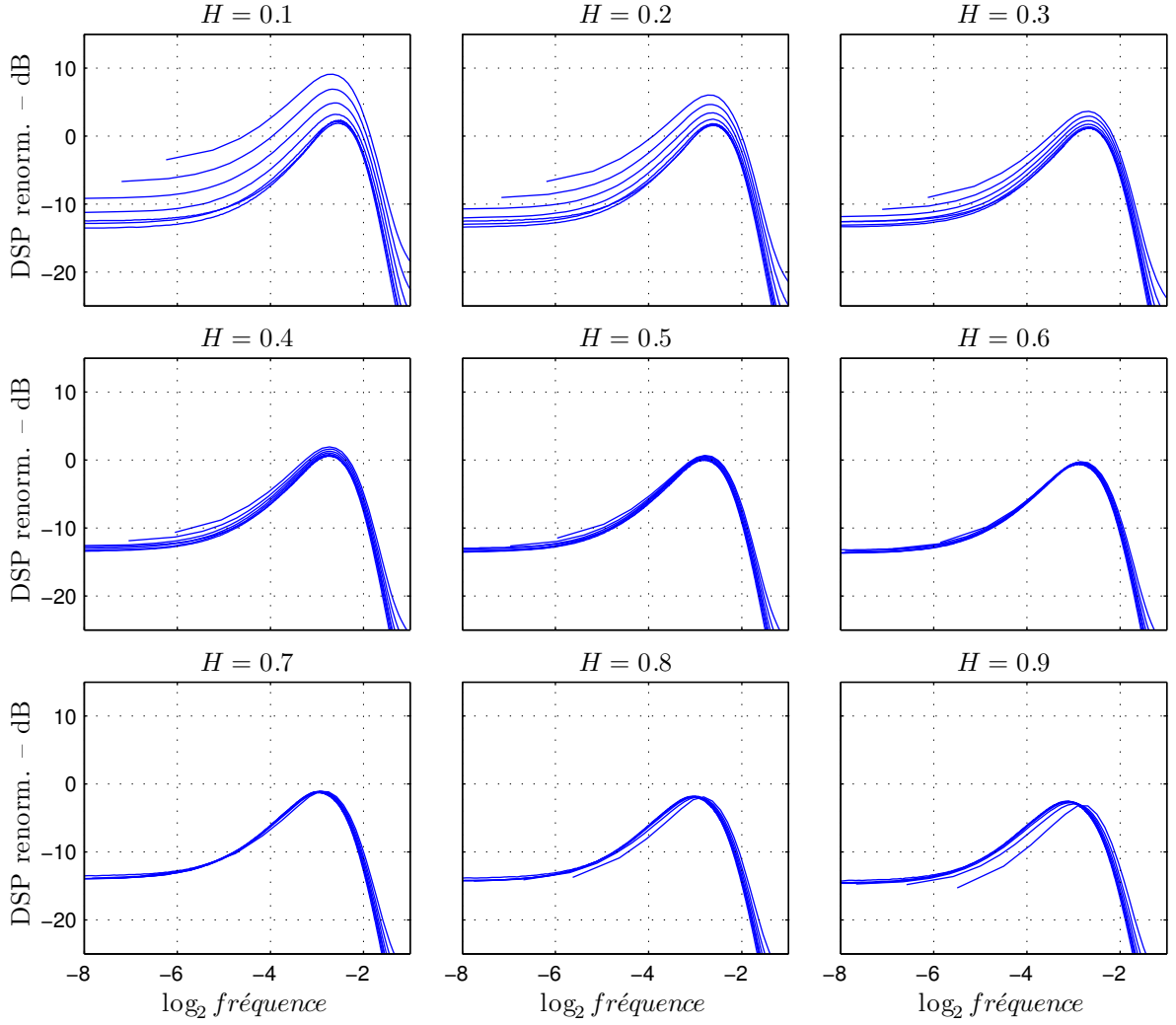


FIGURE 2.69 – Spectres des IMFs renormalisés selon l'équation (2.187). Le paramètre $\rho_{H,n}$ est estimé à partir de la relation (2.196). Les spectres renormalisés se superposent moins bien pour les petites valeurs de H , mais on peut constater que seule la renormalisation en amplitude est mauvaise.

On peut également établir des relations similaires pour les spectres et filtres équivalents des approximations $S_{k,H,n}^a(f)$ et $G_{k,H,n}^a(f)$:

$$S_{k',H,n}^a(f) = \rho_{H,n}^{\alpha_{H,n}(k'-k)} S_{k,H,n}^a(\rho_{H,n}^{k'-k} f), \quad (2.189)$$

$$G_{k',H,n}^a(f) = G_{k,H,n}^a(\rho_{H,n}^{k'-k} f), \quad (2.190)$$

où les paramètres $\rho_{H,n}$ et $\alpha_{H,n}$ doivent être les mêmes que pour les IMFs.

À partir des relations (2.187) et (2.189), on peut montrer que les variances des IMFs et des approximations $V_{H,n}^d[k]$ et $V_{H,n}^a[k]$ vérifient :

$$V_{H,n}^d[k] = C_d \rho_{H,n}^{2(H-1)k}, \quad (2.191)$$

$$V_{H,n}^a[k] = C_a \rho_{H,n}^{2(H-1)k}. \quad (2.192)$$

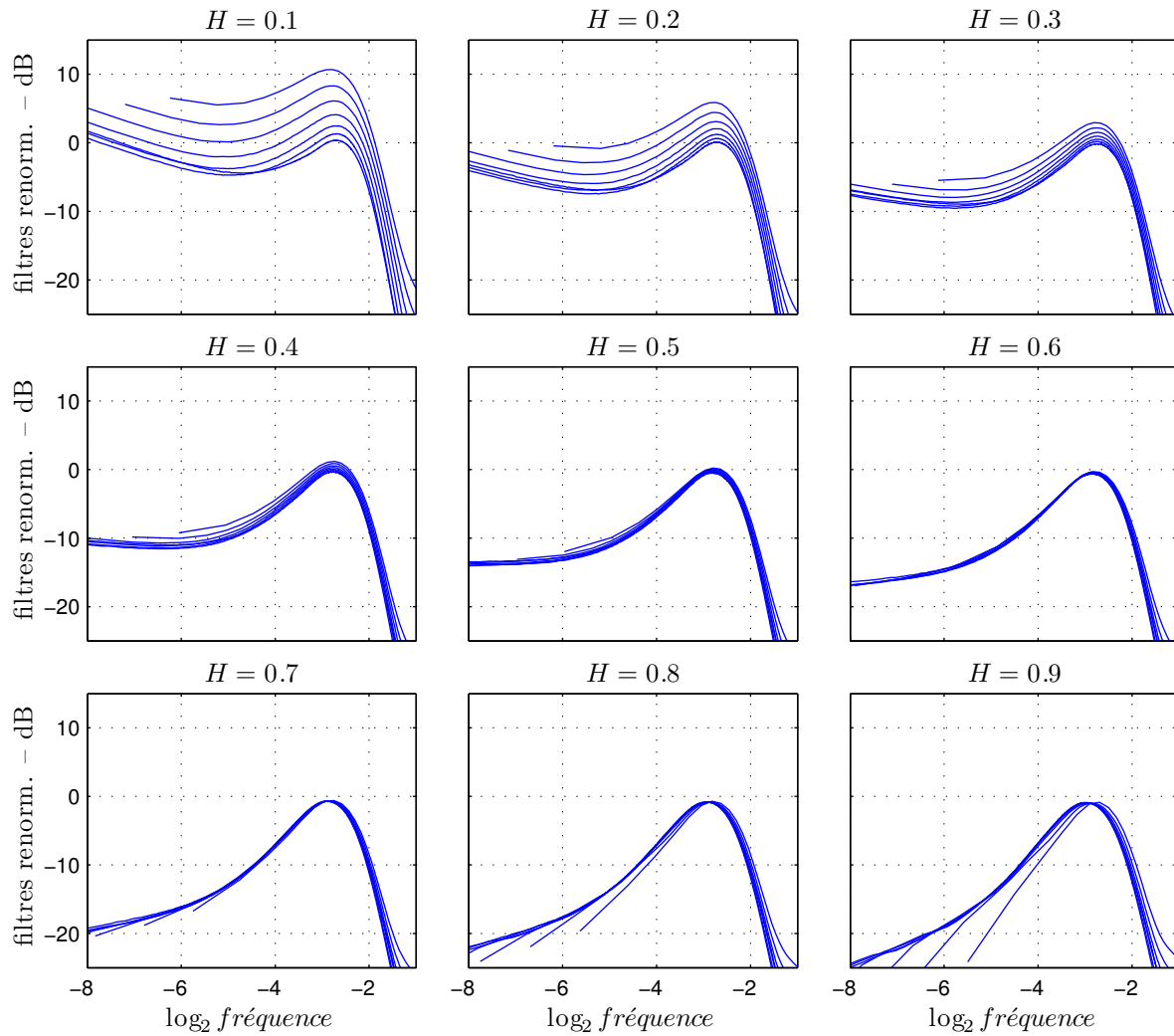


FIGURE 2.70 – Filtres équivalents des IMFs renormalisés selon l'équation (2.188). Le paramètre $\rho_{H,n}$ est estimé à partir de la relation (2.196). Tout comme pour les spectres, les filtres renormalisés se superposent assez mal pour les petites valeurs de H .

En effet, dans les deux cas,

$$V_{H,n}[k'] = \int_{-\frac{1}{2}}^{\frac{1}{2}} S_{k',H,n}(f) df, \quad (2.193)$$

$$= \rho_{H,n}^{(2H-1)(k'-k)} \int_{-\frac{1}{2}}^{\frac{1}{2}} S_{k,H,n}(\rho_{H,n}^{k'-k} f) df, \quad (2.194)$$

$$= \rho_{H,n}^{2(H-1)(k'-k)} V_{H,n}[k]. \quad (2.195)$$

Si l'on définit une période moyenne pour chaque mode $\bar{T}_H[k]$, la relation (2.187) permet d'écrire la relation suivante entre les périodes moyennes pour $k, k' > 1$:

$$\bar{T}_H[k'] = \rho_H^{k'-k} \bar{T}_H[k]. \quad (2.196)$$

On pourrait la démontrer à partir de (2.187) dans le cas où la période moyenne est définie à l'aide d'une moyenne pondérée de la densité spectrale de puissance comme par exemple

$$T = \frac{\int S(f) df}{\int f S(f) df}, \quad (2.197)$$

ou encore

$$T = \frac{\int \frac{1}{f} S(f) df}{\int S(f) df}. \quad (2.198)$$

Dans le cas d'IMFs une solution efficace consiste à définir la période moyenne comme deux fois l'écart moyen entre deux passages à zéro consécutifs :

$$T = 2 \cdot \frac{\text{nombre de passages à zéro} - 1}{\text{distance entre le premier et le dernier passage à zéro}}. \quad (2.199)$$

L'intérêt de cette définition est notamment de permettre une estimation efficace de la période moyenne d'un signal sans avoir besoin de connaître sa densité spectrale de puissance. De plus, il s'avère que cet estimateur est assez nettement meilleur que les deux autres proposés (2.197) et (2.198), du moins pour les IMFs issus d'un fGn (cf Fig. 2.71).

En combinant les relations (2.191), (2.192) et (2.196), on obtient les relations entre variances des IMFs ou approximations et périodes moyennes des IMFs :

$$V_H^d[k] = C_d' (T_H[k])^{2(H-1)}, \quad (2.200)$$

$$V_H^a[k] = C_d' (T_H[k])^{2(H-1)}. \quad (2.201)$$

Les diverses relations peuvent être vérifiées empiriquement, Fig. 2.72 pour (2.196) et Fig. 2.73 pour les relations (2.191), (2.192), (2.200) et (2.201). Bien entendu, les relations (2.187) et (2.188) n'étant pas vérifiées en pratique pour les petites valeurs de H , il en est de même des relations (2.191) et (2.200). En revanche, la relation (2.196) est très bien vérifiée pour toutes les valeurs de H , ce qui est cohérent avec le fait que les spectres renormalisés de la figure Fig. 2.69 ne diffèrent pratiquement que par leur amplitude.

On a vu précédemment que le nombre d'itérations avait une influence sur le facteur d'auto-similarité entre les spectres des IMFs. Plus précisément, la valeur du facteur d'auto-similarité $\rho_{H,n}$ (cf Fig. 2.74) dépend du nombre d'itérations n et dans une moindre mesure de la valeur de H , donc du signal analysé. De ≈ 2.5 pour 1 itération à ≈ 1.7 pour 100 itérations, $\rho_{H,n}$ décroît en fonction de n suivant à peu près une loi de puissance $\rho_{H,n} \propto n^{-0.08}$ au-delà des quelques premières itérations. $\rho_{H,n}$ croît aussi légèrement avec H , de ≈ 1.95 à ≈ 2.1 entre $H = 0.1$ et $H = 0.9$ pour 10 itérations. De plus, le rapport $\rho_{H,n}/\rho_{H',n}$ ne dépend pratiquement pas de n passées les 2-3 premières itérations.

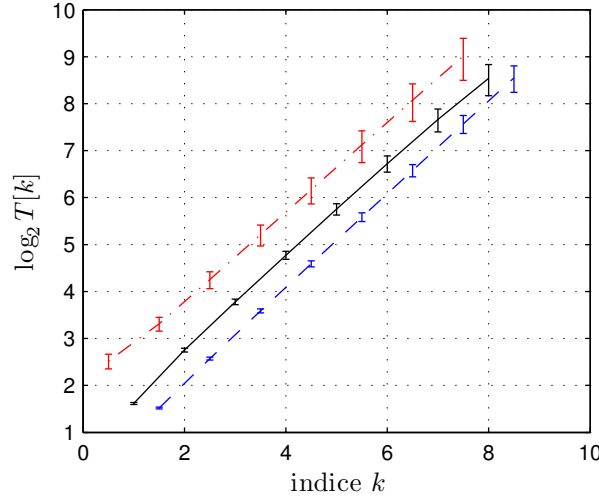


FIGURE 2.71 – Performances de trois estimateurs de la période moyenne dans le cas des IMFs issus d'un bruit blanc gaussien : (2.199) en tirets, (2.197) en trait plein et (2.198) en tiret-point. Les trois courbes ont été légèrement décalées en abscisse pour améliorer la lisibilité. Les trois estimateurs ont une espérance satisfaisante mais l'estimateur basé sur les passages à zéro (2.199) présente systématiquement une meilleure variance que les deux autres. On observe des résultats similaires pour des fGn d'exposants de Hurst différents.

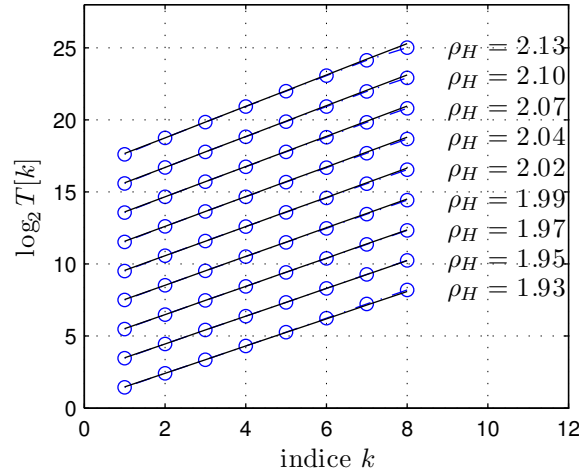


FIGURE 2.72 – Vérification empirique de la relation (2.196) entre les périodes moyennes des IMFs (obtenus avec 10 itérations de tamisage par IMF). Les différentes courbes ont été décalées en ordonnée pour améliorer la lisibilité. La valeur de ρ_H à côté de chaque courbe est estimée par régression linéaire sur la courbe correspondante.

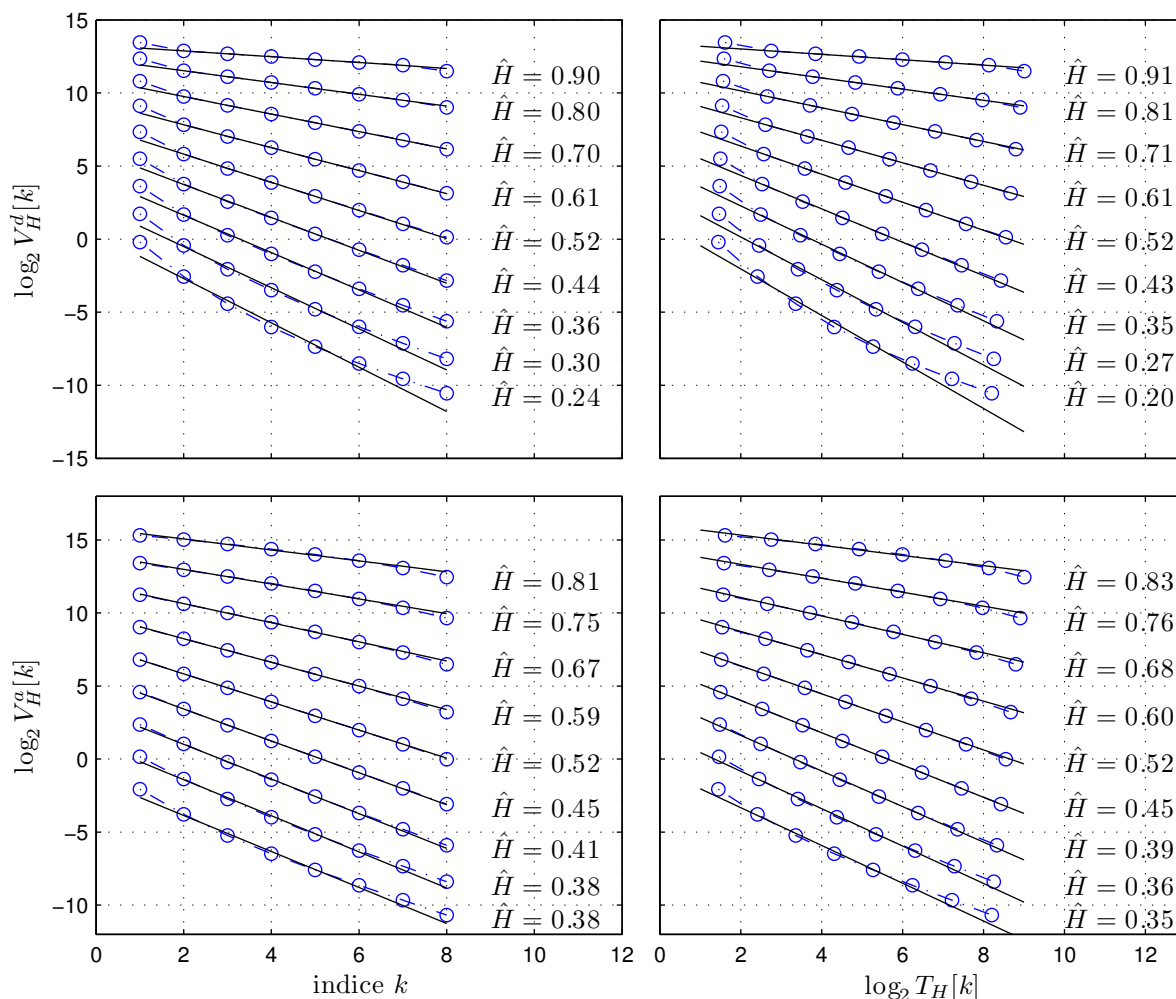


FIGURE 2.73 – Vérification empirique des relations (2.191), (2.192), (2.200) et (2.201). Les variances et périodes moyennes représentées sont des estimations des variances et périodes moyennes des IMFs et approximations (obtenus avec 10 itérations de tamisage par IMF) réalisées à partir des 100000 réalisations de bruit. Les différentes courbes ont été artificiellement décalées en ordonnée pour améliorer la lisibilité. La valeur $\hat{H}$ à côté de chaque courbe correspond à $1 + p/2$ où p est la pente de la courbe correspondante estimée par régression linéaire sur tous les points sauf celui correspondant au premier IMF ou à la première approximation.

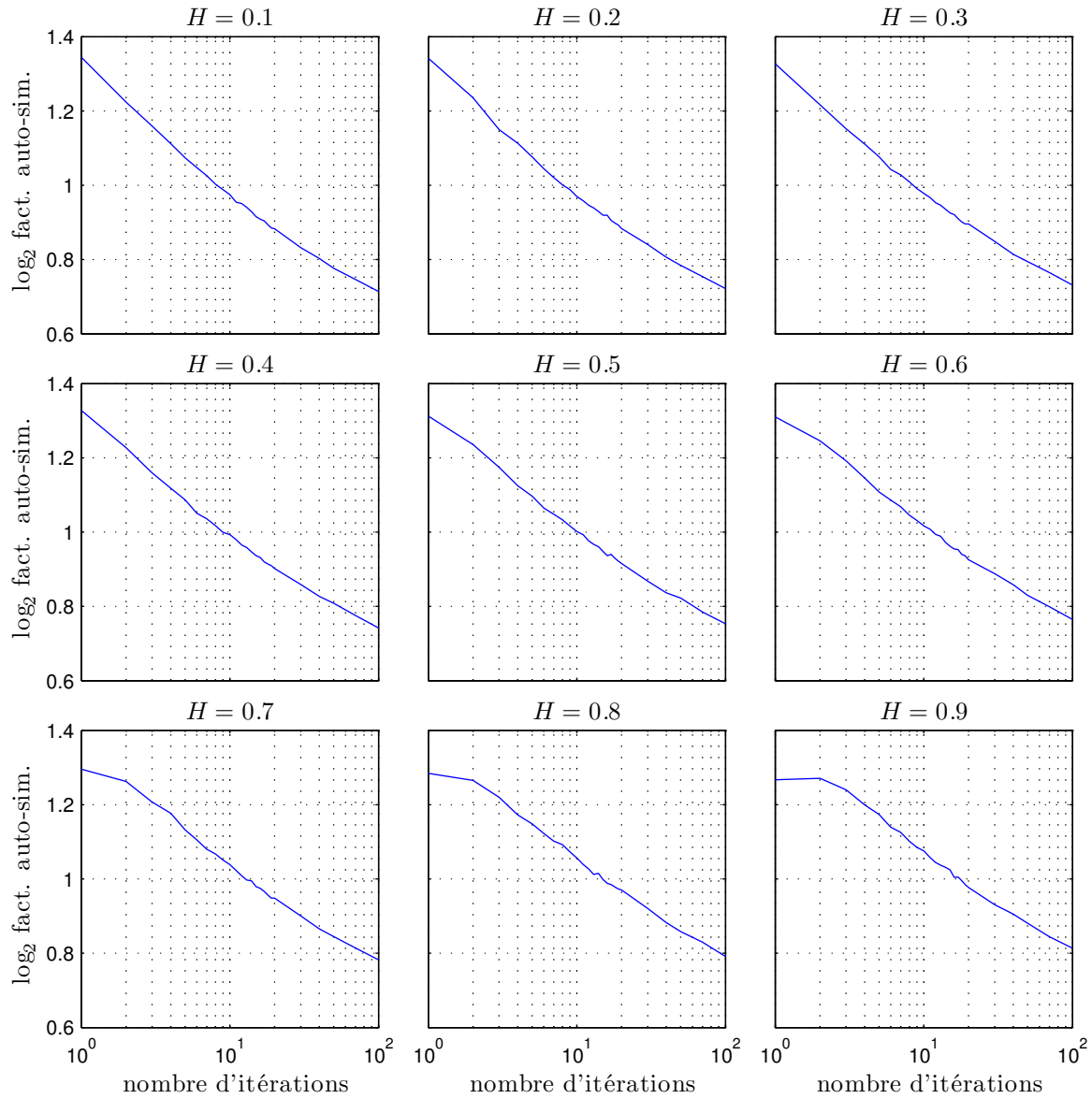


FIGURE 2.74 – Évolution du facteur d'auto-similarité du banc de filtres $\rho_{H,n}$ en fonction du nombre d'itérations n et de l'exposant de Hurst H .

2.2.5 Application à l'estimation de l'exposant d'une loi d'échelle

L'étude précédente a montré que pour divers cas de bruits large bande, l'EMD se comporte pratiquement comme un banc de filtres quasi-dyadique, de manière similaire à la transformée en ondelettes dyadique. On se propose maintenant d'utiliser cette parenté avec la transformée en ondelettes discrète pour imiter l'application de cette dernière à l'estimation de l'exposant d'une loi d'échelle.

2.2.5.1 Processus présentant une loi d'échelle Un processus aléatoire stationnaire X présente un comportement de loi d'échelle ssi sa densité spectrale de puissance $S_X(f)$ prend la forme :

$$S_X(f) \simeq c|f|^{-\alpha}, \quad \alpha \in]-1, 1[. \quad (2.202)$$

Les fGns présentés précédemment, du fait de l'approximation $S_H(f) \sim C\sigma^2|f|^{1-2H}$, font partie de cette famille de processus. Dans la suite, on identifiera l'exposant de la loi d'échelle α avec la fonction de l'exposant de Hurst $1 - 2H$.

Lorsque $\alpha \in]0, 1[$, on parle de processus à longue mémoire du fait de la divergence de type “ $1/f$ ” de leur densité spectrale de puissance.

2.2.5.2 Principe Dans le cas d'un fGn, on a vu que les variances des IMFs et leurs périodes moyennes vérifiaient les relations (2.191) et (2.200). De plus, on a également constaté que le paramètre ρ dans (2.191) dépend essentiellement du nombre d'itérations de tamisage et vaut en particulier environ 2 pour 10 itérations, ce qui fournit l'approximation de (2.191) :

$$V_{H,n}^d[k] \simeq C_d 2^{2(H-1)k}, \quad (2.203)$$

valable quand le nombre d'itérations n est proche de 10.

Si on considère maintenant une réalisation donnée de fGn, on peut calculer son EMD à l'aide de 10 itérations de tamisage, et estimer les variances des IMFs $\hat{V}^d[k]$ ainsi que leurs périodes moyennes $\hat{T}[k]$. De là, les équations (2.203) et (2.200) permettent a priori de remonter à son exposant de Hurst à l'aide d'une modélisation par régression linéaire sur $\log_2 \hat{V}^d[k]$ vs. k ou sur $\log_2 \hat{V}^d[k]$ vs. $\log_2 \hat{T}[k]$. Si l'on suppose de plus ces relations valables dès lors que le signal présente un comportement de loi d'échelle, on dispose alors de méthodes générales pour estimer l'exposant de loi d'échelle d'un signal à l'aide de l'EMD. On pourrait également proposer des méthodes analogues utilisant les variances des approximations à la place des variances des IMFs. En pratique, les relations étant moins bien vérifiées (cf Fig. 2.73) pour les approximations que pour les IMFs, les performances de tels estimateurs basés sur les approximations sont systématiquement nettement moins intéressantes.

2.2.5.3 Statistiques Pour estimer proprement l'exposant de la loi d'échelle, il faut tout d'abord étudier les statistiques des estimées des variances et des périodes moyennes. Celles-ci ont été estimées empiriquement à l'aide d'histogrammes sur 100000 réalisations de fGn d'exposant de Hurst variant de 0.1 à 0.9 par pas de 0.1.

Les résultats montrent que les statistiques des variances (cf Fig. 2.75, Fig. 2.76 et Fig. 2.77) sont assez bien modélisées par des lois Gamma dont la forme générale est :

$$\mathcal{P}_{Gamma}(x; \alpha, \beta) = \frac{1}{\beta^\alpha \Gamma(\alpha)} x^{\alpha-1} e^{-\frac{x}{\beta}}, \quad (2.204)$$

où $\Gamma(x)$ est la fonction Gamma d'Euler. Ces modélisations sont par ailleurs évaluées par des tests de Kolmogorov–Smirnov (cf table Tab. 2.1). Les résultats montrent que les lois Gamma ne modélisent vraiment bien que les distributions des variances des premiers IMFs et que les distributions empiriques s'éloignent d'autant plus des modélisations que H est grand. En pratique, il semble que

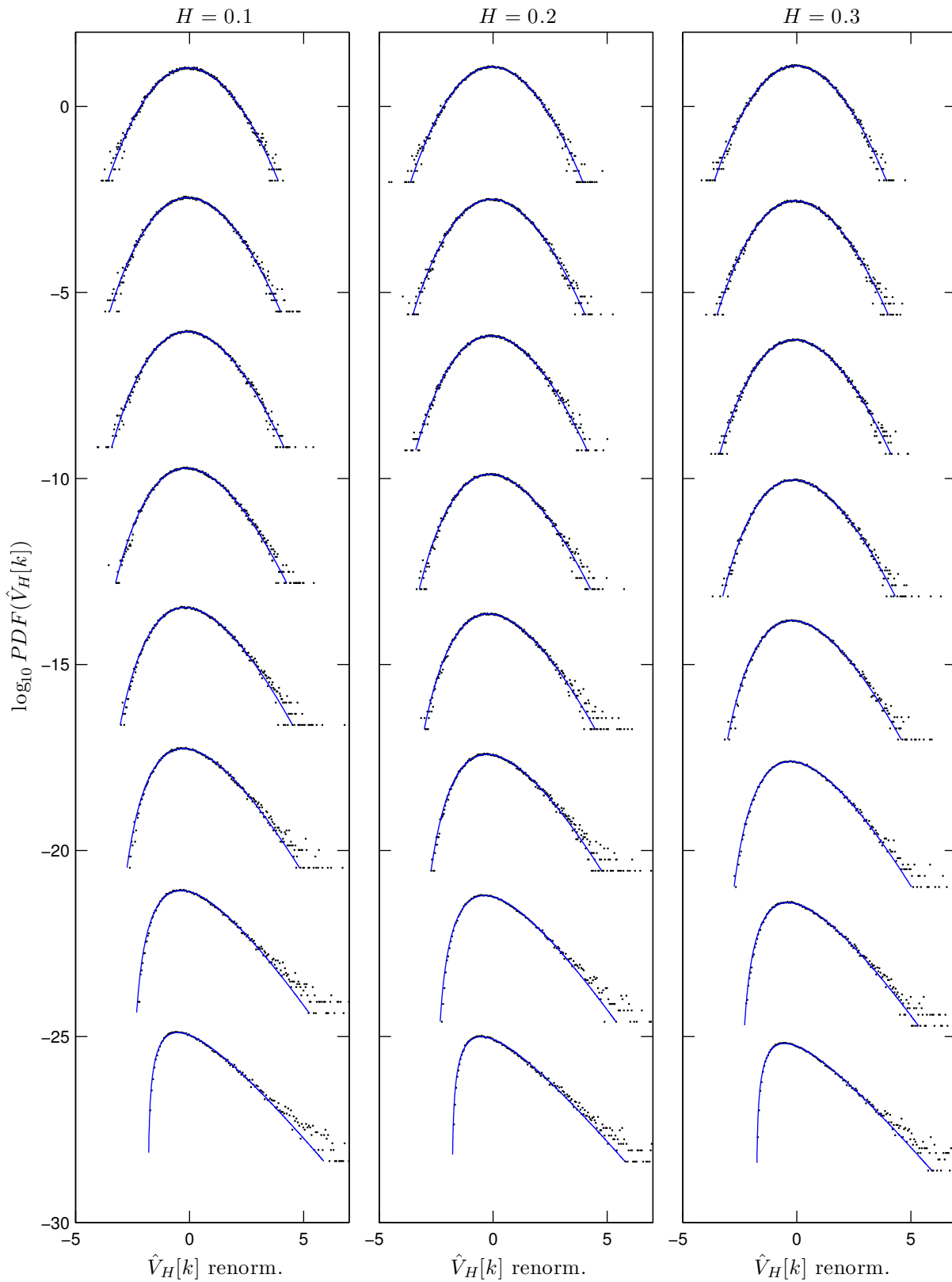


FIGURE 2.75 – Modélisations des distributions des variances des IMFs par des lois Gamma. Les distributions sont déformées en abscisse de manière à être centrées en zéro et avoir une variance unitaire.

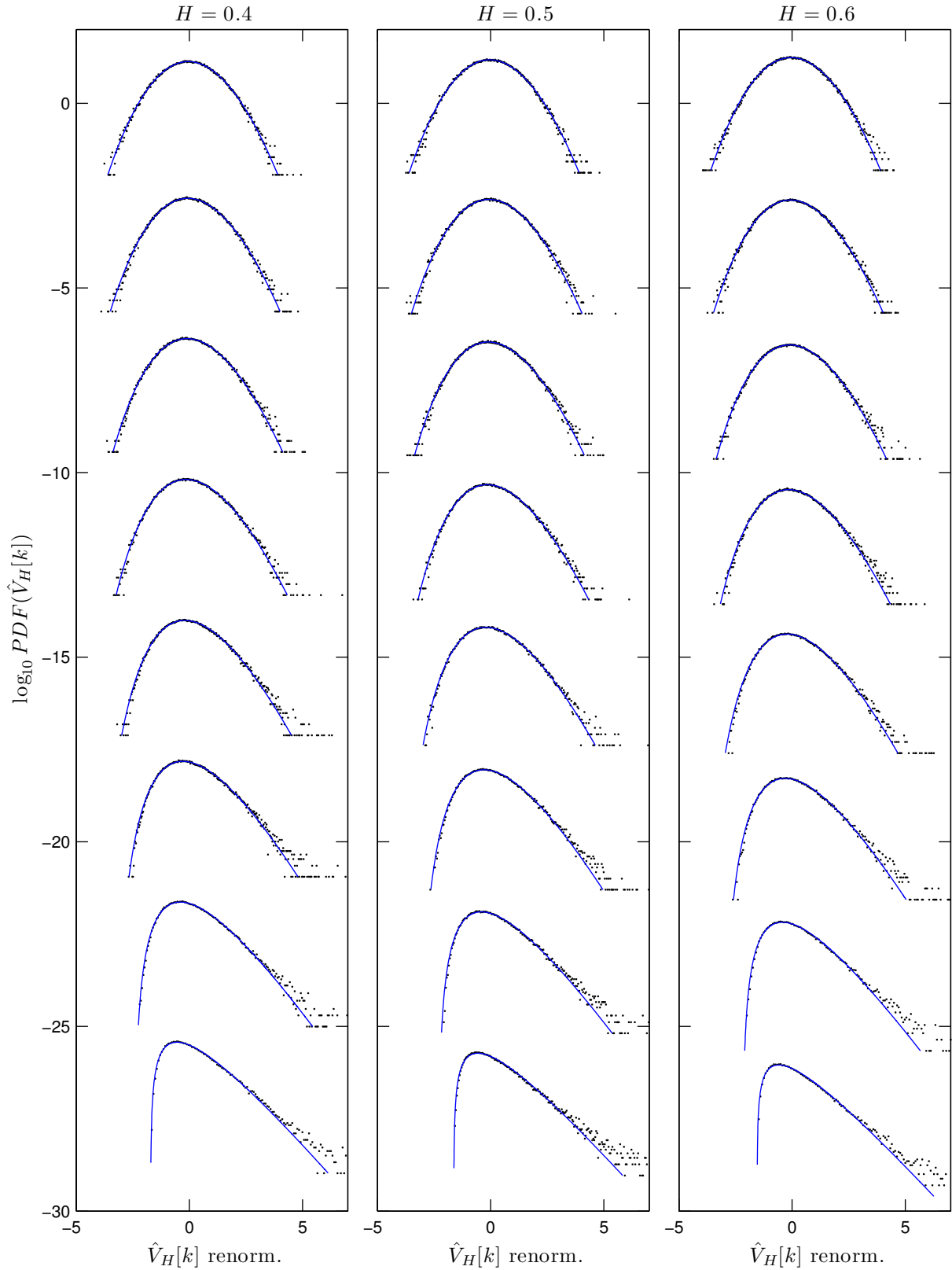


FIGURE 2.76 – Modélisations des distributions des variances des IMFs par des lois Gamma. Les distributions sont déformées en abscisse de manière à être centrées en zéro et avoir une variance unitaire.

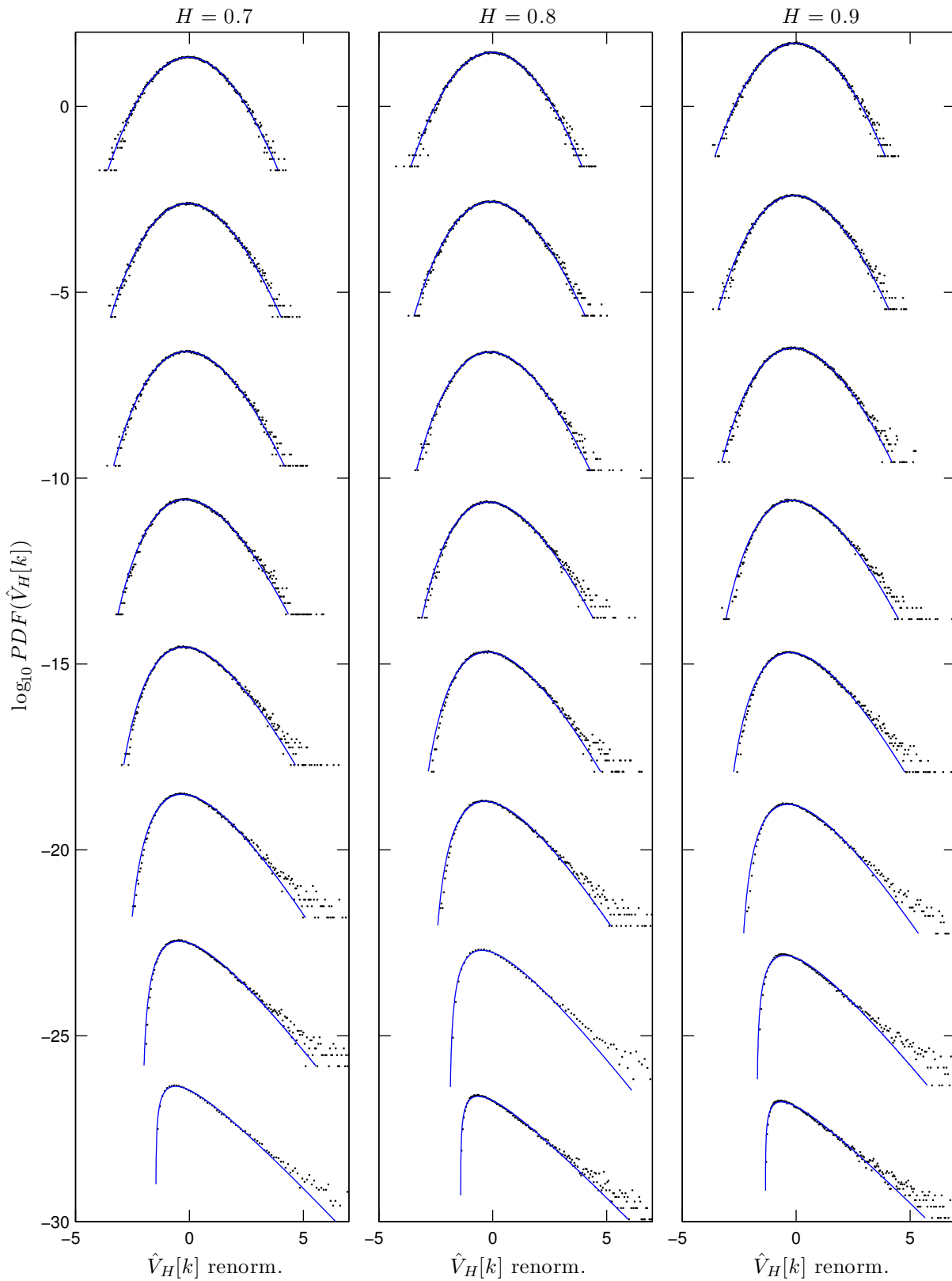


FIGURE 2.77 – Modélisations des distributions des variances des IMFs par des lois Gamma. Les distributions sont déformées en abscisse de manière à être centrées en zéro et avoir une variance unitaire.

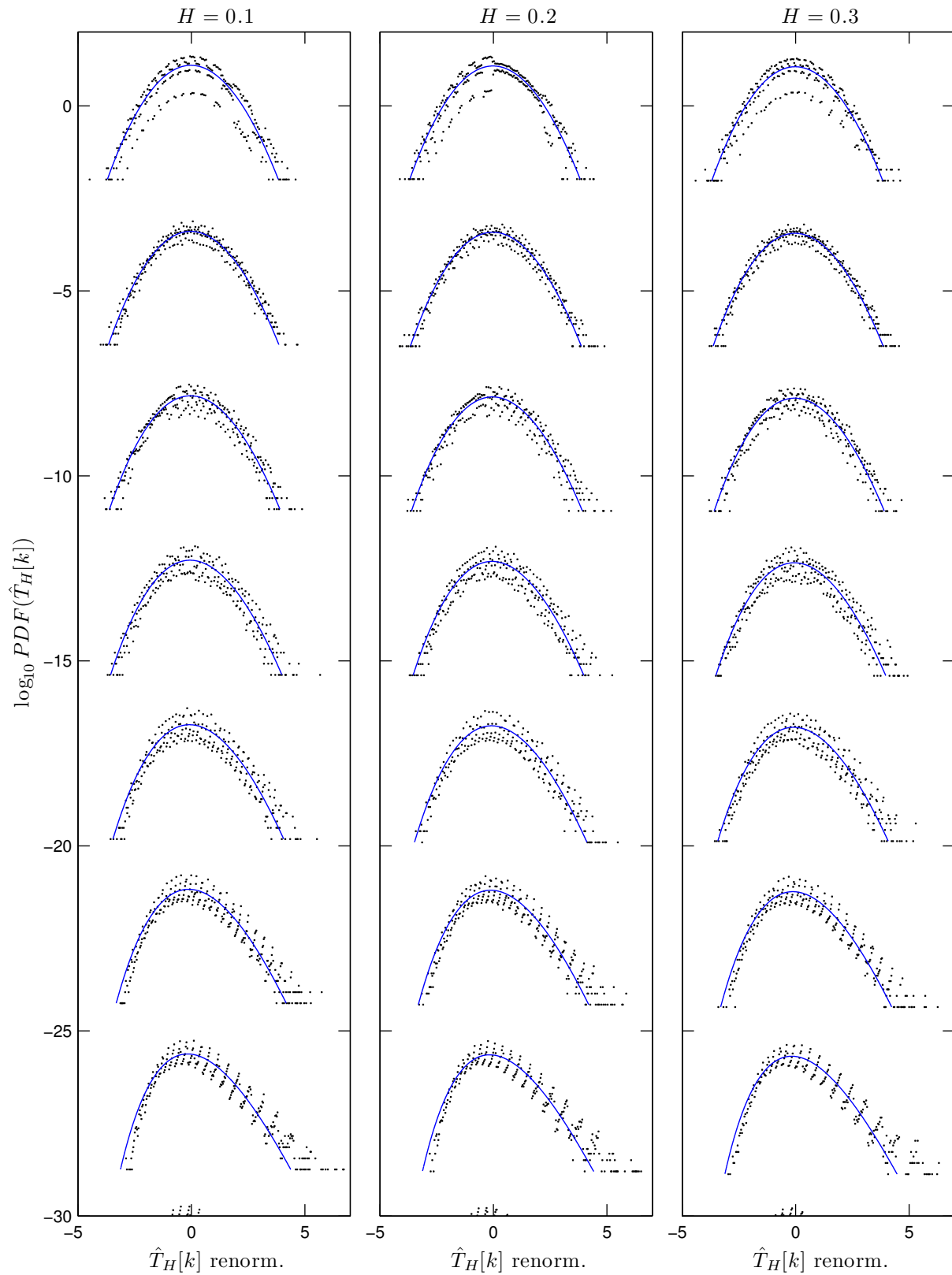


FIGURE 2.78 – Modélisations des distributions des périodes moyennes (2.199) des IMFs par des lois log-normales. Les distributions sont déformées en abscisse de manière à être centrées en zéro et avoir une variance unitaire.

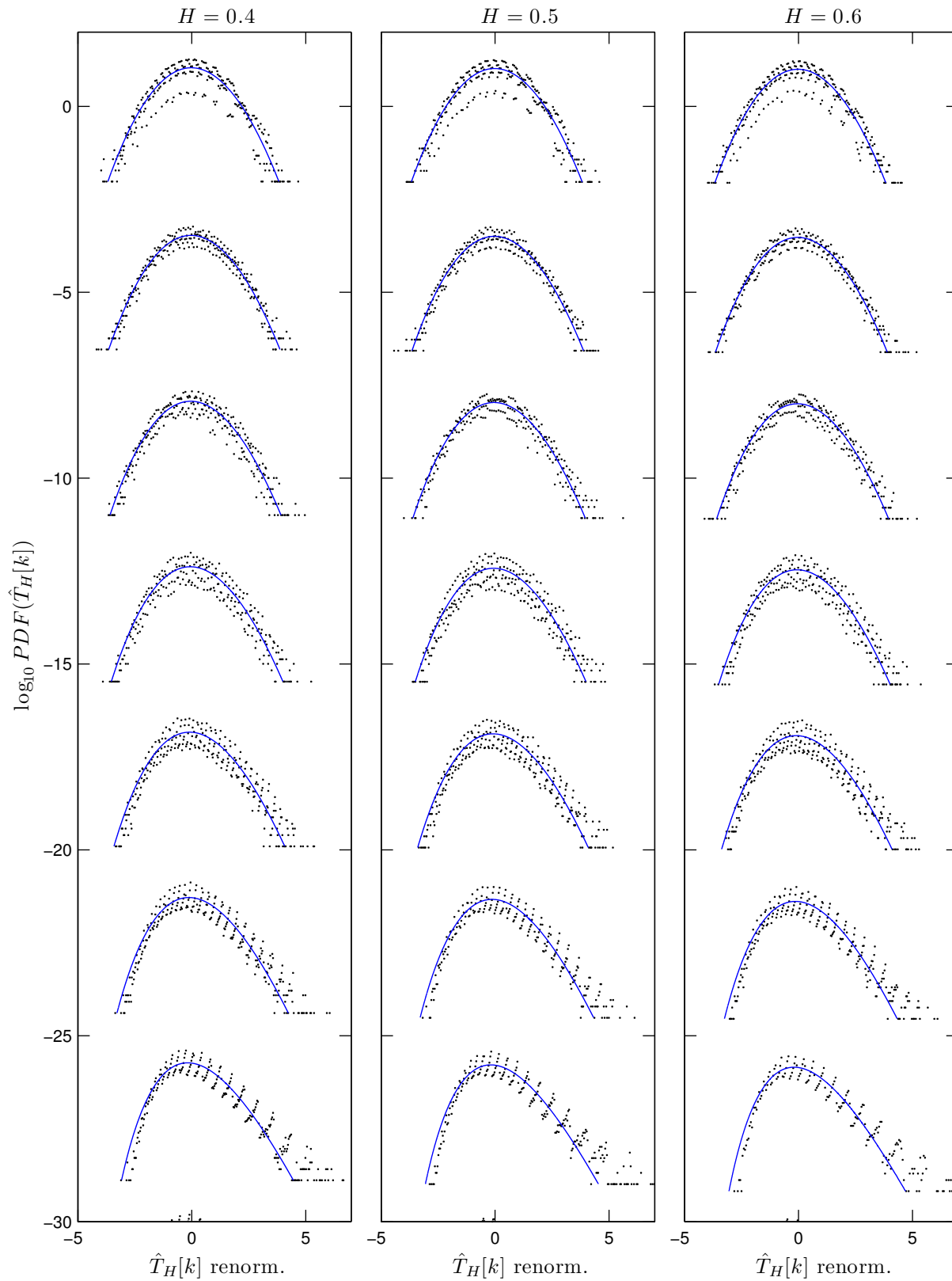


FIGURE 2.79 – Modélisations des distributions des périodes moyennes (2.199) des IMFs par des lois log-normales. Les distributions sont déformées en abscisse de manière à être centrées en zéro et avoir une variance unitaire.

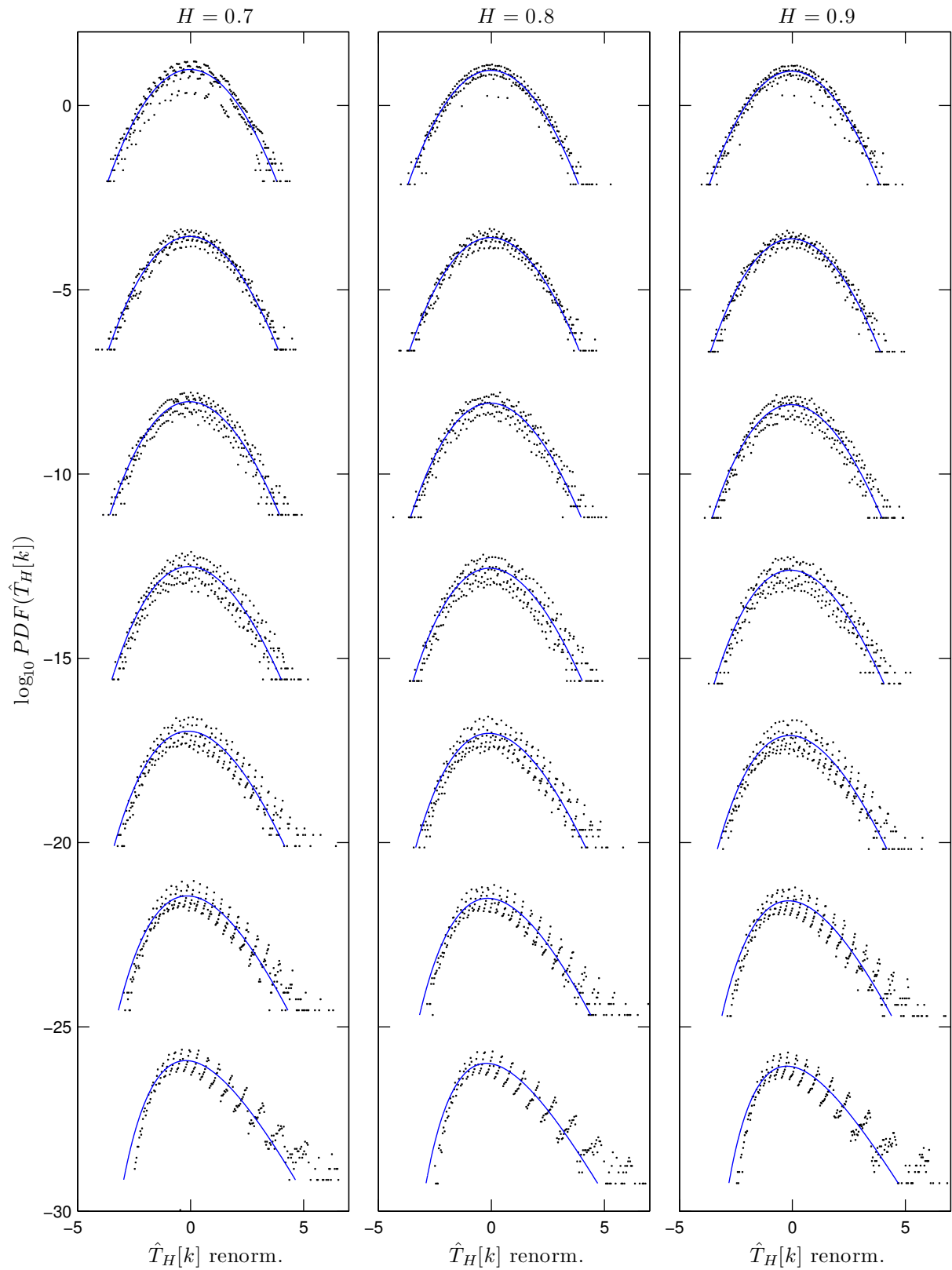


FIGURE 2.80 – Modélisations des distributions des périodes moyennes (2.199) des IMFs par des lois log-normales. Les distributions sont déformées en abscisse de manière à être centrées en zéro et avoir une variance unitaire.

# IMF	H=0.1	H=0.2	H=0.3	H=0.4	H=0.5	H=0.6	H=0.7	H=0.8	H=0.9
1	✓	✓	✓	✓	✓	✓	✓	✓	✓
2	✓	×	✓	✓	✓	×	✓	×	×
3	✓	×	✓	×	✓	×	×	×	×
4	✓	✓	✓	×	×	×	×	×	×
5	×	×	✓	×	×	×	×	×	×
6	×	×	×	×	×	×	×	×	×
7	×	×	×	×	×	×	×	×	×
8	×	×	×	×	×	×	×	×	×

TABLE 2.1 – Résultats des tests de Kolmogorov–Smirnov mesurant l’adéquation des distributions des variances des IMFs et de leurs modélisations par des lois Gamma. Les tests ont été réalisés à partir d’une population de 100000 réalisations de fGns indépendante de celle utilisée pour les modélisations. Le niveau de confiance des tests est de 95 %.

les distributions des variances des IMFs aient souvent des queues plus lourdes que les lois Gamma. Par la suite, on supposera tout de même que les variances des IMFs ont des distributions Gamma puisque c’est la meilleure approximation dont on dispose.

Concernant les périodes moyennes, leurs distributions n’ont pas pu être identifiées aussi bien que celles des variances. La tâche est en fait un peu délicate dans la mesure où ces distributions sont intrinsèquement discrètes puisque l’estimateur (2.199) ne peut de fait prendre que des valeurs rationnelles. Dans le contexte de l’estimation d’exposant de loi d’échelle cependant, la structure fine de ces distributions n’a pas vraiment d’intérêt par rapport à leur allure globale que l’on peut tenter de modéliser à l’aide des distributions continues usuelles. Parmi celles-ci, on observe que les lois log-normales fournissent les meilleures approximations même si elles ne modélisent pas très bien les queues des distributions empiriques (cf Fig. 2.78, Fig. 2.79 et Fig. 2.80).

Enfin, périodes moyennes et variances des IMFs dépendent généralement légèrement les unes des autres comme on peut le constater Fig. 2.81 et Fig. 2.82 où sont représentées les distributions conjointes des logarithmes en base 2 des deux quantités.

2.2.5.4 Estimateurs À partir des relations (2.203) et (2.200), on peut construire directement deux estimateurs simples de l’exposant de Hurst à l’aide d’une régression linéaire sur $\log_2 \hat{V}[k]$ vs. k ou sur $\log_2 \hat{V}[k]$ vs. $\log_2 \hat{T}[k]$:

$\hat{H}_1$: $\hat{H}_1 = 1 + a_{min}/2$, avec a_{min} la valeur de a qui minimise la fonction de coût moindres carrés :

$$C_1(a, b) = \sum_{k=2}^K \left(\log_2 \hat{V}[k] - ak - b \right)^2 \quad (2.205)$$

$\hat{H}_2$: $\hat{H}_2 = 1 + a_{min}/2$, avec a_{min} la valeur de a qui minimise la fonction de coût moindres carrés :

$$C_2(a, b) = \sum_{k=2}^K \left(\log_2 \hat{V}[k] - a \log_2 \hat{T}[k] - b \right)^2 \quad (2.206)$$

Connaissant les distributions des estimateurs des variances et périodes moyennes, ces estimations peuvent être améliorées sur deux aspects :

- les variances de $\log_2 \hat{V}[k]$ et $\log_2 \hat{T}[k]$ croissent avec k . Par conséquent, les valeurs correspondants aux petits indices sont plus fiables que celles correspondant aux grands indices, ce dont les

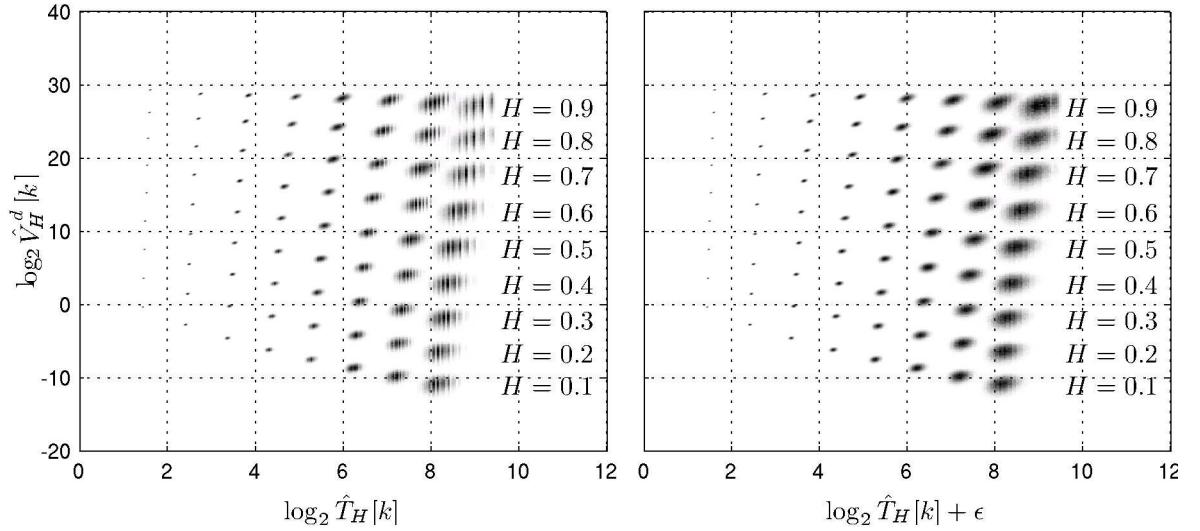


FIGURE 2.81 – Distributions des logarithmes en base 2 des estimées des variances et périodes moyennes mesurées sur 100000 réalisations de fGn pour chaque valeur de H . A gauche, distributions brutes (histogrammes). A droite, distributions lissées en ajoutant aux logarithmes des estimées des périodes moyennes une variable aléatoire uniforme d'écart type égal à l'espacement entre les mailles de l'histogramme.

estimateurs précédents ne tiennent absolument pas compte. Pour tenir compte de cette information, on peut utiliser une régression linéaire pondérée, obtenue en minimisant une fonction de coût de type moindres carrés pondérés. On définit ainsi les estimateurs $\hat{H}_1^w$ et $\hat{H}_2^w$, identiques à $\hat{H}_1$ et $\hat{H}_2$ avec les fonctions de coût

$$C_1^w(a, b) = \sum_{k=2}^K \frac{(\log_2 \hat{V}[k] - ak - b)^2}{\sigma_V[k]^2}, \quad (2.207)$$

$$C_2^w(a, b) = \sum_{k=2}^K \frac{(\log_2 \hat{V}[k] - a \log_2 \hat{T}[k] - b)^2}{\sigma_V[k]^2 + a^2 \sigma_T[k]^2}, \quad (2.208)$$

où $\sigma_V[k]^2$ et $\sigma_T[k]^2$ sont les variances de $\log_2 \hat{V}[k]$ et $\log_2 \hat{T}[k]$. L'inconvénient des nouveaux estimateurs ainsi définis est évidemment qu'ils nécessitent de connaître les valeurs $\sigma_V[k]^2$ et $\sigma_T[k]^2$. Sachant de plus que ces valeurs dépendent de H , il faudrait théoriquement pour faire au mieux utiliser des valeurs moyennées par rapport à une certaine distribution a priori de H . Heureusement, il semble que quand les IMFs sont calculés avec un nombre d'itérations de tamisage par IMF fixé a priori, les distributions empiriques de $\sigma_{V_H}[k]^2$ et $\sigma_{T_H}[k]^2$ dépendent essentiellement de H sous la forme d'un préfacteur indépendant de k ⁷(cf Fig. 2.83) :

$$\sigma_{V_H}[k]^2 \simeq f_V(H)g_V(k) \quad (2.209)$$

$$\sigma_{T_H}[k]^2 \simeq f_T(H)g_T(k). \quad (2.210)$$

Par conséquent, les valeurs (a, b) minimisant les critères (2.207) et (2.208) avec $\sigma_V[k]^2 = \sigma_{V_H}[k]^2$ et $\sigma_T[k]^2 = \sigma_{T_H}[k]^2$ ne dépendent pratiquement pas de H .

En revanche, si les IMFs sont calculés avec des nombres d'itérations contrôlés par le critère d'arrêt « local » (cf 5.2), la dépendance de $\sigma_{V_H}[k]^2$ et $\sigma_{T_H}[k]^2$ par rapport à H est plus compliquée

7. Ce résultat n'a été effectivement vérifié que pour 10 et 20 itérations de tamisage.

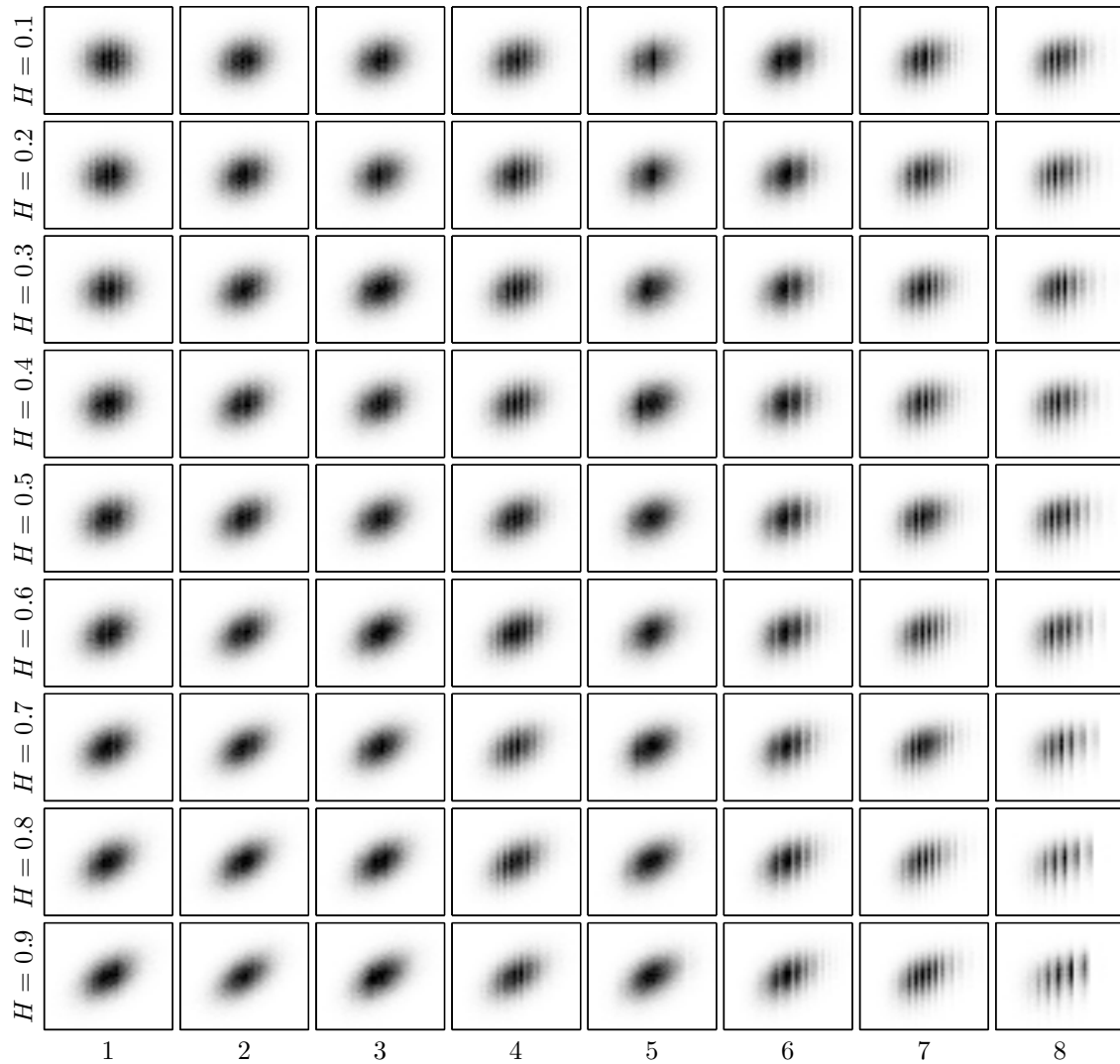


FIGURE 2.82 – Agrandissements des distributions des logarithmes en base 2 des estimées des variances (en ordonnée) et périodes moyennes (en abscisse) mesurées sur 100000 réalisations de fGn pour chaque valeur de H . On observe des corrélations positives entre variances et périodes moyennes des IMFs autres que le premier qui croissent généralement avec H et dans une moindre mesure décroissent avec l'indice de l'IMF à l'exception du cas $H = 0.1$ où la variation en fonction de l'indice est inversée. La corrélation pour le premier IMF est toujours inférieure à celle du second.

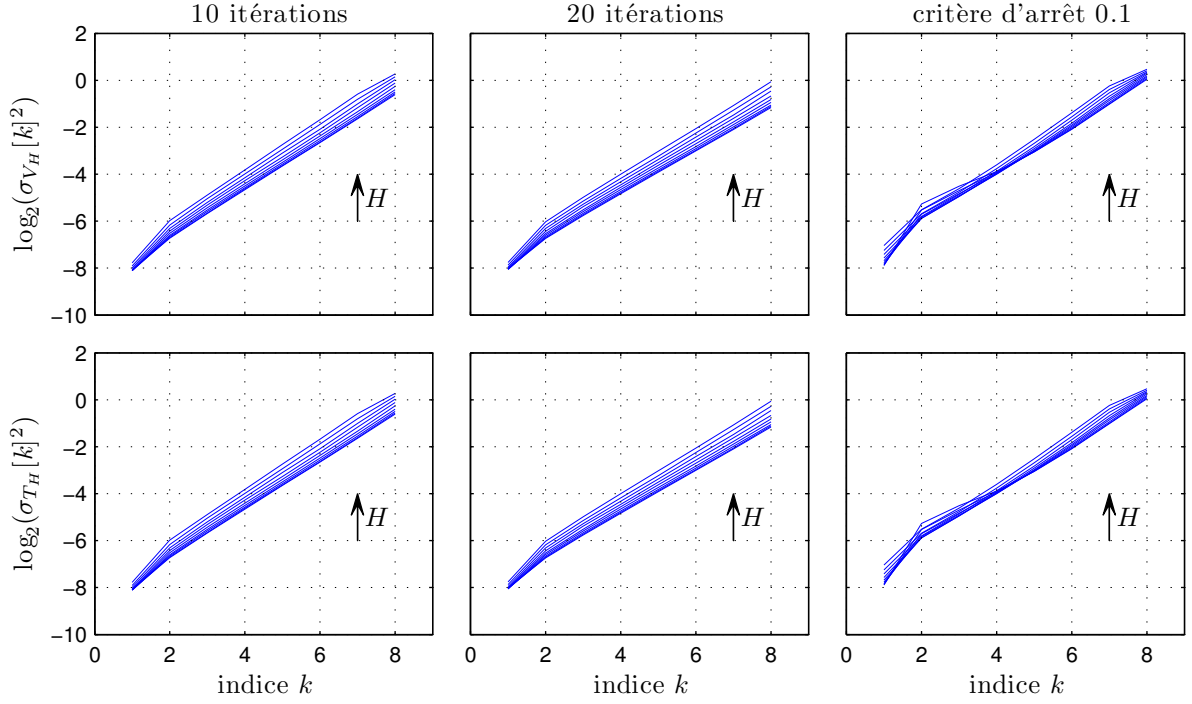


FIGURE 2.83 – Évolutions de $\sigma_{V_H}[k]^2$ et $\sigma_{T_H}[k]^2$ en fonction k et H pour différents nombres d'itérations ou critères d'arrêt du processus de tamisage. Pour des nombres d'itérations fixés a priori, il semble que la dépendance par rapport à H se résume à un préfacteur.

(cf Fig. 2.83) et les valeurs (a, b) minimisant les critères (2.207) et (2.208) avec $\sigma_V[k]^2 = \sigma_{V_H}[k]^2$ et $\sigma_T[k]^2 = \sigma_{T_H}[k]^2$ dépendent de H . Ce problème peut cependant être contourné en utilisant une procédure itérative :

1. réaliser une première estimation grossière avec $\sigma_V[k]^2$ et $\sigma_T[k]^2$ des valeurs moyennées par rapport à H de $\sigma_{V_H}[k]^2$ et $\sigma_{T_H}[k]^2$
 2. réaliser une deuxième estimation plus fine avec $\sigma_V[k]^2$ et $\sigma_T[k]^2$ les valeurs de $\sigma_{V_H}[k]^2$ et $\sigma_{T_H}[k]^2$ correspondant à la valeur de H obtenue lors de la première estimation.
 3. éventuellement itérer l'étape 2 si la nouvelle valeur de H est trop différente de la première estimation, l'inconvénient étant que la convergence n'est pas garantie.
- la régression linéaire — simple ou pondérée — sur un ensemble de points (x_i, y_i) est un estimateur non biaisé des paramètres du modèle linéaire

$$\mathbb{E}y_i = a\mathbb{E}x_i + b. \quad (2.211)$$

Or, à supposer que $\hat{V}[k]$ et $\hat{T}[k]$ sont des estimateurs non biaisés, les relations (2.203) et (2.200) donnent plutôt

$$\log_2 \mathbb{E}V_H[k] \simeq 2(H-1)k + \log_2 C, \quad (2.212)$$

$$\log_2 \mathbb{E}V_H[k] = 2(H-1) \log_2 \mathbb{E}T_H[k] + \log_2 C', \quad (2.213)$$

qui diffèrent du modèle (2.211) par le fait que $\log_2 \mathbb{E}\hat{V}_H[k] \neq \mathbb{E}\log_2 \hat{V}_H[k]$ et $\log_2 \mathbb{E}\hat{T}_H[k] \neq \mathbb{E}\log_2 \hat{T}_H[k]$. De ce fait, les estimateurs $\hat{H}_1$ et $\hat{H}_2$ comportent un certain biais. En pratique, le problème ne se pose en fait vraiment que pour les estimateurs des variances puisque les différences relatives entre $\log_2 \mathbb{E}\hat{V}_H[k]$ et $\log_2 \hat{V}_H[k]$ peuvent atteindre pratiquement 10% alors que les différences relatives entre $\log_2 \mathbb{E}\hat{T}_H[k]$ et $\log_2 \hat{T}_H[k]$ ne dépassent pas 0.5%.

Pour corriger le biais dû à la différence entre $\log_2 \mathbb{E} \hat{V}_H[k]$ et $\mathbb{E} \log_2 \hat{V}_H[k]$, on peut utiliser le fait qu'on connaît les distributions des $\hat{V}_H[k]$. Sachant que pour une variable aléatoire x distribuée selon une loi $\text{Gamma}(\alpha, \beta)$, on a

$$\mathbb{E} \log_2 x = \log_2 \mathbb{E} x + \frac{\Gamma'(\alpha)}{\ln 2\Gamma(\alpha)} - \log_2 \alpha, \quad (2.214)$$

on peut obtenir des estimations a priori moins biaisées de H en remplaçant $\log_2 \hat{V}[k]$ dans les régressions linéaires par $\log_2 \hat{V}[k] + c(\alpha)$, où

$$c(\alpha) = -\frac{\Gamma'(\alpha)}{\ln 2\Gamma(\alpha)} + \log_2 \alpha. \quad (2.215)$$

Si l'ajout du terme correctif $c(\alpha)$ permet a priori de réduire le biais des estimateurs, il présente l'inconvénient qu'il dépend théoriquement de la valeur de H qu'on cherche à estimer. Comme précédemment, cet inconvénient peut heureusement être contourné à l'aide d'une procédure itérative. Enfin, si l'on ajoute cette correction aux estimateurs définis précédemment, on obtient les quatre nouveaux estimateurs $\hat{H}_1^{cor}$, $\hat{H}_2^{cor}$, $\hat{H}_1^{w,cor}$ et $\hat{H}_2^{w,cor}$.

2.2.5.5 Performances Les performances des 8 estimateurs définis précédemment ont été évaluées sur 100000 réalisations de fGn de 2048 points pour H allant de 0.1 à 0.9 par pas de 0.1. Pour chaque réalisation, les estimateurs utilisent les variances et périodes moyennes des IMFs 2 à 7. Les résultats sont proposés dans les tableaux Tab. 2.2, Tab. 2.3 et Tab. 2.4. Les réalisations utilisées pour ces évaluations sont différentes de celles utilisées pour les modélisations. À titre de référence, les tableaux de résultats contiennent aussi les performances de l'estimateur de Abry et Veitch (noté $\hat{H}_{WT}$ dans la suite) fondé sur la transformée en ondelettes discrète [1]. Le principe de ce dernier est très semblable à celui utilisé pour l'estimateur $\hat{H}_1^{w,cor}$ puisque le comportement de loi d'échelle se traduit par les mêmes types de propriétés sur les variances des coefficients d'ondelettes que celles observées sur les variances des IMFs. En fait, les estimateurs basés sur l'EMD s'inspirent très largement des stratégies d'estimation à base de transformée en ondelettes discrète, telle que celle utilisée pour l'estimateur utilisé comme référence. L'estimation de l'exposant de Hurst à l'aide de la transformée en ondelettes est réalisée à partir des échelles 2 à 7 qui correspondent approximativement aux échelles des IMFs 2 à 7.

On constate de manière générale que les estimateurs basés sur l'EMD ont des performances voisines de celles de l'estimateur de référence. Ce dernier est toutefois en général meilleur, essentiellement parce qu'il présente systématiquement une meilleure variance. Pour ce qui est du biais, les estimateurs à base d'EMD ont des biais du même ordre de grandeur que l'estimateur de référence, parfois légèrement plus importants, parfois légèrement meilleurs. Les différents estimateurs à base d'EMD présentent des biais variables d'un estimateur à l'autre mais généralement du même ordre de grandeur. Le biais est, à l'exception de l'estimateur $\hat{H}_1^{cor}$, systématiquement positif pour les petites valeurs de H , et négatif pour les grandes avec un minimum aux alentours de $H = 0.6 - 0.7$. On peut constater de plus que les corrections pour tenir compte de la distribution Gamma des variances des IMFs tendent systématiquement à augmenter la valeur du biais de l'ordre de environ 0.01, réduisant ainsi sa valeur absolue pour les grandes valeurs de H au prix d'une augmentation pour les petites valeurs de H . Les corrections n'affectant pas les variances des estimateurs, on peut conclure de manière générale que leur effet n'améliore pas significativement les estimateurs mais déplacent plutôt les valeurs de H pour lesquelles ils ont les meilleures performances. En moyenne par rapport à H , l'intervalle sur lequel le biais est positif étant plus grand que celui sur lequel le biais est négatif, les corrections ont tendance à légèrement dégrader les performances.

L'apport des pondérations est en revanche très net puisqu'elles permettent de réduire systématiquement les variances des estimateurs d'un facteur légèrement supérieur à 2. Elles ont également un

estimateur	$\hat{H}_1$	$\hat{H}_1^w$	$\hat{H}_1^{cor}$	$\hat{H}_1^{w,cor}$	$\hat{H}_2$	$\hat{H}_2^w$	$\hat{H}_2^{cor}$	$\hat{H}_2^{w,cor}$	$\hat{H}_{WT}$
$H = 0.1$	0,2	0,12	0,21	0,12	0,17	0,066	0,18	0,071	-0,11
$H = 0.2$	0,12	0,08	0,14	0,086	0,1	0,044	0,11	0,049	-0,049
$H = 0.3$	0,072	0,05	0,083	0,056	0,055	0,026	0,066	0,032	-0,018
$H = 0.4$	0,036	0,027	0,047	0,034	0,027	0,013	0,038	0,019	-0,00024
$H = 0.5$	0,012	0,011	0,024	0,018	0,01	0,0032	0,023	0,01	0,01
$H = 0.6$	-0,0043	-0,00074	0,0093	0,0071	0,00093	-0,004	0,014	0,0036	0,016
$H = 0.7$	-0,014	-0,0079	0,00092	0,00084	-0,0057	-0,011	0,0095	-0,0026	0,02
$H = 0.8$	-0,021	-0,012	-0,0036	-0,0018	-0,012	-0,019	0,0055	-0,0093	0,022
$H = 0.9$	-0,026	-0,012	-0,0066	-0,0015	-0,018	-0,025	8,8e-05	-0,014	0,022

TABLE 2.2 – Biais des estimateurs. Pour chaque valeur de H , l'estimateur dont le biais est minimal est mis en gras.

petit effet sur les biais : ceux des estimateurs $\hat{H}_1$ sont réduits en moyenne d'un facteur un peu inférieur à 2 alors que ceux des estimateurs $\hat{H}_2$ ne sont réduits que pour $H < 0.7 - 0.8$ et au contraire augmentés pour les grandes valeurs de H . En moyenne par rapport à H , les erreurs quadratiques moyennes sont très nettement améliorées par la pondération. La seule exception à cette amélioration concerne les estimateurs $\hat{H}_2^w$ et $\hat{H}_2^{w,cor}$ dont les variances pour $H = 0.9$ sont apparemment détériorées. En fait, cette détérioration n'est due qu'à un très faible nombre de réalisations (environ 5 sur 100000) pour lesquels ces deux estimateurs aboutissent à des valeurs totalement aberrantes (par exemple de l'ordre de 10^{12}) dont on a artificiellement limité l'impact sur les performances en les remplaçant par des valeurs moins aberrantes (typiquement 10 lorsque la valeur retournée par l'estimateur est positive, -10 sinon). Ces valeurs aberrantes s'expliquent par le fait que les estimateurs $\hat{H}_2^w$ et $\hat{H}_2^{w,cor}$ utilisent une régression linéaire pondérée avec incertitudes sur les deux coordonnées dont la solution ne peut être obtenue que par une étape d'optimisation numérique. Les cas où les estimateurs retournent des valeurs aberrantes correspondent simplement aux cas où l'algorithme d'optimisation numérique (la fonction de MATLAB `fminsearch`) n'a pas convergé. Si on supprime ces valeurs aberrantes des statistiques, les variances des estimateurs $\hat{H}_2^w$ et $\hat{H}_2^{w,cor}$ ne sont pas moins bonnes pour $H = 0.9$ que pour $H = 0.8$.

Enfin, on n'observe pas de différences très significatives dans les performances des deux familles d'estimateurs à base d'EMD basés sur $\hat{H}_1$ et $\hat{H}_2$. Les estimateurs basés sur $\hat{H}_2$ offrent généralement un biais plus faible aux petites valeurs de H et parfois un peu plus important aux grandes valeurs. En contrepartie, les estimateurs basés sur $\hat{H}_1$ ont systématiquement une variance un peu meilleure d'un facteur environ 1.5. Au final, les erreurs quadratiques moyennes des deux familles d'estimateurs sont assez proches : les estimateurs de la famille $\hat{H}_2$ sont un peu meilleurs pour les petites valeurs de H et un peu moins bons pour les grandes.

2.3 EMD bivariée d'un bruit « blanc » complexe

2.3.1 Modèles de bruits

On s'est intéressé pour cette étude à deux types de bruits complexes. Le premier est un bruit blanc gaussien complexe ayant pour parties réelles et imaginaires deux bruits blancs gaussiens réels indépendants de moyenne nulle et de variance unité. Ses échantillons, considérés comme des couples (partie réelle, partie imaginaire) sont donc indépendants et identiquement distribués selon une loi normale bivariée centrée dont la matrice de covariance est l'identité. Par la suite, ce type de bruit blanc sera appelé « bruit blanc simple ».

Le deuxième type de bruit auquel on s'est intéressé, qu'on appellera par la suite « bruit blanc analytique », est obtenu à partir du bruit blanc précédent en lui appliquant un filtre analytique dont

estimateur	$\hat{H}_1$	$\hat{H}_1^w$	$\hat{H}_1^{cor}$	$\hat{H}_1^{w,cor}$	$\hat{H}_2$	$\hat{H}_2^w$	$\hat{H}_2^{cor}$	$\hat{H}_2^{w,cor}$	$\hat{H}_{WT}$
$H = 0.1$	0,0022	0,001	0,0022	0,001	0,0035	0,0017	0,0035	0,0017	0,00098
$H = 0.2$	0,0022	0,001	0,0022	0,001	0,0033	0,0016	0,0033	0,0016	0,00094
$H = 0.3$	0,0023	0,001	0,0024	0,001	0,0033	0,0015	0,0033	0,0015	0,00092
$H = 0.4$	0,0025	0,0011	0,0025	0,0011	0,0033	0,0015	0,0033	0,0015	0,00091
$H = 0.5$	0,0027	0,0012	0,0027	0,0012	0,0033	0,0015	0,0034	0,0015	0,0009
$H = 0.6$	0,0029	0,0012	0,0031	0,0013	0,0034	0,0015	0,0035	0,0015	0,00091
$H = 0.7$	0,0033	0,0014	0,0034	0,0014	0,0035	0,0016	0,0036	0,0016	0,00092
$H = 0.8$	0,0037	0,0015	0,0038	0,0016	0,0036	0,0017	0,0038	0,0017	0,00094
$H = 0.9$	0,0041	0,0017	0,0042	0,0017	0,0038	0,0072	0,0038	0,011	0,00097

TABLE 2.3 – Variances des estimateurs. Pour chaque valeur de H , l'estimateur dont la variance est minimale est mis en gras.

estimateur	$\hat{H}_1$	$\hat{H}_1^w$	$\hat{H}_1^{cor}$	$\hat{H}_1^{w,cor}$	$\hat{H}_2$	$\hat{H}_2^w$	$\hat{H}_2^{cor}$	$\hat{H}_2^{w,cor}$	$\hat{H}_{WT}$
$H = 0.1$	0,041	0,015	0,045	0,017	0,031	0,006	0,035	0,0068	0,013
$H = 0.2$	0,018	0,0075	0,021	0,0085	0,013	0,0035	0,016	0,004	0,0034
$H = 0.3$	0,0076	0,0035	0,0093	0,0042	0,0062	0,0022	0,0076	0,0025	0,0012
$H = 0.4$	0,0038	0,0018	0,0048	0,0022	0,004	0,0017	0,0048	0,0019	0,00091
$H = 0.5$	0,0028	0,0013	0,0033	0,0015	0,0034	0,0015	0,0039	0,0016	0,001
$H = 0.6$	0,003	0,0012	0,0031	0,0013	0,0034	0,0015	0,0037	0,0016	0,0012
$H = 0.7$	0,0035	0,0014	0,0034	0,0014	0,0035	0,0017	0,0037	0,0016	0,0013
$H = 0.8$	0,0041	0,0017	0,0039	0,0016	0,0038	0,0021	0,0038	0,0018	0,0014
$H = 0.9$	0,0048	0,0018	0,0042	0,0017	0,0041	0,0078	0,0038	0,011	0,0015

TABLE 2.4 – Erreurs quadratiques moyennes pour les différents estimateurs (définies comme $EQM = \text{biais}^2 + \text{variance}$). Pour chaque valeur de H , le meilleur estimateur est mis en gras.

la réponse en fréquence H_a est

$$H_a(f) = \begin{cases} 1 & \text{si } f \geq 0, \\ 0 & \text{si } f < 0. \end{cases} \quad (2.216)$$

Le résultat est un bruit gaussien mais n'est pas à proprement parler un bruit blanc puisque ses échantillons ne sont pas indépendants. Toutefois, on conservera abusivement l'appellation bruit blanc parce que sa densité spectrale de puissance est constante pour les fréquences positives.

2.3.2 Statistiques d'ordre 1

Comme pour le cas de l'EMD classique de bruits blancs gaussiens réels, on s'est intéressé aux densités de probabilité marginales des IMFs et approximations (cf Fig. 2.84 à Fig. 2.99). On observe pour les deux algorithmes que celles des IMFs sont presque gaussiennes à partir du deuxième IMF pour le bruit blanc simple, et que celles des approximations sont toujours gaussiennes. Les densités marginales des IMFs sont similaires à celles observées dans le cas classique dans la mesure où elles ressemblent à des « gaussiennes pointues » mais le caractère pointu est encore plus prononcé ici, surtout dans le cas des IMFs issus du bruit blanc analytique et plus particulièrement ceux calculés par l'algorithme Algo. 4. De plus, les densités de probabilité des approximations sont invariantes par rotation, comme le signal analysé, quel que soit le nombre de directions utilisé pour calculer la moyenne de l'enveloppe (cf Fig. 2.102, Fig. 2.103, Fig. 2.106 et Fig. 2.107). On peut considérer que celles des IMFs autres que le premier le sont approximativement dès que le nombre de directions est supérieur à 16 environ.

Le cas du premier IMF (cf Fig. 2.100, Fig. 2.101, Fig. 2.104 et Fig. 2.105) est encore une fois particulier puisqu'il n'est jamais gaussien et que de plus la structure de sa densité de probabilité marginale dépend fortement du nombre de directions utilisées pour calculer le centre de l'enveloppe. On observe ainsi que la densité de probabilité du premier IMF obtenu à l'aide de N directions est symétrique par rapport à $2N$ axes correspondant à ces directions ainsi qu'à leurs bissectrices. Il n'y a là rien de surprenant puisqu'il s'agit très exactement de la structure de symétrie de l'algorithme lui-même. Plus en détail, on peut justifier les positions des N directions où la densité de probabilité est la plus faible en faisant appel au même raisonnement que celui utilisé pour justifier l'aspect bimodal des densités de probabilité du premier IMF dans le cas de l'EMD classique d'un bruit blanc. Considérons pour cela la formulation de l'algorithme Algo. 6

$$\mathcal{S}^{B2}[x](t) = \frac{4}{N} \sum_{k=1}^{N/2} u_{\varphi_k} \mathcal{S}[p_{\varphi_k}](t) \quad (2.217)$$

où $p_{\varphi_k}(t)$ est la projection de $x(t)$ sur la direction représentée par le vecteur unitaire u_{φ_k} . Dans le cas particulier où $N = 4$, cet algorithme revient à appliquer l'EMD classique indépendamment (mais avec les mêmes nombres d'itérations) sur les parties réelles et imaginaires du signal. L'explication de l'aspect quadrimodal de la densité de probabilité découle alors directement de l'aspect bimodal des parties réelles et imaginaires. Plus généralement, pour un nombre de directions pair, on peut considérer que la contribution $u_{\varphi_k} \mathcal{S}[p_{\varphi_k}](t)$ a tendance à éloigner de zéro les extrema de la projection $p_{\varphi_k}(t)$, ce qui se traduit par une « direction creuse » orthogonale à u_{φ_k} dans la densité de probabilité marginale du premier IMF. Le même raisonnement reste en partie valable pour l'algorithme Algo. 4. Si on reprend la définition, l'opérateur de tamisage pour cet algorithme est défini par

$$\mathcal{S}^{B1}[x](t) = x(t) - \frac{1}{N} \sum_{k=1}^4 e_{\varphi_k}(t), \quad (2.218)$$

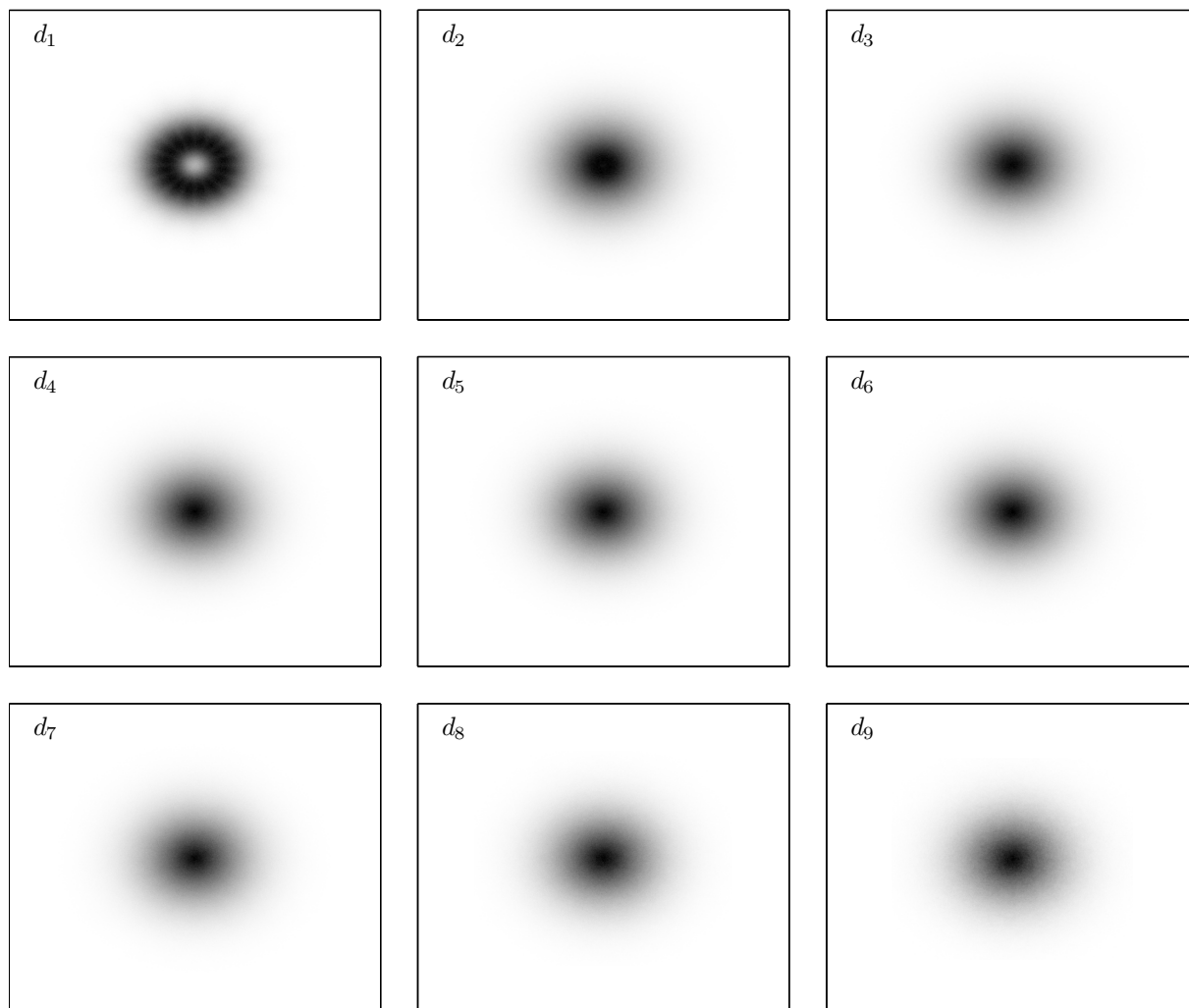


FIGURE 2.84 – Densités marginales des IMFs issus du bruit blanc simple calculés par l'algorithme Algo. 4 avec 16 directions. (échelles normalisées)

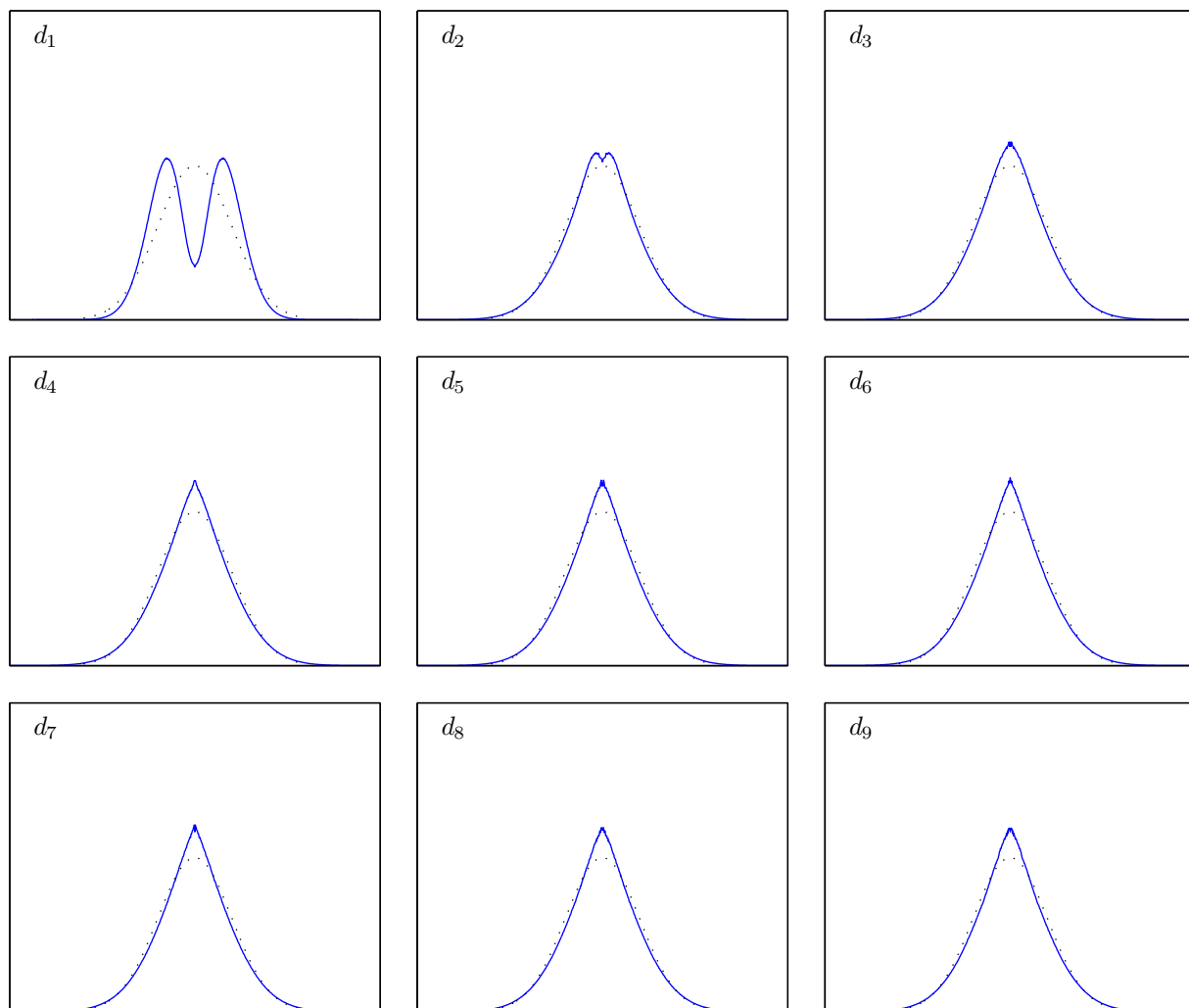


FIGURE 2.85 – Coupes des densités marginales des IMFs issus du bruit blanc simple calculés par l’algorithme Algo. 4 avec 16 directions. Il ne s’agit pas des densités marginales des parties réelles (ou imaginaires) mais d’une coupe des densités bivariées Fig. 2.84. Si $\mathcal{P}_r(r)$ est la densité du module d’un IMF, ce qui est représenté est $\mathcal{P}_r(r)/r$ avec une partie négative rajoutée par symétrie. La courbe en pointillés est la gaussienne de référence.

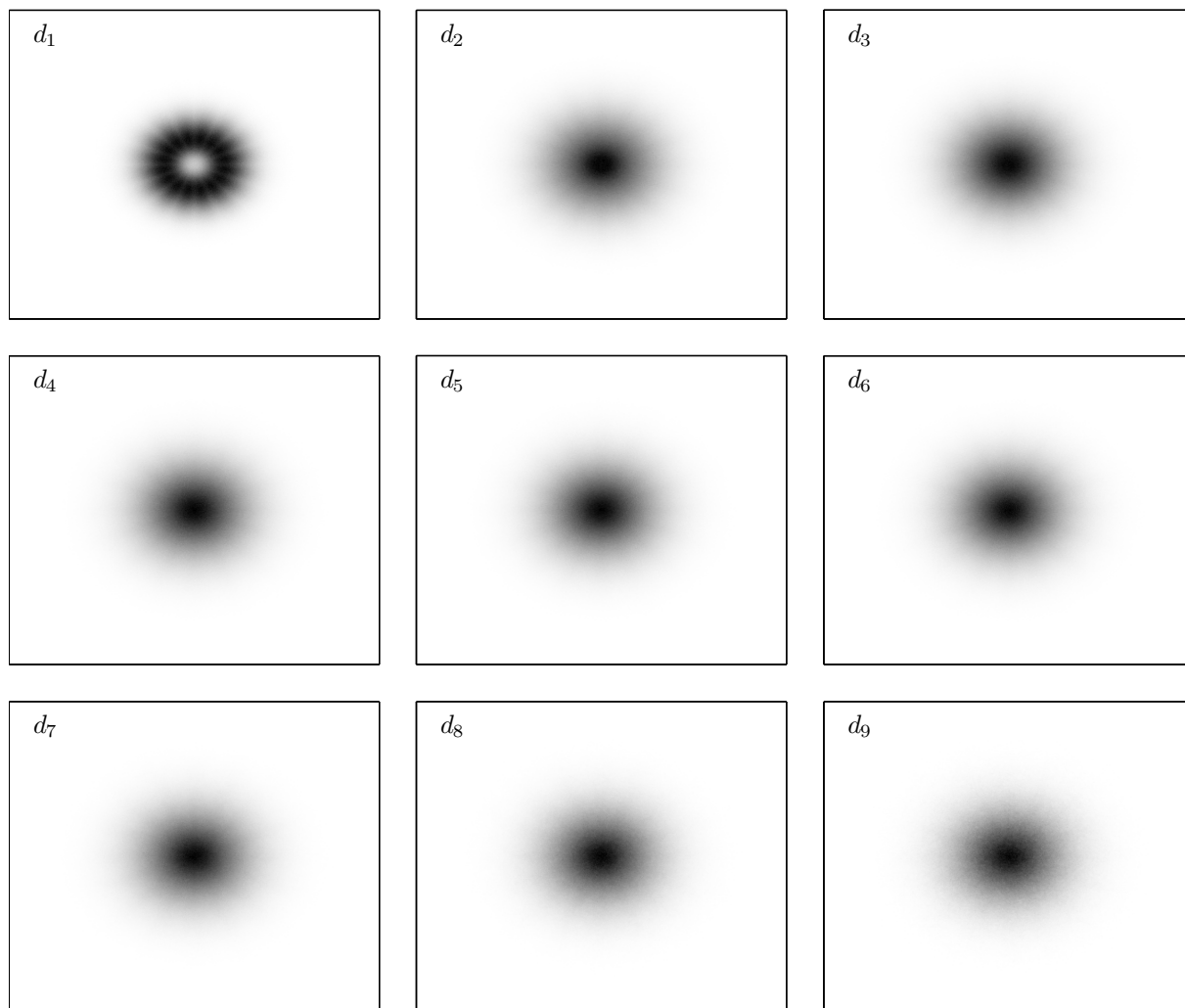


FIGURE 2.86 – Densités marginales des IMFs issus du bruit blanc simple calculés par l'algorithme Algo. 5 avec 16 directions. (échelles normalisées)

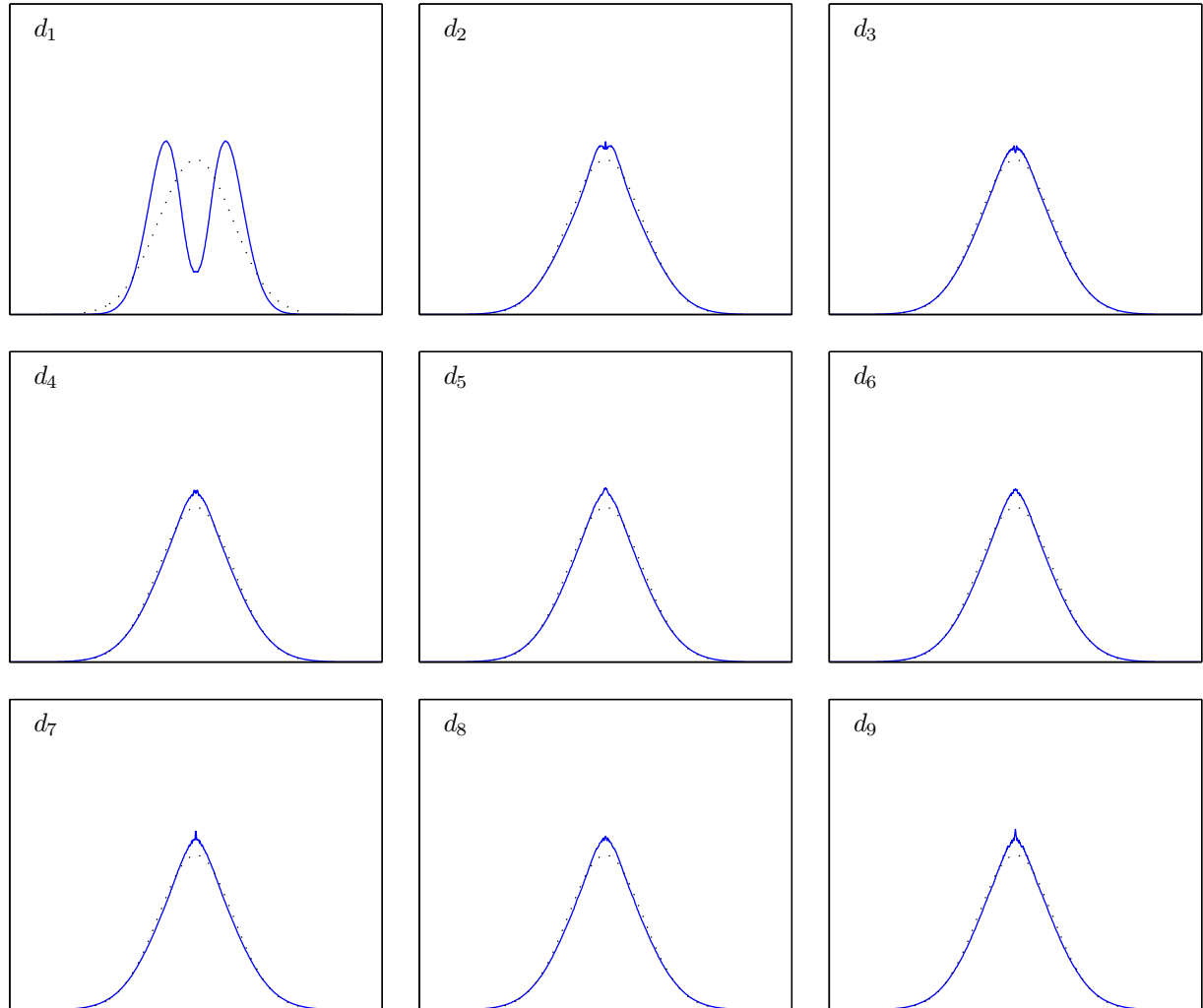


FIGURE 2.87 – Coupes des densités marginales des IMFs issus du bruit blanc simple calculés par l’algorithme Algo. 5 avec 16 directions. Il ne s’agit pas des densités marginales des parties réelles (ou imaginaires) mais d’une coupe des densités bivariées Fig. 2.86. Si $\mathcal{P}_r(r)$ est la densité du module d’un IMF, ce qui est représenté est $\mathcal{P}_r(r)/r$ avec une partie négative rajoutée par symétrie. La courbe en pointillés est la gaussienne de référence.

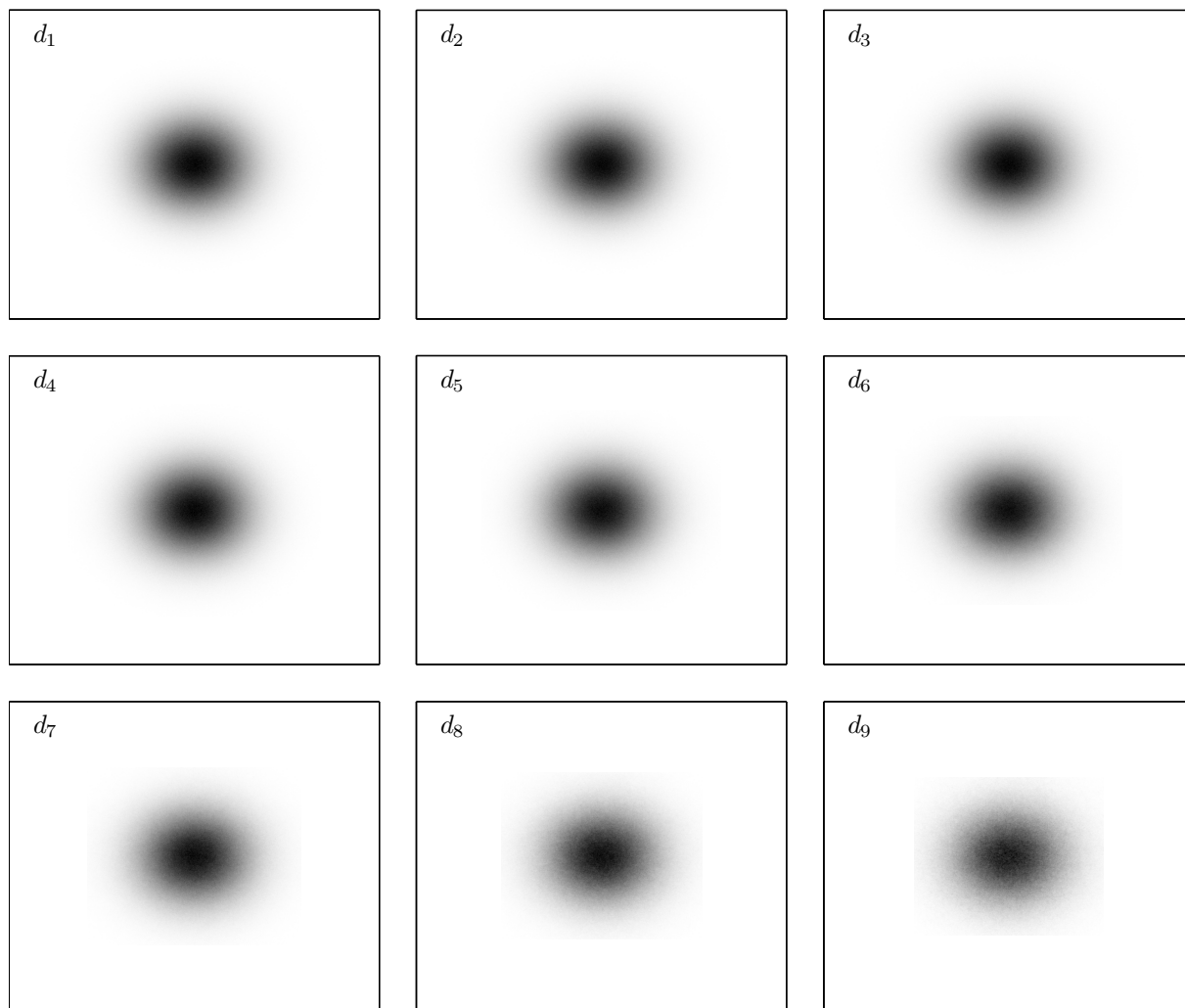


FIGURE 2.88 – Densités marginales des approximations issues du bruit blanc simple calculées par l'algorithme Algo. 4 avec 16 directions. (échelles normalisées)

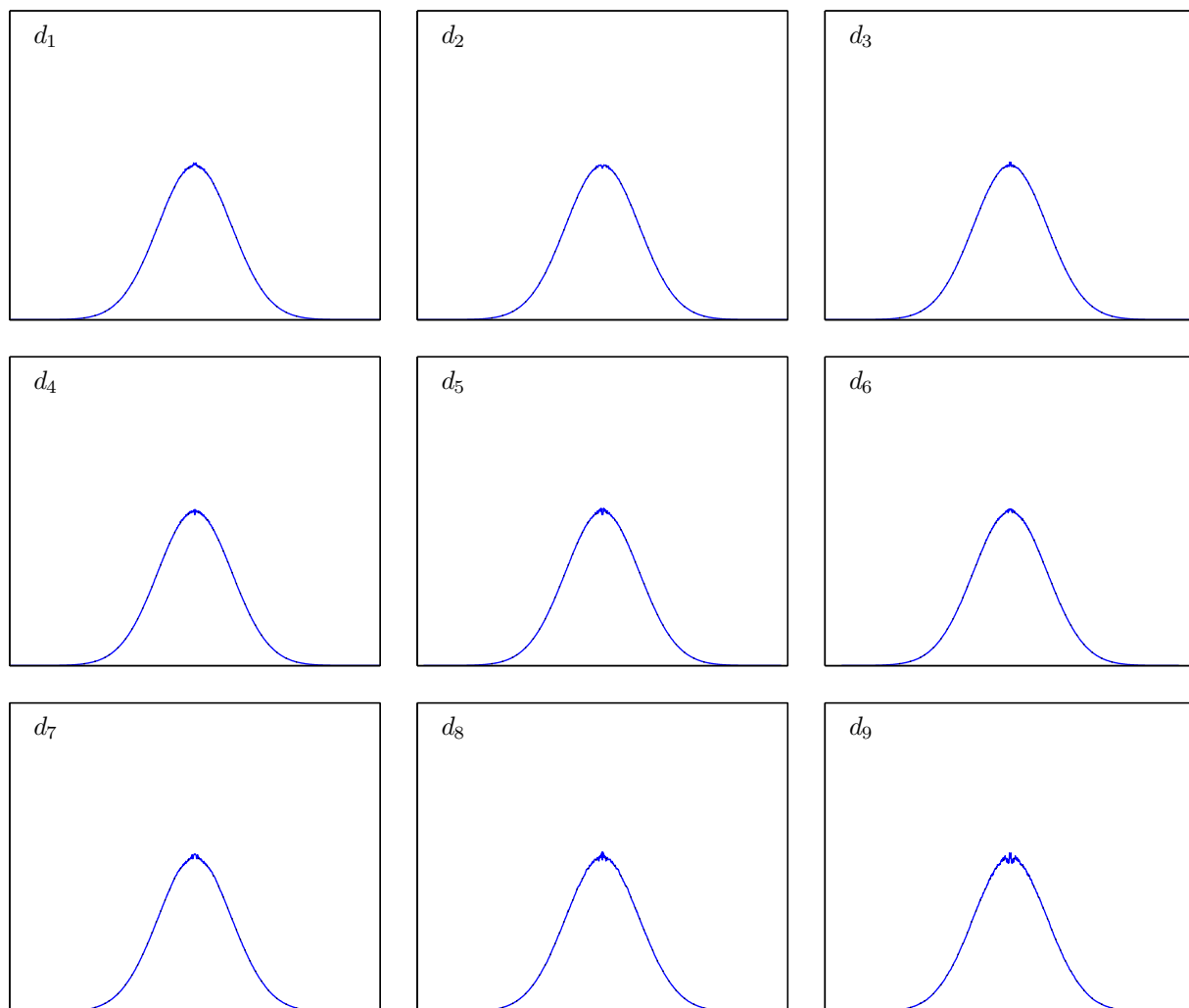


FIGURE 2.89 – Coupes des densités marginales des approximations issues du bruit blanc simple calculées par l'algorithme Algo. 4 avec 16 directions. Il ne s'agit pas des densités marginales des parties réelles (ou imaginaires) mais d'une coupe des densités bivariées Fig. 2.88. Si $\mathcal{P}_r(r)$ est la densité du module d'une approximation, ce qui est représenté est $\mathcal{P}_r(r)/r$ avec une partie négative rajoutée par symétrie. La courbe en pointillés est la gaussienne de référence.

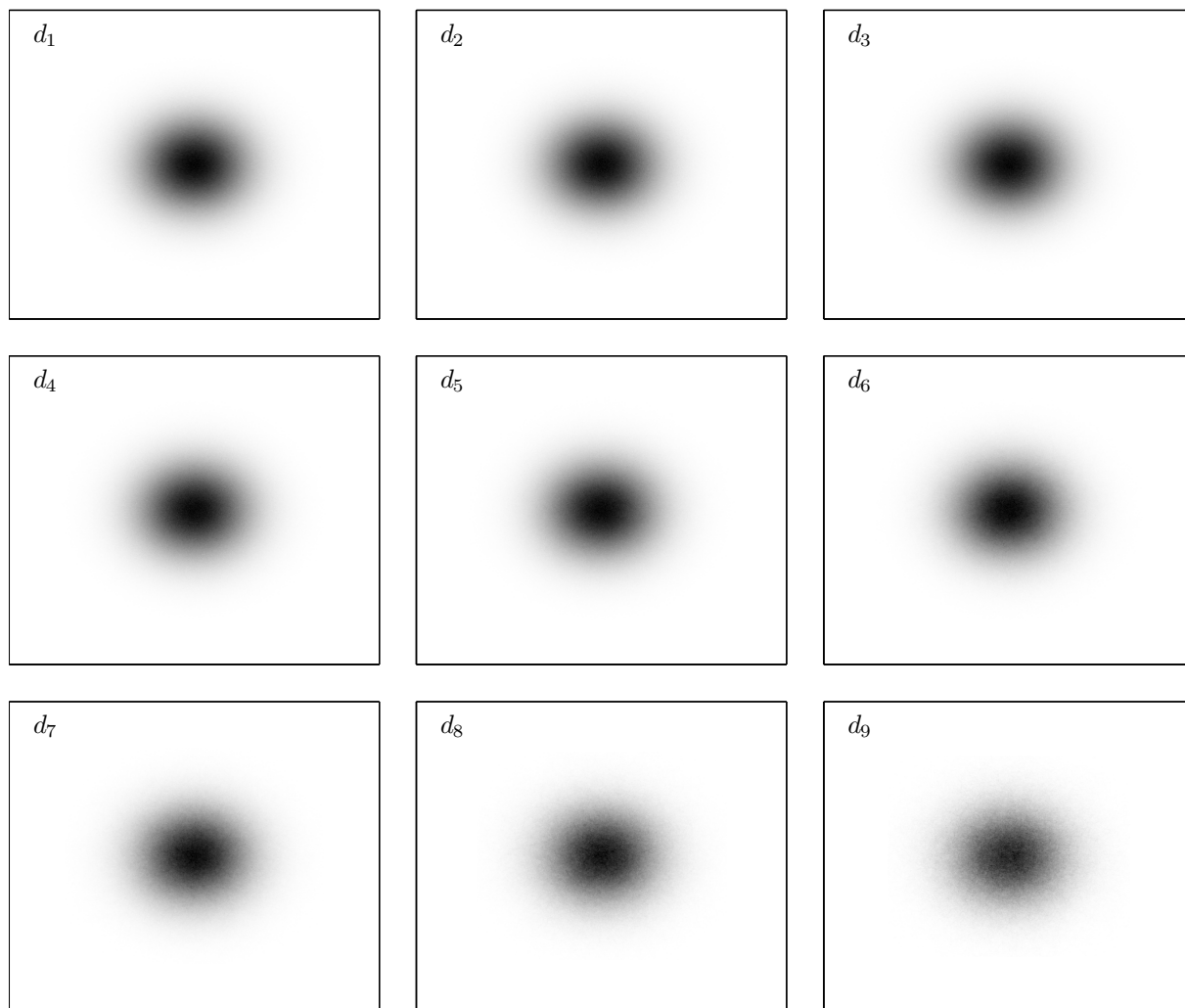


FIGURE 2.90 – Densités marginales des approximations issues du bruit blanc simple calculées par l’algorithme Algo. 5 avec 16 directions. (échelles normalisées)

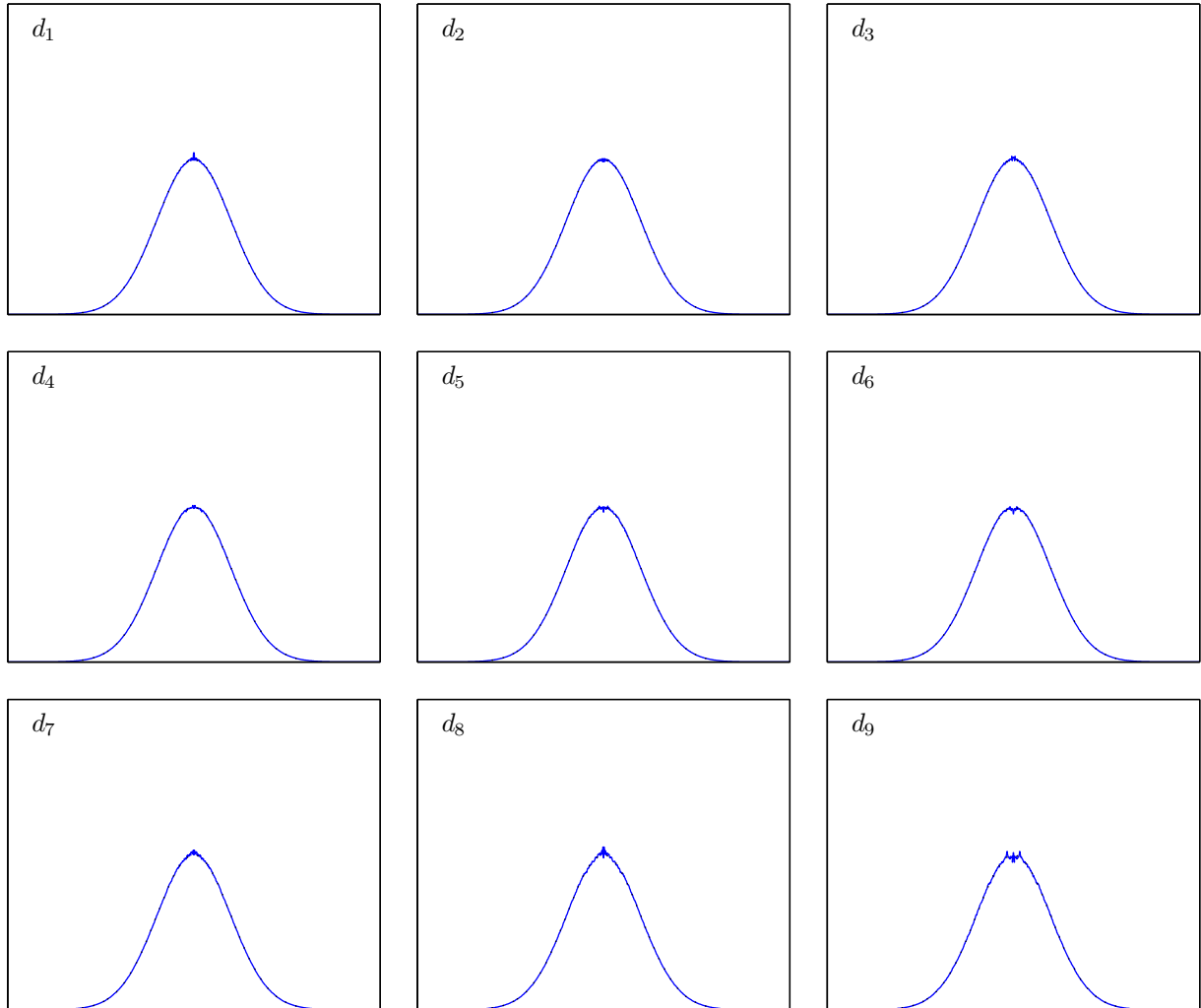


FIGURE 2.91 – Coupes des densités marginales des approximations issues du bruit blanc simple calculées par l'algorithme Algo. 5 avec 16 directions. Il ne s'agit pas des densités marginales des parties réelles (ou imaginaires) mais d'une coupe des densités bivariées Fig. 2.90. Si $\mathcal{P}_r(r)$ est la densité du module d'une approximation, ce qui est représenté est $\mathcal{P}_r(r)/r$ avec une partie négative rajoutée par symétrie. La courbe en pointillés est la gaussienne de référence.

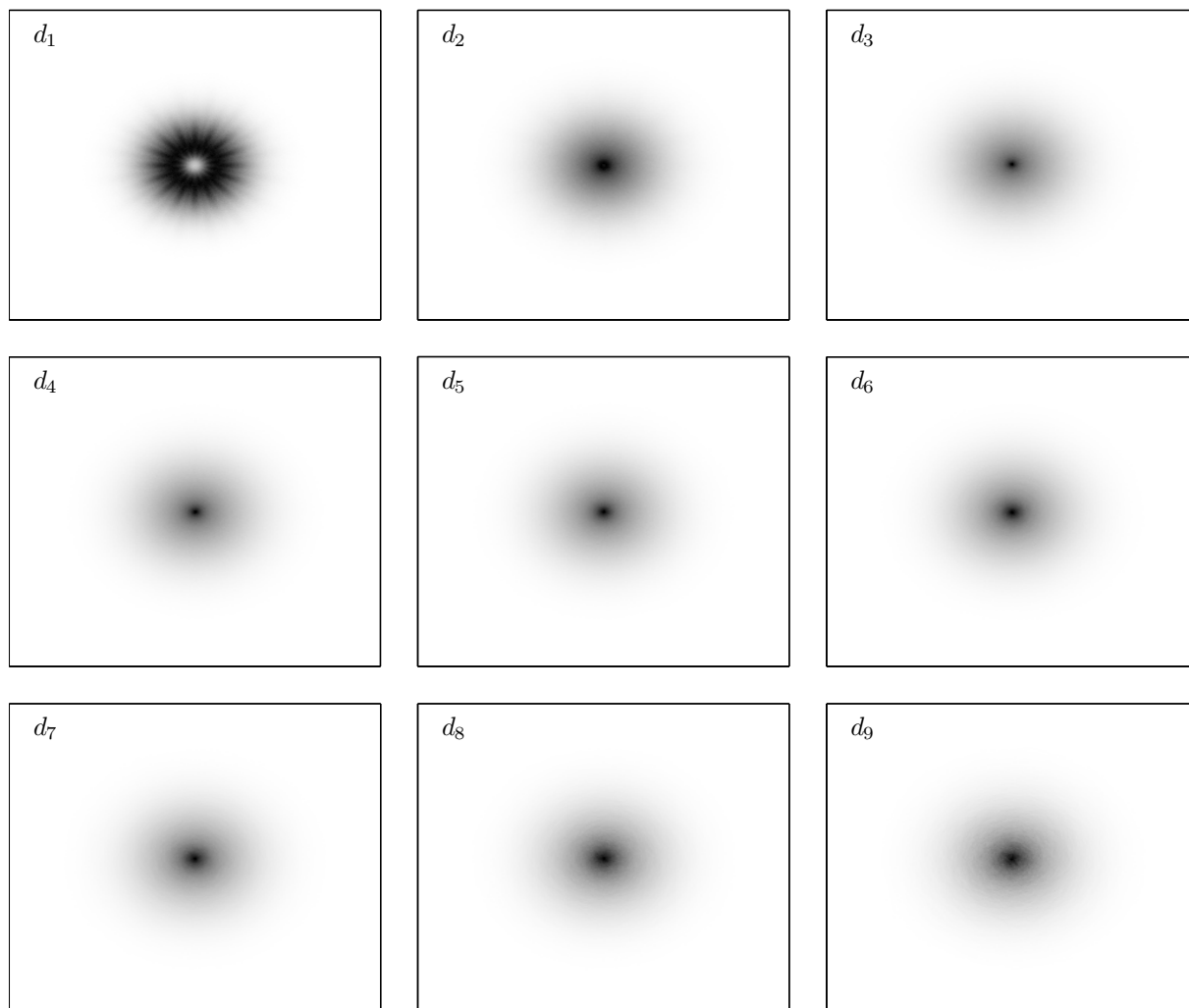


FIGURE 2.92 – Densités marginales des IMFs issus du bruit blanc analytique calculés par l'algorithme Algo. 4 avec 16 directions. (échelles normalisées)

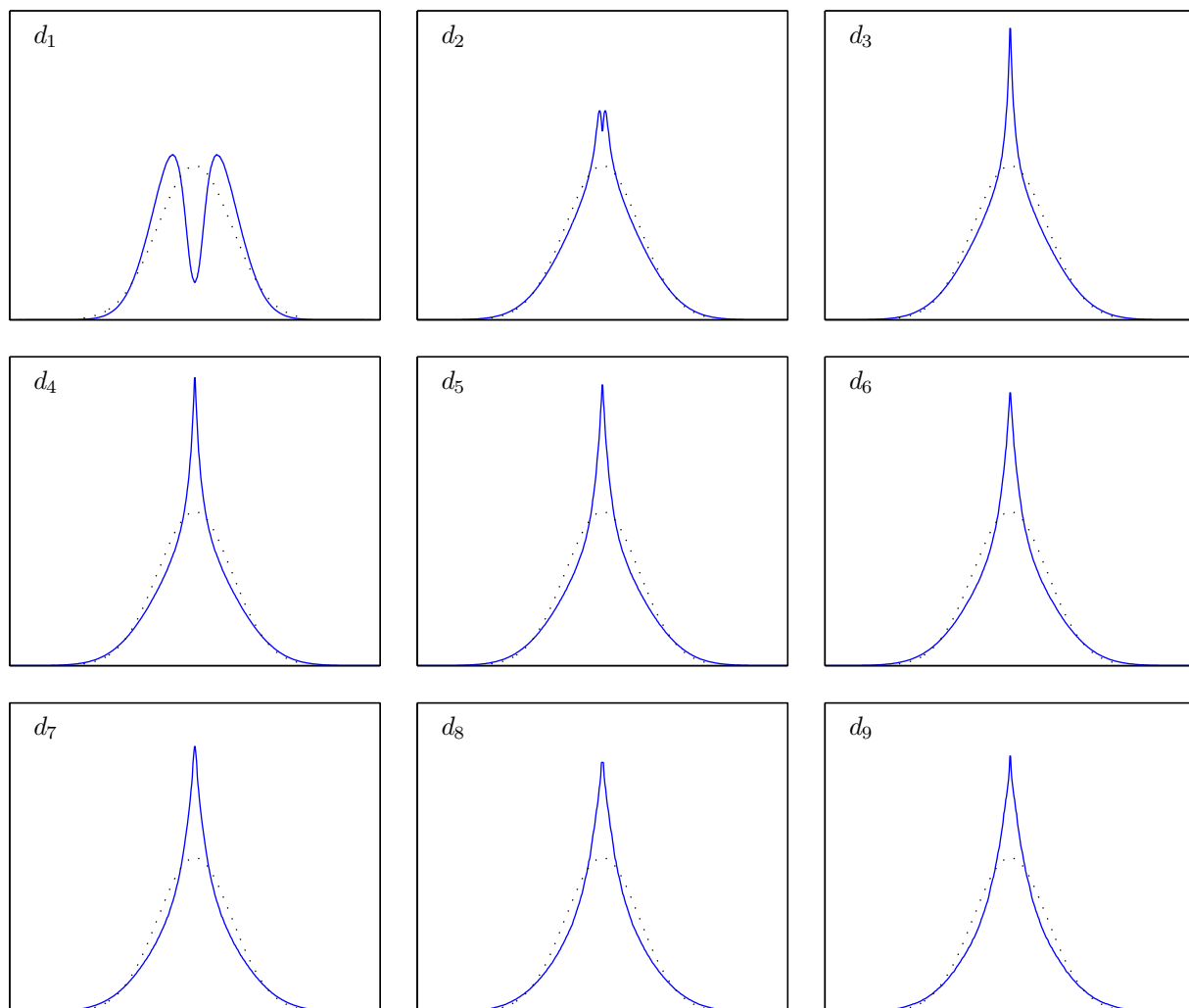


FIGURE 2.93 – Coupes des densités marginales des IMFs issues du bruit blanc analytique calculées par l'algorithme Algo. 4 avec 16 directions. Il ne s'agit pas des densités marginales des parties réelles (ou imaginaires) mais d'une coupe des densités bivariées Fig. 2.92. Si $\mathcal{P}_r(r)$ est la densité du module d'un IMF, ce qui est représenté est $\mathcal{P}_r(r)/r$ avec une partie négative rajoutée par symétrie. La courbe en pointillés est la gaussienne de référence.

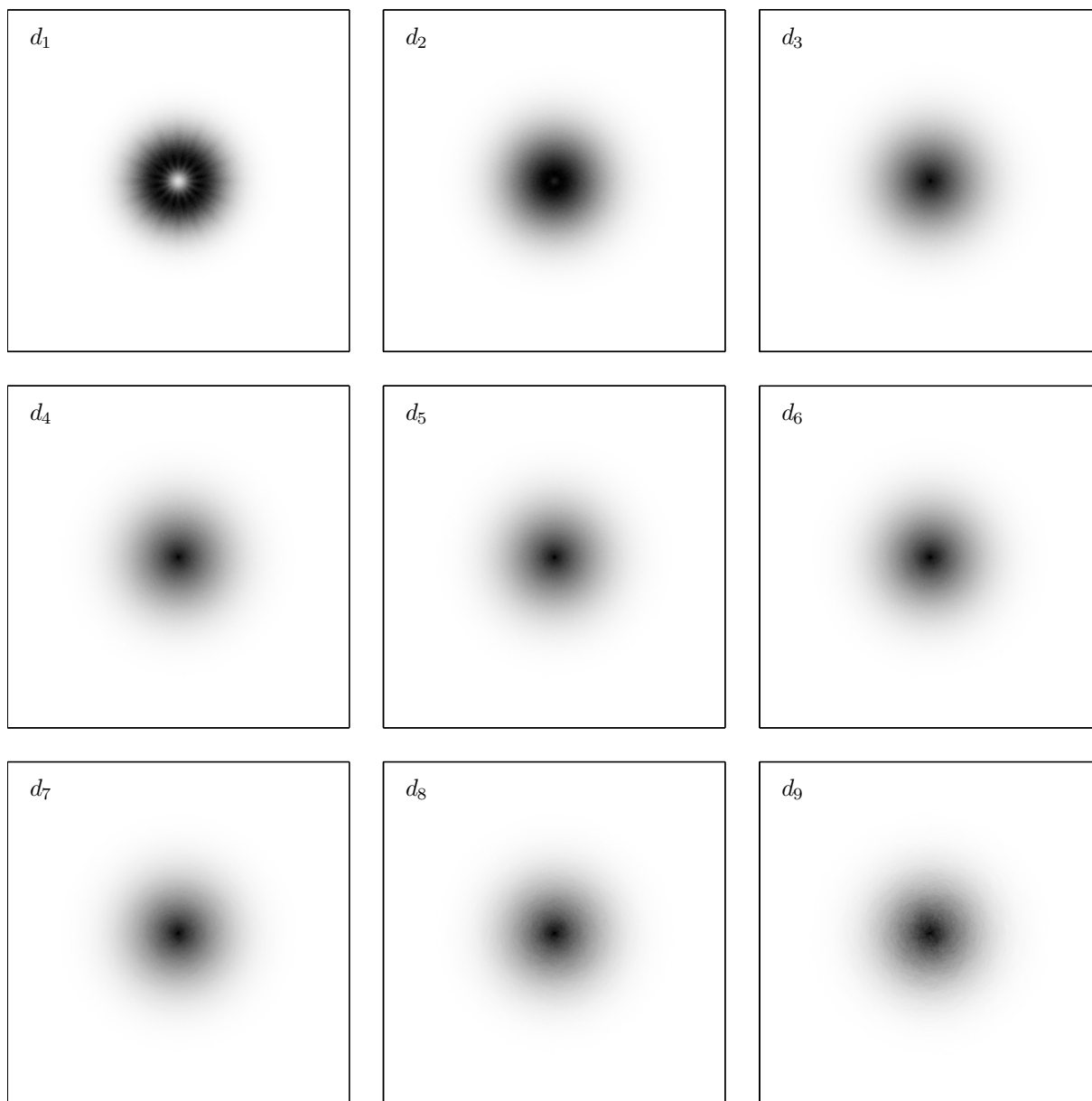


FIGURE 2.94 – Densités marginales des IMFs issus du bruit blanc analytique calculés par l'algorithme Algo. 5 avec 16 directions. (échelles normalisées)

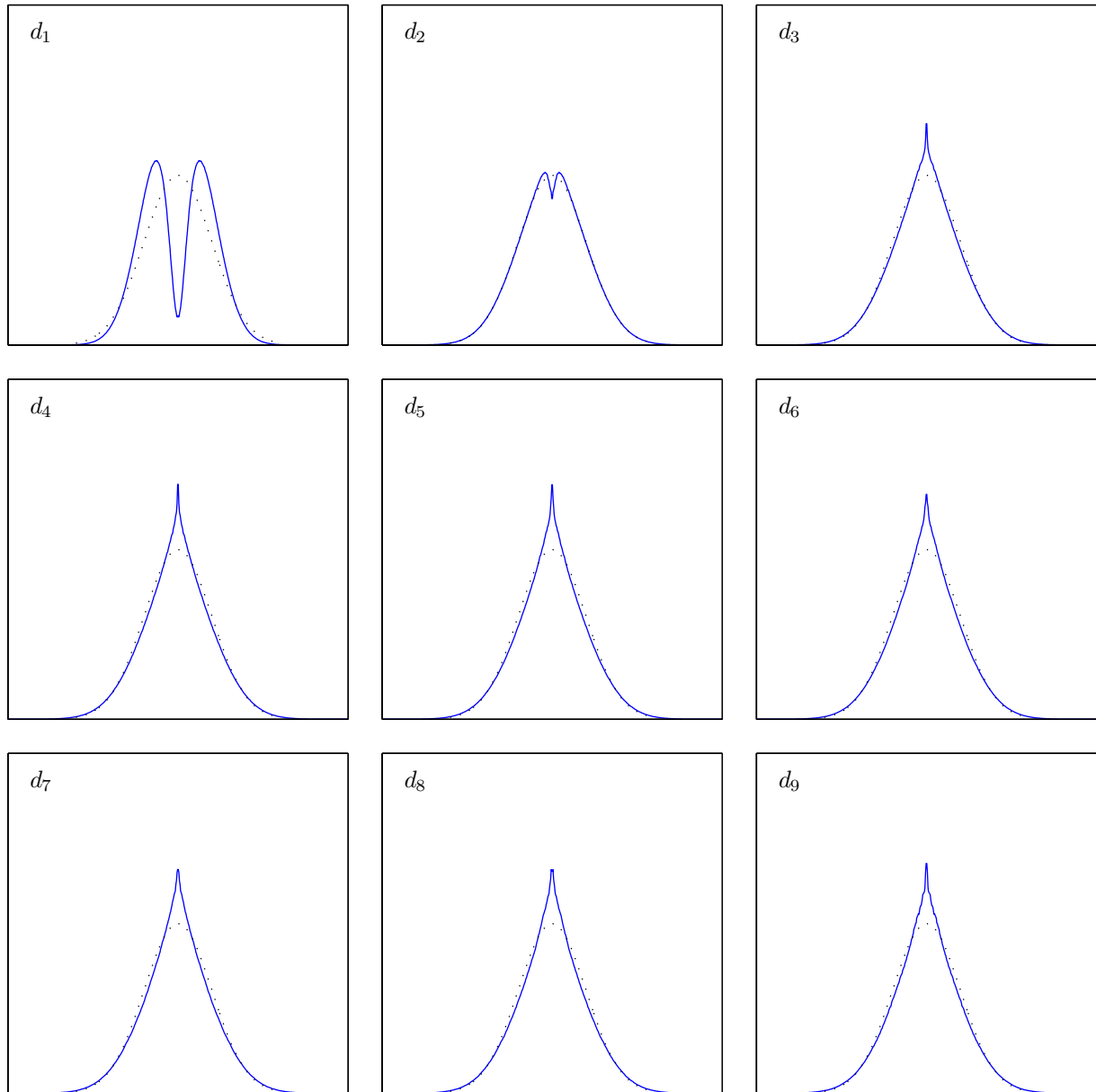


FIGURE 2.95 – Coupes des densités marginales des IMFs issues du bruit blanc analytique calculées par l'algorithme Algo. 5 avec 16 directions. Il ne s'agit pas des densités marginales des parties réelles (ou imaginaires) mais d'une coupe des densités bivariées Fig. 2.94. Si $\mathcal{P}_r(r)$ est la densité du module d'un IMF, ce qui est représenté est $\mathcal{P}_r(r)/r$ avec une partie négative rajoutée par symétrie. La courbe en pointillés est la gaussienne de référence.

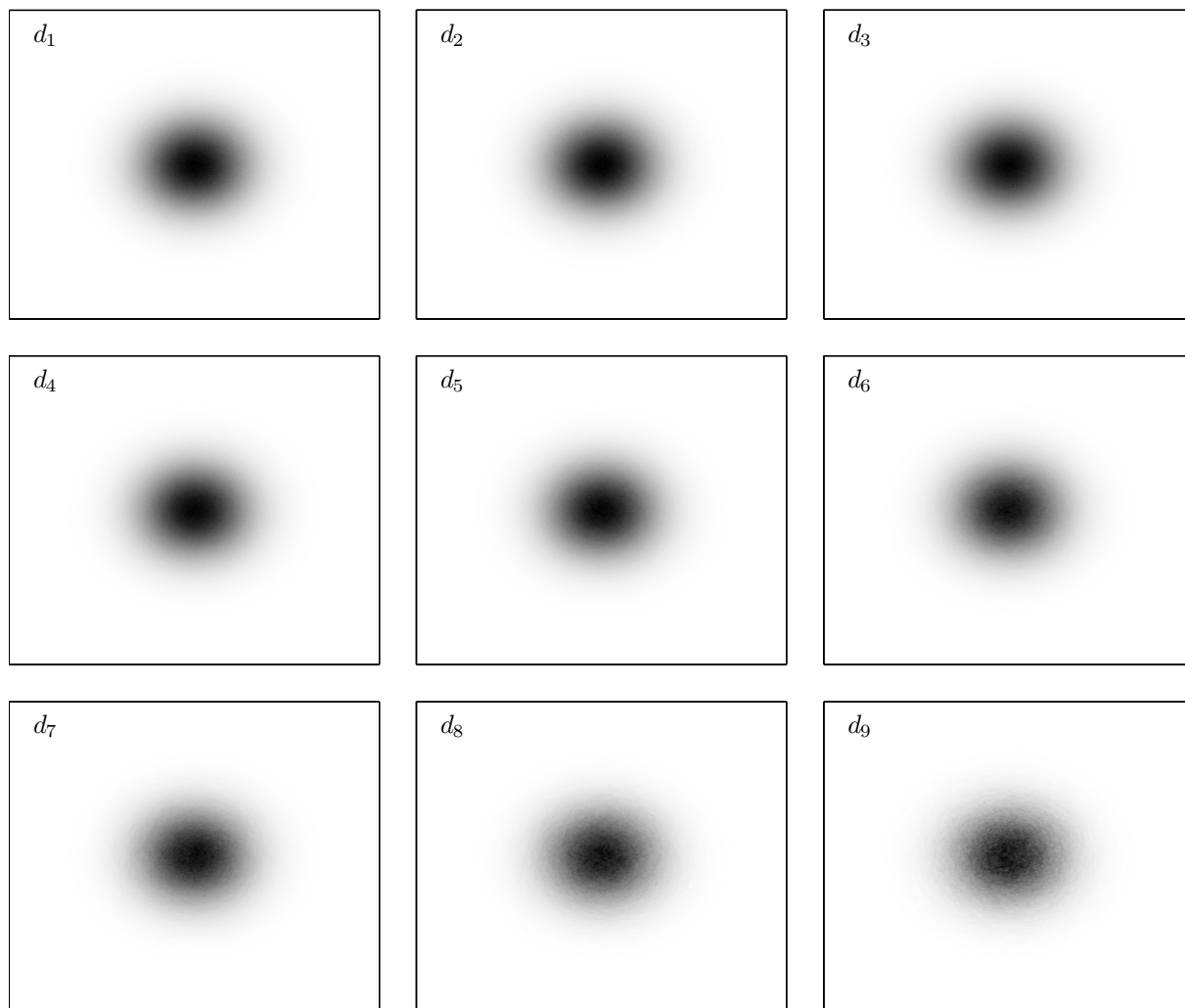


FIGURE 2.96 – Densités marginales des approximations issues du bruit blanc analytique calculées par l’algorithme Algo. 4 avec 16 directions. (échelles normalisées)

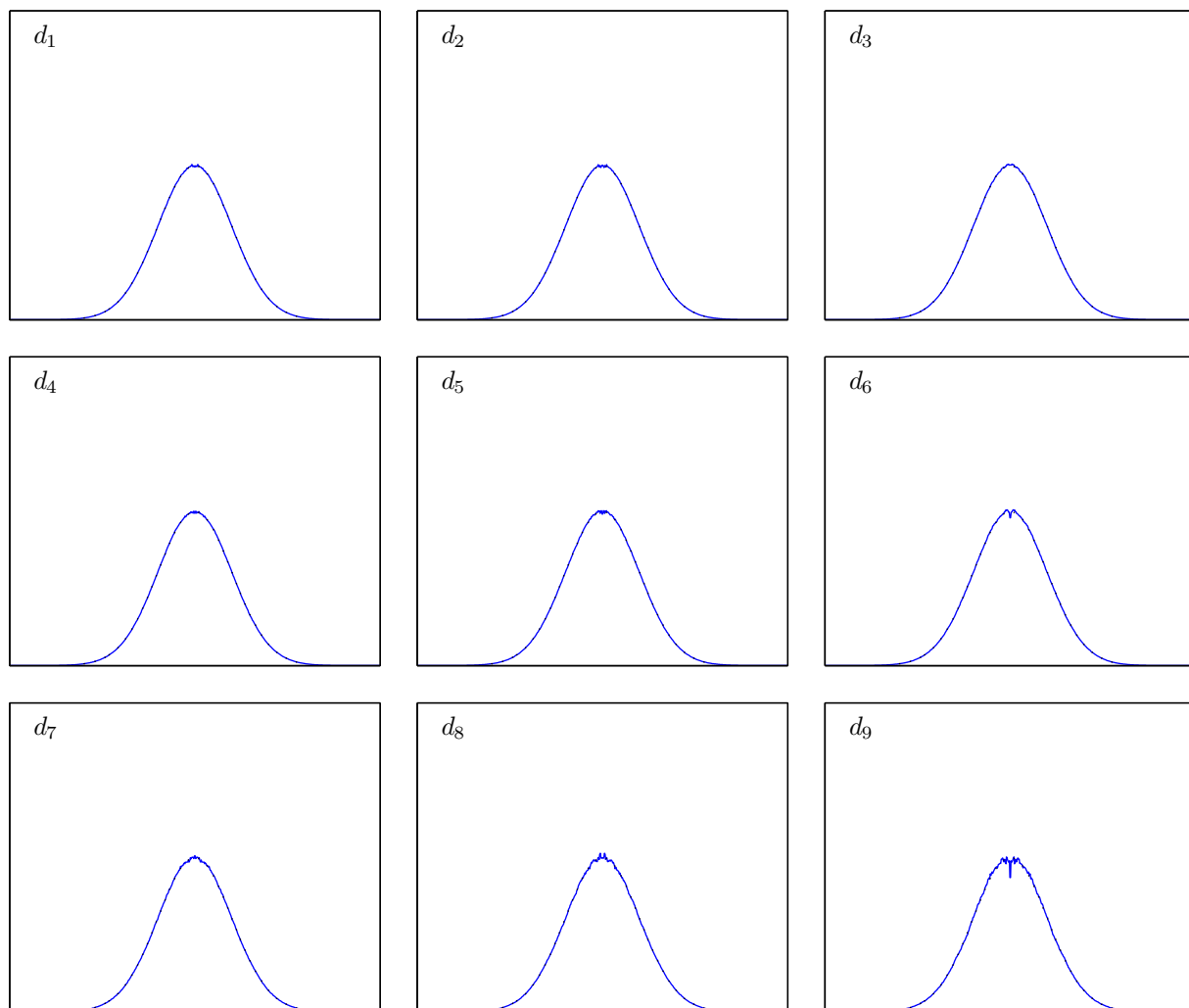


FIGURE 2.97 – Coupes des densités marginales des approximations issues du bruit blanc analytique calculées par l'algorithme Algo. 4 avec 16 directions. Il ne s'agit pas des densités marginales des parties réelles (ou imaginaires) mais d'une coupe des densités bivariées Fig. 2.96. Si $\mathcal{P}_r(r)$ est la densité du module d'un approximation, ce qui est représenté est $\mathcal{P}_r(r)/r$ avec une partie négative rajoutée par symétrie. La courbe en pointillés est la gaussienne de référence.

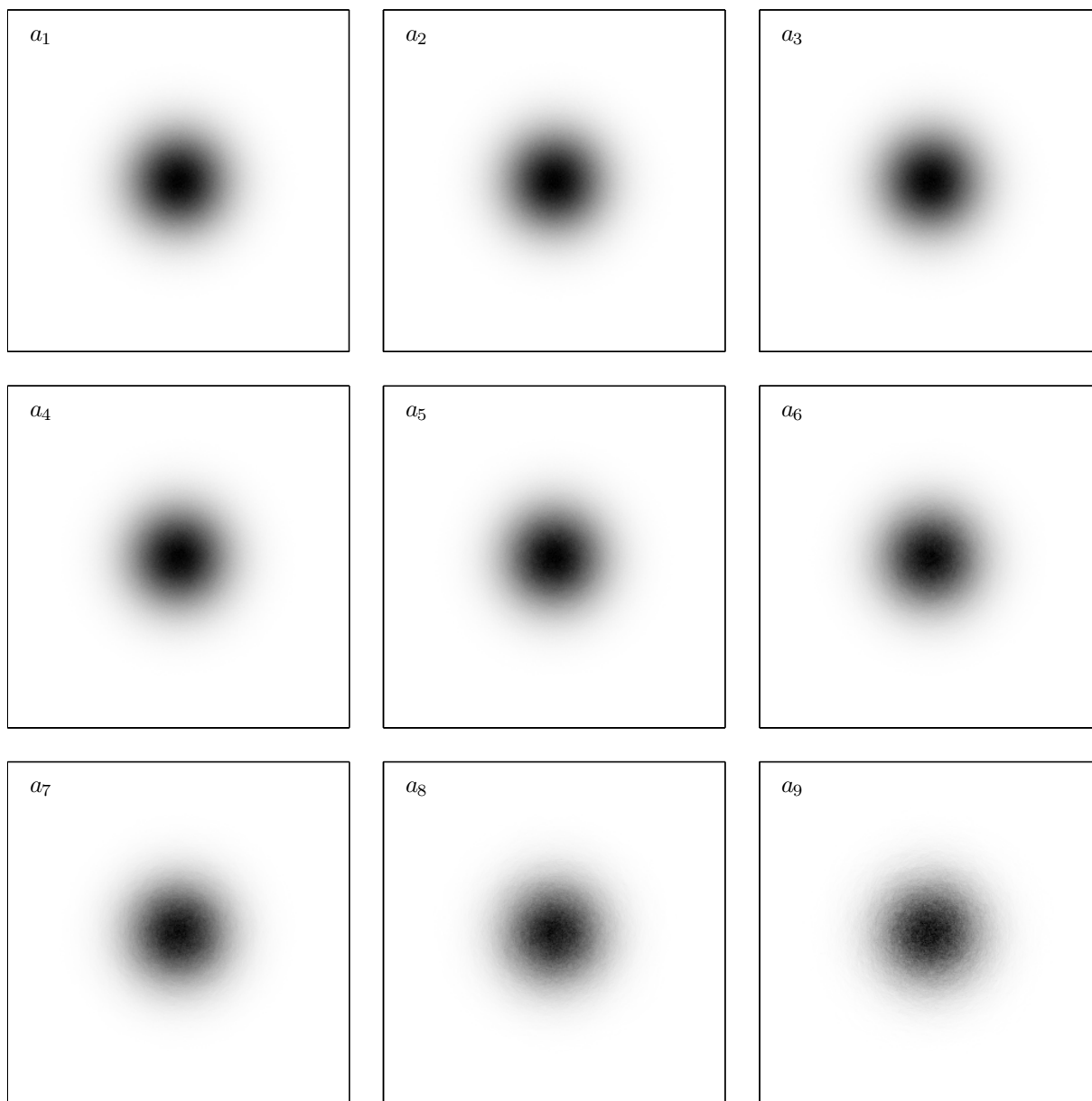


FIGURE 2.98 – Densités marginales des approximations issues du bruit blanc analytique calculées par l'algorithme Algo. 5 avec 16 directions. (échelles normalisées)

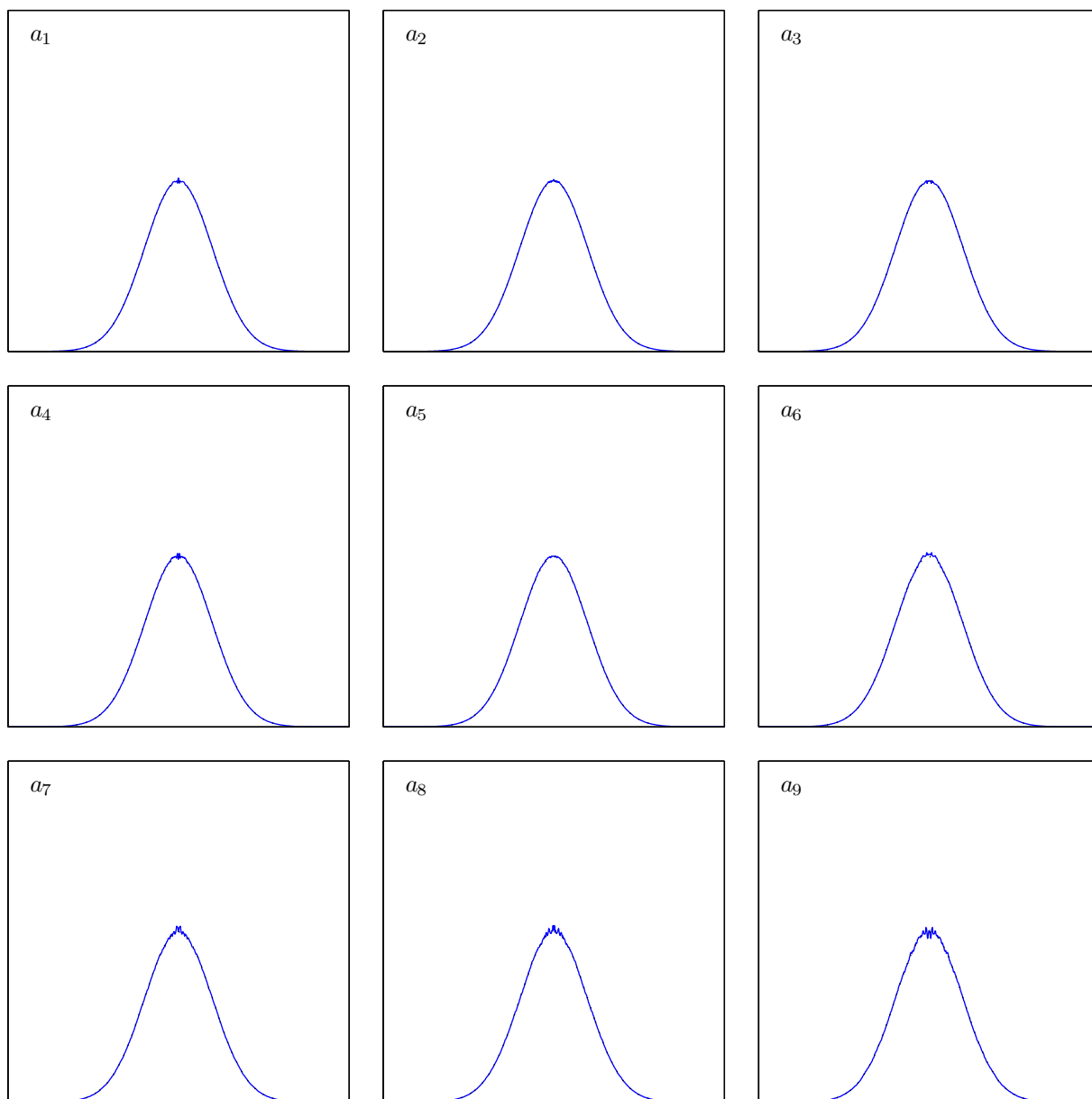


FIGURE 2.99 – Coupes des densités marginales des approximations issues du bruit blanc analytique calculées par l'algorithme Algo. 5 avec 16 directions. Il ne s'agit pas des densités marginales des parties réelles (ou imaginaires) mais d'une coupe des densités bivariées Fig. 2.98. Si $\mathcal{P}_r(r)$ est la densité du module d'un approximation, ce qui est représenté est $\mathcal{P}_r(r)/r$ avec une partie négative rajoutée par symétrie. La courbe en pointillés est la gaussienne de référence.

où $e_{\varphi_k}(t)$ est l'armature latérale associée à la direction u_{φ_k} . Si on divise chacune des armatures latérales en ses composantes colinéaire et orthogonale au vecteur u_{φ_k} , on obtient

$$\mathcal{S}^{B1}[x](t) = x(t) - \frac{1}{N} \sum_{k=1}^N e_{\varphi_k}^{\parallel} + \frac{1}{N} \sum_{k=1}^N e_{\varphi_k}^{\perp} \quad (2.219)$$

où la première somme est en fait la moitié de la moyenne telle que définie dans l'algorithme Algo. 5. On peut alors écrire l'opérateur de tamisage de l'algorithme Algo. 4 en fonction de celui de l'algorithme Algo. 5 :

$$\mathcal{S}^{B1}[x](t) = \frac{1}{2} \mathcal{S}^{B2}[x](t) + \frac{x(t)}{2} - \frac{1}{N} \sum_{k=1}^N e_{\varphi_k}^{\perp}. \quad (2.220)$$

Si l'on se base sur cette expression et qu'on suppose de plus que la deuxième partie de celle-ci est relativement invariante par rotation, on obtient que la densité de probabilité du premier IMF pour l'algorithme Algo. 4 doit présenter les mêmes « directions creuses » que dans le cas de l'algorithme Algo. 5 mais en moins accentué, ce qui correspond aux résultats expérimentaux. De plus, le fait que les lobes de la distribution soient légèrement plus resserrés que dans le cas de l'algorithme Algo. 5 est également cohérent avec ce raisonnement.

Lorsque le nombre N de directions augmente, on constate que la densité de probabilité tend à être invariante par rotation (cf Fig. 2.84, Fig. 2.86, Fig. 2.92, Fig. 2.94). La densité de probabilité asymptotique est le pendant à deux dimensions de la densité bimodale observée dans le cas réel : elle s'apparente à une gaussienne avec un creux important en zéro. L'évolution de l'invariance par rotation de cette densité de probabilité en fonction de N permet de plus d'évaluer la vitesse à laquelle les algorithmes convergent vers des algorithmes covariants par rapport aux rotations du plan, ou du moins une borne supérieure de cette vitesse. Les observations laissent entrevoir que pour 30 directions environ, les deux algorithmes seraient covariants par rapport aux rotations. Cependant, un tel résultat est à prendre avec précautions parce que la résolution des mesures ne permet pas vraiment d'observer ce qui se passe au-delà d'une trentaine de directions.

2.3.3 Statistiques d'ordre 2

Concernant, les statistiques d'ordre 2 on pourrait comme on l'a fait pour les bruits réels s'intéresser aux densités de probabilité à l'ordre 2. Celles-ci ont sans doute des informations à révéler mais le fait qu'elles soient de dimension 4 les rend difficiles à représenter et on s'abstiendra donc de les étudier ici.

Pour ce qui est des spectres et interspectres des IMFs et approximations, on observe des comportements très similaires à ceux de l'EMD classique d'un bruit blanc.

Dans tous les cas, les parties positives des spectres des IMFs (cf Fig. 2.108, Fig. 2.109, Fig. 2.110 et Fig. 2.111) ont la même allure que ceux des IMFs et approximations issus de l'EMD d'un bruit blanc réel : en première approximation, tout se passe comme si les algorithmes d'EMD bivariée agissaient comme des bancs de filtres. Les parties négatives sont quant à elles symétriques des parties positives pour le bruit blanc simple et relativement faibles mais non négligeables pour le bruit blanc analytique. Si on s'intéresse plus en détail aux parties positives des spectres, on peut schématiser les dépendances vis à vis des différents paramètres ainsi :

algorithme les deux algorithmes ont du point de vue des spectres des comportement très similaires qui diffèrent essentiellement par le facteur d'auto-similarité entre les spectres des IMFs (autres que le premier), et encore assez peu dans le cas du bruit blanc analytique.

modèle de bruit les spectres des IMFs issus des deux types de bruits diffèrent à la fois par le facteur d'auto-similarité (surtout pour l'algorithme Algo. 4) et par la forme des spectres des

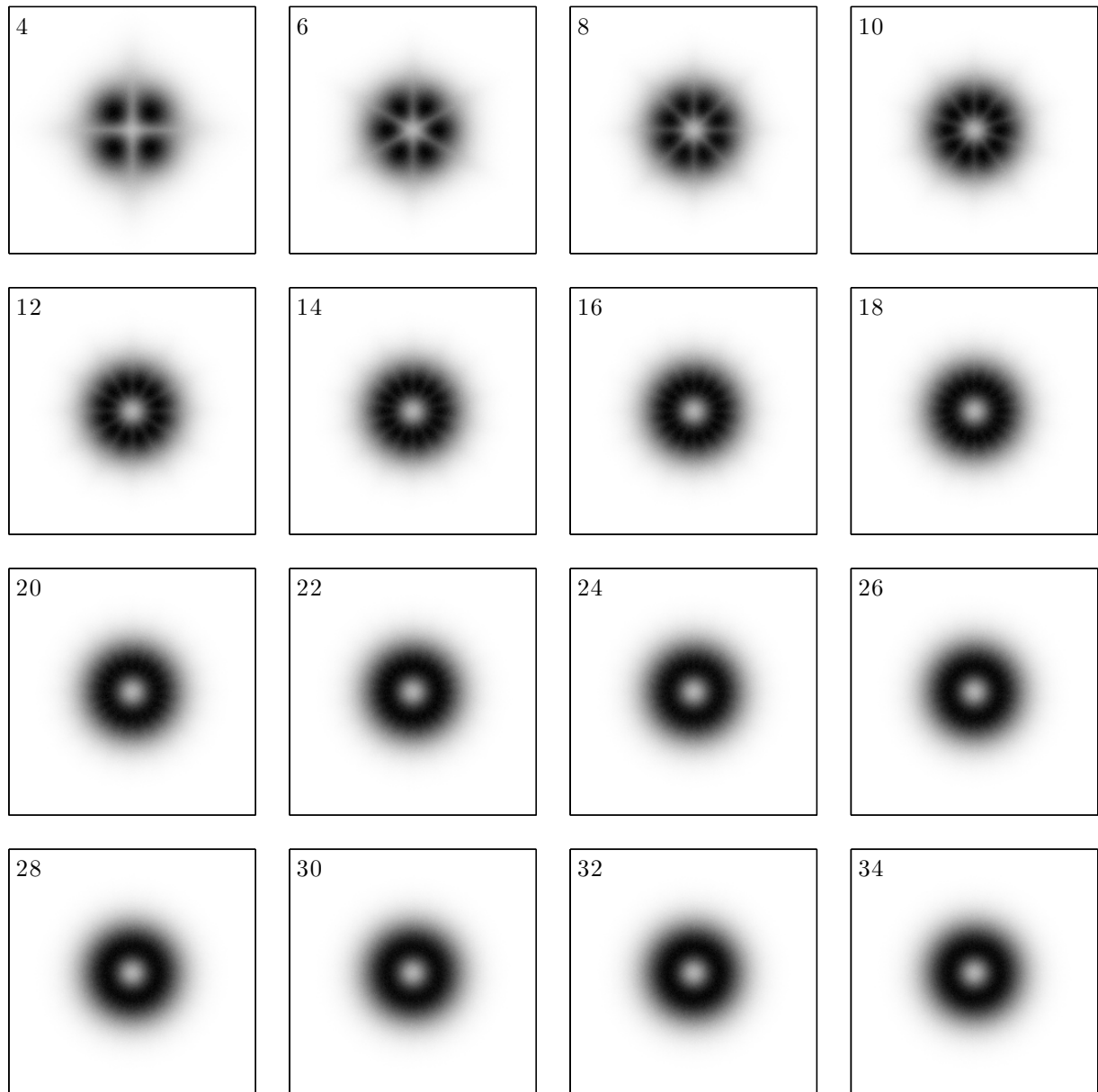


FIGURE 2.100 – Densité marginale du premier IMF issu du bruit blanc simple calculé par l'algorithme Algo. 4 avec de 2 à 34 directions.

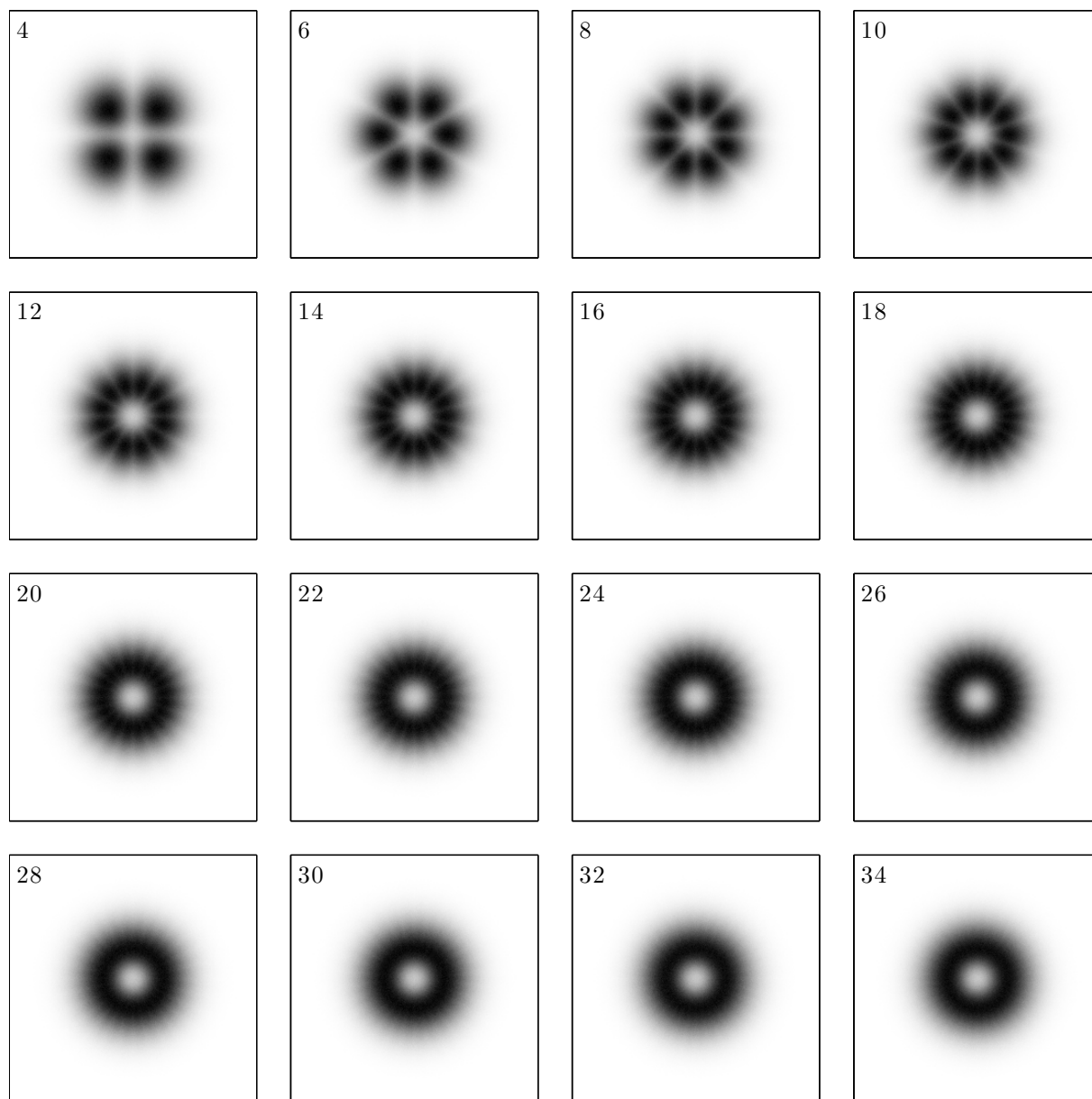


FIGURE 2.101 – Densité marginale du premier IMF issu du bruit blanc simple calculé par l'algorithme Algo. 5 avec de 2 à 34 directions.

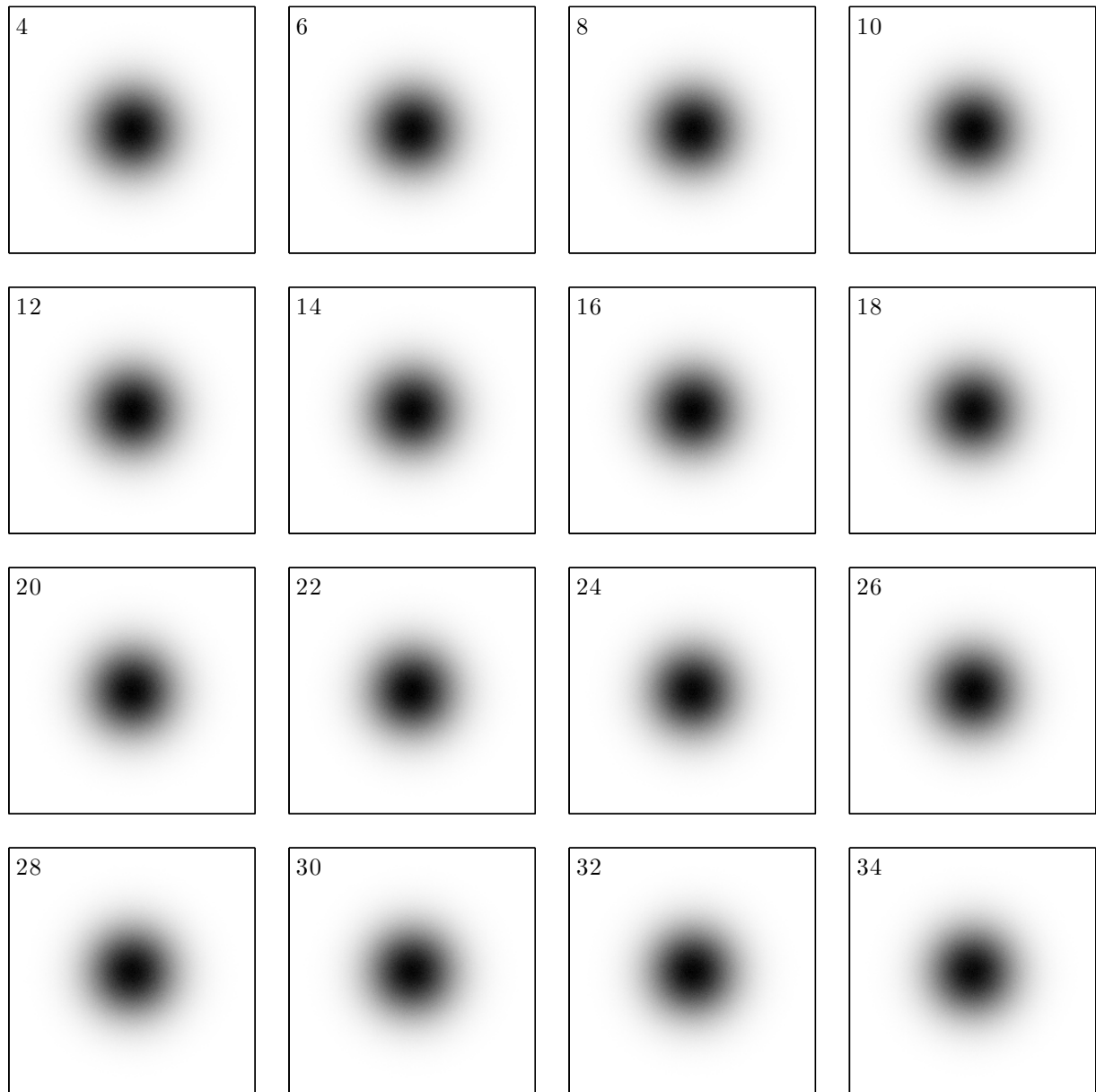


FIGURE 2.102 – Densité marginale de la première approximation issue du bruit blanc simple calculée par l'algorithme Algo. 4 avec de 2 à 34 directions.

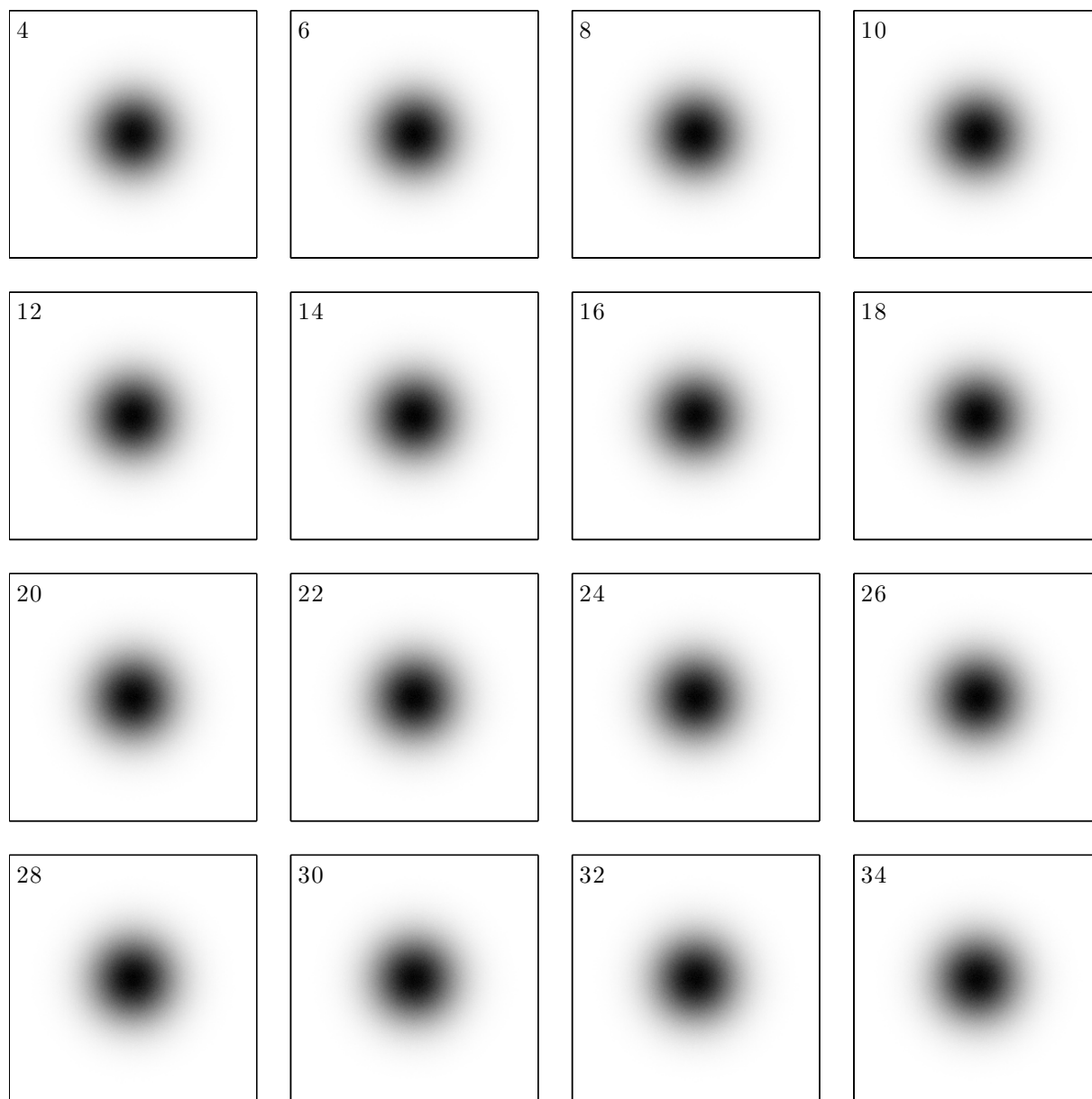


FIGURE 2.103 – Densité marginale de la première approximation issue du bruit blanc simple calculée par l'algorithme Algo. 5 avec de 2 à 34 directions.

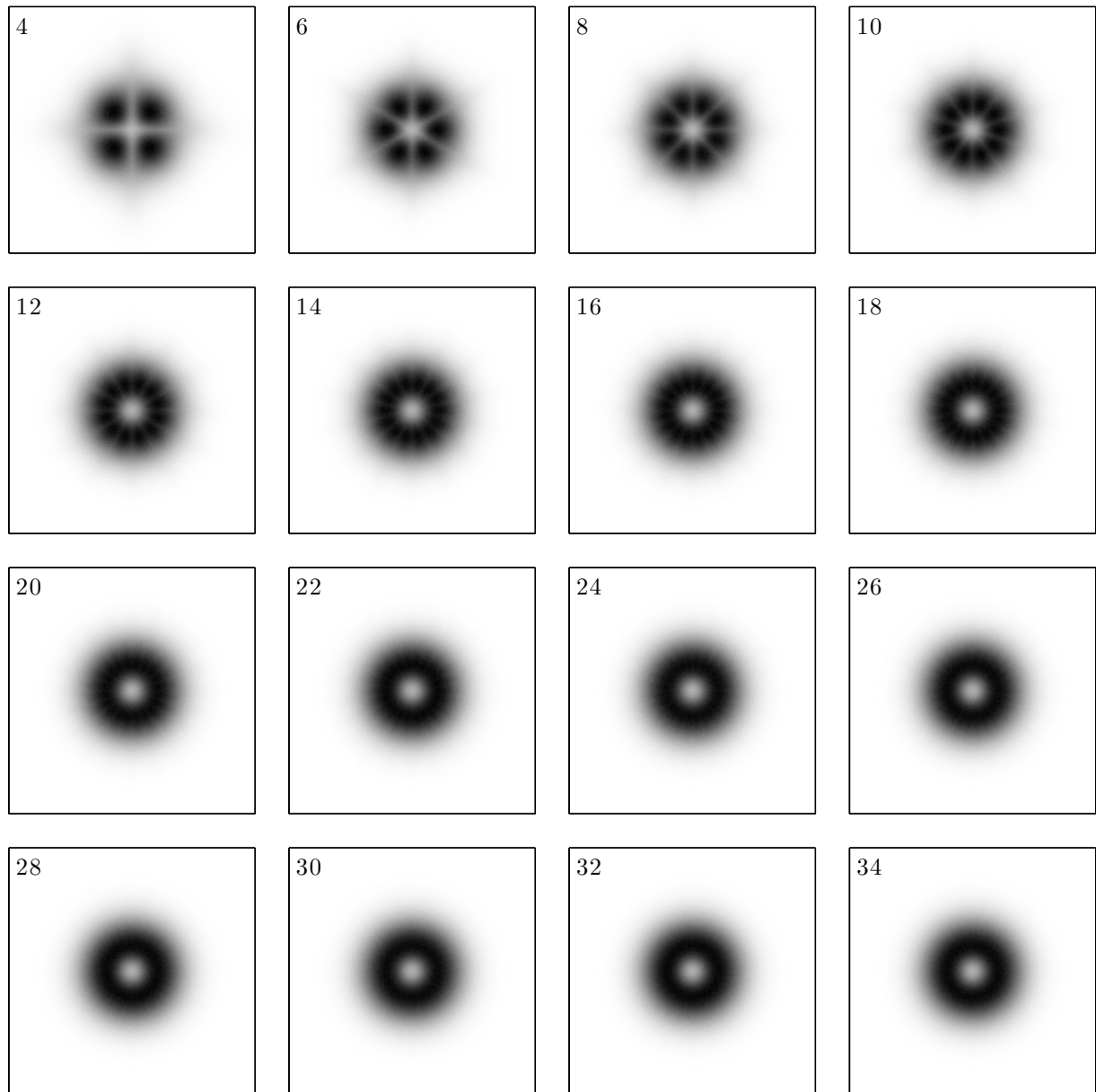


FIGURE 2.104 – Densité marginale du premier IMF issu du bruit blanc simple calculé par l'algorithme Algo. 4 avec de 2 à 34 directions.

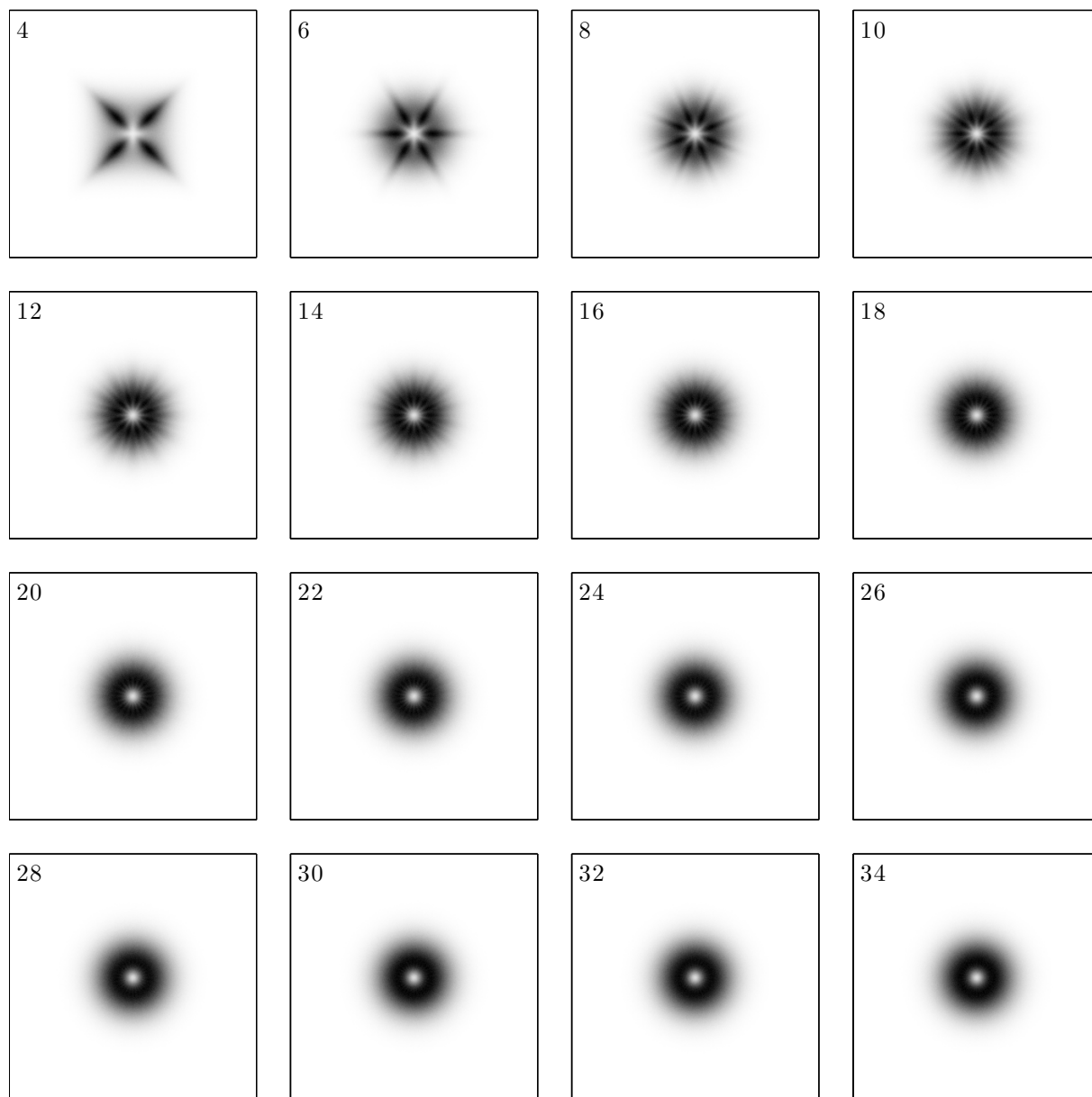


FIGURE 2.105 – Densité marginale du premier IMF issu du bruit blanc analytique calculé par l'algorithme Algo. 5 avec de 2 à 34 directions.

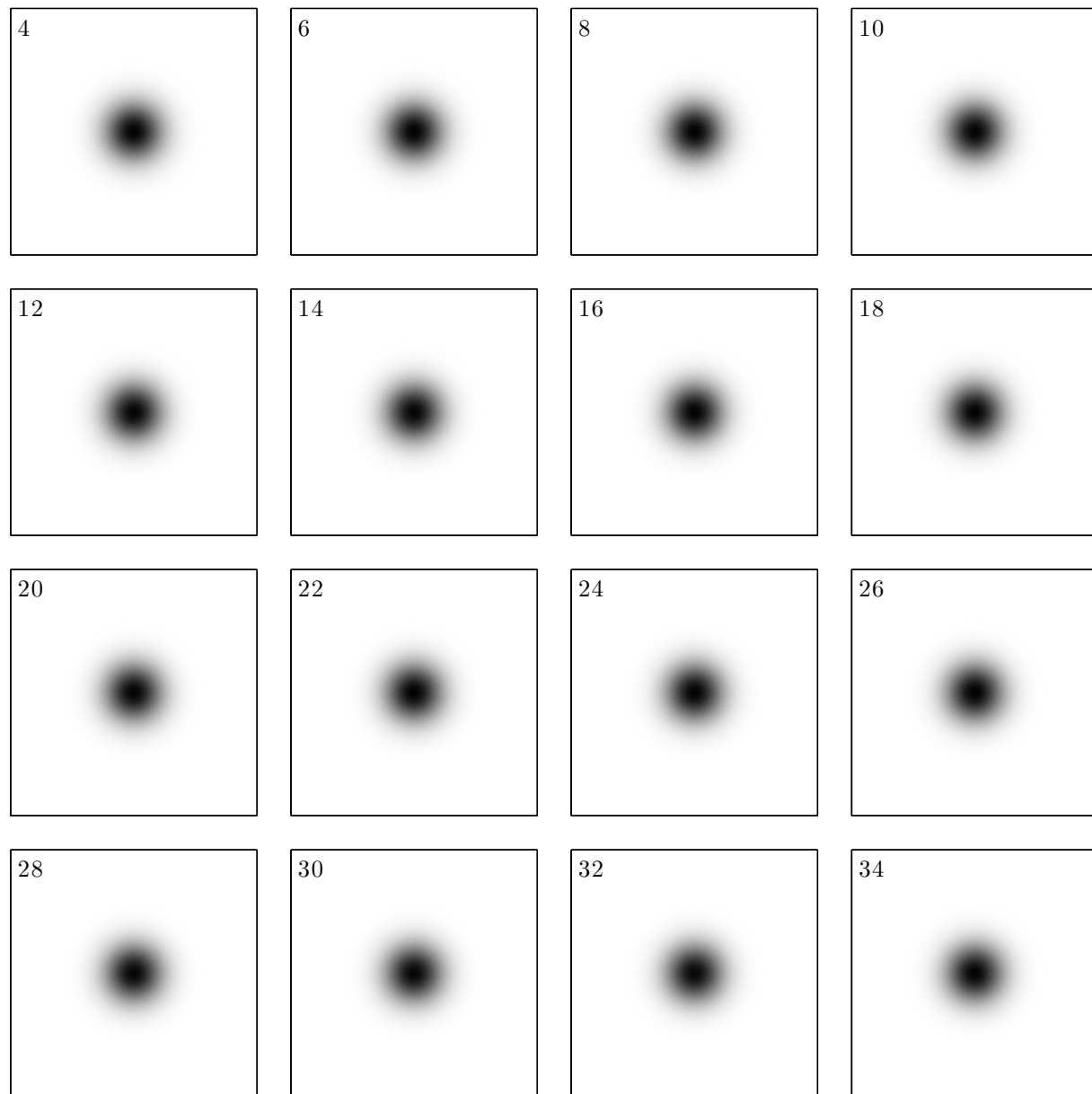


FIGURE 2.106 – Densité marginale de la première approximation issue du bruit blanc analytique calculée par l'algorithme Algo. 4 avec de 2 à 34 directions.

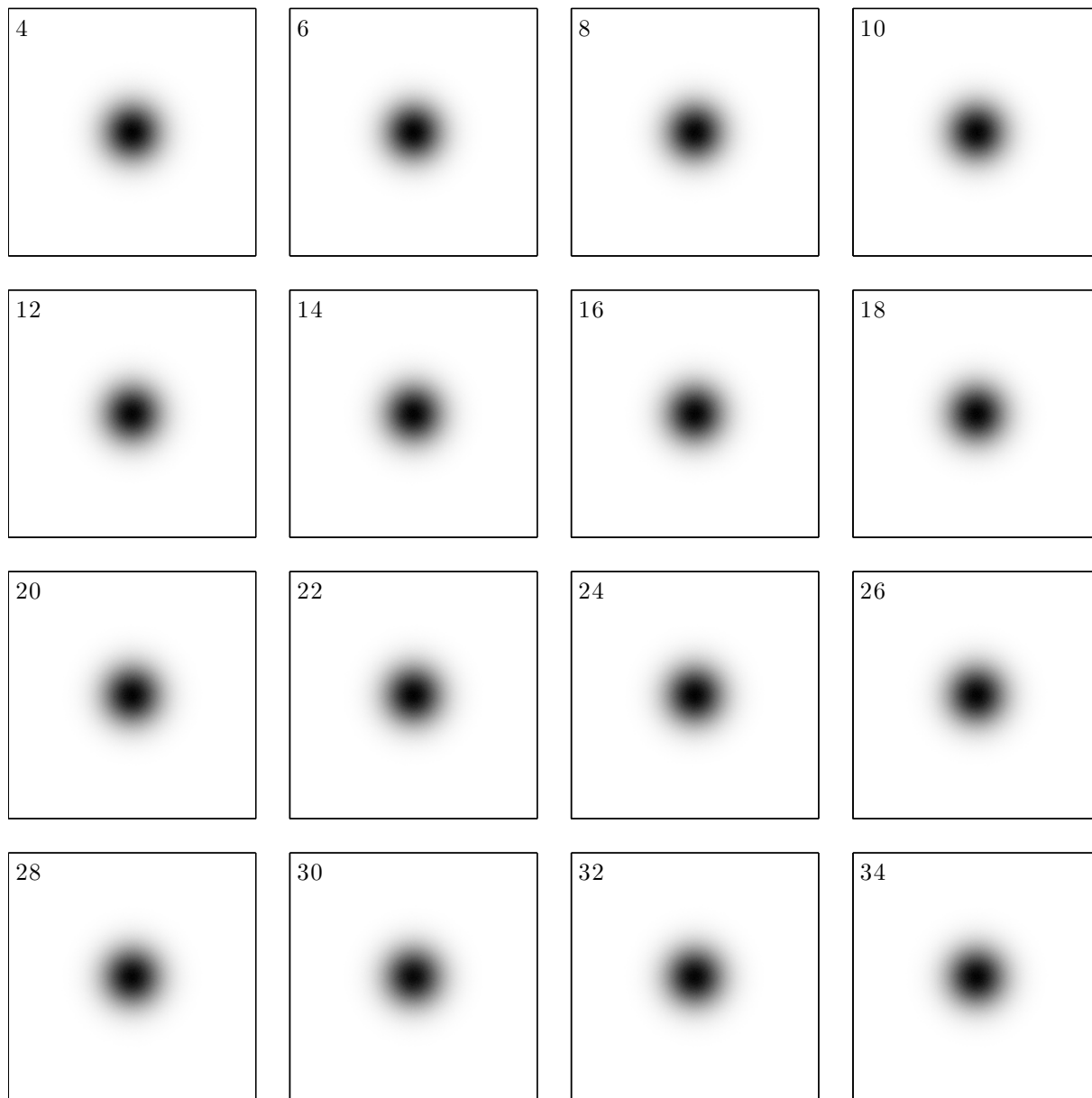


FIGURE 2.107 – Densité marginale de la première approximation issue du bruit blanc analytique calculée par l'algorithme Algo. 5 avec de 2 à 34 directions.

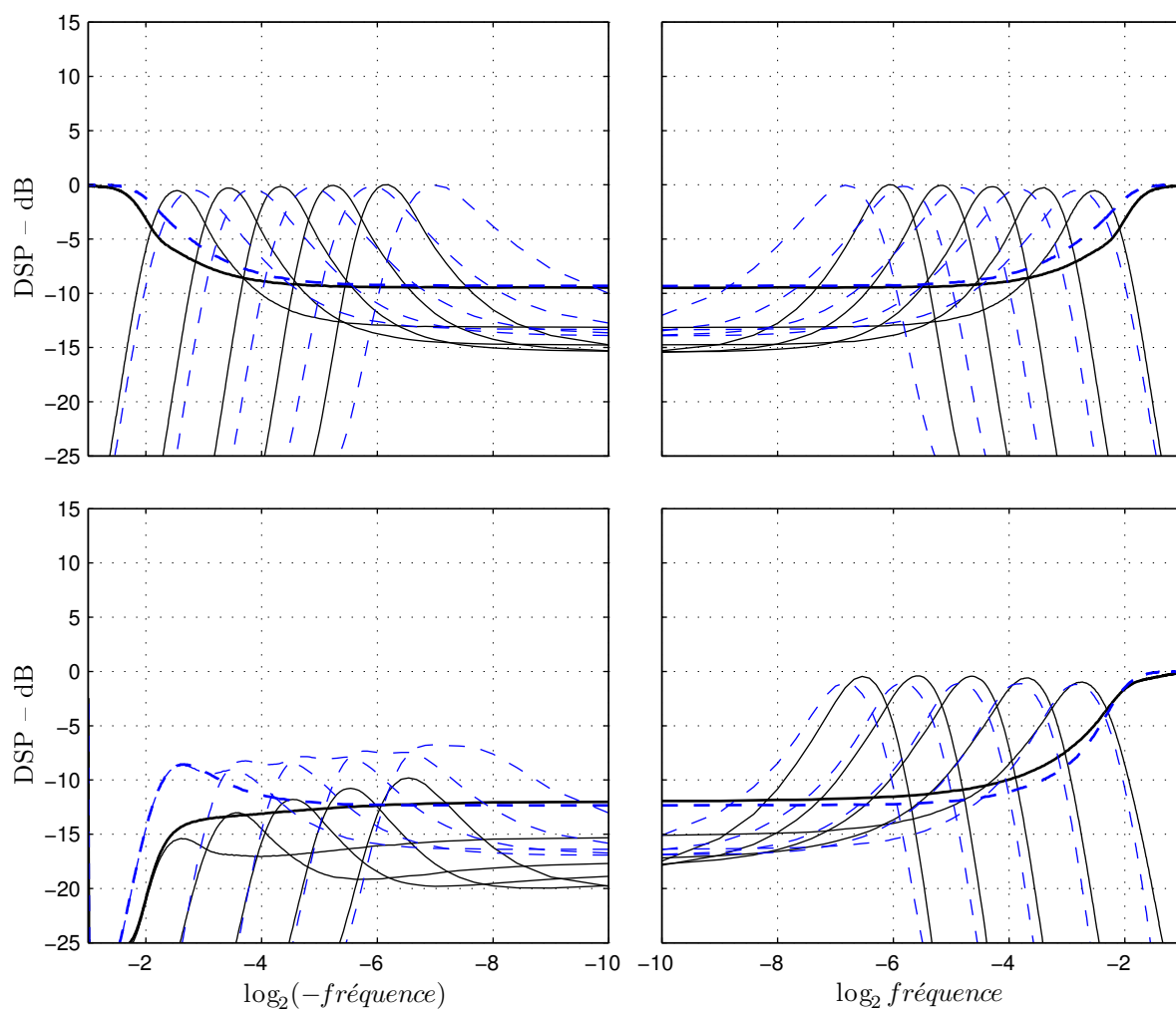


FIGURE 2.108 – Spectres des IMFs obtenus à l'aide de 10 itérations de tamisage par composante et 4 directions. En haut le bruit blanc simple, en bas le bruit blanc analytique. Les spectres tracés en trait plein noir correspondent à l'algorithme Algo. 4. Ceux en tirets bleus correspondent à l'algorithme Algo. 5.

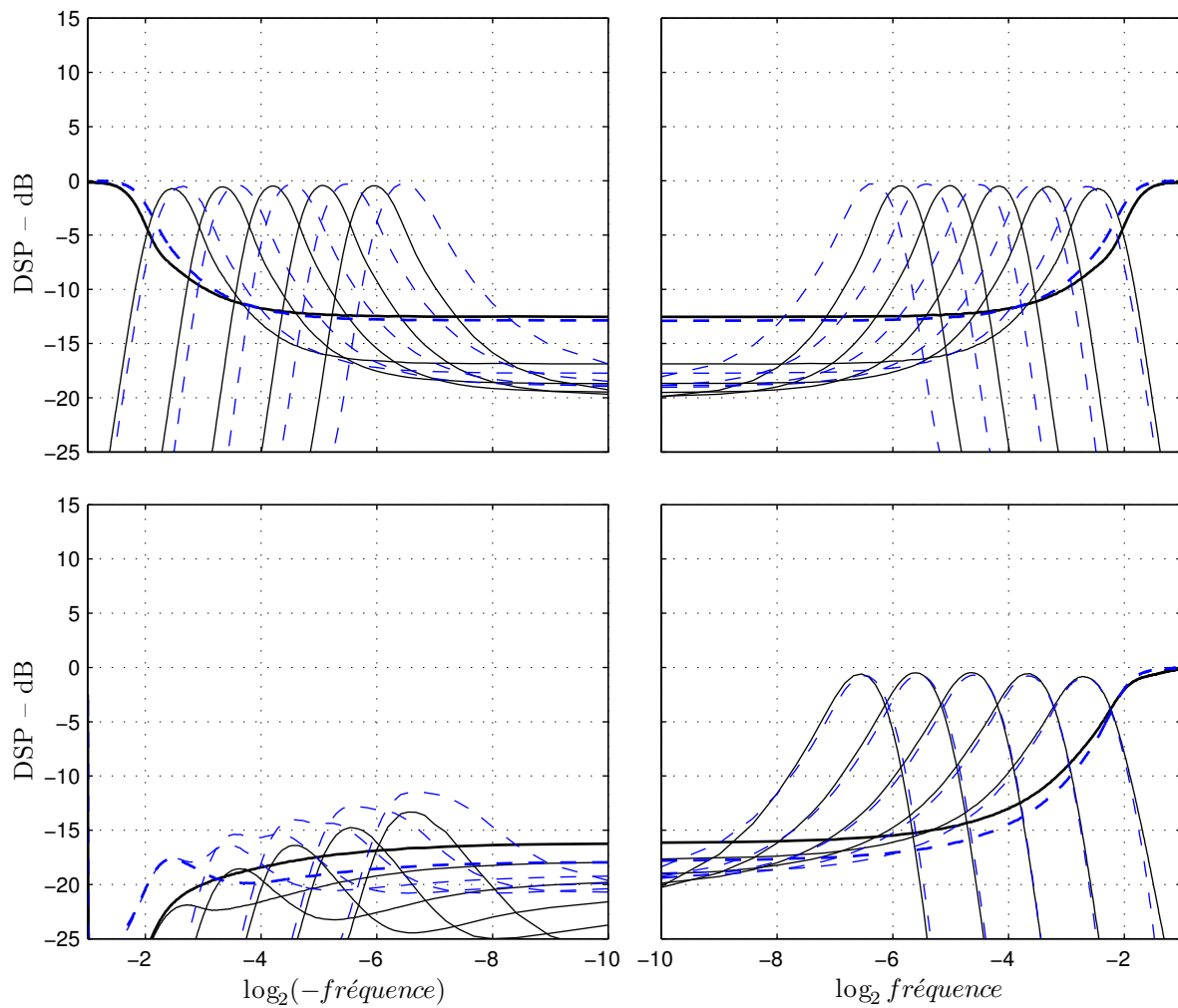


FIGURE 2.109 – Spectres des IMFs obtenus à l’aide de 10 itérations de tamisage par composante et 8 directions. En haut le bruit blanc simple, en bas le bruit blanc analytique. Les spectres tracés en trait plein noir correspondent à l’algorithme Algo. 4. Ceux en tirets bleus correspondent à l’algorithme Algo. 5.

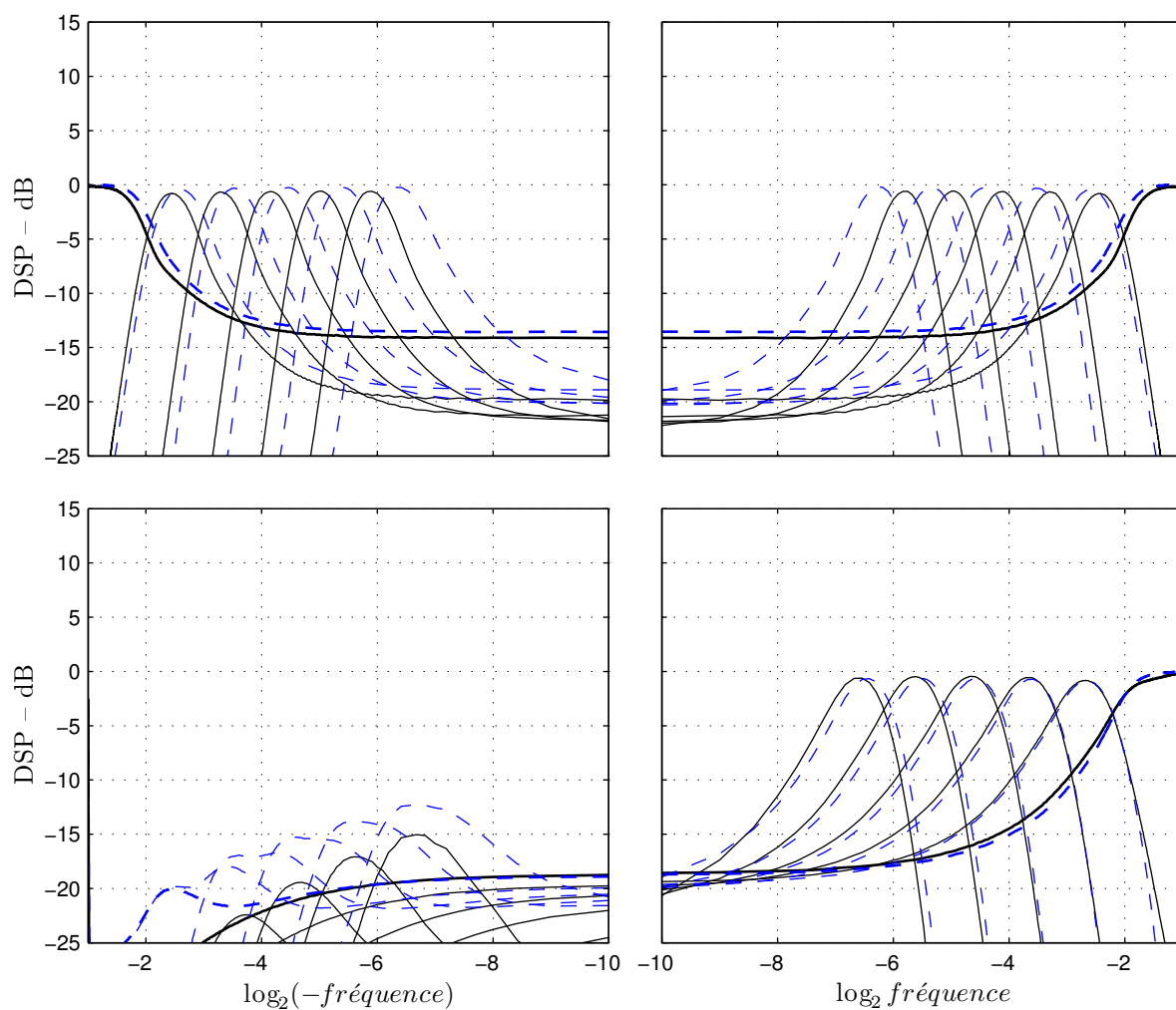


FIGURE 2.110 – Spectres des IMFs obtenus à l’aide de 10 itérations de tamisage par composante et 32 directions. En haut le bruit blanc simple, en bas le bruit blanc analytique. Les spectres tracés en trait plein noir correspondent à l’algorithme Algo. 4. Ceux en tirets bleus correspondent à l’algorithme Algo. 5.

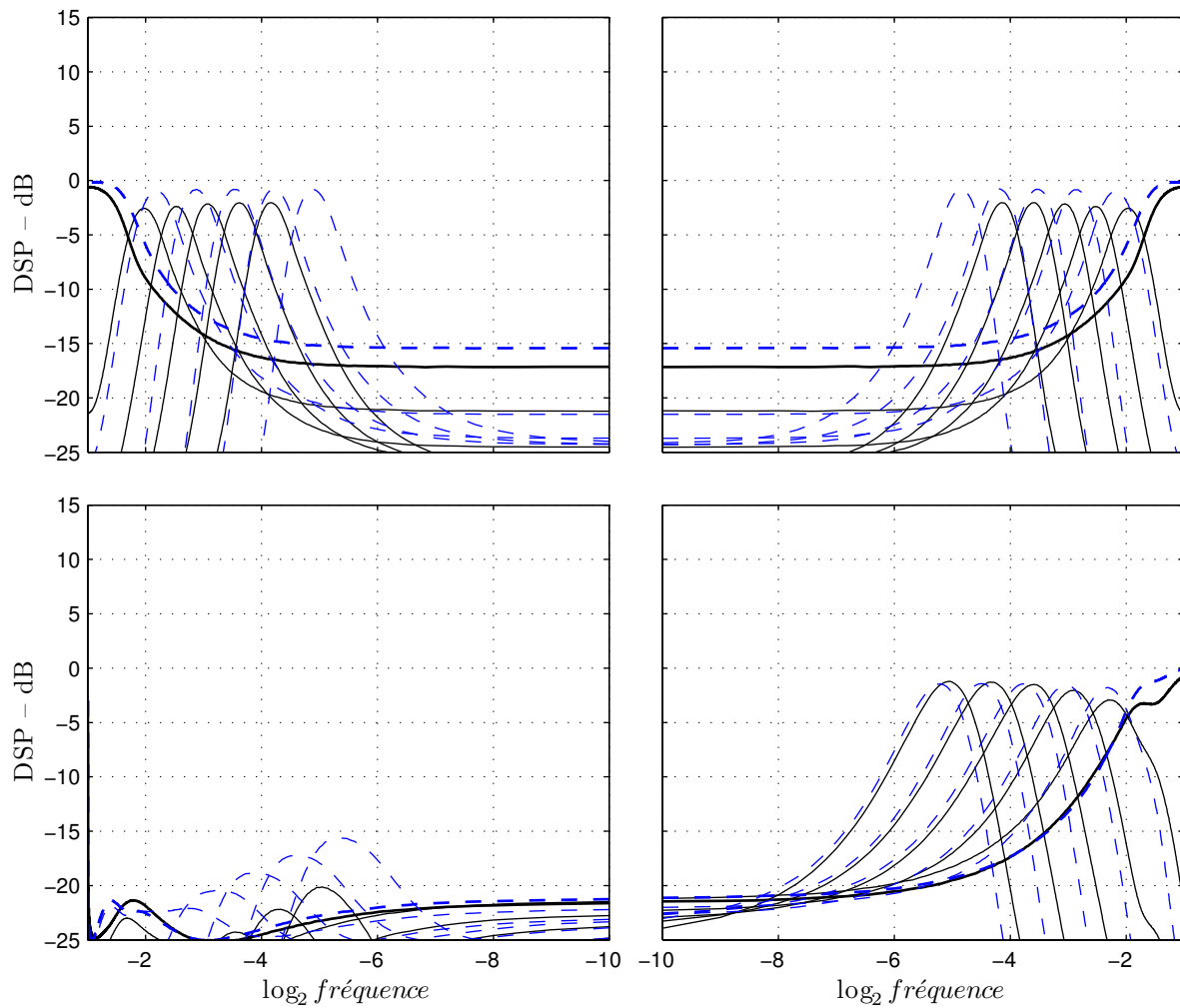


FIGURE 2.111 – Spectres des IMFs obtenus à l’aide de 50 itérations de tamisage par composante et 32 directions. En haut le bruit blanc simple, en bas le bruit blanc analytique. Les spectres tracés en trait plein noir correspondent à l’algorithme Algo. 4. Ceux en tirets bleus correspondent à l’algorithme Algo. 5.

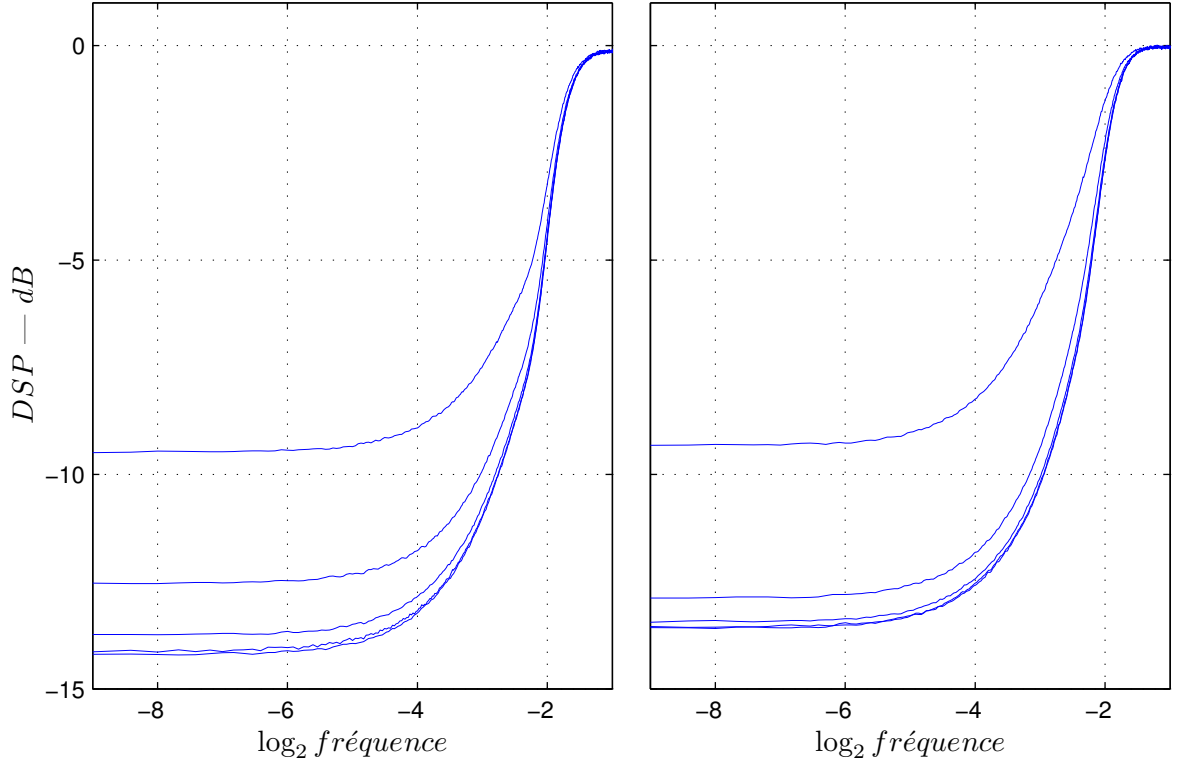


FIGURE 2.112 – Spectre du premier IMF pour le bruit blanc simple pour 4, 8, 16, 32 et 64 directions. À gauche algorithme Algo. 4. À droite algorithme Algo. 5. L'amplitude de la partie basse fréquence du premier IMF décroît quand le nombre de directions augmente mais cette décroissance semble se stabiliser à partir de 32 directions.

IMFs. Ce dernier aspect est surtout marqué sur le premier IMF : celui-ci a une composante basse fréquence nettement plus importante dans le cas du bruit blanc simple.

nombre d'itérations de tamisage le nombre d'itérations de tamisage influence globalement la forme des spectres des IMFs mais pour des nombres d'itérations pas trop grands, son influence se résume, comme dans le cas de l'EMD classique d'un bruit blanc réel, essentiellement au facteur d'auto-similarité entre les spectres des IMFs (cf Fig. 2.113).

nombre de directions le nombre de directions influe à la fois sur les facteurs d'auto-similarité et sur l'importance de la partie basse fréquence du spectre du premier IMF qui décroissent tous deux quand le nombre de directions augmente. Toutefois, il semble que cette décroissance soit limitée puisque le nombre de directions semble ne plus avoir d'influence au-delà de 16 environ sur les facteurs d'auto-similarité et au-delà de 32 environ pour les parties basses fréquences des spectres des premiers IMFs (cf Fig. 2.112).

Les interspectres entre les IMFs sont réels tout comme dans le cas de l'EMD classique d'un bruit blanc. L'explication est à nouveau liée au fait que l'EMD ne privilégie pas de direction temporelle (covariance par rapport aux inversions du temps) à quoi s'ajoute le fait que les algorithmes bivariés n'ont pas non plus de sens de rotation privilégié (covariance par rapport à la conjugaison). Grâce à ces deux propriétés, on peut montrer que l'intercorrélation entre deux IMFs est à symétrie hermitienne, ce qui implique que l'interspectre est réel. Considérons deux IMFs d'indices k et k' issus du signal $x(t)$. L'intercorrélation pour un décalage τ s'écrit

$$r_{k,k'}(\tau) = \mathbb{E} d_k[x](t) d_{k'}^*[x](t + \tau). \quad (2.221)$$

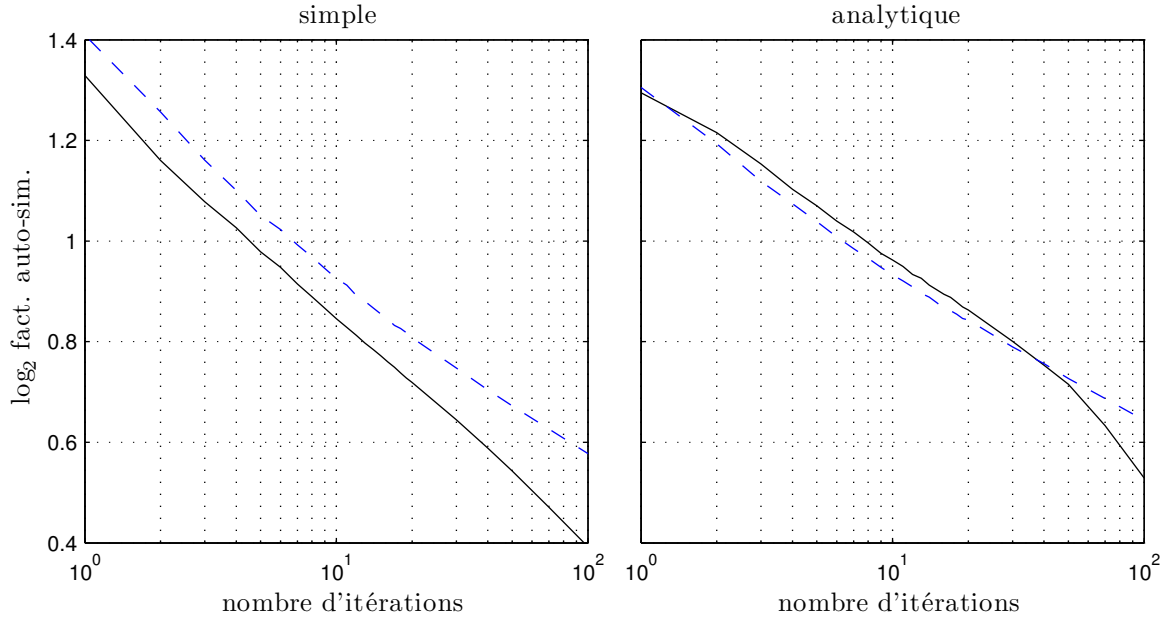


FIGURE 2.113 – Évolution du facteur d'auto-similarité du banc de filtres en fonction du nombre d'itérations pour les deux types de bruit. Les courbes correspondant à l'algorithme Algo. 4 sont en trait plein noir. Celles correspondant à l'algorithme Algo. 5 sont en tirets bleus.

En utilisant la covariance par rapport à la conjugaison, on obtient

$$\mathbb{E}d_k[x](t)d_{k'}^*[x](t+\tau) = \mathbb{E}d_k^*[x^*](t)d_{k'}[x^*](t+\tau). \quad (2.222)$$

Or le processus aléatoire $x(t)$ étant stationnaire gaussien et centré, $x^*(t)$ est identique à $x(-t)$ et la covariance par rapport aux inversions du temps donne

$$\mathbb{E}d_k[x](t)d_{k'}^*[x](t+\tau) = \mathbb{E}d_k^*[x](-t)d_{k'}[x](-t-\tau), \quad (2.223)$$

ce qui permet de conclure que l'intercorrélacion est bien à symétrie hermitienne.

Plus en détail, les allures des interspectres sont très similaires à celles observées dans le cas réel (cf Fig. 2.114, Fig. 2.115, Fig. 2.116 et Fig. 2.117). On retrouve ainsi que les IMFs dont les indices se succèdent sont positivement corrélés dans la bande de fréquence entre celles où ces deux IMFs ont leur contributions maximales puis négativement dans la partie plus basse fréquence de la bande de l'IMF de plus grand indice. De même les IMFs éloignés de plus d'un indice sont négativement corrélés dans la bande correspondant à l'IMF de plus grand indice. Enfin, au-delà de ces ressemblances morphologiques, les amplitudes des corrélacions, mesurées par les normes des interspectres, sont également très semblables aux cas de bruits blancs réels (cf Fig. 2.118).

2.4 Liens avec le modèle déterministe

L'étude du comportement de l'EMD sur des bruits large bande a montré que dans bien des cas, celui-ci s'apparente à celui d'un banc de filtres. Ce comportement est loin d'être évident à expliquer mais il n'est pas inintéressant de le rapprocher de celui de l'EMD sur des sommes de sinusoides et autres composantes périodiques simples, où on a pu mettre en évidence des effets de filtrage linéaire. Le modèle proposé alors étant valable dès lors que le signal a des extrema uniformément répartis, il est clair qu'on ne peut l'appliquer directement aux cas de bruits. Cependant, on a également observé pour les sommes de composantes périodiques que le comportement de l'EMD était bien décrit par le

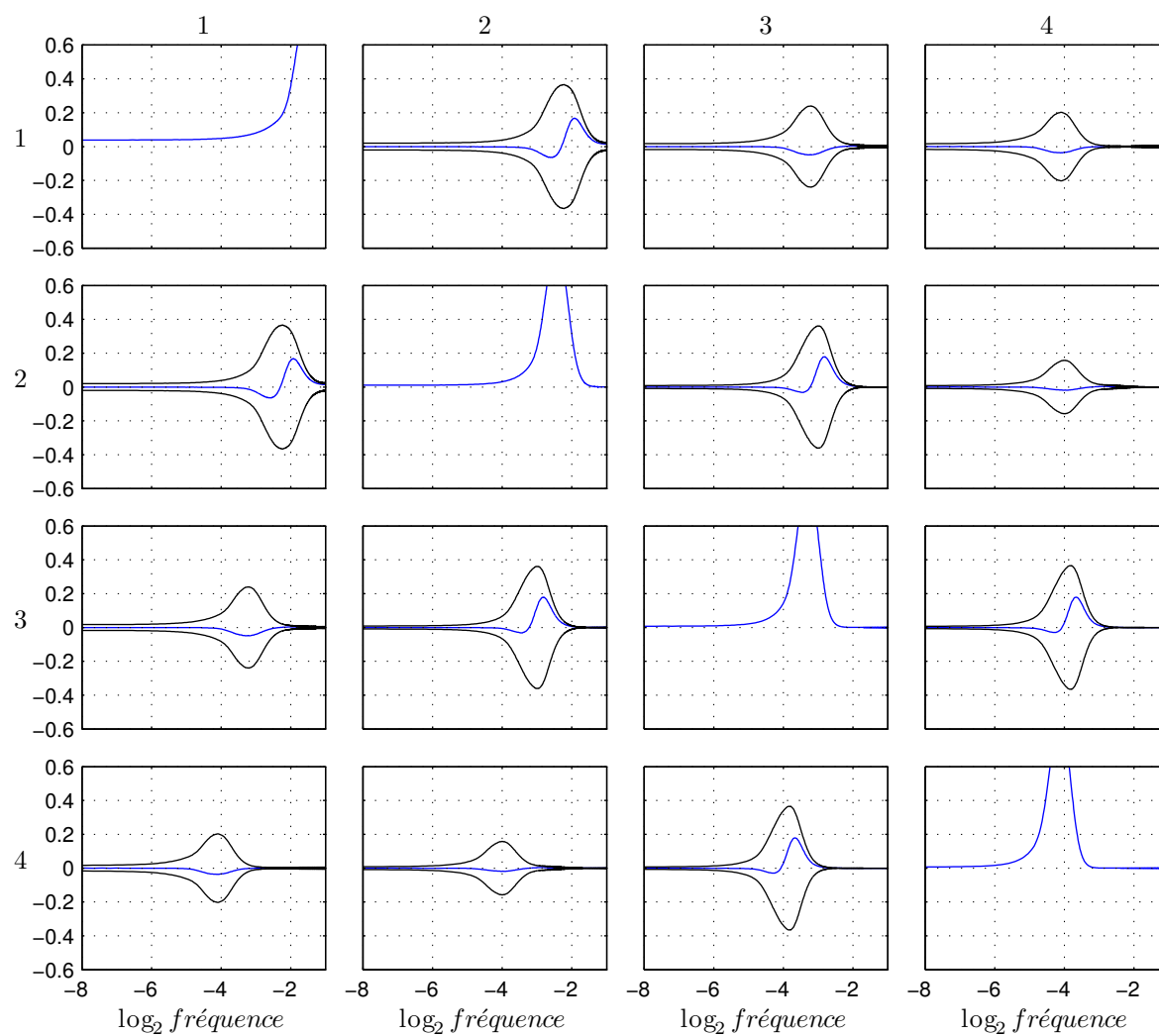


FIGURE 2.114 – Interspectres entre les IMFs issus du bruit blanc simple calculés par l'algorithme Algo. 4 avec 32 directions.

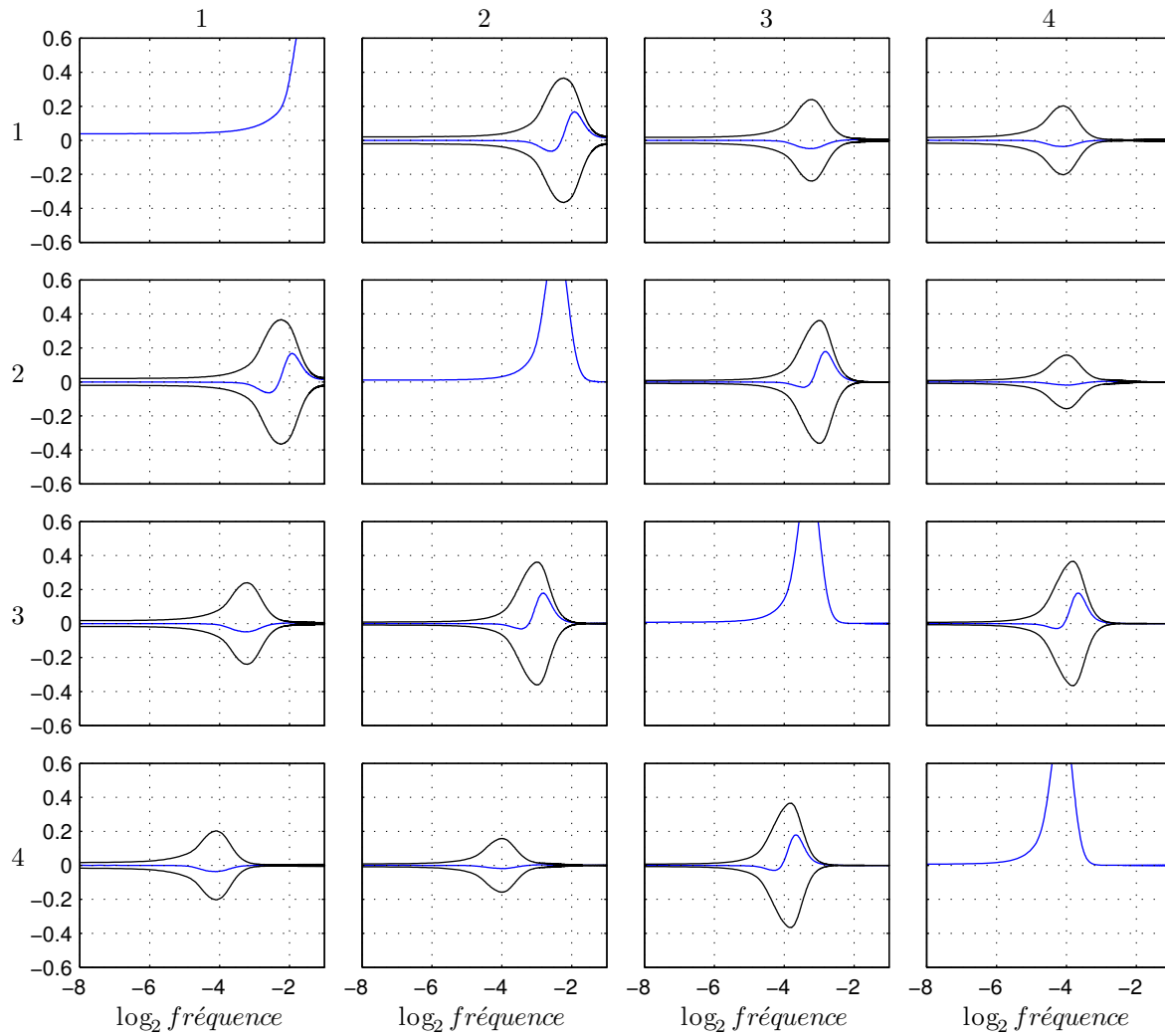


FIGURE 2.115 – Interspectres entre les IMFs issus du bruit blanc simple calculés par l'algorithme Algo. 5 avec 32 directions.

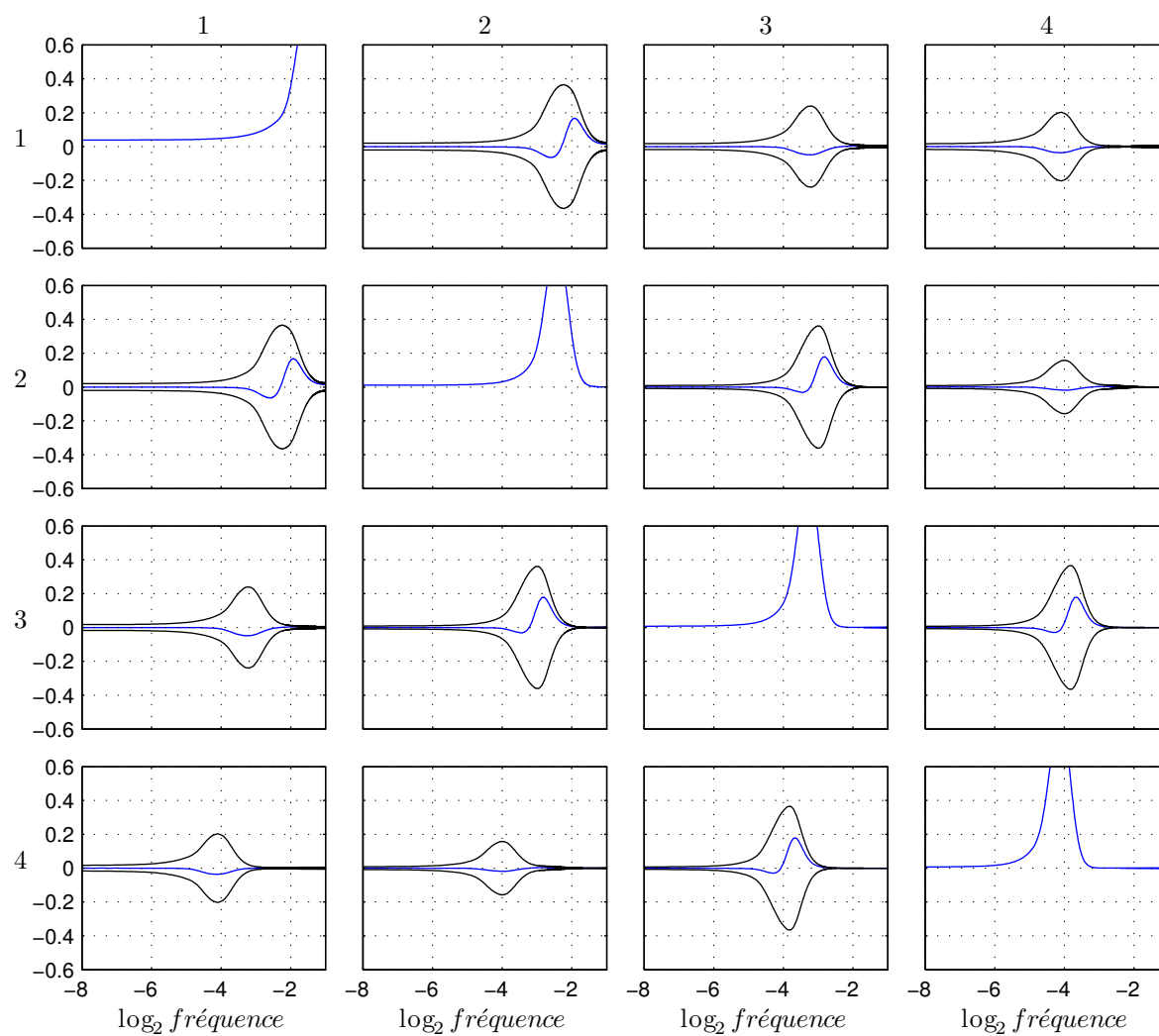


FIGURE 2.116 – Interspectres entre les IMFs issus du bruit blanc analytique calculés par l'algorithme Algo. 4 avec 32 directions.

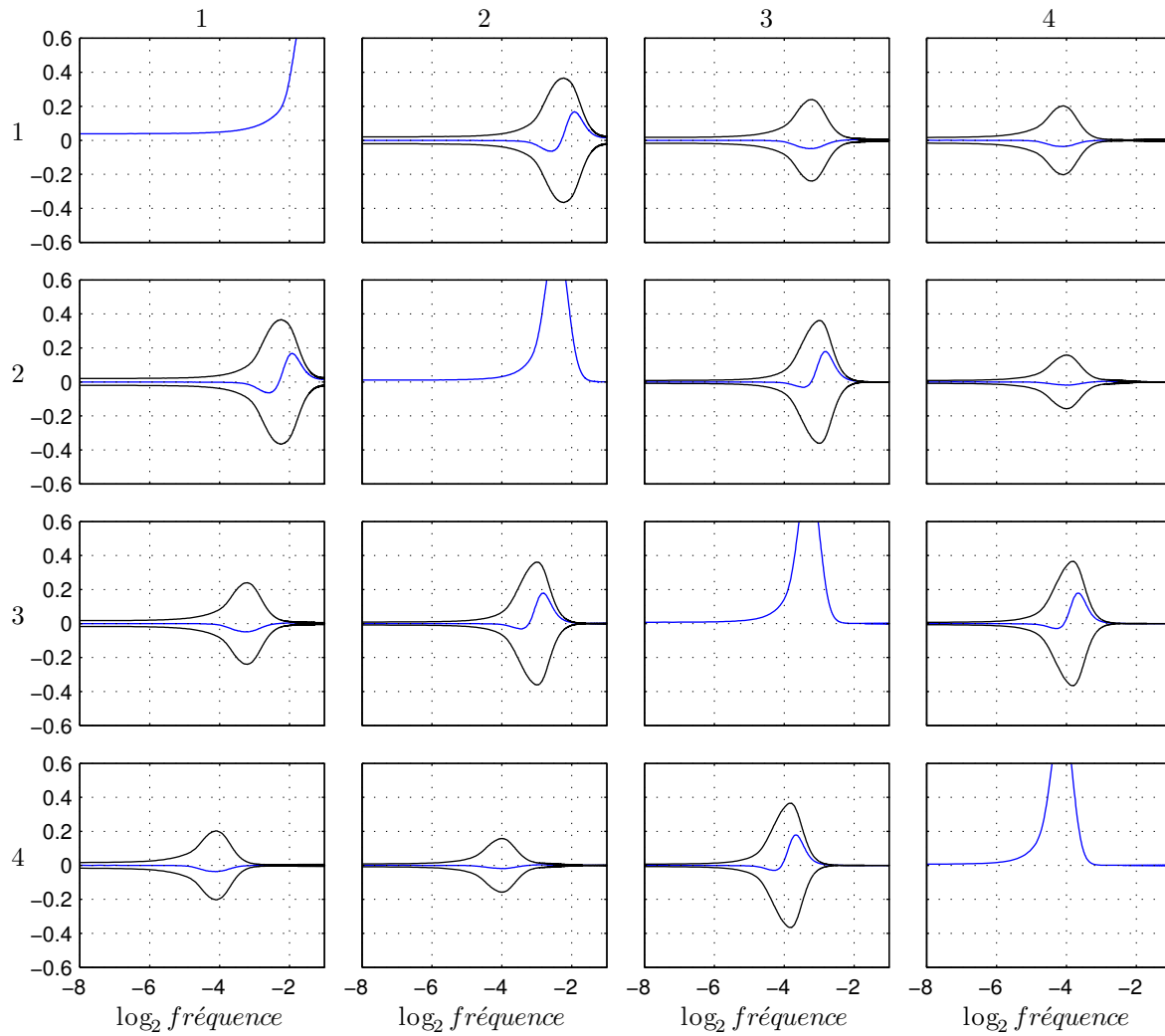


FIGURE 2.117 – Interspectres entre les IMFs issus du bruit blanc analytique calculés par l'algorithme Algo. 5 avec 32 directions.

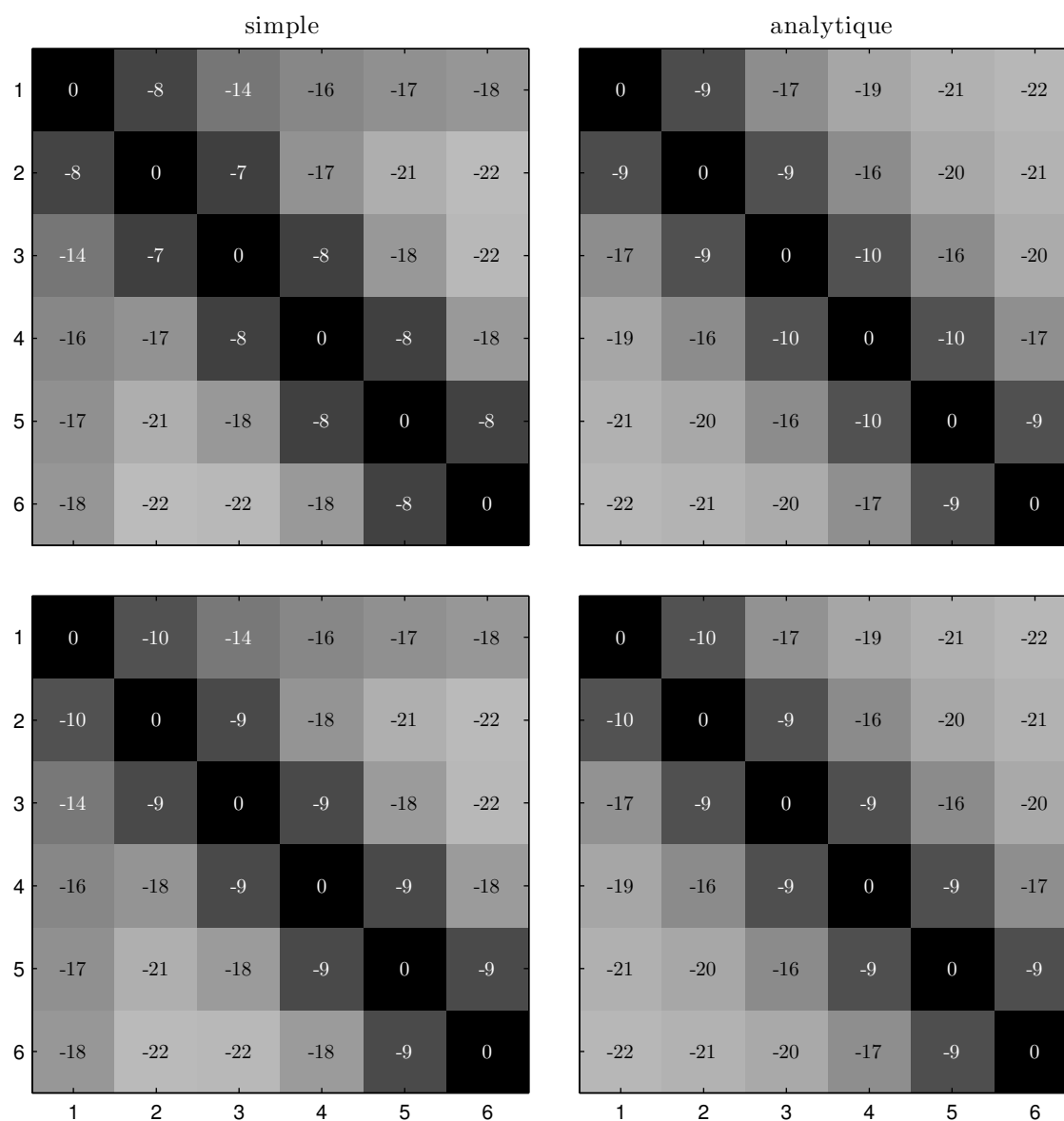


FIGURE 2.118 – Normes des interspectres entre les IMFs. La colonne de gauche correspond au bruit blanc simple, celle de droite au bruit blanc analytique. La rangée du haut correspond à l'algorithme Algo. 4, celle du bas à l'algorithme Algo. 5.

modèle dès lors que la densité d'extrema du signal était la même que dans la situation asymptotique où les extrema sont effectivement uniformément espacés. Dans le cas des bruits, les extrema ne sont jamais uniformément espacés mais on peut toujours définir une densité d'extrema moyenne, ce qui laisse penser que la situation n'est peut-être pas si différente de celle d'une somme de sinusoides dans le cas où la densité d'extrema est identique à celle de l'une des deux composantes mais où les extrema ne sont pas uniformément espacés. En réalité, les situations sont suffisamment différentes pour qu'on ne puisse appliquer simplement le modèle déterministe mais on peut cependant proposer un modèle simpliste qui reprend les ingrédients du modèle déterministe, échantillonnage et filtrage linéaire, et permet de retrouver qualitativement les comportements de type banc de filtres et d'apporter un nouvel éclairage sur certaines de ses caractéristiques.

2.4.1 Adaptation du modèle déterministe aux cas de bruits

Le modèle établi précédemment pour le cas où les extrema sont uniformément espacés est constitué de deux étapes : (sous-)échantillonnage et filtrage linéaire. Pour modéliser le comportement de l'EMD sur des bruits large bande, on se propose de reprendre ces deux ingrédients et de les adapter en s'appuyant sur les résultats de simulation. De manière générale, on peut exprimer le modèle développé précédemment de la manière suivante

$$\mathcal{S}x = x - \mathcal{I}(x + \mathcal{A}x), \quad (2.224)$$

où $\mathcal{I}$ représente un opérateur de filtrage linéaire passe-bas et $\mathcal{A}$ représente les distortions introduites par l'échantillonnage (repliement). La moyenne des enveloppes du signal $x(t)$ s'exprime donc comme

$$m = \mathcal{I}(x + \mathcal{A}x). \quad (2.225)$$

Pour définir plus précisément les opérateurs $\mathcal{I}$ et $\mathcal{A}$, on se propose de formuler tout d'abord un certain nombre d'hypothèses raisonnables et de compléter par la suite les informations manquantes à l'aide de simulations.

2.4.1.1 Opérateur de filtrage linéaire Dans le cadre du modèle développé précédemment pour les sommes de signaux périodiques, on a vu que l'opérateur de filtrage linéaire prenait la forme d'un filtre passe-bas dont la fréquence de coupure était typiquement le quart de la densité d'extrema du signal. Dans le cas de bruits blancs, on sait que la probabilité que le signal présente un maximum (resp. minimum) local en un point donné est exactement $1/3$, ce qui fait qu'on peut lui associer une densité de maxima (resp. minima) de $1/3$ et une densité d'extrema de $2/3$. De là, il semble raisonnable de proposer que l'opérateur de filtrage appliqué au bruit blanc prenne la forme d'un filtre passe-bas de fréquence de coupure $1/6$. De plus, on peut supposer en première approximation que la forme du filtre est la même que dans le cas des sommes de signaux périodiques, c'est-à-dire que sa fonction de transfert est $I(3\nu)$ avec

$$I(\nu) = \left(\frac{\sin \pi \nu}{\pi \nu} \right)^4 \frac{3}{2 + \cos(2\pi \nu)}. \quad (2.226)$$

Plus généralement, pour un signal stationnaire x quelconque, on définira l'opérateur $\mathcal{I}$ comme le filtre linéaire de fonction de transfert $I(2\nu/d_e[x])$ avec $d_e[x]$ la densité d'extrema de x . Dans le cas où x est un processus gaussien, celle-ci peut être obtenue par une formule dérivée de la formule de Rice [48, 34] qui relie la densité de passages à zéro d'un processus aléatoire stationnaire gaussien au comportement de son autocorrélation en zéro. En appliquant une version discrète de la formule de Rice à la dérivée de x , on peut aboutir à la formule suivante [34] qui relie sa densité d'extrema à sa densité spectrale de puissance

$$d_e[x] = \frac{\int_0^1 \cos(2\pi \nu) S_x(\nu) \sin^2(\pi \nu) d\nu}{\int_0^1 S_x(\nu) \sin^2(\pi \nu) d\nu}. \quad (2.227)$$

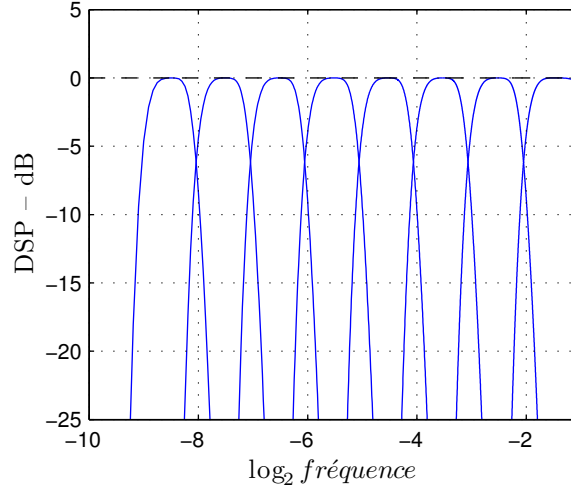


FIGURE 2.119 – Spectres des IMFs calculés par le modèle simplifié avec $\mathcal{A} = 0$ dans le cas d'un bruit blanc avec 10 itérations de tamisage par IMF.

Avant de chercher à modéliser l'opérateur $\mathcal{A}$ représentant les effets dus à l'échantillonnage, on peut d'ores et déjà s'intéresser au comportement du modèle avec $\mathcal{A} = 0$. Le modèle d'une itération de tamisage dans ce cas est simplement décrit par la procédure en deux étapes

1. calculer la densité d'extrema du signal (supposé gaussien) à l'aide de la formule (2.227) : $d_e[x]$
2. appliquer le filtre linéaire : $(\widehat{\mathcal{S}x})(\nu) = \left(1 - I\left(\frac{2\nu}{d_e[x]}\right)\right) \hat{x}(\nu)$

Le résultat de l'application de ce modèle à un bruit blanc est donné Fig. 2.119. Pour chaque IMF on a appliqué exactement 10 itérations de la procédure précédente. On observe que ce modèle très simple permet de retrouver le comportement de type banc de filtres et même de retrouver un facteur d'auto-similarité entre les spectres proche de 2 pour 10 itérations de tamisage. En revanche, il ne permet pas de modéliser correctement la forme des spectres des IMFs, notamment le fait qu'ils aient une partie basses fréquences relativement importante entre -10 et -15 dB.

2.4.1.2 Modélisation des distortions introduites par l'échantillonnage La modélisation précédente peut être améliorée en tenant compte des effets d'échantillonnage à travers l'opérateur $\mathcal{A}$. Celui-ci est plus délicat à modéliser puisque l'échantillonnage par les maxima/minima d'un bruit blanc est a priori assez différent d'un échantillonnage régulier ou quasi-régulier tel qu'il était dans le cadre des sommes de composantes périodiques. On va donc plutôt s'inspirer des résultats de simulations pour déterminer quelle forme on peut lui donner pour le modèle. Si $\mathcal{I}$ se comporte bien comme un opérateur de filtrage linéaire, le modèle suggère que la densité spectrale de puissance de la moyenne des enveloppes S_m d'un bruit blanc ainsi que l'interspectre $S_{x,m}$ entre cette dernière et le signal prennent la forme

$$S_m(\nu) = I(3\nu)^2 (S_x(\nu) + S_{\mathcal{A}x}(\nu) + 2\text{Re}(S_{x,\mathcal{A}x}(\nu))), \quad (2.228)$$

$$S_{x,m}(\nu) = I(3\nu) (S_x(\nu) + S_{x,\mathcal{A}x}(\nu)). \quad (2.229)$$

En particulier, sachant que $S_x(\nu) = 1$ et $I(\nu)$ est constant égal à 1 lorsque ν tend vers zéro, on doit avoir au voisinage de zéro $S_m(\nu) = 1 + S_{\mathcal{A}x}(\nu) + 2\text{Re}(S_{x,\mathcal{A}x}(\nu))$ et $S_{x,m}(\nu) = 1 + S_{x,\mathcal{A}x}(\nu)$. La comparaison de ces dernières expressions aux résultats de simulations présentés Fig. 2.120, permet

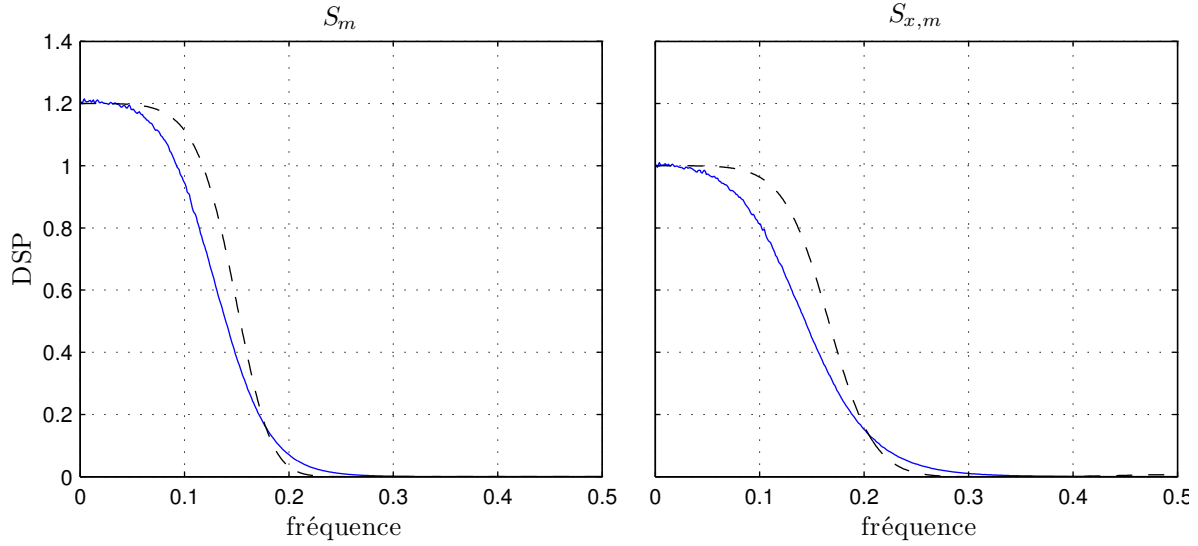


FIGURE 2.120 – Calibration de l'opérateur $\mathcal{A}$. À gauche, le spectre de la moyenne des enveloppes S_m en trait plein et sa modélisation en tirets. À droite, l'interspectre entre le signal et la moyenne de ses enveloppes $S_{x,m}$ en trait plein et sa modélisation en tirets.

d'identifier les différents termes, du moins pour ν tendant vers zéro :

$$S_{x,\mathcal{A}x}(\nu \rightarrow 0) \simeq 0, \quad (2.230)$$

$$S_{\mathcal{A}x}(\nu \rightarrow 0) \simeq 0.2. \quad (2.231)$$

Tout se passe donc comme si $\mathcal{A}x$ avait une densité spectrale constante à basse fréquence et était décorrélé de x également à basse fréquence. En fait, compte tenu de la forme de la fonction de transfert du filtre $\mathcal{I}$, on peut même supposer que ces comportements sont valables sur toute la gamme de fréquences sans pour autant trop s'éloigner des résultats des simulations (cf Fig. 2.120).

En s'inspirant de ce résultat observé dans le cas où x est un bruit blanc gaussien, on propose de modéliser de manière plus générale $\mathcal{A}x$ par un bruit blanc indépendant de x . On supposera de plus $\mathcal{A}x$ gaussien de manière à rester dans un cadre gaussien pour pouvoir appliquer la formule (2.227). Il reste à savoir quelle puissance lui attribuer. Classiquement, lorsqu'on échantillonne un signal stationnaire $x(t)$ à une fréquence f , la puissance des distortions dues à l'échantillonnage observées dans la bande $[-f/2, f/2]$ est la puissance du signal en dehors de cette bande. Dans le cas de l'EMD d'un signal aléatoire, l'équivalent de la fréquence d'échantillonnage f est la densité d'extrema d_e qui est la fréquence moyenne à laquelle le signal est échantillonné si l'on considère les deux enveloppes conjointement. Dans la mesure où $\mathcal{A}$ représente des distortions dues au sous-échantillonnage, on peut donc supposer que la puissance de $\mathcal{A}x$ dans la bande $[-d_e/2, d_e/2]$ est égale à la puissance de x en dehors de cette même bande. En réalité, comme on l'a vu pour les sommes de signaux périodiques simples, le fait que l'échantillonnage ne soit pas régulier fait que des composantes de fréquences inférieures à $d_e/2$ peuvent également contribuer aux distortions dues à l'échantillonnage. Cependant, comme leur influence est a priori moindre que celles des composantes de fréquences supérieures à $d_e/2$, on se contentera de l'approximation précédente.

Selon ce modèle, on définirait donc $\mathcal{A}x$ comme un bruit blanc gaussien de puissance

$$\frac{2}{d_e[x]} \int_{d_e[x]/2}^{0.5} S_x(\nu) d\nu. \quad (2.232)$$

Le problème est qu'une telle définition est en fait relativement loin des résultats observés puisque dans le cas d'un bruit blanc, elle suggère une densité spectrale de puissance de 0.5 pour $\mathcal{A}x$ alors qu'on

observe en réalité plutôt 0.2. Dans la mesure où l'échantillonnage par les extrema est très différent d'un échantillonnage régulier, il n'est pas vraiment surprenant d'observer un tel écart, d'autant plus que l'estimation de sa puissance par (2.232) est très grossière. On peut par exemple observer que les maxima (ou les minima) d'un signal ont une variance différente de celle du signal lui-même, typiquement de l'ordre de la moitié pour un bruit blanc gaussien, ce qui va dans le sens de la puissance des effets d'aliasing réellement observée. Pour tenir compte de cet écart difficile à modéliser plus précisément, on propose de simplement ajouter un facteur de proportionnalité de 0.4 à l'expression de la puissance de $\mathcal{A}x$ (2.232). Il s'agit certes d'un ajustement grossier mais pas forcément si gênant dans la mesure où l'intérêt du modèle est plus de justifier qualitativement la forme des spectres des IMFs à partir d'un modèle simple que d'expliquer quantitativement toutes les caractéristiques observées.

2.4.1.3 Algorithme du modèle Partant d'un signal stationnaire gaussien x de densité spectrale de puissance S_x , le modèle ainsi défini permet de calculer itérativement les densités spectrales de puissance du premier IMF et de la première approximations obtenus à l'aide de n itérations de tamisage selon la procédure Algo. 7. En appliquant à nouveau cette procédure sur la première

Algorithme 7 : Modèle du comportement de l'EMD sur les bruits large bande

1 $S_{d_1^{(0)}} \leftarrow S_x$

2 **pour** $1 \leq k \leq n$ **faire**

3 Calculer la densité d'extrema :

$$d_e^{(k)} = \frac{\int_0^1 \cos(2\pi\nu) S_{d_1^{(k)}}(\nu) \sin^2(\pi\nu) d\nu}{\int_0^1 S_{d_1^{(k)}}(\nu) \sin^2(\pi\nu) d\nu}$$

4 Déterminer la puissance du bruit blanc gaussien $\mathcal{A}d_e^{(k)}$:

$$P^{(k)} = \frac{2}{d_e^{(k)}} \int_{d_e^{(k)}/2}^{0.5} S_{d_1^{(k)}}(\nu) d\nu.$$

5 Calculer la densité spectrale de puissance de $d_1^{(k+1)}$:

$$S_{d_1^{(k+1)}}(\nu) = \left| 1 - I\left(\frac{2\nu}{d_e^{(k)}}\right) \right|^2 S_{d_1^{(k)}} + 0.4P^{(k)} \left| I\left(\frac{2\nu}{d_e^{(k)}}\right) \right|^2$$

et l'interspectre avec x :

$$S_{d_1^{(k+1)},x}(\nu) = \left(1 - I\left(\frac{2\nu}{d_e^{(k)}}\right) \right) S_{d_1^{(k)},x}$$

6 Calculer la densité spectrale de puissance de l'approximation $a_1^{(n)} = x - d_1^{(n)}$:

$$S_{a_1^{(n)}}(\nu) = S_{d_1^{(n)}}(\nu) + S_x(\nu) - 2S_{d_1^{(n)},x}(\nu)$$

approximation $a_1^{(n)}$ on peut obtenir les densités spectrales de puissance du deuxième IMF et de la deuxième approximation, puis récursivement celles de tous les IMFs et approximations.

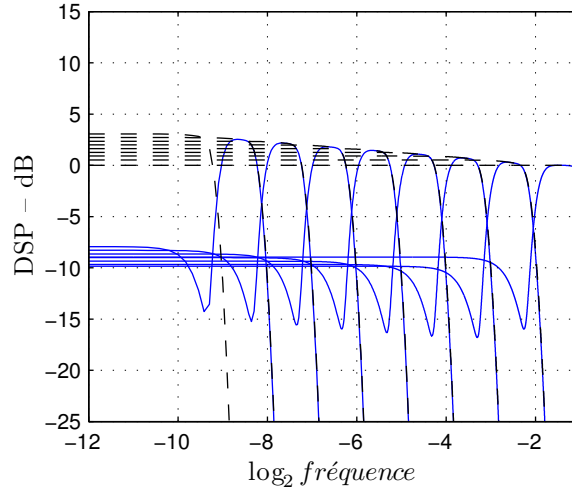


FIGURE 2.121 – Application du modèle au cas d’un bruit blanc avec 10 itérations de tamisage par IMF. Modèles de spectres des IMFs (trait plein) et modèles de spectres des approximations (tirets).

2.4.2 Résultats

2.4.2.1 Cas d’un bruit blanc Les spectres des IMFs et approximations d’un bruit blanc obtenus par cette méthode avec 10 itérations de tamisage par IMF, sont représentés Fig. 2.121. On observe que l’ajout de l’opérateur $\mathcal{A}$ au modèle permet de retrouver le fait que les IMFs ont une composante basses fréquences relativement importante. Ce résultat n’est en fait pas vraiment probant dans la mesure où $\mathcal{A}x$ a en réalité été créé à partir de cette observation dans le cas particulier d’une unique itération de tamisage sur un bruit blanc. En revanche, l’ajout de $\mathcal{A}$ permet aussi de retrouver le fait que les densités spectrales de puissance des approximations en zéro augmentent avec leur indice, ce qui tend à valider l’idée selon laquelle les composantes basses fréquences des IMFs seraient liées à des effets de sous-échantillonnage.

Plus en détail, on observe que les composantes basses fréquences des IMFs autres que le premier prévues par le modèle sont plus importantes que celles observées en réalité. En particulier l’importance de celles-ci augmente avec l’indice de l’IMF alors qu’en réalité, les densités spectrales de puissance des IMFs en 0 sont à peu près constantes à l’exception du premier IMF. Du point de vue de l’interprétation, le fait que le modèle surévalue ainsi les composantes basses fréquences pourrait indiquer que les effets de sous-échantillonnage sont plus importants pour le premier IMF que pour les autres, sans doute parce que l’échelle de variation du premier IMF est proche du pas de discrétisation.

2.4.2.2 Cas de bruits gaussiens fractionnaires Au-delà du bruit blanc, on s’est également intéressé au comportement du modèle sur des fGns (cf Fig. 2.122). On observe que le modèle permet dans tous les cas de retrouver qualitativement les comportements observés en réalité : l’EMD se comporte grossièrement comme un banc de filtres pour $H > 0.5$; pour $H < 0.5$, les allures des spectres sont similaires mais leur amplitude diminue avec leur indice moins vite que l’amplitude de la densité spectrale du signal aux mêmes fréquences. Quantitativement, on retrouve le défaut observé dans le cas du bruit blanc qui est que la composante basses fréquences des IMFs d’indice supérieur à un est surévaluée. Dans le cas où $H < 0.5$, il en découle que l’amplitude des spectres des IMFs selon le modèle diminue encore moins vite qu’en réalité, voire augmente à partir d’un certain indice.

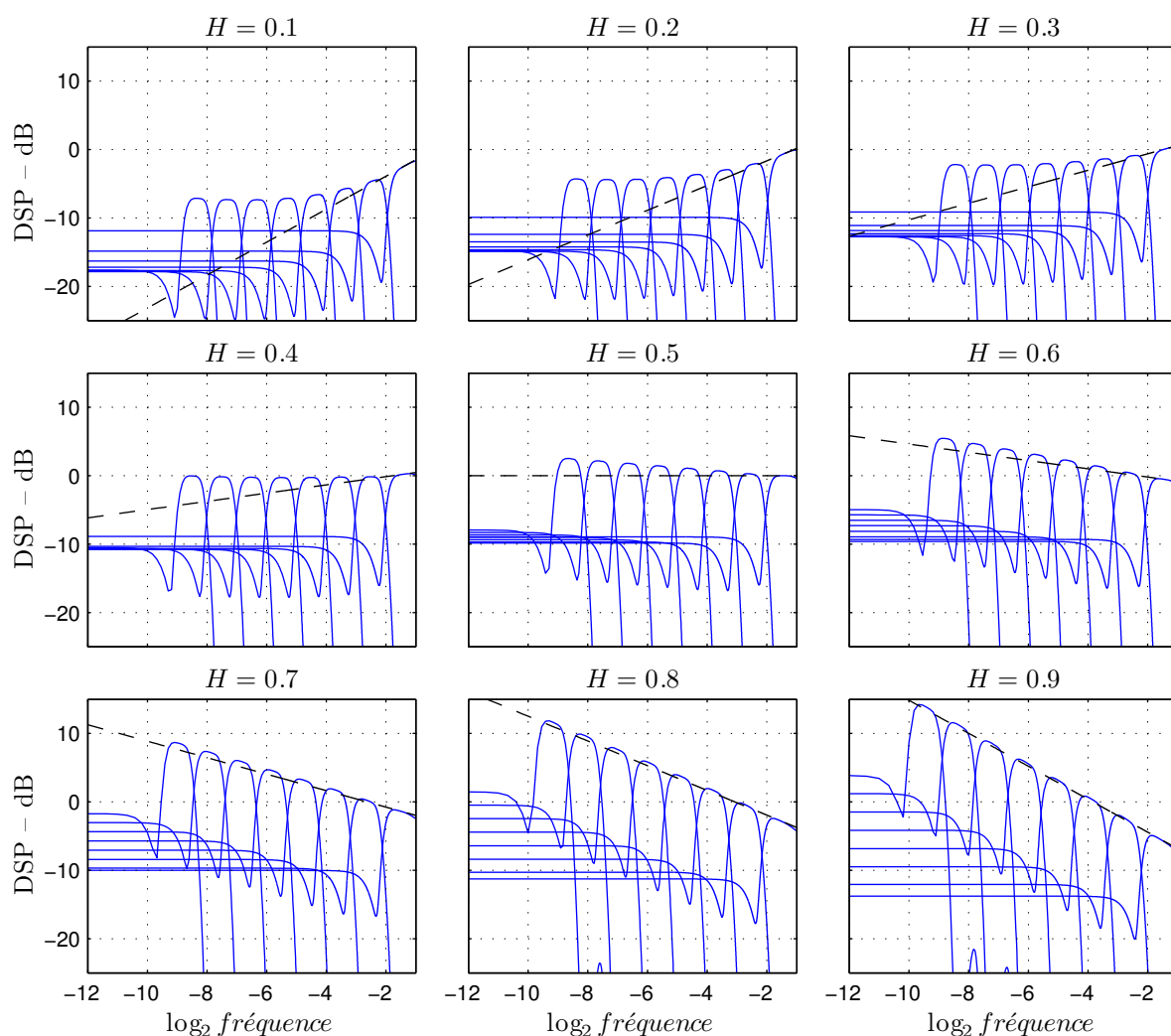


FIGURE 2.122 – Application du modèle au cas de fGns avec 10 itérations de tamisage par IMF. Seuls les modèles des spectres des IMFs sont représentés pour améliorer la lisibilité.

Synthèse et ouverture

Arrivant au terme de ce manuscrit, on peut dresser le bilan des contributions présentées et ouvrir quelques pistes pour de futurs travaux. Par rapport à l'objectif initial qui était d'améliorer la compréhension de l'EMD, on a dans un premier temps évalué ses performances dans les cas simples où le signal en entrée de l'EMD est soit une somme de deux composantes périodiques simples soit un bruit large bande. Dans un deuxième temps on a proposé un modèle de son comportement dans ces cas simples. Au passage, on s'est intéressé aux choix algorithmiques nécessaires à la réalisation d'une implantation et on s'est également penché sur la question de l'influence de l'échantillonnage dont les effets étaient a priori difficiles à prévoir du fait de la non-linéarité de l'algorithme. Enfin, une autre contribution importante de cette thèse est l'introduction d'algorithmes permettant d'appliquer l'EMD à des signaux bivariés, ou complexes.

La mise en place d'un algorithme d'EMD a été l'occasion d'aborder trois petits problèmes pour lesquels soit la contribution originale propose des solutions perfectibles, soit elle n'en propose pas du tout. Ces problèmes sont ceux des conditions aux bords pour définir les enveloppes aux extrémités du signal, celui de l'arrêt du processus de tamisage, et celui de l'interpolation. À l'exception du cas de l'interpolation, l'étude de ces problèmes a permis de proposer de nouvelles solutions qui remplacent avantageusement celles proposées initialement, même si elles sont a priori loin d'être optimales. En plus de ces solutions, ces études ont permis de mieux cerner les difficultés liées à chaque problème.

Au-delà de l'algorithme en lui-même, on s'est intéressé aux conditions d'une bonne utilisation de ce dernier à travers la problématique de l'influence de l'échantillonnage. Le bilan de notre étude sur ce sujet est mitigé. On arrive en effet à contrôler l'ordre de grandeur des variations dues à l'échantillonnage sur une itération de tamisage mais les résultats se dégradent rapidement lorsqu'on réalise plusieurs itérations et qu'on extrait plusieurs IMFs. Dans le cas d'une unique itération, le résultat le plus important de notre étude est sans doute le fait que l'écart au continu se comporte généralement comme l'inverse de la fréquence d'échantillonnage quand celle-ci tend vers l'infini et plutôt comme le carré de son inverse quand elle tend vers zéro, du moins tant que tous les extrema du signal à temps continu ont leur contrepartie dans le signal échantillonné. L'intérêt de ce résultat est cependant limité dans la mesure où dès qu'on se place dans les conditions d'utilisation réelles de l'EMD, le nombre d'itérations de tamisage devient variable, ce qui crée des variations plus importantes que celles dues à l'échantillonnage sur une seule itération, sauf éventuellement pour les très basses fréquences d'échantillonnage. De plus, le fait que les IMFs soient extraits successivement a pour conséquence que les variations observées sur un IMF se transmettent au suivant ce qui fait qu'à mesure qu'on extrait des IMFs, on accumule de plus en plus de variations. Au final, de petites variations dues à l'échantillonnage sur le signal peuvent mener à des variations nettement plus importantes de la décomposition. La problématique de l'échantillonnage permet ainsi d'observer des propriétés relatives à la stabilité ou à la robustesse de la décomposition. Il apparaît en effet que le fait que le nombre d'itérations soit variable constitue une cause d'instabilité dans la mesure où l'écart entre deux itérations successives est non nul lors de l'arrêt du processus de tamisage. Cette instabilité s'explique par le fait que le choix de poursuivre ou non le processus de tamisage est binaire et qu'il ne peut donc toujours évoluer de manière infinitésimale sous l'action d'une perturbation infinitésimale du signal, à moins de ne pas dépendre du tout de ce dernier. L'EMD est donc généralement instable à

moins d'utiliser des nombres d'itérations fixés a priori. Et encore, même dans ce cas, le phénomène d'apparition d'extrema peut être une cause d'instabilité dans la mesure où de petites perturbations peuvent influencer l'apparition ou non d'une paire d'extrema à une itération donnée du processus de tamisage.

En dehors de ces observations relatives à la stabilité de l'EMD on s'est intéressé à ses performances à travers l'étude de son comportement sur des sommes de signaux périodiques simples et des bruits large bande. On a observé en particulier dans ces deux situations que le tamisage avait un comportement apparenté à un filtrage linéaire et que donc l'EMD pouvait être globalement raisonnablement bien décrite par une succession d'opérations de filtrage linéaire. De ce fait, la capacité de l'EMD à décomposer des sommes de signaux non linéaires périodiques n'est a priori pas meilleure que celle d'un filtrage linéaire. En fait, elle semble même moins bonne si on ajoute à cela les effets de repliement spectral apparaissant du fait que les enveloppes sont construites uniquement à partir des extrema du signal. Toutefois, on peut remarquer que les deux situations étudiées sont stationnaires et que par conséquent les capacités de l'EMD à s'adapter aux variations temporelles de la structure du signal ne sont pas utilisées. Dans des situations non stationnaires, ces capacités d'adaptation sont clairement un avantage par rapport à un filtrage linéaire. On ne les a finalement pratiquement pas abordées au cours de cette thèse, mais leur étude est sans doute aussi importante pour la compréhension de l'EMD que celles que nous avons menées dans les situations stationnaires. Les capacités d'adaptation de l'algorithme ne sont cependant pas indépendantes de ses performances dans les situations stationnaires. En particulier, on a observé que la résolution fréquentielle de l'EMD est relativement faible, deux fréquences devant typiquement être dans un rapport supérieur à deux pour être séparées. En contrepartie, la résolution temporelle est généralement très bonne. Si on prend l'exemple d'une modulation d'amplitude sinusoïdale, la résolution fréquentielle relative de 0.5 suggère que l'EMD serait capable de suivre des variations d'amplitude jusqu'à une fréquence de l'ordre du tiers de celle de la porteuse, ce qui correspond assez bien aux observations même si le raisonnement est très grossier. De plus, l'augmentation de la résolution fréquentielle avec le nombre d'itérations est en accord avec la diminution concomitante de la résolution temporelle qui se traduit par le fait que les variations des enveloppes se font plus douces. En revanche, on n'a par cette voie aucune information sur la capacité de l'EMD à suivre les évolutions de la fréquence du signal. Le cas d'une modulation de fréquence pure suggère que celle-ci serait excellente et indépendante du nombre d'itérations contrairement aux variations d'amplitude. Toutefois il est a priori difficile d'évaluer dans quelle mesure ces capacités évoluent dans des situations plus compliquées.

Au-delà de ces évaluations de performances, la principale avancée de cette thèse en matière de compréhension de l'EMD est la mise en évidence d'un modèle décrivant son comportement dans le cas où le signal présente des extrema espacés de manière relativement régulière. Ce modèle permet de rapprocher l'EMD d'outils plus classiques du traitement du signal dans la mesure où il est constitué de deux ingrédients simples, échantillonnage et filtrage linéaire, et il explique efficacement le comportement de l'EMD dans une grande partie des cas où le signal est la somme de deux composantes périodiques faiblement non linéaires. Plus précisément, il donne de bons résultats dès que les extrema du signal sont proches de ceux de l'une des deux composantes, ce qui est généralement le cas quand le rapport de leurs amplitudes est soit grand soit petit en valeur absolue. De plus, les premiers résultats concernant une adaptation au cas de composantes modulées en amplitude et en fréquence sont très encourageants et on a également montré que le modèle permettait de retrouver les principales caractéristiques du comportement de l'EMD sur des bruits large bande, moyennant quelques ajustements. Ces adaptations du modèle n'ont été qu'ébauchées ici mais elles méritent sans doute une étude plus approfondie. En dépit de ces résultats prometteurs, le modèle souffre d'une limitation importante dans la mesure où il ne peut expliquer la décomposition d'un signal constitué de deux composantes périodiques lorsque le rapport des amplitudes n'est ni suffisamment petit ni suffisamment grand pour que les extrema du signal puissent être considérés comme proches de ceux

de l'une des deux composantes. Il y a cependant certains cas où il pourrait être adapté, notamment celui où les extrema peuvent localement sur des intervalles pas trop courts être associés à l'une des deux composantes. On a pu voir en particulier que ce cas se produisait pour les sommes de deux sinusoïdes. Si le modèle peut alors être appliqué localement et modéliser correctement le résultat d'une itération de tamisage, l'extension à plusieurs itérations pose un problème dans la mesure où des extrema peuvent apparaître au cours du processus de tamisage ce qui a pour conséquence de déplacer ou supprimer les zones où ces derniers peuvent être associés à l'une ou l'autre composante. Pour adapter le modèle à ces situations il faudrait donc être en mesure de prévoir les apparitions d'extrema. Plus généralement, le phénomène d'apparition d'extrema au cours du processus de tamisage n'a pratiquement pas été abordé dans cette thèse mais son étude constitue clairement un objectif important pour améliorer la compréhension de l'EMD. Ce phénomène accroît en effet de manière importante les capacités de décomposition de l'EMD et constitue l'intérêt principal de l'aspect itératif du processus de tamisage. Aussi, une application importante de son étude pourrait être de prévoir les apparitions d'extrema ce qui permettrait d'inclure les futurs extrema dès la première itération de tamisage, avec éventuellement la possibilité de ramener ainsi le processus de tamisage à une seule itération.

Enfin, outre les apports à l'EMD originale, une contribution importante de cette thèse est l'introduction de nouveaux algorithmes étendant l'EMD au cas de signaux bivariés. Malgré le côté très récent de ces algorithmes, on s'est efforcé de transposer autant que possible les analyses effectuées sur la méthode originale au cas bivarié. Les premiers résultats montrent que les propriétés de ces algorithmes semblent prolonger celles de l'EMD mais il est encore difficile d'évaluer leurs capacités dans des situations plus générales en raison du très faible nombre de cas d'étude, qu'il s'agisse de simulations ou de signaux expérimentaux. De plus, les deux algorithmes proposés ont des comportements proches mais différents et, si l'un apparaît moins sensible à l'échantillonnage, il paraît prématuré d'exclure totalement l'autre qui peut encore se révéler intéressant dans des situations très bien échantillonnées.

Bibliographie

- [1] P. ABRY, P. FLANDRIN, M. S. TAQQU et D. VEITCH : Wavelets for the analysis, estimation and synthesis of scaling data. In K. PARK et W. WILLINGER, édés : *Self-Similar Network Traffic and Performance Evaluation*, p. 39–88. Wiley, 2000.
- [2] H. AKIMA : A new method of interpolation and smooth curve fitting based on local procedures. *Journal of the ACM*, 17(4):589–602, 1970.
- [3] F. AUGER et P. FLANDRIN : Improving the readability of time-frequency and time-scale representations by reassignment methods. *IEEE Transactions on Signal Processing*, 43(5):1068–1089, 1995.
- [4] E. BEDROSIAN : A product theorem for hilbert transforms. *Proceedings of the IEEE*, 51:868–869, 1963.
- [5] B. BOASHASH, éd. *Time Frequency signal analysis and processing*. Elsevier, 2003.
- [6] R. CARMONA, W.-L. HWANG et B. TORRÉSANI : *Practical time-frequency analysis*. Academic Press, 1998.
- [7] M. CASDAGLI et S. EUBANK, édés. *Nonlinear modeling and forecasting*. Addison-Wesley, 1992.
- [8] J. C. CEXUS : *Analyse des signaux non stationnaires par Transformation de Huang, Opérateur de Teager–Kaiser et transformation de Huang–Teager*. Thèse de doctorat, Université de Rennes 1, 2005.
- [9] Q. CHEN, N. HUANG, S. RIEMENSCHNEIDER et Y. XU : A B-spline approach for empirical mode decompositions. *Advances in Computational Mathematics*, 24(1):171–195, jan 2006.
- [10] L. COHEN : *Time-frequency analysis*. Prentice Hall Signal Processing Series, 1995.
- [11] C. DAMERVAL, S. MEIGNEN et V. PERRIER : A fast algorithm for bidimensional EMD. *IEEE Signal Processing Letters*, 12(10):701–704, 2005.
- [12] M. DÄTIG et T. SCHLURMANN : Performance and limitations of the hilbert-huang transformation (hht) with an application to irregular water waves. *Ocean Engineering*, 31(14-15):1783–1834, October 2004.
- [13] C. de BOOR : *A practical guide to splines*. Springer Verlag, New York, 1978.
- [14] R. DEERING et J. KAISER : The use of a masking signal to improve empirical mode decomposition. In *IEEE International Conference on Acoustics Speech and Signal Processing*, vol. 4, p. 485–488, March 2005.
- [15] E. DELÉCHELLE, J. LEMOINE et O. NIANG : Empirical mode decomposition : an analytical approach for sifting process. *IEEE Signal Processing Letters*, 12(11):764–767, november 2005.

- [16] Y. DENG, W. WANG, C. QIAN, Z. WANG et D. DAI : Boundary-processing-technique in EMD method and Hilbert transform. *Chinese Science Bulletin*, 46(11):954–961, june 2001.
- [17] P. G. DRAZIN : *Nonlinear systems*. Cambridge University Press, 1992.
- [18] P. FLANDRIN : *Temps-fréquence*. Hermes, 1993.
- [19] P. FLANDRIN et P. GONÇALVES : Empirical Mode Decompositions as a data-driven wavelet-like expansions. *International Journal of Wavelets, Multiresolution and Information Processing*, 2(4):477–496, 2004.
- [20] P. FLANDRIN, P. GONÇALVES et G. RILLING : EMD equivalent filter banks, from interpretations to applications. In [27], p. 57–74.
- [21] P. FLANDRIN, G. RILLING et P. GONÇALVES : Empirical Mode Decomposition as a filter bank. *IEEE Signal Processing Letters*, 11(2):112–114, fév. 2004.
- [22] F. N. FRITSCH et R. E. CARLSON : Monotone piecewise cubic interpolation. *SIAM Journal on Numerical Analysis*, 17:238–246, 1980.
- [23] D. GUÉGAN : *Séries chronologiques non linéaires à temps discret*. Economica, 1994.
- [24] F. HLAWATSCH et F. AUGER, éd. *Temps-fréquence*. Hermes, 2005.
- [25] N. HUANG, Z. SHEN et S. R. LONG : A new view of nonlinear water waves : The hilbert spectrum. *Annual Review of Fluid Mechanics*, 31:417–457, 1999.
- [26] N. E. HUANG : Introduction to the Hilbert–Huang Transform and its related mathematical problems. In [27], p. 1–26.
- [27] N. E. HUANG et S. S. P. SHEN, éd. *Hilbert–Huang Transform and Its Applications*. World Scientific, 2005.
- [28] N. E. HUANG, Z. SHEN, S. R. LONG, M. L. WU, H. H. SHIH, Q. ZHENG, N. C. YEN, C. C. TUNG et H. H. LIU : The Empirical Mode Decomposition and Hilbert spectrum for nonlinear and non-stationary time series analysis. *Proceedings of the Royal Society London Series A*, 454:903–995, 1998.
- [29] N. E. HUANG, M.-L. C. WU, S. R. LONG, S. S. P. SHEN, W. QU, P. GLOERSEN et K. L. FAN : A confidence limit for the empirical mode decomposition and Hilbert spectral analysis. *Proceedings of the Royal Society London Series A*, 459:2317–2345, 2003.
- [30] N. E. HUANG, Z. WU et S. R. LONG : On instantaneous frequency. *Advances in Adaptive Data Analysis*, 1(2):177–229, avr. 2009.
- [31] R. N. M. JR. : HHT sifting and filtering. In [27], p. 75–105.
- [32] C. JUNSHENG, Y. DEJIE et Y. YU : Research on the intrinsic mode function (imf) criterion in emd method. *Mechanical Systems and Signal Processing*, 20(4):817–824, May 2006.
- [33] C. JUNSHENG, Y. DEJIE et Y. YU : Application of support vector regression machines to the processing of end effects of Hilbert–Huang transform. *Mechanical Systems and Signal Processing*, 21(3):1197–1211, April 2007.
- [34] B. KEDEM : Spectral analysis and discrimination by zero-crossings. *Proceedings of the IEEE*, 74(11):1477–1493, November 1986.

- [35] Y. KOPSINIS et S. MCCLAUGHLIN : Enhanced empirical mode decomposition using a novel sifting-based interpolation points detection. *In IEEE Workshop on Statistical Signal Processing*, Madison (USA) 2007.
- [36] Y. KOPSINIS et S. MCCLAUGHLIN : Investigation and performance enhancement of the empirical mode decomposition method based on a heuristic search optimization approach. *IEEE Transactions on Signal Processing*, 2007. (à paraître).
- [37] J. M. LILLY et J.-C. GASCARD : Wavelet ridge diagnosis of time-varying elliptical signals with application to an oceanic eddy. *Nonlinear Processes in Geophysics*, 13:467–483, 2006.
- [38] A. LINDERHED : 2D empirical mode decompositions in the spirit of image compression. *In Wavelet and Independent Component Analysis Applications IX , SPIE Proceedings*, vol. 4738, p. 1–8, April 2002.
- [39] A. LINDERHED : *Adaptive image compression with wavelet packets and empirical mode decomposition*. Thèse de doctorat, Linköping University, Image Coding Group, 2004.
- [40] Z. LIU et S. PENG : Boundary processing of bidimensional EMD using texture synthesis. *IEEE Signal Processing Letters*, 12(1):33–36, jan. 2004.
- [41] Z. LIU, H. WANG et S. PENG : Texture segmentation using directional empirical mode decomposition. *In International Conference on Image Processing*, vol. 1, p. 279–282, 2004.
- [42] S. MALLAT : *A wavelet tour of signal processing*. Academic Press, 1998.
- [43] J. C. NUNES : *Analyse multiéchelle d'image, Application à l'angiographie rétinienne et à la DMLA*. Thèse de doctorat, Paris 12 Val de Marne : Créteil, 2003.
- [44] J. C. NUNES, Y. BOUAOUNE, E. DELECELLE, O. NIANG et P. BUNEL : Image analysis by bidimensional empirical mode decomposition. *Image and Vision Computing*, 21(12), 2003.
- [45] A. PAPANDREOU-SUPPAPPOLA, éd. *Applications in Time-Frequency analysis*. CRC Press, 2003.
- [46] K. QI, Z. HE et Y. ZI : Cosine window-based boundary processing method for emd and its application in rubbing fault diagnosis. *Mechanical Systems and Signal Processing*, 21(7):2750–2760, October 2007.
- [47] S. QIN et Y. ZHONG : A new envelope algorithm of hilbert-huang transform. *Mechanical Systems and Signal Processing*, 20(8):1941–1952, November 2006.
- [48] S. O. RICE : Statistical properties of a sine wave plus random noise. *Bell System Technical Journal*, 27:109–157, jan 1948.
- [49] P. RICHARDSON, D. WALSH, L. ARMI, M. SCHRÖDER et J. F. PRICE : Tracking three Meddies with SOFAR floats. *Journal of Physical Oceanography*, 19:371–383, 1989.
- [50] G. RILLING et P. FLANDRIN : Sur la Décomposition Modale Empirique des signaux échantillonnés. *In Actes du 20^e Colloque GRETSI sur le Traitement du Signal et des Images*, p. 755–758, Louvain-la-Neuve (B) 2005.
- [51] G. RILLING et P. FLANDRIN : On the influence of sampling on the Empirical Mode Decomposition. *In IEEE International Conference on Acoustics Speech and Signal Processing*, Toulouse (F) 2006. DOI : 10.1109/ICASSP.2006.1660686.

- [52] G. RILLING et P. FLANDRIN : Une extension bivariée pour la décomposition modale empirique – application à des bruits blancs complexes. *In Actes du 21^e Colloque GRETSI sur le Traitement du Signal et des Images*, Troyes (F) 2007.
- [53] G. RILLING et P. FLANDRIN : One or two frequencies? the empirical mode decomposition answers. *IEEE Transactions on Signal Processing*, 2007. (à paraître).
- [54] G. RILLING et P. FLANDRIN : Sampling effects on the empirical mode decomposition. *Advances in Adaptive Data Analysis*, 1(1):43–59, jan. 2009.
- [55] G. RILLING, P. FLANDRIN et P. GONÇALVES : On Empirical Mode Decomposition and its algorithms. *In IEEE-EURASIP Workshop on Nonlinear Signal and Image Processing*, Grado (I) 2003.
- [56] G. RILLING, P. FLANDRIN et P. GONÇALVES : Empirical Mode Decomposition, fractional Gaussian noise and Hurst exponent estimation. *In IEEE International Conference on Acoustics Speech and Signal Processing*, vol. 4, p. iv/489– iv/492, mars 2005.
- [57] G. RILLING, P. FLANDRIN, P. GONÇALVES et J. M. LILLY : Bivariate Empirical Mode Decomposition. *IEEE Signal Processing Letters*, 14(12):936–939, déc. 2007.
- [58] B. SCHÖLKOPF et A. J. SMOLA : *Learning with kernels*. MIT Press, 2002.
- [59] R. C. SHARPLEY et V. VATCHEV : Analysis of the Intrinsic Mode Functions. *Constructive Approximation*, 24:17–47, 2006.
- [60] T. TANAKA et D. P. MANDIC : Complex Empirical Mode Decomposition. *IEEE Signal Processing Letters*, 14(2):101–104, 2007.
- [61] H. TONG : *Non-linear time series analysis*. Oxford science publications, 1990.
- [62] M. UNSER : Splines : A perfect fit for signal processing. *IEEE Signal Processing Magazine*, 16(6):22–38, 1999.
- [63] W. WANG, X. LI et R. ZHANG : Boundary processing of hht using support vector regression machines. *In International Conference on Computer Science*, p. 147–177, Beijing (China) 2007.
- [64] Z. WU et N. E. HUANG : Statistical significance test of intrinsic mode functions. *In [27]*, p. 107–127.
- [65] Z. WU et N. E. HUANG : Ensemble empirical mode decomposition : a noise-assisted data analysis method. Rap. tech. Tech. Rep. No.173., Centre for Ocean-Land-Atmosphere Studies, 2004.
- [66] Y. XU, B. LIU, J. LIU et S. RIEMENSCHNEIDER : Two-dimensional empirical mode decomposition by finite elements. *Proceedings of the Royal Society London Series A*, 462(2074):3081–3096, 2006.

Résumé.

Décompositions Modales Empiriques : Contributions à la théorie, l'algorithmie et l'analyse de performances

La Décomposition Modale Empirique (EMD pour « Empirical Mode Decomposition ») est un outil récent de traitement du signal dévolu à l'analyse de signaux non stationnaires et/ou non linéaires. L'EMD produit pour tout signal une décomposition multi-échelles pilotée par les données. Les composantes obtenues sont des formes d'onde oscillantes potentiellement non harmoniques dont les caractéristiques, forme, amplitude et fréquence peuvent varier au cours du temps. L'EMD étant une méthode encore jeune, elle n'est définie que par la sortie d'un algorithme inhabituel, comportant de multiples degrés de liberté et sans fondement théorique solide.

Nous nous intéressons dans un premier temps à l'algorithme de l'EMD. Nous étudions d'une part les questions soulevées par les choix de ses degrés de liberté afin d'en établir une implantation. Nous proposons d'autre part des variantes modifiant légèrement ses propriétés et une extension permettant de traiter des signaux à deux composantes.

Dans un deuxième temps, nous nous penchons sur les performances de l'EMD. L'algorithme étant initialement décrit dans un contexte de temps continu, mais systématiquement appliqué à des signaux échantillonnés, nous étudions la problématique des effets d'échantillonnage sur la décomposition. Ces effets sont modélisés dans le cas simple d'un signal sinusoïdal et une borne de leur influence est obtenue pour des signaux quelconques.

Enfin nous étudions le mécanisme de la décomposition à travers deux situations complémentaires, la décomposition d'une somme de sinusoïdes et celle d'un bruit large bande. Le premier cas permet de mettre en évidence un modèle simple expliquant le comportement de l'EMD dans une très grande majorité des cas de sommes de sinusoïdes. Ce modèle reste valide pour des sinusoïdes faiblement modulées en amplitude et en fréquence ainsi que dans certains cas de sommes d'ondes non harmoniques périodiques. La décomposition de bruits large bande met quant à elle en évidence un comportement moyen de l'EMD proche de celui d'un banc de filtres auto-similaire, analogue à ceux correspondant aux transformées en ondelettes discrètes. Les propriétés du banc de filtres équivalent sont étudiées en détail en fonction des paramètres clés de l'algorithme de l'EMD. Le lien est également établi entre ce comportement en banc de filtres et le modèle développé dans le cas des sommes de sinusoïdes.

Mots-clés : Empirical Mode Decomposition, décomposition adaptative, multi-résolution, temps-fréquence, temps-échelle, banc de filtres

Abstract.

Empirical Mode Decompositions : Contributions to theory, algorithms and performance analysis

The Empirical Mode Decomposition (EMD) is a novel signal processing tool dedicated to the analysis of nonstationary and/or nonlinear signals. The EMD provides for any signal a data-driven multi-scale decomposition. The components are oscillatory signals, not necessarily harmonic, whose characteristics, waveform, amplitude and frequency may be time-varying. The EMD being rather recent, it is only defined as the output of an unusual algorithm, with many degrees of freedom and no sound theoretical basis.

We first focus on the algorithm of the EMD. The questions raised by the possibilities for the degrees of freedom are studied in order to propose an implementation. We also propose some slight variations on the original algorithm and an extension to process bivariate signals.

Motivated by the fact that the algorithm is initially presented in a continuous time framework, but systematically applied to sampled signals, we study the effect of sampling on the decomposition. A model of sampling effects is proposed for the simple case of a sinusoidal input signal and a bound on the magnitude of these effects is derived for arbitrary input signals.

Finally the mechanism underlying the decomposition is studied by means of the analysis of two complementary situations, the decomposition of sums of two sinusoids and that of broadband noise signals. The first case allows the derivation of a simple model explaining the behavior of the EMD for a vast majority of the possibilities of sums of sinusoids. This model remains valid for slightly amplitude and frequency modulated sinusoids and also for some cases of sums of non harmonic periodic waveforms. As for broadband noise signals, we observe that the behavior of the EMD is close to that of an autosimilar filter bank, analogous to those corresponding to discrete wavelet transforms. The properties of the equivalent filter bank are studied in detail as a function of the key parameters of the EMD algorithm. A link is also established between this filter bank behavior and the model developed for the sums of sinusoids.

Key words : Empirical Mode Decomposition, adaptive decomposition, multi-resolution, time-frequency, time-scale, filter bank