

Problèmes inverses

Ali Mohammad-Djafari

Cours du DEA Automatique et Traitement du Signal

Laboratoire des Signaux et Systèmes (CNRS-ESE-UPS)

École Supérieure d'Électricité

Plateau de Moulon, 91192 Gif-sur-Yvette Cedex, France.

Version provisoire du: 1^{er} mai 2001¹

1. Rapport interne LSS: fichier `poly.tex`

Table des matières

1	Introduction	9
2	Exemples de problèmes inverses	13
2.1	Instrumentation	14
2.2	Tomographie à rayons X	15
2.3	Imagerie à ondes diffractées : cadre général	20
2.4	Imagerie micro ondes : cas 2D	23
2.5	Imagerie micro ondes : cas 3D	27
2.6	Imagerie par courants de FOUCAULT	28
2.7	Imagerie à émission de positons	35
2.8	Imagerie à résonance magnétique nucléaire RMN	37
2.9	Imagerie radar à ouverture synthétique : SAR	40
2.10	Imagerie en géophysique	42
2.11	Imagerie en radio astronomie	43
3	Cadre général et unification	45
3.1	Présentation du cadre général	46
3.2	Exemples classiques de problèmes inverses linéaires	49
3.2.1	Déconvolution de signaux 1-D	49
3.2.2	Débruitage ou filtrage	50
3.2.3	Restauration d'image	51
3.2.4	Reconstruction d'image en tomographie X	52
3.2.5	Radiographie des objets cylindriques	54
3.2.6	Synthèse d'ouverture en radio astronomie ou imagerie radar	56
3.2.7	Synthèse de FOURIER en reconstruction tomographique d'image	57
3.2.8	Synthèse de FOURIER-LAPLACE	59
3.3	Exemples de problèmes inverses non linéaires	60
3.3.1	Imagerie par ondes diffractées	60
3.3.2	Tomographie d'impédance	62
4	Problèmes mal-posés	63
4.1	Rappels d'analyse fonctionnelle	64
4.2	Définitions	68
4.3	Exemples de problèmes mal-posés	69
4.4	Déconvolution : un problème mal-posé	72
4.4.1	Passage dans le domaine de FOURIER	73
4.4.2	Discrétisation	74
4.4.3	Décomposition spectrale	76

4.5	Changer la notion de solution	77
5	Moindres carrés et inversion généralisée	79
5.1	Définition en dimension infinie	80
5.2	Définition en dimension finie	81
5.3	Algorithmes d'inversion généralisée	83
5.4	Décomposition en valeurs singulières	83
5.4.1	Méthode de BIALY	85
5.4.2	Méthode de LANDWEBER	85
5.4.3	Méthode de VAN-CITTERT	85
5.4.4	Méthode de KACZMARZ	86
6	Régularisation	87
6.1	Définition d'un opérateur de régularisation	87
6.2	Décomposition tronquée en valeurs singulières	88
6.3	Régularisation par méthodes itératives	89
6.4	Introduction de contraintes supplémentaires	89
6.5	Régularisation avec <i>a priori</i> de douceur	90
6.6	Algorithmes de régularisation	91
6.6.1	Méthode de HUNT (approximation circulante)	92
6.6.2	Méthodes itératives	95
6.6.3	Méthodes récursives	97
7	Classification des méthodes d'inversion	107
7.1	Méthodes analytiques	108
7.2	Méthodes de développement en série	112
7.3	Méthodes de décomposition sur une base	115
7.4	Méthodes algébriques déterministes	119
7.5	Méthodes probabilistes	121
8	Approches probabilistes	123
8.1	Approche n'utilisant que les moments	124
8.2	Estimation au sens du maximum de vraisemblance	127
8.3	Approche du maximum d'entropie	129
8.3.1	Définition de l'entropie	129
8.3.2	Lois à maximum d'entropie	130
8.3.3	Lois à minimum d'entropie relative	131
8.3.4	Extension au cas continu	132
8.3.5	Extension au cas multivariable	134
8.4	Utilisation pour les problèmes inverses	136
8.4.1	Maximum d'entropie classique	136
8.4.2	Maximum d'entropie sur la moyenne	136
9	Approche bayésienne	141
9.1	Introduction	142
9.2	Cas linéaire gaussien	145
9.3	Calcul de la loi <i>a posteriori</i>	147
9.3.1	Cas gaussien	147
9.3.2	Cas général	147

9.4	Choix des fonctions de coût	148
9.5	Choix de la loi <i>a priori</i>	149
9.5.1	Maximum d'entropie	150
9.5.2	Modèles markoviens	155
9.6	Quelques exemples de construction à la main	156
9.7	Estimation des paramètres d'une loi à partir des observations directes	163
9.7.1	Méthode des moments	163
9.7.2	Méthode du maximum de vraisemblance	164
9.8	Estimation des hyperparamètres	167
10	Modélisation markovienne	169
10.1	Introduction et notations	170
10.1.1	Site, voisinage et cliques	170
10.1.2	Champ aléatoire	172
10.1.3	Champ aléatoire markovien	172
10.2	Champs de Gibbs et équivalence Gibbs-Markov	173
10.2.1	Energie et Potentiel	173
10.2.2	Mesure de Gibbs	173
10.2.3	Équivalence Gibbs-Markov	173
10.3	Exemples de modèles classiques	174
10.3.1	Auto-modèles	174
10.3.2	Auto-logistique	174
10.3.3	Modèle d'Ising	174
10.3.4	Multi-Level Logistic Model (MLL)	175
10.4	Tentative de classification	176
10.5	Modèles couplés ou hiérarchiques	179
10.6	Application en segmentation d'image	180
10.7	Application en restauration d'image	182
10.8	Outils de simulation et d'optimisation	183
10.8.1	Echantillonneur de Gibbs	183
10.8.2	Recuit simulé	184
10.8.3	Algorithme d'ICM (Iterated Conditional Modes)	184
10.9	Algorithmes de relaxations déterministes	185
10.9.1	Principe de base du GNC	185
10.10	Problèmes et questions ouverts	189
11	Choix des hyperparamètres	191
11.1	Position du problème	192
11.2	Choix du coefficient de régularisation	195
11.2.1	Adéquation aux données	196
11.2.2	Risque moyen	197
11.2.3	Validation croisée	198
11.2.4	Validation croisée généralisée	199
11.3	Un autre regard	200
11.4	Algorithme EM	203
11.5	Exemples et exercices:	204

12 Étude de cas et exemples d'applications	215
12.1 Déconvolution des signaux	216
12.1.1 Un problème d'instrumentation	216
12.1.2 Méthodes naïves	216
12.1.3 Filtre de Wiener pour l'inversion	219
12.1.4 Régularisation quadratique pour l'inversion	221
12.1.5 Régularisation quadratique pour l'estimation de la réponse impul- sionnelle	224
12.1.6 Déconvolution aveugle	226
12.2 Restauration d'image	231
12.3 Reconstruction d'image en tomographie X	236
12.4 Reconstruction d'image en tomographie microondes	237
12.5 Reconstruction d'image en CND par courant de FOUCAULT	238
A Algorithmes d'optimisation pour un critère de moindre carrés	241
A.1 Introduction	242
A.2 Cas des critères moindres carrés	243
A.3 Cas des problèmes inverses avec un bruit additif	245
A.4 Comparaison entre les méthodes	247
B Algorithme de Gauss-Seidel pour optimisation d'un critère quadratique	249
B.1 Introduction	250
B.2 Le cas d'un critère régularisé	252
C Expressions des différentes fonctions potentiels	255
C.1 Potentiels convexes	256
C.2 Potentiels non convexes	261
D Quelques exemples simples d'inférence bayésienne	265
D.1 Introduction et rappel des notations	266
E Quelques rappels sur les matrices	275
E.1 Introduction	275
E.1.1 Définition	275
E.1.2 Opérations élémentaires	276
E.1.3 Matrices élémentaires	277
E.1.4 Rang d'une matrice	277
E.1.5 Matrice inverse	278
E.1.6 Déterminant d'une matrice	278
E.1.7 Matrice singulière	278
E.1.8 Trace d'une matrice	278
E.1.9 Vecteurs propres et valeurs propre d'une matrice	279
E.1.10 Valeurs singulières d'une matrice	280
E.1.11 Matrices et les opérations linéaires	280
E.1.12 Décomposition tronquée en valeurs singulière	282
E.1.13 Inverse généralisée d'une matrice	283
E.2 Quelques matrices célèbres	283
E.3 Matrices blocs	293
E.4 Matrices de Toeplitz et de Hankel	295

E.4.1	Matrices de Toeplitz	295
E.4.2	Matrices de Hankel	296
E.5	Matrices circulantes	297
E.6	Matrices bloc-Toeplitz et bloc-circulantes	299
E.6.1	Matrices bloc-Toeplitz	299
E.6.2	Matrices bloc-circulantes	300
E.7	Inversion récursive des matrices	301
E.8	Résolution des systèmes d'équations linéaires	305
E.8.1	Conditionnement d'un problème de résolution des systèmes linéaires	306
E.8.2	Conditionnement d'un problème de calcul des valeurs singulières d'une matrice	308
E.9	Méthodes récursives de résolution des systèmes d'équations linéaires	309
E.10	Diverses propriétés	313
F	Quelques transformations utiles	315
F.1	Transformations 1-D	316
F.2	Transformations 2-D	318
F.3	Transformations n -D	320
F.4	Transformations linéaires dans le cas discret	321
G	Classement bibliographique	327

Chapitre 1

Introduction

En sciences expérimentales, rares sont les situations où l'on peut mesurer directement une quantité f à laquelle on s'intéresse, pour étudier sa distribution spatiale $f(\mathbf{r})$ à un instant donné ou sa variation temporelle $f(t)$ à une position donnée ou encore sa variation spatio-temporelle $f(\mathbf{r}, t)$.

Souvent, cette quantité est mesurée par l'intermédiaire d'un simple capteur (ou d'un système de mesure plus complexe et rarement idéal) qui a pour rôle de fournir une grandeur mesurable g liée à f par une relation simple. Prenons comme exemple la mesure d'une température par un thermomètre classique. La grandeur réellement mesurée est une longueur qui, idéalement, est proportionnelle à la température. Mais, lorsque l'on souhaite étudier la variation temporelle de cette température les difficultés apparaissent : le capteur ne reproduit pas fidèlement la variation de la température lorsque celle-ci est un peu trop rapide (limitation de la bande passante du capteur). On essaie ensuite d'étudier le lien entre la sortie du capteur $g(t)$ et son entrée $f(t)$ en fournissant un *modèle*, souvent linéaire et invariant par translation—*modèle de convolution*—et donc caractérisé par sa *réponse impulsionnelle* $h(t)$ ou sa *réponse fréquentielle* $H(\omega)$. Ensuite on *analyse* ce modèle pour prescrire la plage de son utilisation.

L'établissement d'un lien entre la sortie $g(t)$ et l'entrée $f(t)$ est ce qu'on appelle *modélisation directe*. La détermination de la réponse impulsionnelle ou de la réponse fréquentielle est le problème de l'*identification*. Étant donné le modèle, l'étude de la sortie pour différentes formes de l'entrée est l'*analyse du problème directe*. Et finalement, étant donné le modèle et la sortie, la détermination de l'entrée est ce qu'on appelle l'*inversion* ou encore la *résolution du problème inverse*.

Lorsque le modèle liant la sortie $g(t)$ et l'entrée $f(t)$ est linéaire et invariant par translation on a un modèle de *convolution* :

$$g(t) = \int f(t')h(t - t') dt'.$$

Étant données l'entrée $f(t)$ et la réponse impulsionnelle $h(t)$, en général, on ne rencontre pas de problème particulier pour déterminer la sortie $g(t)$. Mais, malheureusement, les problèmes réels se formulent ainsi :

1. Ayant mesuré $f(t)$ et $g(t)$, déterminer la réponse impulsionnelle $h(t)$. C'est le problème de l'*identification de la réponse impulsionnelle*.
2. Connaissant $h(t)$ et ayant mesuré $g(t)$, déterminer l'entrée $f(t)$. C'est le problème inverse de la *déconvolution simple*.

3. Ayant seulement mesuré $g(t)$, déterminer à la fois la réponse impulsionnelle $h(t)$ et l'entrée $f(t)$. C'est le problème inverse de la *déconvolution aveugle*.

Prenons un deuxième exemple: l'imagerie tomographique, où on souhaite étudier la distribution spatiale de la densité de matière $f(x,y,z)$ à l'intérieur d'un corps. Pour la mesurer, nous avons besoin d'un intermédiaire: des rayons X par exemple. On illumine ce corps par ces rayons et on mesure les résultats de son interaction avec ces ondes: la perte d'énergie des rayons lors de la traversé du corps, c'est-à-dire les radiographies suivant plusieurs directions $p(r,s,\phi)$. On essaie ensuite d'étudier le lien entre ces mesures $p(r,s,\phi)$ et l'inconnue $f(x,y,z)$ en fournissant un *modèle* liant ces quantités. Le problème inverse consiste alors à déterminer la fonction $f(x,y,z)$ à partir de ces *projections* $p(r,s,\phi)$. Lorsqu'on s'intéresse à une section du corps étudié $f(x,y,z_0)$ on parle de la *reconstruction tomographique* ou de la *reconstruction d'image 2D* et dans le cas général on parle de la *reconstruction 3D*. Dans le cas 2D, une modélisation simple entre $f(x,y)$ et $p(r,\phi)$ est la transformée de Radon (TR):

$$p(r,\phi) = \int_L f(x,y) dl = \iint f(x,y) \delta(r - x \cos \phi - y \sin \phi) dx dy$$

L'objet de ce texte est de fournir un cadre général et compréhensible pour la description et la résolution de ces problèmes.

- Dans le chapitre 2, nous montrons à travers des exemples que les problèmes inverses se trouvent au cœur d'un très grand nombre de domaines des sciences expérimentales et plus particulièrement des sciences pour l'ingénieur.
- Dans le chapitre 3, nous essayons de fournir un cadre général pour les problèmes inverses et de montrer que les différents exemples d'applications que nous avons présentés au chapitre précédent sont des cas particulier de ce cadre général.
- Dans le chapitre 4¹, nous essayons de définir les notions des problèmes bien-posés et mal-posés et de montrer pourquoi la résolution des problèmes inverses est une tâche difficile.
- Dans le chapitre 5, nous introduisons les notions de pseudo-solutions, solutions au sens des moindres carrés (MC) et solution inverse généralisée (IG). Nous verrons ensuite que, malheureusement, bien que ces notions permettent, sous certaines conditions, de définir une solution unique pour un problème inverse mal-posé, elles ne permettent pas de fournir une solution satisfaisante à ces problèmes.
- Le chapitre 6 est consacré à la notion de régularisation et à l'étude des différentes techniques de régularisation pour la résolution des problèmes inverses linéaires.
- Dans le chapitre 7, nous essaierons de classifier les différentes méthodes d'inversion. Ceci nous permettra d'avoir une vue plus synthétique de ces différentes méthodes et des liens qui existent entre elles.
- Dans le chapitre 8, nous présentons la résolution des problèmes inverses comme un problème d'estimation dans un cadre probabiliste. Nous distinguerons ensuite quatre approches:

1. Approche estimation au sens des moindres carrés où estimation linéaire en moyenne quadratique où seuls les moments d'ordre un et deux sont utilisés,

1. Les contenus des chapitres 4, 5 et 6, restent pour l'instant très sommaires et sont inspirés du polycopié du Guy Demoment intitulé "Déconvolution des signaux". Ce polycopié est disponible à la bibliothèque de Supélec.

ce qui revient implicitement à faire l'hypothèse que toutes les variables sont gaussiennes. Pour illustrer cette approche nous prendrons comme exemple le filtrage optimal ;

2. Approche estimation au sens du maximum de vraisemblance où nous verrons rapidement que cette approche, bien que plus générale, dans le sens où l'on peut prendre en compte explicitement la loi des mesures, est peu satisfaisante pour la résolution des problèmes inverses ;
 3. Approche maximum d'entropie où, après un rappel sur le principe du maximum d'entropie et son utilité pour l'attribution d'une loi de probabilité, nous verrons que cette approche permet de fournir une nouvelle interprétation de la notion de régularisation et permet de répondre à un certain nombre de questions concernant le choix des fonctionnelles de régularisation ;
 4. Approche estimation bayésienne où nous verrons comment elle peut fournir un cadre fédérateur pour la définition d'une solution régularisée englobant tous les critères précédents. Nous verrons ensuite que cette approche est plus cohérente, plus générale, plus simple et plus efficace pour aborder les problèmes inverses.
- Le chapitre 9 est entièrement réservé à l'approche bayésienne. Dans ce chapitre, nous détaillerons cette approche et présenterons un certain nombre de cas particuliers qui illustreront la généralité et l'universalité de cette approche.
 - Dans le chapitre 11, nous donnons un aperçu général de la modélisation markovienne des signaux et des images. Les modèles markoviens et particulièrement les champs de Gibbs sont devenus les outils indispensables pour la traduction d'une information *a priori* plus élaborée que la douceur ou la positivité en une loi de probabilité. En effet, la modélisation markovienne permet de construire des modèles qui peuvent prendre en compte des discontinuités dans un signal ou une image. Cette outil prend particulièrement son importance en traitement d'image où l'on peut construire les modèles composites (bords, contours, intensité) pour les images et les utiliser lors de la segmentation, restauration ou reconstruction d'image.
 - Dans le chapitre 12, nous aborderons le problème de l'estimation des hyperparamètres, et en particulier, la détermination du paramètre de régularisation. Dans un premier temps, nous décrirons brièvement les méthodes classiques du type : adéquation aux données, la courbe en L, la validation croisée, et la validation croisée généralisée. Ensuite, nous présenterons le problème général de l'estimation des hyperparamètres dans le cadre bayésien.
 - Dans les chapitres 13, 14 et 15 nous présenterons un certain nombre d'*études de cas* comme la déconvolution des signaux, la restauration d'image, la reconstruction d'images,

Un certain nombre d'annexes sont présentées à la fin du document qui ont pour but soit de compléter certains détails ou d'apporter des éléments et des outils qui sont nécessaires lorsqu'un ingénieur veut aborder la résolution de problèmes inverses.

Et, finalement, une bibliographie classifiée conclura ce document.

Chapitre 2

Exemples de problèmes inverses

L'objet de ce chapitre est de montrer que les problèmes inverses se trouvent au cœur d'un très grand nombre de domaines de la science expérimentale et plus particulièrement de la science pour l'ingénieur. Bien entendu, la description de chaque application ne peut qu'être très brève.

2.1 Instrumentation

Un capteur ou un appareil de mesure serait idéal si sa sortie $g(t)$ était proportionnelle à son entrée $f(t)$ pour tout t , c'est-à-dire lorsque

$$\forall t \quad g(t) = af(t) + b,$$

a et b étant deux constantes fixées.

Ceci est rare. Le mieux que l'on puisse demander est que sa sortie soit une fonction linéaire de son entrée :

$$g(t) = \int f(\tau)h(t, \tau) d\tau.$$

Mieux encore, c'est le cas où sa réponse $h(t, \tau)$ ne varie pas en cours du temps, ce qui veut dire $h(t, \tau) = h(t - \tau)$. On a ainsi

$$g(t) = \int f(\tau)h(t - \tau) d\tau.$$

Il s'agit d'une équation de convolution. La fonction $h(t)$ est appelée la réponse impulsionnelle. Lorsque le système est causal, i.e; $h(t) = 0, \forall t < 0$, on a

$$g(t) = \int_0^t f(\tau)h(t - \tau) d\tau.$$

Un capteur serait parfait si

$$h(t) \simeq a\delta(t) \quad \longrightarrow \quad f(t) = \frac{1}{a}g(t)$$

Lorsque ceci n'est pas le cas, les ingénieurs ont l'habitude de travailler avec la réponse indicielle du capteur :

$$s(t) = \int_0^t h(t - \tau) d\tau$$

sur laquelle ils mesurent, le *temps de montée*, le *dépassement*, le *déphasage*, la *résolution maximale*, le critère de Raleigh, etc. Et lorsque ces grandeurs ne sont pas satisfaisantes ils changent le capteur. Mais cela peut être coûteux et même parfois techniquement impossible.

Que peut-on faire alors? Il y a vingt ans, cette question n'avait pas de réponse pratique. Aujourd'hui, avec les techniques numériques de l'inversion, on peut proposer la *déconvolution*. L'idée de base est de placer derrière le capteur un microprocesseur qui, partant des échantillons de la réponse impulsionnelle $h(t)$ et ceux de la sortie du capteur $g(t)$ fournirais $\hat{f}(t)$ qui est beaucoup plus proche de $f(t)$ que $g(t)$, dépassant ainsi le critère de Raleigh. En effet, $\hat{f}(t)$ aura une meilleure résolution que $g(t)$. On peut même par des techniques adaptatives compenser les effets d'une dégradation des caractéristiques du capteur en réestimant un certain nombre de paramètres de sa réponse impulsionnelle au cours du temps.

Ainsi, au lieu d'exiger du capteur d'avoir un temps de montée de telle valeur, un dépassement de telle autre valeur, etc., on lui imposerait d'être *linéaire*, *fidèle*, et *précis*. Le traitement numérique permettra alors de compenser le manque de résolution ou la limitation de la bande passante, etc.

$$\left\{ \begin{array}{l} \text{linéarité} \\ \text{fidélité} \\ \text{précision} \end{array} \right. + \text{Déconvolution} = \left\{ \begin{array}{l} \text{Amélioration de la résolution} \\ \text{Dépassement du critère de Raleigh} \end{array} \right.$$

2.2 Tomographie à rayons X

Si un objet de matière homogène est illuminé par des rayons X et qu'un rayon mono-énergétique de faible largeur traverse cet objet, l'intensité du rayon mesuré à la sortie s'obtient par

$$I = I_0 \exp[-\mu L]$$

où I_0 est l'intensité du rayon à l'entrée, L est la distance parcourue par le rayon, et μ est la constante d'atténuation linéaire de l'objet qui dépend de la densité de matière et de sa composition nucléaire.

Considérons maintenant une section transversale d'un objet non homogène perpendiculaire à l'axe oz , et supposons que les rayons traversent cette section. En caractérisant l'objet par la distribution de sa constante d'atténuation linéaire μ , fonction continue des deux variables d'espace; $\mu(x,y) = f(x,y)$; on a alors

$$\frac{I}{I_0} = \exp \left[- \int_L f(x,y) dl \right]$$

ou bien

$$-\ln \left(\frac{I}{I_0} \right) = \int_L f(x,y) dl$$

où dl est l'élément de longueur sur le parcours L . Ainsi, si on déplace en parallèle l'émetteur (source S) et le récepteur (détecteur D) sur une ligne droite faisant un angle ϕ avec l'axe ox on obtient une projection $p(r,\phi)$:

$$p(r,\phi) = \int_L f(x,y) dl$$

En se référant aux figures 2.2 et 2.2, on voit que cette équation peut aussi s'écrire sous la forme

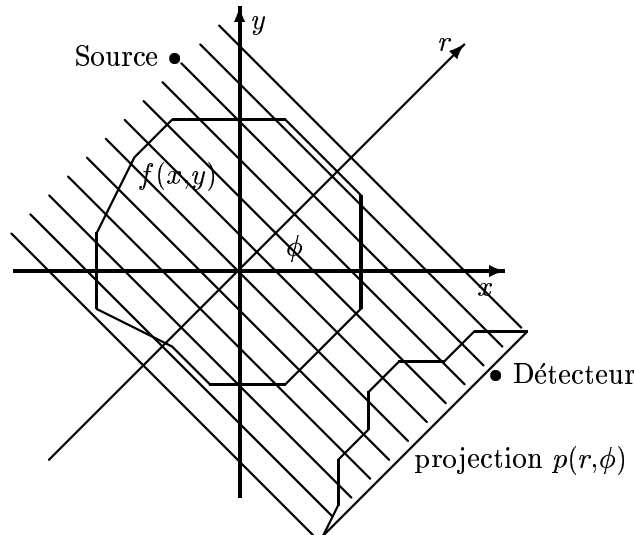


FIG. 2.1 – Tomographie à rayons X

$$p(r,\phi) = \iint_D f(x,y) \delta(r - x \cos \phi - y \sin \phi) dx dy$$

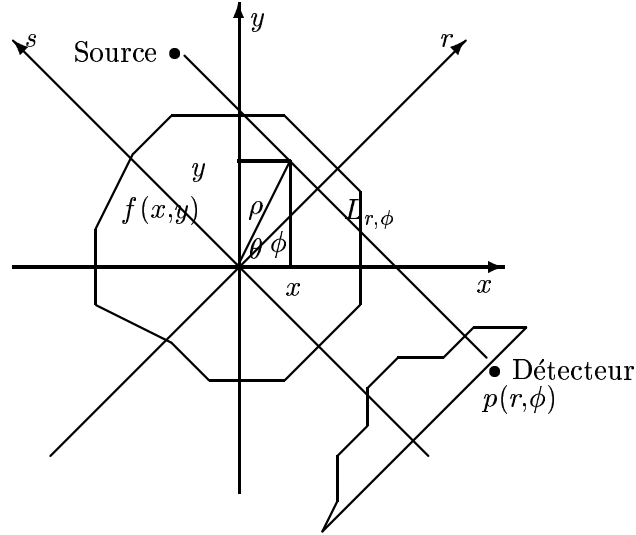


FIG. 2.2 – Notations utilisées pour la définition de la projection $p(r, \phi)$ d'une fonction $f(x, y)$

Pour des variables continues r et ϕ , $p(r, \phi)$ est, par définition, la transformée de Radon de la fonction $f(x, y)$. J. RADON a montré l'existence et l'unicité de cette transformée et de son inverse pour des fonctions continues à supports compacts. Ceci peut être résumé ainsi :

- Si $f(x, y)$ est continue et à support compact, alors $p(r, \phi) = \mathcal{R}\{f(x, y)\}$ est déterminée d'une façon unique par le résultat de l'intégration sur toutes les lignes droites $L_{r, \phi}$ dans le plan $(x, y) \in \mathbb{R}^2$, $r \in \mathbb{R}$, $\phi \in [0, \pi]$.
- La connaissance de $p(r, \phi)$ en tout point $r \in \mathbb{R}$ et $\phi \in [0, \pi]$ permet de déterminer d'une façon unique la fonction $f(x, y)$ par :

$$f(x, y) = \left(-\frac{1}{2\pi^2}\right) \int_0^\pi \int_{-\infty}^{+\infty} \frac{\frac{\partial}{\partial r} p(r, \phi)}{(r - x \cos \phi - y \sin \phi)} dr d\phi$$

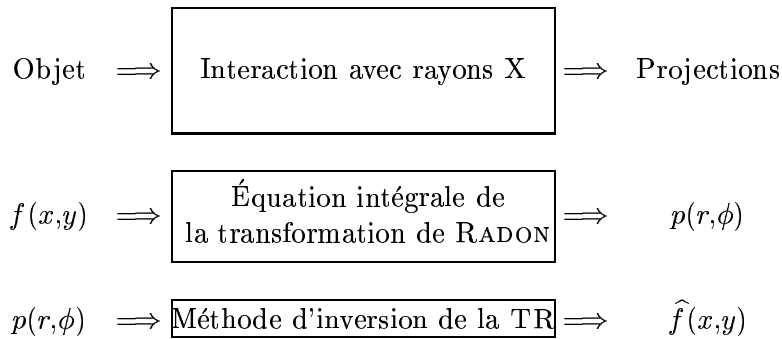


FIG. 2.3 – Reconstruction d'image en tomographie à rayons X

Voici ici un exemple de problème inverse pour lequel nous avons directement une expression analytique pour sa solution et on peut s'imaginer qu'il n'y a plus de problème.

En pratique cependant, les mesures sont discrètes. En particulier les projections sont obtenues pour les valeurs discrètes de ϕ . Il est alors plus convenable de dire que chaque projection $p(r, \phi_i)$ nous fournit un échantillonnage de la TR de la fonction $f(x, y)$. La multiplication des projections obtenues dans des conditions différentes nous permet de “remplir” progressivement le domaine de la transformée. Si ce remplissage pouvait être dense et uniforme alors l’inversion pourrait être effectuée par la transformation inverse pour récupérer l’objet. Mais la connaissance de tous les échantillons n’est malheureusement pas possible dans une situation réelle et on ne peut mesurer qu’un sous-ensemble de ces échantillons. Nous verrons que c’est là que se situe la difficulté de la résolution du problème.

Pour mieux comprendre les fondements des méthodes de reconstruction d’image il est utile de noter qu’il existe une relation fondamentale entre la TR et la TF d’une fonction. Elle se résume ainsi :

Si on note par $P(\Omega, \phi)$ la TF de $p(r, \phi)$ par rapport à la variable r :

$$P(\Omega, \phi) = \int p(r, \phi) \exp[-j\Omega r] dr$$

et par $F(\omega_x, \omega_y)$ la TF 2D de $f(x, y)$:

$$F(\omega_x, \omega_y) = \iint f(x, y) \exp[-j\omega_x x - j\omega_y y] dx dy$$

alors, il est facile de montrer la relation suivante plus connue sous le nom du *Théorème de projection* :

$$F(\omega_x, \omega_y) = P(\Omega, \phi) \quad \text{pour} \quad \omega_x = \Omega \cos \phi \quad \text{et} \quad \omega_y = \Omega \sin \phi$$

Cette relation permet de transformer le problème de la reconstruction d’image $f(x, y)$ à partir de ses projections $p(r, \phi)$ en un problème de *synthèse de FOURIER*, qui consiste à déterminer la fonction $f(x, y)$ à partir d’une connaissance partielle de sa transformée de FOURIER $F(\omega_x, \omega_y)$.

Encore une fois, on peut constater que si l’on pouvait connaître parfaitement toutes les projections on pourrait alors remplir le domaine de FOURIER de l’objet (connaître $F(\omega_x, \omega_y)$) et ensuite par une transformation de FOURIER inverse on pourrait déterminer la fonction $f(x, y)$. Mais, en pratique, ceci est rarement possible et le passage dans le domaine de FOURIER ne change rien.

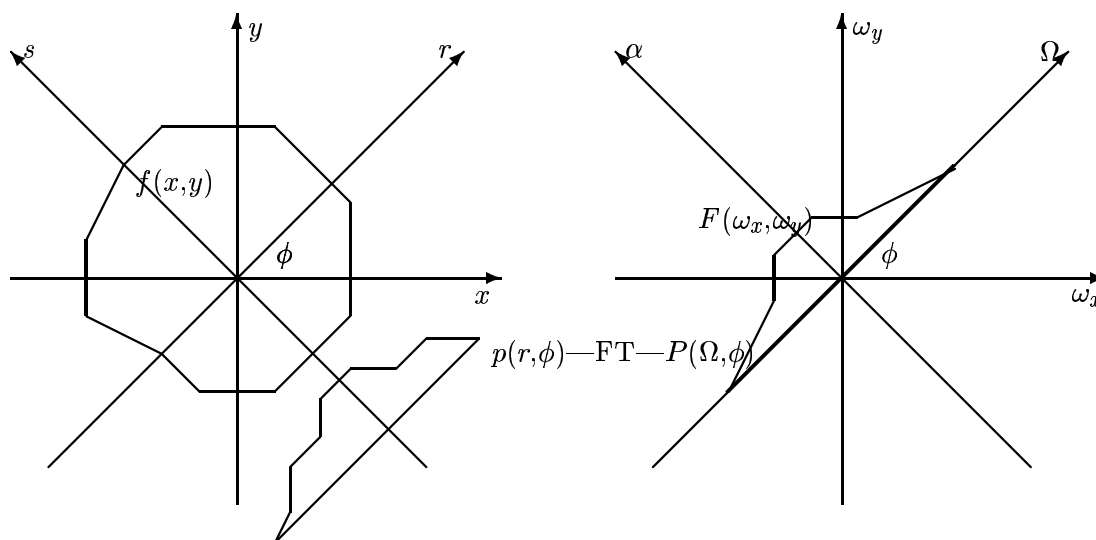


FIG. 2.4 – *Théorème de projection : Relation entre la TF 1-D de la projection $p(r,\phi)$ et la TF 2-D de $f(x,y)$*

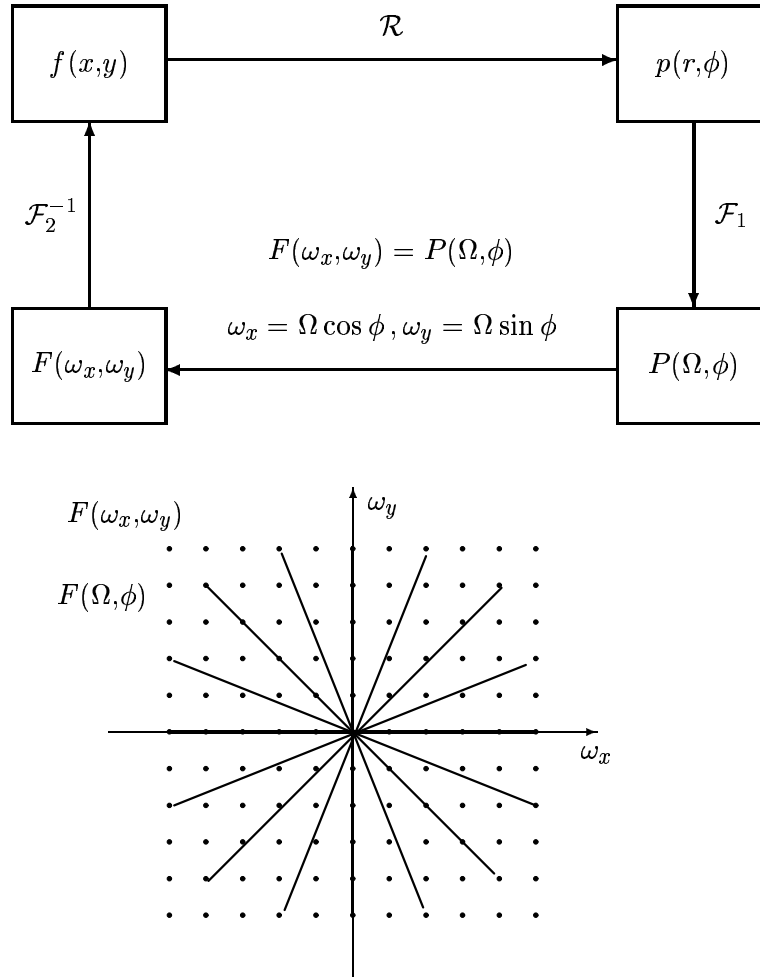


FIG. 2.5 – *Reconstruction d'image en tomographie X en passant dans le domaine de FOURIER.*

2.3 Imagerie à ondes diffractées : cadre général

D'une manière générale lorsqu'on cherche à *voir* l'intérieur d'un objet (sans le couper), on peut envisager deux procédés :

- soit l'illuminer par une onde (rayons X, ultrasons ou micro ondes) de l'extérieur et mesurer son interaction avec celle-ci; c'est l'imagerie active;
- soit profiter éventuellement de son propre rayonnement lorsque cela est possible (infra rouge, radioactivité) et mesurer les effets de ces rayons à l'extérieur de l'objet; c'est l'imagerie passive.

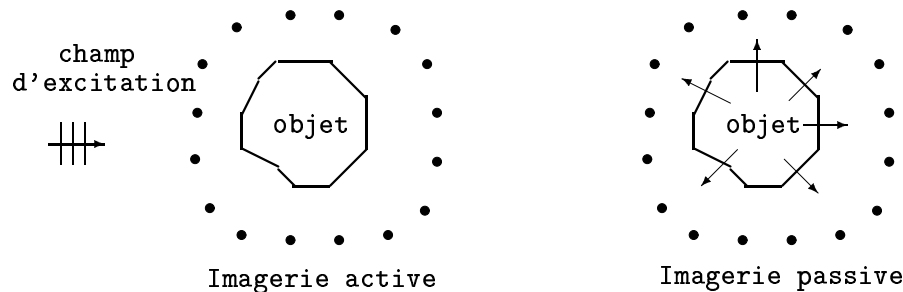


FIG. 2.6 – *Imagerie active et passive.*

Dans le cas d'imagerie active que nous considérons ici, on peut envisager trois cas :

- Imagerie par transmission où la ligne de mesure se trouve du côté opposé du champ incident ;
- Imagerie par réflexion où la ligne de mesure se trouve du même côté que celle du champ incident ;
- Imagerie mixte où les points de mesure se trouvent tout autour de l'objet.

Dans tous les cas, il faut d'abord établir une relation mathématique (modèle) entre ces mesures et une propriété interne de l'objet qui le caractérise, et ensuite inverser cette relation.

D'une manière plus générale, dans toute technique d'imagerie, il y a trois problèmes fondamentaux à résoudre :

- un problème direct, qui consiste à établir un modèle mathématique suffisamment simple décrivant le mieux possible la relation entre les grandeurs mesurées et les grandeurs inconnues auxquelles nous nous intéressons ;
- un problème d'instrumentation, qui consiste à construire un système d'imagerie : générer les ondes, concevoir les capteurs et mesurer les résultats des interactions de l'onde avec l'objet ;
- un problème inverse, qui consiste à développer des méthodes d'inversion qui, partant des données mesurées, reconstruisent ces images qui caractérisent l'objet étudié. Bien entendu, nous voulons que ces résultats soient, à la fois, robustes vis-à-vis de l'inévitable bruit de mesure et des erreurs d'approximation, de la discrétisation et de la quantification qui sont nécessaires pour le calcul numérique, et qu'ils aient en même temps la plus grande résolution spatiale possible.

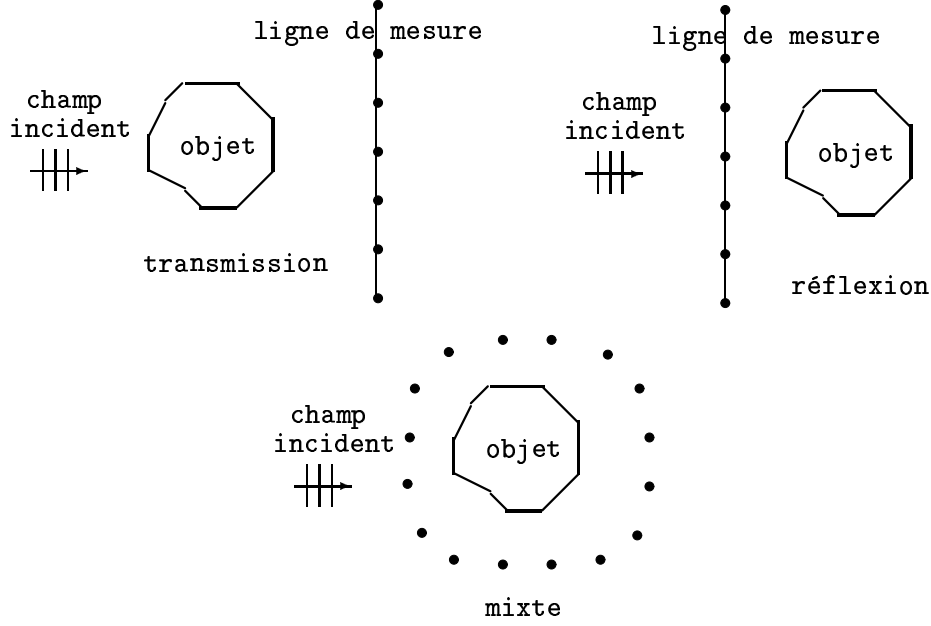


FIG. 2.7 – *Différents modes d'imagerie : par réflexion, par transmission et mode mixte.*

Dans la section précédente, nous avons vu comment les rayons X peuvent servir pour obtenir une carte de distribution de la densité de matière à l'intérieur d'un objet. Ici, nous allons voir qu'il est possible de remplacer les rayons X par d'autres types d'ondes, par exemple les ultrasons ou les micro ondes. Mais, en général, contrairement aux rayons X et due au fait que les longueurs d'ondes utilisées sont de l'ordre de dimensions de l'objet, l'établissement d'une relation entre les mesures (champs diffractés) et l'objet devient plus complexe.

Ici, sans entrer trop dans le détail, nous allons décrire le cadre général qui nous permettra d'identifier les problèmes inverses associés.

Lorsqu'un objet $f(\mathbf{r})$ est illuminé par une onde incidente $\phi_0(\mathbf{r})$ le champ total $\phi(\mathbf{r})$ peut être modélisé par la somme du champ incident $\phi_0(\mathbf{r})$ et du champ diffracté $\phi_s(\mathbf{r})$ qui dépend à la fois de l'objet et du champ incident :

$$\phi(\mathbf{r}) = \phi_0(\mathbf{r}) + \int_D G_o(\mathbf{r}, \mathbf{r}') \phi(\mathbf{r}') f(\mathbf{r}') d\mathbf{r}'$$

et le champ diffracté à l'extérieur de l'objet $g(\mathbf{u})$ que l'on mesure est une fonction de l'objet et du champ total à l'intérieur de l'objet :

$$g(\mathbf{r}') = \int_D G_m(\mathbf{r}, \mathbf{r}') f(\mathbf{r}) \phi(\mathbf{r}) d\mathbf{r}.$$

Dans ces deux équations G_o et G_m sont des fonctions de GREEN. Notons que la première équation est une équation implicite en ϕ et que la deuxième équation est bilinéaire en f et en ϕ .

Le problème de l'imagerie dans ce contexte est de fournir une estimation de $f(\mathbf{r})$ à partir d'un nombre fini de mesures $g(\mathbf{r}'_i), i = 1, \dots, M$. Comme on peut constater, la

relation entre g et f n'est pas linéaire, ce qui nous donne un exemple type d'un problème inverse non linéaire.

Le cas de l'approximation de BORN étant la situation de faibles perturbations où l'on peut remplacer $\phi(\mathbf{r})$ par $\phi_0(\mathbf{r})$ dans cette dernière équation. Cette approximation permet d'établir une relation linéaire entre f et g et ainsi transformer le problème inverse non linéaire à un problème inverse linéaire :

$$g(\mathbf{r}') = \int_D G_m(\mathbf{r}, \mathbf{r}') f(\mathbf{r}) \phi_0(\mathbf{r}) \, d\mathbf{r}.$$

Par la suite nous allons décrire les techniques d'imagerie basées sur cette approximation.

2.4 Imagerie micro ondes : cas 2D

Pour comprendre les principales idées de l'imagerie à ondes diffractées, nous allons considérer le cas particulier de l'imagerie micro ondes. Nous nous placerons dans le cas d'un problème 2D. Pour cela considérer le cas où un objet cylindrique est illuminé par une onde plane polarisée rectiligne selon l'axe oz .

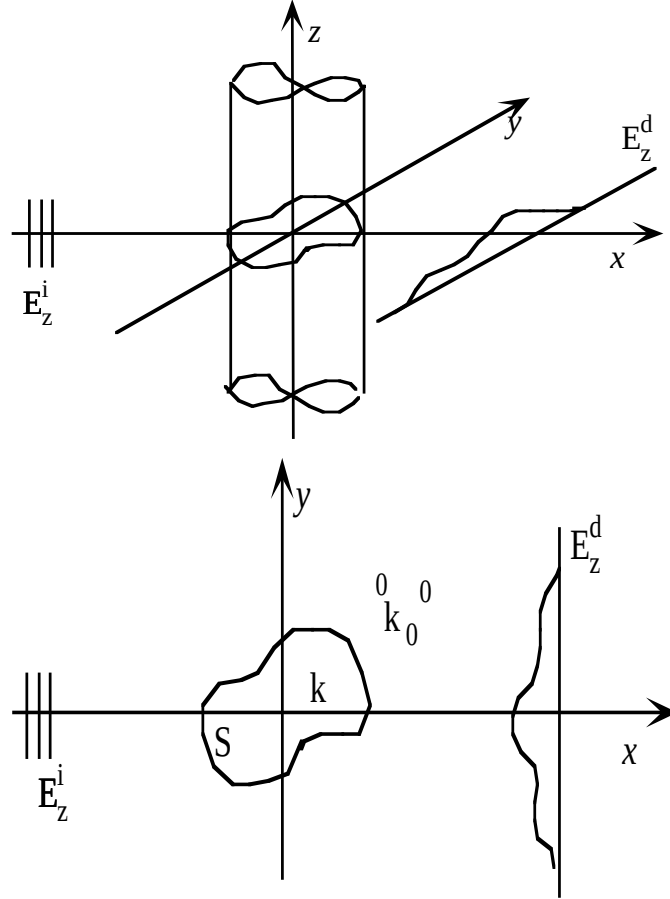


FIG. 2.8 – Géométrie de l'imagerie micro ondes.

Si on fait l'hypothèse que l'objet est homogène suivant l'axe oz ($\frac{\partial}{\partial z} = 0$) alors les champs incident et diffracté n'auront qu'une seule composante et on a les relations suivantes :

$$E_z(\mathbf{r}) = E_z^i(\mathbf{r}) + E_z^d(\mathbf{r})$$

$$E_z^d(\mathbf{r}) = \int G(\mathbf{r} - \mathbf{r}') J(\mathbf{r}') d\mathbf{r}'$$

où

$$J(\mathbf{r}') = [k^2(\mathbf{r}') - k_0^2] E_z(\mathbf{r}') \quad \text{et} \quad k(\mathbf{r}') = \sqrt{\omega^2 \mu_0 \epsilon(\mathbf{r}') - j\omega \mu_0 \sigma(\mathbf{r}')}$$

Deux cas peuvent alors être considérés :

- On s'intéresse à la quantité $J_n(\mathbf{r}') = J(\mathbf{r}')/E_z^i(\mathbf{r}')$ qui représente la densité des courants induits dans l'objet (normalisé par rapport au champ incident) et qui dépend à la fois de l'objet et du champ incident, mais qui a la propriété

$$f(\mathbf{r}') = [k^2(\mathbf{r}') - k_0^2] = 0 \longrightarrow J(\mathbf{r}') = 0.$$

On a alors

$$E_z^d(\mathbf{r}) = \int G(\mathbf{r} - \mathbf{r}') J_n(\mathbf{r}') E_z^i(\mathbf{r}') d\mathbf{r}'$$

- On s'intéresse à la quantité $o(\mathbf{r}') = [k^2(\mathbf{r}') - k_0^2]$ qui représente directement l'objet et on a alors

$$E_z^d(\mathbf{r}) = \int G(\mathbf{r} - \mathbf{r}') o(\mathbf{r}') E_z^i(\mathbf{r}') d\mathbf{r}'$$

Mais en se mettant dans le cadre de l'approximation de BORN, c'est-à-dire $E_z(\mathbf{r}') \simeq E_z^i(\mathbf{r}')$, alors on a

$$E_z^d(\mathbf{r}) = \int G(\mathbf{r} - \mathbf{r}') o(\mathbf{r}') E_z^i(\mathbf{r}') d\mathbf{r}'$$

Dans les deux cas, si on note par $\phi(\mathbf{r}) = E_z^d(\mathbf{r})$, $\phi_0(\mathbf{r}') = E_z^i(\mathbf{r}')$ et par $f(\mathbf{r}')$ soit $J_n(\mathbf{r}')$ soit $o(\mathbf{r}')$, alors les deux problèmes se confondent à

$$\phi(\mathbf{r}) = \int G(\mathbf{r} - \mathbf{r}') f(\mathbf{r}') \phi_0(\mathbf{r}') d\mathbf{r}'$$

Maintenant, remplaçant l'expression de la fonction de GREEN dans le cas 2D :

$$G(\mathbf{r} - \mathbf{r}') = \frac{j}{4} k_0^2 H_0(k_0 |\mathbf{r} - \mathbf{r}'|)$$

Regardant de plus près cette équation on constate qu'il s'agit d'une opération de convolution entre $G(\mathbf{r})$ et le produit $f(\mathbf{r}') \phi_0(\mathbf{r}')$. Passant dans le domaine de FOURIER on a

$$\text{TF}\{\phi\} = \text{TF}\{G\} \quad [\text{TF}\{f\} * \text{TF}\{\phi_0\}]$$

ou encore

$$\widehat{\phi}(\boldsymbol{\omega}) = \widehat{G}(\boldsymbol{\omega}) \left[\widehat{f}(\boldsymbol{\omega}) * \widehat{\phi}_0(\boldsymbol{\omega}) \right]$$

et où

$$\widehat{G}(\boldsymbol{\omega}) = \frac{k_0^2}{k_0^2 - \omega^2} \quad \text{avec} \quad \omega^2 = |\boldsymbol{\omega}|^2.$$

Supposons maintenant que le champ incident est une onde plane se propageant en direction de l'axe ox : $\phi_0 = \exp[j\mathbf{k}_0 \mathbf{r}]$ avec $\mathbf{k}_0 = [k_0, 0, 0]$, on a alors :

$$\widehat{\phi}_0(\boldsymbol{\omega}) = \delta(\boldsymbol{\omega} - \mathbf{k}_0)$$

et

$$\widehat{\phi}(\boldsymbol{\omega}) = \frac{k_0^2}{k_0^2 - \omega^2} \widehat{f}(\boldsymbol{\omega} - \mathbf{k}_0)$$

Ainsi, il y a une relation entre la TF spatiale de l'objet et la TF spatiale du champ diffracté. Mais pour être plus précis, mettons nous dans le cas 2D où l'on mesure le champ diffracté sur une ligne droite perpendiculaire à la direction de propagation du champ incident. On peut alors montrer la relation suivante :

$$\widehat{f}(\omega_x = -k_0 + \sqrt{k_0^2 - \omega_y^2}, \omega_y) = \widehat{\phi}(\omega_y)$$

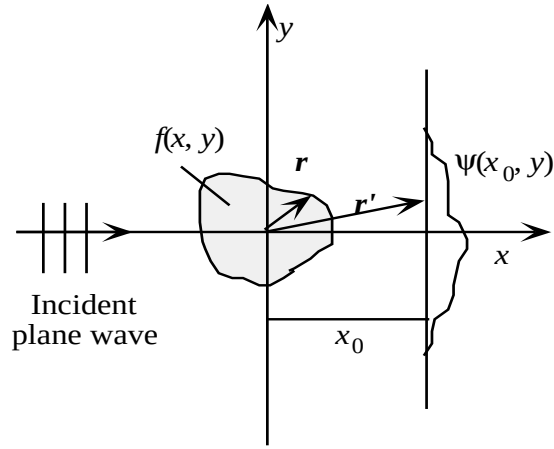
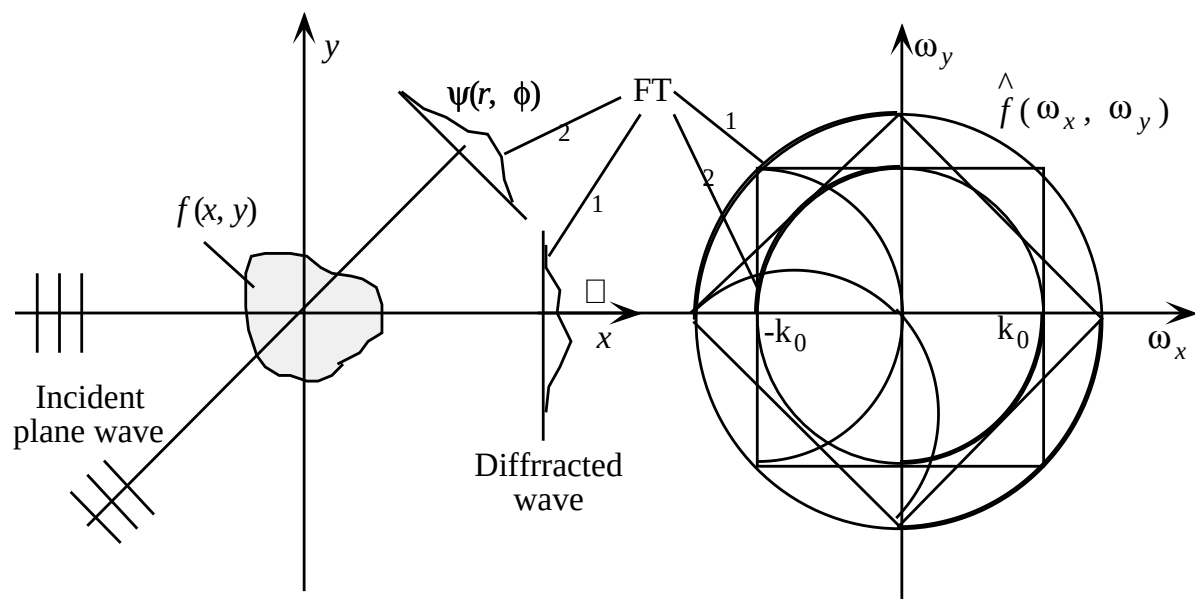


FIG. 2.9 – Géométrie de la tomographie de diffraction 2D.

En d'autres termes, les valeurs de la TF 1D du champ diffracté mesuré sur une ligne droite perpendiculaire à la direction de propagation d'une onde plane, sont égales aux valeurs de la TF 2D de l'objet sur un demi-cercle de rayon k_0 centré au point $\mathbf{k} = (k_0, 0)$ (voir figure). videmment, lorsque la direction de propagation fait un angle θ avec l'axe ox , on a la même relation, sauf que le demi-cercle aussi sera tourné d'un angle θ .

La mise en évidence de cette relation permet de transformer le problème de la reconstruction d'objet $f(x, y)$ à partir de mesures des champs diffractés $\phi(r, \theta_i)$ en un problème de synthèse de FOURIER qui consiste en la détermination de $f(x, y)$ à partir d'une connaissance partielle de sa TF. De plus, on peut faire l'analogie avec le problème de la reconstruction d'image en tomographie X. En effet en tomographie X, les projections nous fournissaient les valeurs de la TF de l'objet sur des lignes droites, alors qu'ici, les champs diffractés nous fournissent les valeurs de la TF de l'objet sur des demi-cercles dont le rayon est inversement proportionnel à la longueur d'onde utilisée.

FIG. 2.10 – *Synthèse de FOURIER en tomographie de diffraction 2D.*

2.5 Imagerie micro ondes : cas 3D

Les mêmes raisonnements et les mêmes équations peuvent être obtenues dans le cas de l'imagerie 3D, si l'onde incidente est plane et si le plan sur lequel les mesures du champ diffracté sont effectuées, on aura la même relation dans le cas 3D.

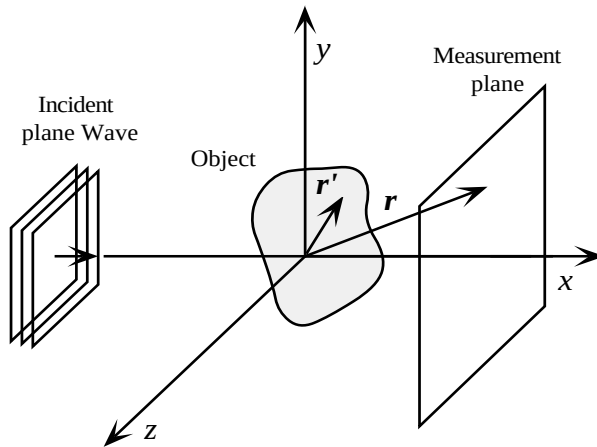


FIG. 2.11 – Géométrie de la tomographie de diffraction 3D.

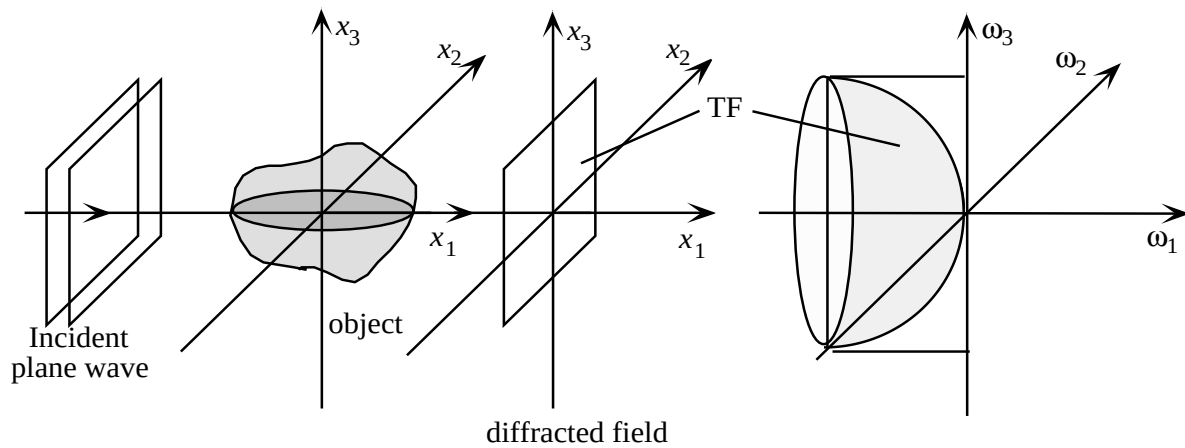


FIG. 2.12 – Synthèse de FOURIER en tomographie de diffraction 3D.

2.6 Imagerie par courants de Foucault

Les techniques de contrôle non destructif utilisent couramment les courants de FOUCAULT pour détecter une anomalie dans un milieu conducteur. L'imagerie par courants de FOUCAULT est plus récente. Son objectif est, une fois une anomalie détectée, de fournir une image qui représente la variation de la conductivité à l'intérieur d'un objet conducteur. Cette image peut alors être utilisée pour fournir un diagnostic plus précis sur le défaut.

Bien que les fréquences utilisées pour exciter l'objet par cette technique sont de l'ordre du kHz à MHz, on peut encore utiliser le cadre général développé en imagerie micro ondes pour obtenir des équations liant l'objet aux mesures. Trois différences importantes cependant existent :

- En tomographie de diffraction, en général, l'objet étudié est dans un milieu homogène où se trouvent les capteurs (émetteurs et récepteurs). En imagerie par courants de FOUCAULT le défaut se trouve à l'intérieur d'un milieu homogène (bloc métallique), qui est différent du milieu où se trouvent les capteurs (air).
- En tomographie de diffraction, en général, on peut négliger l'atténuation des ondes, alors qu' en imagerie par courants de FOUCAULT cela n'est pas le cas. En effet, les ondes se propagent dans un milieu conducteur où les deux phénomènes de diffusion et de propagation sont d'importance égale (la constante de propagation est un nombre complexe dont la partie réelle et la partie imaginaire sont de même ordre de grandeur).
- En tomographie de diffraction, en général, on peut tourner autour de l'objet, ce qui permet de remplir le domaine de FOURIER de l'objet. En imagerie par courants de FOUCAULT, en général, il n'est pas possible de tourner autour de l'objet. C'est pourquoi, ici, on modifie la fréquence temporelle du champ d'excitation pour obtenir des données complémentaires, ce qui permet dans un sens de remplir partiellement le domaine de FOURIER de l'objet.

Pour donner une description simple de cette technique considérons le cas 2D de la figure (2.6).

Supposons que le milieu conducteur est suffisamment épais et que le défaut se trouve près de la surface pour ne considérer que deux milieux D_1 (air) et D_2 (métal) avec une interface que nous supposerons plane. Supposons aussi que la région D_2 est amagnétique. Partant des équations de Helmholtz, on peut obtenir une équation intégrale liant la composante normale du champ électrique E_z ou les composantes tangentielles du champ magnétique (H_x et H_y) aux sources de courants induites par le défaut. Pour cela, nous noterons :

- $\vec{E}_0, \vec{H}_0$, les champs électrique et magnétique en l'absence du défaut ;
- $\vec{E}_f, \vec{H}_f$, les champs en présence du défaut ;
- $\vec{E}_s = \vec{E}_0 - \vec{E}_f, \vec{H}_s = \vec{H}_0 - \vec{H}_f$, les différences des champs en l'absence et en présence du défaut ;
- $k_1^2 = \omega^2 \mu_0 \epsilon_0$, la constante de propagation dans la région R_1 ;
- $k_2^2 = \omega^2 \mu_0 \epsilon_0 + j\omega \mu_0 \sigma_0 \simeq j\omega \mu_0 \sigma_0$, la constante de propagation dans la région R_2 ;
- $k_\Omega^2(x,y) = \omega^2 \mu_0 \epsilon_0 + j\omega \mu_0 \sigma(x,y)$, la constante de propagation dans la région du défaut ;
- $f(x,y) = \frac{k_2^2 - k_\Omega^2(x,y)}{k_2^2}$, la fonction du contraste qui caractérise le défaut ;

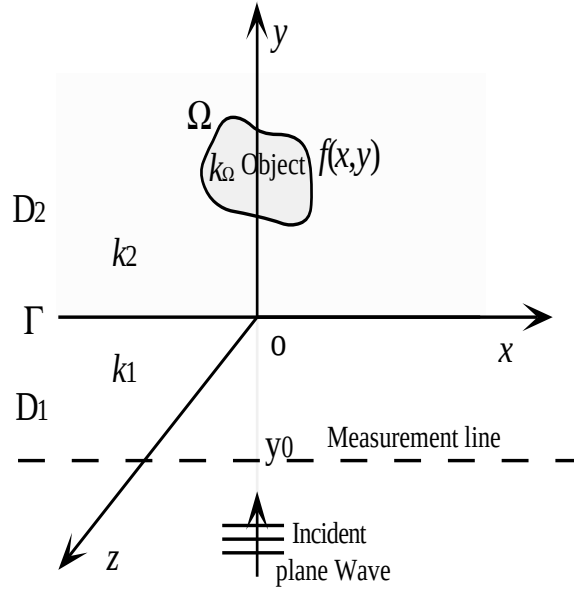


FIG. 2.13 – Géométrie de l'imagerie par courants de FOUCAULT.

- $J_{\Omega}(x, y) = k_2^2 f(x, y) E(x, y)$, les sources de courants équivalentes induites dans le défaut ;
- $J_{\Omega_n}(x, y) = \frac{J_{\Omega}(x, y)}{E_0(x, y)}$, une quantité normalisée par rapport au champ incident E_0 .

On peut alors montrer que le champ d'anomalie mesuré sur des capteurs placés sur une ligne droite $y = y_0$ parallèle à la surface de l'interface Γ s'écrit :

$$E_s(x, y_0) = \iint J_{\Omega}(x', y') G(x - x', y_0 - y') dx' dy'$$

Supposant aussi que le champ incident est une onde plane se propageant dans la direction perpendiculaire de l'interface Γ , on peut écrire :

$$E_0(x, y) = E_0(y) = \begin{cases} \exp[+jk_1 y] + R \exp[-jk_1 y] & y > 0 \\ T \exp[+jk_2 y] & y < 0 \end{cases}$$

où R et T sont des coefficients de réflexion et de transmission.

Par la suite, pour être bref, remplaçant l'expression de $E_0(x, y)$ et l'expression de la fonction de GREEN dans l'équation du champ d'anomalie, on peut montrer la relation suivante :

$$\hat{E}_s(u, y_0) = \frac{2T}{\delta^2} \frac{\exp[-j\beta_1 y_0]}{\beta_1 + \beta_2} \iint_{\Omega} f(x', y') \exp[-j[ux' - (k_2 + \beta_2)y']] dx' dy',$$

ou d'une manière équivalente

$$\frac{\delta^2 (\beta_1 + \beta_2)}{2T \exp[-j\beta_1 y_0]} \hat{E}_s(u, y_0) = \iint_{\Omega} f(x', y') \exp[-j[ux' - (k_2 + \beta_2)y']] dx' dy'$$

où :

$$\hat{E}(u, y_0) = \text{TF1D} \{E_s(x, y_0)\} = \int E_s(x, y_0) \exp[-jux] dx$$

et

$$\begin{aligned} T &= \frac{2k_1}{k_1+k_2} \simeq \frac{2k_1}{k_2}, & \delta &= \sqrt{\frac{2}{\omega\mu\sigma}} \\ k_1^2 &= \omega^2\mu\epsilon, & k_2^2 &= \omega^2\mu\epsilon + j\omega\mu\sigma \simeq j\omega\mu\sigma, \quad k_2 = \frac{1}{\delta}(1+j) \\ \beta_1 &= \sqrt{k_1^2 - u^2}, & \beta_2 &= \sqrt{k_2^2 - u^2} \\ f(x,y) &= \frac{k_2^2 - k_1^2(x,y)}{k_2^2} \simeq -\frac{\sigma(x,y) - \sigma_0}{\sigma_0} \end{aligned}$$

Pour mieux interpréter ces équations, notons par

$$s = k_2 + \beta_2 = k_2 + \sqrt{k_2^2 - u^2}$$

et par

$$\hat{f}(u,s) = \text{TFL}\{f(x,y)\} = \iint_D f(x,y) \exp[-jux - sy] \, dx \, dy.$$

la transformée de FOURIER-LAPLACE (TFL) de $f(x,y)$. On a alors :

$$\hat{E}_s(u,y_0) = \frac{2T \exp[-j\beta_1 y_0]}{\delta^2 \beta_1 + \beta_2} \hat{f}(u,s)$$

Ainsi, les différentes mesures du champs $E_s(x,y_0)$ à des fréquences temporelles ω_i différentes, nous fournit des valeurs de la TFL de l'objet $\hat{f}(u,s)$ pour les différentes valeurs de u et de s . Le problème devient ainsi celui de la *synthèse de FOURIER-LAPLACE* qui consiste en la détermination de la fonction $f(x,y)$ à partir d'une connaissance partielle des valeurs de sa TFL.

Le travail avec la TFL n'étant pas facile, on peut envisager des approximations :

1. Négliger entièrement l'atténuation :

Il s'agit de négliger la partie imaginaire de k_2 et la partie réelle de s , c'est à dire :

$$\begin{aligned} k_2 &= \frac{1}{\delta}(1+j) \simeq \frac{1}{\delta} \\ s &= k_2 + \beta_2 = k_2 + \sqrt{k_2^2 - u^2} \simeq \frac{1}{\delta} + \sqrt{\frac{1}{\delta^2} - u^2} \simeq \frac{1}{\delta} \left(1 + \sqrt{1 - (\delta u)^2}\right) \end{aligned}$$

ce qui nous permet d'écrire :

$$\hat{E}_s(u,y_0) = \frac{2T \exp[-j\beta_1 y_0]}{\delta^2 \beta_1 + \beta_2} \hat{f}\left(u, v = -\frac{1}{\delta} \left(1 + \sqrt{1 - (\delta u)^2}\right)\right)$$

où $\hat{f}(u,v)$ est la TF2D de $f(x,y)$.

2. Négliger partiellement l'atténuation :

Il s'agit de négliger la partie imaginaire de $\beta_2 = \sqrt{k_2^2 - u^2}$, c'est à dire :

$$\begin{aligned} s \simeq k_2 + \text{Im}(\beta_2) &= k_2 + \text{Im}\left(\sqrt{k_2^2 - u^2}\right) \\ &= \frac{1}{\delta}(1+j) + \text{Im}\left(\sqrt{j\frac{2}{\delta^2} - u^2}\right) \\ &= \frac{1}{\delta} \left[1 + \frac{1}{\sqrt{2}} \sqrt{\sqrt{4 + (\delta u)^4} - (\delta u)^2}\right] + j\frac{1}{\delta} \end{aligned}$$

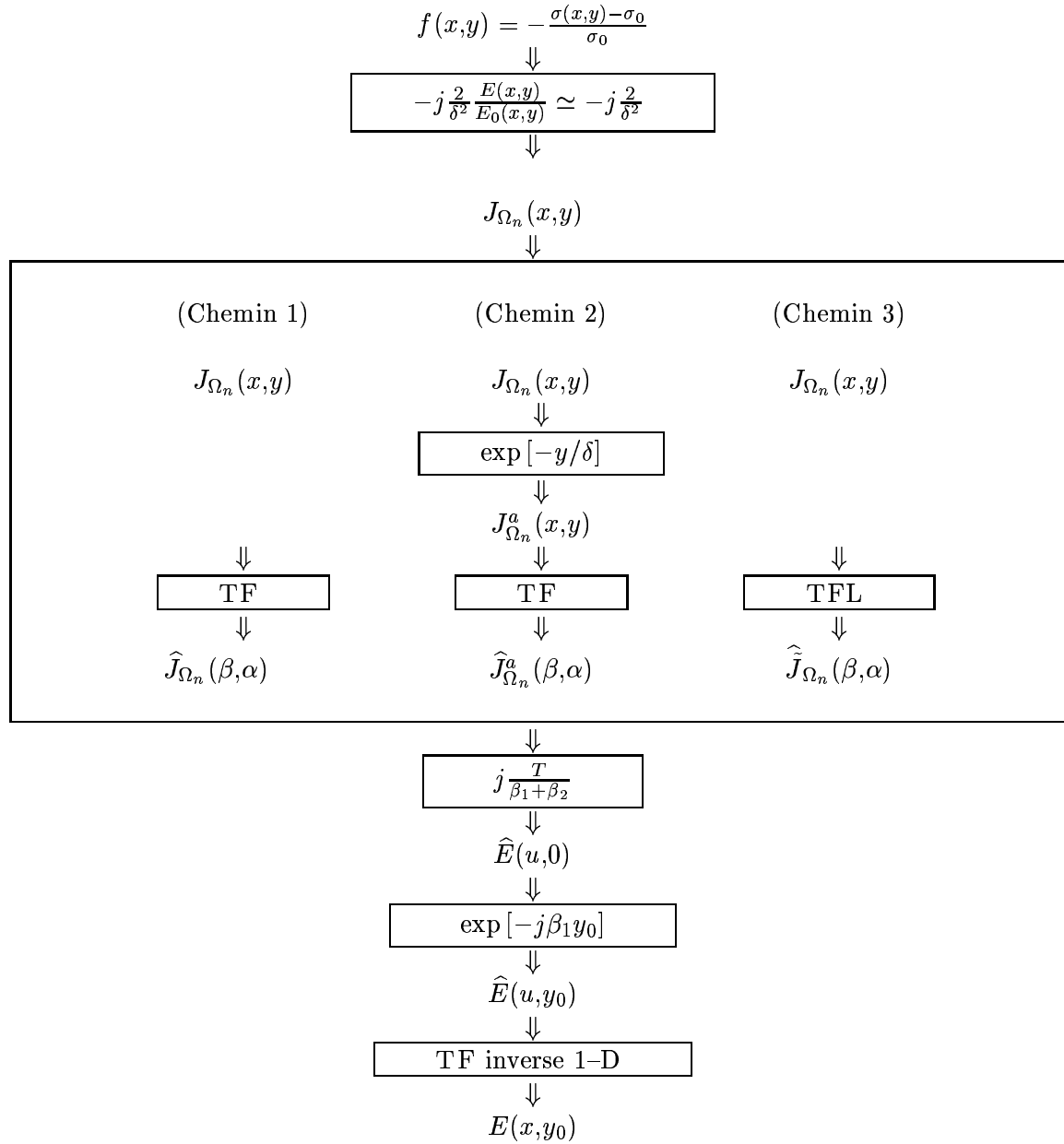
ce qui nous permet d'écrire :

$$\hat{E}_s(u, y_0) = \frac{2T}{\delta^2} \frac{\exp[-j\beta_1 y_0]}{\beta_1 + \beta_2} \hat{f}\left(u, v = -\frac{1}{\delta} \left(1 + \frac{1}{\sqrt{2}} \sqrt{\sqrt{4 + (\delta u)^4} - (\delta u)^2}\right)\right)$$

où $\hat{f}(u, v)$ est la TF2D de $f(x, y)$.

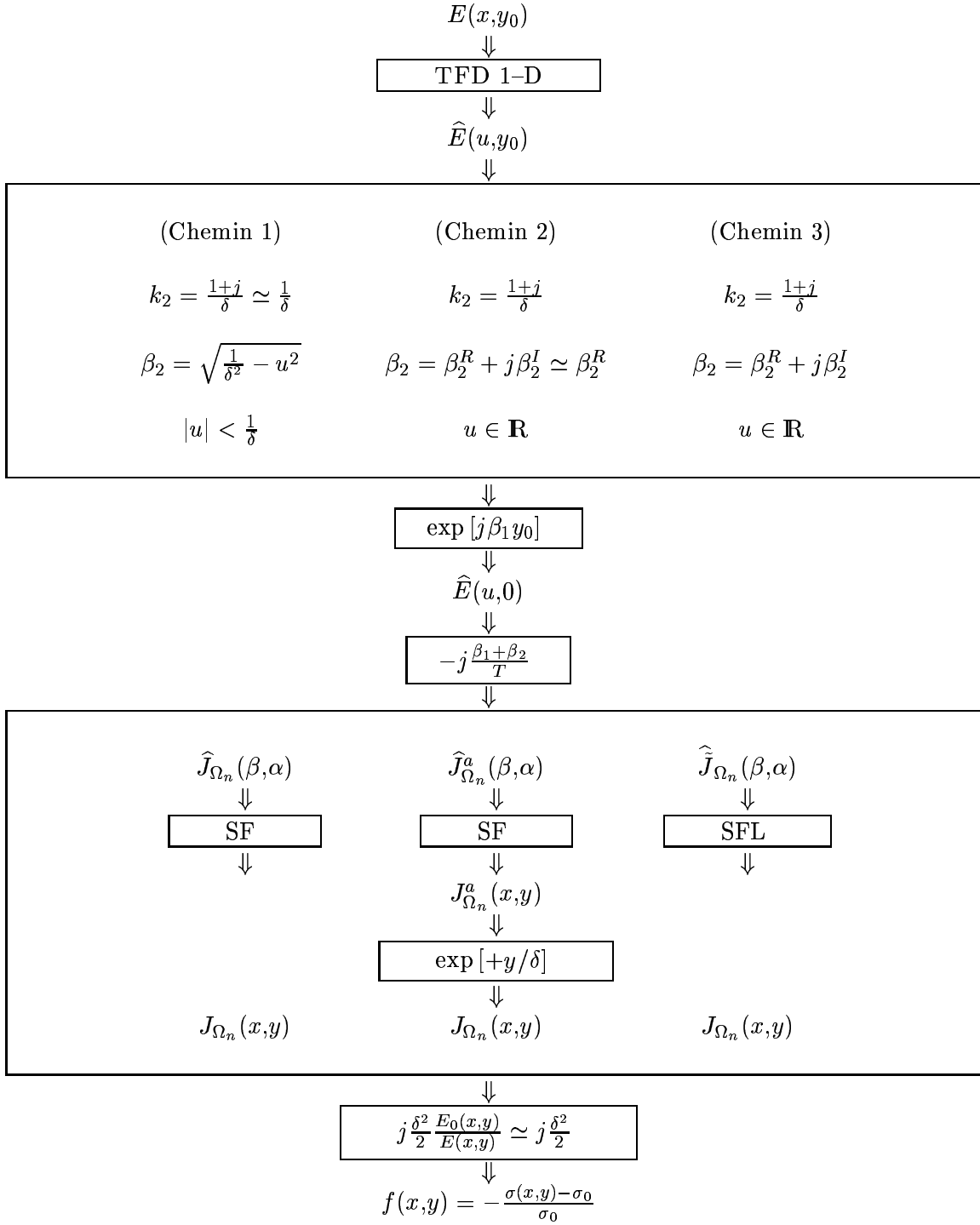
On peut résumer ces différentes situations dans le schéma bloc suivant:

Le schéma bloc de différentes approximations du problème direct
en imagerie par courants de FOUCAULT.



Lors de l'inversion on doit effectuer le trajet inverse. Ceci est résumé dans le schéma bloc suivant :

Le schéma bloc des différentes étapes de l'inversion
en imagerie par courant de FOUCAULT.



Ainsi les étapes qui nous permettent de simuler le problème direct de l'imagerie par courants de FOUCAULT sont :

- Calculer la source de courant normalisée

$$J_{\Omega_n}(x,y) = -j \frac{2}{\delta^2} \frac{E(x,y)}{E_0(x,y)} f(x,y)$$

- Suivant les trois chemins :

- calculer la TF de $J_{\Omega_n}(x,y)$ sur des contours algébriques définis par :

$$\begin{cases} u &= \alpha \\ v &= -\frac{1}{\delta} \left[1 + \sqrt{1 - (\delta\alpha)^2} \right] \end{cases} ,$$

- calculer la TF de $J_{\Omega_n}(x,y) \exp \left[-\frac{y}{\delta} \right]$ sur des contours algébriques définis par :

$$\begin{cases} u &= \alpha \\ v &= -\frac{1}{\delta} \left[1 + \frac{1}{\sqrt{2}} \sqrt{\sqrt{(\delta u)^4 + 4} - (\delta u)^2} \right] \end{cases} ,$$

- calculer la TFL de $J_{\Omega_n}(x,y)$ sur des contours algébriques définis par :

$$\begin{cases} u &= \alpha \\ v &= -j \left[k_2 + \sqrt{k_2^2 - \alpha^2} \right] \end{cases}$$

- Calculer $\hat{E}(u,0)$ à partir de l'une des quantités calculées dans l'étape précédente :
- Calculer $E(x,0)$ et finalement $E(x,y_0)$.

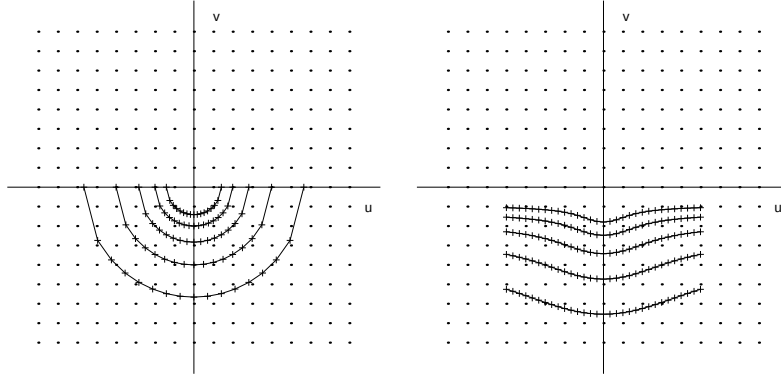


FIG. 2.14 – *Synthèse de FOURIER en imagerie par courants de FOUCAULT.*

2.7 Imagerie à émission de positons

En médecine nucléaire l'objectif de la tomographie d'émission est de déterminer la distribution d'un isotope radioactif d'une substance existante ou injectée dans le corps en mesurant le nombre de photons émis par ces substances.

Il existe deux types de tomographie à émission dépendant du type d'isotope utilisé:

- Tomographie à émission de positons (TEP): le radioélément est un émetteur de positons, qui après un parcours moyen très restreint va interagir avec la matière, de manière à émettre deux rayonnements gamma à 180 degrés l'un de l'autre. Le système de détection devra mesurer le nombre de photons qui coïncident, c'est à dire estimer la présence d'un positon par la détection simultanée des deux rayonnement gamma sur deux détecteurs opposés.
- Tomographie à émission d'un seul photon (SPECT): le radioélément est un émetteur gamma ordinaire, dit mono photonique, qui émet son rayonnement dans tous les sens. À la détection, l'ensemble des collimateurs mesurent dans une direction donnée le rayonnement émis. Ce type d'imagerie est appelé SPECT, de l'anglais (Single Photon Emitter Computed Tomographie).

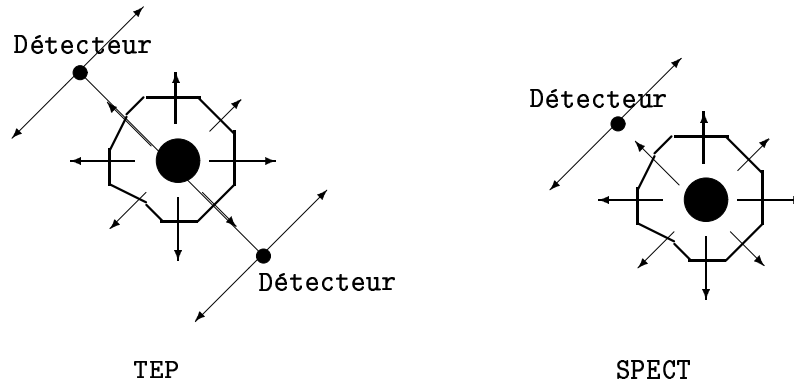


FIG. 2.15 – *Principe de la tomographie d'émission : TEP et SPECT.*

Dans la tomographie à émission de positons, le positon éjecté perd une très grande partie de son énergie dans une distance de quelques millimètres et forme un rayon gamma en rencontrant un électron se trouvant sur son chemin. Cette rencontre se produisant presque au repos, chaque rayon gamma a une énergie de 0.511 MEV et les photons émis repartent dans des directions opposées sur une ligne droite. Si un ensemble de détecteurs en anneau est placé autour de l'objet, deux détecteurs se trouvant face à face sur cette ligne, détecteront simultanément une énergie de 0.511 MEV. Ainsi, la direction des émissions est connue et les nombres de photons par unité de temps mesurés par chaque détecteur est une mesure de la concentration totale des substances radioactives le long de la ligne reliant les deux détecteurs.

Ainsi un modèle simple reliant la distribution de la densité de la matière radioactive $f(x,y)$ aux mesures est :

$$p(r,\phi) = \int f(x,y) dl$$

Bien entendu, dans des situations réelles, on ne peut pas négliger l'atténuation et un

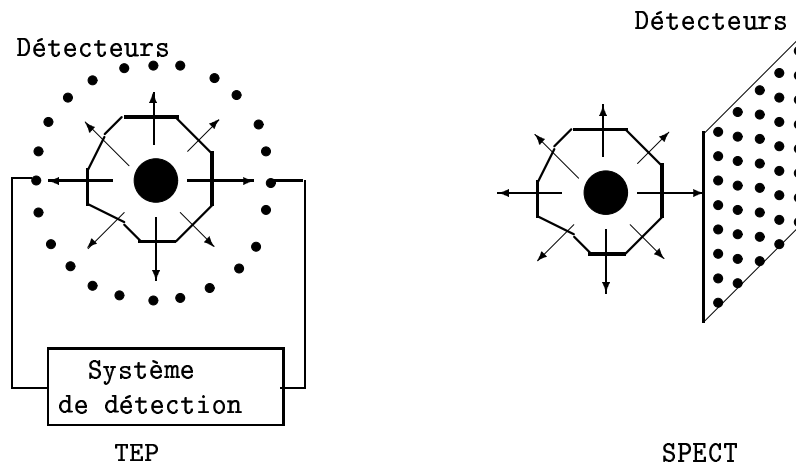


FIG. 2.16 – *Système de mesures en tomographie d'émission : TEP et SPECT.*

modèle plus approprié est :

$$p(r, \phi) = \int f(x, y) a(x, y) dl$$

où $a(x, y)$ est une fonction dépendant de la densité de la matière $f(x, y)$. Le problème est donc non linéaire. Une manière de traiter ce problème est, de supposer dans un premier temps $a(x, y)$ constant et estimer alors $f(x, y)$, puis déterminer $a(x, y)$ par un filtrage passe-bas de $f(x, y)$ et recommencer (linéarisations successives).

Une autre différence entre la tomographie à rayons X et la TEP est que dans la première les détecteurs utilisés sont des photomultiplicateurs qui, sous l'action du flux de photons importants que représente le rayonnement X de la source, fournissent sur un temps court une réponse intégrée, c'est à dire, un courant. Dans ces conditions, le rôle des fluctuations statistiques est faible et le rapport signal à bruit est élevé. Alors qu'en tomographie d'émission les détecteurs sont des compteurs à scintillation qui fournissent des impulsions individualisées liées à des flux de photons faibles, dont le taux est soumis à des fluctuations statistiques importantes.

2.8 Imagerie à résonance magnétique nucléaire RMN

L'énergie magnétique du moment nucléaire $\vec{\mu}$ dans un champ magnétique extérieur $\vec{H}$ est de la forme :

$$E = - \vec{\mu} \cdot \vec{H}$$

A l'équilibre thermique, l'aimantation macroscopique $\vec{\mu}$ d'un ensemble de N spins nucléaires identiques est parallèle au champ extérieur. Si $\vec{\mu}$ est écarté de son orientation d'équilibre (en le soumettant à une impulsion de radiofréquence) d'un angle θ , son évolution ultérieure consistera en une précession autour du $\vec{H}$ de fréquence $\omega_0 = \gamma |\vec{H}|$ où γ est le rapport gyromagnétique et ω_0 est la fréquence de Larmor. Du fait de la relaxation transversale l'amplitude du signal de précession libre décroît exponentiellement au cours du temps.

Le principe de base de l'imagerie est d'observer le signal de précession libre en présence d'un gradient statique du champ magnétique, superposé au champ statique homogène principal. Comme la fréquence de Larmor en chaque point est proportionnelle au champ en ce point, l'utilisation d'un gradient réalise un codage de la position par la fréquence.

Considérons un échantillon macroscopique entouré d'une bobine de résonance. Après une impulsion de $\theta = \pi/2$ modifiant la direction du champ magnétique d'un angle de $\pi/2$, le signal de précession libre que l'on observe provient de tous les spins de cette espèce dans l'échantillon. L'objectif de l'imagerie RMN est d'obtenir, séparément, la contribution à ce signal des divers points de l'échantillon.

A la suite d'une impulsion de $\theta = \pi/2$ à la fréquence ω , les aimantations transversales de chaque point de l'échantillon sont toutes parallèles à une direction ox du référentiel tournant à la fréquence ω . Appliquant, ensuite, un gradient $\mathbf{G}(t)$ dans le référentiel tournant à la fréquence ω , on mesure un signal de précession libre $s(t)$ qui nous fournit des renseignements sur la distribution spatiale de la densité de spin magnétique dans l'objet.

Le cas idéal le plus simple est celui d'un objet linéaire (la distribution de la densité de spin dans l'objet variant linéairement le long de l'axe) : le gradient étant connu, il suffit de calculer la transformée de FOURIER du signal de précession libre, dont l'amplitude à chaque fréquence donne la contribution de chaque point au signal.

Dans le cas plus général d'un objet à deux ou trois dimensions, on peut montrer qu'il y a un lien entre le signal temporel de précession libre $s(t)$ et la distribution de la densité de spin magnétique dans l'objet $f(\mathbf{r})$ qui est donnée par la relation suivante :

$$s(t) = \iiint f(\mathbf{r}) \exp \left[j \left(\int \mathbf{G}(t') \cdot \mathbf{r} dt' \right) \right] \exp \left[\frac{-t}{T} \right] dx dy dz$$

Le terme $\exp \left[\frac{-t}{T} \right]$ est due à l'effet de la relaxation transversale.

Cette relation peut être résumée par le schéma suivant :

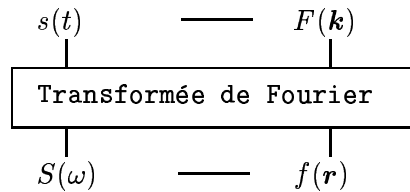


FIG. 2.17 – Principe de l'imagerie RMN.

où $S(\omega)$ est la transformée de FOURIER temporelle de $s(t)$ (une distribution de fréquences temporelles) et $F(\mathbf{k})$ est la transformée de FOURIER spatiale de $f(\mathbf{r})$.

En coordonnées cartésiennes on a :

$$s(t) = \iiint f(x,y,z) \exp [j(k_x(t)x + k_y(t)y + k_z(t)z)] \exp \left[\frac{-t}{T} \right] dx dy dz$$

où

$$k_x(t) = \gamma \int_0^t G_x(t') dt', \quad k_y(t) = \gamma \int_0^t G_y(t') dt', \quad k_z(t) = \gamma \int_0^t G_z(t') dt'$$

et où $G_x(t)$, $G_y(t)$ et $G_z(t)$ sont les composantes du gradient du champ magnétique au point (x,y,z) .

Ainsi, à un instant donné $s(t)$ est proportionnelle à la valeur de la TF spatiale 3D de $f(x,y,z)$ au point correspondant $\mathbf{k} = (k_x, k_y, k_z)$. Ainsi, en changeant les trois composantes $G_x(t)$, $G_y(t)$ et $G_z(t)$ du gradient du champ magnétique on peut remplir le domaine de FOURIER de l'objet.

Simplifions ces relations pour le cas d'un problème en 2D et négligeons pour l'instant le terme exponentiel. On a ainsi

$$s(t) = \iint f(x,y) \exp [j(k_x(t)x + k_y(t)y)] dx dy$$

$$F(u,v) = \iint f(x,y) \exp [j(ux + vy)] dx dy$$

$$s(t) = F(u,v), \forall u = k_x(t), \quad v = k_y(t)$$

L'organisation des séquences d'observation doit donc être agencée de façon à fournir une carte complète de l'espace réciproque de l'objet étudié. Deux méthodes sont actuellement d'usage.

Dans la première, on applique un gradient de champ constant et on répète l'expérience en changeant chaque fois son orientation. Le signal complexe de précession libre $s(t)$ à chaque instant d'échantillonnage t nous fournit donc deux points de l'espace réciproques de l'objet dans la direction du gradient. La carte de l'espace de FOURIER de l'objet ainsi obtenue est de la même forme que dans la tomographie à rayons X, mais ici la grandeur mesurée nous fournit directement ces renseignements.

On constate que ces points divergent le long des rayons. L'obtention d'une bonne densité des points aux grandes valeurs de $|\mathbf{k}|$ s'accompagnent d'une densité bien plus grande aux faibles valeurs de $|\mathbf{k}|$, autrement dit, l'obtention d'une bonne résolution à courte distance s'accompagne d'une sur-information à grande distance et vis versa.

Cette méthode est actuellement supplantée par une méthode dite de codage de phase dont le principe est le suivant : Après l'impulsion de $\pi/2$, on applique pendant un temps t un gradient G_y , sans effectuer d'observation; on annule, ensuite, le gradient G_y et on applique un gradient G_x . On observe alors le signal en fonction du temps t , durant un intervalle du temps T , suivant cette commutation du gradient. Le signal observé est:

$$s(t) = \iint f(x,y) \exp [jG_x x + G_y y] dx dy = F(-G_x t, -G_y t).$$

En modifiant la valeur du gradient G_y on remplit l'espace de FOURIER directement sur des coordonnées cartésiennes. Cependant à cause de la limitation de la valeur du gradient G_y , et l'amplitude décroissante de $s(t)$, due à l'effet de la relaxation, le remplissage reste incomplet.

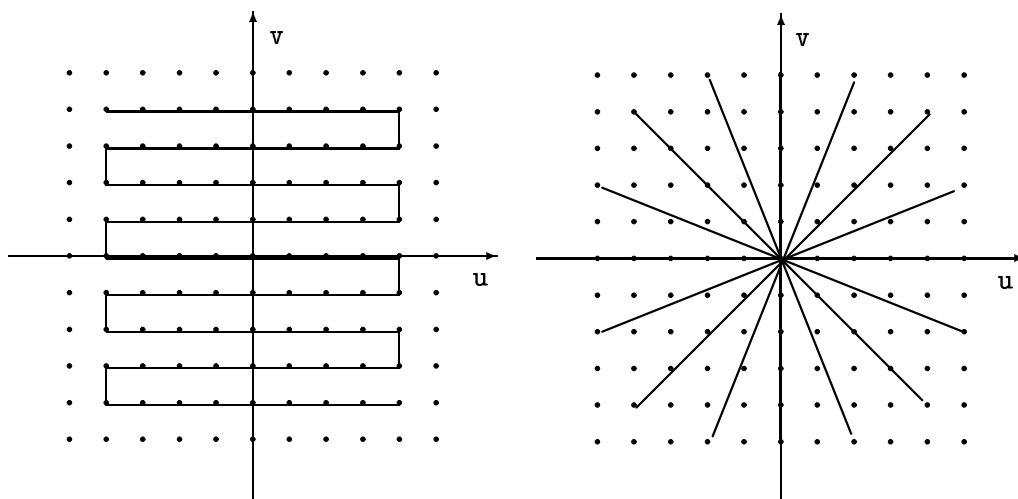


FIG. 2.18 – *Modes de remplissage du domaine de FOURIER en imagerie RMN.*

2.9 Imagerie radar à ouverture synthétique : SAR

L'imagerie radar à focalisation synthétique est basée sur le principe de l'effet DOPPLER. On peut montrer que le signal DOPPLER après détection synchrone et filtrage nous renseigne sur l'intersection avec une ligne droite de la TF spatiale en 2-D de l'objet illuminé par des ondes radar. Par exemple si l'on se place dans la géométrie de la figure (2.9), en supposant que h est petit, et si on modélise la réflectivité de l'objet par une fonction complexe $f(x,y) = |f(x,y)| \exp[j\phi(x,y)]$ et si le signal émis $s(t)$ est une impulsion modulée linéairement en fréquence de la forme

$$s(t) = \begin{cases} \exp[j(\omega_0 t + at^2)] & |t| \leq T/2, \\ 0 & \text{ailleurs} \end{cases}$$

alors le signal reçu est de la forme

$$r_\theta(t) = A \int_{-L}^L p_\theta(x') s\left(t - \frac{2(R+x')}{c}\right) dx'$$

où

$$p_\theta(x') = \int f(x' \cos \theta - y' \sin \theta, x' \sin \theta + y' \cos \theta) dy'.$$

En remplaçant cette dernière équation dans l'équation précédente et en notant $\tau_0 = \frac{2R}{c}$, le signal reçu, après la détection synchrone et filtrage passe-bas, est de la forme :

$$c_\theta(t) = \frac{A}{2} \int_{-L}^L p_\theta(x') \exp\left[-j\frac{2}{c}(\omega_0 + 2a(t - \tau_0))x'\right] dx'$$

Cette équation nous montre que le signal mesuré $c_\theta(t)$ est relié à la TF de la projection $p_\theta(x')$ par

$$c_\theta(t) = \frac{A}{2} P_\theta\left(\frac{2}{c}(\omega_0 + 2a(t - \tau_0))\right).$$

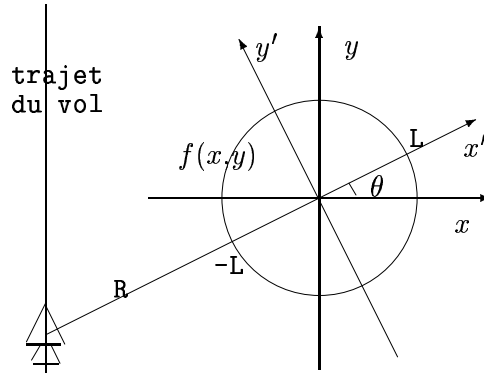


FIG. 2.19 – La géométrie de l'imagerie radar.

Ce signal est disponible pour $-\frac{T}{2} + \frac{2(R-L)}{c} \leq t \leq \frac{T}{2} + \frac{2(R+L)}{c}$, ce qui veut dire que l'on dispose pour chaque projection de données sur la TF spatiale de l'objet dans un intervalle $[k_{x_1}, k_{x_2}]$ donné par

$$k_{x_1} = \frac{2}{c}(\omega_0 - aT + \frac{4aL}{c}) \leq k_x \leq k_{x_2} = \frac{2}{c}(\omega_0 + aT - \frac{4aL}{c})$$

De plus l'angle θ est en général limité à $|\theta| \leq \theta_M$. Ainsi l'ensemble des signaux mesurés nous fournit des renseignements sur $F(k_x, k_y)$ sur un anneau comme le montre la figure 2.20. Dans des applications pratiques on a en général $\omega_0 aT \gg \frac{4aL}{c}$, ce qui veut dire que k_{x_1} et k_{x_2} sont proportionnels respectivement à la fréquence la plus petite et à la fréquence la plus grande dans le signal émis, et on a

$$k_{x_1} = \frac{2}{c}(\omega_0 - aT), \quad \text{et} \quad k_{x_2} = \frac{2}{c}(\omega_0 + aT)$$

Comme on peut le constater sur la figure 2, le remplissage du domaine de FOURIER de l'objet est très incomplet. De plus, les données se trouvent ici aussi sur des contours qui ne correspondent pas à un maillage cartésien.

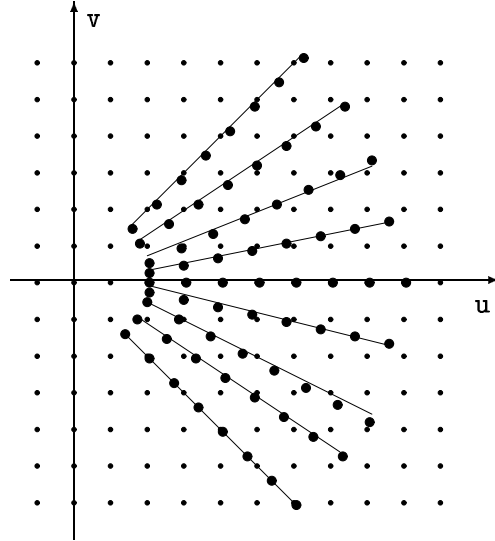


FIG. 2.20 – Remplissage du domaine de FOURIER en imagerie radar.

2.10 Imagerie en géophysique

L'équation différentielle de la propagation d'une onde dans un milieu non-homogène peut être transformée en une équation intégrale (équation de Lippmann-Schwinger) de la forme:

$$P_s(\mathbf{r}, \omega) = k^2 \int_V f(\mathbf{r}') P(\mathbf{r}', \omega) G_0(\mathbf{r}, \mathbf{r}', \omega) d\mathbf{r}'$$

dans laquelle P_s et P sont respectivement le champ total et le champ diffracté, G_0 est la fonction de GREEN, $f(\mathbf{r}) = n^2(\mathbf{r}) - 1$ avec $n(\mathbf{r}) = \frac{c}{v(\mathbf{r})}$ et $c = \frac{\omega}{k}$. Du fait que $P(\mathbf{r}', \omega)$ est fonction des inconnues $v(\mathbf{r})$ ou $n(\mathbf{r})$, on ne peut pas résoudre cette équation directement. Mais si on se place dans le cadre de l'approximation de BORN, c'est-à-dire un objet peu diffractant, on peut remplacer $P(\mathbf{r}', \omega)$ par l'onde incidente $P_0(\mathbf{r}', \omega)$.

Considérons maintenant le problème suivant : un objet (supposé en 2-D par simplicité) est illuminé par une onde plane et on mesure le champ diffracté pour toutes les fréquences sur une ligne contournant l'objet (figure 1). Dans ce cas, si on suppose que l'onde plane reste plane en traversant l'objet, et si on néglige les diffractions multiples, on a les relations suivantes :

$$P_s(\mathbf{r}', \omega) = k^2 S(\omega) \iint_S P_0(\mathbf{r}', \omega) G(\mathbf{r}, \mathbf{r}', \omega) d\mathbf{r}'$$

$$A(\mathbf{k}, \omega) = - \int_R [P_s(\mathbf{r}, \omega) \nabla G(\mathbf{r}) \nabla P_s(\mathbf{r}, \omega) G(\mathbf{r})] \mathbf{n} dl$$

où $S(\omega)$ est la TF temporelle de la source, $P_0(\mathbf{r}', \omega) = \exp[j\mathbf{k}_0 \cdot \mathbf{r}']$ est l'onde plane incidente, $A(\mathbf{k}, \omega)$ est la composante onde plane du champ diffracté. Et on montre facilement que l'on a

$$A(\mathbf{k}, \omega) = -k^2 S(\omega) \iint_S f(\mathbf{r}') \exp[-j(\mathbf{k} - \mathbf{k}_0) \cdot \mathbf{r}'] d\mathbf{r}'$$

Pour ω et $\mathbf{k}$ fixes, $A(\mathbf{k}, \omega)$ nous donne ainsi la valeur en un point de la TF spatiale de l'objet $F(\mathbf{k} - \mathbf{k}_0)$ et en faisant varier ω et $\mathbf{k}$ on remplit le domaine de FOURIER. Ce remplissage est cependant incomplet car la source n'a en général que des composantes en fréquences entre 0 et $\omega_{\max}$, et k varie pour $\alpha_1 \leq \theta \leq \alpha_2$ (figure 2).

D'autre part, le calcul de $A(\mathbf{k}, \omega)$ nécessite la connaissance de $P_s(\mathbf{r}, \omega)$ et de son gradient $\nabla P_s(\mathbf{r}, \omega)$. Mais si on considère le cas où les récepteurs se trouvent sur une ligne droite (figure 3) on a

$$A(\mathbf{k}, \omega) = 2\mathbf{n} \cdot \mathbf{k} \int_R jk P_s(x, \omega) \exp[-jk(x_0 + x)] dx$$

et on n'a plus besoin de $\nabla P_s(\mathbf{r}, \omega)$ pour le calcul de $A(\mathbf{k}, \omega)$. Ainsi en mesurant le champ diffracté $P_s(\mathbf{r}, \omega)$ pour $x_1 \leq x \leq x_2$ on calcule $A(\mathbf{k}, \omega)$ pour $0 \leq \omega \leq \omega_{\max}$ et k tel que $\alpha_1 \leq \theta \leq \alpha_2$ et on remplit ainsi le domaine de FOURIER sur des lignes droites.

Figure 1: Imagerie en géophysique.

Figure 2: Remplissage du domaine de FOURIER en géophysique.

A COMPLETER

2.11 Imagerie en radio astronomie

La synthèse d'ouverture en radio-astronomie consiste à reconstruire la distribution de brillance émise, à une certaine longueur d'onde, par un champ continu de sources célestes. Pour cela, on dispose d'un ensemble de capteurs, ou antennes (des radio-télescopes ou des télescopes optiques), permettant d'échantillonner spatialement ce champ et de recueillir des données. Ce problème est similaire à ceux rencontrés en traitement spatial (radar ou sonar).

L'ensemble des intercorrélations inter-capteurs, appelées aussi *visibilité complexes*, correspondent à des versions dégradées et bruitées de certains coefficients de FOURIER de la distribution de brillance recherchée (théorème de Van Cittert-Zernicke). Les visibilités sont souvent affectées d'une sévère distorsion de phase, liée aux conditions météorologiques et aux turbulences atmosphériques. De plus, la structure lacunaire du réseau ne permet pas d'accéder à toutes les données de FOURIER. Le problème comporte ainsi deux aspects : un problème de synthèse de FOURIER, et, simultanément, un problème de calibration de phase. Il s'agit là d'un problème inverse dont la résolution nécessite des méthodes de traitement spécifiques.

A COMPLETER

Chapitre 3

Cadre général et unification

Dans ce chapitre nous essayons de fournir un cadre général aux problèmes inverses et d'établir les principales notions que nous utiliserons par la suite.

3.1 Présentation du cadre général

Nous avons vu que, dans un grand nombre de problèmes réels en science expérimentale, nous sommes souvent amenés à déterminer une grandeur non observable $f(\mathbf{u})$ à partir d'un ensemble fini de *mesures* d'une grandeur observable $g(\mathbf{u})$ reliées par un *modèle* :

$$\mathcal{H}(g(\mathbf{u}), f(\mathbf{r})) = 0,$$

où $\mathcal{H}$ est un opérateur— linéaire ou non— décrivant la relation théorique entre $g(\mathbf{u})$ et $f(\mathbf{u})$. En général $\mathbf{r} \in D_1 \subset \mathbf{R}^k$ et $\mathbf{u} \in D_2 \subset \mathbf{R}^{k'}$, et k n'est pas forcément égale à k' .

Parfois la description du modèle peut contenir d'autres grandeurs intermédiaires, par exemple

$$\begin{cases} g(\mathbf{u}) &= \mathcal{H}_1(f(\mathbf{r}), z(\mathbf{r})) \\ 0 &= \mathcal{H}_2(f(\mathbf{r}), z(\mathbf{r})) \end{cases},$$

où $\mathcal{H}_1$ et $\mathcal{H}_2$ sont deux opérateurs (souvent $\mathcal{H}_1$ un opérateur bilinéaire en f et en z et $\mathcal{H}_2$ un opérateur non linéaire), et où $z(\mathbf{r})$ est une grandeur inconnue qui intervient dans le modèle mais qui n'est pas notre centre d'intérêt.

Parfois ce modèle peut être *explicite*:

$$g(\mathbf{u}) = \mathcal{H}(f(\mathbf{r}), z(\mathbf{r})),$$

et parfois, il est même possible d'éliminer les variables intermédiaires et d'obtenir une relation explicite encore plus simple

$$g(\mathbf{u}) = \mathcal{H}(f(\mathbf{r})).$$

La description des équations qui lient ces différentes grandeurs est appelée *modélisation*. Le modèle $\mathcal{H}$ est, en général, caractérisé par un certain nombre de paramètres $\boldsymbol{\theta}$. Connaissant $g(\mathbf{r})$ et $f(\mathbf{r})$, la détermination de $\boldsymbol{\theta}$ est ce que l'on appelle le problème de *l'identification* ou *l'estimation des paramètres* du modèle.

Connaissant le modèle $\mathcal{H}$ et l'objet $f(\mathbf{r})$, le calcul de $g(\mathbf{u})$ est ce que l'on appelle le *problème direct*. Inversement, le calcul de $f(\mathbf{r})$ connaissant $\mathcal{H}$ et $g(\mathbf{u})$ est ce que l'on appelle le *problème inverse*.

Autant, la résolution des problèmes directs, en général, ne pose pas de difficulté particulière, la résolution des autres problèmes est souvent plus délicate, car ceux-ci sont souvent des *problèmes mal posés*. En effet, la résolution de ces problèmes présente, entre autre, la caractéristique structurelle désagréable d'être très sensible aux erreurs sur les données. Or aucun dispositif expérimental n'est complètement affranchi d'une incertitude, dont l'origine la plus simple est la précision finie des mesures. De même, aucun modèle mathématique ne peut exactement refléter la réalité. Il est donc plus réaliste de considérer que l'objet recherché et les mesures sont reliés par une équation de la forme

$$\mathcal{H}(g(\mathbf{u}), f(\mathbf{r}), b(\mathbf{u})) = 0,$$

où $b(\mathbf{u})$ représente les erreurs, appelé communément *bruit*.

Lorsqu'on a un modèle explicite et simple, et qu'on peut faire l'hypothèse que les erreurs n'interviennent qu'à la sortie, on obtient :

$$g(\mathbf{u}) = \mathcal{H}(f(\mathbf{r})) \diamond b(\mathbf{u}).$$

Si de plus on peut supposer que les erreurs sont additives on a :

$$g(\mathbf{u}) = \mathcal{H}(f(\mathbf{r})) + b(\mathbf{u}),$$

Lorsque $\mathcal{H}$ est un opérateur linéaire cette relation sous sa forme générale s'écrit

$$g(\mathbf{u}) = \int f(\mathbf{r}) h(\mathbf{r}, \mathbf{u}) d\mathbf{r} + b(\mathbf{u}),$$

où $h(\mathbf{r}, \mathbf{u})$, le noyau de cette équation intégrale de première espèce, est appelé la fonction d'appareil. Alors, dans différents domaines d'applications on peut se trouver devant les situations suivantes :

- Lorsque les variables $\mathbf{r}$ et $\mathbf{u}$ sont homogènes et

$$h(\mathbf{r}, \mathbf{u}) = h(\mathbf{r} - \mathbf{u}),$$

i.e. lorsque le modèle est *invariant par translation* on a une équation de *convolution*:

$$g(\mathbf{u}) = \int f(\mathbf{r}) h(\mathbf{r} - \mathbf{u}) d\mathbf{r} + b(\mathbf{u}).$$

La fonction $h(\mathbf{r})$ est alors appelée la *réponse impulsionnelle* et le problème inverse correspondant *déconvolution* ou encore *restauration d'image*.

- $h(\mathbf{r}, \mathbf{u})$ peut être séparable :

$$h(\mathbf{r}, \mathbf{u}) = \prod_k h_k(r_k, u_k)$$

C'est, par exemple, le cas de la *synthèse de FOURIER*

$$g(\mathbf{u}) = \int f(\mathbf{r}) \exp[-j \langle \mathbf{r}, \mathbf{u} \rangle] d\mathbf{r},$$

où

$$h(\mathbf{r}, \mathbf{u}) = \exp[-j \langle \mathbf{r}, \mathbf{u} \rangle] = \exp\left[-j \sum_k r_k u_k\right].$$

- $h(\mathbf{r}, \mathbf{u})$ peut être à la fois séparable et invariant par translation :

$$h(\mathbf{r}, \mathbf{u}) = \prod_k h_k(r_k - u_k)$$

C'est, par exemple, le cas de la restauration d'image lorsque la réponse impulsionnelle est séparable.

- $h(\mathbf{r}, \mathbf{u})$ peut être de la forme :

$$h(\mathbf{r}, \mathbf{u}) = \delta(c(\mathbf{r}, \mathbf{u}))$$

où $\delta(\cdot)$ est la fonction DIRAC et $c(\mathbf{r}, \mathbf{u}) = 0$ décrit un sous-espace dans le domaine de la fonction $f(\mathbf{r})$. C'est par exemple le cas de la *reconstruction d'image* en tomographie X où on a $\mathbf{u} = (r, \phi)$, $\mathbf{r} = (x, y)$ et

$$g(r, \phi) = \iint f(x, y) \delta(r - x \cos \phi - y \sin \phi) dx dy$$

et où $r - x \cos \phi - y \sin \phi = 0$ décrit une ligne droite dans l'espace (x, y) du domaine de la fonction $f(x, y)$.

– $h(\mathbf{r}, \mathbf{u})$ peut être de la forme :

$$h(\mathbf{r}, \mathbf{u}) = \exp [-j < \mathbf{r}, \mathbf{c}(\mathbf{u}) >]$$

où $< \cdot, \cdot >$ représente un produit scalaire et $\mathbf{c}(\mathbf{u})$ un ensemble de fonctions telles que les équations $\mathbf{u} = \mathbf{c}(\mathbf{u})$ décrivent un sous-espace dans le domaine de FOURIER de la fonction $f(\mathbf{r})$. À titre d'exemple, prenons le cas où $\mathbf{u} = (\Omega, \phi)$, $\mathbf{r} = (x, y)$, $c_1(\mathbf{u}) = c_1(\Omega, \phi) = \Omega \cos \phi$, $c_2(\mathbf{u}) = c_2(\Omega, \phi) = \Omega \sin \phi$ et

$$g(\Omega, \phi) = \iint f(x, y) \exp [-j [\Omega \cos \phi x + \Omega \sin \phi y]] \, dx \, dy$$

et $u = \Omega \cos \phi$ et $v = \Omega \sin \phi$ décrivent une ligne droite dans l'espace (u, v) faisant un angle ϕ avec l'axe des u .

3.2 Exemples classiques de problèmes inverses linéaires

Parmi les problèmes inverses linéaires que l'on rencontre souvent en traitement du signal et de l'image, nous allons en citer un certain nombre qui sont les plus classiques.

3.2.1 Déconvolution de signaux 1-D

$$\begin{aligned} g(t) &= f(t) * h(t) + b(t) \\ &= \int h(\tau) f(t - \tau) d\tau + b(t) \\ &= \int f(\tau) h(t - \tau) d\tau + b(t) \end{aligned}$$

où $g(t)$ est le signal mesuré, $f(t)$ est le signal recherché, et $h(t)$ est la réponse impulsionnelle du système de mesure. Ici l'opérateur $\mathcal{H}$ est un opérateur de convolution.

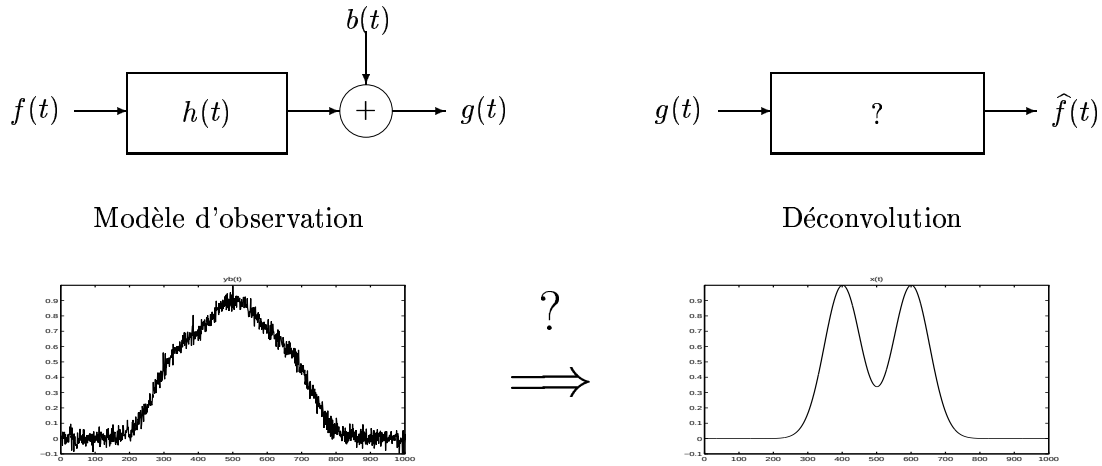


FIG. 3.1 – *Déconvolution de signaux.*

Sous forme discrète cette relation s'écrit :

$$g(m) = \sum_{k=-q}^p h(k) f(m - k) + b(m), \quad m = 0, \dots, M$$

ou encore

$$\mathbf{g} = \mathbf{H} \mathbf{f} + \mathbf{b}$$

où $\mathbf{g} = [g_0, \dots, g_M]$, $\mathbf{f} = [f_{-p}, \dots, f_{M+q}]$ et où $\mathbf{H}$ est une matrice de TOEPLITZ dont les éléments sont déterminés par les échantillons de sa réponse impulsionnelle $\mathbf{h} = [h_{-q}, \dots, 0, \dots, h_p]$.

Le caractère TOEPLITZ de la matrice $\mathbf{H}$ est due à la propriété d'invariance par translation de l'opérateur.

3.2.2 Débruitage ou filtrage

$$g(t) = f(t) + b(t),$$

où $g(t)$ est le signal mesuré, $f(t)$ est le signal recherché, et $b(t)$ est le bruit. Lorsque $g(t)$ est observée aux instants réguliers $t_i = (i - 1)\Delta T, i = 1, \dots, n$ on peut considérer ce problème comme un cas particulier de déconvolution où la réponse impulsionnelle $h(t) = \delta(t)$.

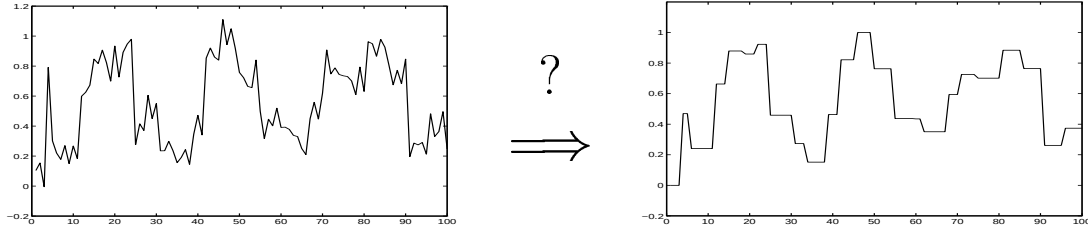


FIG. 3.2 – *Débruitage de signaux.*

Lorsque $g(t)$ est observée aux instants irréguliers t_i il s'agit de problèmes d'interpolation et d'extrapolation des signaux.

De même on peut considérer le débruitage ou le filtrage des images ::

$$g(x,y) = f(x,y) + b(x,y).$$

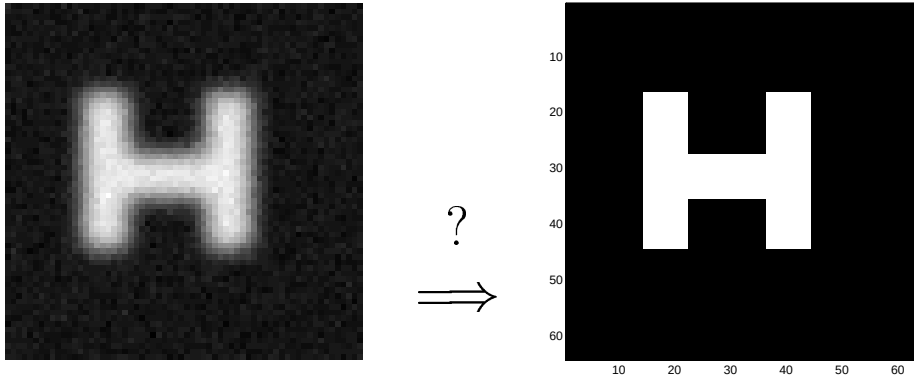


FIG. 3.3 – *Débruitage d'images.*

3.2.3 Restauration d'image

$$\begin{aligned}
 g(x,y) &= f(x,y) * h(x,y) + b(x,y) \\
 &= \iint_D f(x',y') h(x-x',y-y') dx' dy' + b(x,y) \\
 &= \iint_D h(x',y') f(x-x',y-y') dx' dy' + b(x,y),
 \end{aligned}$$

où $g(x,y)$ est l'image observée, $f(x,y)$ est l'image recherchée, et $h(x,y)$ est la réponse impulsionnelle du système d'imagerie. Ici $\mathbf{u} = (x,y)$, $\mathbf{r} = (x',y')$, $h(\mathbf{r},\mathbf{u}) = h(\mathbf{r} - \mathbf{u})$ et l'opérateur $\mathcal{H}$ est un opérateur de convolution bivariable.

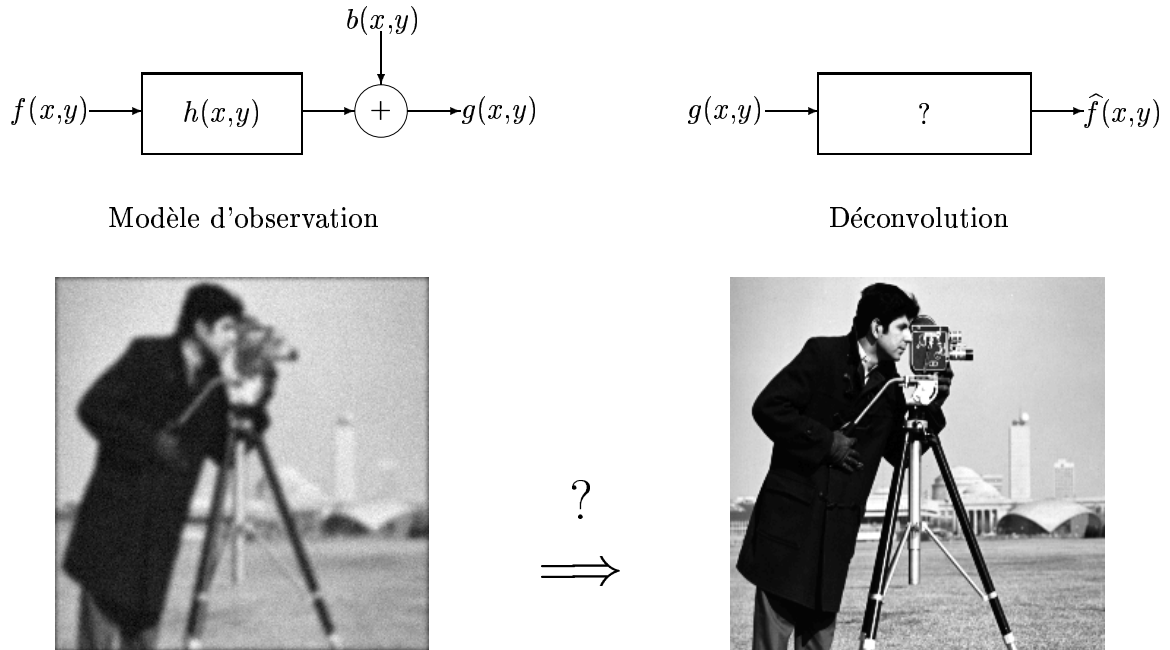


FIG. 3.4 – *Restauration d'image.*

Sous forme discrète cette relation s'écrit :

$$g(m,n) = \sum_{k=-p}^{p'} \sum_{l=-q}^{q'} h(k,l) f(m-k, n-l) + b(m,n)$$

ou encore

$$\mathbf{g} = \mathbf{H}\mathbf{f} + \mathbf{b}$$

où $\mathbf{g}$ et $\mathbf{f}$ sont deux vecteurs contenant les lignes (ou les colonnes, dans un ordre lexicographique) de l'image dégradée et de l'image initiale et $\mathbf{H}$ est une matrice qui a une structure de TOEPLITZ-BLOC-TOEPLITZ dont les éléments sont déterminés par les échantillons de la réponse impulsionnelle $h(k,l)$.

Le caractère TOEPLITZ-BLOC-TOEPLITZ de la matrice $\mathbf{H}$ est propre à la propriété de l'opérateur de convolution.

3.2.4 Reconstruction d'image en tomographie X

$$g(r, \phi) = \iint_D f(x, y) \delta(r - x \cos \phi - y \sin \phi) dx dy + b(r, \phi),$$

où $g(r, \phi)$ représente la projection suivant l'angle ϕ , et $f(x, y)$ est l'objet à reconstruire. Ici $\mathbf{u} = (r, \phi)$, $\mathbf{r} = (x, y)$, $h(\mathbf{r}, \mathbf{u}) = \delta(r - x \cos \phi - y \sin \phi)$ et l'opérateur $\mathcal{H}$ est l'opérateur de la transformée de RADON (TR).

Notons que, théoriquement, cette transformée admet une inversion qui est

$$f(x, y) = \frac{1}{2\pi^2} \int_0^\pi \int_0^\infty \frac{\frac{\partial p(r, \phi)}{\partial r}}{(r - x \cos \phi - y \sin \phi)} dr d\phi$$

mais, contrairement à ce qu'on peut imaginer, l'existence d'une expression analytique pour l'inversion de cette transformée ne réduit en rien les difficultés de son inversion dans des situations réelles. Nous reviendrons sur cette remarque dans les chapitres suivants.

La discrétisation de cette équation est plus délicate que celle de la restauration d'image. On peut cependant, imaginer facilement que cette relation sous forme discrète, puisse aussi s'écrire de la forme

$$\mathbf{g} = \mathbf{H}\mathbf{f} + \mathbf{b}$$

où $\mathbf{g}$ est un vecteur concaténant l'ensemble des projections $p(r_i, \phi_j)$ et $\mathbf{f}$ est un vecteur contenant les pixels de l'image. La matrice $\mathbf{H}$ dont les éléments sont entièrement déterminés par la géométrie de l'instrumentation (la discrétisation de l'objet et les positions de capteurs). Cette matrice n'a pas, en général, de structure particulière, cependant elle est sûrement très creuse. En effet, si on partage l'espace de (x, y) en pixels et si on fait l'hypothèse que $f(x, y)$ est constante à l'intérieure de chaque pixel, alors l'éléments H_{ij} correspond à la longueur du segment du rayons i traversant le pixel j .

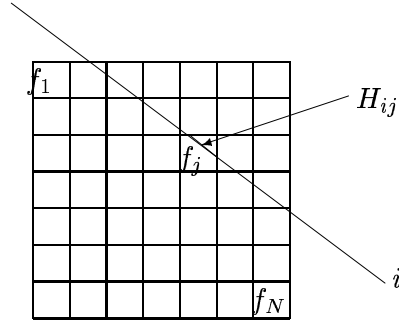


FIG. 3.5 – Discrétisation de la transformée de RADON.

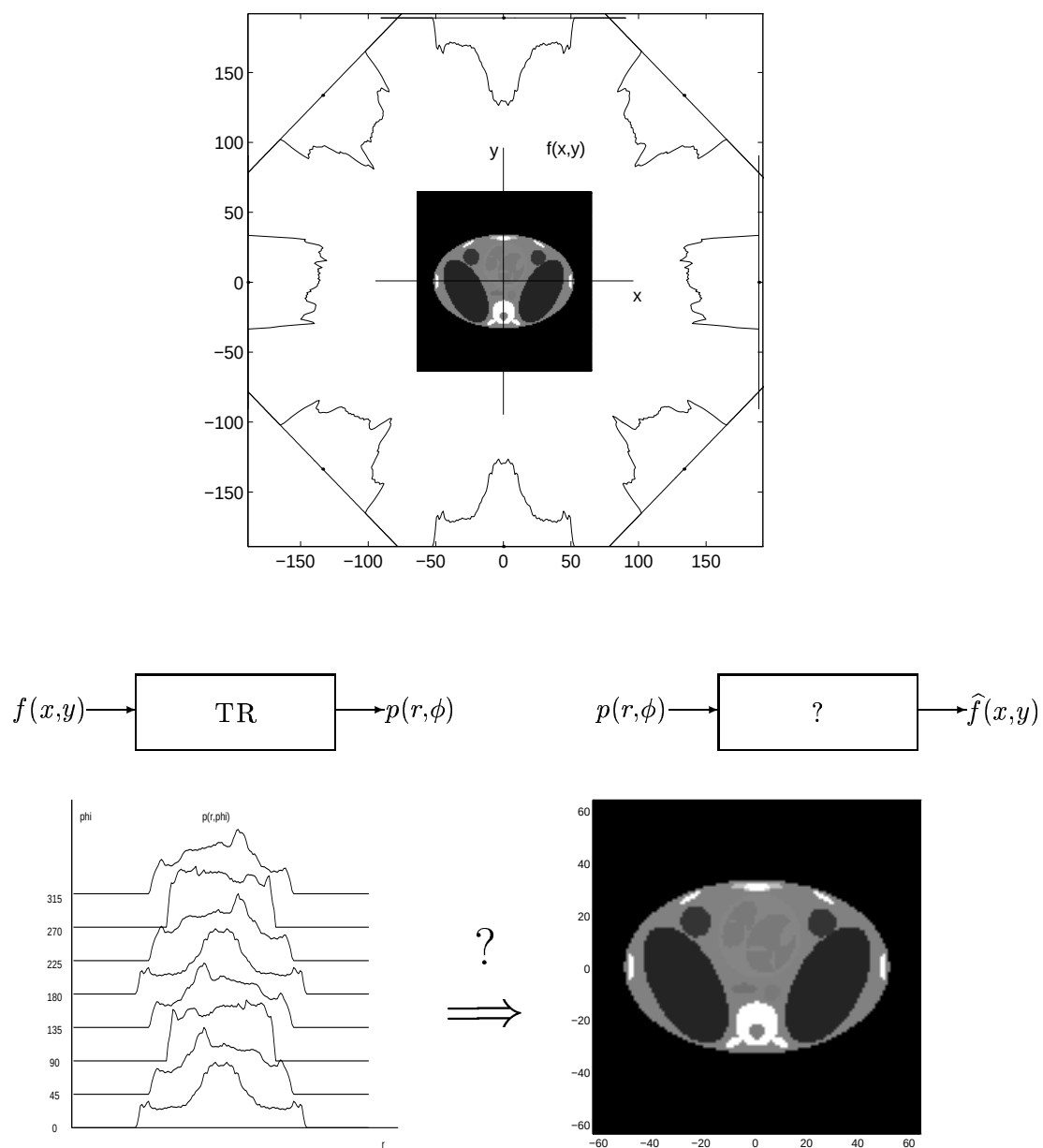


FIG. 3.6 – Reconstruction d'image en tomographie X.

3.2.5 Radiographie des objets cylindriques

Dans le problème de reconstruction d'image en tomographie X, lorsque l'objet $f(x,y)$ a une symétrie de révolution, la fonction ne dépend pas de l'angle ϕ et peut être exprimée seulement en fonction de ρ en coordonnées cylindriques: $f(x,y) = f(\rho)$. De même, ses projections $g(r,\phi)$ suivant les différents angles ϕ sont identiques et ne dépendent pas de ϕ . On peut alors établir une relation entre $f(\rho)$ et $g(r)$ qui est :

$$g(r) = \int_r^R f(\rho) \frac{\rho}{\sqrt{\rho^2 - r^2}} d\rho + b(r),$$

où $g(r)$ est l'intensité des rayons traversés (mesures), $f(\rho)$ est la fonction densité de la matière qui ne dépend que de la coordonnée radiale ρ et R est le rayon extérieur de l'objet. Remarquons que cette équation intégrale correspond à une transformée d'ABEL.

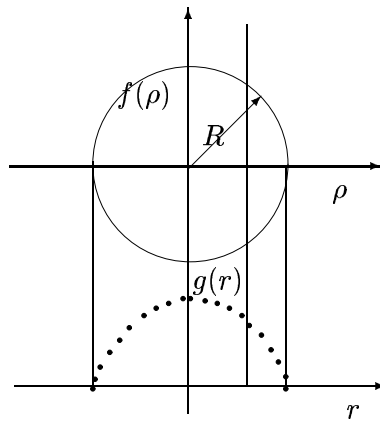
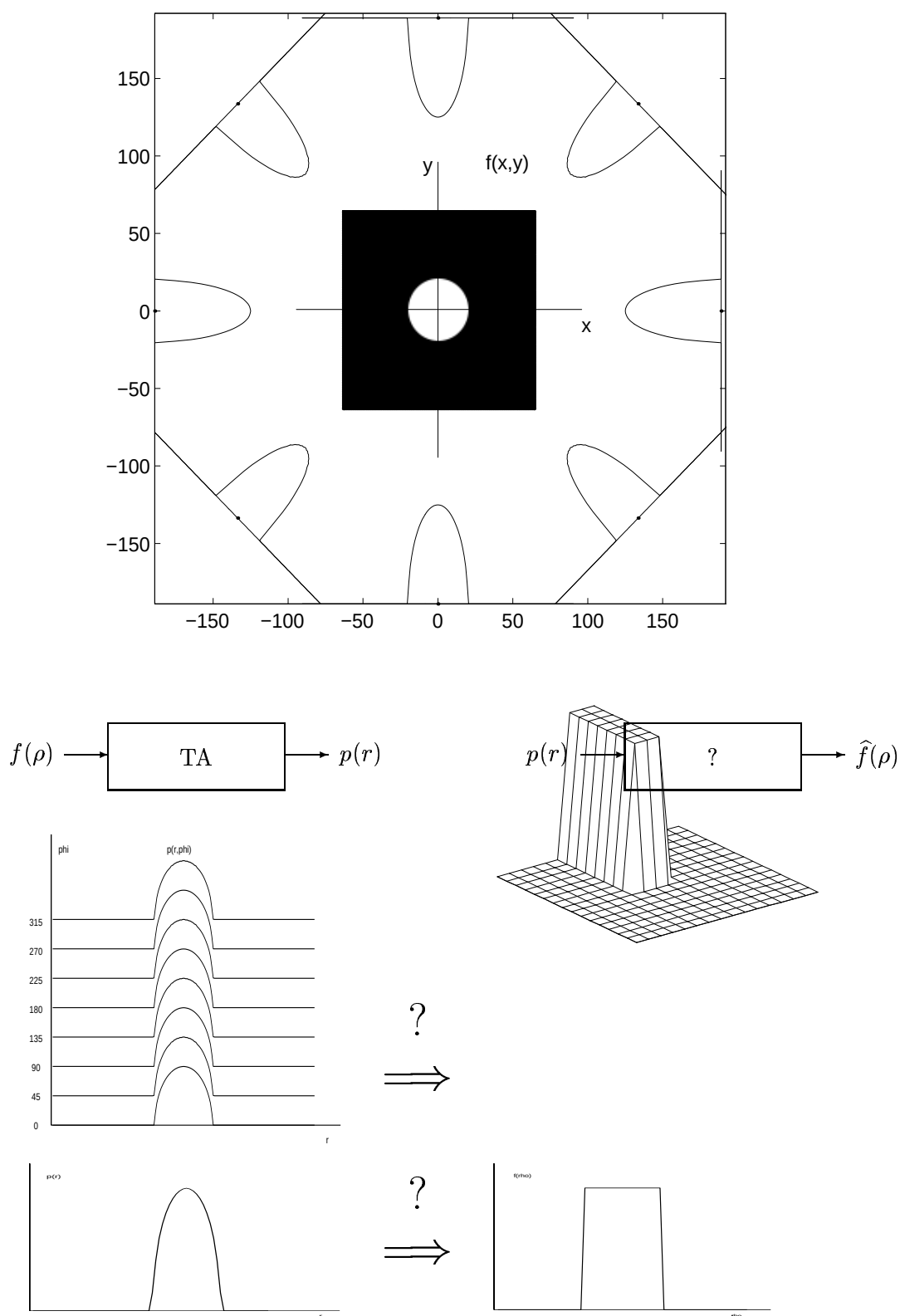


FIG. 3.7 – Radiographie des objets cylindriques.

FIG. 3.8 – *Reconstruction d'image en tomographie X pour des objets cylindriques.*

3.2.6 Synthèse d'ouverture en radio astronomie ou imagerie radar

$$g(u,v) = \iint_D f(x,y) \exp[-j(ux + vy)] dx dy + b(u,v),$$

où $f(x,y)$ est l'image à reconstruire, et $g(u,v)$ représente, soit une mesure de la corrélation entre les signaux recueillis par deux antennes en synthèse d'ouverture en radio astronomie, soit la transformée de FOURIER (TF) du signal d'écho en imagerie radar. Dans le cas de la radio astronomie, $f(x,y)$ représente la brillance du ciel tandis que dans le cas de l'imagerie radar elle représente la variation de l'indice de réfraction de l'objet (la cible) par rapport à son environnement.

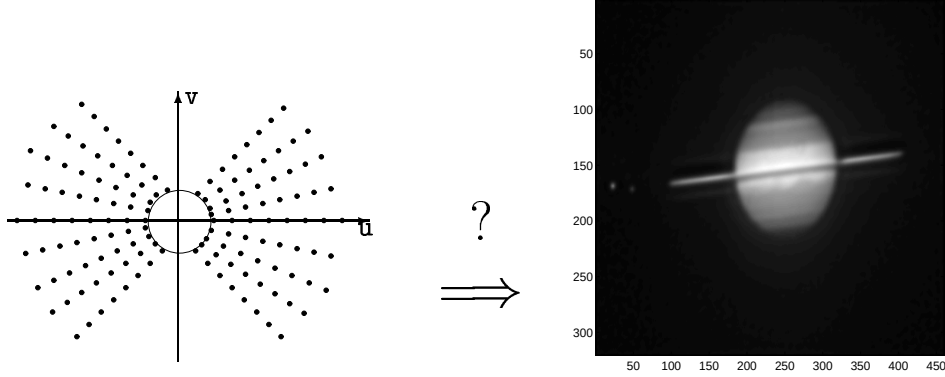


FIG. 3.9 – *Synthèse d'ouverture en radioastronomie.*

Ici $\mathbf{u} = (u,v)$, $\mathbf{r} = (x,y)$, $h(\mathbf{r},\mathbf{u}) = \exp[-j \langle \mathbf{r}, \mathbf{u} \rangle]$ et l'opérateur $\mathcal{H}$ est l'opérateur de la transformée de FOURIER bivariable.

Notons qu'ici aussi, théoriquement, la TF inverse existe et on a :

$$f(x,y) = \iint g(u,v) \exp[+j(ux + vy)] du dv,$$

mais, là encore, l'existence de cette formule ne réduit en rien la difficulté du problème inverse de la synthèse de FOURIER. En effet, si $g(u,v)$ était connue précisément et pour tout $(u,v) \in \mathbb{R}^2$, alors $f(x,y)$ pourrait être calculée à partir de cette intégrale. Mais, en pratique, $g(u,v)$ peut être mesurée sur un nombre fini de points (u_i, v_i) non nécessairement distribués d'une manière uniforme dans l'espace (u,v) , et c'est exactement là la difficulté de l'inversion.

3.2.7 Synthèse de Fourier en reconstruction tomographique d'image

$$g(\Omega, \phi) = \iint_D f(x, y) \exp[-j(T_1(\Omega, \phi)x + T_2(\Omega, \phi)y)] dx dy + b(\Omega, \phi)$$

où

$$\begin{cases} u &= T_1(\Omega, \phi) \\ v &= T_2(\Omega, \phi) \end{cases}$$

définissent un ensemble de contours algébriques dans l'espace de FOURIER (u, v) .

Comme nous l'avons vu dans le chapitre précédent, la synthèse de FOURIER se trouve au cœur d'un grand nombre d'applications en imagerie. Parmi celles-ci, on trouve :

- Tomographie à rayons X :

$$\begin{pmatrix} u = T_1(\Omega, \phi) \\ v = T_2(\Omega, \phi) \end{pmatrix} = \begin{pmatrix} \cos \phi & -\sin \phi \\ \sin \phi & \cos \phi \end{pmatrix} \begin{pmatrix} \Omega \\ 0 \end{pmatrix},$$

où les contours algébriques sont des lignes droites concentriques, et où $g(\Omega, \phi_i)$ représente la TF spatiale de $p(r, \phi_i)$ par rapport à la variable r .

- Imagerie par résonance magnétique nucléaire (RMN) :

Cette formulation du problème se trouve aussi dans le cas de l'imagerie RMN. Dans ce cas $g(\Omega, \phi_i) = s(t, \phi_i)$, où $s(t, \phi_i)$ est directement le signal de précession libre mesuré quand le gradient du champ magnétique est dans la direction ϕ_i .

- Tomographie à ondes diffractées :

$$\begin{pmatrix} u = T_1(\Omega, \phi) \\ v = T_2(\Omega, \phi) \end{pmatrix} = \begin{pmatrix} \cos \phi & -\sin \phi \\ \sin \phi & \cos \phi \end{pmatrix} \begin{pmatrix} -k_0 + \sqrt{k_0^2 - \Omega^2} \\ \Omega \end{pmatrix},$$

où k_0 est la constante de propagation des ondes, et où les $g(\Omega, \phi_i)$ représentent les TF spatiales des champ diffractés mesurés. Ici, les contours algébriques sont des demi-cercles concentriques.

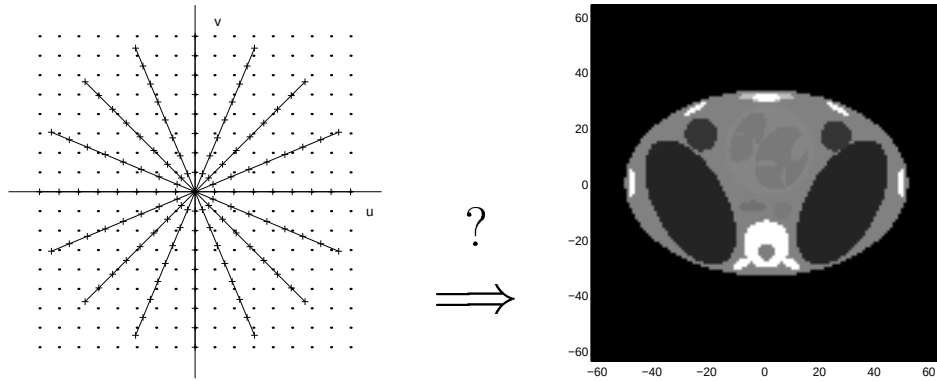
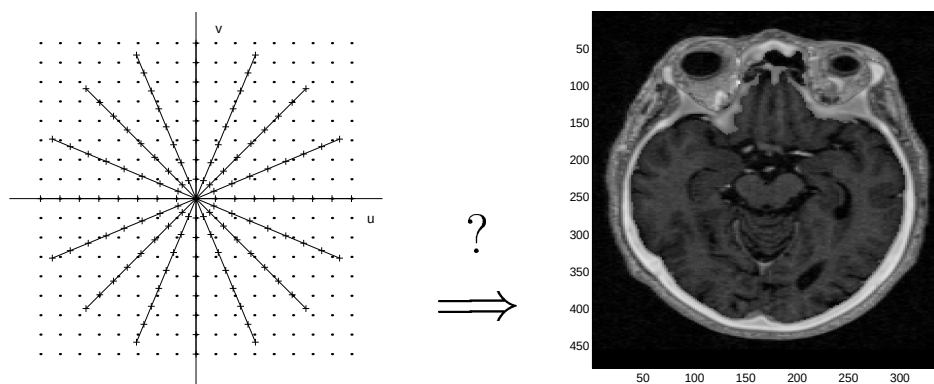
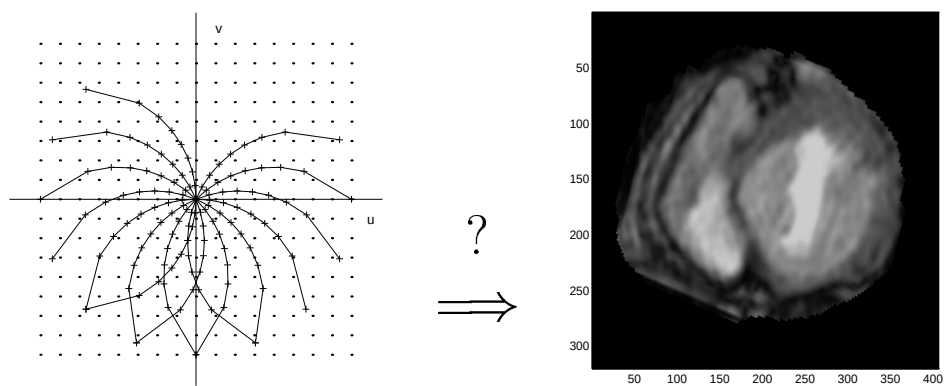


FIG. 3.10 – Synthèse de FOURIER en tomographie X.

FIG. 3.11 – *Synthèse de FOURIER en imagerie RMN.*FIG. 3.12 – *Synthèse de FOURIER en tomographie par ondes diffractées.*

3.2.8 Synthèse de Fourier-Laplace

$$g(u,s) = \iint_D f(x,y) \exp[-j(ux + sy)] \, dx \, dy + b(u,s),$$

où u est une variable réelle, tandis que $s = \gamma + jv$ est une variable complexe.

Cette équation intervient en tomographie de diffraction dans des objets dissipatifs ou bien en imagerie tomographique par courants de FOUCAULT en contrôle non destructif.

(u,v) représente un point dans le domaine de FOURIER spatial de l'objet $f(x,y)$. Ceci correspond au phénomène de propagation dans les deux directions, tandis que γ représente un phénomène de diffusion ou d'atténuation.

3.3 Exemples de problèmes inverses non linéaires

Parmi les problèmes inverses non linéaires classiques nous allons citer deux exemples.

3.3.1 Imagerie par ondes diffractées

Lorsqu'un objet $f(\mathbf{r})$ est illuminé par une onde incidente $\phi_0(\mathbf{r})$ le champ total $\phi(\mathbf{r})$ peut être modélisé par la somme du champ incident et un champ diffracté qui dépend à la fois de l'objet et du champ incident :

$$\phi(\mathbf{r}) = \phi_0(\mathbf{r}) + \int_D G_o(\mathbf{r}, \mathbf{r}') \phi(\mathbf{r}') f(\mathbf{r}') d\mathbf{r}'$$

et le champ diffracté à l'extérieur de l'objet $g(\mathbf{u})$ que l'on mesure est une fonction de l'objet et du champ total à l'intérieur de l'objet :

$$g(\mathbf{u}) = \int_D G_m(\mathbf{r}, \mathbf{u}) \phi(\mathbf{r}) f(\mathbf{r}) d\mathbf{r} + b(\mathbf{u}).$$

Le problème de l'imagerie dans ce contexte est de fournir une estimation de $f(\mathbf{r})$ à partir d'un nombre fini de mesures $g(\mathbf{u}_i), i = 1, \dots, M$.

Le cas de l'approximation de BORN étant la situation de faibles perturbations où on peut remplacer $\phi(\mathbf{r})$ par $\phi_0(\mathbf{r})$ dans cette dernière équation. Nous avons déjà présentés et détaillés ces équations dans le cas linéaire. Ici, l'objectif est de considérer le problème non linéaire.

Travailler dans le cadre fonctionnel ne permet pas de voir facilement les difficultés du problème. C'est pourquoi, nous allons discrétiser ces équations, en utilisant par exemple la méthode des moments. Il n'est alors pas difficile de voir que la forme discrétisée de ces deux équations devient :

$$\begin{aligned} \phi &= \phi_0 + \mathbf{G}_o \mathbf{F} \phi, \\ g &= \mathbf{G}_m \mathbf{F} \phi + b, \end{aligned}$$

où

- ϕ_0 et ϕ sont deux vecteurs représentant respectivement le champ incident et le champ total,
- $\mathbf{G}_m$ et $\mathbf{G}_o$ sont deux matrices reliées aux fonctions de GREEN,
- $\mathbf{F}$ est une matrice diagonale dont les éléments diagonaux sont formés par les éléments du vecteur $\mathbf{f}$ qui représente l'objet (les pixels de l'image à reconstruire concatenés ligne par ligne ou colonne par colonne) et,
- $\mathbf{g}$ est un vecteur contenant les mesures et $\mathbf{b}$ un vecteur représentant les erreurs de mesures.

Écrite sous cette forme, pour mieux illustrer le caractère non linéaire du problème, il est maintenant possible d'expliciter la relation liant les mesures $\mathbf{g}$ aux inconnus $\mathbf{f}$, en éliminant les variables intermédiaires ϕ . En effet, calculant ϕ à l'aide de la première équation et la remplaçant dans la deuxième on obtient :

$$\mathbf{g} = \mathbf{G}_m \mathbf{F} (\mathbf{I} - \mathbf{G}_o \mathbf{F})^{-1} \phi_0 + \mathbf{b} = \mathbf{H}(\mathbf{f}) + \mathbf{b},$$

ce qui permet de montrer d'une manière plus explicite le caractère non linéaire de l'opérateur $\mathbf{H}$ liant l'objet aux mesures.

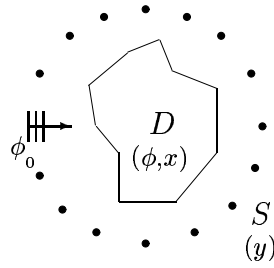


FIG. 3.13 – Géométrie de l'imagerie par ondes diffractées.

Evidemment, il n'existe pas de relation explicite pour l'inversion de ces équations et l'obtention d'une solution ne peut se faire que d'une manière itérative, ceci quelque soit le cadre choisi qu'il soit en dimension infinie ou finie.

Dans ce problème non linéaire aussi, la résolution du problème directe, c'est à dire étant donnée $f(\mathbf{r})$ calculer $g(\mathbf{r})$, ne pose pas de grandes difficultés. En effet, on peut montrer qu'avec un choix judicieux pour les fonctions de bases et la discrétisation de ces équations, on peut calculer $g(\mathbf{r})$ aussi précisément que l'on souhaite si on connaît l'objet $f(\mathbf{r})$. Par contre, la résolution du problème inverse, c'est à dire le calcul de $f(\mathbf{r})$ à partir de $g(\mathbf{r})$ est plus délicat. En effet, même lorsqu'on suppose pouvoir mesurer $g(\mathbf{r})$ partout autour de l'objet avec une précision infinie, l'existence et l'unicité de la solution ne sont pas garanties pour n'importe quel objet $f(\mathbf{r})$. De plus, même lorsqu'on impose à l'objet $f(\mathbf{r})$ dans un ensemble de fonctions telles que ces deux propriétés soient garanties, la solution du problème reste très sensible aux erreurs de mesures.

A COMPLETER

3.3.2 Tomographie d'impédance

$$p(\mathbf{u}_i) = \int_V \sigma(\mathbf{r}) |\nabla \phi_i(\mathbf{r})|^2 d\mathbf{r} + b_i$$

où $p(\mathbf{u}_i)$ est la puissance dissipée mesurée entre le point de coordonnée $(\mathbf{u}_i)$ et un point de référence, $\sigma(\mathbf{r})$ est la distribution des conductivité électriques dans l'objet, $\phi_i(\mathbf{r})$ est la distribution du potentiel électrique dans l'objet. Notons que $\phi_i(\mathbf{r})$ dépend du $\sigma(\mathbf{r})$, ce qui rend cette équation non linéaire.

A COMPLETER

Chapitre 4

Problèmes mal-posés

Pour mieux comprendre les difficultés de la résolution des problèmes inverses, un cadre classique est l'étude des problèmes inverses linéaires dans des espaces de Hilbert où l'on définit deux ensembles F et G et un opérateur $H : F \mapsto G$ tel que :

$$Hf = g \quad f \in F \quad g \in G \quad \text{espaces de Hilbert}$$

Souvent l'opérateur H est borné et compact et nous verrons que c'est exactement ces caractères qui sont à l'origine des principales difficultés de la résolution des problèmes inverses.

Ce cadre topologique et abstrait nous permettra de définir les notions de problèmes bien posés et mal posés. C'est pourquoi, avant de donner ces définitions un rappel sur la terminologie de l'analyse fonctionnelle est nécessaire. Après ce rappel les notions de problèmes *bien posés* et *mal posés* seront définies et un certain nombre d'exemples détailleront ces notions.

Finalement, une analyse plus détaillée du problème de la déconvolution nous permettra de comprendre d'une manière plus concrète le caractère mal posé du problème et les difficultés de sa résolution.

4.1 Rappels d'analyse fonctionnelle

Considérons deux ensembles F et G et un opérateur $H : F \mapsto G$.

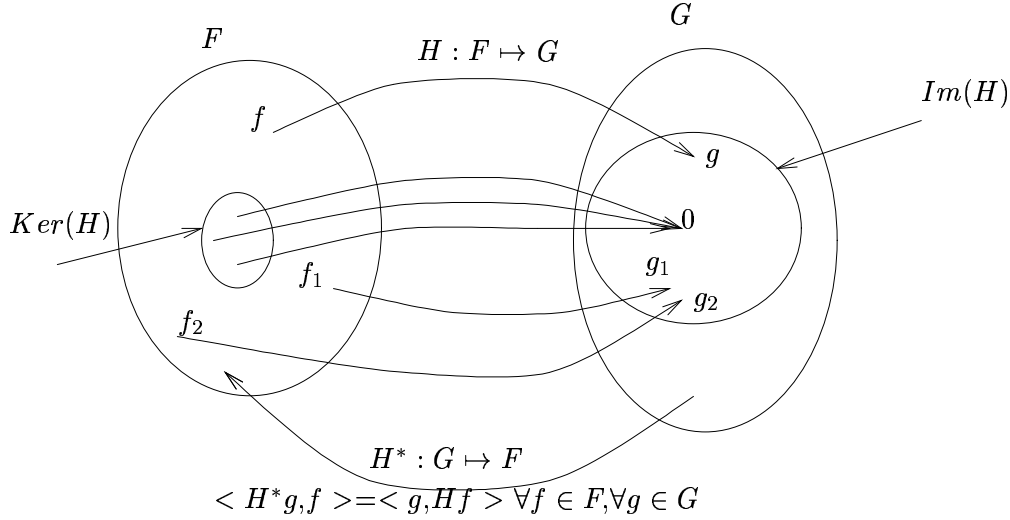


FIG. 4.1 – Analyse fonctionnelle et topologie

- Topologie sur un ensemble F :
Un système de sous-ensembles de F , dits **ouverts**, tels que la réunion d'ensembles ouverts est un ouvert, et l'intersection d'un nombre fini d'ensembles ouverts est un ouvert.
- $F_1 \subset F$ est dit **fermé** si son complémentaire est un ouvert.
- Pour un ouvert M , on définit :
 - sa fermeture : $\overline{M}$ intersection de tous les fermés qui contiennent M
 - sa frontière : $\overline{M} - M$
- $\mathcal{I}m(\mathbf{H}) = \bigcup \{\mathbf{H}f\} \subset G$ est l'image de $\mathbf{H}$
- $\mathcal{K}er(\mathbf{H}) = \{f, \mathbf{H}f = 0\} \subset F$ est le noyau de $\mathbf{H}$
- $\mathbf{H}^* : G \mapsto F$ est l'opérateur adjoint de $\mathbf{H}$

$$\langle \mathbf{H}^*g, f \rangle = \langle g, \mathbf{H}f \rangle \quad \forall f \in F, \forall g \in G$$

- $\mathcal{K}er(\mathbf{H}) = \mathcal{I}m(\mathbf{H}^*)^\perp$
- $\mathcal{K}er(\mathbf{H}^*) = \mathcal{I}m(\mathbf{H})^\perp$
- $\mathcal{K}er(\mathbf{H})^\perp = \overline{\mathcal{I}m(\mathbf{H}^*)}$
- $\mathcal{K}er(\mathbf{H}^*)^\perp = \overline{\mathcal{I}m(\mathbf{H})}$

- $\mathbf{H}$ est **borné** s'il existe un réel $c \geq 0$ tel que $\|\mathbf{H}\mathbf{f}\| \leq c\|\mathbf{f}\|$, $\forall \mathbf{f} \in F$
- Un ensemble $B \subset \mathbf{R}^n$ est dit **borné** s'il est entièrement situé dans une boule ouverte dont le centre est à l'origine et le rayon est fini.
- Un ensemble fermé et borné est dit **compact**.
- Un opérateur $\mathbf{H}$ est dit **compact** s'il applique tout ensemble borné $B \subset F$ dans un ensemble $A(B)$ dont la fermeture $\overline{A(B)}$ est compacte.

Exemple 1:

$$\begin{pmatrix} g_1 \\ g_2 \end{pmatrix} = \begin{pmatrix} 1 & 2 \\ 3 & 6 \end{pmatrix} \begin{pmatrix} f_1 \\ f_2 \end{pmatrix}$$

$$\mathbf{H} = \begin{pmatrix} 1 & 2 \\ 3 & 6 \end{pmatrix}, \quad \mathbf{H}^* = \begin{pmatrix} 1 & 3 \\ 2 & 6 \end{pmatrix}$$

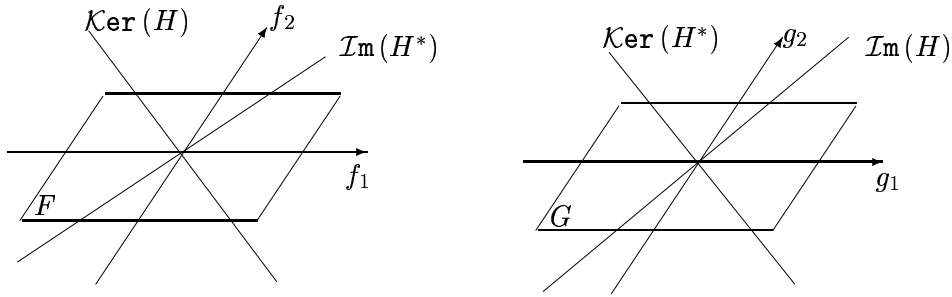


FIG. 4.2 – Image et noyau d'un opérateur et de son adjoint : $F = G = \mathbf{R}^2$

$$\begin{aligned} \text{Ker}(\mathbf{H}) &= \{(f_1, f_2) : f_1 = -2f_2\} \\ \text{Im}(\mathbf{H}^*) &= \left\{ (f_1, f_2) : f_1 = \frac{1}{2}f_2 \right\} \\ \text{Ker}(\mathbf{H}^*) &= \{(g_1, g_2) : g_1 = -3g_2\} \\ \text{Im}(\mathbf{H}) &= \left\{ (g_1, g_2) : g_1 = \frac{1}{3}g_2 \right\} \end{aligned}$$

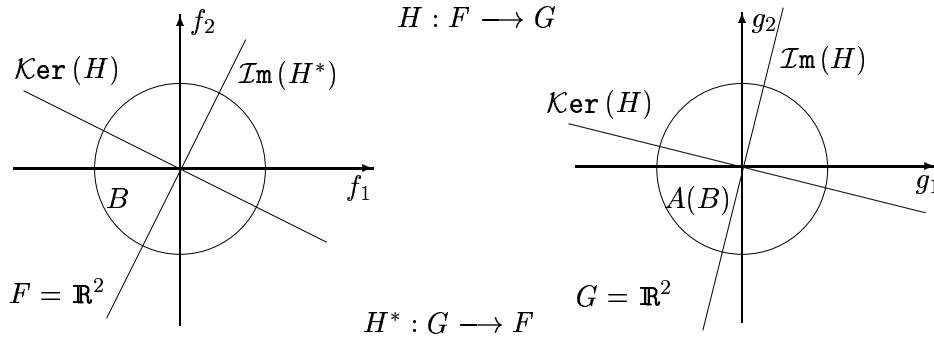
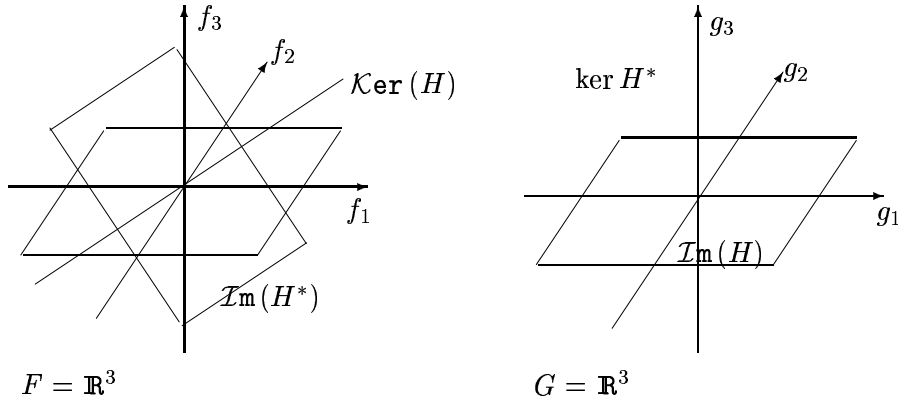


FIG. 4.3 – Ensemble et opérateur compacts

Exemple 2:

$$\begin{pmatrix} g_1 \\ g_2 \\ g_3 \end{pmatrix} = \begin{pmatrix} 1 & 1 & 1 \\ 1 & -1 & 0 \\ 0 & 0 & 0 \end{pmatrix} \begin{pmatrix} f_1 \\ f_2 \\ f_3 \end{pmatrix}$$

$$\mathbf{H} = \begin{pmatrix} 1 & 1 & 1 \\ 1 & -1 & 0 \\ 0 & 0 & 0 \end{pmatrix}, \quad \mathbf{H}^* = \begin{pmatrix} 1 & 1 & 0 \\ 1 & -1 & 0 \\ 1 & 0 & 0 \end{pmatrix}, \quad \mathbf{H}^* \mathbf{H} = \begin{pmatrix} 2 & 0 & 1 \\ 0 & 2 & 1 \\ 1 & 1 & 1 \end{pmatrix}, \quad \mathbf{H} \mathbf{H}^* = \begin{pmatrix} 3 & 0 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & 0 \end{pmatrix}$$

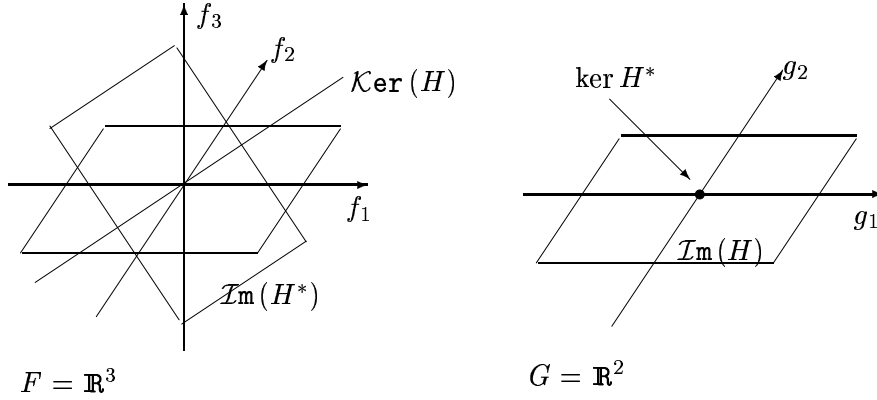
FIG. 4.4 – Image et noyau d'un opérateur et de son adjoint : $F = G = \mathbb{R}^3$

$$\begin{aligned} \text{Ker}(\mathbf{H}) &= \{(f_1, f_2, f_3) : f_1 = f_2, f_3 = 2f_2\} \\ \text{Im}(\mathbf{H}^*) &= \{(f_1, f_2, f_3) : f_1 + f_2 - 2f_3 = 0\} \\ \text{Ker}(\mathbf{H}^*) &= \{(g_1, g_2, g_3) : g_1 = g_2 = 0\} \\ \text{Im}(\mathbf{H}) &= \{(g_1, g_2, g_3) : g_3 = 0\} \end{aligned}$$

Exemple 3:

$$\begin{pmatrix} g_1 \\ g_2 \end{pmatrix} = \begin{pmatrix} 1 & 1 & 1 \\ 1 & -1 & 0 \end{pmatrix} \begin{pmatrix} f_1 \\ f_2 \\ f_3 \end{pmatrix}$$

$$\mathbf{H} = \begin{pmatrix} 1 & 1 & 1 \\ 1 & -1 & 0 \end{pmatrix}, \quad \mathbf{H}^* = \begin{pmatrix} 1 & 1 \\ 1 & -1 \\ 1 & 0 \end{pmatrix}, \quad \mathbf{H}^* \mathbf{H} = \begin{pmatrix} 2 & 0 & 1 \\ 0 & 2 & 1 \\ 1 & 1 & 1 \end{pmatrix}, \quad \mathbf{H} \mathbf{H}^* = \begin{pmatrix} 3 & 0 \\ 0 & 2 \end{pmatrix}$$

FIG. 4.5 – Image et noyau d'un opérateur et de son adjoint : $F = \mathbb{R}^3, G = \mathbb{R}^2$

$$\begin{aligned} \text{Ker}(\mathbf{H}) &= \{(f_1, f_2, f_3) : f_1 = f_2, f_3 = 2f_2\} \\ \text{Im}(\mathbf{H}^*) &= \{(f_1, f_2, f_3) : f_1 + f_2 - 2f_3 = 0\} \\ \text{Ker}(\mathbf{H}^*) &= \{(g_1, g_2) : g_1 = g_2 = 0\} \\ \text{Im}(\mathbf{H}) &= \{(g_1, g_2) : g_1 = g_2 \in \mathbb{R}\} \end{aligned}$$

4.2 Définitions

D'après HADAMARD, un problème est dit *bien posé* si le problème admet une solution (Existence), si elle unique (Unicité) et si elle est stable (Stabilité). Le problème est dit *mal posé* si au moins une de ces trois conditions n'est pas vérifiée.

• **Existence** $\mathbf{g} \in G = \mathcal{I}m(\mathbf{H}) \quad \forall \mathbf{g} \in G \exists \mathbf{f} : \mathbf{H}\mathbf{f} = \mathbf{g}$

• **Unicité** $\mathcal{K}er(\mathbf{H}) = \{0\} \quad \mathbf{H}\mathbf{f} = 0 \longrightarrow \mathbf{f} = 0$

• **Stabilité** $\mathcal{I}m(\mathbf{H}) = \overline{\mathcal{I}m(\mathbf{H})} = \mathcal{K}er(\mathbf{H}^*)^\perp$
faible perturbation de $\mathbf{g}$ donne
faible perturbation de la solution $\mathbf{f}$

$$F = \mathcal{K}er(\mathbf{H}) + \mathcal{K}er(\mathbf{H})^\perp, \quad G = \mathcal{K}er(\mathbf{H}^*) + \mathcal{K}er(\mathbf{H}^*)^\perp$$

Quatre situations mutuellement exclusives :

1. $\mathcal{I}m(\mathbf{H})$ est dense dans G : $(\mathcal{K}er(\mathbf{H}) = \{0\} \text{ et } \mathbf{g} \in \mathcal{I}m(\mathbf{H}))$
Il existe une solution au sens classique.
2. $\mathcal{I}m(\mathbf{H})$ n'est pas dense dans G : $(\mathcal{K}er(\mathbf{H}) = \{0\} \text{ et } \mathbf{g} \notin \mathcal{I}m(\mathbf{H}))$
Il n'existe pas de solution au sens classique.
3. $\overline{\mathcal{I}m(\mathbf{H})} \subset G$ et $\mathbf{g} \in \mathcal{I}m(\mathbf{H}) + \mathcal{I}m(\mathbf{H})^\perp$
On peut définir une solution au sens de l'inversion généralisée.
4. $\overline{\mathcal{I}m(\mathbf{H})} \subset G$ et $\mathbf{g} \notin \mathcal{I}m(\mathbf{H}) + \mathcal{I}m(\mathbf{H})^\perp$
Il n'existe pas de solution au sens classique.

$\mathbf{H}$: opérateur linéaire borné et compact $\longrightarrow \mathcal{I}m(\mathbf{H})$ ouvert

La symétrisation de $\mathbf{H}$ n'apporte rien.

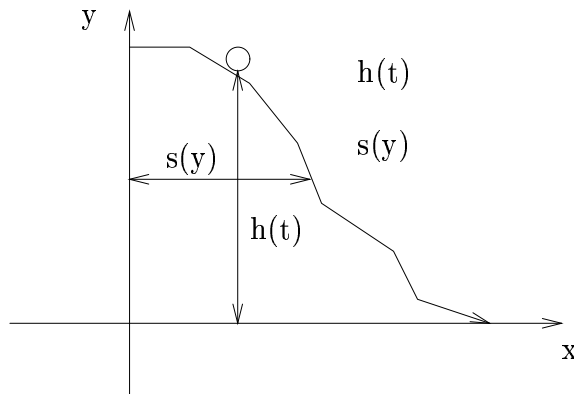
$$\mathcal{I}m(\mathbf{H}^*\mathbf{H}) = \mathcal{I}m(\mathbf{H}^*), \text{ si } \mathcal{K}er(\mathbf{H}) = \{0\} \rightarrow \mathcal{K}er(\mathbf{H}\mathbf{H}^*) = \mathcal{K}er(\mathbf{H}^*)$$

$$\mathbf{H}^*\mathbf{g} = \mathbf{H}^*\mathbf{H}\mathbf{f}$$

est aussi mal-posé. Les problèmes directs sont généralement bien posés. Les problèmes inverses sont souvent mal posés.

4.3 Exemples de problèmes mal-posés

Exemple 1: Problème d'Abel



- $s(y)$ profil de trajectoire d'un point pesant, partant de l'altitude h .
- $t(h)$ temps d'arrivée à l'altitude zéro

$$t(h) = \frac{1}{\sqrt{2g}} \int_0^h \frac{s'(y)}{\sqrt{(h-y)}} dy$$

inversement

$$s(y) = -\frac{\sqrt{2g}}{\pi} \int_0^y \frac{t(h)}{\sqrt{(y-h)}} dh$$

- s_i = résultat d'une mesure de $s(y_i)$,
- $d_s = \sup |s_i - s(y_i)|$ mesure "naturelle",
- t_e = valeur de t permettant de coller exactement sur une donnée expérimentale,
- $d_t = \sup |t_e - t|$ mesure "naturelle",
- L'application $s \mapsto t$ n'est pas continue pour ces distances!

Exemple 2: Diffraction inverse ou Synthèse d'ouverture

- $f(r)$ forme d'un objet,
- $g(u)$ champ diffracté,

$$g(u) = \int_{-\infty}^{\infty} f(r) \exp[-jur] \, dr$$

$$f(r) = \frac{1}{2\pi} \int_{-\infty}^{\infty} g(u) \exp[jur] \, du$$

$$g_1(u) = g_0(u) + \epsilon \exp[-\alpha u^2]$$

$$f_1(r) = f_0(r) + \frac{\epsilon}{(4\pi\alpha)^{1/2}} \exp[-r^2/4\alpha]$$

- norme sup:

$$\max_u |g_1(u) - g_0(u)| \leq \alpha\epsilon$$

$$\max_r |f_1(r) - f_0(r)| \leq \alpha \frac{\epsilon}{(4\pi\alpha)^{1/2}}$$

- norme L^2 :

$$\left\{ \int_{-\infty}^{\infty} [g_1(u) - g_0(u)]^2 \, du \right\}^{1/2} = \epsilon(\pi/2\alpha)^{1/4}$$

$$\left\{ \int_{-\infty}^{\infty} [f_1(r) - f_0(r)]^2 \, dr \right\}^{1/2} = \frac{1}{2\pi} \epsilon(\pi/2\alpha)^{1/4}$$

- Un problème mal-posé peut être rendu bien-posé juste en changeant l'espace des solutions.

Exemple 3: Opérateurs linéaires et compacts

$$g(s) = \int_a^b h(s,t)f(t)dt \quad f \in F, \quad g \in G, \quad F = G = L^2[a,b],$$

et

$$\int_a^b \int_a^b |h(s,t)|^2 ds dt < \infty$$

$\mathbf{H}$ est compact et admet une décomposition spectrale

$$h(s,t) = \sum_{k=1}^{\infty} \lambda_k \phi_k(s) \psi_k(t)$$

$$\int_a^b h(s,t) \psi_k(t) dt = \lambda_k \phi_k(s), \quad \int_a^b h(s,t) \phi_k(s) ds = \lambda_k \psi_k(t)$$

$\{\psi_k\}$ et $\{\phi_k\}$, fonctions orthonormales singulières de $\mathbf{H}$, ou bien fonctions propres de $\mathbf{H}^* \mathbf{H}$ et $\mathbf{H} \mathbf{H}^*$, forment une base de $L^2[a,b]$.

$$g(s) = \sum_{k=1}^{\infty} a_k \phi_k(s)$$

$$f(t) = \sum_{k=1}^{\infty} b_k \psi_k(t) \quad \text{avec} \quad b_k = \frac{\langle g, \phi_k \rangle}{\lambda_k}$$

Pour que cette solution appartienne à $L^2[a,b]$, il faut imposer

$$\|f\|^2 = \sum_{k=1}^{\infty} b_k^2 < \infty$$

$$\sum_{k=1}^{\infty} b_k^2 = \sum_{k=1}^{\infty} \frac{1}{\lambda_k^2} a_k^2 \longrightarrow 0$$

G doit être réduit par rapport à $L^2[a,b]$

sinon, F peut contenir des solutions qui peuvent être inacceptables.

4.4 Déconvolution : un problème mal-posé

Pour montrer que la déconvolution est un problème mal-posé nous pouvons utiliser les outils suivants :

- La transformée de FOURIER qui permet de transformer l'opérateur de la convolution en une multiplication ;
- La discrétisation du problème et l'utilisation de l'algèbre matricielle ;
- L'analyse fonctionnelle et la décomposition spectrale des opérateurs dans l'espace de Hilbert.

Mais avant d'essayer d'expliquer pourquoi regarder sur la figure qui suit

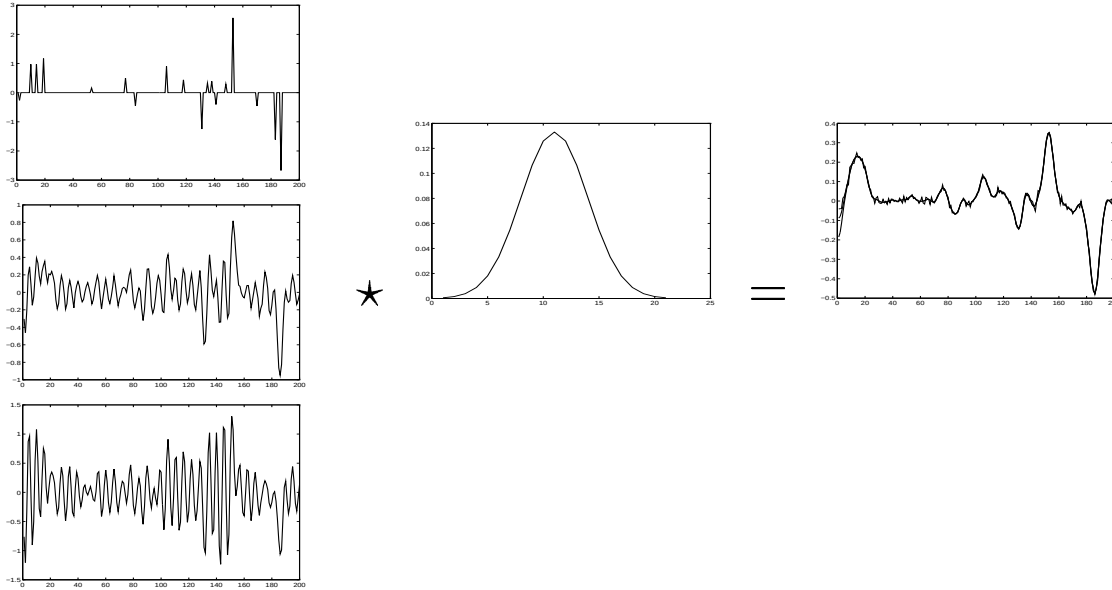


FIG. 4.6 – *Déconvolution : un problème mal-posé.*

On constate que trois signaux $f_1(t)$, $f_2(t)$ et $f_3(t)$ très différents à l'entrée d'un instrument ayant pour réponse impulsionnelle $h(t)$, ont fourni pratiquement les mêmes signaux $g_1(t)$, $g_2(t)$ et $g_3(t)$ de sortie (les trois signaux sont tracés sur un seul graphe). Ceci signifie qu'une méthode d'inversion se basant seulement sur une distance de sortie ne peut distinguer ces trois signaux.

4.4.1 Passage dans le domaine de Fourier

Le schéma suivant résume l'essentiel de cette approche :

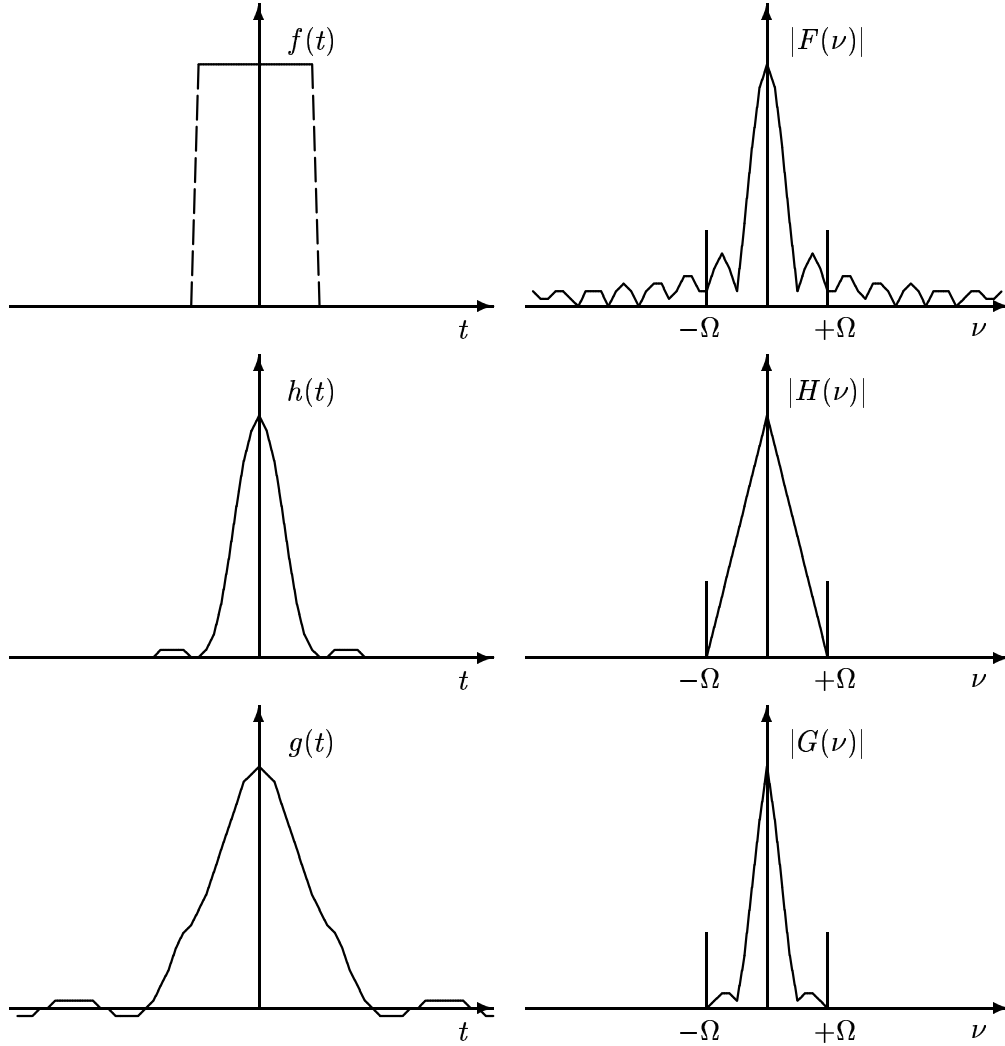


FIG. 4.7 – Déconvolution : un problème mal posé

En effet, on peut constater que tout signal $f(t)$ qui a même contenu spectral dans la bande passante de l'opérateur (l'intervalle $[-\Omega, \Omega]$) fournit la même sortie $g(t)$.

La déconvolution est ainsi un problème inverse mal posé. Un moyen pour le rendre bien posé est de restreindre l'espace des solutions aux fonctions à spectre limité dans la bande $[-\Omega, \Omega]$:

$$F(\omega) = 0 \quad \forall \omega \notin [-\Omega, +\Omega]$$

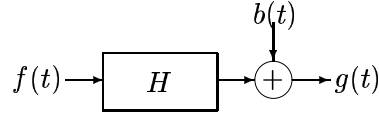
En faisant ceci, on retrouve le critère de Rayleigh classique, on assure l'unicité de la solution mais on n'a pas tout résolu. Une autre difficulté qui persiste est la sensibilité de la solution au bruit. En effet, lorsque les signaux mesurés sont bruités on a :

$$G(\omega) = H(\omega)F(\omega) + B(\omega) \longrightarrow F(\omega) = \frac{G(\omega)}{H(\omega)} + \frac{B(\omega)}{H(\omega)}$$

et il faut étudier le comportement de $B(\omega)/H(\omega)$ qui conditionne cette sensibilité.

4.4.2 Discrétisation

Considérons le modèle :



$$g(t) = \int f(t')h(t-t') dt' + b(t) = \int h(t')f(t-t') dt' + b(t)$$

et faisons les hypothèses suivantes :

- on discrétise les signaux $f(t)$, $g(t)$, $h(t)$ avec le même pas de d'échantillonnage $\Delta T = 1$,
- la réponse impulsionnelle est à support limité : $h(t) = 0, \forall t < -q\Delta T$, et $\forall t > p\Delta T$.

On peut alors écrire :

$$g(m) = \sum_{k=-q}^p h(k)f(m-k) + b(m), \quad m = 0, \dots, M$$

ou sous forme matricielle :

$$\mathbf{g} = \mathbf{H}\mathbf{b} + \mathbf{b}$$

$$\begin{pmatrix} g(0) \\ g(1) \\ \vdots \\ \vdots \\ \vdots \\ g(M) \end{pmatrix} = \begin{pmatrix} h(p) & \cdots & h(0) & \cdots & h(-q) & 0 & \cdots & \cdots & \cdots & 0 \\ 0 & \ddots & & \ddots & & \ddots & & & & \vdots \\ \vdots & & & & & & & & & \vdots \\ \vdots & & & h(p) & \cdots & h(0) & \cdots & h(-q) & & \vdots \\ \vdots & & & & & & & & & \vdots \\ \vdots & & & & & \ddots & & \ddots & & 0 \\ 0 & \cdots & \cdots & \cdots & 0 & h(p) & \cdots & h(0) & \cdots & h(-q) \end{pmatrix} \begin{pmatrix} f(-p) \\ \vdots \\ f(0) \\ f(1) \\ \vdots \\ f(M) \\ f(M+1) \\ \vdots \\ f(M+q) \end{pmatrix}$$

Notons que $\mathbf{g}$ est de dimension $M+1$, $\mathbf{f}$ de dimension $M+p+q+1$, $\mathbf{h} = [h(p), \dots, h(0), \dots, h(-q)]$ de dimension $(p+q+1)$ et la matrice $\mathbf{H}$ de dimensions $(M+1) \times (M+p+q+1)$.

Maintenant si on fait l'hypothèse que le système est causal, on a $q = 0$ et

$$\begin{pmatrix} g(0) \\ g(1) \\ \vdots \\ \vdots \\ \vdots \\ g(M) \end{pmatrix} = \begin{pmatrix} h(p) & \cdots & h(0) & 0 & \cdots & \cdots & \cdots & 0 \\ 0 & & & & & & & \vdots \\ \vdots & & & & & & & \vdots \\ \vdots & & h(p) & \cdots & h(0) & & & \vdots \\ \vdots & & & & & & & \vdots \\ \vdots & & & & & & 0 & \vdots \\ 0 & \cdots & \cdots & \cdots & 0 & h(p) & \cdots & h(0) \end{pmatrix} \begin{pmatrix} f(-p) \\ \vdots \\ f(0) \\ f(1) \\ \vdots \\ \vdots \\ f(M) \end{pmatrix}$$

et si de plus, on fait l'hypothèse que le signal d'entrée est aussi causal, on obtient :

$$\begin{pmatrix} g(0) \\ g(1) \\ \vdots \\ \vdots \\ \vdots \\ g(M) \end{pmatrix} = \begin{pmatrix} h(0) & & & & & \\ h(1) & \ddots & & & & \\ \vdots & & & & & \\ h(p) & \cdots & h(0) & & & \\ 0 & \ddots & & \ddots & & \\ \vdots & & & & & \\ 0 & \cdots & 0 & h(p) & \cdots & h(0) \end{pmatrix} \begin{pmatrix} f(0) \\ f(1) \\ \vdots \\ \vdots \\ \vdots \\ f(M) \end{pmatrix}$$

et, finalement, si $p = M$, on obtient :

$$\begin{pmatrix} g(0) \\ g(1) \\ \vdots \\ g(M) \end{pmatrix} = \begin{pmatrix} h(0) & & & \\ h(1) & h(0) & & \\ \vdots & & \ddots & \\ h(M) & \cdots & h(1) & h(0) \end{pmatrix} \begin{pmatrix} f(0) \\ f(1) \\ \vdots \\ f(M) \end{pmatrix}$$

On peut remarquer que dans tous les cas, la matrice $\mathbf{H}$ est TOEPLITZ et, sauf dans ces deux derniers cas, elle est aussi circulante.

Dans le cas où la matrice n'est pas circulante, il est possible de la rendre artificiellement circulante en ajoutant des zéros à la suite du vecteur $\mathbf{f}$ (bourrage de zéros). En effet on peut toujours écrire :

$$\begin{pmatrix} g(0) \\ g(1) \\ \vdots \\ \vdots \\ \vdots \\ g(M) \end{pmatrix} = \begin{pmatrix} h(0) & & & h(p) & \cdots & h(1) \\ h(1) & \ddots & & & \ddots & \vdots \\ \vdots & & & & & h(p) \\ h(p) & \cdots & h(0) & & & 0 \\ 0 & \ddots & & \ddots & & \vdots \\ \vdots & & & & & f(M) \\ 0 & \cdots & 0 & h(p) & \cdots & h(0) \\ & & & 0 & \cdots & 0 \end{pmatrix} \begin{pmatrix} f(0) \\ f(1) \\ \vdots \\ \vdots \\ \vdots \\ f(M) \\ 0 \\ \vdots \\ 0 \end{pmatrix}$$

Mettons nous ensuite dans la situation idéale sans bruit. On pourrait alors s'imaginer qu'il suffit de choisir $p = M$ pour que la matrice $\mathbf{H}$ soit carrée et définir la solution par $\hat{\mathbf{f}} = \mathbf{H}^{-1}\mathbf{g}$ ou même dans les autres cas, de définir la solution soit par $\hat{\mathbf{f}} = (\mathbf{H}^t \mathbf{H})^{-1} \mathbf{H}^t \mathbf{g}$ ou par $\hat{\mathbf{f}} = \mathbf{H}^t (\mathbf{H} \mathbf{H}^t)^{-1} \mathbf{g}$. Là encore il ne faut oublier que la matrice $\mathbf{H}$ (ou $\mathbf{H}^t \mathbf{H}$ ou $\mathbf{H} \mathbf{H}^t$) peut être singulière ou du moins très mal conditionnée. En effet, le conditionnement de ces matrices dépend du pas d'échantillonnage. Pour comprendre cela, il suffit remarquer que les lignes successives de $\mathbf{H}$ sont très proches (due au caractère TOEPLITZ), et cela d'autant plus que le pas d'échantillonnage est petit. On peut aussi expliquer ceci à l'aide des propriétés des valeurs propres et des vecteurs propres des matrices de TOEPLITZ:

A COMPLETER

4.4.3 Décomposition spectrale

- $\mathbf{H} : F \mapsto G$ linéaire borné et Hilbertien
 $\{\psi_k, \phi_k; \lambda_k^2\}$ système singulier de $\mathbf{H}$:

$$\mathbf{H}\psi_k = \lambda_k \phi_k, \quad \mathbf{H}^* \phi_k = \lambda_k \psi_k$$

$$\mathbf{H}\mathbf{H}^* \phi_k = \lambda_k^2 \phi_k, \quad \mathbf{H}^* \mathbf{H} \psi_k = \lambda_k^2 \psi_k$$

- Résolution de $\mathbf{H}^* \mathbf{g} = \mathbf{H}^* \mathbf{H} \mathbf{f}$ par projection :
 $\{\psi_k\}$ base de G , $\{\phi_k\}$ base de F

$$x = \sum_{k=1}^{\infty} \frac{1}{\lambda_k} \langle y, \phi_k \rangle \psi_k$$

- **Théorème de Hilbert-Schmidt :**

Si $\mathbf{H}$ est compact, alors 0 est le seul point d'accumulation de $\{\lambda_k^2\}$:

$$\lim_{k \rightarrow \infty} \lambda_k^2 = 0 \quad \exists k_0 : \forall k > k_0, \lambda_k = 0$$

$\Downarrow$

$$x + \delta x = \sum_{k=1}^{\infty} \frac{1}{\lambda_k} \langle y, \phi_k \rangle \left[1 + \frac{\langle \delta y, \phi_k \rangle}{\langle y, \phi_k \rangle} \right] \psi_k$$

◊ Changer la notion de la solution

4.5 Changer la notion de solution

- Bien définir le problème direct $\mathbf{H} : F \mapsto G$
- Pour tout $\mathbf{g} \in G$, étudier l'existence des sous ensembles non vides de $F : \mathbf{H}\mathbf{f} = \mathbf{g}$
- Rechercher et définir des applications $G \mapsto F$
- étudier le problème des erreurs après avoir donné à F et G des structures d'espaces métriques.

Les propriétés de la solution ne sont pas entièrement contenues dans $\mathbf{g} = \mathbf{H}\mathbf{f}$

- Solutions approchées $\hat{\mathbf{f}} : \|\mathbf{H}\hat{\mathbf{f}} - \mathbf{g}\| \leq \epsilon, \epsilon > 0$ fixé
- Quasi-solutions $\hat{\mathbf{f}} : \|\mathbf{g} - \mathbf{H}\hat{\mathbf{f}}\| \leq \|\mathbf{g} - \mathbf{H}\mathbf{f}\|, \quad \forall \mathbf{f} \in F$
- Changement d'espace ou de topologie
- Opérateurs régularisants
- Concepts probabilistes

Chapitre 5

Moindres carrés et inversion généralisée

Nous avons vu qu'un problème est dit mal-posé s'il ne satisfait pas à une des trois conditions : existence, unicité et stabilité. Lorsque la difficulté vient de la non unicité de la solution l'inversion généralisée permet d'y remédier. En effet, la résolution de l'équation $\mathbf{g} = \mathbf{H}\mathbf{f}$ présente trois possibilités :

- Il y a une solution au problème ;
- Il y a une infinité de solutions au problème ;
- Il n'y pas de solution au problème.

Dans le premier cas, il ne reste qu'à calculer l'inverse. Dans le second cas, le problème devient celui d'un choix parmi l'ensemble des solutions suivant un critère (par exemple, la norme minimale). Dans le troisième cas, on souhaiterait aussi pouvoir définir une solution, cela est possible en définissant un critère du type distance dans l'espace G entre $\mathbf{g}$ et $\mathbf{H}\mathbf{f}$, (par exemple la distance euclidienne), et en choisissant comme solution celle qui la minimise. Dans le cas où la distance choisie est une distance euclidienne la solution est dite *au sens des moindres carrés* et nous verrons qu'il y a un lien entre la solution au sens des MC et celle de l'inverse généralisée.

5.1 Définition en dimension infinie

Soit $\mathbf{H} : F \mapsto G$,

$\mathbf{P}$ le projecteur orthogonal de F sur $\text{Ker}(\mathbf{H})^\perp = \overline{\text{Im}(\mathbf{H}^*)}$ et

$\mathbf{Q}$ le projecteur orthogonal de G sur $\text{Ker}(\mathbf{H}^*)^\perp = \overline{\text{Im}(\mathbf{H})}$. On a alors

$$\mathbf{H}\mathbf{f} = \mathbf{H}\mathbf{P}\mathbf{f} \quad \mathbf{H}^*\mathbf{g} = \mathbf{H}^*\mathbf{Q}\mathbf{g}$$

Définition 1 *L'inverse généralisée $\mathbf{H}^+$ de $\mathbf{H}$ est l'opérateur dont le domaine est $\mathcal{D}(\mathbf{H}^+) = \text{Im}(\mathbf{H}) \oplus \text{Im}(\mathbf{H})^\perp$, tel que $\mathbf{H}^+\mathbf{g} = \mathbf{f}_0 \in F$, où $\mathbf{f}_0 \in S = \{\mathbf{f} \in F : \mathbf{H}\mathbf{f} = \mathbf{Q}\mathbf{g}\}$ et $\|\mathbf{f}_0\| \leq \|\mathbf{f}\|$, $\forall \mathbf{f} \in S, \mathbf{f} \neq \mathbf{f}_0$.*

Notons que

$$\|\mathbf{g} - \mathbf{H}\mathbf{f}\|^2 = \|\mathbf{Q}\mathbf{g} - \mathbf{H}\mathbf{f}\|^2 + \|\mathbf{g} - \mathbf{Q}\mathbf{g}\|^2$$

Si $\mathbf{g} \in \text{Im}(\mathbf{H}) \oplus \text{Im}(\mathbf{H})^\perp$, l'ensemble des solutions des MC est non-vide, fermé et convexe. Il existe donc un élément unique de norme minimale, qui est $\mathbf{H}^+\mathbf{g}$.

Limitations :

- $\forall \mathbf{g} \in G$, il existe une solution inverse généralisée ssi $\text{Im}(\mathbf{H})$ est fermée.
- En général $\mathbf{H}$ est compact, $\text{Im}(\mathbf{H})$ n'est pas fermé, sauf si $\mathbf{H}$ est de dimension finie.
- Existence : Non $\text{Im}(\mathbf{H}) \oplus \text{Im}(\mathbf{H})^\perp \neq G$
- Unicité : Oui
- Stabilité : Non $\text{Im}(\mathbf{H}) \neq \overline{\text{Im}(\mathbf{H})}$

5.2 Définition en dimension finie

Considérons l'équation

$$\mathbf{g} = \mathbf{H}\mathbf{f}$$

où $\mathbf{H}$ est une matrice de dimension $[m, n]$.

Définition 2 (Quasi-solution) On appelle $\hat{\mathbf{f}}$ une quasi-solution de l'équation $\mathbf{g} = \mathbf{H}\mathbf{f}$ une solution qui minimise une distance $\Delta(\mathbf{g}, \mathbf{H}\mathbf{f})$ entre $\mathbf{g}$ et $\mathbf{H}\mathbf{f}$:

$$\hat{\mathbf{f}} = \arg \min_{\mathbf{f}} \{\Delta(\mathbf{g}, \mathbf{H}\mathbf{f})\}$$

Lorsque la distance choisie est une distance euclidienne $\Delta(\mathbf{g}, \mathbf{H}\mathbf{f}) = \|\mathbf{g} - \mathbf{H}\mathbf{f}\|^2$, la solution ainsi définie est dite "au sens des moindres carrés".

Définition 3 (Moindre carré) On appelle $\hat{\mathbf{f}}$ une solution au sens des moindres carrés de l'équation $\mathbf{g} = \mathbf{H}\mathbf{f}$ une solution qui minimise la norme euclidienne $\|\mathbf{g} - \mathbf{H}\mathbf{f}\|^2$:

$$\hat{\mathbf{f}} = \arg \min_{\mathbf{f}} \{\|\mathbf{g} - \mathbf{H}\mathbf{f}\|^2\}.$$

Cet ensemble coïncide avec l'ensemble des solutions de l'équation normale

$$[\mathbf{H}^t \mathbf{H}] \hat{\mathbf{f}} = \mathbf{H}^t \mathbf{g}.$$

- Si $M = N$ et $\text{rang}\{\mathbf{H}\} = N$, alors $\hat{\mathbf{f}} = \mathbf{f} = \mathbf{H}^{-1}\mathbf{g}$.
- Si $N \leq M$ et $\text{rang}\{\mathbf{H}\} = N$, alors $[\mathbf{H}^t \mathbf{H}]$ est de rang plein et $\hat{\mathbf{f}} = [\mathbf{H}^t \mathbf{H}]^{-1} \mathbf{H}^t \mathbf{g}$.
- Si $M < N$, alors $[\mathbf{H}^t \mathbf{H}]$ est forcément singulière. Elle a un rang au maximum égal à M , et l'équation (1) n'a pas de solution unique (Elle a une infinité de solutions).

Définition 4 (Inverse généralisée) On appelle $\mathbf{f}^+ = \mathbf{H}^+ \mathbf{g}$ la solution inverse généralisée de l'équation $\mathbf{g} = \mathbf{H}\mathbf{f}$ la solution MC de norme minimale de l'équation normale $[\mathbf{H}^t \mathbf{H}] \hat{\mathbf{f}} = \mathbf{H}^t \mathbf{g}$. Si $\mathbf{g} \in \text{Im}(\mathbf{H})$ alors cette solution existe et elle est unique. Elle est donnée par :

$$\mathbf{f}^+ = \mathbf{H}^+ \mathbf{g}.$$

où $\mathbf{H}^+$ est l'inverse généralisée de $\mathbf{H}$.

Si $\text{rang}\{\mathbf{H}\} = N$ alors $\mathbf{H}^+ = [\mathbf{H}^t \mathbf{H}]^{-1} \mathbf{H}^t$ et si $\text{rang}\{\mathbf{H}\} = M$ alors $\mathbf{H}^+ = \mathbf{H}^t [\mathbf{H}^t \mathbf{H}]^{-1}$.

Pour être complet nous mentionnons deux autres définitions pour l'inverse généralisée d'une matrice.

Définition 5 (Moore) Une matrice $\mathbf{G}$ de dimension $[N, M]$ est dite inverse généralisée de la matrice $\mathbf{H}$ si

$$\mathbf{H}\mathbf{G} = \mathbf{P}_A \quad \text{et} \quad \mathbf{G}\mathbf{H} = \mathbf{P}_G$$

où $\mathbf{P}_A$ est l'opérateur (matrice) de projection des vecteurs de l'espace $\mathbb{R}^m$ sur $\mu(\mathbf{H})$ et $\mathbf{P}_G$ est l'opérateur (matrice) de projection des vecteurs de l'espace $\mathbb{R}^n$ sur $\mu(\mathbf{G})$.

Définition 6 (Penrose) Une matrice $\mathbf{G}$ de dimension $[N, M]$ est dite inverse généralisée de la matrice $\mathbf{H}$ si

$$\mathbf{H}\mathbf{G}\mathbf{H} = \mathbf{H}, \quad (\mathbf{H}\mathbf{G})^t = \mathbf{H}\mathbf{G}, \quad \mathbf{G}\mathbf{H}\mathbf{G} = \mathbf{G} \quad \text{et} \quad (\mathbf{G}\mathbf{H})^t = \mathbf{G}\mathbf{H}$$

Ces définitions sont identiques si les normes associées aux espaces $\mathbb{R}^m$ et $\mathbb{R}^n$ sont du type euclidien $\|\mathbf{f}\| = \langle \mathbf{f}, \mathbf{f} \rangle^{1/2}$.

Supposons que $\mathbf{H}\mathbf{f} = \mathbf{g}$ et calculons l'erreur quadratique moyenne entre une estimation $\hat{\mathbf{f}}$ linéaire quelconque $\hat{\mathbf{f}} = \mathbf{G}^+\mathbf{g}$ de la solution et $\mathbf{f}$:

$$EQM = [\mathbf{f} - \hat{\mathbf{f}}]^t [\mathbf{f} - \hat{\mathbf{f}}] = \text{Tr} \left\{ [\mathbf{f} - \hat{\mathbf{f}}][\mathbf{f} - \hat{\mathbf{f}}]^t \right\} = \text{Tr} \left\{ \mathbf{f}\mathbf{f}^t [\mathbf{I} - \mathbf{H}^+\mathbf{H}]^t \right\}$$

Il est alors facile de vérifier que cette erreur est minimale si $\hat{\mathbf{f}} = \mathbf{H}^+\mathbf{g}$ et qu'elle est donnée par

$$EQM = [\mathbf{f} - \hat{\mathbf{f}}]^t [\mathbf{f} - \hat{\mathbf{f}}] = \text{Tr} \left\{ \mathbf{f}\mathbf{f}^t [\mathbf{I} - \mathbf{H}^+\mathbf{H}]^t \right\}$$

l'estimation est parfaite si $\mathbf{H}^+\mathbf{H} = \mathbf{I}$.

En résumé, lorsque nous voulons calculer une solution au sens inverse généralisée de l'équation $\mathbf{g} = \mathbf{H}\mathbf{f}$ avec $\mathbf{H}$ une matrice de dimensions $[M, N]$, quatre situations se présentent :

- si $M = N$ et $\text{rang} \{\mathbf{H}\} = N$ alors $\mathbf{H}^+ = \mathbf{H}^{-1}$
- si $M > N$ et $\text{rang} \{\mathbf{H}\} = N$ alors $\mathbf{H}^+ = (\mathbf{H}^t \mathbf{H})^{-1} \mathbf{H}^t$
- si $M < N$ et $\text{rang} \{\mathbf{H}\} = M$ alors $\mathbf{H}^+ = \mathbf{H}^t (\mathbf{H} \mathbf{H}^t)^{-1}$
- si $\text{rang} \{\mathbf{H}\} = K < \inf(M, N)$ alors $\left\{ \begin{array}{l} \text{Décomp. en valeurs singulières} \\ \text{Méthodes itératives} \end{array} \right.$

5.3 Algorithmes d'inversion généralisée

5.4 Décomposition en valeurs singulières

Considérons l'équation

$$\mathbf{g} = \mathbf{H}\mathbf{f}$$

et notons $\{\mathbf{v}_j, j = 1, \dots, K\}$ les vecteurs propres de $\mathbf{H}^t\mathbf{H}$, $\{\mathbf{u}_i, i = 1, \dots, K\}$ les vecteurs propres de $\mathbf{H}\mathbf{H}^t$ et λ_j^2 les valeurs propres correspondantes. Nous avons alors :

$$\begin{aligned} \mathbf{H}^t\mathbf{H}\mathbf{v}_j &= \lambda_j^2\mathbf{v}_j, & j = 1, \dots, K \\ \mathbf{H}\mathbf{H}^t\mathbf{u}_i &= \lambda_i^2\mathbf{u}_i, & i = 1, \dots, K \\ \mathbf{H}\mathbf{v}_j &= \lambda_j\mathbf{u}_j, & j = 1, \dots, K \\ \mathbf{H}^t\mathbf{u}_i &= \lambda_i\mathbf{v}_i, & i = 1, \dots, K \end{aligned}$$

$$\mathbf{H} = \mathbf{U}\mathbf{\Lambda}\mathbf{V}^t$$

$$\mathbf{U} = (\mathbf{u}_1 : \mathbf{u}_2 : \dots : \mathbf{u}_K), \quad \mathbf{V} = (\mathbf{v}_1 : \mathbf{v}_2 : \dots : \mathbf{v}_K), \quad \mathbf{\Lambda} = \text{diag}(\lambda_1, \lambda_2, \dots, \lambda_K)$$

$$\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_K, \quad K \leq \inf(M, N).$$

La solution inverse généralisée au sens du Moore Penrose est définie par

$$\mathbf{f}^+ = \mathbf{H}^+\mathbf{g}, \quad \text{avec} \quad \mathbf{H}^+ = \mathbf{V}\mathbf{\Lambda}^+\mathbf{U}^t$$

et où

$$\mathbf{\Lambda}^+ = \text{diag}\{\alpha_i\}, \quad \begin{cases} \alpha_i = \frac{1}{\lambda_i} & \text{si } \lambda_i \neq 0 \\ \alpha_i = 0 & \text{si } \lambda_i \simeq 0 \end{cases}$$

On a aussi les relations suivantes :

$$\mathcal{Ker}(\mathbf{H}) = \mathbf{k}(z) = \mathbf{V}(\mathbf{I} - \Lambda^+ \Lambda)z, \quad \forall z \in \mathbb{R}^n$$

Les composantes d'indices 0 à K des vecteurs $\mathbf{k}(z)$ sont nulles. On peut alors écrire :

$$\mathbf{k}(z) = \sum_{k=K+1}^N z_k \mathbf{v}_k \quad \forall z \in \mathbb{R}^n$$

et définir ainsi l'ensemble S des solutions de l'équation $\mathbf{g} = \mathbf{H}\mathbf{f}$ par

$$S = \{\mathbf{f} : \mathbf{f} = \mathbf{f}^+ + \mathbf{k}(z) \quad \forall z \in \mathbb{R}^n\}$$

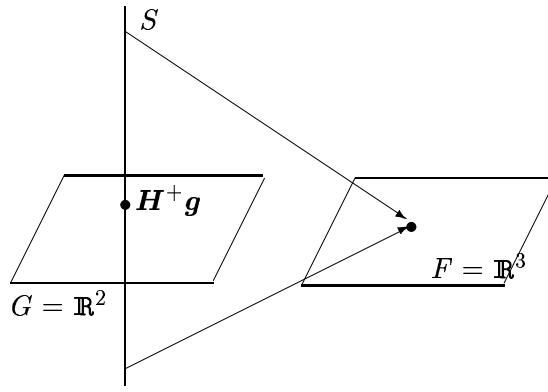


FIG. 5.1 – Inversion généralisée et espace des solutions.

5.4.1 Méthode de Bialy

Cette méthode est basée sur le théorème suivant :

Théorème 1 Soit $\mathbf{H} : H \mapsto H$ un opérateur borné et non négatif. Soit $\mathbf{P}$ l'opérateur de projection orthogonale sur $\text{Ker}(\mathbf{H})$. Alors, s'il existe une solution pour $\mathbf{H}\mathbf{f} = \mathbf{g}$, pour tout $\mathbf{g} \in H, \mathbf{f}_0 \in H$ la suite des itérations

$$\mathbf{f}_{n+1} = \mathbf{f}_n + \alpha(\mathbf{g} - \mathbf{H}\mathbf{f}_n), \quad \text{avec } 0 < \alpha < 2/\|\mathbf{H}\|$$

converge vers $\mathbf{P}\mathbf{f}_0 + \mathbf{f}_{IG}$, où $\mathbf{f}_{IG}$ est la solution IG de $\mathbf{H}\mathbf{f} = \mathbf{g}$.

- si $\mathbf{f}_0 = \mathbf{0}$, alors $\lim_{n \rightarrow \infty} \mathbf{f}_n = \mathbf{f}_{IG}$
- $\mathbf{f}_{IG}$ est la solution de norme minimale de $\mathbf{H}^t \mathbf{H}\mathbf{f} = \mathbf{H}^t \mathbf{g}$.

5.4.2 Méthode de Landweber

Cette méthode, qui est une extension de la méthode précédente, est basée sur le théorème suivant :

Théorème 2 Soient $\mathbf{H} : H_1 \mapsto H_2$ un opérateur borné, $\mathbf{g} \in \text{Im}(\mathbf{H}) + \text{Im}(\mathbf{H})^\perp$. Alors la suite des itérations

$$\begin{aligned} \mathbf{f}_0 &= \mathbf{0} \\ \mathbf{f}_{n+1} &= \mathbf{f}_n + \alpha \mathbf{H}^t(\mathbf{g} - \mathbf{H}\mathbf{f}_n), \quad \text{avec } 0 < \alpha < 2/\|\mathbf{H}^t \mathbf{H}\| \end{aligned}$$

converge vers la solution IG de $\mathbf{H}\mathbf{f} = \mathbf{g}$.

Notons que cet algorithme peut être considéré comme un algorithme du type gradient pour la minimisation du critère des MC $J(\mathbf{f}) = \|\mathbf{g} - \mathbf{H}\mathbf{f}\|^2$. En effet, on peut remarquer que le gradient de ce critère est $\nabla J(\mathbf{f}) = -2\mathbf{H}^t(\mathbf{g} - \mathbf{H}\mathbf{f})$.

5.4.3 Méthode de Van-Cittert

Cette méthode est une extension de la méthode précédente et permet d'apporter des contraintes supplémentaires sur la solution. En effet, si $\mathbf{T}$ est un opérateur traduisant ces contraintes on peut envisager l'algorithme suivant :

$$\begin{cases} \mathbf{f}_0 = \mathbf{0} \\ \mathbf{f}_{n+1} = \mathbf{f}_n + \alpha \mathbf{H}^t(\mathbf{g} - \mathbf{H}\mathbf{T}\mathbf{f}_n), \quad \text{avec } 0 < \alpha < 2/\|\mathbf{T}^t \mathbf{H}^t \mathbf{H}\mathbf{T}\| \end{cases}$$

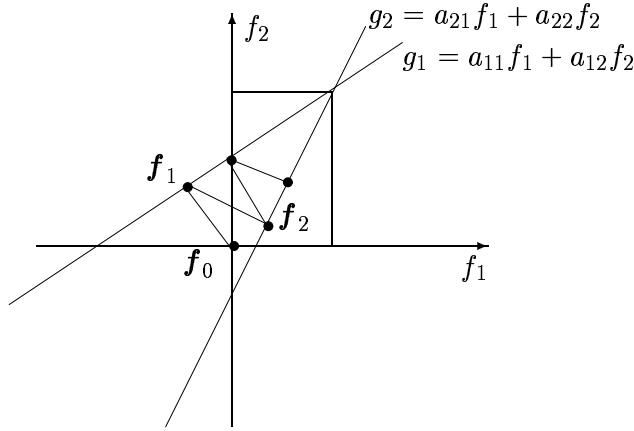
5.4.4 Méthode de Kaczmarz

L'idée de base dans cette méthode est d'utiliser les données g_i une par une. Ainsi, une mise en œuvre en ligne serait possible. Dans une relation linéaire $\mathbf{g} = \mathbf{H}\mathbf{f}$ chaque équation $g_i = \sum_{j=1}^N h_{ij} f_j = \langle \mathbf{h}_i, \mathbf{f} \rangle$ définit un sous-espace dans l'espace des solutions. Cet algorithme recherche une solution à l'itération $k+1$ en projetant la solution à l'itération k sur le sous-espace créé par l'équation $g_k = \sum_{j=1}^N h_{kj} f_j$.

$$\begin{cases} \mathbf{f}_0 = \mathbf{0} \\ \mathbf{f}_{n+1} = \mathbf{f}_n + \frac{g_i - \langle \mathbf{f}_n, \mathbf{h}_i \rangle}{\langle \mathbf{h}_i, \mathbf{h}_i \rangle} \mathbf{h}_i \quad i = 1, 2, \dots, M, 1, 2, \dots, M, 1, 2 \end{cases}$$

où $\mathbf{h}_i$ est la $i^{\text{ème}}$ ligne de la matrice $\mathbf{H}$.

Pour bien comprendre l'essentiel de cet algorithme, considérons le cas d'une matrice $\mathbf{H}$ de dimensions (2×2) . Dans ce cas, chaque équation décrit une ligne droite dans l'espace des solutions (f_1, f_2) et la figure ci-dessous montre alors le principe de cet algorithme.



Évidemment, lorsqu'on veut imposer d'autres contraintes, comme par exemple la positivité, on peut introduire au cours de ces itérations. Ainsi, si on note par $p_i(\mathbf{f}), i = 1, \dots, m$ les m contraintes supplémentaires que l'on veut imposer à la solution, l'algorithme s'écrit :

$$\begin{cases} \mathbf{f}_0 = \mathbf{0} \\ \mathbf{f}_{n+1} = P(\mathbf{f}_n) = p_m(\dots p_2(p_1(\mathbf{f})) \dots) \\ p_i(\mathbf{f}) = \mathbf{f}_n + \frac{g_i - \langle \mathbf{f}_n, \mathbf{h}_i \rangle}{\langle \mathbf{h}_i, \mathbf{h}_i \rangle} \mathbf{h}_i \quad i = 1, 2, \dots, M, 1, 2, \dots, M, 1, 2 \end{cases}$$

Chapitre 6

Régularisation

La régularisation d'un problème **mal-posé** consiste à le transformer en un problème **bien-posé**, c'est-à-dire à définir **une solution unique** pour **toutes les mesures possibles** dans l'espace des observations, et à **assurer la stabilité** de cette solution vis-à-vis des **erreurs** sur ces mesures.

6.1 Définition d'un opérateur de régularisation

Un **régulariseur** de $\mathbf{g} = \mathbf{H}\mathbf{f}$ est une famille d'opérateurs $\{\mathbf{R}_\lambda; \lambda \in \Lambda \subset \mathbf{R}^+\}$ tels que :

- $\forall \mathbf{f} \in F, \lim_{\lambda \rightarrow 0} \|\mathbf{R}_\lambda \mathbf{H}\mathbf{f} - \mathbf{f}\|^2 = 0$
- $\forall \lambda \in \Lambda, \mathbf{R}_\lambda$ est un opérateur continu de G dans F

$\Downarrow$
 Solution Approchée

$$\mathbf{f}_\epsilon = \mathbf{R}_\lambda \mathbf{g}_\epsilon, \quad \text{avec} \quad \mathbf{g}_\epsilon = \mathbf{H}\mathbf{f} + \mathbf{b}$$

$$\|\mathbf{f}_\epsilon - \mathbf{H}^{-1}\mathbf{g}\| \leq \underbrace{\|\mathbf{R}_\lambda \mathbf{g} - \mathbf{H}^{-1}\mathbf{g}\|}_{\text{erreur de régularisation}} + \underbrace{\|\mathbf{R}_\lambda(\mathbf{g}_\epsilon - \mathbf{g})\|}_{\text{erreur due au bruit}}$$

- Utilisation d'E.H.N.R.
- Décomposition TSVD
- Régularisation P.T.T.

6.2 Décomposition tronquée en valeurs singulières

Considérons l'équation

$$\mathbf{g} = \mathbf{H} \mathbf{f}$$

et notons par $\{\mathbf{v}_j, j = 1, \dots, K\}$ les vecteurs propres de $\mathbf{H}^t \mathbf{H}$ et par $\{\mathbf{u}_i, i = 1, \dots, K\}$ les vecteurs propres de $\mathbf{H} \mathbf{H}^t$ et par λ_j^2 les valeurs propres correspondantes. Nous avons alors :

$$\begin{aligned} \mathbf{H}^t \mathbf{H} \mathbf{v}_j &= \lambda_j^2 \mathbf{v}_j, \quad j = 1, \dots, K \\ \mathbf{H} \mathbf{H}^t \mathbf{u}_i &= \lambda_i^2 \mathbf{u}_i, \quad i = 1, \dots, K \\ \mathbf{H} \mathbf{v}_j &= \lambda_j \mathbf{u}_j, \quad j = 1, \dots, K \\ \mathbf{H}^t \mathbf{u}_i &= \lambda_i \mathbf{v}_i, \quad i = 1, \dots, K \end{aligned}$$

$$\mathbf{H} = \mathbf{U} \mathbf{\Lambda} \mathbf{V}^t$$

$$\mathbf{U} = (\mathbf{u}_1 : \mathbf{u}_2 : \dots : \mathbf{u}_K), \quad \mathbf{V} = (\mathbf{v}_1 : \mathbf{v}_2 : \dots : \mathbf{v}_K), \quad \mathbf{\Lambda} = \text{diag}(\lambda_1, \lambda_2, \dots, \lambda_K)$$

$$\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_K, \quad K \leq \inf(M, N)$$

- On décide qu'à partir du rang $r < K$, les valeurs singulières ordonnées ne sont plus significatives.
- Dans l'inversion, on utilise la pseudo-inverse $\mathbf{H}^+$.

$$\mathbf{H}^+ = \mathbf{V} \mathbf{\Lambda}^+ \mathbf{U}^t$$

$$\mathbf{\Lambda}^+ = \text{diag}\{\alpha_i\}, \quad \begin{cases} \alpha_i = \frac{1}{\lambda_i} & \text{si } \lambda_i \neq 0 \text{ (significatif)} \\ \alpha_i = 0 & \text{si } \lambda_i \simeq 0 \text{ (non-significatif)} \end{cases}$$

$$\mathbf{g} + \delta \mathbf{g} = \mathbf{H}(\mathbf{f} + \delta \mathbf{f}) \longrightarrow \frac{\|\delta \mathbf{f}\|}{\|\mathbf{f}\|} \leq \left(\frac{\lambda_{\max}}{\lambda_{\min}} \right)^2 \frac{\|\delta \mathbf{g}\|}{\|\mathbf{g}\|}$$

6.3 Régularisation par méthodes itératives

- Méthodes itératives du type inversion généralisée (BIALY, LANDWEBER,...)
- établir une règle d'arrêt fondée sur l'évolution du résidu

$$\|\mathbf{H}\hat{\mathbf{f}}^k - \mathbf{g}_\epsilon\| > \delta \text{ pour } i < k$$

$$\|\mathbf{H}\hat{\mathbf{f}}^k - \mathbf{g}_\epsilon\| < \delta \text{ pour } i > k$$

- $\frac{1}{k}$ joue le rôle du paramètre de régularisation.

6.4 Introduction de contraintes supplémentaires

$\mathbf{C}$ un opérateur de contrainte

$$\mathbf{f} = \mathbf{C}\mathbf{f} \quad \text{ssi } \mathbf{f} \text{ satisfait la contrainte}$$

Exemple : Positivité

$$x_j = \begin{cases} x_j & \text{si } x_j > 0 \\ 0 & \text{sinon} \end{cases}, \quad j = 1, \dots, n$$

Le problème initial $\mathbf{g} = \mathbf{H}\mathbf{f}$ devient alors

$$\mathbf{g} = \mathbf{H}\mathbf{C}\mathbf{f}$$

L'itération générale est modifiée en remplaçant l'opérateur $\mathbf{H}$ par $\mathbf{H}\mathbf{C}$.

6.5 Régularisation avec *a priori* de douceur

Tikhonov (régularisation continue) :

$$g_i = \int f(t) h_i(t) dt + b_i$$

$$\begin{aligned} J(f) &= Q(f) + \lambda \Omega(f) \\ &= \sum_i \left| g_i - \int f(t) h_i(t) dt \right|^2 + \lambda \int \sum_k p_k(t) \left| \frac{\partial^k f}{\partial t^k} \right|^2 dt, \quad \lambda > 0, \end{aligned}$$

où les $p_k(t)$ sont des fonctions positives, continues, et continûment dérivables.

PHILIPPS et TWOMEY (régularisation discrète) :

$$J_\lambda(\mathbf{f}) = \|\mathbf{g} - \mathbf{H}\mathbf{f}\|^2 + \lambda \Omega(\mathbf{f}) \quad \lambda > 0$$

Exemple : $k = 1, p_k(t) = 1$

$$J_\lambda(f) = \sum_i \left| g_i - \int f(t) h_i(t) dt \right|^2 + \lambda \int |f'(t)|^2 dt$$

Cas discret :

$$J_\lambda(\mathbf{f}) = \|\mathbf{g} - \mathbf{H}\mathbf{f}\|^2 + \lambda \|\mathbf{D}_k \mathbf{f}\|^2 = [\mathbf{g} - \mathbf{H}\mathbf{f}]^t [\mathbf{g} - \mathbf{H}\mathbf{f}] + \lambda [\mathbf{D}_k \mathbf{f}]^t [\mathbf{D}_k \mathbf{f}]$$

avec

$$\mathbf{D}_0 = \mathbf{I}, \quad \mathbf{D}_1 = \begin{pmatrix} 1 & 0 & \cdots & \cdots & 0 \\ -1 & 1 & \ddots & & \vdots \\ 0 & -1 & 1 & \ddots & \vdots \\ \vdots & \ddots & & & 0 \\ 0 & \cdots & 0 & -1 & 1 \end{pmatrix}, \quad \mathbf{D}_2 = \begin{pmatrix} 1 & 0 & \cdots & \cdots & \cdots & 0 \\ -2 & 1 & \ddots & & & \vdots \\ 1 & -2 & 1 & \ddots & & \vdots \\ 0 & 1 & -2 & 1 & & \vdots \\ \vdots & \ddots & & & & \vdots \\ 0 & \cdots & 0 & 1 & -2 & 1 \end{pmatrix}$$

$$\begin{aligned} \nabla J_\lambda &= 2\mathbf{H}^t[\mathbf{g} - \mathbf{H}\mathbf{f}] + 2\lambda \mathbf{D}^t \mathbf{D} \mathbf{f} = 0 \longrightarrow \\ [\mathbf{H}^t \mathbf{H} + \lambda \mathbf{D}^t \mathbf{D}] \hat{\mathbf{f}} &= \mathbf{H}^t \mathbf{g} \longrightarrow \hat{\mathbf{f}} = [\mathbf{H}^t \mathbf{H} + \lambda \mathbf{D}^t \mathbf{D}]^{-1} \mathbf{H}^t \mathbf{g} \end{aligned}$$

6.6 Algorithmes de régularisation

$$\text{minimiser } J(\mathbf{f}) = Q(\mathbf{f}) + \lambda\Omega(\mathbf{f})$$

$$\text{avec } Q(\mathbf{f}) = [\mathbf{g} - \mathbf{H}\mathbf{f}]^t [\mathbf{g} - \mathbf{H}\mathbf{f}]$$

$$=$$

$$\text{minimiser } \Omega(\mathbf{f}) \text{ sous la contrainte}$$

$$e = [\mathbf{g} - \mathbf{H}\mathbf{f}]^t [\mathbf{g} - \mathbf{H}\mathbf{f}] < \epsilon$$

Information *a priori*:

- Douceur

$$\Omega(\mathbf{f}) = [\mathbf{D}\mathbf{f}]^t [\mathbf{D}\mathbf{f}] = \mathbf{f}^t \mathbf{D}^t \mathbf{D} \mathbf{f}$$

$$\hat{\mathbf{f}} = [\mathbf{H}^t \mathbf{H} + \lambda \mathbf{P}]^{-1} \mathbf{H}^t \mathbf{g} \quad \text{avec } \mathbf{P} = \mathbf{D}^t \mathbf{D}$$

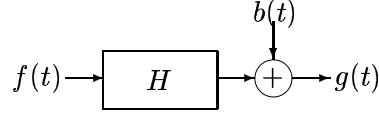
- Positivité $\Omega(\mathbf{f}) =$ fonction non quadratique de $\mathbf{f}$
Pas de solution explicite

Difficulté: Le calcul de l'expression $[\mathbf{H}^t \mathbf{H} + \lambda \mathbf{D}^t \mathbf{D}]^{-1}$

- approximation circulante
- méthodes récursives
- méthodes itératives

6.6.1 Méthode de Hunt (approximation circulante)

Pour illustrer cette méthode considérons le problème de la déconvolution des signaux :



$$g(t) = \int f(t')h(t-t') dt' + b(t) = \int h(t')f(t-t') dt' + b(t)$$

Nous avons vu que lorsqu'on discrétise cette équation on obtient

$$g(n) = \sum_{i=0}^p h(i)f(n-i) + b(n)$$

ou sous forme matricielle :

$$\mathbf{g} = \mathbf{H}\mathbf{f} + \mathbf{b}$$

où $\mathbf{H}$ est une matrice de TOEPLITZ. Notons que si on définit $\mathbf{h} = [h_0, \dots, h_p]$ on peut noter que $\mathbf{H}\mathbf{f} = \text{conv}(\mathbf{h}, \mathbf{f})$.

La solution régularisée au sens de PTT est définie par

$$\hat{\mathbf{f}} = \arg \min_{\mathbf{f}} \{J(\mathbf{f})\}$$

avec

$$\begin{aligned} J(\mathbf{f}) &= Q(\mathbf{f}) + \lambda\Omega(\mathbf{f}) \\ Q(\mathbf{f}) &= [\mathbf{g} - \mathbf{H}\mathbf{f}]^t [\mathbf{g} - \mathbf{H}\mathbf{f}] = \|\mathbf{g} - \mathbf{H}\mathbf{f}\|^2 \\ \Omega(\mathbf{f}) &= [\mathbf{D}\mathbf{f}]^t [\mathbf{D}\mathbf{f}] = \|\mathbf{D}\mathbf{f}\|^2 \end{aligned}$$

où $\mathbf{D}$ est une matrice correspondant à une opération de filtrage linéaire, donc de convolution par une réponse impulsionnelle, par exemple :

$$\mathbf{d} = [1, -2, 1], \quad \mathbf{D}\mathbf{f} = \text{conv}(\mathbf{d}, \mathbf{f})$$

et on a

$$\Omega(\mathbf{f}) = \sum_{j=1}^N (f(i+1) - 2f(i) + f(i-1))^2 = \mathbf{f}^t \mathbf{D}^t \mathbf{D} \mathbf{f}.$$

La solution régularisée a pour expression $\hat{\mathbf{f}} = [\mathbf{H}^t \mathbf{H} + \lambda \mathbf{D}^t \mathbf{D}]^{-1} \mathbf{H}^t \mathbf{g}$ et la principale difficulté est le coût de calcul de l'inverse de la matrice $[\mathbf{H}^t \mathbf{H} + \lambda \mathbf{D}^t \mathbf{D}]$. Cette matrice est symétrique, mais, dans le cas général n'a pas d'autres propriétés remarquables. Dans le cas particulier où $\mathbf{H}$ et $\mathbf{D}$ sont des matrices de TOEPLITZ, cette matrice est aussi TOEPLITZ.

L'idée de base est d'étendre les vecteurs $\mathbf{f}$, $\mathbf{h}$ et $\mathbf{g}$ par des zéros :

$$\mathbf{f}_e(i) = \begin{cases} f(i) & i = 1, \dots, Nf \\ 0 & i = Nf + 1, \dots, P \geq Nf + Nh - 1 \end{cases}$$

$$\mathbf{g}_e(i) = \begin{cases} g(i) & i = 1, \dots, Ng \\ 0 & i = Ng + 1, \dots, P \end{cases}$$

$$h_e(i) = \begin{cases} h(i) & i = 1, \dots, Nh \\ 0 & i = Nh + 1, \dots, P \end{cases}$$

pour obtenir une relation $\mathbf{g}_e = \mathbf{H}_e \mathbf{f}_e$ avec $\mathbf{H}_e$ une matrice circulante que l'on peut diagonaliser par transformée de FOURIER discrète (TFD). En effet, pour une matrice circulante $\mathbf{H}_e$ on a

$$\mathbf{H}_e = \mathbf{F} \mathbf{\Lambda} \mathbf{F}^{-1}, \quad \mathbf{\Lambda} = \text{diag}[\lambda_1, \dots, \lambda_P]$$

avec

$$\mathbf{F}[k, l] = \exp \left[j 2 \pi \frac{kl}{P} \right] \quad \mathbf{F}^{-1} = \frac{1}{P} \exp \left[-j 2 \pi \frac{kl}{P} \right]$$

$$[\lambda_1, \dots, \lambda_P] = \text{TFD} [h_1, \dots, h_{Nh}, 0, \dots, 0]$$

De même, on peut étendre le vecteur $\mathbf{d} = [1, -2, 1]$

$$d_e(i) = \begin{cases} d(i) & i = 1, \dots, 3 \\ 0 & i = 4, \dots, P \end{cases}$$

de telle sorte que l'on puisse écrire les relations suivantes :

$$\hat{\mathbf{f}} = [\mathbf{H}^t \mathbf{H} + \lambda \mathbf{D}^t \mathbf{D}]^{-1} \mathbf{H}^t \mathbf{g}$$

$$\mathbf{F} \hat{\mathbf{f}}_e = [\Lambda_h^t \Lambda_h + \lambda \Lambda_c^t \Lambda_c]^{-1} \Lambda_h^t \mathbf{F} \mathbf{g}$$

$$\text{TFD} \{\mathbf{f}_e\} = [\Lambda_h^t \Lambda_h + \lambda \Lambda_c^t \Lambda_c]^{-1} \Lambda_h^t \text{TFD} \{\mathbf{g}\}$$

$$\hat{F}(\omega) = \frac{1}{H(\omega)} \frac{|H(\omega)|^2}{|H(\omega)|^2 + \lambda |C(\omega)|^2} G(\omega)$$

Lien avec le filtre de Wiener :

$$C(\omega) = S_{bb}(\omega) / S_{ff}(\omega)$$

Restauration d'image :

D matrice de convolution par une réponse impulsionnelle

$$H_1 = \begin{pmatrix} 0 & 1 & 0 \\ 1 & -4 & 1 \\ 0 & 1 & 0 \end{pmatrix}$$

$$\Omega(\mathbf{f}) = \sum \sum (f(i+1,j) + f(i-1,j) + f(i,j+1) + f(i,j-1) - 4f(i,j))^2$$

$$f_e(k,l) = \begin{cases} f(k,l) & k = 1, \dots, K \quad l = 1, \dots, L \\ 0 & k = K+1, \dots, P \quad l = L+1, \dots, P \end{cases}$$

$$F(\omega_x, \omega_y) = \frac{1}{H(\omega_x, \omega_y)} \frac{|H(\omega_x, \omega_y)|^2}{|H(\omega_x, \omega_y)|^2 + \lambda |C(\omega_x, \omega_y)|^2} G(\omega_x, \omega_y)$$

6.6.2 Méthodes itératives

Nous avons vu que la solution régularisée est définie comme l'argument qui minimise un critère $J(\mathbf{f})$:

$$\hat{\mathbf{f}} = \arg \min_{\mathbf{f}} \{J(\mathbf{f}) = Q(\mathbf{f}) + \lambda \Omega(\mathbf{f})\}$$

Lorsque le critère est unimodal, on peut alors utiliser les algorithmes d'optimisation classique du type gradient pour calculer la solution. Notant par :

$$\mathbf{g} = \nabla J(\mathbf{f}) = \left[\frac{\partial J}{\partial f_1}, \dots, \frac{\partial J}{\partial f_n} \right]^t \quad \text{le gradient de } J, \quad \mathbf{g}_k = \nabla J(\mathbf{f}_k)$$

et par

$$\mathbf{H} = \nabla^2 J(\mathbf{f}) = \left\{ \frac{\partial^2 J}{\partial f_i \partial f_j} \right\} \quad \text{la matrice Hessienne de } J, \quad \mathbf{H}_k = \nabla^2 J(\mathbf{f}_k)$$

La majorité des méthodes d'optimisation locale calcule la solution d'une manière itérative suivante :

$$\begin{aligned} \mathbf{f}_{k+1} &= \mathbf{f}_k - \alpha_k \mathbf{D}_k \nabla J(\mathbf{f}_k) \\ &= \mathbf{f}_k - \alpha_k \mathbf{D}_k \mathbf{g}_k \\ &= \mathbf{f}_k + \alpha_k \mathbf{d}_k \end{aligned}$$

Suivant les différents choix pour α_k et $\mathbf{D}_k$ on peut distinguer les algorithmes suivants :

1. Méthodes de premier ordre ou de gradient

- gradient à pas fixe : $\mathbf{D}_k = \mathbf{I}$, $\alpha_k = \alpha$ et $\mathbf{d}_k = -\mathbf{g}_k$
- gradient à pas variable : $\mathbf{D}_k = \mathbf{I}$ et

$$\alpha_k = -\frac{\mathbf{g}^{(k)t} \mathbf{g}_k}{\mathbf{g}^{(k)t} \mathbf{H}_k \mathbf{g}_k} = \frac{\|\mathbf{g}_k\|^2}{\|\mathbf{g}_k\|_H^2}$$

- gradient conjugué :

$$\mathbf{f}_{k+1} = \mathbf{f}_k + \alpha_k \mathbf{d}_k \quad \alpha_k = -\frac{\mathbf{d}^{(k)t} \mathbf{g}_k}{\mathbf{d}^{(k)t} \mathbf{H}_k \mathbf{d}_k}$$

$$\mathbf{d}_{k+1} = \mathbf{d}_k + \beta_k \mathbf{g}_k \quad \beta_k = -\frac{\mathbf{g}^{(k)t} \mathbf{g}_k}{\mathbf{g}^{(k-1)t} \mathbf{g}^{(k-1)}}$$

2. Méthodes de second ordre

- Méthode de Newton : $\mathbf{D}_k = [\mathbf{H}_k]^{-1}$

$$\mathbf{f}_{k+1} = \mathbf{f}_k - \alpha_k [\mathbf{H}_k]^{-1} \mathbf{g}_k$$

- Méthode de Newton approchée : $\mathbf{D}_k = \text{diag}\{d_{1_k}, \dots, d_{n_k}\}$ avec

$$d_{i_k} = \mathbf{D}_{ii} = \left(\frac{\partial^2 J}{\partial f_i^2} \right)^{-1}$$

les éléments diagonaux de la matrice Hessienne.

- Méthodes de Newton modifiées :

$$\mathbf{D}_k = \mathbf{D}_0 = [\mathbf{H}_0]^{-1} = \mathbf{D}$$

ou

$$\mathbf{D}_k = \mathbf{D}_0 = \text{diag}\{d_{1_0}, \dots, d_{n_0}\} \quad \text{avec} \quad d_{i_0} = \left(\frac{\partial^2 J}{\partial f_i^2}(\mathbf{f}_0) \right)^{-1}$$

- Avantage : $\Omega(\mathbf{f})$ peut être quelconque
- Inconvénient : Coût de calcul

6.6.3 Méthodes récursives

Considérons l'équation $\mathbf{g} = \mathbf{H}\mathbf{f} + \mathbf{b}$ et sa solution régularisée :

$$\hat{\mathbf{f}} = [\mathbf{H}^t \mathbf{H} + \lambda \mathbf{D}]^{-1} \mathbf{H}^t \mathbf{g}$$

Notons par $\mathbf{H}_i$ la matrice $\mathbf{H}$ lorsqu'on a observé les données jusqu'à l'indice i ; $\mathbf{g}_i$ et par $\mathbf{H}_{i+1}$ la matrice $\mathbf{H}$ lorsqu'on a observé les données jusqu'à l'indice i ; $\mathbf{g}_{i+1}$. On a alors

$$\mathbf{g}_i = \begin{pmatrix} g_1 \\ \vdots \\ g_i \end{pmatrix} = \mathbf{H}_i \mathbf{f} + \mathbf{b},$$

$$\mathbf{g}_{i+1} = \begin{pmatrix} g_1 \\ \vdots \\ g_i \\ g_{i+1} \end{pmatrix} = \begin{pmatrix} \mathbf{g}_i \\ \dots \\ g_{i+1} \end{pmatrix} = \begin{pmatrix} \mathbf{H}_i \\ \dots \\ \mathbf{h}_{i+1}^t \end{pmatrix} \mathbf{f} + \mathbf{b}$$

L'idée de base consiste à exprimer la solution $\mathbf{f}_{i+1}$

$$\begin{aligned} \mathbf{f}_{i+1} &= (\mathbf{H}_{i+1}^t \mathbf{H}_{i+1} + \lambda \mathbf{D})^{-1} \mathbf{H}_{i+1}^t \mathbf{g}_{i+1} \\ &= \left[\begin{pmatrix} \mathbf{H}_i^t & \vdots & \mathbf{h}_{i+1}^t \end{pmatrix} \begin{pmatrix} \mathbf{H}_i \\ \dots \\ \mathbf{h}_{i+1}^t \end{pmatrix} + \lambda \mathbf{D} \right]^{-1} \begin{pmatrix} \mathbf{H}_i^t & \vdots & \mathbf{h}_{i+1}^t \end{pmatrix} \begin{pmatrix} \mathbf{g}_i \\ \dots \\ g_{i+1} \end{pmatrix} \end{aligned}$$

en fonction de $\mathbf{f}_i$:

$$\mathbf{f}_i = (\mathbf{H}_i^t \mathbf{H}_i + \lambda \mathbf{D})^{-1} \mathbf{H}_i^t \mathbf{g}_i$$

Utilisant la lemme d'inversion des matrices on obtient facilement :

$$\mathbf{f}_{i+1} = (\mathbf{H}_i^t \mathbf{H}_i + \mathbf{h}_{i+1}^t \mathbf{h}_{i+1} + \lambda \mathbf{D})^{-1} (\mathbf{H}_i^t \mathbf{g}_i + \mathbf{h}_{i+1} g_{i+1})$$

qui peut être réécrite d'une manière plus familière en notant :

$$\mathbf{P}_i = (\mathbf{H}_i^t \mathbf{H}_i + \lambda \mathbf{D})^{-1}$$

$$\mathbf{P}_{i+1}^t = \mathbf{P}_i^t + \mathbf{H}_{i+1} \mathbf{H}_{i+1}^t$$

$$\Downarrow$$

$$\mathbf{f}_{i+1} = \mathbf{f}_i + \mathbf{P}_{i+1} \mathbf{H}_{i+1} (\mathbf{g}_{i+1} - \mathbf{H}_{i+1} \mathbf{f}_i)$$

$$\mathbf{P}_{i+1} = \mathbf{P}_i - \mathbf{P}_i \mathbf{H}_{i+1} (\mathbf{H}_{i+1}^t \mathbf{P}_i \mathbf{H}_{i+1} + \lambda)^{-1} \mathbf{H}_{i+1}^t \mathbf{P}_i$$

Filtre de Kalman

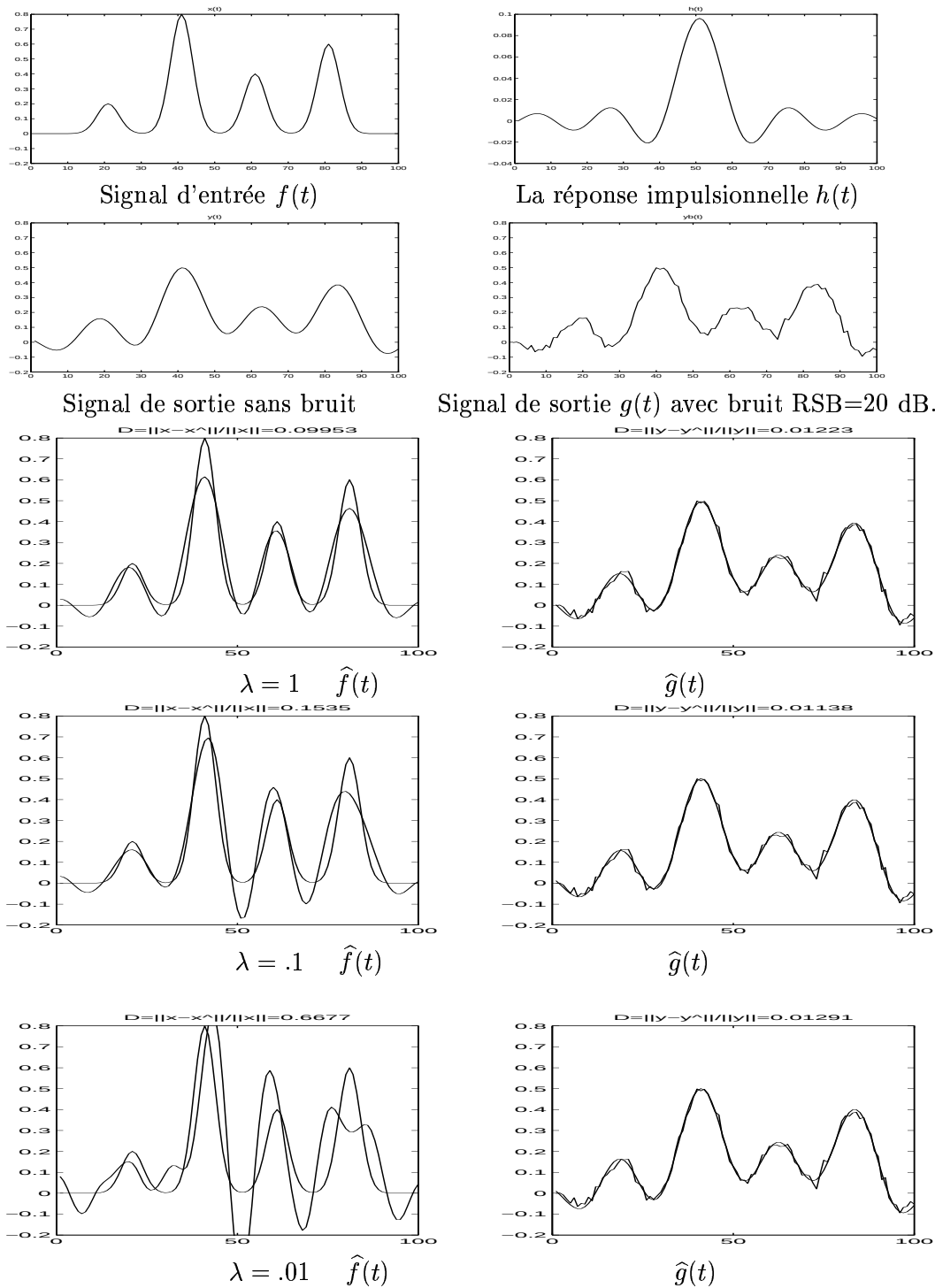


FIG. 6.1 – Déconvolution par la méthode de Hunt : Exemple 1

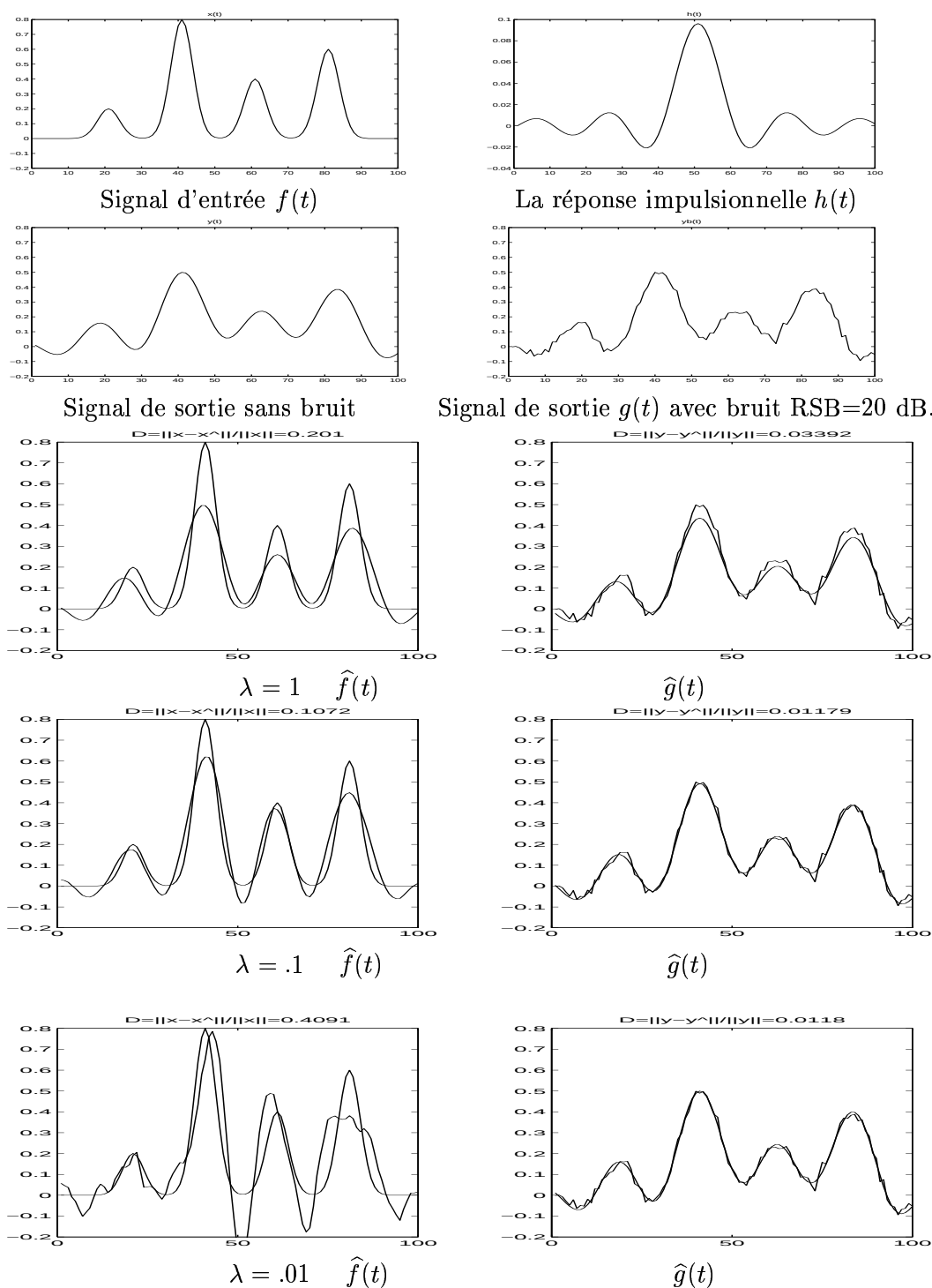


FIG. 6.2 – Déconvolution par filtre de Wiener : Exemple 1

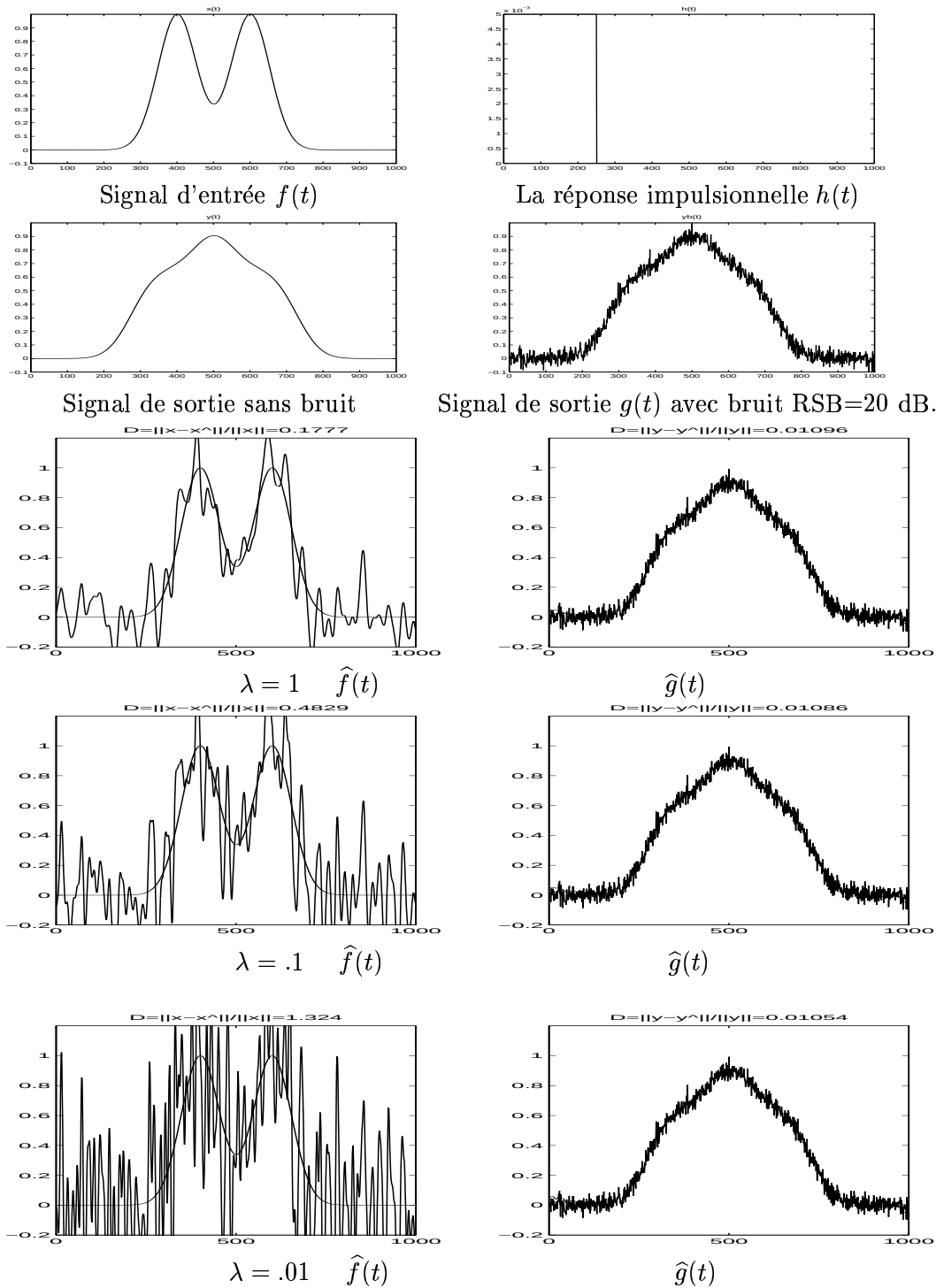


FIG. 6.3 – Déconvolution par la méthode de Hunt : Exemple 2

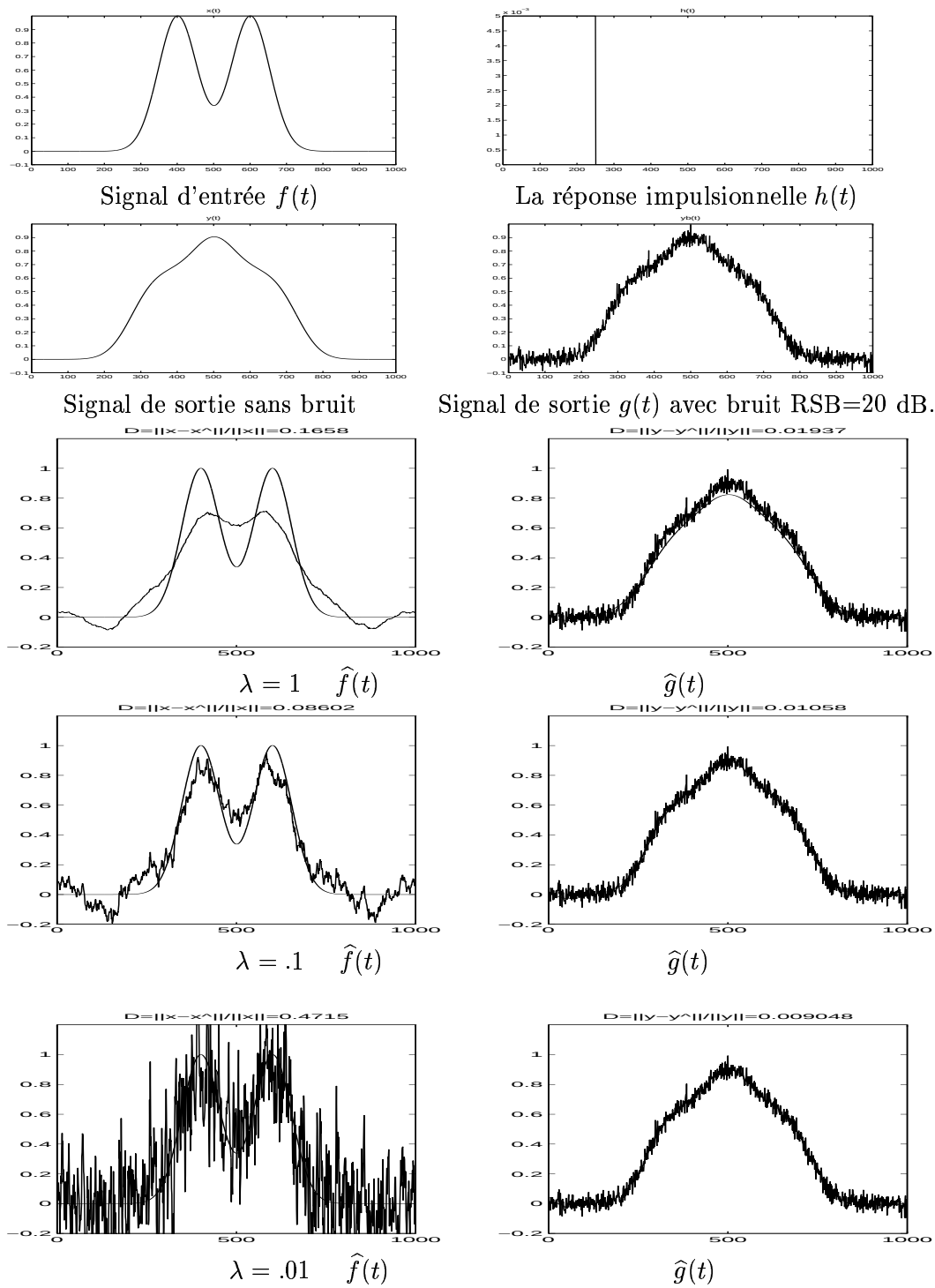


FIG. 6.4 – Déconvolution par filtre de Wiener : Exemple 2

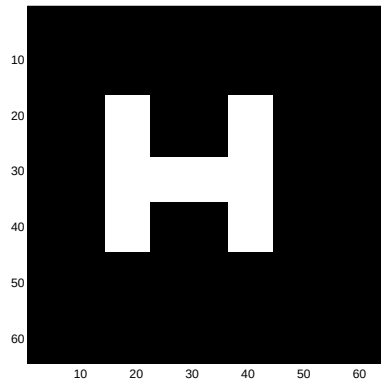
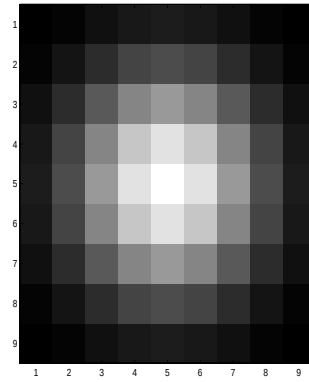
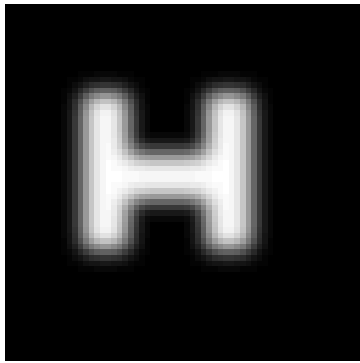
Image d'entrée $f(x,y)$ La réponse impulsionnelle $h(x,y)$ 

Image de sortie sans bruit

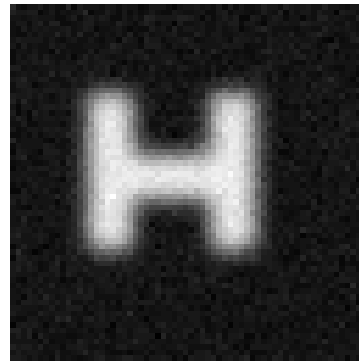
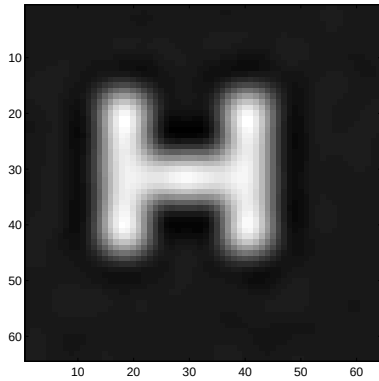
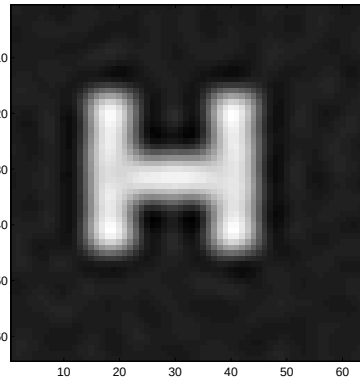
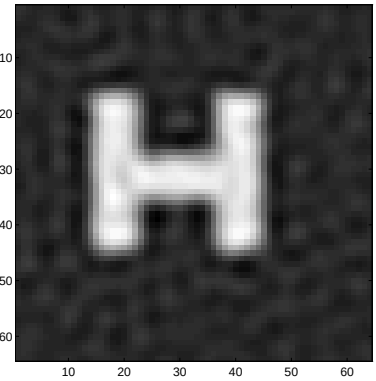
Image de sortie $g(x,y)$ avec bruit RSB=20 dB. $\lambda = 1 \quad \hat{f}(x,y)$  $\lambda = .1 \quad \hat{f}(x,y)$  $\lambda = .01 \quad \hat{f}(x,y)$

FIG. 6.5 – Restauration par la méthode de Hunt: Exemple 1

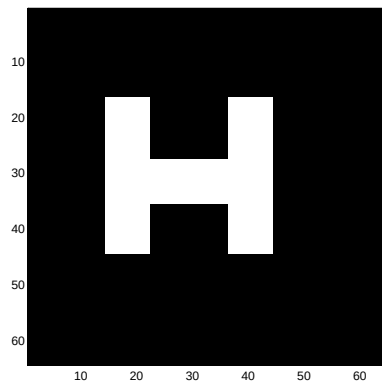
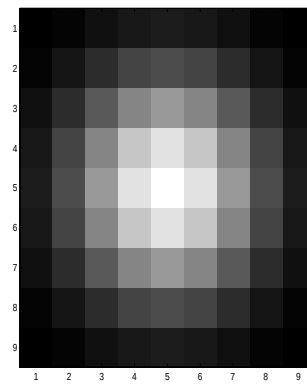
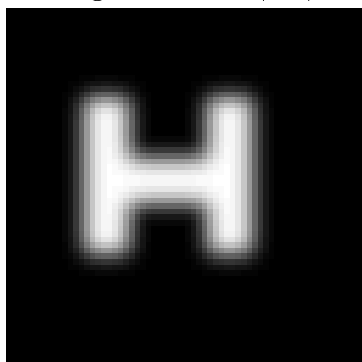
Image d'entrée $f(x,y)$ La réponse impulsionnelle $h(x,y)$ 

Image de sortie sans bruit

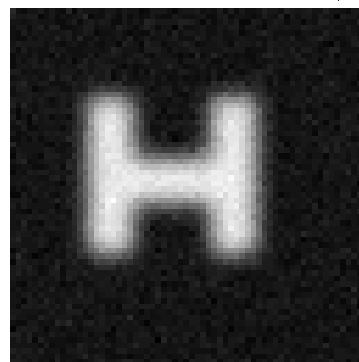
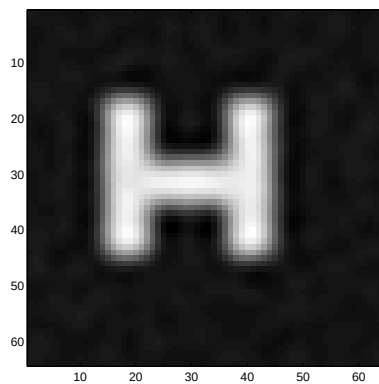
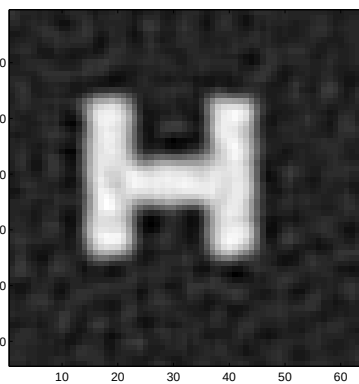
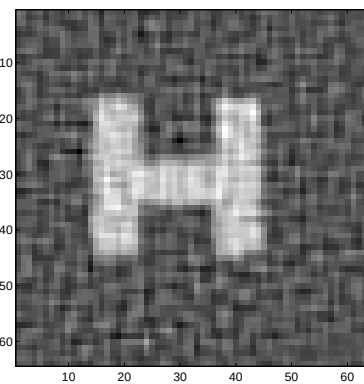
Image de sortie $g(x,y)$ avec bruit RSB=20 dB. $\lambda = 1 \quad \hat{f}(x,y)$  $\lambda = .1 \quad \hat{f}(x,y)$  $\lambda = .01 \quad \hat{f}(x,y)$

FIG. 6.6 – Restauration par filtre de Wiener : Exemple 1

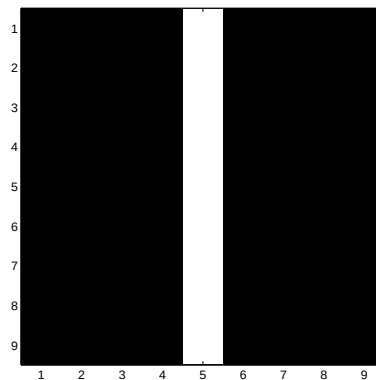
Image d'entrée $f(x,y)$ La réponse impulsionnelle $h(x,y)$ 

Image de sortie sans bruit

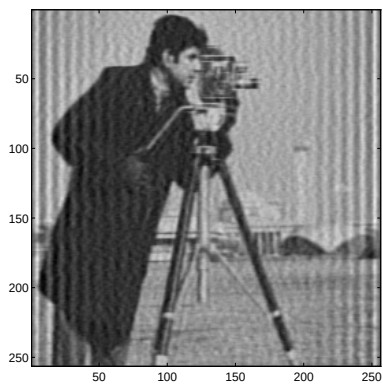
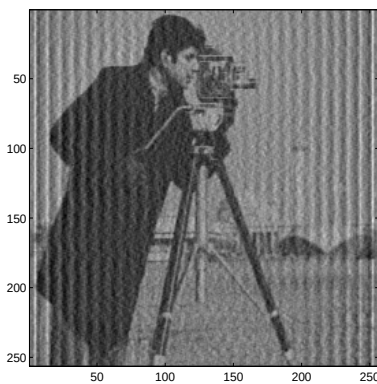
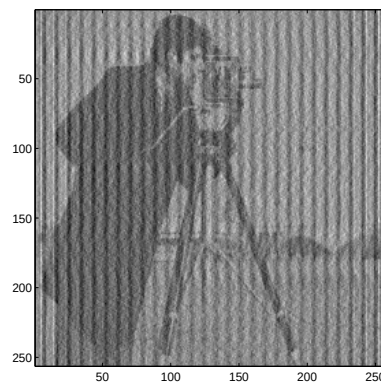
Image de sortie $g(x,y)$ avec bruit RSB=20 dB. $\lambda = 1 \quad \hat{f}(x,y)$  $\lambda = .1 \quad \hat{f}(x,y)$  $\lambda = .01 \quad \hat{f}(x,y)$

FIG. 6.7 – Restauration par la méthode de Hunt: Exemple 2

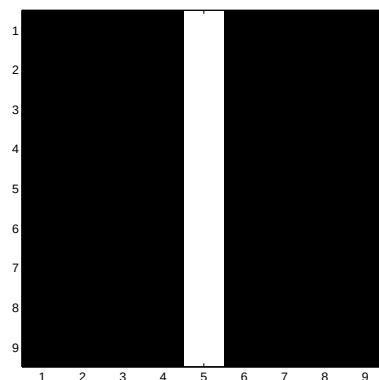
Image d'entrée $f(x,y)$ La réponse impulsionnelle $h(x,y)$ 

Image de sortie sans bruit

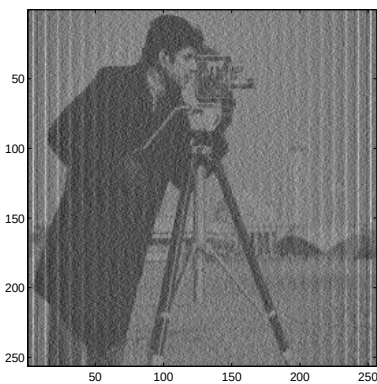
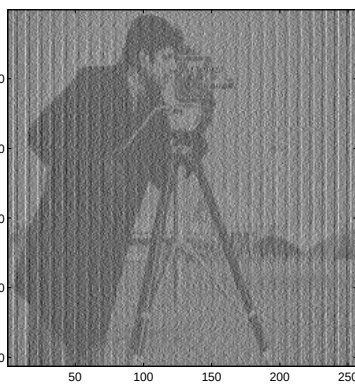
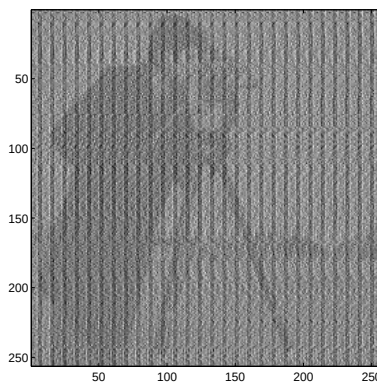
Image de sortie $g(x,y)$ avec bruit RSB=20 dB. $\lambda = 1 \quad \hat{f}(x,y)$  $\lambda = .1 \quad \hat{f}(x,y)$  $\lambda = .01 \quad \hat{f}(x,y)$

FIG. 6.8 – Restauration par filtre de Wiener : Exemple 2

Chapitre 7

Classification des méthodes d'inversion

L'objet de ce chapitre est de donner un aperçu général des différentes approches pour la résolution des problèmes inverses, ce qui nous permettra de mieux situer ces différentes méthodes. Dans le cas général, on peut distinguer deux approches : dans la première, on travaille directement dans le cadre fonctionnelle et en dimension infinie (fonctions et opérateurs) ; dans la deuxième, en revanche, dès le départ on discrétise les équations et on travaille en dimensions finie.

Dans ce chapitre, j'ai voulu faire une classification un peu plus fine en distinguant quatre approches :

- *Méthodes analytiques* où, soit directement soit après passage dans un espace dual, on peut obtenir une expression analytique explicite de la solution en fonction de la grandeur observable.
- *Méthodes de développement en série* où l'on décompose la grandeur observable et inconnue en une série de fonctions bien choisie pour pouvoir trouver une relation explicite, simple et inversible entre les coefficients de ces deux développements.
- *Méthode de décomposition sur une base* où, l'on projette la grandeur inconnue sur une base de fonctions bien choisie en fonction de l'opérateur la liant aux mesures, de telle sorte que l'on puisse obtenir une relation explicite et inversible entre les mesures et ces coefficients.
- *Méthodes algébriques* où on projette aussi la grandeur inconnue sur une base de fonctions bien choisie mais indépendantes de l'opérateur. De plus, ici on ne cherche pas forcément une relation inversible entre les mesures et ces coefficients, mais on travaille directement dans cet espace de dimension finie tout en étant conscient qu'il faut apporter de l'information *a priori* pour obtenir une solution satisfaisante au problème.

L'objet de ce chapitre est de détailler ces approches pour montrer les limites de chacune.

7.1 Méthodes analytiques

Dans ces méthodes, on fait l'hypothèse que les mesures sont des échantillons d'une fonction qui est liée à l'objet par une équation intégrale. On cherche alors une relation analytique explicite pour l'inverser. L'outil ici est l'analyse fonctionnelle. Une fois cette équation analytique trouvée, on la calcule numériquement. Evidemment, cette approche ne peut être utilisée qu'au cas par cas et en fonction de la nature de l'opérateur liant l'objet à la grandeur observable.

Pour illustrer le principe de ces méthodes considérons par exemple, le cas de la reconstruction d'image en tomographie à rayons X. Si l'on modélise le lien entre les mesures (projections) et l'objet par l'équation intégrale de la transformée de RADON $\mathcal{R}$:

$$p(r, \phi) = \iint f(x, y) \delta(r - x \cos \phi - y \sin \phi) dx dy$$

alors on a une expression analytique pour son inverse :

$$f(x, y) = \frac{1}{2\pi^2} \int_0^\pi \int_0^\infty \frac{\frac{\partial p(r, \phi)}{\partial r}}{(r - x \cos \phi - y \sin \phi)} dr d\phi$$

qui est classiquement décomposée en trois opérateurs :

$$\begin{aligned} \text{Dérivation } \mathcal{D}: \quad \bar{p}(r, \phi) &= \frac{\partial p(r, \phi)}{\partial r} \\ \text{Transformée de Hilbert } \mathcal{H}: \quad g(r', \phi) &= \frac{1}{\pi} \int_0^\infty \frac{\bar{p}(r, \phi)}{(r - r')} dr \\ \text{Rétroprojection } \mathcal{B}: \quad f(x, y) &= \frac{1}{2\pi} \int_0^\pi g(r' = x \cos \phi + y \sin \phi, \phi) d\phi \end{aligned}$$

ce qui permet d'écrire :

$$f = \mathcal{B} \mathcal{H} \mathcal{D} \mathcal{R} f$$

On peut alors s'imaginer que le problème est résolu et qu'il suffit de calculer cette expression intégrale d'une manière numérique en approximant l'équation intégrale par une somme appropriée. En effet, si l'on connaissait parfaitement $p(r, \phi)$ pour tout r et tout ϕ alors cela pourrait fournir une solution satisfaisante. Mais, en pratique on connaît $p(r, \phi)$ pour des valeurs discrètes de r et de ϕ et cela avec une précision finie. Alors, deux difficultés surviennent : la première est le calcul de la dérivée $\frac{\partial p(r, \phi)}{\partial r}$ et la seconde, l'intégration par rapport à r due à la singularité du noyau de l'intégrale.

C'est pourquoi, dans cet exemple précis de reconstruction d'image en tomographie X, bien que l'expression de l'inversion était connue depuis longtemps, rares sont les méthodes de reconstruction qui effectuent l'inversion directement de cette manière. Il existe cependant un très grand nombre de travaux basés sur cette équation, mais où l'opération de dérivation $\mathcal{D} = \frac{\partial p(r, \phi)}{\partial r}$ est approximée par un filtrage passe haut et l'opération d'intégration par rapport à la variable r qui correspond à une transformation de Hilbert $\mathcal{H}$ est aussi effectuée dans le domaine de FOURIER, et finalement, l'intégration par rapport à la variable ϕ est reconnue comme une opération de *rétroprojection* $\mathcal{B}$.

En effet, notant que la dérivation par rapport à la variable r dans le domaine spatial correspond à une multiplication par Ω dans l'espace dual et la transformation de Hilbert correspond à un changement de signe dans l'espace dual, les deux opérations peuvent s'effectuer dans cet espace par un filtre dont la fonction de transfert (réponse fréquentielle)

est $|\Omega|$. D'autre part, en dimension infinie, tous ces opérateurs sont linéaires et symétriques et peuvent commuter, ce qui permet, par exemple écrire

$$f = \mathcal{B} \mathcal{H} \mathcal{D} \mathcal{R} f = \mathcal{H} \mathcal{D} \mathcal{B} \mathcal{R} f$$

Ainsi, on peut obtenir d'autres implantations algorithmiques pour calculer cette solution. Pour pouvoir résumer les différentes méthodes qui peuvent être utilisées pour le calcul de la solution de la TR inverse, notons les opérateurs suivants :

Transformée de RADON $\mathcal{R}$:	$p(r, \phi) = \iint f(x, y) \delta(r - x \cos \phi - y \sin \phi) dx dy$
Dérivation $\mathcal{D}$:	$\bar{p}(r, \phi) = \frac{\partial p(r, \phi)}{\partial r}$
Transformée de Hilbert $\mathcal{H}$:	$g(r', \phi) = \frac{1}{\pi} \int_0^\infty \frac{\bar{p}(r, \phi)}{(r - r')} dr$
Rétroprojection $\mathcal{B}$:	$f(x, y) = \frac{1}{2\pi} \int_0^\pi g(r' = x \cos \phi + y \sin \phi, \phi) d\phi$
Transformée de FOURIER 1D $\mathcal{F}_1$:	$P(\Omega, \phi) = \int p(r, \phi) \exp[-j\Omega r] dr$
TF inverse 1D $\mathcal{F}_1^{-1}$:	$p(r, \phi) = \int P(\Omega, \phi) \exp[+j\Omega r] d\Omega$
TF 2D $\mathcal{F}_2$:	$F(\omega_x, \omega_y) = \iint f(x, y) \exp[-j(\omega_x x + \omega_y y)] dx dy$
TF inverse 2D $\mathcal{F}_2^{-1}$:	$f(x, y) = \iint F(\omega_x, \omega_y) \exp[-j(\omega_x x + \omega_y y)] d\omega_x d\omega_y$

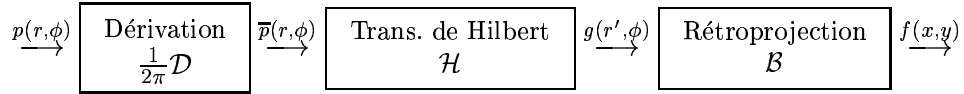
On peut alors résumer les différents algorithmes qui peuvent être utilisés pour le calcul de la solution de la TR inverse :

Malheureusement, bien que ces méthodes aient été utilisées avec succès dans certaines applications, en particulier en imagerie médicale, leur utilisation pratique reste plus limitée dans d'autres applications par exemple en contrôle non destructif. En effet, souvent en imagerie médicale on peut tourner autour de l'objet et on peut disposer d'un très grand nombre de mesures réparties d'une manière uniforme autour de l'objet, mais dans les applications telles qu'en contrôle non destructif, les mesures ne sont pas équiréparties sur l'espace de la transformée, et elles sont très limitées en nombre et entachées de bruit. Le résultat est que : les solutions analytiques ainsi calculées sont rarement satisfaisantes.

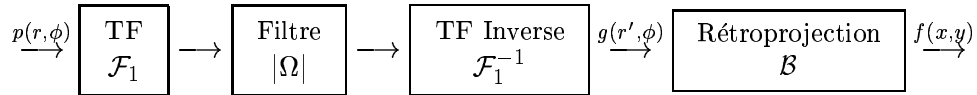
Les figures de la page suivante illustrent ces difficultés.

De plus, modéliser la relation entre les données et les inconnues par une transformation telle que l'on puisse trouver une transformation inverse analytique nécessite souvent des approximations qui deviennent rapidement assez irréalistes. Par exemple, vouloir prendre en compte la largeur du détecteur dans cet exemple met en cause le modèle simple de la transformée de RADON entre les projections et l'objet.

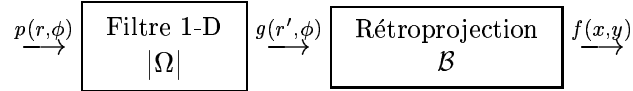
- Inversion directe :



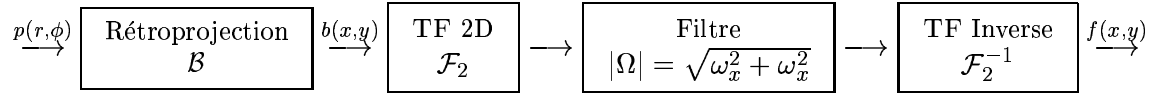
- Rétroprojection des projections filtrées :



- Filtrage par convolution et rétroprojection :



- Rétroprojection suivie de filtrage en 2D :



- Rétroprojection suivie d'un filtrage par convolution en 2D :

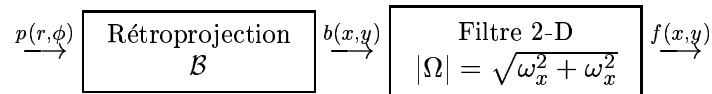


FIG. 7.1 – Méthodes analytiques pour l'inversion de la TR.

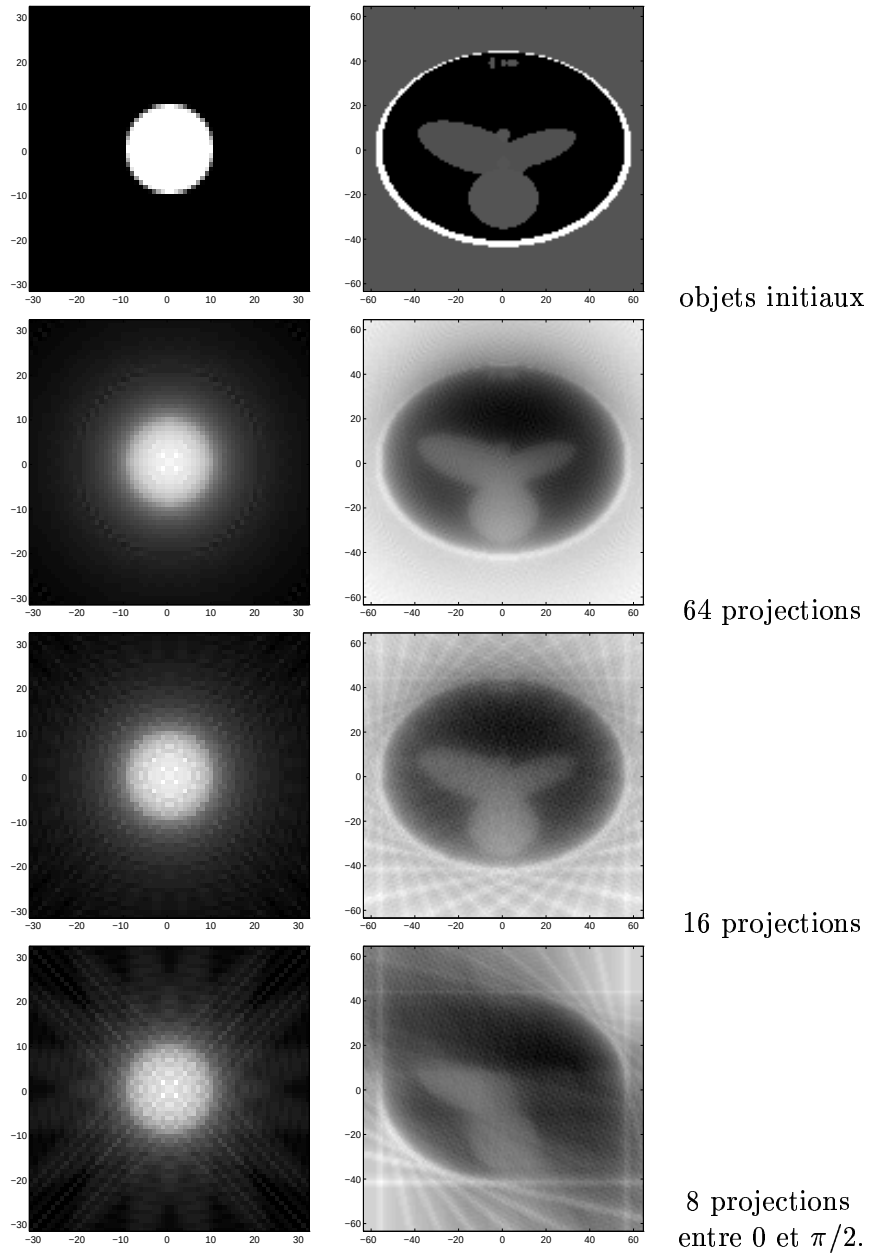


FIG. 7.2 – *Reconstruction par rétroprojection filtrée : Trois cas sont considérés : 64 projections, 16 projections réparties entre 0 et π et 8 projections réparties entre 0 et $\pi/2$.*

7.2 Méthodes de développement en série

Dans ces méthodes, on développe en série les fonctions $f(\mathbf{r})$ et $g(\mathbf{u})$

$$f(\mathbf{r}) = \sum_{k=1}^{\infty} f_k b_k(\mathbf{r}), \quad g(\mathbf{u}) = \sum_{k=1}^{\infty} g_k b_k(\mathbf{u}),$$

en choisissant les fonctions b_k telles que l'on puisse trouver une relation analytique simple et inversible entre les coefficients de ces développements. Ensuite, à partir des mesures $g(\mathbf{u}_i), i = 1, \dots, K$ on détermine K coefficients $g_k, k = 1, \dots, K$ et on en déduit les K coefficients $f_k, k = 1, \dots, K$. Finalement, on reconstruit $\hat{f}(\mathbf{r})$ par

$$\hat{f}(\mathbf{r}) = \sum_{k=1}^K f_k b_k(\mathbf{r}).$$

Le choix des $b_k(\mathbf{r})$ est crucial pour la mise en œuvre de cette méthode. Il dépend du contexte.

Pour illustrer le principe de ces méthodes nous allons considérer deux exemples : le premier est le problème de la déconvolution et le deuxième encore celui de la reconstruction d'image en tomographie à rayons X.

Déconvolution :

Considérons le problème de la déconvolution des signaux :

$$g(r) = \int f(r') h(r, r') \, dr'$$

Sachant que les fonctions exponentielles sont les fonctions propres d'un tel opérateur, on décide de développer les deux fonctions $f(r)$ et $g(r)$ en une série de ces fonctions :

$$\begin{aligned} f(r) &= \sum_k f_k \exp[-jk\omega_0 r] \\ g(r) &= \sum_k g_k \exp[-jk\omega_0 r] \end{aligned}$$

Il est alors facile de montrer qu'il y a une relation très facile entre les coefficients f_k et g_k qui est :

$$g_k = H_k(\omega_0) f_k$$

où

$$H_k(\omega_0) = \int h(r) \exp[+jk\omega_0 r] \, dr$$

On peut alors envisager une méthode d'inversion en trois étapes :

1. À partir des mesures $g(r)$ on estime les K coefficients g_k par

$$g_k = \frac{1}{2\pi} \int g(r) \exp[-jk\omega_0 r] \, dr, \quad k = 1 \dots, K$$

2. On calcule les K coefficients f_k en utilisant

$$f_k = \frac{g_k}{H_k(\omega_0)}$$

3. On calcule la solution par

$$f(r) = \sum_k^K f_k \exp[-jk\omega_0 r]$$

On peut facilement voir que lorsque $\omega_0 = \frac{2\pi}{T}$ avec T la période d'échantillonnage des signaux $f(r)$, $g(r)$ et $h(r)$, l'on retrouve la notion de filtrage inverse. En effet, f_k , g_k sont, respectivement, les coefficients de la décomposition en série de FOURIER des fonctions $f(r)$ et $g(r)$ et $H_k(\omega_0)$ est la réponse fréquentielle du système. La principale difficulté, comme on le sait, se trouve dans l'étape 2 lorsque certains coefficients $H_k(\omega_0)$ deviennent proche de zéro (amplification du bruit).

Reconstruction d'image en tomographie à rayons X :

Nous avons vu que la relation entre un objet $f(x, y)$ et ses projections $p(r, \phi)$ en tomographie à rayons X est donnée par la TR. Pour illustrer l'utilisation de la méthode de développement en série dans ce cas, écrivons la TR en coordonnées polaires (ρ, θ) :

$$p(r, \phi) = \int_0^\infty \int_0^{2\pi} f(\rho, \theta) \delta(r - \rho \cos(\phi - \theta)) \rho d\rho d\theta$$

Notant que $p(r, \phi)$ est périodique en ϕ et $f(\rho, \theta)$ est périodique en θ , on peut les développer en série de FOURIER :

$$p(r, \phi) = \sum_k p_k(r) \exp[-jk\phi], \quad f(\rho, \theta) = \sum_k f_k(\rho) \exp[-jk\theta],$$

et on peut établir les relations suivantes entre les coefficients $p_k(r)$ et $f_k(\rho)$:

$$\begin{aligned} p_k(r) &= 2 \int_r^\infty f_k(\rho) T_k\left(\frac{r}{\rho}\right) \left(1 - \frac{r^2}{\rho^2}\right)^{-1/2} d\rho \\ g_k(\rho) &= \frac{-1}{\pi\rho} \int_\rho^\infty \frac{\partial p_k(r)}{\partial r} T_k\left(\frac{r}{\rho}\right) \left(\frac{r^2}{\rho^2} - 1\right)^{-1/2} dr \end{aligned}$$

avec $T_k(t) = \cos[k \cos^{-1} t]$.

On peut alors envisager une méthode d'inversion en trois étapes :

1. À partir des mesures $p(r, \phi)$ on estime les K coefficients $p_k(r)$ par

$$p_k(r) = \frac{1}{2\pi} \int_{-\pi}^{\pi} p(r, \phi) \exp[-jk\phi] d\phi, \quad k = 1 \dots, K$$

2. On calcule les K coefficients $f_k(\rho)$ en utilisant

$$g_k(\rho) = \frac{-1}{\pi\rho} \int_\rho^\infty \frac{\partial p_k(r)}{\partial r} T_k\left(\frac{r}{\rho}\right) \left(\frac{r^2}{\rho^2} - 1\right)^{-1/2} dr, \quad k = 1 \dots, K$$

3. On calcule la solution par

$$f(\rho, \theta) = \sum_k^K f_k(\rho) \exp[-jk\theta]$$

Notez que l'étape la plus délicate dans ce cas est la deuxième, car elle nécessite le calcul d'une intégrale avec un noyau ayant des singularités.

Dans le cas général, malheureusement, ces méthodes ont aussi une utilisation pratique limitée car :

1. Dans la plupart des applications réelles on ne dispose pas d'expression analytique pour le noyau de l'équation intégrale. De plus on n'a pas, en général, de symétrie de telle sorte que l'on puisse choisir les fonctions $b_k(\mathbf{r})$ satisfaisant à une relation analytique, simple et inversible avec g_k ;
2. La solution ainsi obtenue est très sensible d'une part au choix du seuil de troncature K de la série et d'autre part à la manière d'approximer le calcul de l'intégrale qui nous donne ces coefficients ;
3. La solution ainsi obtenue peut être très sensible aux erreurs de modélisation comme par exemple l'erreur géométrique dans les applications d'imagerie.

7.3 Méthodes de décomposition sur une base

Dans ces méthodes, on fait l'hypothèse que la solution du problème, c'est à dire la fonction objet, peut être décomposée en un nombre fini de fonctions de base $b_j(\mathbf{r})$:

$$f(\mathbf{r}) = \sum_{n=1}^N f_n b_n(\mathbf{r})$$

Ces fonctions sont choisies en rapport avec les propriétés du noyau $h(\mathbf{r}, \mathbf{r}')$ de l'équation intégrale, de telle sorte qu'une fois la fonction objet décomposée sur cette base et remplacée dans l'équation d'observation :

$$\begin{aligned} g(\mathbf{r}_m) &= \int f(\mathbf{r}') h(\mathbf{r}_m, \mathbf{r}') d\mathbf{r}' = \int \sum_{n=1}^N f_n b_n(\mathbf{r}') h(\mathbf{r}_m, \mathbf{r}') d\mathbf{r}' \\ &= \sum_{n=1}^N f_n \int b_n(\mathbf{r}') h(\mathbf{r}_m, \mathbf{r}') d\mathbf{r}' = \sum_{n=1}^N f_n H_{ij}, \quad i = 1, \dots, M, \end{aligned}$$

on puisse calculer analytiquement les coefficients H_{ij} .

Notons que cette relation peut s'écrire sous la forme :

$$\mathbf{g} = \mathbf{H} \mathbf{f}$$

où $\mathbf{H}$ est une matrice avec les éléments H_{ij} , $\mathbf{g} = [g_1, \dots, g_M]$ et $\mathbf{f} = [f_1, \dots, f_N]$.

Ensuite, deux approches sont courantes :

1. Choisir $N = M$ et espérer inverser directement la matrice $\mathbf{H}$ et résoudre ainsi ce système d'équations pour obtenir les coefficients de la base f_j ;
2. Choisir $N < M$ et essayer de résoudre ce système au sens des moindres carrés ou au sens de l'inversion généralisée.

Pour illustrer cette approche, considérons un exemple de problème inverse que l'on rencontre en imagerie par microscopie électronique pour imager la distribution $\rho(r)$ de la densité de la charge électrique du noyau d'un atome en mesurant une quantité $F_c(q)$ appelée le facteur de forme. Le lien entre ces deux quantités est décrit par la relation intégrale suivante :

$$F_c(q) = 4\pi \int_0^\infty r^2 J_0(qr) \rho(r) dr \quad (7.1)$$

où J_0 est la fonction de Bessel d'ordre zéro et q est la valeur absolue du moment de transfert. Lors de l'expérimentation on peut mesurer $F_c(q)$ pour un nombre très réduit de $\mathbf{q} = \{q_1, \dots, q_M\}$ et l'objectif de l'imagerie est de déterminer la fonction $\rho(r)$.

Dans cet exemple, supposant que $\rho(r)$ s'annule au delà d'un rayon R_c , un choix judicieux de fonctions de base est :

$$\rho(r) = \begin{cases} \sum_{n=1}^N f_n J_0(q_n r) & r \leq R_c \\ 0 & r > R_c \end{cases} \quad (7.2)$$

Avec ce choix, et en utilisant une des propriétés d'orthogonalité des fonctions de Bessel :

$$\int_0^{R_c} r^2 J_l(q_n r) J_l(q_m r) dr = \frac{R_c^3}{2} J_{l+1}^2(q_n R_c) \delta_{n,m} \quad (7.3)$$

il n'est pas difficile de montrer que :

$$\mathbf{F}_c = \mathbf{A}\mathbf{f} \quad (7.4)$$

avec $\mathbf{q} = \{q_1, \dots, q_M\}$, $\mathbf{F}_c = \{F_c(q_1), \dots, F_c(q_M)\}$, $\mathbf{f} = \{f_1, \dots, f_N\}$ et

$$A_{mn} = \frac{4\pi R_c^2}{q_m} \frac{(-1)^n}{(q_m R_c)^2 - (n\pi)^2} \sin(q_m R_c). \quad (7.5)$$

En effet, remplaçant (7.2) dans (7.1) on a

$$\begin{aligned} F_c(q_m) &= 4\pi \int_0^{R_c} r^2 J_0(q_m r) \sum_{n=1}^N f_n J_0(q_n r) dr \\ &= 4\pi \sum_{n=1}^N f_n \int_0^{R_c} r^2 J_0(q_m r) J_0(q_n r) dr \\ &= \sum_{n=1}^N f_n A_{mn} \end{aligned}$$

avec A_{mn} définie par (7.5).

Notons que, si $M = N$ et si les mesures pouvaient être faites exactement aux points $q_n = \frac{n\pi}{R_c}$, alors la matrice $\mathbf{A}$ serait une matrice diagonale avec les éléments diagonaux $A_{nn} = \frac{R_c^3}{2} J_1^2(q_n R_c)$. On pourrait alors obtenir directement les coefficients f_n par

$$f_n = \frac{F_c(q_n)}{2\pi R_c^3 J_1^2(q_n R_c)}. \quad (7.6)$$

Ceci est due au fait que si $q_n = \frac{n\pi}{R_c}$ la base choisie est une base orthogonale.

En général cependant, les mesures sont faites aux points qui ne sont pas équirépartis et la matrice $\mathbf{A}$ n'est plus diagonale. De plus, il n'y a aucune raison de choisir $M = N$, car on peut toujours essayer d'estimer $\mathbf{f}$ au moins par MC :

$$\begin{aligned} \hat{\mathbf{f}} &= (\mathbf{A}^t \mathbf{A})^{-1} \mathbf{A}^t \mathbf{F}_c \quad \text{si } M > N \\ &= \mathbf{A}^t (\mathbf{A} \mathbf{A}^t)^{-1} \mathbf{F}_c \quad \text{si } M < N \end{aligned}$$

ou par MC de norme minimale (MCNM) :

$$\begin{aligned} \hat{\mathbf{f}} &= (\mathbf{A}^t \mathbf{A} + \lambda \mathbf{I})^{-1} \mathbf{A}^t \mathbf{F}_c \quad \text{si } M > N \\ &= \mathbf{A}^t (\mathbf{A} \mathbf{A}^t + \lambda \mathbf{I})^{-1} \mathbf{F}_c \quad \text{si } M < N \end{aligned}$$

ou encore par d'autres méthodes.

Les figures qui suivent montrent quelques exemples de simulation pour ce problème.

Ces méthodes ont une portée un peu plus générale que les méthodes analytiques précédentes; elles peuvent être utilisées lorsqu'on a une expression analytique du noyau de l'équation intégrale qui lie les mesures à l'objet et lorsqu'il y a une base convenable telle que la fonction désirée puisse se décomposer sur cette base avec un nombre faible de coefficients.

Cependant, ces méthodes aussi ont une utilisation pratique limitée car :

1. dans la plupart des applications réelles on ne dispose pas d'expression analytique pour le noyau de l'équation intégrale (prenez par exemple le cas de la déconvolution des signaux ou celui de la restauration d'image);
2. Il est assez difficile de trouver une base orthogonale et complète;
3. La solution ainsi obtenue est, en général, très sensible au choix du seuil de troncature N de la série.
4. Lorsque N est grand, la matrice $\mathbf{A}$ peut être très mal conditionnée.
5. Lorsqu'on a une information *a priori* sur la solution (par exemple la positivité) il est difficile de transmettre cette information sur les coefficients.

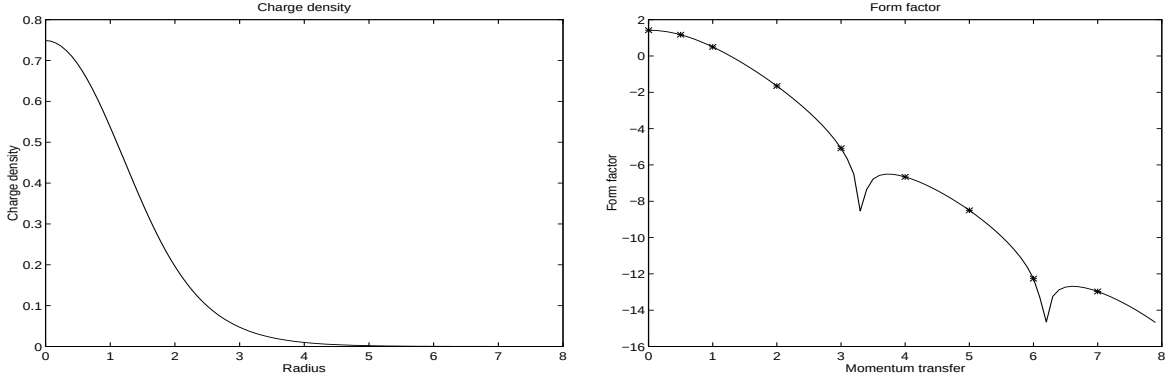


Figure a : La densité de charge théorique $\rho(r)$ [gauche], le facteur de forme théorique $\log \|F_c(q)\|$ et les données utilisées pour les simulations [droite].

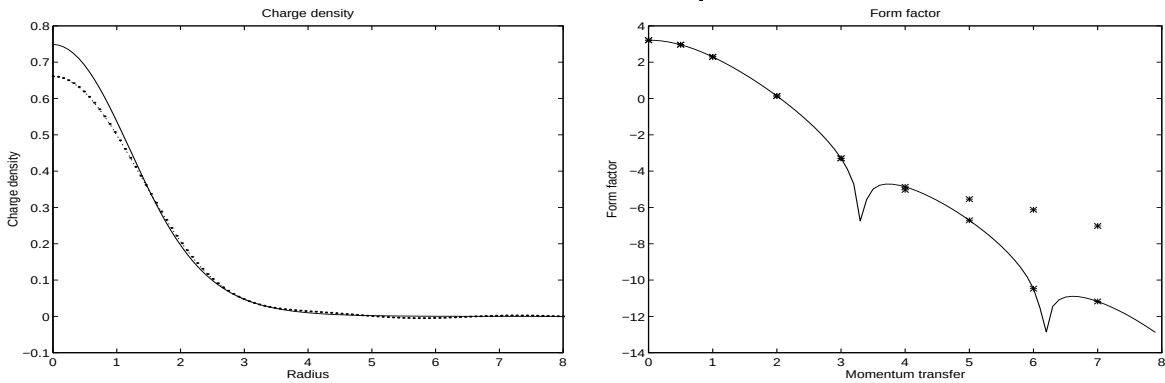


Figure b. Résultat de reconstruction de $\hat{\rho}(r)$ obtenus par MC (points) et par MCNM (tirets) pour $N = 5$ [gauche] et les facteurs de formes correspondants $\log \|\hat{F}_c(q)\|$ [droite]. Dans ce cas les deux solutions sont pratiquement identiques car la matrice $\mathbf{A}^t \mathbf{A}$ est bien conditionnée.

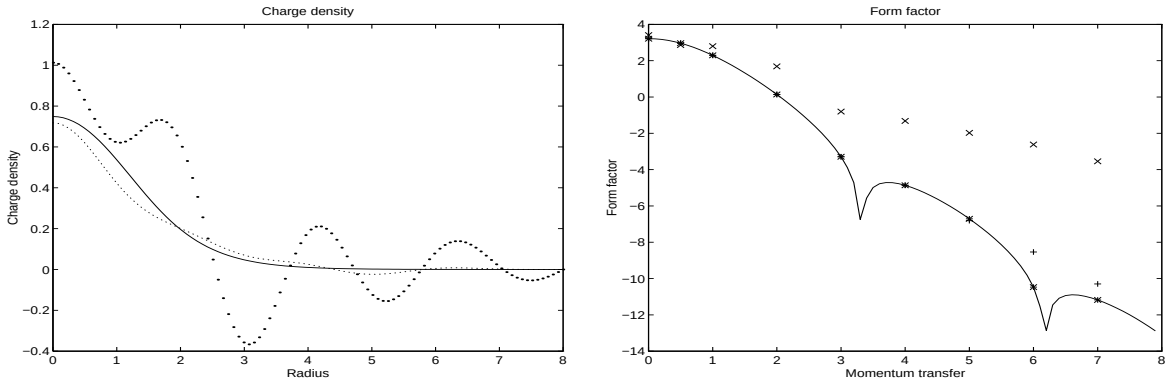


Figure c. Résultat de reconstruction de $\hat{\rho}(r)$ obtenus par MC (points) et par MCNM (tirets) pour $N = 10$ [gauche] et les facteurs de formes correspondants $\log \|\hat{F}_c(q)\|$ [droite]. Dans ce cas la matrice $\mathbf{A} \mathbf{A}^t$ est mal conditionnée, alors que la matrice $\mathbf{A} \mathbf{A}^t + \lambda \mathbf{I}$ est bien conditionnée. C'est pourquoi, la solution MC n'est plus satisfaisante, alors que la solution MCNM est encore satisfaisante.

FIG. 7.3 – Illustration de la méthode de décomposition sur une base.

7.4 Méthodes algébriques déterministes

Dans ces méthodes on commence par discrétiser la fonction objet $f(\mathbf{r})$ et par conséquent l'équation intégrale avec la résolution maximale souhaitée. Une manière de faire cette discrétisation est, comme dans le cas précédent, de décomposer la fonction $f(\mathbf{r})$ sur une base :

$$f(\mathbf{r}) = \sum_{n=1}^N f_n b_n(\mathbf{r}),$$

sauf qu'ici les fonctions de base sont choisies parmi

$$b_n(\mathbf{r}) = \delta(\mathbf{r} - n\mathbf{\Delta}) \quad (7.7)$$

ou

$$b_n(\mathbf{r}) = \begin{cases} 1 & \text{si } \mathbf{r} \in [(n-1)\mathbf{\Delta}, n\mathbf{\Delta}] \\ 0 & \text{ailleurs} \end{cases} \quad (7.8)$$

ou encore (pour un échantillonnage non régulier)

$$b_n(\mathbf{r}) = \begin{cases} 1 & \text{si } \mathbf{r} \in D_n \\ 0 & \text{ailleurs} \end{cases} \quad (7.9)$$

avec D_n un ensemble de régions dans le domaine de la fonction telle que $D_i \cap D_j = 0, \forall i \neq j$ et $\cup D_j = S$.

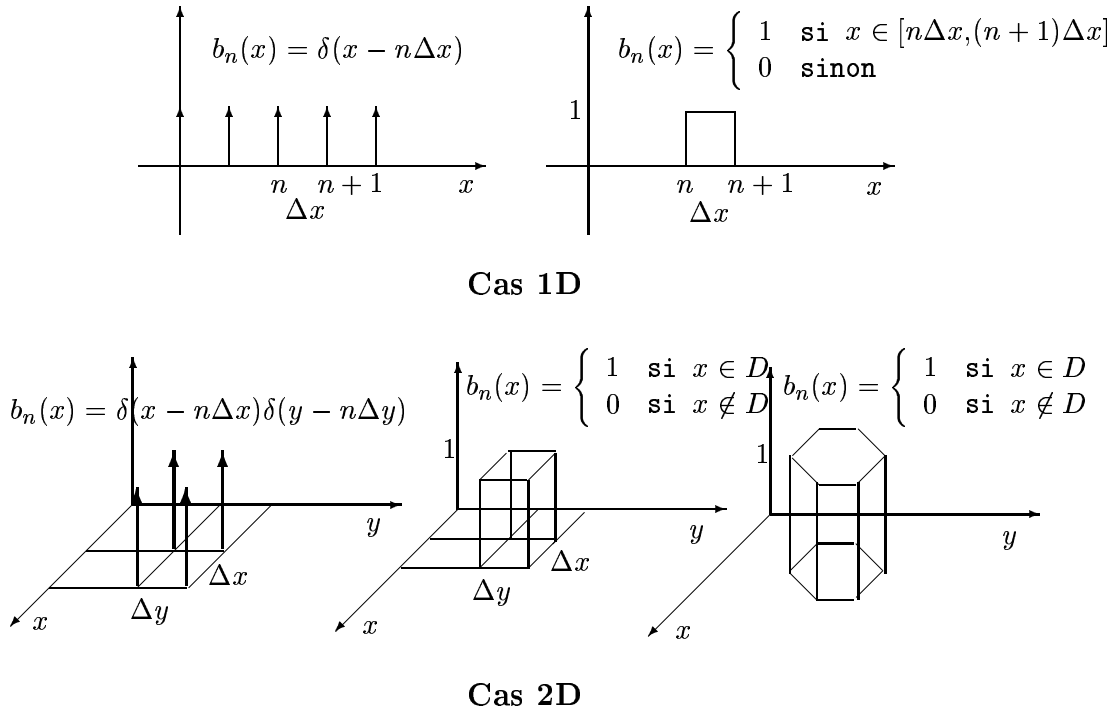


FIG. 7.4 – *chantillonnage et fonctions bases.*

Les pas d'échantillonnage $\mathbf{\Delta} = \{\Delta_1, \Delta_2, \dots\}$ sont choisis en faisant un compromis entre la résolution maximale souhaitée et le coût de calcul.

On obtient ainsi, comme dans le cas précédent, un système d'équations linéaires de la forme :

$$\mathbf{g} = \mathbf{H} \mathbf{f},$$

où $\mathbf{f}$ est un vecteur de dimension N contenant les coefficients f_n . La différence essentielle est dans la signification de ces coefficients. En choisissant les fonctions de base (7.7), ces coefficients représentent les échantillons de la fonction f , alors qu'en choisissant les fonctions de base (7.8) ou (7.9), ces coefficients représentent les valeurs moyenne de la fonction f dans des régions $[(n-1)\Delta, n\Delta]$ ou D_j .

Ainsi, contrairement aux méthodes décrites précédemment, dans ces méthodes cette discrétisation se fait sans trop se soucier du choix des fonctions de base, du nombre de ces fonctions (on suppose en général que l'objet est à support borné et on choisit un pas de discrétisation de l'espace le plus petit correspondant à la résolution souhaitée), du mauvais conditionnement de la matrice $\mathbf{H}$ et des erreurs de discrétisation. En effet, la résolution de ces problèmes est laissée à la charge de la méthode de régularisation qui suit.

Il serait en effet illusoire de vouloir se ramener à une équation de la forme $\mathbf{g} = \mathbf{H} \mathbf{f}$ avec $\mathbf{H}$ une matrice carrée et inversible et vouloir proposer la solution $\hat{\mathbf{f}} = \mathbf{H}^{-1} \mathbf{g}$.

Il serait de même illusoire de se contenter des solutions du type pseudo-solutions

$$\hat{\mathbf{f}} = \arg \min_{\mathbf{f}} \{ \Delta(\mathbf{g}, \mathbf{H} \mathbf{f}) \}$$

où $\Delta(.,.)$ est une mesure de distance dans l'espace des mesures, ou les solutions au sens des moindres carrés :

$$\hat{\mathbf{f}} = \arg \min_{\mathbf{f}} \{ \|\mathbf{g} - \mathbf{H} \mathbf{f}\|^2 \}$$

qui satisfont l'équation normale

$$\mathbf{H}^t \mathbf{H} \hat{\mathbf{f}} = \mathbf{H}^t \mathbf{g}$$

ou encore, la solution inverse généralisée (IG) qui est la solution de norme minimale de cette équation normale. Ces solutions restent très sensibles aux erreurs de discrétisation et de mesure $\mathbf{b}$. Il faut introduire une information *a priori* de manière explicite sur la solution et sur les erreurs de mesure ; il faut régulariser.

Parmi les méthodes de régularisation déterministes, on peut citer les méthodes basées sur la décomposition tronquée des valeurs singulières, les méthodes itératives de minimisation d'un critère du type $J(\mathbf{f}) = \|\mathbf{g} - \mathbf{H} \mathbf{f}\|^2$, avec un nombre limité d'itérations ou avec des contraintes supplémentaires du type positivité par exemple, ou encore des méthodes qui minimisent un critère mixte du type

$$J(\mathbf{f}) = \|\mathbf{g} - \mathbf{H} \mathbf{f}\|^2 + \lambda H(\mathbf{f}),$$

où $H(\mathbf{f})$ est une fonctionnelle régularisante et λ est le coefficient de régularisation.

De manière encore plus générale on peut définir la solution régularisée comme celle qui minimise un critère de la forme

$$J(\mathbf{f}) = \Delta_1(\mathbf{g}, \mathbf{H} \mathbf{f}) + \lambda \Delta_2(\mathbf{f}, \mathbf{f}_0),$$

où Δ_1 et Δ_2 sont deux distances et $\mathbf{f}_0$ est une solution par défaut (*a priori*).

Parmi ces méthodes, la dernière a été utilisée avec succès dans un grand nombre d'applications. Cependant elle souffre des limitations suivantes :

1. le choix de la fonctionnelle $H(\mathbf{f})$ est assez arbitraire et on manque d'arguments pour le choix du paramètre de régularisation λ ;

2. le bruit $\mathbf{b}$ est pris en compte mais de manière implicite ; la prise en compte explicite des caractéristiques du bruit n'est pas aisée ;
3. il n'est pas aisé non plus de prendre en compte explicitement des informations statistiques sur la solution.

Pour remédier à ces difficultés il n'y a pas d'autres solutions que de considérer la résolution du problème inverse comme un problème d'estimation probabiliste.

7.5 Méthodes probabilistes

L'idée de base dans ces méthodes est de dire que lorsqu'on cherche à résoudre d'une manière satisfaisante un problème inverse il faut

- prendre en compte explicitement les erreurs : non seulement celles de la discrétisation mais aussi celles due aux mesures et celles due à modélisation ;
- prendre en compte les informations *a priori* ; et
- caractériser l'incertitude qui reste dans la solution proposée.

L'outil mathématique *théorie des probabilités* est alors l'outil idéal. On utilise cet outil pour représenter l'incertitude sur les mesures et/ou le manque de connaissance parfaite sur une grandeur par une loi de probabilité, et ensuite on utilise la notion d'estimation pour déterminer les grandeurs inconnues. On peut distinguer alors plusieurs approches ou classes de méthodes :

1. L'approche ou les méthodes qui n'utilisent que les moments – jusqu'à l'ordre deux, ce qui implicitement revient à faire l'hypothèse gaussienne – d'ordre plus élevé. On cherche à lier les moments de la grandeur observée aux moments de la grandeur inconnue afin de les déterminer en inversant ce lien. On peut citer en exemple le filtrage optimal en restauration des signaux. Le principal avantage de ces méthodes est la facilité de mise en œuvre et le fait que la formulation du problème peut se faire en continu ou en discret.
2. L'approche estimation *statistique classique* où l'on prend en compte explicitement le caractère incertain des mesures en la caractérisant par une loi de probabilité conditionnelle $p(\mathbf{g}|\mathbf{f})$. On considère ensuite $p(\mathbf{g}|\mathbf{f})$ comme une fonction de $\mathbf{f}$ que l'on nomme la fonction *vraisemblance* $V(\mathbf{f})$. On définit alors un critère basé sur cette fonction, par exemple le maximum de vraisemblance (MV), qui permet alors d'estimer $\mathbf{f}$.

Mais malheureusement, ces méthodes ne peuvent pas fournir de solutions satisfaisantes à des problèmes inverses, où souvent le nombre de paramètres à estimer est aussi grand, sinon plus que le nombre de points de mesures.

3. L'approche estimation *statistique bayésienne* où non seulement on prend en compte l'incertitude sur les mesures, mais aussi on attribue une loi *a priori* aux inconnues $\mathbf{f}$ qui traduit notre information *a priori* sur ces inconnues. Ainsi, on peut combiner l'information contenue dans les données et cette information *a priori* si nécessaire dans des problèmes inverses.

On retrouve en effet la principale idée de la régularisation qui consiste à combiner l'information contenue dans les données et celle contenue dans l'*a priori*.

4. Enfin, on peut aussi noter l'approche basée sur le principe du maximum d'entropie. L'idée de base peut se résumer dans les étapes suivantes :

- on définit une loi *a priori* $p_0(\mathbf{f})$ ou plus généralement une mesure de référence

pour traduire notre information *a priori* sur ces inconnues.

- on considère les mesures comme des contraintes sur la loi *a posteriori* $p(\mathbf{f})$, mais ces mesures n'étant pas suffisantes pour déterminer d'une manière unique $p(\mathbf{f})$, on choisirait la loi $p(\mathbf{f})$ qui est la plus proche de $p_0(\mathbf{f})$ au sens de la mesure de Kulback, et finalement,
- on définit comme solution du problème inverse la moyenne de cette loi.

Nous allons détailler ces approches dans le chapitre suivant.

Chapitre 8

Approches probabilistes

Nous avons distingué quatre approches probabilistes différentes pour la résolution des problèmes inverses. L'objet de ce chapitre est de détailler un peu plus ces approches.

Dans ce chapitre, nous aborderons les trois premières approches, en montrant par des exemples simples leur limites. Nous réservons une place privilégiée à l'approche bayésienne qui sera détaillée dans un chapitre séparé.

8.1 Approche n'utilisant que les moments

Dans cette approche, on considère que les grandeurs observées et inconnues sont bien des fonctions aléatoires, on se limite à caractériser les lois de probabilités de ces grandeurs que par leurs moments, souvent jusqu'à l'ordre deux. Ensuite on établit un lien entre les moments des grandeurs observées et ceux des grandeurs inconnues puis on cherche à l'inverser.

Mais, un point qui est souvent laissé pour compte est : l'estimation des moments des grandeurs observées. Souvent on fait l'hypothèse que l'on peut estimer ces moments par les moyennes empiriques. Ceci peut être raisonnable lorsqu'on se limite aux moments d'ordre un et deux, mais cela devient inefficace pour les moments d'ordre supérieur. C'est d'ailleurs dans le premier cas que ces méthodes trouvent particulièrement leur intérêt : c'est l'estimation en moyenne quadratique.

Pour illustrer cette approche, nous allons considérer le cas du filtrage optimal de Wiener pour la résolution du problème de déconvolution de signaux.

Filtrage optimal de Wiener :

Considérons le problème de la déconvolution des signaux :

$$g(t) = h(t) * f(t) + b(t)$$

supposons que $f(t)$, $b(t)$ et $g(t)$ sont des fonctions aléatoires et cherchons à lier les moments d'ordre un :

$$E[g(t)], \quad E[b(t)] \quad \text{et} \quad E[f(t)]$$

et les moments d'ordre deux :

$$\begin{aligned} R_{gg}(\tau) &= E[g(t)g(t+\tau)], \\ R_{ff}(\tau) &= E[f(t)f(t+\tau)], \\ R_{bf}(\tau) &= R_{fb}(-\tau) = E[b(t)f(t+\tau)] \\ R_{gf}(\tau) &= R_{fg}(-\tau) = E[g(t)f(t+\tau)] \end{aligned}$$

de ces grandeurs. Faisant l'hypothèse que $f(t)$ et $b(t)$ sont indépendantes. On obtient :

$$\begin{aligned} E[g(t)] &= h(t) * E[f(t)] + E[b(t)] \\ R_{gg}(\tau) &= h(t) * h(t) * R_{ff}(\tau) + R_{bb}(\tau) \\ R_{gf}(\tau) &= h(t) * R_{ff}(\tau) \end{aligned}$$

Passant dans le domaine de FOURIER on a :

$$\begin{aligned} S_{gg}(\omega) &= |H(\omega)|^2 S_{ff}(\omega) + R_{bb}(\omega) \\ S_{gf}(\omega) &= H(\omega) S_{ff}(\omega), \\ S_{fg}(\omega) &= H^*(\omega) S_{ff}(\omega), \end{aligned}$$

L'objectif du filtrage optimal est de fournir une solution $\hat{f}(t)$ qui s'obtient par un filtrage linéaire de $g(t)$:

$$\hat{f}(t) = w(t) * g(t)$$

et la réponse impulsionnelle du filtre $w(t)$ est telle que l'erreur quadratique moyenne

$$EQM = E \left[\left(f(t) - \hat{f}(t) \right)^2 \right]$$

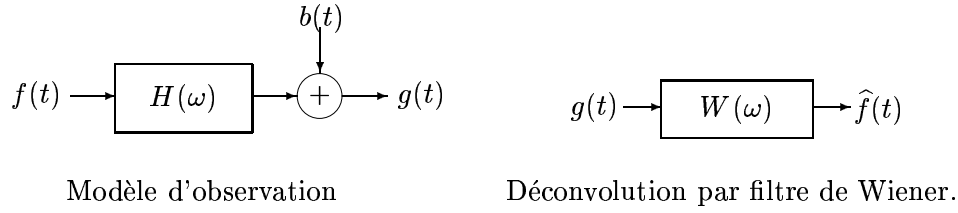


FIG. 8.1 – Filtrage de WIENER.

soit minimale.

En utilisant le principe d'orthogonalité :

$$E[(f(t) - w(t) * g(t)) g(t - \tau)] = 0 \quad \forall t, \tau \longrightarrow R_{fg}(\tau) = w(t) * R_{gg}(\tau).$$

Passant dans le domaine de FOURIER on a :

$$W(\omega) = \frac{S_{fg}(\omega)}{S_{gg}(\omega)} = \frac{H^*(\omega) S_{ff}(\omega)}{|H(\omega)|^2 S_{ff}(\omega) + S_{bb}(\omega)}.$$

Remplaçant pour les expressions de $S_{fg}(\omega)$ et $S_{gg}(\omega)$ on obtient :

$$W(\omega) = \frac{H^*(\omega) S_{ff}(\omega)}{|H(\omega)|^2 S_{ff}(\omega) + S_{bb}(\omega)} = \frac{1}{H(\omega)} \frac{|H(\omega)|^2}{|H(\omega)|^2 + \frac{S_{bb}(\omega)}{S_{ff}(\omega)}}$$

La difficulté dans cette approche commence lors de la mise en œuvre réelle de la méthode. En effet, il faut pouvoir estimer $S_{gg}(\omega)$ et $S_{fg}(\omega)$ ce qui nécessite la connaissance *a priori* de $S_{ff}(\omega)$ et de $S_{bb}(\omega)$.

Souvent, en pratique on fait l'hypothèse que

$$\frac{S_{bb}(\omega)}{S_{ff}(\omega)} = r,$$

avec r une constante représentant l'inverse du rapport signal à bruit. Avec cette hypothèse on a

$$W(\omega) = \frac{1}{H(\omega)} \frac{|H(\omega)|^2}{|H(\omega)|^2 + r}.$$

Cette hypothèse n'est cependant pas très réaliste car elle considère que le signal et le bruit ont la même forme de spectre, ce qui est rarement le cas. En effet, on peut souvent considérer le bruit blanc, mais le signal que l'on cherche n'a pas le même spectre.

Notons que lorsque $r \mapsto 0$ on retrouve le filtrage inverse.

Une autre hypothèse plus réaliste consiste à supposer

$$\frac{S_{bb}(\omega)}{S_{ff}(\omega)} = |D(\omega)|^2,$$

où $D(\omega)$ représente la fonction de transfert d'un filtre passe-haut. Avec cette hypothèse on a

$$W(\omega) = \frac{1}{H(\omega)} \frac{|H(\omega)|^2}{|H(\omega)|^2 + |D(\omega)|^2}.$$

Nous verrons dans le chapitre suivant qu'il existe un lien entre cette approche et celle de l'estimation bayésienne dans le cas linéaire avec des lois gaussiennes.

8.2 Estimation au sens du maximum de vraisemblance

Dans cette approche on prend en compte explicitement le caractère incertain des mesures en les caractérisant par une loi de probabilité $p(\mathbf{g}|\mathbf{f})$.

L'idée de base ensuite est de considérer $p(\mathbf{g}|\mathbf{f})$ comme une fonction $V(\mathbf{f}) = p(\mathbf{g}|\mathbf{f})$ appelée fonction *vraisemblance*. On définit alors la solution $\hat{\mathbf{f}}$ du problème inverse comme l'argument qui maximise cette fonction :

$$\hat{\mathbf{f}} = \arg \max_{\mathbf{f}} \{p(\mathbf{g}|\mathbf{f})\} = \arg \min_{\mathbf{f}} \{-\ln p(\mathbf{g}|\mathbf{f})\}$$

Exemple 1: Loi Gaussienne :

Considérons le modèle $\mathbf{g} = \mathbf{H}\mathbf{f} + \mathbf{b}$ et supposons que $\mathbf{b}$ puisse être modélisé par un vecteur aléatoire centrée, blanc et gaussien :

$$b_i \sim \mathcal{N}(0, \sigma_b^2) \longrightarrow \mathbf{b} \sim \mathcal{N}(\mathbf{0}, \sigma_b^2 \mathbf{I}) \longrightarrow \mathbf{g}|\mathbf{f} \sim \mathcal{N}(\mathbf{H}\mathbf{f}, \sigma_b^2 \mathbf{I}).$$

On en déduit alors :

$$p(\mathbf{g}|\mathbf{f}) = K \exp \left[\frac{-1}{2\sigma_b^2} [\mathbf{g} - \mathbf{H}\mathbf{f}]^t [\mathbf{g} - \mathbf{H}\mathbf{f}] \right].$$

L'estimé au sens du MV devient alors :

$$\hat{\mathbf{f}} = \arg \min_{\mathbf{f}} \{-\ln p(\mathbf{g}|\mathbf{f})\} = \arg \min_{\mathbf{f}} \{\|\mathbf{g} - \mathbf{H}\mathbf{f}\|^2\}$$

et on retrouve l'estimation au sens des moindres carrés.

Si nous connaissons un peu plus sur les caractéristiques des b_i , par exemple si les b_i ont des variances différentes, alors

$$b_i \sim \mathcal{N}(0, \sigma_i^2)$$

et il n'est pas difficile de montrer que

$$p(\mathbf{g}|\mathbf{f}) = K \exp \left[\frac{-1}{2} (\mathbf{g} - \mathbf{H}\mathbf{f})^t \mathbf{W}^{-1} (\mathbf{g} - \mathbf{H}\mathbf{f}) \right]$$

avec $\mathbf{W} = \text{diag} \{\sigma_1^2, \dots, \sigma_m^2\}$. Alors,

$$\hat{\mathbf{f}} = \arg \min_{\mathbf{f}} \{-\ln p(\mathbf{g}|\mathbf{f})\} = \arg \min_{\mathbf{f}} \{\|\mathbf{g} - \mathbf{H}\mathbf{f}\|_{\mathbf{W}}^2\}$$

et on retrouve l'estimation au sens des moindres carrés généralisés.

L'extension au cas d'un modèle non linéaire se fait aussi très facilement.

Exemple 2: Loi Gaussienne généralisée :

Supposons maintenant que les b_i suivent une loi gaussienne généralisée :

$$b_i \sim p(b_i) \propto \exp [-\beta |b_i|^p], \quad \beta > 0, \quad 1 < p \leq 2.$$

Si on fait l'hypothèse que $\mathbf{b}$ est blanc alors on a :

$$\mathbf{b} \sim p(\mathbf{b}) \propto \exp \left[-\beta \sum_{i=1}^M |b_i|^p \right] \longrightarrow -\ln p(\mathbf{b}) = K + \beta \sum_{i=1}^M |b_i|^p$$

On en déduit alors :

$$-\ln p(\mathbf{g}|\mathbf{f}) = K + \beta \sum_{i=1}^M |g_i - [\mathbf{H}\mathbf{f}]_i|^p.$$

L'estimé au sens du MV devient alors :

$$\hat{\mathbf{f}} = \arg \min_{\mathbf{f}} \{-\ln p(\mathbf{g}|\mathbf{f})\} = \arg \min_{\mathbf{f}} \left\{ \sum_{i=1}^M |g_i - [\mathbf{H}\mathbf{f}]_i|^p \right\} = \arg \min_{\mathbf{f}} \{\|\mathbf{g} - \mathbf{H}\mathbf{f}\|^p\}$$

et on retrouve l'estimation au sens des moindres carrés.

Exemple 3: Loi Gamma :

Supposons maintenant que les b_i suivent une loi de gamma :

$$b_i \sim \mathcal{G}(\alpha, \beta) \longrightarrow p(b_i) = K_i b_i^{-\alpha} \exp[-\beta b_i]$$

et si on fait l'hypothèse que $\mathbf{b}$ est blanc (b_i et b_j) indépendantes $\forall i, j, i \neq j$, on a

$$p(\mathbf{b}) = \prod_{i=1}^M p(b_i) \longrightarrow \ln p(\mathbf{b}) = K + \sum_{i=1}^M \alpha \log(b_i) + \beta b_i.$$

On en déduit :

$$\ln p(\mathbf{g}|\mathbf{f}) = K + \sum_{i=1}^M \alpha \log(g_i - [\mathbf{H}\mathbf{f}]_i) + \beta(g_i - [\mathbf{H}\mathbf{f}]_i)$$

et

$$\hat{\mathbf{f}} = \arg \min_{\mathbf{f}} \{-\ln p(\mathbf{g}|\mathbf{f})\}.$$

Il n'y a pas d'expression analytique pour cette solution, mais elle peut être calculé numériquement par un algorithme itératif.

En conclusion :

- L'approche MV permet de mieux prendre en compte la nature du bruit en lui attribuant une loi de probabilité $p(\mathbf{b})$, ou, d'une manière plus générale, l'incertaine sur les mesures en leur attribuant une loi de probabilité $p(\mathbf{g}|\mathbf{f})$;
- Dans le cas gaussien on retrouve l'estimation au sens des moindres carrés ;
- Cette approche peut fournir des résultats satisfaisants lorsqu'elle est utilisée en combinaison avec une modélisation paramétrique des inconnues avec un nombre de paramètres très inférieure au nombre de données indépendantes. Malheureusement, cette approche donne rarement des résultats satisfaisants pour la résolution des problèmes inverses dans le cadre algébrique où le nombre d'inconnues est du même ordre de grandeur voir plus grand que le nombre des mesures. C'est pourquoi, certains auteurs ont proposé une approche dite *Maximum de vraisemblance pénalisée* qui consiste à définir la solution par :

$$\hat{\mathbf{f}} = \arg \min_{\mathbf{f}} \{-\ln p(\mathbf{g}|\mathbf{f}) + \phi(\mathbf{f})\}$$

où $\phi(\mathbf{f})$ est une fonction de pénalisation choisie d'une manière adhoc pour permettre d'obtenir une solution satisfaisante. Nous verrons alors que cette approche peut mieux être interprétée dans le cadre de l'inférence bayésienne où $\phi(\mathbf{f}) = -\ln p(\mathbf{f})$ avec $p(\mathbf{f})$ la loi *a priori*.

8.3 Approche du maximum d'entropie

Avant de voir comment le principe du maximum d'entropie peut être utilisé pour la résolution des problèmes inverses, nous allons rappeler un certain nombre de définitions et de notions qui seront essentielles pour la compréhension de cette approche.

8.3.1 Définition de l'entropie

Le principe du ME peut être approché de différentes manières. L'approche de la théorie de l'information [SW48] est sans doute la mieux adaptée à notre problème. JAYNES ([Jay68, Jay82, Jay85]) est parmi les premiers auteurs modernes à avoir introduit le formalisme du ME par l'approche de la théorie de l'information. Cette notion d'entropie est introduite de la manière suivante : considérons une variable aléatoire discrète X produisant des réalisations $\mathbf{x} = \{x_1, x_2, \dots, x_n\}$ et attribuons les probabilités $\mathbf{p} = \{p_1, p_2, \dots, p_n\}$ à ces réalisations pour représenter notre information partielle sur cette variable. On définit la quantité $I_i = \ln(1/p_i)$ comme la quantité d'information apportée par la réalisation x_i .

Le raisonnement intuitif conduisant à cette expression est le suivant : plus un événement est rare, plus le gain d'information obtenu par sa réalisation est grand. L'utilisation du logarithme rend additif le gain total d'information obtenu par la réalisation de plusieurs événements indépendants. On définit alors l'entropie d'un processus par la somme pondérée des informations individuelles de chaque réalisation. C'est la définition de l'entropie donnée par SHANNON :

$$S(\mathbf{p}) = \sum_{i=1}^n p_i \ln \frac{1}{p_i} = - \sum_{i=1}^n p_i \ln p_i. \quad (8.1)$$

L'entropie S est une mesure d'incertitude de la distribution $\mathbf{p} = \{p_1, p_2, \dots, p_n\}$, déterminée uniquement par certaines règles élémentaires de cohérence logique et d'additivité ([Jay82, Jay85, Bal82]).

Généralisant ce concept, on définit $-\ln(q_i/p_i)$ comme le gain d'information, sur une probabilité *a priori* p_i , apporté par la connaissance de la probabilité q_i de réalisation de l'événement x_i . On définit alors :

$$S(\mathbf{q}, \mathbf{p}) = - \sum_{i=1}^n q_i \ln \frac{p_i}{q_i} = \sum_{i=1}^n q_i \ln \frac{q_i}{p_i}, \quad (8.2)$$

appelée entropie croisée ou entropie relative de la distribution q_i par rapport à la distribution p_i . Notons ici qu'il y a un changement de signe pour des raisons historiques et il est clair que la minimisation de l'entropie croisée $S(\mathbf{q}, \mathbf{p})$ se réduit à la maximisation de l'entropie $S(\mathbf{q})$ si l'*a priori* $\mathbf{p}$ est uniforme.

Sous réserve de quelques précautions, on peut généraliser ce qui précède au cas de distributions continues, et définir l'entropie [SW48] par :

$$S(p) = - \int p(x) \ln p(x) dx, \quad (8.3)$$

et l'entropie croisée par :

$$S(q, p) = - \int q(x) \ln \frac{p(x)}{q(x)} dx. \quad (8.4)$$

8.3.2 Lois à maximum d'entropie

Voyons maintenant ce que signifie le choix d'une distribution de probabilité à maximum d'entropie contenant une information *a priori*, ou qui soit compatible avec des contraintes connues sur cette distribution. Pour cela, précisons tout d'abord la signification de l'information contenue dans une distribution de probabilité $\{p_1, p_2, \dots, p_n\}$. À l'évidence il faut que l'on puisse extraire cette information de cette distribution. Considérons une variable aléatoire discrète X qui peut prendre des valeurs $\mathbf{x} = \{x_1, x_2, \dots, x_n\}$ avec une distribution de probabilité $\mathbf{p} = \{p_1, p_2, \dots, p_n\}$. Maintenant si on nous demande quelle est la meilleure estimée $\hat{\phi}$ d'une fonction $\phi(X)$ au sens du minimum de l'erreur quadratique moyenne. La solution est immédiate :

$$\hat{\phi} = E[\phi] = \sum_{i=1}^n p_i \phi(x_i).$$

C'est un problème direct et bien-posé. Inversement, si l'on se pose la question d'ajuster une distribution $\mathbf{p}$ pour incorporer une information donnée sur la fonction $\phi(X)$, il faut se poser la question suivante :

connaissant une règle d'estimation précise, par exemple la minimisation de l'erreur quadratique moyenne, comment choisir $\mathbf{p}$ pour que l'on ait $\hat{\phi} = E[\phi]$?

Le problème est qu'il existe, en général, beaucoup de distributions qui satisfont cette contrainte. Il s'agit d'un problème inverse mal posé au sens où la solution n'est pas unique. Le principe du maximum d'entropie nous permet alors de choisir une solution.

Soit maintenant le cas plus général où nous considérons les m fonctions $\{\phi_1(X), \phi_2(X), \dots, \phi_m(X)\}$ pour lesquelles nous avons un ensemble de données $\{d_1, d_2, \dots, d_m\}$ pouvant s'exprimer sous la forme des m contraintes simultanées suivantes :

$$E[\phi_k(X)] = \sum_{i=1}^n p_i \phi_k(x_i) = d_k, \quad k = 1, \dots, m.$$

Dans chaque problème, ces données $\{d_1, d_2, \dots, d_m\}$ peuvent avoir des interprétations physiques différentes et la difficulté consiste à incorporer ces données (contraintes) dans notre distribution de probabilité. Nous voulons ajuster la distribution de probabilité $\{p_1, p_2, \dots, p_n\}$ à nos données. En général, le nombre des contraintes m est inférieur à n et il y a une infinité de solutions à ce problème. En d'autres termes, on peut trouver une infinité de distributions de probabilité $\{p_1, p_2, \dots, p_n\}$ qui satisfont ces contraintes. C'est ici que le principe du maximum d'entropie est mis en œuvre pour choisir, parmi ces solutions possibles, celle qui a l'entropie maximale, c'est-à-dire la loi qui satisfait toutes les contraintes (toute l'information connue) et qui est la moins compromettante vis-à-vis de toute autre information inconnue. Le mot *information* devant être pris au sens de la définition de l'information moyenne ou de l'entropie de Shannon (équation 3).

Le problème se formule ainsi :

Problème P1 :

$$\begin{aligned} &\text{maximiser} \quad S(\mathbf{p}) = - \sum_{i=1}^n p_i \ln p_i, \\ &\text{sous les contraintes} \quad \sum_{i=1}^n p_i \phi_k(x_i) = d_k, \quad k = 1, \dots, m. \end{aligned}$$

La solution est obtenue par une technique variationnelle de multiplicateurs de Lagrange et est donnée par :

$$p_i = \frac{1}{Z(\lambda_1, \dots, \lambda_m)} \exp \left[- \sum_{k=1}^m \lambda_k \phi_k(x_i) \right], \quad i = 1, \dots, n, \quad (8.5)$$

où

$$Z(\lambda_1, \dots, \lambda_m) = \sum_{i=1}^n \exp \left[- \sum_{k=1}^m \lambda_k \phi_k(x_i) \right] \quad (8.6)$$

est la fonction de partition, et les $\{\lambda_1, \lambda_2, \dots, \lambda_m\}$ sont déterminés par le système d'équations :

$$-\frac{\partial \ln Z(\lambda_1, \dots, \lambda_m)}{\partial \lambda_k} = d_k, \quad k = 1, \dots, m, \quad (8.7)$$

avec m données d_k et m inconnues λ_k . Notons que ce système d'équations n'est autre que le système d'équations des contraintes constituées par les données.

La valeur maximale de l'entropie est :

$$S_{\max} = \ln Z + \sum_{k=1}^m \lambda_k d_k.$$

Si on note $\lambda_0 = \ln Z$ on peut vérifier facilement les propriétés suivantes :

$$-\frac{\partial \lambda_0}{\partial \lambda_k} = E[\phi_k(x_i)] = d_k, \quad (8.8)$$

$$-\frac{\partial^2 \lambda_0}{\partial \lambda_k^2} = E[\phi_k^2(x_i)] - d_k^2 = \text{Var}[\phi_k(x_i)], \quad (8.9)$$

$$-\frac{\partial^2 \lambda_0}{\partial \lambda_k \partial \lambda_l} = E[\phi_k(x_i) \phi_l(x_i)] - d_k d_l = \text{Cov} \{ \phi_k(x_i), \phi_l(x_i) \}. \quad (8.10)$$

8.3.3 Lois à minimum d'entropie relative

L'extension de ce qui précède au cas de l'entropie croisée se formule ainsi :

Problème P2 :

Étant donné une distribution de probabilité *a priori* $\mathbf{p} = \{p_1, p_2, \dots, p_n\}$, déterminer la distribution de probabilité *a posteriori* $\mathbf{q} = \{q_1, q_2, \dots, q_n\}$ qui minimise

$$S(\mathbf{q}, \mathbf{p}) = - \sum_i q_i \ln \frac{p_i}{q_i} = \sum_i q_i \ln \frac{q_i}{p_i},$$

et qui satisfait les m contraintes

$$\sum_{i=1}^n q_i \phi_k(x_i) = d_k, \quad k = 1, \dots, m.$$

La solution est obtenue, là aussi, par une technique variationnelle de multiplicateurs de Lagrange et donnée par :

$$q_i = \frac{1}{Z} p_i \exp \left[- \sum_{k=1}^m \lambda_k \phi_k(x_i) \right], \quad i = 1, \dots, n, \quad (8.11)$$

où

$$Z(\lambda_1, \dots, \lambda_m) = \sum_{i=1}^n p_i \exp \left[- \sum_{k=1}^m \lambda_k \phi_k(x_i) \right] \quad (8.12)$$

est la fonction de partition, et les $\{\lambda_1, \lambda_2, \dots, \lambda_n\}$ sont déterminés par le système d'équations (9).

8.3.4 Extension au cas continu

Dans le cas continu, les définitions (8.3) et (8.4) conduisent à :

Problème P3 :

$$\begin{aligned} &\text{maximiser} \quad S(p) = - \int p(x) \ln p(x) \, dx, \\ &\text{sous les contraintes} \quad \int \phi_k(x) p(x) \, dx = d_k, \quad k = 1, \dots, m. \end{aligned}$$

La solution est

$$p(x) = \frac{1}{Z} \exp \left[- \sum_{k=1}^m \lambda_k \phi_k(x) \right], \quad (8.13)$$

où

$$Z(\lambda_1, \dots, \lambda_m) = \int \exp \left[- \sum_{k=1}^m \lambda_k \phi_k(x) \right] \, dx \quad (8.14)$$

est la fonction de partition, et les $\{\lambda_1, \lambda_2, \dots, \lambda_n\}$ sont déterminés par le système d'équations (9).

Problème P4 :

Étant donné une densité de probabilité *a priori* $p(x)$, déterminer la densité de probabilité *a posteriori* $q(x)$ qui minimise l'entropie croisée

$$S(q, p) = - \int q(x) \ln \frac{p(x)}{q(x)} \, dx,$$

et qui satisfait les m contraintes

$$\int \phi_k(x) q(x) \, dx = d_k, \quad k = 1, \dots, m.$$

La solution est

$$q(x) = \frac{1}{Z} p(x) \exp \left[- \sum_{k=1}^m \lambda_k \phi_k(x) \right], \quad (8.15)$$

où

$$Z(\lambda_1, \dots, \lambda_m) = \int p(x) \exp \left[- \sum_{k=1}^m \lambda_k \phi_k(x) \right] \, dx \quad (8.16)$$

est la fonction de partition, et les $\{\lambda_1, \lambda_2, \dots, \lambda_n\}$ sont déterminés par (9).

Notons également que le vecteur $\hat{\lambda} = (\lambda_1, \dots, \lambda_m)$ dans tous ces problèmes peut être considéré comme la solution du problème d'optimisation suivant :

Problème dual des problèmes P1-P4 :

$$\hat{\lambda} = \arg \min_{\lambda} \left\{ D(\lambda) = \lambda^t d + \ln Z(\lambda) \right\} \quad (8.17)$$

où $\mathbf{d} = (d_1, \dots, d_m)$. Ce problème d'optimisation est appelé *problème dual* des problèmes P1-P4. Notons que l'expression de $Z(\lambda)$ est différente suivant le problème.

Une présentation légèrement différente de ces relations peut être obtenue si on définit

$$\phi_0(x) = 1, \quad \text{et} \quad d_0 = 1,$$

ce qui permet d'inclure la contrainte de normalisation de $q(x)$ et de formuler le problème précédent de la manière suivante :

$$\begin{aligned} \text{minimiser} \quad & S(q,p) = - \int q(x) \ln \frac{p(x)}{q(x)} dx, \\ \text{sous les contraintes} \quad & \int \phi_k(x) q(x) dx = d_k, \quad k = 0, \dots, m. \end{aligned}$$

La solution peut être écrite sous la forme

$$q(x) = p(x) \exp \left[- \sum_{k=0}^m \lambda_k \phi_k(x) \right], \quad (8.18)$$

et les $\{\lambda_0, \lambda_1, \lambda_2, \dots, \lambda_n\}$ sont déterminés par :

$$G(\lambda_0, \lambda_1, \lambda_2, \dots, \lambda_n) = - \int \phi_k(x) p(x) \exp \left[- \sum_{l=0}^m \lambda_l \phi_l(x) \right] dx = d_k, \quad k = 0, \dots, m. \quad (8.19)$$

On remarque que λ_0 est relié à la fonction de partition Z par $Z = -\exp[\lambda_0]$, et on a :

$$\lambda_0 = -\ln Z = -\ln \int p(x) \exp \left[- \sum_{k=1}^m \lambda_k \phi_k(x) \right] dx. \quad (8.20)$$

Les autres coefficients λ_k sont déterminés par le système d'équations :

$$-\frac{\partial}{\partial \lambda_k} \lambda_0(\lambda_1, \dots, \lambda_n) = d_k, \quad k = 1, \dots, m. \quad (8.21)$$

Notons aussi que la valeur minimale $S_{\min}(q,p)$ peut être exprimée en fonction des λ_k par :

$$S_{\min}(q,p) = -\lambda_0 - \sum_{k=1}^m \lambda_k \phi_k(x). \quad (8.22)$$

Il n'est en général pas possible d'obtenir une relation explicite pour les coefficients λ_k . On résoud alors le système d'équations (8.19) numériquement par des méthodes itératives. Cependant, dans certaines situations simples on peut résoudre le problème d'une façon analytique [Jay82, JS84]. Le tableau 1 montre quelques exemples de lois, que l'on obtient sous forme analytique, en résolvant le problème P3 :

$$\begin{aligned} \text{maximiser} \quad & S(p) = - \int p(x) \ln p(x) dx, \\ \text{sous les contraintes} \quad & \int \phi_k(x) p(x) dx = d_k, \quad k = 1, \dots, m. \end{aligned}$$

Tableau 1: Exemples de lois à maximum d'entropie.		
contraintes	domaine de x	loi de probabilité
$\phi_1(x) = x$	$x \in \mathbf{R}_+$	$p(x) = \frac{1}{\mu} \exp[-\mu x]$ loi exponentielle
$\phi_1(x) = x $	$x \in \mathbf{R}$	$p(x) = \frac{1}{2\mu} \exp[-\mu x]$ loi de LAPLACE
$\begin{cases} \phi_1(x) = x \\ \phi_2(x) = x^2 \end{cases}$	$x \in \mathbf{R}$	$p(x) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left[-\frac{(x-m)^2}{2\sigma^2}\right]$ loi de Gausse
$\begin{cases} \phi_1(x) = x \\ \phi_2(x) = \ln x \end{cases}$	$x \in \mathbf{R}_+$	$p(x) \propto \exp[-\lambda \ln x - \mu x]$ $\propto x^{-\lambda} \exp[-\mu x]$ loi Gamma
$\begin{cases} \phi_1(x) = \ln x \\ \phi_2(x) = \ln(1-x) \end{cases}$	$x \in]0, 1[$	$p(x) \propto \exp[-\lambda \ln x - \mu \ln(1-x)]$ $\propto x^{-\lambda} (1-x)^{-\mu}$ loi Béta
$\begin{cases} \phi_1(x) = \ln x \\ \phi_2(x) = x^2 \end{cases}$	$x > 0$	$p(x) \propto \exp[-\lambda \ln x - \mu x^2]$ $\propto x^{-\lambda} \exp[-\mu x^2]$ loi de Rayleigh

8.3.5 Extension au cas multivariable

Toutes ces relations peuvent être généralisées au cas multivariable. Par exemple le problème P4 devient :

Problème P5 :

Étant donné une densité de probabilité *a priori* $p(\mathbf{x})$, déterminer la densité de probabilité *a posteriori* $q(\mathbf{x})$ qui minimise l'entropie croisée

$$S(q, p) = - \int q(\mathbf{x}) \ln \frac{p(\mathbf{x})}{q(\mathbf{x})} d\mathbf{x},$$

et qui satisfait les m contraintes

$$\int \phi_k(\mathbf{x}) q(\mathbf{x}) d\mathbf{x} = d_k, \quad k = 1, \dots, m.$$

La solution est

$$q(\mathbf{x}) = \frac{1}{Z} p(\mathbf{x}) \exp \left[- \sum_{k=1}^m \lambda_k \phi_k(\mathbf{x}) \right], \quad (8.23)$$

où

$$Z(\lambda_1, \dots, \lambda_m) = \int p(\mathbf{x}) \exp \left[- \sum_{k=1}^m \lambda_k \phi_k(\mathbf{x}) \right] d\mathbf{x} \quad (8.24)$$

et les $\{\lambda_1, \lambda_2, \dots, \lambda_n\}$ sont déterminés par :

$$-\frac{\partial \ln Z(\lambda_1, \dots, \lambda_m)}{\partial \lambda_k} = d_k, \quad k = 1, \dots, m. \quad (8.25)$$

Ici aussi, dans certains cas on peut obtenir une relation explicite pour la solution. Par exemple, si $p(\mathbf{x})$ est une distribution exponentielle multivariée de la forme séparable

$$p(\mathbf{x}) = \prod_{i=1}^n \frac{1}{\alpha_i} \exp \left[-\frac{x_i}{\alpha_i} \right],$$

et si les contraintes sont des contraintes sur les moments d'ordre un

$$\mathbb{E}[X_i] = \int x_i q(x_i) dx_i = m_i, \quad i = 1, \dots, n,$$

alors la solution $q(\mathbf{x})$ reste une fonction exponentielle multivariée séparable :

$$q(\mathbf{x}) = \prod_{i=1}^n \frac{1}{m_i} \exp \left[-\frac{x_i}{m_i} \right].$$

De même si $p(\mathbf{x})$ est une fonction gaussienne multivariée séparable :

$$p(\mathbf{x}) = \prod_{i=1}^n \frac{1}{\sqrt{2\pi}\alpha_i} \exp \left[-\frac{1}{2} \left(\frac{x_i - m_i}{\alpha_i} \right)^2 \right],$$

et si les contraintes sont des contraintes du second ordre :

$$\mathbb{E}[(X_i - m_i)^2] = \int (x_i - m_i)^2 q(x_i) dx_i = \sigma_i^2,$$

alors la solution $q(\mathbf{x})$ reste une fonction gaussienne multivariée séparable :

$$q(\mathbf{x}) = \prod_{i=1}^n \frac{1}{\sqrt{2\pi}\sigma_i} \exp \left[-\frac{1}{2} \left(\frac{x_i - m_i}{\sigma_i} \right)^2 \right].$$

Ce dernier résultat peut être généralisé au cas où $p(\mathbf{x})$ est une fonction gaussienne multivariée non dégénérée :

$$p(\mathbf{x}) = \frac{|\mathbf{\Sigma}|^{-1/2}}{(2\pi)^{n/2}} \exp \left[-\frac{1}{2} (\mathbf{x} - \mathbf{m})^t \mathbf{\Sigma}^{-1} (\mathbf{x} - \mathbf{m}) \right],$$

et si les contraintes sont des contraintes sur les moments d'ordre un et deux non singulières ($\mathbf{\Sigma}'$ inversible) :

$$\begin{aligned} \mathbb{E}[\mathbf{x}] &= \int \mathbf{x} q(\mathbf{x}) d\mathbf{x} = \mathbf{m}', \\ \mathbb{E}[(\mathbf{x} - \mathbf{m}')(\mathbf{x} - \mathbf{m}')^t] &= \int (\mathbf{x} - \mathbf{m}')(\mathbf{x} - \mathbf{m}')^t q(\mathbf{x}) d\mathbf{x} = \mathbf{\Sigma}', \end{aligned}$$

alors la solution $q(\mathbf{x})$ reste une fonction gaussienne multivariée non dégénérée :

$$q(\mathbf{x}) = \frac{|\mathbf{\Sigma}'|^{-1/2}}{(2\pi)^{n/2}} \exp \left[-\frac{1}{2} (\mathbf{x} - \mathbf{m}')^t \mathbf{\Sigma}'^{-1} (\mathbf{x} - \mathbf{m}') \right].$$

8.4 Utilisation pour les problèmes inverses

Arrivés à ce stade, nous avons les outils nécessaires pour comprendre comment le principe du maximum d'entropie peut être utilisé dans la résolution des problèmes inverses.

Il existe au moins deux catégories de méthodes qui utilisent le principe du maximum d'entropie pour la résolution des problèmes inverses :

8.4.1 Maximum d'entropie classique

Dans cette approche qui ne s'applique qu'aux objets positifs, l'idée de base consiste à assimiler l'objet $\mathbf{f}$ à une distribution de probabilité et les données $\mathbf{g}$ à des contraintes linéaires sur celle-ci. Sachant qu'en général, ces contraintes ne sont pas suffisantes pour définir une unique solution au problème, on utilise le PME pour choisir une solution parmi l'ensemble des solutions admissibles qui peuvent être définies soit par

$$\{\mathbf{f} : \mathbf{H}\mathbf{f} = \mathbf{g}\},$$

ou par

$$\{\mathbf{f} : \|\mathbf{g} - \mathbf{H}\mathbf{f}\|^2 \leq c\}.$$

On choisit alors dans cet ensemble la solution qui maximise l'entropie

$$-\sum_{j=1}^N f_j \log f_j,$$

ou d'une manière plus générale celle qui minimise

$$-\sum_{j=1}^N \left[f_j \log \left(\frac{f_j}{m_j} \right) + (f_j - m_j) \right],$$

où $\mathbf{m}$ est une solution par défaut (*a priori*).

Notons que dans cette approche on assimile $\mathbf{f}$ à une distribution de probabilité et les données $\mathbf{g}$ à des contraintes linéaires sur celle-ci.

8.4.2 Maximum d'entropie sur la moyenne

Dans cette approche

- On fait l'hypothèse que $\mathbf{f} \in \mathcal{C}$ où $\mathcal{C}$ est un ensemble convexe muni d'une mesure de référence $\mu(\mathbf{f})$ (ou d'une densité de probabilité *a priori* $d\mu(\mathbf{f}) = p_0(\mathbf{f}) d\mathbf{f}$ et on suppose que

$$\begin{aligned} \mathbf{m} &= \int_{\mathcal{C}} \mathbf{f} d\mu(\mathbf{f}) \\ &= \int_{\mathcal{C}} \mathbf{f} p_0(\mathbf{f}) d\mathbf{f}. \end{aligned}$$

- On considère $\mathbf{f}$ comme la moyenne d'un vecteur aléatoire $\mathbf{F}$ pour lequel on définit une loi de probabilité $P(\mathbf{f})$ (ou une densité de probabilité $dP(\mathbf{f}) = p(\mathbf{f}) d\mathbf{f}$:

$$\begin{aligned} \mathbf{f} = \mathbf{E}[\mathbf{F}] &= \int_{\mathcal{C}} \mathbf{f} dP(\mathbf{f}) \\ &= \int_{\mathcal{C}} \mathbf{f} p(\mathbf{f}) d\mathbf{f}, \end{aligned}$$

et les données $\mathbf{g} = \mathbf{H}\mathbf{f}$ comme des contraintes linéaires sur la loi de probabilité $p(\mathbf{f})$:

$$\begin{aligned}\mathbf{g} = \mathbf{H}\mathbf{f} &= \mathbf{H}\mathbf{E}[\mathbf{F}] = \mathbf{E}[\mathbf{H}\mathbf{F}] = \int_{\mathcal{C}} \mathbf{H}\mathbf{f} \, dP(\mathbf{f}) \\ &= \int_{\mathcal{C}} \mathbf{H}\mathbf{f} p(\mathbf{f}) \, d\mathbf{f}.\end{aligned}$$

- On cherche alors parmi l'ensemble des lois de probabilités satisfaisant ces contraintes celle qui minimise la distance de Kulback entre $dP(\mathbf{f})$ et $d\mu(\mathbf{f})$:

$$- \int \log \frac{dP(\mathbf{f})}{d\mu(\mathbf{f})} dP(\mathbf{f}),$$

ou, d'une manière équivalente, celle qui minimise l'entropie relative

$$- \int p(\mathbf{f}) \log \left(\frac{p(\mathbf{f})}{p_0(\mathbf{f})} \right) d\mathbf{f}.$$

La solution, nous l'avons vu, peut être obtenue via le Lagrangien :

$$\begin{aligned}\mathcal{L}(\mathbf{f}, \boldsymbol{\lambda}) &= \int_{\mathcal{C}} \left[\log \frac{dP(\mathbf{f})}{d\mu(\mathbf{f})} - \sum_{i=1}^M \lambda_i (g_i - [\mathbf{H}\mathbf{f}]_i) \right] dP(\mathbf{f}) \\ &= \int_{\mathcal{C}} \left[\log \frac{dP(\mathbf{f})}{d\mu(\mathbf{f})} - \boldsymbol{\lambda}^t (\mathbf{g} - \mathbf{H}\mathbf{f}) \right] dP(\mathbf{f}) \\ &= \int_{\mathcal{C}} \left[\log \frac{p(\mathbf{f})}{p_0(\mathbf{f})} - \boldsymbol{\lambda}^t (\mathbf{g} - \mathbf{H}\mathbf{f}) \right] p(\mathbf{f}) \, d\mathbf{f}\end{aligned}$$

et donnée par

$$dP(\mathbf{f}, \boldsymbol{\lambda}) = \exp \left[\boldsymbol{\lambda}^t [\mathbf{H}\mathbf{f}] - \log Z(\boldsymbol{\lambda}) \right] d\mu(\mathbf{f}),$$

ou d'une manière équivalente

$$p(\mathbf{f}, \boldsymbol{\lambda}) = \exp \left[\boldsymbol{\lambda}^t [\mathbf{H}\mathbf{f}] - \log Z(\boldsymbol{\lambda}) \right] p_0(\mathbf{f}) \, d\mathbf{f},$$

où

$$Z(\boldsymbol{\lambda}) = \int_{\mathcal{C}} \exp \left[\boldsymbol{\lambda}^t [\mathbf{H}\mathbf{f}] \right] d\mu(\mathbf{f}) = \int_{\mathcal{C}} \exp \left[\boldsymbol{\lambda}^t [\mathbf{H}\mathbf{f}] \right] p_0(\mathbf{f}) \, d\mathbf{f}.$$

Les paramètres de Lagrange sont obtenus en cherchant la solution du système d'équations non linéaires suivant :

$$\frac{\partial \log Z(\boldsymbol{\lambda})}{\partial \lambda_i} = g_i, \quad i = 1, \dots, M.$$

- Une fois $p(\mathbf{f})$ déterminée on définit la solution du problème par

$$\hat{\mathbf{f}} = \mathbf{E}[\mathbf{f}] = \int \mathbf{f} \, dP(\mathbf{f}) = \int \mathbf{f} p(\mathbf{f}) \, d\mathbf{f}.$$

On peut montrer que, cette solution $\hat{\mathbf{f}}$, qui dépend ainsi de $\boldsymbol{\lambda}$, peut être obtenue directement en minimisant un critère $\Omega(\mathbf{f}, \mathbf{m})$ sous les contraintes $\mathbf{g} = \mathbf{H}\mathbf{f}$.

En effet, en notant

$$\mathbf{s} = \mathbf{H}^t \boldsymbol{\lambda}, \quad G^*(\mathbf{s}) = \log Z(\mathbf{s}) = \log \int_{\mathcal{C}} \exp \left[\mathbf{s}^t \mathbf{f} \right] d\mu(\mathbf{f}),$$

et

$$\Omega(\mathbf{f}) = \inf_{\mathbf{s}} \{\mathbf{s}^t \mathbf{f} - G^*(\mathbf{s})\}, \quad D(\boldsymbol{\lambda}) = \boldsymbol{\lambda}^t \mathbf{g} - G^*(\mathbf{H}^t \boldsymbol{\lambda}),$$

on peut alors montrer que :

- Le vecteur de $\boldsymbol{\lambda}$ optimal peut être calculé en cherchant l'optimum de $D(\boldsymbol{\lambda})$:

$$\hat{\boldsymbol{\lambda}} = \arg \max_{\boldsymbol{\lambda}} \{D(\boldsymbol{\lambda})\}$$

Le critère $D(\boldsymbol{\lambda})$ est appelé *Critère dual*;

- La solution du problème $\hat{\mathbf{f}}$ peut être calculée en cherchant l'optimum d'un critère $\Omega(\mathbf{f}, \mathbf{m})$ appelé *critère primal* :

$$\hat{\mathbf{f}} = \arg \min_{\mathbf{f} \in \mathcal{C}} \{H(\mathbf{f}, \mathbf{m})\} \quad \text{sous les contraintes} \quad \mathbf{g} = \mathbf{H} \mathbf{f};$$

- Il existe une relation explicite pour calculer la solution en fonction de $G^*(\mathbf{s})$:

$$\hat{\mathbf{f}}(\mathbf{s}) = \frac{dG^*(\mathbf{s})}{d\mathbf{s}},$$

et où :

- les fonctions G et Ω dépendent de la mesure de référence $\mu(\mathbf{f})$;
- $D(\boldsymbol{\lambda})$ est appelé le *critère dual* et dépend des données $\mathbf{g}$ et de la fonction G ;
- $\Omega(\mathbf{f}, \mathbf{m})$ est appelé le *critère primal* et peut être considéré comme une mesure de distance entre $\mathbf{f}$ et $\mathbf{m}$, ce qui signifie qu'elle satisfait les conditions suivantes :
 - $\Omega(\mathbf{f}, \mathbf{m}) \geq 0$, et $\Omega(\mathbf{f}, \mathbf{m}) = 0$ ssi $\mathbf{f} = \mathbf{m}$;
 - $\Omega(\mathbf{f}, \mathbf{m})$ est continue et convexe sur $\mathcal{C}$;
 - $\Omega(\mathbf{f}, \mathbf{m}) = \infty$ si $\mathbf{f} \notin \mathcal{C}$.

Lorsque la mesure de référence est séparable

$$\mu(\mathbf{f}) = \prod_{j=1}^N \mu_j(f_j),$$

on a aussi :

$$dP(\mathbf{f}, \boldsymbol{\lambda}) = \prod_{j=1}^N dP_j(f_j, \lambda_j),$$

et

$$G(\mathbf{s}) = \sum_j g_j(s_j), \quad \Omega(\mathbf{f}, \mathbf{m}) = \sum_j h_j(f_j, m_j), \quad \hat{x}_j = g'_j(s_j).$$

En remplaçant $\mathbf{s} = \mathbf{H}^t \boldsymbol{\lambda}$ on obtient :

$$G(\boldsymbol{\lambda}) = \sum_j g_j([\mathbf{H}^t \boldsymbol{\lambda}]_j), \quad \Omega(\mathbf{f}, \mathbf{m}) = \sum_j h_j(f_j, m_j), \quad \hat{x}_j = g'_j([\mathbf{H}^t \hat{\boldsymbol{\lambda}}]_j).$$

On peut voir que $\Omega(\mathbf{f}, \mathbf{m})$ est aussi séparable et que h_j et g_j dépendent de la mesure de référence $\mu_j(f_j)$:

- g_j est la transformée de Cramer (log de la transformée de LAPLACE) de μ_j :

$$g_j(s) = \log \int \exp[s f_j] d\mu_j(f_j);$$

– h_j est la convexe conjuguée de g_j : $h_j(f) = \inf_s \{sf - g_j(s)\}$.

Prenons quelques exemples :

	$\mu_j(f)$	$g_j(s)$	$h_j(f, m)$
Gaussien :	$\exp \left[-\frac{1}{2}(f - m)^2 \right]$	$\frac{1}{2}(s - m)^2$	$\frac{1}{2}(f - m)^2$
Poisson :	$\frac{m^f}{f!} \exp [-m]$	$\exp [m - s]$	$-f \log \frac{f}{m} + m - f$
Gamma :	$f^{\alpha-1} \exp \left[-\frac{f}{m} \right]$	$\log(s - m)$	$-\log \frac{f}{m} + \frac{f}{m} - 1$

Ainsi, si on considère le critère primal :

$$\text{maximiser } \Omega(\mathbf{f}, \mathbf{m}) = \sum_{j=1}^N h_j(f_j, m_j) \text{ sous contraintes } \mathbf{H}\mathbf{f} = \mathbf{g},$$

on peut par exemple remarquer que lorsque la mesure de référence est une mesure poissonnienne on a $h_j(f_j, m_j) = -f_j \ln f_j/m_j + (m_j - f_j)$. On peut ainsi faire le parallèle avec la méthode du ME classique.

Cependant, lorsque la mesure de référence n'est pas séparable ou, lorsqu'il y a de l'incertitude sur les données $\mathbf{g}$, il n'est plus aussi facile d'utiliser cette approche. Mentionnons quand même que cette approche est récente, et des travaux tout à fait récents ont essayé de prendre en compte le bruit. Ici nous en présentons une approche qui consiste à remplacer $\mathbf{g} = \mathbf{H}\mathbf{f} + \mathbf{n}$ par

$$\mathbf{g} = [\mathbf{H}|\mathbf{I}] \begin{bmatrix} \mathbf{f} \\ \mathbf{n} \end{bmatrix} \longrightarrow \mathbf{g} = \tilde{\mathbf{H}} \tilde{\mathbf{f}}$$

et faire l'hypothèse d'indépendance entre $\mathbf{b}$ et $\mathbf{f}$:

$$\mu(\tilde{\mathbf{f}}) = \mu_x(\mathbf{f})\mu_n(\mathbf{n})$$

Ensuite, suivant l'approche, on montre que la solution de l'optimisation du critère primal devient :

$$\hat{\mathbf{f}} = \arg \min_{\mathbf{f} \in \mathcal{C}} \{ \mathcal{Q}(\mathbf{g} - \mathbf{H}\mathbf{f}) + \alpha \Omega(\mathbf{f}, \mathbf{m}) \}$$

avec

$$\Omega(\mathbf{f}, \mathbf{m}) = \sum_{j=1}^N h_j(f_j, m_j), \quad \text{et} \quad \mathcal{Q}(\mathbf{z}) = \sum_{i=1}^M q_i(z_i).$$

Là encore, les fonctions $h_j(f_j)$ et $q_i(z_i)$ dépendent des mesures de références $\mu_f(\mathbf{f})$ et $\mu_n(\mathbf{f})$.

Chapitre 9

Approche bayésienne

L'approche bayésienne est une approche cohérente pour la résolution des problèmes inverses car elle permet de prendre en compte et de traiter de la même manière l'information *a priori* sur la solution et celle sur les données.

En effet, cette approche permet non seulement de prendre en compte le caractère incertain des erreurs au travers de l'attribution d'une loi de probabilité aux mesures, mais aussi l'information *a priori* au travers d'une loi de probabilité *a priori* que l'on attribue aux inconnues du problème.

La fusion de ces deux sources d'information s'effectue par la règle de Bayes. On peut alors dire que cette approche généralise la régularisation déterministe.

9.1 Introduction

L'approche bayésienne correspond formellement au cheminement suivant :

1. On explicite un ensemble d'hypothèses $\mathcal{H}$ sur le problème concernant le modèle d'observation, les connaissances *a priori* sur les inconnues $\mathbf{f}$ et sur le bruit $\mathbf{b}$.
2. On attribue une loi de probabilité *a priori* $p(\mathbf{f}|\boldsymbol{\theta}_1, \mathcal{H})$ aux inconnues du problème pour traduire la connaissance initiale sur ces inconnues. Cette loi peut dépendre d'un certain nombre de paramètres $\boldsymbol{\theta}_1$.
3. On attribue une loi de probabilité $p(\mathbf{g}|\mathbf{f}, \boldsymbol{\theta}_2; \mathcal{H})$ aux grandeurs mesurées pour traduire l'incertitude (due au bruit, aux erreurs de discrétisation et de quantification, au manque de précision de l'appareil de mesure, etc.) sur ces données. Cette loi aussi peut dépendre d'un certain nombre de paramètres $\boldsymbol{\theta}_2$. L'ensemble des paramètres $\boldsymbol{\theta} = (\boldsymbol{\theta}_1, \boldsymbol{\theta}_2)$ sont appelés les hyperparamètres du problème.
4. On utilise la règle de Bayes pour combiner l'information contenue dans les données et celle contenue dans la loi *a priori* et calculer ainsi la loi de probabilité *a posteriori*

$$p(\mathbf{f}|\mathbf{g}, \boldsymbol{\theta}; \mathcal{H}) = \frac{p(\mathbf{g}|\mathbf{f}, \boldsymbol{\theta}_2; \mathcal{H}) p(\mathbf{f}|\boldsymbol{\theta}_1, \mathcal{H})}{p(\mathbf{g}|\boldsymbol{\theta}; \mathcal{H})}$$

où

$$p(\mathbf{g}|\boldsymbol{\theta}; \mathcal{H}) = \int p(\mathbf{g}|\mathbf{f}, \boldsymbol{\theta}_2; \mathcal{H}) p(\mathbf{f}|\boldsymbol{\theta}_1, \mathcal{H}) d\mathbf{f}.$$

Cette loi contient toute l'information disponible sur les inconnues $\mathbf{f}$ après cette fusion.

5. On peut alors calculer n'importe quelle quantité. Par exemple, la moyenne *a posteriori* de $\mathbf{f}$:

$$\mathbb{E}[\mathbf{f}] = \int \mathbf{f} p(\mathbf{f}|\mathbf{g}, \boldsymbol{\theta}; \mathcal{H}) d\mathbf{f},$$

ou la moyenne de n'importe quelle autre fonction $h(\mathbf{f})$:

$$\mathbb{E}[h(\mathbf{f})] = \int h(\mathbf{f}) p(\mathbf{f}|\mathbf{g}, \boldsymbol{\theta}; \mathcal{H}) d\mathbf{f},$$

ou encore calculer la probabilité pour que $\underline{\mathbf{f}} < \mathbf{f} \leq \overline{\mathbf{f}}$:

$$P(\underline{\mathbf{f}} < \mathbf{f} \leq \overline{\mathbf{f}}) = \int_{\underline{\mathbf{f}}}^{\overline{\mathbf{f}}} p(\mathbf{f}|\mathbf{g}, \boldsymbol{\theta}; \mathcal{H}) d\mathbf{f}$$

ou seulement s'intéresser à une des composantes f_j de $\mathbf{f}$ et calculer

$$P(\underline{f}_j < f_j \leq \overline{f}_j) = \int_{\underline{f}_j}^{\overline{f}_j} p(f_j|\mathbf{g}, \boldsymbol{\theta}; \mathcal{H}) df_j,$$

où

$$p(f_j|\mathbf{g}, \boldsymbol{\theta}; \mathcal{H}) = \int \cdots \int p(\mathbf{f}|\mathbf{g}, \boldsymbol{\theta}; \mathcal{H}) df_1 \cdots df_{j-1} df_{j+1} \cdots df_n$$

est la loi *a posteriori marginale*.

Lorsque $\mathbf{f}$ est un scalaire ou un vecteur avec seulement deux composantes, on peut même tracer $p(\mathbf{f}|\mathbf{g}, \boldsymbol{\theta}; \mathcal{H})$ ou $p(f_1, f_2|\mathbf{g}, \boldsymbol{\theta}; \mathcal{H})$, ce qui permettra d'avoir une idée plus précise sur la forme de ces distribution.

Mais en pratique, la manipulation d'une densité de probabilité n'est pas très comode dès lors qu'on a plus de deux variables. C'est pourquoi, on se contente souvent de fournir une description sommaire de la loi *a posteriori* en indiquant

– son mode :

$$M[\mathbf{f}] = \arg \max_{\mathbf{f}} \{p(\mathbf{f}|\mathbf{g}, \boldsymbol{\theta}; \mathcal{H})\},$$

– sa moyenne :

$$E[\mathbf{f}] = \int \mathbf{f} p(\mathbf{f}|\mathbf{g}, \boldsymbol{\theta}; \mathcal{H}) d\mathbf{f},$$

– les modes de ses marginales :

$$M[f_j] = \arg \max_{f_j} \{p(f_j|\mathbf{g}, \boldsymbol{\theta}; \mathcal{H})\},$$

– les α -quantiles de ses marginales :

$$Q_\alpha(f_j) : P(f_j \leq Q_\alpha(f_j)) = \alpha,$$

– les médianes de ses marginales :

$$\text{Méd}[f_j] = Q_{\frac{1}{2}}(f_j) : P(f_j \leq \text{Méd}[f_j]) = \frac{1}{2},$$

– les régions de plus haute probabilité ou les supports α -interquantiles :

$$[a, b] = [Q_{(1-\alpha)/2}(f_j), Q_{(1+\alpha)/2}(f_j)] : P(a \leq f_j \leq b) = 1 - \alpha,$$

– les variances :

$$\text{Var}[f_j] = \int (f_j - \bar{f}_j)^2 p(f_j|\mathbf{g}, \boldsymbol{\theta}; \mathcal{H}),$$

où

$$\bar{f}_j = E[f_j] = \int f_j p(f_j|\mathbf{g}, \boldsymbol{\theta}; \mathcal{H}),$$

– les covariances :

$$\text{Cov}[f_j, f_k] = \iint (f_j - \bar{f}_j)(f_k - \bar{f}_k) p(f_j, f_k|\mathbf{g}, \boldsymbol{\theta}; \mathcal{H}) df_j df_k.$$

6. Une autre approche fondée sur la théorie de la décision consiste en la définition d'un estimateur $\hat{\mathbf{f}}$ qui minimiserait la moyenne d'une fonction de coût $C(\hat{\mathbf{f}}, \mathbf{f})$. Nous reviendront sur cette approche et les choix de ces fonctions coût, ainsi que sur les liens qu'existe entre ces estimateurs et les caractéristiques de la loi *a posteriori*.
7. Finalement, une fois une solution choisie, il est impératif d'y attribuer un degré de confiance ainsi que de déterminer la sensibilité de cette solution vis-à-vis des erreurs de mesures et celles de la modélisation. Heureusement, en principe, cette approche bayésienne nous fournit tout ce qui est nécessaire pour pouvoir satisfaire cette exigence. Car disposant de la loi *a posteriori* on peut calculer, par exemple, la covariance de l'erreur de l'estimation, ce qui nous permet de mettre des barres d'erreurs sur les solutions estimées.

Voyons maintenant quelles sont les difficultés de l'utilisation de cette approche dans un problème réel et quelles sont les solutions existantes pour les résoudre. On peut classer ces difficultés dans les rubriques qui suivent :

1. choix de la forme de la loi *a priori* $p(\mathbf{f}|\boldsymbol{\theta}_1; \mathcal{H})$;

2. choix ou estimation de ses paramètres θ_1 ;
3. choix de la forme de la loi $p(\mathbf{g}|\boldsymbol{\beta}, \theta_2; \mathcal{H})$;
4. choix ou estimation de ses paramètres θ_2 ;
5. calcul effectif de la loi *a posteriori* $p(\mathbf{f}|\mathbf{g}, \theta; \mathcal{H})$;
6. choix d'une fonction de coût $C(\hat{\mathbf{f}}, \mathbf{f})$ et calcul effectif de l'estimateur qu'en découle ;
7. détermination ou estimation *a posteriori* des hyperparamètres du problème, *i.e.* les paramètres $\theta = (\theta_1, \theta_2)$ des lois de probabilités directes $p(\mathbf{f}|\theta_1; \mathcal{H})$ et $p(\mathbf{g}|\mathbf{f}, \theta_2; \mathcal{H})$, à partir des données, c'est-à-dire en même temps que la solution ;

L'objectif principal de ce chapitre est de faire état de l'art concernant les solutions proposées pour résoudre ces problèmes et de fournir quelques exemples concrets (étude de cas) pour illustrer les difficultés et les performances des solutions proposées.

Mais, avant d'aller plus loin prenons un exemple, désormais classique, qui est le cas linéaire gaussien qui nous permettra de comprendre l'essentiel de ces différentes étapes.

9.2 Cas linéaire gaussien

S'il y a un cas où les calculs peuvent se faire facilement jusqu'au bout et où l'on peut facilement comprendre l'approche bayésienne est le cas où le modèle reliant les inconnues $\mathbf{f}$ et les données $\mathbf{g}$ est linéaire et où on peut attribuer des lois de probabilités gaussiennes à $\mathbf{f}$ et aux mesures $\mathbf{g}$.

Nous allons suivre les étapes présentées en introduction pour la résolution du problème linéaire

$$\mathbf{g} = \mathbf{H}\mathbf{f} + \mathbf{b} \quad (9.1)$$

- Information *a priori* sur $\mathbf{f}$:

Faisons l'hypothèse que nous puissions *a priori* connaître seulement la moyenne $E[\mathbf{f}] = \mathbf{f}_0$ et la matrice de covariance $E[(\mathbf{f} - \mathbf{f}_0)(\mathbf{f} - \mathbf{f}_0)^t] = \mathbf{R}_f = \sigma_f^2 \mathbf{P}_0$ de $\mathbf{f}$.

Nous pouvons alors faire l'hypothèse que $p(\mathbf{f}) = \mathcal{N}(\mathbf{f}_0, \sigma_f^2 \mathbf{R}_f)$:

$$p(\mathbf{f}) = A \exp \left[\frac{-1}{2} (\mathbf{f} - \mathbf{f}_0)^t \mathbf{R}_f^{-1} (\mathbf{f} - \mathbf{f}_0) \right] = A \exp \left[\frac{-1}{2\sigma_f^2} (\mathbf{f} - \mathbf{f}_0)^t \mathbf{P}_0^{-1} (\mathbf{f} - \mathbf{f}_0) \right] \quad (9.2)$$

avec $A = (2\pi)^{-n/2} |\mathbf{R}_f|^{-1/2} = (2\pi\sigma_f^2)^{-n/2} |\mathbf{P}_0|^{-1/2}$.

D'ailleurs, cette hypothèse peut être validée par le principe du maximum d'entropie.

- Information *a priori* sur $\mathbf{b}$:

Faisons l'hypothèse que $E[\mathbf{b}] = \mathbf{0}$ et que sa matrice de covariance $E[\mathbf{b}\mathbf{b}^t] = \mathbf{R}_b = \sigma_b^2 \mathbf{I}$. Cette hypothèse est bien raisonnable. En effet, supposer que $E[\mathbf{b}] = \mathbf{0}$ signifie que nous faisons l'hypothèse qu'il n'y a pas d'erreur systématique et l'hypothèse $E[\mathbf{b}\mathbf{b}^t] = \sigma_b^2 \mathbf{I}$ signifie que le bruit est supposé blanc (non corrélé).

Avec les mêmes arguments que dans le cas précédent on fait alors l'hypothèse que $p(\mathbf{b}) = \mathcal{N}(\mathbf{0}, \sigma_b^2 \mathbf{R}_b)$. Maintenant, partant du modèle (9.1) et faisant l'hypothèse que le bruit $\mathbf{b}$ est aussi indépendant du $\mathbf{f}$ on peut facilement déduire

$$p(\mathbf{g}|\mathbf{f}) = B \exp \left[\frac{-1}{2} (\mathbf{g} - \mathbf{H}\mathbf{f})^t \mathbf{R}_b^{-1} (\mathbf{g} - \mathbf{H}\mathbf{f}) \right] = B \exp \left[\frac{-1}{2\sigma_b^2} (\mathbf{g} - \mathbf{H}\mathbf{f})^t (\mathbf{g} - \mathbf{H}\mathbf{f}) \right] \quad (9.3)$$

avec $B = (2\pi)^{-m/2} |\mathbf{R}_b|^{-1/2} = (2\pi\sigma_b^2)^{-m/2}$.

- Règle de Bayes :

Utilisant la règle de Bayes, il est facile de voir que

$$p(\mathbf{f}|\mathbf{g}) = C \exp \left[\frac{-1}{2} J(\mathbf{f}) \right], \quad (9.4)$$

avec

$$\begin{aligned} J(\mathbf{f}) &= (\mathbf{g} - \mathbf{H}\mathbf{f})^t \mathbf{R}_b^{-1} (\mathbf{g} - \mathbf{H}\mathbf{f}) + (\mathbf{f} - \mathbf{f}_0)^t \mathbf{R}_f^{-1} (\mathbf{f} - \mathbf{f}_0) \\ &= \frac{1}{\sigma_b^2} \left[(\mathbf{g} - \mathbf{H}\mathbf{f})^t (\mathbf{g} - \mathbf{H}\mathbf{f}) + \lambda (\mathbf{f} - \mathbf{f}_0)^t \mathbf{P}_0^{-1} (\mathbf{f} - \mathbf{f}_0) \right], \end{aligned}$$

où $\lambda = \frac{\sigma_b^2}{\sigma_f^2}$.

Avec un peu de calcul il n'est pas difficile de montrer que cette loi *a posteriori* est aussi gaussienne et on a

$$p(\mathbf{f}|\mathbf{g}) = \mathcal{N}(\hat{\mathbf{f}}, \hat{\mathbf{P}})$$

où

$$\begin{cases} \hat{\mathbf{f}} &= \mathbf{R}_f \mathbf{H}^t (\mathbf{H} \mathbf{R}_f \mathbf{H}^t + \mathbf{R}_b)^{-1} \mathbf{g} = \hat{\mathbf{P}} \mathbf{H}^t \mathbf{R}_b^{-1} \mathbf{g}, \\ \hat{\mathbf{P}} &= \mathbf{R}_f - \mathbf{R}_f \mathbf{H}^t (\mathbf{H} \mathbf{R}_f \mathbf{H}^t + \mathbf{R}_b)^{-1} \mathbf{H} \mathbf{R}_f = (\mathbf{R}_f^{-1} + \mathbf{H}^t \mathbf{R}_b^{-1} \mathbf{H})^{-1}, \end{cases}$$

ou encore

$$\begin{cases} \hat{\mathbf{f}} &= (\mathbf{H}^t \mathbf{H} + \lambda \mathbf{P}_0^{-1})^{-1} \mathbf{H}^t \mathbf{g}, \\ \hat{\mathbf{P}} &= \sigma_b^2 (\mathbf{H}^t \mathbf{H} + \lambda \mathbf{P}_0^{-1})^{-1}. \end{cases}$$

- Caractéristique de la solution :

Nous avons ainsi une expression analytique pour la loi *a posteriori* et nous pouvons facilement en déduire tout ce que l'on souhaite. Par exemple nous savons que la moyenne, la médiane et la mode sont identiques et égaux à $\hat{\mathbf{f}}$.

Ainsi tous ces estimateurs : le maximum *a posteriori* (MAP) ou la moyenne *a posteriori* (MP) sont identiques et égaux à $\hat{\mathbf{f}}$ qui peut aussi être calculé par

$$\hat{\mathbf{f}} = \arg \max_{\mathbf{f}} \{p(\mathbf{f}|\mathbf{g})\} = \arg \min_{\mathbf{f}} \{J(\mathbf{f}) = Q(\mathbf{f}) + \lambda \Omega(\mathbf{f})\}$$

avec

$$Q(\mathbf{f}) = (\mathbf{g} - \mathbf{H}\mathbf{f})^t (\mathbf{g} - \mathbf{H}\mathbf{f})$$

et

$$\Omega(\mathbf{f}) = (\mathbf{f} - \mathbf{f}_0)^t \mathbf{P}_0^{-1} (\mathbf{f} - \mathbf{f}_0),$$

ou encore, si on note par $\mathbf{P}_0^{-1} = \mathbf{D}^t \mathbf{D}$ on a

$$\begin{aligned} J(\mathbf{f}) &= (\mathbf{g} - \mathbf{H}\mathbf{f})^t (\mathbf{g} - \mathbf{H}\mathbf{f}) + \lambda (\mathbf{f} - \mathbf{f}_0)^t \mathbf{D}^t \mathbf{D} (\mathbf{f} - \mathbf{f}_0) \\ &= \|\mathbf{g} - \mathbf{H}\mathbf{f}\|^2 + \lambda \|\mathbf{D}(\mathbf{f} - \mathbf{f}_0)\|^2. \end{aligned}$$

On retrouve alors la notion de solution régularisée et la régularisation quadratique avec les avantages suivantes :

- Dans l'approche régularisation déterministe le choix des distances $\Delta_1(\mathbf{g}, \mathbf{H}\mathbf{f}) = \|\mathbf{g} - \mathbf{H}\mathbf{f}\|^2$ et de $\Delta_2(\mathbf{f}, \mathbf{f}_0) = \|\mathbf{D}(\mathbf{f} - \mathbf{f}_0)\|^2$ est plutôt descriptif et un peu arbitraire, alors que dans l'approche bayésienne, ces termes sont, respectivement, les conséquences des hypothèses faites sur la loi du bruit et sur la loi *a priori*.
- Dans régularisation déterministe, le coefficient de régularisation λ est un paramètre qui est choisi empiriquement, alors qu'ici, $\lambda = \frac{\sigma_b^2}{\sigma_f^2}$. Si l'on connaît les deux variances σ_b^2 et σ_f^2 , sa valeur est fixée. Mais, bien sûr, en pratique on ne connaît pas ces deux valeurs et on est encore amené à les fixer par expérience. Cependant, rien que l'expression de λ nous permet de mieux comprendre son comportement : plus le niveau du bruit est important, plus il faut choisir sa valeur grande pour obtenir un résultat satisfaisant. Nous verrons que l'approche bayésienne nous donne les outils nécessaires pour la déterminer (voir le chapitre sur l'estimation des hyperparamètres).
- Dans le cadre de la régularisation déterministe, si on se demande qu'elle est le degré de confiance que l'on peut avoir dans la solution $\hat{\mathbf{f}}$ proposée, nous n'avons pas d'outil convenable pour répondre. Dans l'approche bayésienne, nous pouvons répondre à cette question, par exemple en calculant la matrice de covariance *a posteriori* $\mathbf{P} = \mathbf{E}[(\mathbf{f} - \hat{\mathbf{f}})^2]$. Sachant que les éléments diagonaux $P_{jj} = \mathbf{E}[(f_j - \hat{f}_j)^2]$ sont les variances *a posteriori*, nous pouvons les utiliser pour mettre des barres d'erreur sur la solution. Les éléments non-diagonale $P_{ij} = \mathbf{E}[(f_i - \hat{f}_i)(f_j - \hat{f}_j)]$ peuvent aussi nous renseigner sur le lien (corrélation) qui peut exister entre l'élément f_i et l'élément f_j .

9.3 Calcul de la loi *a posteriori*

9.3.1 Cas gaussien

Nous avons vu que dans le cas d'un modèle linéaire gaussien, la loi *a posteriori* est aussi gaussienne. Sachant qu'une loi gaussienne est entièrement définie par ses deux premiers moments, le problème devient assez simple. Nous disposons des expressions analytiques pour la moyenne et la matrice de covariance *a posteriori*, expressions dont les calculs nécessitent soit l'inversion des matrices soit l'optimisation d'un critère quadratique. Il reste cependant le coût de calcul. Souvent, on peut profiter de la structure des matrices à inverser ou de la forme quadratique du critère à optimiser pour fournir des algorithmes spécifiques pour effectuer ces calculs.

9.3.2 Cas général

Dans le cas général, le calcul complet de la loi *a posteriori* peut devenir plus délicat, et souvent, on se contente, soit de l'approximer par une loi gaussienne et de calculer la moyenne et la matrice de covariance, soit de définir un estimateur ponctuel à partir de cette loi, comme par exemple le *maximum a posteriori* (MAP) :

$$\hat{\mathbf{f}}_{\text{MAP}} = \arg \max_{\mathbf{f}} \{p(\mathbf{f}|\mathbf{g})\} = \arg \min_{\mathbf{f}} \{-\ln p(\mathbf{f}|\mathbf{g})\},$$

ou la *moyenne a posteriori* (en Anglais posterior mean PM) :

$$\hat{\mathbf{f}}_{\text{PM}} = \int \mathbf{f} p(\mathbf{f}|\mathbf{g}) d\mathbf{f},$$

ou encore le *maximum a posteriori marginal* (MAPmarginal) :

$$\hat{f}_j = \arg \max_{f_j} \{p(f_j|\mathbf{g})\},$$

où

$$p(f_j|\mathbf{g}, \boldsymbol{\theta}; \mathcal{H}) = \int \cdots \int p(\mathbf{f}|\mathbf{g}, \boldsymbol{\theta}; \mathcal{H}) df_1 \cdots df_{j-1} df_{j+1} \cdots df_n.$$

Notons aussi que le calcul des estimateurs PM et MAPmarginal nécessite le calcul des intégrales de dimension très élevée, alors que le calcul de l'estimateur MAP ne nécessite qu'une optimisation.

Notons aussi que dans le cas gaussien tous ces estimateurs sont identiques. En revanche, dans les autres cas, ils peuvent fournir des résultats très différents. Nous reviendrons sur les propriétés de ces estimateurs et les outils qui existent pour les calculer dans les chapitres suivants.

9.4 Choix des fonctions de coût

Dans l'approche théorie de la décision nous avons vu que, partant de la loi *a posteriori*, on peut définir une solution au problème, par l'intermédiaire d'une fonction de coût $C(\hat{\mathbf{f}}, \mathbf{f})$. En effet, dans cette approche on définit l'estimateur optimal $\hat{\mathbf{f}}$ comme l'argument qui minimise le coût moyen :

$$\hat{\mathbf{f}} = \arg \min_{\mathbf{z}} \{E[C(\mathbf{z}, \mathbf{f})]\} = \arg \min_{\mathbf{z}} \left\{ \int C(\mathbf{z}, \mathbf{f}) p(\mathbf{f}|\mathbf{g}) d\mathbf{f} \right\}.$$

Nous allons voir que suivant le choix de la fonction de coût on obtient des estimateurs différents, et en particulier, les estimateurs MAP, PM et MAPmarginal :

- Coût quadratique et estimateur PM :

$$C(\mathbf{f}, \hat{\mathbf{f}}) = (\mathbf{f} - \hat{\mathbf{f}})^t \mathbf{Q} (\mathbf{f} - \hat{\mathbf{f}})^t \longrightarrow \hat{\mathbf{f}} = \int \mathbf{f} p(\mathbf{f}|\mathbf{g}) d\mathbf{f}.$$

- Coût dirac et estimateur MAP :

$$C(\mathbf{f}, \hat{\mathbf{f}}) = 1 - \delta(\mathbf{f} - \hat{\mathbf{f}}) \longrightarrow \hat{\mathbf{f}} = \arg \max_{\mathbf{f}} \{p(\mathbf{f}|\mathbf{g})\}.$$

- Coût produit de diracs et estimateur MAPmarginal :

$$C(\mathbf{f}, \hat{\mathbf{f}}) = \prod_j 1 - \delta(x_j - \hat{x}_j) \longrightarrow \hat{f}_j = \arg \max_{x_j} \{p(x_j|\mathbf{g})\},$$

D'autres fonctions de coût peuvent aussi être utilisées, mais leur usage reste plus spécifique.

9.5 Choix de la loi *a priori*

Le problème de la conversion d'une information *a priori* en une loi de probabilité est un problème difficile et encore largement ouvert.

La principale difficulté réside dans le fait que rarement, en pratique, l'information *a priori* se présente directement sous une forme probabiliste. Par exemple, on nous dit que la grandeur recherchée est *positive ou bornée entre $[0, 1]$* , de variation *douce, croissante ou décroissante*, à *bande limitée ou à spectre limitée*, de variation *douce ou constante par morceaux ou par région*, de variation *impulsionnelle*, etc. La tâche, ensuite, est de construire des lois de probabilités qui puissent incorporer ces qualificatifs.

Les méthodes existantes peuvent être schématiquement regroupées en trois grandes classes.

Certaines reposent sur la théorie des *groupes de transformation* pour déterminer la mesure de référence “naturelle” pour la grandeur recherchée. On utilise ces méthodes surtout lorsqu'on sait peu de choses sur la grandeur recherchée, c'est à dire lorsque notre information *a priori* se résume à une connaissance qualitative sur la nature de la grandeur recherchée. Par exemple, nous savons que la grandeur recherchée représente un paramètre de *localisation* pouvant prendre une valeur quelconque sur l'axe des réels $\mathbf{R}$ ou un paramètre *d'échelle* qui est par nature positif et pouvant prendre une valeur quelconque sur la demi-droite des réels positifs $\mathbf{R}_+$.

Mais en pratique, cette approche a permis de justifier après coup l'emploi de la mesure de Lebesgue pour les paramètres de localisation (fournissant ainsi une extension au cas continu de la distribution uniforme résultant de l'application du “Principe d'indifférence” de Bernoulli dans le cas discret) et de la mesure de Jeffreys dans le cas des paramètres d'échelle.

D'autres méthodes reposent sur des principes informationnel. Il s'agit principalement des méthodes dites “à maximum d'entropie” dans lesquelles on recherche une distribution qui soit la plus proche (au sens d'une distance de Kullback) d'une distribution de référence (souvent choisie par l'approche précédente) tout en vérifiant une information incomplète connue *a priori* sous la forme des moments (des contraintes linéaires) sur la loi recherchée. Mais là encore, cette approche a surtout permis de justifier après coup certains choix. De plus, elle n'est véritablement praticable que lorsque cette information *a priori* est faite de contraintes linéaires sur la distribution recherchée. On trouve alors la famille des *distributions exponentielles* comme nous l'avons vu dans l'étude du principe du maximum d'entropie.

Il existe enfin une dernière classe très importante, celle des constructions faites “à la main”. C'est dans cette catégorie qu'entrent par exemple les modèles markoviens qui ont connu un développement spectaculaire depuis 1984 en traitement d'image, et qui permettent d'incorporer dans une distribution *a priori* des propriétés locales essentielles que doit posséder l'objet. La construction de ces modèles demande beaucoup de savoir-faire.

Bien entendu, il est hors de question de vouloir détailler toutes ces méthodes dans le cadre de ce document. Nous nous limiterons, dans la section suivante, à une description succincte des deux premières approches, plutôt à travers de quelques exemples que d'une manière formelle. Mais, en ce qui concerne la modélisation markovienne, nous lui réservons un chapitre tout entier.

9.5.1 Maximum d'entropie

Nous avons vu dans la section (8.3) que le principe du maximum d'entropie peut être utilisé pour attribuer une loi de probabilité à une grandeur traduisant une information *a priori* sous forme de moments. Nous allons illustrer ceci à travers quelques exemples.

Supposons que nous ayons un voltmètre et que nous cherchions à mesurer la tension de la ville $X(t)$ à différents instants. Supposons que notre instrument soit un voltmètre idéal et ce que nous mesurons $Y(t)$ peut être modélisé par

$$Y(t) = X(t) + B(t)$$

où $B(t)$ représente à la fois les erreurs de lecture et le bruit externe. Nous voudrions appliquer l'approche bayésienne pour fournir une estimation $\hat{X}(t)$ de $X(t)$. Considérons tout d'abord que nous avons fait juste une mesure y_1 à partir de laquelle nous voudrions estimer X . Nous avons vu que la première étape avant d'appliquer la règle de Bayes est d'attribuer les lois de probabilité $p(x)$ et $p(y|x)$.

Concernant X , nous savons que sa moyenne est de 220 volts et que sa variation autour de cette moyenne est ± 5 volts (ceci dépend sûrement de pays et de la qualité des installations, mais supposons que nous soyons en France!). Comment traduire alors cette information en une loi de probabilité? C'est exactement dans ce cadre que le principe du ME peut fournir une réponse satisfaisante :

$$E[X] = x_0 = 220 \text{ volts}, \quad \text{Var}[X] = \sigma_x^2 = 25 \text{ volts}^2 \quad \xrightarrow{\text{ME}} \quad X \simeq \mathcal{N}(x_0, \sigma_x^2)$$

Concernant B , il est naturel de faire l'hypothèse que sa moyenne est nulle (il n'y a pas d'erreur systématique) et supposons que l'on puisse avoir une idée de sa puissance. En langage d'ingénieur : "le bruit est d'environ 10 dB". Là aussi, le principe du ME peut fournir une réponse satisfaisante :

$$E[B] = 0, \quad \text{Var}[B] = \sigma_b^2 = 1 \text{ volts}^2 \quad \xrightarrow{\text{ME}} \quad B \simeq \mathcal{N}(0, \sigma_b^2).$$

Nous avons maintenant tous les ingrédients pour calculer la loi *a posteriori* :

$$X|Y \simeq \mathcal{N}(\hat{x}, \hat{\sigma}^2),$$

et il n'est pas difficile de calculer $\hat{x}$ et $\hat{\sigma}^2$:

$$\hat{x} = \frac{\sigma_x^2}{\sigma_x^2 + \sigma_b^2} y + \frac{\sigma_b^2}{\sigma_x^2 + \sigma_b^2} x_0$$

$$\hat{\sigma}^2 = \frac{\sigma_x^2 \sigma_b^2}{\sigma_x^2 + \sigma_b^2}.$$

Notons que

$$\begin{aligned} \text{si } \sigma_b^2 = 0 &\longrightarrow \hat{x} = y, \quad \hat{\sigma}^2 = 0 \quad \text{et} \\ \text{si } \sigma_x^2 = 0 &\longrightarrow \hat{x} = x_0, \quad \hat{\sigma}^2 = 0. \end{aligned}$$

Considérons maintenant que la grandeur X que nous voulons mesurer est l'énergie ou la puissance consommée dans un circuit et l'appareil de mesure est un wattmètre. Mais, avant même de choisir notre instrument nous savons que la grandeur à mesurer est positive et que son ordre de grandeur est d'environ 200 Watts :

$$X \geq 0, \quad E[X] = x_0 = 200 \text{ Watts} \quad \xrightarrow{\text{ME}} \quad X \simeq \mathcal{E}(x_0)$$

ce qui signifie :

$$p(x) = x_0 \exp [-x/x_0], \quad x > 0.$$

Si nous gardons les mêmes hypothèses sur B alors nous avons

$$p(y|x) \propto \exp \left[-\frac{1}{2\sigma_b^2} (y-x)^2 \right]$$

ce qui nous donne pour la loi *a posteriori* :

$$p(x|y) \propto \exp \left[-\frac{1}{2\sigma_b^2} \left((y-x)^2 + \frac{2\sigma_b^2}{x_0} x \right) \right], \quad x > 0.$$

En réarrangeant le terme en exponentiel, on peut facilement voir que cette loi est une loi gaussienne tronquée :

$$p(x|y) \propto \exp \left[-\frac{1}{2\sigma_b^2} (x - \hat{x})^2 \right], \quad x > 0, \quad \text{avec} \quad \hat{x} = y - \frac{\sigma_b^2}{x_0}$$

Revenons maintenant sur le premier exemple et considérons que nous avons fait M mesures aux intervalles réguliers $\mathbf{Y} = \{y_1, \dots, y_M\}$ et que nous voulons estimer $\mathbf{X} = \{X_1, \dots, X_M\}$.

Une première hypothèse consiste à dire qu'*a priori*, il n'y a pas de raison de faire l'hypothèse que la loi de X_i soit différente de la loi de X_j et que X_i soit dépendante de X_j (hypothèse dite i.i.d.) et d'attribuer alors une loi de probabilité jointe :

$$p(\mathbf{x}) = \prod_{j=1}^M p(x_j) \longrightarrow \mathbf{X} \simeq \mathcal{N}(\mathbf{x}_0, \sigma_x^2 \mathbf{I}).$$

La même hypothèse concernant les éléments de $\mathbf{B} = \{B_1, \dots, B_M\}$ nous permet d'écrire :

$$p(\mathbf{b}) = \prod_{j=1}^M p(b_j) \longrightarrow \mathbf{B} \simeq \mathcal{N}(\mathbf{0}, \sigma_b^2 \mathbf{I}).$$

Il n'est pas difficile de montrer que la loi *a posteriori* est :

$$p(\mathbf{x}|\mathbf{y}) = \prod_{j=1}^M p(x_j|y_j) \longrightarrow \mathbf{X}|\mathbf{Y} \simeq \mathcal{N}(\mathbf{Y}, \sigma^2 \mathbf{I})$$

Une deuxième hypothèse plus réaliste est de dire qu'*a priori*, X_j dépend de X_{j-1} (la valeur de tension à un instant ne peut pas être très différente de sa valeur à un pas d'échantillonnage plus loin). Comment peut-on alors traduire cette information? Supposons que nous puissions avoir une idée *a priori* sur le coefficient de corrélation entre X_j et X_{j-1} . Ainsi, si on résume tout ce que nous connaissons sur $\mathbf{X}$:

$$\mathbb{E}[X_j] = x_0, \quad \text{Var}[X_j] = \sigma_x^2, \quad \mathbb{E}[(X_j - x_0)(X_{j-1} - x_0)] = \beta \sigma_x^2.$$

Utilisant de nouveau le principe du ME, nous obtenons :

$$p(\mathbf{x}) \propto \exp \left[\frac{1}{2} (\mathbf{x} - \mathbf{x}_0)^t \mathbf{P}^{-1} (\mathbf{x} - \mathbf{x}_0) \right], \quad \text{avec} \quad \mathbf{P} = \sigma_x^2 \begin{pmatrix} 1 & \beta & & & \\ \beta & 1 & \beta & & \\ & \ddots & \ddots & \ddots & \\ & & & \beta & 1 & \beta \\ & & & & \beta & 1 \end{pmatrix}$$

Ainsi, d'une manière générale, l'utilisation du principe du maximum d'entropie nous conduit vers les lois exponentielles généralisées. En effet, si on suppose que notre information *a priori* est de la forme :

$$\mathbb{E}[\phi_k(\mathbf{X})] = d_k, \quad k = 1, \dots, K,$$

nous avons vu que la loi à entropie maximale qui satisfait ces contraintes est de la forme :

$$p(\mathbf{x}) = \frac{1}{Z} \exp \left[- \sum_{k=1}^K \lambda_k \phi_k(\mathbf{x}) \right],$$

où

$$Z(\lambda_1, \dots, \lambda_m) = \int p(\mathbf{x}) \exp \left[- \sum_{k=1}^K \lambda_k \phi_k(\mathbf{x}) \right] d\mathbf{x}$$

et où les $\{\lambda_1, \lambda_2, \dots, \lambda_n\}$ sont déterminés par :

$$-\frac{\partial \ln Z(\lambda_1, \dots, \lambda_m)}{\partial \lambda_k} = d_k, \quad k = 1, \dots, K.$$

Un cas particulier intéressant est lorsque les fonctions $\phi_k(x)$ sont séparables :

$$\phi_k(\mathbf{x}) = \sum_{j=1}^N \phi_{k_j}(x_j)$$

on a alors

$$p(\mathbf{x}) = \prod_{j=1}^N p(x_j), \quad \text{avec} \quad p(x_j) = \frac{1}{Z_j} \exp \left[- \sum_{k=1}^K \lambda_k \phi_{k_j}(x_j) \right],$$

ce qui signifie que les x_j sont indépendants, et si de plus $\phi_{k_j} = \phi_k$ on a

$$p(\mathbf{x}) = \prod_{j=1}^N p(x_j), \quad \text{avec} \quad p(x_j) = \frac{1}{z} \exp \left[- \sum_{k=1}^K \lambda_k \phi_k(x_j) \right],$$

ce qui signifie que les x_j sont i.i.d (indépendants et de même loi).

On trouve parfois l'appellation *lois entropiques* pour les lois appartenant à ces deux derniers cas. Quelques cas particuliers de ces lois utilisées fréquemment sont :

– Gaussienne généralisée :

$$\phi(x_j) = |x_j|^r, 1 < r < 2, \quad \longrightarrow \quad p(x_j) \propto \exp[-\alpha |x_j|^r].$$

La loi conjointe est alors de la forme :

$$p(\mathbf{x}) \propto \exp \left[-\alpha \sum_{j=1}^N |x_j|^r \right] \quad \longrightarrow \quad -\ln p(\mathbf{x}) = cte + \alpha \sum_{j=1}^N |x_j|^r.$$

– Loi Gamma :

$$\phi(x_j) = -\alpha \ln\left(\frac{x_j}{m_j}\right) + \left(\frac{x_j}{m_j}\right), \quad \longrightarrow \quad p(x_j) \propto \left(\frac{x_j}{m_j}\right)^{-\alpha} \exp \left[\frac{x_j}{m_j} \right].$$

et

$$p(\mathbf{x}) = \prod_{j=1}^N p(x_j) \propto \left(\prod_{j=1}^N \left(\frac{x_j}{m_j} \right)^{-\alpha} \right) \exp \left[\sum_{j=1}^N \frac{x_j}{m_j} \right]$$

et par conséquent

$$-\ln p(\mathbf{x}) = cte - \alpha \sum_{j=1}^N \ln \left(\frac{x_j}{m_j} \right) + \sum_{j=1}^N \frac{x_j}{m_j}.$$

– Loi Béta :

$$\phi(x_j) = \alpha \ln(x_j) + \beta \ln(1 - x_j), \quad \longrightarrow \quad p(x_j) \propto x_j^\alpha (1 - x_j)^\beta.$$

et la loi conjointe est :

$$p(\mathbf{x}) \propto \prod_{j=1}^N x_j^\alpha (1 - x_j)^\beta \quad \longrightarrow \quad -\ln p(\mathbf{x}) = cte + \alpha \sum_{j=1}^N \ln(x_j) + \beta \sum_{j=1}^N \ln(1 - x_j).$$

Nous avons vu que dans le cas i.i.d., la forme générale de ces lois est :

$$p(\mathbf{x}) \propto \exp \left[- \sum_{k=1}^K \lambda_k \phi_k(\mathbf{x}) \right], \quad \text{avec} \quad \phi_k(\mathbf{x}) = \sum_{j=1}^N \phi_k(x_j).$$

On peut se poser alors une question sur le choix des fonctions $\phi_k(x_j)$. Par exemple, on peut se demander s'il existe des familles de fonctions ϕ_k de telle sorte que la lois $p(x_j)$ et par conséquent la loi $p(\mathbf{x})$ aient une propriété, par exemple d'invariance par changement d'échelle. Dans une étude récente, certains auteurs se sont posés de telles questions et ont essayé d'y répondre. Ici, nous résumons quelques uns de ces résultats.

L'idée de base se trouve dans la recherche des fonctions ϕ_k de telle sorte que la famille de lois $p(x_j)$ correspondantes soit fermée pour un groupe de transformations comme le changement d'échelle par exemple.

Prenons le cas du changement d'échelle pour illustrer cette idée. Si x_j représente une quantité et $p(x_j)$ sa loi de probabilité, on voudrait que lors d'un changement d'échelle αx_j avec $\alpha > 0$, la loi $p(\alpha x_j)$ reste dans la même famille que $p(x_j)$. Autrement dit :

$$p(x_j) \propto \exp \left[- \sum_{k=1}^K \lambda_k \phi_k(x_j) \right] \longrightarrow p(\alpha x_j) \propto \exp \left[- \sum_{k=1}^K \lambda'_k \phi_k(x_j) \right]$$

où λ'_k ne doivent dépendre que des λ_k et de facteur d'échelle α .

Il est alors facile de montrer que pour satisfaire à cette propriété il suffit d'avoir :

$$\forall x_j \text{ \& } \forall \alpha > 0, \quad \sum_{k=1}^K \lambda_k \phi_k(x_j) = cte + \sum_{k=1}^K \lambda'_k \phi_k(x_j)$$

L'étude de cette propriété pour $K = 1$ et $K = 2$ donne les familles suivantes :

Lois à un paramètre $K = 1$:

$$\left\{ \phi(x) \right\} = \left\{ x^r, \ln x \right\}, \quad r > 0.$$

Lois à deux paramètres $K = 2$:

$$\left\{ (\phi_1(x), \phi_2(x)) \right\} = \left\{ (x^{r_1}, x^{r_2}), (x^{r_1}, \ln x), (x^{r_1}, x^{r_1} \ln x), (\ln x, \ln^2 x) \right\},$$

où r , r_1 et r_2 sont des réels positifs.

On retrouve quelques cas particuliers intéressants en choisissant $\phi_2(x) = x$ et pour différents choix de $\phi_1(x)$:

– $\phi_1(x) = x^2$; on retrouve la loi gaussienne :

$$p(x) \propto \exp \left[-\lambda x^2 - \mu x \right] \propto \exp \left[-\lambda \left(x + \frac{\mu}{2\lambda} \right)^2 \right] = \mathcal{N} \left(m = \frac{-\mu}{\lambda}, \sigma^2 = \frac{1}{2\lambda} \right).$$

– $\phi_1(x) = \ln x$; on retrouve la loi gamma :

$$p(x) \propto \exp \left[-\lambda \ln x - \mu x \right] = x^{-\lambda} \exp \left[-\mu x \right].$$

– $\phi_1(x) = x \ln x$; on retrouve une loi dite de la forme entropie de Shannon :

$$p(x) \propto \exp \left[-\lambda x \ln x - \mu x \right].$$

9.5.2 Modèles markoviens

La modélisation markovienne a pris une importance sans égale en traitement d'image depuis les travaux de Geman et Geman. L'objet de ce document n'est pas de donner une présentation complète de la modélisation markovienne, mais plutôt une présentation différente.

Nous avons vu dans la section précédente que, d'une manière générale, l'utilisation du principe du maximum d'entropie avec une information *a priori* de la forme :

$$E[\phi_k(\mathbf{x})] = d_k, \quad k = 1, \dots, K,$$

nous conduit aux lois de la forme :

$$p(\mathbf{x}) = \frac{1}{Z} \exp \left[- \sum_{k=1}^K \lambda_k \phi_k(\mathbf{x}) \right].$$

Considérons maintenant le cas où les fonctions $\phi_k(\mathbf{x})$ sont de la forme

$$\phi_k(\mathbf{x}) = \sum_j \sum_{i \in V_j} \phi_k(x_j, x_i)$$

où V_j signifie les (échantillons ou les pixels) voisins de j . Les fonctions ϕ_k sont alors appelées *potentiels* et

$$U(\mathbf{x}) = \sum_j \sum_{i \in V_j} \phi_k(x_j, x_i)$$

l'énergie.

Un tel choix pour ces fonctions signifie que nous faisons l'hypothèse qu'il y a une corrélation locale entre les valeurs du signal aux instants successifs ou de l'image sur des pixels voisins.

Un cas particulier très courant est lorsque $k = 1$ et $\phi_k(x_j, x_i) = \phi_k(x_j - x_i)$. On a alors

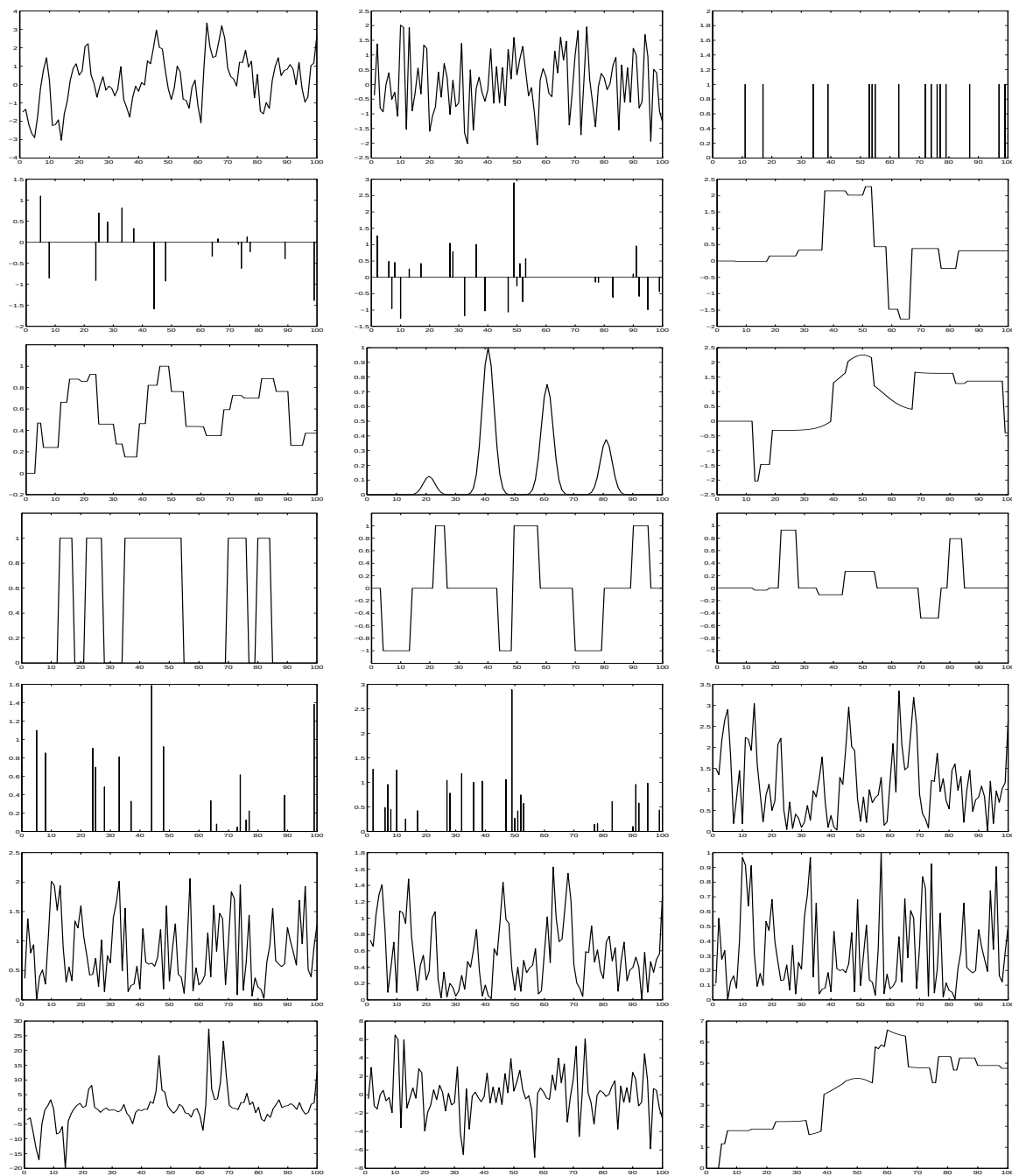
$$p(\mathbf{x}) \propto \exp \left[-\lambda \sum_j \sum_{i \in V_j} \phi(x_j - x_i) \right]$$

Le choix de la fonction potentiel ϕ est crucial : les fonctions convexes permettront de modéliser des signaux et des champs continus et les fonctions non convexes permettront de modéliser des signaux et des champs qui contiennent des discontinuités. Une liste de ces fonctions classiquement utilisées en traitement du signal et des images est fournie en annexe D.

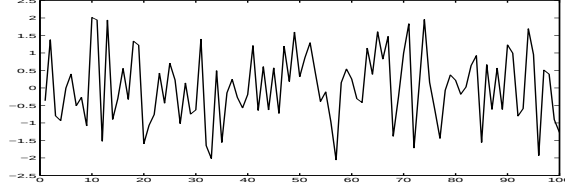
9.6 Quelques exemples de construction à la main

Dans cette section, nous allons montrer comment on peut choisir une loi de probabilité *a priori* à la main et avec du bon sens. La figure qui suit montre un ensemble de signaux très différents.

Nous allons, pour chacune de ces formes de signaux proposer une famille de lois convenable.

FIG. 9.1 – *Différents types de signaux.*

1. Modèle gaussien blanc pour des signaux continus à variation rapide



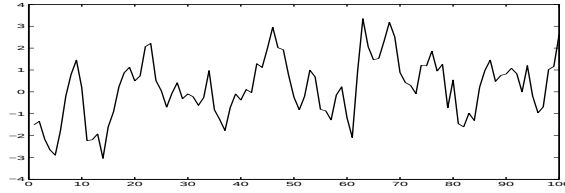
En regardant la forme de ce signal et en constatant que les échantillons du signal sont réparties d'une manière symétrique autour de zéro, on peut proposer une loi gaussienne centrée :

$$\mathbf{x} \sim \mathcal{N}(\mathbf{0}, \mathbf{R}_x = \sigma_x^2 \mathbf{I}) \longrightarrow p(\mathbf{x}) \propto \exp \left[-\frac{1}{2\sigma_x^2} \|\mathbf{x}\|^2 \right]$$

ou encore

$$p(\mathbf{x}) \propto \exp \left[-\frac{1}{2\sigma_x^2} \sum_j x_j^2 \right].$$

2. Modèle gaussien coloré pour des signaux continus à variation lente :



En regardant la forme de ce signal, et en la comparant au cas précédent, on peut proposer une loi centrée, gaussienne mais colorée :

$$\mathbf{x} \sim \mathcal{N}(\mathbf{0}, \mathbf{R}_x = \sigma_x^2 \mathbf{P}_0) \longrightarrow p(\mathbf{x}) \propto \exp \left[-\frac{1}{2\sigma_x^2} \mathbf{x}^t \mathbf{P}_0^{-1} \mathbf{x} \right],$$

avec $\mathbf{P}_0$ une matrice symétrique et définie positive, ce qui permet d'écrire $\mathbf{P}_0^{-1} = \mathbf{D}^t \mathbf{D}$ et on a alors

$$p(\mathbf{x}) \propto \exp \left[-\frac{1}{2\sigma_x^2} \|\mathbf{D}\mathbf{x}\|^2 \right].$$

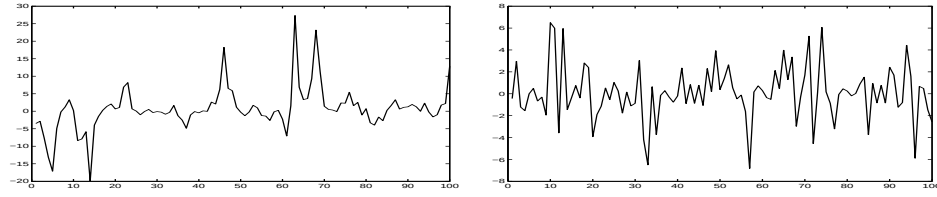
Un choix possible pour $\mathbf{D}$ est la matrice de différences finies d'ordre un :

$$\mathbf{D}_1 = \begin{pmatrix} 1 & 0 & \cdots & \cdots & 0 \\ -1 & 1 & \ddots & & \vdots \\ 0 & -1 & 1 & \ddots & \vdots \\ \vdots & \ddots & & & 0 \\ 0 & \cdots & 0 & -1 & 1 \end{pmatrix}$$

ce qui permet d'écrire :

$$p(\mathbf{x}) \propto \exp \left[-\frac{1}{2\sigma_x^2} \sum_j (x_j - x_{j-1})^2 \right]$$

3. Modèle gaussien généralisé pour des signaux continus plutôt impulsions



En regardant la forme de ce signal, et en la comparant au cas précédent, on peut constater que la proportion des échantillons ayant des valeurs proches de zéro est plus importante que celle des échantillons ayant des valeurs plus grandes et que cette proportion semble décroître très rapidement. On peut alors proposer une loi gaussienne généralisée :

$$p(\mathbf{x}) \propto \exp[-\beta \|\mathbf{D}_k \mathbf{x}\|^p], \quad 1 < p \leq 2,$$

avec $\mathbf{D}_0 = \mathbf{I}$ et $\mathbf{D}_k$ la matrice des différences finies d'ordre k .

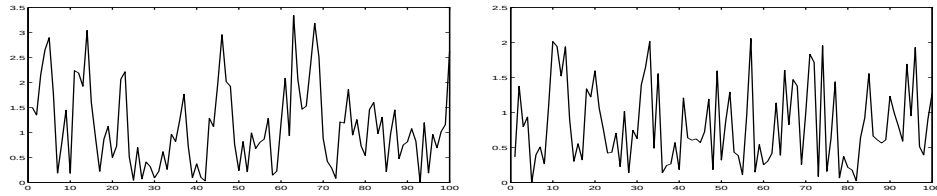
Dans le cas $\mathbf{D}_0 = \mathbf{I}$ on a

$$p(\mathbf{x}) \propto \exp \left[-\beta \sum_j |x_j|^p \right]$$

et dans le cas $\mathbf{D}_1$ on a

$$p(\mathbf{x}) \propto \exp \left[-\beta \sum_j |x_j - x_{j-1}|^p \right].$$

4. Modèle des signaux positifs



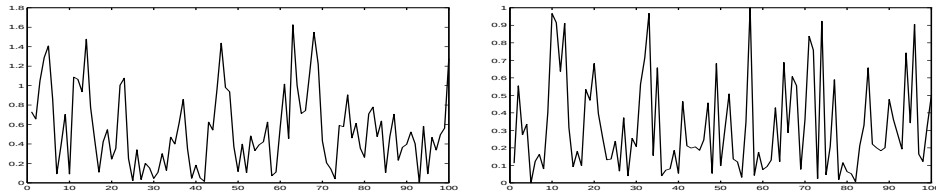
Loi gaussienne généralisée tronquée :

$$p(\mathbf{x}) \propto \exp \left[-\beta \sum_j x_j^p \right], \quad x_j > 0.$$

Loi Gamma :

$$p(\mathbf{x}) \propto \exp \left[-\alpha \sum_j \ln x_j - \beta \sum_j x_j \right], \quad x_j > 0.$$

5. Modèle des signaux continus mais bornés entre 0 et 1 : Loi Béta

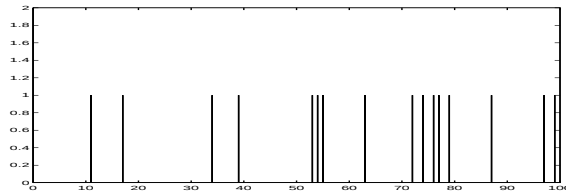


$$p(\mathbf{x}) = \prod_{j=1}^N p(x_j),$$

avec

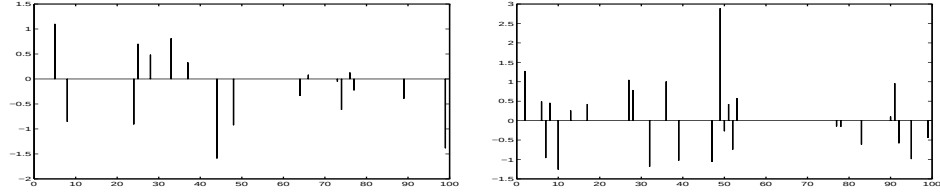
$$\begin{aligned} p(x_j) &\propto x_j^{-\alpha} - (1 - x_j)^{-\beta} \\ &\propto \exp[-\alpha \ln x_j - \beta \ln(1 - x_j)] \\ -\ln p(\mathbf{x}) &= cte + \alpha \ln x_j + \beta \ln(1 - x_j) \end{aligned}$$

6. Modèle de Bernoulli :



$$\begin{cases} P(X_j = 1) = \alpha, \\ P(X_j = 0) = 1 - \alpha \end{cases} \longrightarrow P(\mathbf{X} = \mathbf{x}) = \alpha^s (1 - \alpha)^{1-s} \text{ avec } s = \sum (x_j)$$

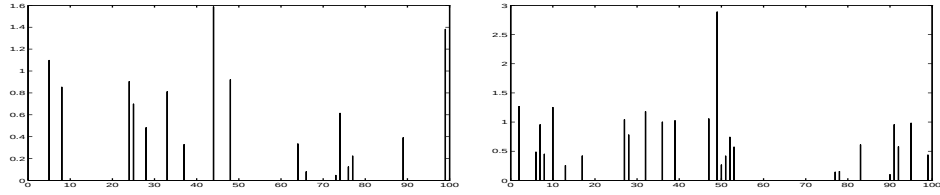
7. Modèle Bernoulli-Gaussien :



$$p(x_j|q_j) = \begin{cases} 0 & \text{si } q_j = 0, \\ \mathcal{N}(0, \sigma_x^2) & \text{si } q_j = 1 \end{cases} \quad \text{et} \quad \begin{cases} P(q_j = 1) = \lambda, \\ P(q_j = 0) = 1 - \lambda \end{cases}$$

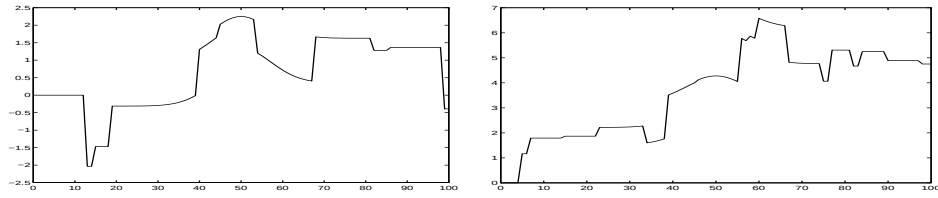
$$p(\mathbf{x}|\mathbf{q}) \propto \exp \left[\frac{-1}{2q_j \sigma_x^2} \sum_j q_j x_j^2 \right]$$

8. Modèle Bernoulli-Gamma :



$$p(x_j|q_j) = \begin{cases} 0 & \text{si } q_j = 0, \\ \text{Gamma}(\alpha, \beta) & \text{si } q_j = 1 \end{cases} \quad \text{et} \quad \begin{cases} P(q_j = 1) = \lambda, \\ P(q_j = 0) = 1 - \lambda \end{cases}$$

9. Modèles pour les signaux continus par morceaux

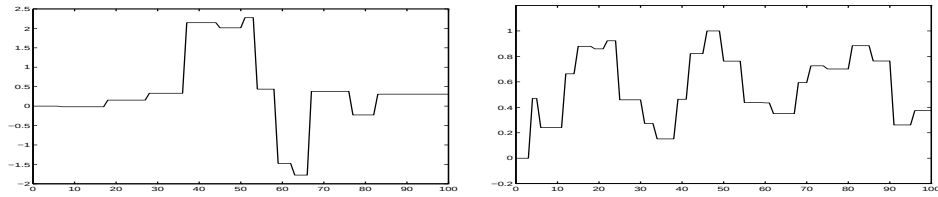


$$p(x_j|x_{j-1}, q_j) = \begin{cases} \mathcal{N}(x_{j-1}, \sigma_1^2) & \text{si } q_j = 0 \\ \mathcal{N}(0, \sigma_2^2) & \text{si } q_j = 1 \end{cases}$$

ou

$$p(\mathbf{x}|\mathbf{q}) \propto \exp \left[-\frac{1}{\sigma_1^2} \sum_j q_j (x_j - x_{j-1}) \right]$$

10. Modèles pour les signaux constants par morceaux



$$p(x_j|x_{j-1}, q_j) = \begin{cases} \delta(x_j - x_{j-1}) & \text{si } q_j = 0 \\ \mathcal{N}(0, \sigma_2^2) & \text{si } q_j = 1 \end{cases}$$

9.7 Estimation des paramètres d'une loi à partir des observations directes

Nous avons vu que le PME peut nous permettre d'attribuer une loi de probabilité $p(x)$ à une quantité X pour traduire une information incomplète sur X si cette information porte sur des moments $E[\phi_k(X)] = d_k$. Malheureusement, en pratique, rarement nous connaissons les valeurs numériques de d_k . C'est pourquoi, en général, l'attribution d'une loi de probabilité se fait en deux étapes :

1. Choix d'une famille de loi qui est équivalent au choix des fonctions ϕ_k ;
2. Détermination des valeurs des paramètres à partir d'un certain nombre d'observations directes ou indirectes du X .

Dans cette section, nous allons nous limiter au cas où nous avons la possibilité d'observer directement X . Supposons maintenant que l'étape 1 est faite et que nous avons choisi une loi $p(x; \theta)$ dépendant des paramètres $\theta = [\theta_1, \dots, \theta_K]$. Considérons alors le cas où nous avons N échantillons $\{x_1, \dots, x_N\}$. Nous cherchons à déterminer θ à partir de ces échantillons. Deux méthodes classiquement utilisées sont : la méthode des moments (MM) et la méthode du maximum de vraisemblance (MV).

9.7.1 Méthode des moments

L'idée de base dans cette méthode est de déterminer $\theta = (\theta_1, \dots, \theta_K)$ en utilisant un système d'équations reliant les moments théoriques $E[X^k]$ et les moments empiriques $\overline{X^k} = \frac{1}{N} \sum_{j=1}^N x_j^k$:

$$E[X^k] = \int x^k p(x; \theta) dx = \frac{1}{N} \sum_{j=1}^N x_j^k, \quad k = 1, \dots, M$$

avec $M \geq K$ et en espérant qu'il existe une solution unique à ce système d'équations (non linéaires en θ).

Pour illustrer cette méthode considérons les deux exemples suivants :

Exemple 1 :

Cas d'une loi gaussienne à deux paramètres $\theta = [m, \sigma^2]$

$$p(x; m, \sigma^2) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp \left[-\frac{1}{2\sigma^2} (x - m)^2 \right].$$

Si on forme le système d'équations :

$$\begin{aligned} E[X] &= m = \frac{1}{N} \sum_{j=1}^N x_j \\ E[X^2] &= \sigma^2 + m^2 = \frac{1}{N} \sum_{j=1}^N x_j^2 \end{aligned}$$

il est alors facile de voir qu'il y a une solution qui est :

$$\begin{aligned} m &= \overline{X} = \frac{1}{N} \sum_{j=1}^N x_j \\ \sigma^2 &= \overline{(X - m)^2} = \frac{1}{N} \sum_{j=1}^N (x_j - m)^2. \end{aligned}$$

Exemple 2 :

Cas d'une loi gamma à deux paramètres $\theta = [\alpha, \beta]$

$$p(x; \alpha, \beta) = \frac{\beta^\alpha}{\Gamma \alpha} x^{\alpha-1} \exp[-\beta x].$$

Si on forme le système d'équations :

$$\begin{aligned} E[X] &= \frac{\alpha}{\beta} = \frac{1}{N} \sum_{j=1}^N x_j, \\ E[X^2] &= \frac{\alpha(1+\alpha)}{\beta^2} = \frac{1}{N} \sum_{j=1}^N x_j^2, \end{aligned}$$

et si on note par $m = \overline{X} = \frac{1}{N} \sum_{j=1}^N x_j$ et par $v = \overline{(X-m)^2} = \frac{1}{N} \sum_{j=1}^N (x_j - m)^2$, il est alors facile de voir qu'il y a une solution qui est :

$$\begin{aligned} \alpha &= \frac{2v - m^2}{v}, \\ \beta &= \frac{m}{v}. \end{aligned}$$

Les inconvenients majeurs de cette méthode sont :

- l'absence de base théorique pour cette méthode ;
- la sensibilité au nombre de données N de ces estimateurs.

9.7.2 Méthode du maximum de vraisemblance

Estimation au sens du MV consiste à définir la fonction de vraisemblance $V(\theta) = p(x_1, \dots, x_n; \theta)$ ou plutôt la fonction log-vraisemblance $L(\theta) = -\ln V(\theta)$ et à calculer l'argument $\hat{\theta}$ qui la maximise :

$$\hat{\theta} = \arg \max_{\theta} \{V(\theta)\} = \arg \min_{\theta} \{L(\theta)\}$$

Le principal avantage de cet estimateur est qu'il est efficace.

Pour illustrer cette méthode et la comparer à la méthode des moments nous allons considérer les deux exemples de la section précédente.

Exemple 1 :

Cas d'une loi gaussienne à deux paramètres $\theta = [m, \sigma^2]$

$$p(x; m, \sigma^2) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp \left[-\frac{1}{2\sigma^2} (x - m)^2 \right]$$

La fonction du log-vraisemblance s'écrit :

$$L(m, \sigma^2) = \frac{N}{2} \ln(2\pi) + \frac{N}{2} \ln \sigma^2 + \frac{1}{2\sigma^2} \sum_{j=1}^N (x_j - m)^2.$$

Pour calculer l'argument qui la minimise, il faut annuler simultanément la dérivée par rapport à m et la dérivée par rapport à σ^2 , ce qui permet d'obtenir un résultat bien connu :

$$\begin{aligned} m &= \frac{1}{N} \sum_{j=1}^N x_j, \\ \sigma^2 &= \frac{1}{N-1} \sum_{j=1}^N (x_j - m)^2. \end{aligned}$$

Notons la seule différence par rapport aux résultats de la méthode des moments apparaît dans l'estimation du σ^2 .

Exemple 2 :

Cas d'une loi gamma à deux paramètres $\theta = [\alpha, \beta]$

$$p(x; \alpha, \beta) = \frac{\beta^\alpha}{\Gamma(\alpha)} x^{\alpha-1} \exp[-\beta x].$$

La fonction du log-vraisemblance s'écrit :

$$L(\alpha, \beta) = N[\alpha \ln \beta - \ln \Gamma(\alpha)] - (\alpha - 1) \sum_{j=1}^N \ln x_j - \beta \sum_{j=1}^N x_j.$$

Pour calculer l'argument qui la minimise, il faut annuler simultanément la dérivée par rapport à α et la dérivée par rapport à β :

$$\begin{aligned} \frac{\partial L}{\partial \alpha} &= N \left[\ln \beta - \frac{\partial \ln \Gamma(\alpha)}{\partial \alpha} \right] - \sum_{j=1}^N \ln x_j = 0 \\ \frac{\partial L}{\partial \beta} &= \frac{\alpha}{\beta} - \sum_{j=1}^N x_j = 0 \end{aligned}$$

Comme on peut le constater, il n'y a pas de solution analytique à ces deux équations, mais on peut la calculer numériquement.

Pour montrer cependant qu'il peut y avoir un lien entre ces deux méthodes, on peut réécrire les relations différemment.

Supposons que la loi $p(x; \theta)$ soit de la forme générale :

$$p(x; \theta) \propto \exp \left[- \sum_{k=1}^K \theta_k \phi_k(x) \right].$$

On a alors

$$L(\theta) = -\ln p(x_1, \dots, x_N; \theta) = cte + \sum_{j=1}^N \sum_{k=1}^K \theta_k \phi_k(x_j).$$

On peut alors montrer facilement que, dans la méthode des moments nous avons à résoudre

$$\int x^k p(x; \theta) dx = \frac{1}{N} \sum_{j=1}^N x_j^k, \quad k = 0, \dots, K,$$

et dans la méthode du MV nous avons à résoudre

$$\int \phi_k(x) p(x; \boldsymbol{\theta}) \, dx = \frac{1}{N} \sum_{j=1}^N \phi_k(x_j), \quad k = 0, \dots, K,$$

où nous avons noté $\phi_0(x) = 1$.

On peut alors remarquer que les deux méthodes sont équivalentes pour la famille des lois exponentielles où $\phi_k(x) = x^k$.

9.8 Estimation des hyperparamètres

Il existe assez peu de méthodes de détermination des hyperparamètres. Dans le cas où ceux-ci se limitent au seul coefficient de régularisation, et où le régulariseur est quadratique, les méthodes de *validation croisée* fournissent des solutions acceptables. Mais il s'agit de méthodes déterministes par essence, qui minimisent un critère de risque dépendant de l'objet inconnu $\mathbf{f}$, ce qui ne peut être effectué qu'asymptotiquement puisque cet objet est évidemment inconnu.

Les méthodes bayésiennes, qui attribuent une distribution de probabilité *a priori* à l'objet, ne présentent pas cette limitation. Les hyperparamètres θ constituent un second niveau de description du problème, indispensable pour “rigidifier” le premier niveau constitué par les paramètres eux-mêmes— c'est-à-dire l'objet $\mathbf{f}$. Dans un problème mal-posé, la valeur des hyperparamètres est importante pour obtenir une solution acceptable, mais ne présente pas d'intérêt en soi. Dans une approche bayésienne, on peut donc distinguer deux niveaux d'inférence. Le premier infère sur $\mathbf{f}$, pour une valeur donnée de θ , au travers de la distribution *a posteriori*. Le second infère sur θ grâce à une relation analogue :

$$p(\theta | \mathbf{g}, \mathbf{H}) = \frac{p(\theta | \mathbf{H}) p(\mathbf{g} | \theta, \mathbf{H})}{p(\mathbf{g} | \mathbf{H})}. \quad (9.5)$$

On retrouve là une caractéristique de l'utilisation de la règle de Bayes : la vraisemblance $p(\mathbf{g} | \theta, \mathbf{H})$ attachée aux données dans le second niveau est le coefficient de normalisation dans le premier.

Si, comme cela est souvent le cas, ce terme est suffisamment “piqué”, l'influence de la distribution *a priori* $p(\theta | \mathbf{H})$ est négligeable, et le second niveau d'inférence peut être résolu par maximisation de cette vraisemblance. Mais il faut pour cela résoudre un problème de marginalisation :

$$p(\mathbf{g} | \theta, \mathbf{H}) = \int p(\mathbf{f}, \mathbf{g} | \theta, \mathbf{H}) d\mathbf{f} = \int p(\mathbf{g} | \mathbf{f}, \theta, \mathbf{H}) p(\mathbf{f} | \theta) d\mathbf{f}. \quad (9.6)$$

Une telle intégrale conduit très rarement à une forme explicite.

Pour contourner cette difficulté, on peut introduire des “variables cachées” $\mathbf{z}$ qui viennent compléter les observations $\mathbf{g}$ de manière à ce que la nouvelle vraisemblance $p(\mathbf{g}, \mathbf{z} | \theta, \mathbf{H})$ soit plus simple à calculer. On est alors conduit à maximiser des espérances conditionnelles par des techniques itératives, déterministes ou stochastiques (algorithmes EM et SEM).

On peut aussi remarquer que la *vraisemblance généralisée*

$$p(\mathbf{g}, \mathbf{f} | \theta, \mathbf{H}) = p(\mathbf{f} | \mathbf{g}, \theta, \mathbf{H}) p(\mathbf{g} | \theta, \mathbf{H}) = p(\mathbf{g} | \mathbf{f}, \theta, \mathbf{H}) p(\mathbf{f} | \theta) \quad (9.7)$$

résume toute l'information propre au premier niveau d'inférence, et vouloir en faire la maximisation conjointe par rapport à $\mathbf{f}$ et θ . Le problème d'intégration soulevé par (9.6) est évidemment évacué. À θ fixé, le maximum de vraisemblance généralisée (MVG) coïncide avec le MAP. Par contre, à $\mathbf{f}$ fixé, la situation est beaucoup moins favorable : la vraisemblance généralisée n'est pas, en général, majorée sur son domaine de définition, et il n'existe même pas toujours un maximum local.

A COMPLETER

Chapitre 10

Modélisation markovienne

Dans ce chapitre nous donnons un aperçu général de la modélisation markovienne des signaux et des images. Les modèles markoviens et particulièrement les champs de Markov et les champs de Gibbs sont devenus des outils indispensables pour la traduction d'une information *a priori* plus élaborée que la douceur ou la positivité en une loi de probabilité. En effet, la modélisation markovienne permet de construire des modèles qui peuvent prendre en compte les discontinuités dans un signal ou une image. Cette outil prend particulièrement son importance en traitement d'image où l'on peut construire les modèles composites (bords, contours, intensité) pour les images et les utiliser lors de la segmentation, restauration ou reconstruction d'image.

10.1 Introduction et notations

10.1.1 Site, voisinage et cliques

Considérons un ensemble de sites $\mathcal{S} = \{s_1, \dots, s_n\}$ et un système de voisinage $\mathcal{V} = V_s, s \in \mathcal{S}$ sur cet ensemble.

- On dit que le site r est voisin du site s si

$$r \in V_s \iff s \in V_r, \quad s \notin V_s$$

$\{\mathcal{S}, \mathcal{V}\}$ est un graphe.

- $c \subset \mathcal{S}$ est une clique de $\{\mathcal{S}, \mathcal{V}\}$ si :

$$|c| = 1 \quad \text{ou si} \quad |c| > 1 \text{ et } \forall s_1, s_2 \in c \longrightarrow s_1 \text{ et } s_2 \text{ sont voisins}$$

Notons par $\mathcal{C} = \{c\}$ l'ensemble de cliques associées au graphe $\{\mathcal{S}, \mathcal{V}\}$.

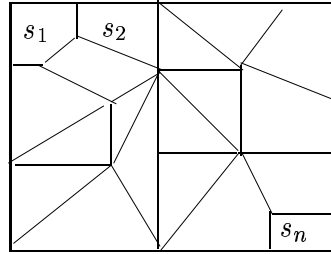
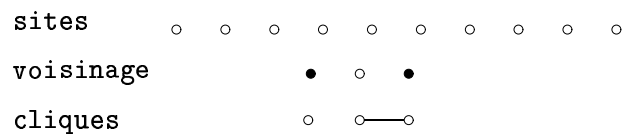


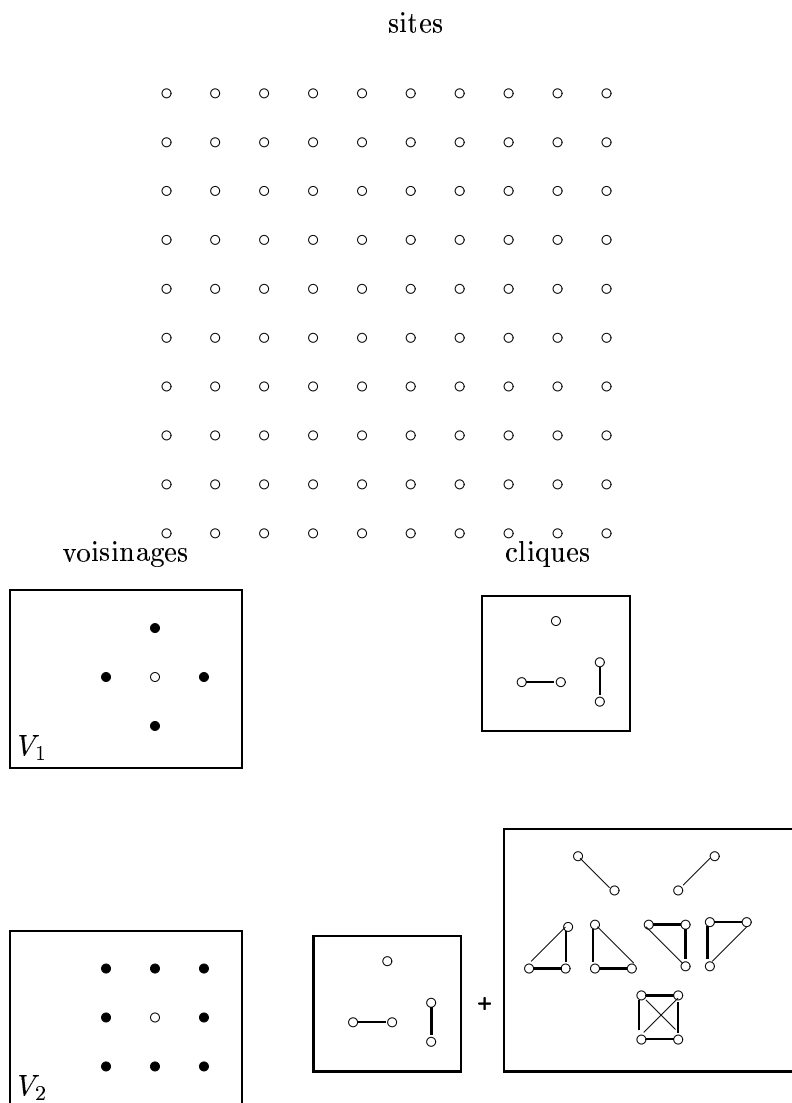
FIG. 10.1 – Sites, voisinage et champ.

Notons aussi que les sites ne sont pas obligés être tous identiques, mais en traitement du signal ces sites peuvent correspondre aux positions des échantillons. En traitement d'image on a l'habitude de parler de *pixels*. Ci-dessous deux exemples

Exemple 1 : $\mathcal{S} = \mathbf{Z}$



Exemple 2 : $\mathcal{S} = \mathbf{Z}^2$



10.1.2 Champ aléatoire

Ensembles des variables aléatoires $\mathcal{X} = \{X_s, s \in \mathcal{S}\}$ est appelé un champ aléatoire sur les sites $\mathcal{S}$. La variable aléatoire X_s prend sa valeur dans l'ensemble $\mathcal{G} = \{g_1, \dots, g_k\}$ qui appelé l'espace des phases (niveau de gris, couleurs, etc.).

Une réalisation du champ aléatoire $\mathcal{X}$ est notée :

$$\{\mathbf{X} = \mathbf{x}\} = \{X_{s_1} = x_{s_1}, \dots, X_{s_N} = x_{s_N}\}$$

et l'ensemble des réalisations (ou configurations) possibles est noté :

$$\Omega = \{\mathbf{X} = (X_{s_1}, \dots, X_{s_N}) : X_{s_n} \in \mathcal{G}, n = 1, \dots, N\}$$

Lorsque l'espace des phases est à dimension finie, (cas discret) le champ est caractérisé par la loi de probabilité conjointe $P(\mathbf{X} = \mathbf{x})$, et lorsque l'espace des phases est à dimension infinie, (cas continu) le champ est caractérisé par sa densité de probabilité conjointe $p(\mathbf{x})$.

On peut aussi définir les lois de probabilités conditionnelles :

$$P(X_s = x_s | X_r = x_r), \text{ cas discret} \quad \text{ou} \quad p(x_s | X_r = x_r), \text{ cas continu}$$

10.1.3 Champ aléatoire markovien

- $\mathcal{X}$ champ aléatoire sur $\mathcal{S}$ à valeurs dans $\mathcal{G}$ est dit markovien relativement au système de voisinage $\mathcal{V}$ si
 - $P(\mathbf{X} = \mathbf{x}) > 0, \quad \forall \mathbf{x} \in \Omega, \quad \text{et}$
 - $\forall X_s \in \mathcal{G}, \forall s \in \mathcal{S}, \quad P(X_s = x_s | X_r = x_r, r \neq s) = P(X_s = x_s | X_r = x_r, r \in V_s)$
- Les fonctions : $\pi_s(x_s | x_{V_s}) = P(X_s = x_s | X_{V_s} = x_{V_s}), s \in \mathcal{S}$ sont les caractéristiques locales du champ markovien $\mathcal{X}$

Deux questions :

Etant donné un système $\{\pi_s(x_s | x_{V_s}), s \in \mathcal{S}\} = \mathcal{P}$,

1. existe-t-il un champ markovien avec $\mathcal{P}$ comme caractéristiques locales ?
2. est-il unique ?

10.2 Champs de Gibbs et équivalence Gibbs-Markov

10.2.1 Energie et Potentiel

- $\{V_a, a \subseteq \mathcal{S}\} : \Omega \longrightarrow \mathbf{R}^+$ telle que $V_a(\mathbf{x})$ ne dépend que de $\mathbf{x}_a$
- $V_c, c \in \mathcal{C}$ est un potentiel de voisinage si $V_c = 0$ lorsque c n'est pas une clique
- $U(\mathbf{x}) = \sum_{a \subseteq \mathcal{S}} V_a(\mathbf{x})$ est l'énergie associée au potentiels V_a

Lorsque U est l'énergie associée au potentiel de voisinage $V_c, c \in \mathcal{C}$ on a

$$U(\mathbf{x}) = \sum_{c \in \mathcal{C}} V_c(\mathbf{x})$$

10.2.2 Mesure de Gibbs

- la mesure de Gibbs associée à l'énergie $U(\mathbf{x})$ et à la température T est la mesure de probabilité sur Ω :

$$\pi_T(\mathbf{x}) = \frac{1}{Z_T} \exp \left[-\frac{U(\mathbf{x})}{T} \right]$$

- la constante

$$Z_T = \sum_{\mathbf{x} \in \Omega} \exp \left[-\frac{U(\mathbf{x})}{T} \right]$$

est la fonction de partition.

- T est la température
- Un champ aléatoire $\mathcal{X}$ défini sur $\mathcal{S}$ est un champ de Gibbs si:

$$P(\mathbf{X} = \mathbf{x}) = \frac{1}{Z} \exp \left[-\frac{U(\mathbf{x})}{T} \right] \quad \text{avec} \quad U(\mathbf{x}) = \sum_{c \in \mathcal{C}} V_c(\mathbf{x})$$

10.2.3 Équivalence Gibbs-Markov

Théorème : Soit $\mathcal{V}$ un système de voisinage sur $\mathcal{S}$. Alors, $\mathcal{X}$ est un champ markovien sur $\mathcal{S}$ à valeurs dans $\mathcal{G}$ relativement au système de voisinage $\mathcal{V}$ ssi $p(\mathbf{x}) = P(\mathbf{X} = \mathbf{x})$ est une distribution de Gibbs relativement à $\mathcal{V}$.

(Hammersley	&	Clifford)
Gibbs	$\iff$	Markov
Potentiels	$\iff$	caracteristiques locales

Gibbs $\longrightarrow$ Markov :

$$P(X_s = x_s | X_{-\mathcal{S}} = x_{-\mathcal{S}}) = \frac{\exp \left[-1/T \sum_{c \in \mathcal{C}} V_c(x_s) \right]}{\sum_{y_s \in \mathcal{G}} \exp \left[-1/T \sum_{c \in \mathcal{C}} V_c(y_s, x_{-\mathcal{S}}) \right]}$$

Markov $\longrightarrow$ Gibbs : plus délicat. Potentiels canoniques

10.3 Exemples de modèles classiques

10.3.1 Auto-modèles

$$U(\mathbf{x}) = \sum_r x_r g_r(x_r) + \sum_r \sum_s \beta_{rs} x_r x_s, \quad r, s \in \mathcal{S}$$

$g_r(\cdot)$ fonctions arbitraires

β_{rs} coefficients d'interaction entre les sites r et s

$\beta_{rs} \neq 0$ ssi les sites r et s sont voisins: $|r - s|^2 \leq d$

10.3.2 Auto-logistique

$$\mathcal{G} = \{0, 1\} \quad \text{ou} \quad \{-1, +1\}$$

$$U(\mathbf{x}) = \sum_r \alpha_r x_r + \sum_r \sum_s \beta_{rs} x_r x_s \quad r, s \in \mathcal{C}$$

Si $\alpha = 0$, et $\beta_{rs} = \begin{cases} \beta & \text{si } |r - s|^2 \leq d \\ 0 & \text{sinon} \end{cases}$, U mesure la longueur d'un contour.

10.3.3 Modèle d'Ising

Auto-logistique avec voisinage d'ordre 1

Deux exemples spécifiques:

– Chaîne de Markov binaire: $\mathcal{G} = \{0, 1\}$, $\mathcal{S} = \mathbf{Z}$,

$$\text{Matrice de transitions} \begin{pmatrix} 0-0 & 0-1 \\ 1-0 & 1-1 \end{pmatrix} \quad P = \begin{pmatrix} s & 1-s \\ 1-t & t \end{pmatrix}$$

$$\pi(\mathbf{x}) = \frac{1}{Z} \exp \left[\bar{V}_0 \sum_{i=0, n} x_i + V_0 \sum_{0 < i < n} x_i + V_1 \sum_{|i-j|=1} x_i x_j \right]$$

$$\bar{V}_0 = \log \frac{1-s}{s}, \quad V_0 = \log \frac{(1-s)(1-t)}{s^2}, \quad V_1 = \log \frac{st}{(1-s)(1-t)}$$

– Images binaires et voisinage d'ordre 1:

$$\mathcal{G} = \{0, 1\}, \quad \mathcal{S} = \mathbf{Z}^2, \quad V_1 = \left\{ \begin{array}{ccc} & \bullet & \\ \bullet & \circ & \bullet \\ & \bullet & \end{array} \right\}$$

$$U(\mathbf{x}) = \alpha \sum x_{ij} + \beta \sum x_{i,j} v_{i,j}$$

avec

$$v_{i,j} = x_{i,j-1} + x_{i-1,j} + x_{i,j+1} + x_{i+1,j}$$

10.3.4 Multi-Level Logistic Model (MLL)**ou Ising généralisé ou encore Champ de labels** $\mathcal{G} = g_1, \dots, g_M$ niveaux des gris dans chaque région

$$U(\mathbf{x}) = \sum_k V_k(x) + \sum_c V_c(\mathbf{x})$$

– pour les cliques simples

$$V_c(x) = \begin{cases} \alpha_k & \text{si } x_{ij} = g_k \\ 0 & \text{sinon} \end{cases}$$

– pour les cliques multiples

$$V_c(\mathbf{x}) = \begin{cases} +\beta & \text{si } x_{ij} \text{ dans } C \text{ ont les mêmes valeurs} \\ -\beta & \text{sinon} \end{cases}$$

10.4 Tentative de classification

- **Champ de markov au sens strict (SSM) :**

(Définition via les probabilités conditionnelles)

Soient $\mathcal{B} \subseteq \mathcal{A} \subseteq \mathcal{S}$ $P(X_{ij} = x_{ij} | X_{\mathcal{B}} = x_{\mathcal{B}}) = P(X_{ij} = x_{ij} | X_{\mathcal{A}} = x_{\mathcal{A}})$

- **Champ de markov au sens large (WSM)**

(Définition via la régression linéaire)

(Estimation linéaire en moyenne quadratique ELMQ)

$X = \{X_{ij}\}, \quad (i, j) \in \mathbf{Z}^2, \quad \mathcal{B} \subseteq \mathcal{A} \subseteq \mathcal{S}$

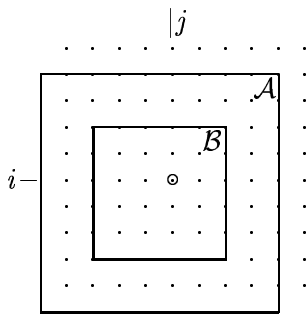
$$ELMQ \{X_{ij} | X_{\mathcal{B}} = x_{\mathcal{B}}\} = ELMQ \{X_{ij} | X_{\mathcal{A}} = x_{\mathcal{A}}\}$$

- Propriétés au sens large = Propriétés au deuxième ordre

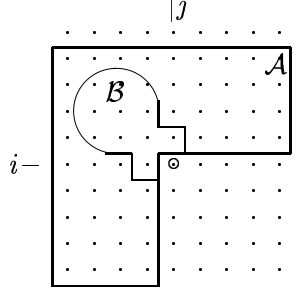
- Représentation équivalente par un modèle AR: $X_s = \sum_r a_r X_r + Z_s$

Champ de markov au sens strict (SSM)

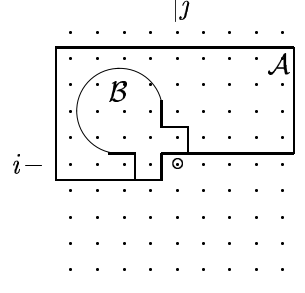
Soient $\mathcal{B} \subseteq \mathcal{A} \subseteq \mathcal{S}$ $P(X_{ij} = x_{ij} | X_{\mathcal{B}} = x_{\mathcal{B}}) = P(X_{ij} = x_{ij} | X_{\mathcal{A}} = x_{\mathcal{A}})$



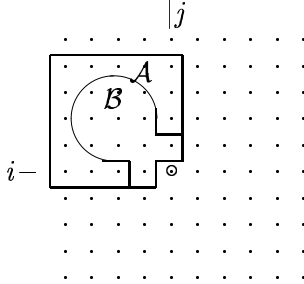
Bilatéral



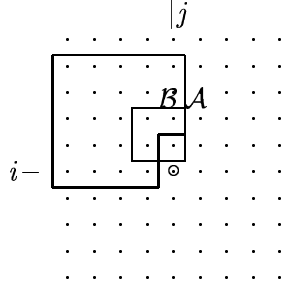
MMRF



Unilateral



Causal



Pickard

Pickard $\subseteq$ Causal
Causal $\subseteq$ Unilatéral
Unilatéral $\subseteq$ MMRF
MMRF $\subseteq$ Bilatéral

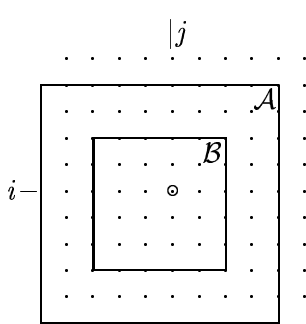
Champ de markov au sens large (WSM)

(Définition via la régression linéaire)

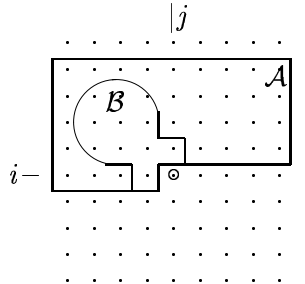
(Estimation linéaire en moyenne quadratique ELMQ)

$X = \{X_{ij}\}, (i,j) \in \mathbb{Z}^2, \mathcal{B} \subseteq \mathcal{A} \subseteq \mathcal{S}$

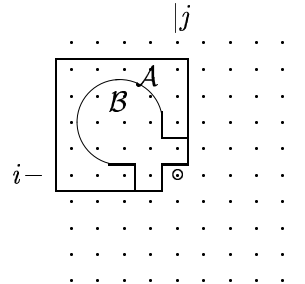
$$ELMQ\{X_{ij} | X_{\mathcal{B}} = x_{\mathcal{B}}\} = ELMQ\{X_{ij} | X_{\mathcal{A}} = x_{\mathcal{A}}\}$$



Bilatéral



Unilateral



Causal

Champ de markov gaussien : sens large = sens stricte

Un exemple classique utilisé en traitement d'image :

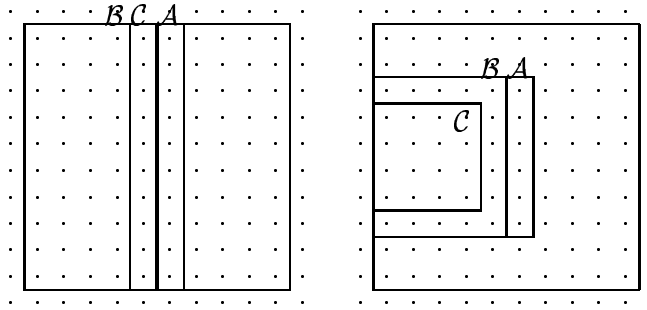
$$\mathcal{S} = \mathbf{Z}_{mn}^2 = \{(i,j) : 1 \leq i \leq m, 1 \leq j \leq n\}, \mathcal{X} = \{X_{i,j}, (i,j) \in \mathcal{S}\}$$

$$\nu_1 = \left\{ \begin{array}{ccc} & \bullet & \\ \bullet & \circ & \bullet \\ & \bullet & \end{array} \right\} \quad \nu_2 = \left\{ \begin{array}{ccc} \bullet & \bullet & \bullet \\ \bullet & \circ & \bullet \\ \bullet & \bullet & \bullet \end{array} \right\}$$

$$P(X_{ij} = x_{ij} | X_{kl} = x_{kl}, (k,l) \in \mathcal{S} - \{i,j\}) = P(X_{ij} = x_{ij} | X_{kl} = x_{kl}, (k,l) \in V_{ij})$$

Quelques propriétés :

$$P(x_a, a \in \mathcal{A} | x_b, b \in \mathcal{B}) = P(x_a, a \in \mathcal{A} | x_b, b \in \mathcal{C})$$



10.5 Modèles couplés ou hiérarchiques

			$\begin{array}{cccccccc} \circ & & \circ & & \circ & & \circ & & \circ \\ \hline \circ & & \circ & & \circ & & \circ & & \circ \\ \hline \circ & & \circ & & \circ & & \circ & & \circ \\ \hline \circ & & \circ & & \circ & & \circ & & \circ \\ \hline \circ & & \circ & & \circ & & \circ & & \circ \\ \hline \circ & & \circ & & \circ & & \circ & & \circ \\ \hline \circ & & \circ & & \circ & & \circ & & \circ \end{array}$
–	$\mathbf{Z}_m$	réseau des entiers	
–	$\mathcal{S}_1 = \mathbf{Z}_m$	sites pixels	$\circ$
–	$\mathcal{S}_2 = \mathbf{Z}_m$	sites contours	–
–	$\mathcal{S} = \mathcal{S}_1 \cup \mathcal{S}_2$	sites image	
	• $\mathcal{F} = \{F_{i,j}, (i,j) \in \mathcal{S}_1\}$ champ intensité • $\mathcal{L} = \{L_{\alpha,\beta}, (\alpha,\beta) \in \mathcal{S}_2\}$ champ des lignes • $\mathcal{X} = (\mathcal{F}, \mathcal{L})$ image		
	$\begin{array}{cccc} \bullet & & \bullet & \bullet & \bullet \\ \bullet & \circ & \bullet & \bullet & \circ & \bullet \\ \bullet & & \bullet & \bullet & \bullet \end{array}$	Voisinage pixels	$\begin{array}{ccccc} \bullet & \square & \bullet & \bullet & \bullet \\ \square & \square & \square & \square & \square \\ \bullet & \bullet & \bullet & \bullet & \bullet \end{array}$
			voisinage contours

Modèle hiérarchique :

$$P(\mathbf{f}, \mathbf{l}) = P(\mathbf{f}|\mathbf{l})P(\mathbf{l})$$

$$P(\mathbf{l}) = \frac{1}{Z_L} \exp[-U_L(\mathbf{l})]$$

$$P(\mathbf{f}|\mathbf{l}) = \frac{1}{Z_{F|L}} \exp[-U_{F|L}(\mathbf{f}|\mathbf{l})]$$

$$\longrightarrow P(\mathbf{f}, \mathbf{l}) = \frac{1}{Z_{F,L}} \exp[-U_{F,L}(\mathbf{f}, \mathbf{l})],$$

avec

$$U_{F,L}(\mathbf{f}, \mathbf{l}) = U_{F|L}(\mathbf{f}|\mathbf{l}) + U_L(\mathbf{l})$$

10.6 Application en segmentation d'image

- $\mathcal{S} = \{(i,j) : 1 \leq i \leq m, 1 \leq j \leq n\}$ sites pixels ◦
- $\mathcal{X} = \{X_{i,j}, (i,j) \in \mathcal{S}\}$ image
- $\mathcal{Y} = \{Y_{i,j}, (i,j) \in \mathcal{S}\}$ image dégradée
- $\mathcal{N} = \{N_{i,j}, (i,j) \in \mathcal{S}\}$ bruit

Modèle de dégradation :

$$\mathcal{Y} = \mathcal{X} + \mathcal{N}$$

$X_{i,j} = g_m$, si $X_{i,j} = m$ avec $\mathcal{G} = \{g_1, \dots, g_M\}$ niveaux des gris dans chaque région

Hypothèses statistiques :

- $\mathcal{X}$ un champ de Markov à valeur dans $\mathcal{G} = \{g_1, \dots, g_M\}$
- $\mathcal{N}$ un champ gaussien blanc
- $N_{i,j}$ v.a. indépendantes et de lois $\mathcal{N}(0, \sigma^2)$
- $\mathcal{N}$ et $\mathcal{X}$ indépendants

Modèle a posteriori et choix de l'estimateur

- $\mathcal{X}$ est un champ de Gibbs d'énergie $U(\mathbf{x})$:

$$P(\mathbf{X} = \mathbf{x}) = \frac{1}{Z} \exp [-U(\mathbf{x})]$$

$$\begin{aligned} U(\mathbf{x}) = & \sum_j \alpha_j x_j \\ & + \beta_1 \sum_{(i,j)0}^j x_i x_j + \beta_2 \sum_{(i,j)90} x_i x_j + \beta_3 \sum_{(i,j)-45} x_i x_j + \beta_4 \sum_{(i,j)45} x_i x_j \\ & + \gamma_1 \sum_{(i,j,k)} x_i x_j x_k + \gamma_2 \sum_{(i,j)90} x_i x_j x_k + \dots \\ & + \xi_1 \sum_{(i,j,k,l)} x_i x_j x_k x_l \end{aligned}$$

avec M paramètres $\alpha_m, m = 1 = \dots, M$, 4 paramètres $\beta_1, \dots, \beta_4$, 4 paramètres $\gamma_1, \dots, \gamma_4$ et un paramètre ξ_1 .

- $\mathcal{N}$ est un champ gaussien, blanc et indépendant de $\mathcal{X}$,

$$P(\mathbf{Y} = \mathbf{y} | \mathbf{X} = \mathbf{x}) = K \exp \left[-\frac{1}{2\sigma_b^2} \|\mathbf{y} - \mathbf{x}\|^2 \right]$$

- Règle de Bayes :

$$P(\mathbf{X} = \mathbf{x} | \mathbf{Y} = \mathbf{y}) = \frac{P(\mathbf{Y} = \mathbf{y} | \mathbf{X} = \mathbf{x}) P(\mathbf{X} = \mathbf{x})}{P(\mathbf{Y} = \mathbf{y})}$$

- Énergie a posteriori : $U^p(\mathbf{x}) = \frac{1}{2\sigma_b^2} \|\mathbf{y} - \mathbf{x}\|^2 + U(\mathbf{x})$

Deux types de problèmes :**1. Classification ou segmentation supervisée :**

On connaît $\theta = \{\{g_1, \dots, g_M\}, \{\beta_1, \dots, \beta_4\}, \{\gamma_1, \dots, \gamma_4\}, \xi_1\}$ et $\mathbf{y} \longrightarrow \hat{\mathbf{x}}$.

2. Segmentation non supervisée : $\mathbf{y} \longrightarrow \hat{\mathbf{x}}, \hat{\theta}$

– Apprentissage supervisée : $\mathbf{x}, \mathbf{y} \longrightarrow \hat{\theta}$

– Apprentissage non supervisée : $\mathbf{y} \longrightarrow \hat{\theta}$

Quelques estimateurs possibles pour le problème 1

- MAP (Maximum A Posteriori) : $\hat{\mathbf{x}} = \arg \max_{\mathbf{x}} \{P(\mathbf{x}|\mathbf{y})\}$
- MPM (Maximum of the Posterior Marginal) :

$$\hat{\mathbf{x}} = \{\hat{x}_{s \in S}\} \quad \text{avec} \quad \hat{x}_s = \arg \max_{x_s} \{p_s(X_s = x_s|\mathbf{y})\}$$

- PM (Posterior Mean) Moyenne *a posteriori* :

$$\hat{\mathbf{x}} = \{\hat{x}_{s \in S}\} \quad \text{avec} \quad \hat{x}_s = \sum_{x_i \in \mathcal{G}} x_s p_s(X_s = x_s|\mathbf{y})$$

Deux algorithmes stochastiques :

- **Le recuit simulé ou relaxation stochastique**
permet d'estimer le MAP
- **L'échantillonneur de Gibbs**
permet de simuler des réalisations de $p(x_s|\mathbf{y}, x_{j \neq s})$, et de $p(\mathbf{x}|\mathbf{y})$ et donc d'estimer le MPM ou la moyenne *a posteriori*.

Quelques estimateurs possibles pour les problèmes 2

- Apprentissage supervisée : $\mathbf{x}, \mathbf{y} \longrightarrow \hat{\theta}$ MAP, PM, MPM
- Apprentissage non supervisée : $\mathbf{y} \longrightarrow \hat{\theta}$ MV (EM)
- Segmentation non supervisée : $\mathbf{y} \longrightarrow \hat{\mathbf{x}}, \hat{\theta}$ MVG

Lorsque $\mathbf{x}$ est modélisé par un champ unilatéral ou par une chaîne de Markov :

- $\mathbf{y}, \theta \longrightarrow \hat{\mathbf{x}}_{\text{MAP}}$ Viterbi (Programmation dynamique)
- $\mathbf{y} \longrightarrow \hat{\theta}_{\text{MV}}$ EM, (Forward-Backward)
- $\mathbf{y} \longrightarrow \hat{\mathbf{x}}_{\text{MAP}}, \hat{\theta}_{\text{MV}}$: Algorithme itératif, par ex. MVG:

$$\hat{\theta}^{(0)}, \mathbf{y} \longrightarrow \hat{\mathbf{x}}_{\text{MAP}}^{(1)}, \mathbf{y} \longrightarrow \hat{\theta}_{\text{MV}}^{(1)}, \mathbf{y} \longrightarrow \hat{\mathbf{x}}_{\text{MAP}}^{(1)}, \mathbf{y} \longrightarrow \dots$$

10.7 Application en restauration d'image

Modèle de Geman-Geman : modèle couplé pixels-contours

- $\mathcal{F} = \{F_{i,j}, (i,j) \in \mathbf{Z}_m\}$ champ intensité $U(\mathbf{f}|\mathbf{l})$
- $\mathcal{L} = \{L_{\alpha,\beta}, (\alpha,\beta) \in \mathcal{D}_m\}$ champ des lignes $U(\mathbf{l})$
- $\mathcal{X} = (\mathcal{F}, \mathcal{L})$ image $U(\mathbf{f}, \mathbf{l}) = U(\mathbf{f}|\mathbf{l}) + U(\mathbf{l})$
- $\mathcal{G} = \{G_{i,j}, (i,j) \in \mathbf{Z}_m\}$ image dégradée

$$G_{i,j} = \sum_{(k,l)} H_{i-k,j-l} F_{k,l} + N_{i,j} \quad \text{ou} \quad \mathbf{g} = \mathbf{H}\mathbf{f} + \mathbf{n}$$

Remarque : le champ des ligne n'est pas dégradée

- $\mathcal{N} = \{N_{i,j}, (i,j) \in \mathbf{Z}_m\}$ champ bruit d'intensité (gaussien, blanc)
- $\mathbf{H} = \{H_{i,j}, (i,j) \in \mathbf{Z}_m\}$ matrice de flou

Estimateur du maximum a posteriori (MAP)

$$\begin{aligned} \mathcal{X} &= (\mathcal{F}, \mathcal{L}), & \mathbf{x} &= (\mathbf{f}, \mathbf{l}) \\ \mathcal{Y} &= (\mathcal{G}, \mathcal{L}), & \mathbf{y} &= (\mathbf{g}, \mathbf{l}) \end{aligned}$$

- $\mathcal{X}$ est un champ de Markov d'énergie $U(\mathbf{f}, \mathbf{l})$:

Voisinage pixels et voisinage contours



- $\mathcal{N}$ est gaussien, blanc et indépendant de $\mathcal{F}$,
- $\mathbf{H}$ est invariant par translation :

$$P(\mathbf{G} = \mathbf{g} | \mathbf{X} = \mathbf{x}) = K \exp \left[-\frac{1}{2\sigma_b^2} \|\mathbf{g} - \mathbf{H}\mathbf{f}\|^2 \right]$$

- Règle de Bayes : $P(\mathbf{X} = \mathbf{x} | \mathbf{G} = \mathbf{g}) = A \exp[-U^p(\mathbf{f}, \mathbf{l})]$ avec

$$U^p(\mathbf{f}, \mathbf{l}) = \frac{1}{2\sigma_b^2} \|\mathbf{g} - \mathbf{H}\mathbf{f}\|^2 + U(\mathbf{f}, \mathbf{l})$$

- Loi *a posteriori* est aussi une mesure de Gibbs. Il reste à déterminer le système de voisinage de cette loi.
- Le champ des lignes $\mathcal{L}$ n'est pas affecté par la dégradation
- Le voisinage du champ d'intensité $\mathcal{X}$ est élargi par la réponse impulsionnelle de la dégradation.

10.8 Outils de simulation et d'optimisation

10.8.1 Echantillonneur de Gibbs

- Soit $P(\mathbf{X} = \mathbf{x})$ une distribution de probabilité.
- Soit $(s_k), k \in \mathcal{N}$ une suite de visite des sites telle que, pour tout $s \in \mathcal{S}$, s est visité un nombre infini de fois.

$$X_s^{(k+1)} = \begin{cases} X_s^{(k)} & \text{si } s \neq s_k \\ \xi & \text{si } s = s_k \end{cases}$$

ξ est une v.a. tirée suivant la loi conditionnelle

$$p_{s_k}(\xi = \lambda) = p_{s_k}(\xi = \lambda | x_j = x_j(k), j \neq s_k)$$

Si l'ensemble des états est fini ou si X est gaussien on a :

Théorème :

Quelle que soit la configuration initiale x_0

$$\lim_{k \rightarrow \infty} P(\mathbf{X}(k) = \mathbf{x} | \mathbf{X}(0) = \mathbf{x}(0)) = P(\mathbf{X} = \mathbf{x})$$

Remarques :

- Si X est un champ de Markov, la loi conditionnelle s'exprime simplement en fonction de la configuration sur un voisinage de s_k , ce qui rend cette simulation raisonnable en temps de calcul.
- Pour la segmentation $X|Y$ a la même taille de voisinage que X , mais
- Pour la restauration la taille de voisinage *a posteriori* dépend à la fois de la taille de voisinage a priori et du support de l'opérateur de flou.

10.8.2 Recuit simulé

$$P(\mathbf{X} = \mathbf{x} | \mathbf{Y} = \mathbf{y}) = \frac{1}{Z} \exp[-U(\mathbf{x})] = \pi(\mathbf{x})$$

- **Physique statistique :** niveau d'intensité d'un pixel=état d'un particule élémentaire
- distribution de Gibbs: $\pi_T(\mathbf{x}) = \frac{1}{Z_T} \exp\left[-\frac{1}{T}U(\mathbf{x})\right]$
- chaque réalisation d'image représente l'état de l'ensemble des particule auquel on associe une énergie $U(\mathbf{x})$
 - $T \rightarrow \infty$ $\pi_T(\mathbf{x}) \rightarrow$ loi uniforme
 - $T = 1$ $\pi_T(\mathbf{x}) \rightarrow \pi(\mathbf{x})$
 - $T \rightarrow 0$ $\pi_T(\mathbf{x}) \rightarrow$ loi concentrée sur l'estimée MAP

Algorithme de recuit simulé :

(avec l'échantillonneur de Gibbs)

$$X_s(k+1) = \begin{cases} X_s^{(k+1)} & \text{si } s \neq s_k \\ \xi & \text{si } s = s_k \end{cases}$$

ξ est une v.a. tirée suivant la loi conditionnelle

$$p_{s_k, T_k}(X = \xi) = \frac{1}{Z_k} \exp\left[-\frac{1}{T_k}U(X_{s_k} = x | x_j, j \neq s_k)\right]$$

Théorème : Si $\lim_{k \rightarrow \infty} T_k \ln k \geq K$ où K est une constante assez grande, alors quelle que soit la configuration initiale $X(0) = x_0$, $P(X(k) = x | X(0) = x_0)$ converge en loi vers une distribution uniforme sur les minima globaux de U

10.8.3 Algorithme d'ICM (Iterated Conditional Modes)

(BESAG)

Il s'agit en effet d'une adaptation déterministe de l'algorithme précédent.

$$X_s^{(k+1)} = \begin{cases} X_s^{(k+1)} & \text{si } s \neq s_k \\ \arg \max_X \{p_i(X_{s_k} = x | x_j, j \neq s_k)\} & \text{si } s = s_k \end{cases}$$

Cet algorithme est rapide mais le résultat dépend du point initial et fournit un minimum local de U .

10.9 Algorithmes de relaxations déterministes

Lors de calcul de la solution au sens du MAP, deux situations peuvent se présenter :

1. La loi $a \text{ posteriori}$ est unimodale.

Alors il n'y a qu'une seule solution et, en général, l'optimisation du critère MAP (l'énergie $a \text{ posteriori}$) ne pose pas de grandes difficultés et on peut utiliser les algorithmes de descentes du type gradient ou autre pour effectuer cette optimisation.

2. La loi $a \text{ posteriori}$ est multimodale.

Dans ce cas le problème devient plus difficile et il faut aller à la recherche de techniques d'optimisation globale. L'algorithme du recuit simulé (RS) peut alors être utilisé, mais son coût de calcul devient très vite prohibitif, surtout lorsque la taille de voisinage de la loi $a \text{ posteriori}$ devient plus grande. En effet, la taille de voisinage de la loi $a \text{ posteriori}$ est directement liée à la taille de voisinage de la loi $a \text{ priori}$, mais aussi et surtout aux caractéristiques de la l'opérateur liant la grandeur observable à la grandeur inconnue. Par exemple, lorsqu'il s'agit de la ségmentation ou la restauration d'image avec un noyau de taille très réduite, alors on peut utiliser la RS. Mais dans d'autres applicatiuons du genre la reconstruction d'image en tomographie X ou par ondes diffractées où le support de l'opérateur devient important, on ne peut plus envisager l'utilisation de tel algorithme. On peut alors faire recours aux techniques de relaxations déterministes qui cherchent à faire un compromis entre l'optimalité au sens de la recherche de l'optimum global et le coût de calcul. Parmi ces techniques on peut citer la technique de non convexité graduelle (Graduated Non Convexity GNC) que nous allons la décrire ici dans un cas particulier.

10.9.1 Principe de base du GNC

Soit $J(\mathbf{f})$ un critère multimodal dont nous cherchons à trouver l'argument $\mathbf{f}$ correspondant à son minimum global. La principale idée du GNC est de construire une suite de critères $J_{c_k}(\mathbf{f})$ telles que :

- $J_{c_0}(\mathbf{f})$ soit convexe et
- $\forall \mathbf{f} \lim_{c \rightarrow c_\infty} J_c(\mathbf{f}) = J(\mathbf{f})$.

Ensuite, on minimise $J_{c_0}(\mathbf{f})$ pour obtenir $\mathbf{f}_0$, et finalement, pour une suite de $\{c_1, \dots, c_\infty\}$, on minimise localement $J_{c_k}(\mathbf{f})$ au voisinage de la solution obtenue à l'itération précédente.

On epère ainsi que la suite des solutions $\hat{\mathbf{f}}_k$ ainsi obtenues converge vers le minimum global du critère. Notons que, la convexité du J_{c_0} assure l'unicité de la solution, mais, il n'y a aucune garantie que l'algorithme converge vers le minimum global. Parmi les points importants concernant cet algorithme, il faut noter que son coût, contrairement aux algorithmes à base du recuits simulé, ne dépends pas de la taille de voisinage de la loi $a \text{ posteriori}$ dans le cas d'un estimateur au sens du MAP.

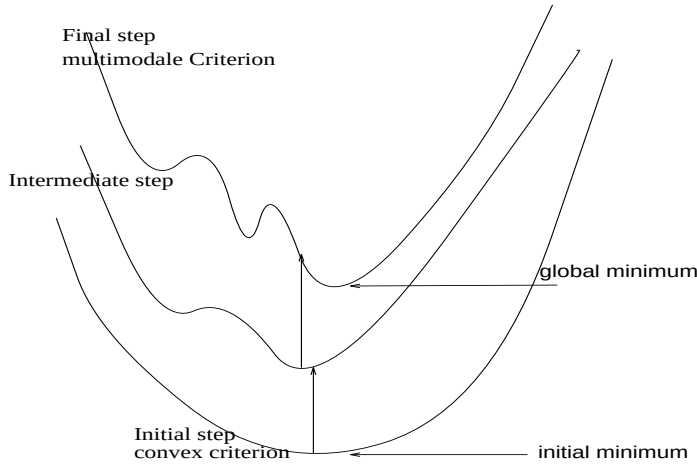


FIG. 10.2 – Principe de base de la technique du GNC.

Pour illustrer la mise en œuvre de cette technique nous allons considérer deux cas :

- le cas du calcul de la solution au sens du MAP d'un problème linéaire avec une loi *a priori* markovienne dont le terme d'énergie est nonconvexe ;
- le cas d'un problème non linéaire où la multimodalité du critère MAP vient plutôt de la non linéarité du problème que du choix de la loi *a priori*.

Cas des problèmes linéaires :

Considérons le cas d'un problème inverse linéaire $\mathbf{g} = \mathbf{H}\mathbf{f} + \mathbf{b}$ en une dimension et supposons que le signal d'entrée est un signal continu par morceau et notons par $\mathbf{f}$ le vecteur contenant les valeurs du signal et par $\mathbf{l}$ le vecteur contenant les positions des discontinuités ($l_j = 0$, pas de discontinuité, $l_j = 1$, il y a une discontinuité). L'objectif de la reconstruction est alors d'estimer à la fois $\mathbf{f}$ et $\mathbf{l}$.

- Modèles avec un processus de ligne binaire

$$p(\mathbf{f}, \mathbf{l} | \mathbf{g}) = p(\mathbf{g} | \mathbf{f}) p(\mathbf{f} | \mathbf{l}) p(\mathbf{l}) \longrightarrow$$

$$U^p(\mathbf{f}, \mathbf{l}) = \|\mathbf{g} - \mathbf{H}\mathbf{f}\|^2 + U(\mathbf{f}, \mathbf{l}) = \|\mathbf{g} - \mathbf{H}\mathbf{f}\|^2 + U(\mathbf{f} | \mathbf{l}) + U(\mathbf{l})$$

- Si le processus de ligne est non-interactif, ce qui signifie $p(\mathbf{l}) = \prod_j p(l_j)$ ou encore $U(\mathbf{l}) = \sum_j \phi_1(l_j)$, alors on peut montrer que la solution

$$(\hat{\mathbf{l}}, \hat{\mathbf{f}}) = \arg \min_{(\mathbf{f}, \mathbf{l})} \{U^p(\mathbf{f}, \mathbf{l} | \mathbf{g})\}$$

peut être calculés en deux étapes :

- Calculer $\hat{\mathbf{f}}$ en minimisant un critère ne dépendant que de $\mathbf{f}$:

$$U^p(\mathbf{f} | \mathbf{g}) = \|\mathbf{g} - \mathbf{H}\mathbf{f}\|^2 + \lambda \sum_{s \in \mathcal{S}} \phi([D\mathbf{f}]_s)$$

avec $D\mathbf{f}$ un opérateur de différences finies et ϕ une fonction non-convexe et ;

- Dédire $\hat{\mathbf{l}}$ par une opération directe à partir de $\hat{\mathbf{f}}$.

- **Exemple :**

$$U(\mathbf{f} | \mathbf{l}) = \lambda^2 \sum_j (f_j - f_{j-1})^2 (1 - l_j), \quad U(\mathbf{l}) = \mu \sum_j l_j,$$

Si

$$(\hat{\mathbf{l}}, \hat{\mathbf{f}}) = \arg \min_{(\mathbf{f}, \mathbf{l})} \{U^p(\mathbf{f}, \mathbf{l} | \mathbf{g})\}$$

alors

$$\hat{\mathbf{f}} = \arg \min_{(\mathbf{f})} \{U^p(\mathbf{f} | \mathbf{g})\} \quad \text{avec} \quad U^p(\mathbf{f} | \mathbf{g}) = \|\mathbf{g} - \mathbf{H}\mathbf{f}\|^2 + \sum_j \phi(x_j - x_{j-1})$$

avec

$$\phi(t) = \begin{cases} (\lambda t)^2 & \text{si } |t| \leq T \\ \alpha = (\lambda T)^2 & \text{si } |t| > T \end{cases}$$

et

$$\hat{l}_j = \begin{cases} 1 & \text{si } |x_j - x_{j-1}| \leq T \\ 0 & \text{si } |x_j - x_{j-1}| > T \end{cases}$$

- $U^p(\mathbf{f} | \mathbf{g})$ est alors multimodale et la recherche de la solution MAP ne peut se faire par une méthode d'optimisation locale.
- Recuit simulé trop coûteux lorsque l'opérateur $\mathbf{H}$ a un support large. En effet le coût de calcul de cet algorithme croît d'une manière exponentielle avec la taille de son support.
- GNC est une technique d'optimisation non-locale dont le coût de calcul n'est pas dépendant du support de l'opérateur $\mathbf{H}$.

Non convexité vient du choix du terme $U^p(\mathbf{f} | \mathbf{g})$ qui est une conséquence du choix de la loi *a priori*. C'est pourquoi pour mettre en œuvre l'algorithme du GNC dans ce cas, on peut envisager de définir une suite de critères $J_c(\mathbf{f})$:

$$J_c(\mathbf{f}) = \|\mathbf{g} - \mathbf{H}\mathbf{f}\|^2 + \sum_j \phi_c(x_j - x_{j-1})$$

avec

$$\phi_c(t) = \begin{cases} (\lambda t)^2 & \text{si } |t| < q_c \\ \alpha - \frac{1}{2}c(|t| - r_c)^2 & \text{si } q_c > |t| > r_c \\ \alpha = (\lambda T)^2 & \text{si } |t| > r_c \end{cases}$$

où

$$\begin{cases} q_c = T(1 + 2\lambda^2/c)^{-1/2} \\ r_c = T(1 + 2\lambda^2/c)^{1/2} \end{cases}$$

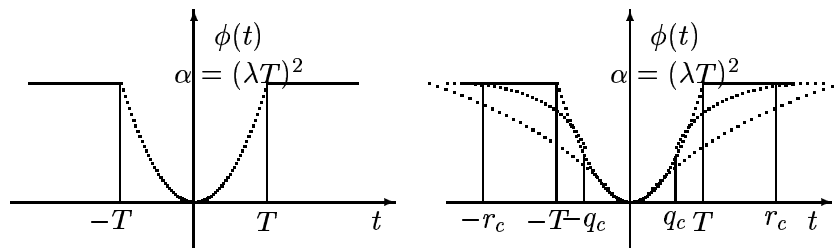


FIG. 10.3 – Fonction quadratique tronquée et sa relaxation

La principale difficulté est ensuite de trouver un c_0 telle que $J_{c_k}(\mathbf{f})$ soit convexe et ensuite choisir la suite $\{c_0, c_1, \dots\}$. La solution est alors calculée par

$$\hat{\mathbf{f}}_{c_k} = \arg \min_{\mathbf{f} \in V(\hat{\mathbf{f}}_{c_{k-1}})} \{J_{c_k}(\mathbf{f})\}$$

- [Blake et Zisserman]:
 - Segmentation d'image $\mathbf{A} = \mathbf{I} \longrightarrow$ Problème bien posé
 - $J_{c_k}(\mathbf{f}) = \|\mathbf{g} - \mathbf{f}\|^2 + \Phi_{c_k}(\mathbf{f}), \quad \exists c_0$ tel que J_{c_0} soit convexe
- [Nikolova, Djafari, Idier]:
 - Extension pour les problème mal posés $\mathbf{A}$
 - $\nexists c_0$ tel que J_{c_0} soit convexe $\longrightarrow$

$$J_{c_k}(\mathbf{f}) = \|\mathbf{g} - \mathbf{A}\mathbf{f}\|^2 + \Phi_{c_k}(\mathbf{f}) + \Psi_{a_k}(\mathbf{f})$$

Double relaxations: $a_k \mapsto 0$ et $c_0 \mapsto c_\infty$

- Extension pour d'autres modèles que celui de la chaîne ou membrane faible

[GG84, Gem90, GR92, Bes74, Bes86, Bes89, BG93, BZ87, NMD96a]

A COMPLETER

Cas d'un problème inverse non linéaire : Dans le cas d'un problème inverse non linéaire $\mathbf{g} = \mathbf{H}(\mathbf{f}) + \mathbf{b}$ si on choisit de définir la solution comme le minimum d'un critère du type :

$$\hat{\mathbf{f}} = \arg \min_{\mathbf{f}} \left\{ J(\mathbf{f}) = \|\mathbf{g} - \mathbf{H}(\mathbf{f})\|^2 + \lambda \Omega(\mathbf{f}) \right\},$$

même si on choisit $\Omega(\mathbf{f})$ convexe, $J(\mathbf{f})$ est en général non convexe et multimodale à cause du terme de l'adéquation aux données $\|\mathbf{g} - \mathbf{H}(\mathbf{f})\|^2$.

Ici, pour illustrer l'usage de l'algorithme du GNC, il faut bien entendu préciser la nature de l'opérateur nonlinéaire $\mathbf{H}(\mathbf{f})$. Pour ceci, nous allons considérer le cas d'un problème de l'imagerie à ondes diffractées :

$$\begin{aligned} g(\mathbf{r}_i) &= \iint_D G_m(\mathbf{r}_i, \mathbf{r}') \phi(\mathbf{r}') f(\mathbf{r}') d\mathbf{r}', \quad \mathbf{r}_i \in S \\ \phi(\mathbf{r}) &= \phi_0(\mathbf{r}) + \iint_D G_o(\mathbf{r}, \mathbf{r}') \phi(\mathbf{r}') f(\mathbf{r}') d\mathbf{r}', \quad \mathbf{r} \in D \end{aligned}$$

La version discrétisée de ces deux équations :

$$\begin{aligned} \mathbf{g} &= \mathbf{G}_m \mathbf{F} \phi + \mathbf{b} \\ \phi &= \phi_0 + \mathbf{G}_o \mathbf{F} \phi \end{aligned}$$

Remplaçant la deuxième équation dans la première on a :

$$\begin{aligned} \mathbf{g} &= \mathbf{H}(\mathbf{f}) + \mathbf{b} \quad \text{avec} \quad \mathbf{H}(\mathbf{f}) = \mathbf{G}_m \mathbf{F} (\mathbf{I} - \mathbf{G}_o \mathbf{F})^{-1} \phi_0 \\ J(\mathbf{x}) &= \|\mathbf{y} - \mathbf{A}(\mathbf{x})\|^2 + \Omega(\mathbf{x}) \end{aligned} \tag{10.1}$$

$$\mathbf{A}(\mathbf{x}) = \mathbf{G}_m \mathbf{X} (\mathbf{I} - \mathbf{G}_o \mathbf{X})^{-1} \phi_0 \tag{10.2}$$

Un choix possible pour créer une suite de critères relaxés est :

$$\mathbf{H}_{c_k}(\mathbf{f}) = \mathbf{G}_m \mathbf{F} (\mathbf{I} - c_k \mathbf{G}_o \mathbf{F})^{-1} \phi_0 \tag{10.3}$$

avec $c_0 = 0$, and $\lim_{k \rightarrow \infty} c_k = 1$.

Notons que pour $c_0 = 0$ on a :

$$\mathbf{H}_0(\mathbf{f}) = \mathbf{G}_m \mathbf{F} \phi_0 \longrightarrow J_0(\mathbf{f}) = \|\mathbf{g} - \mathbf{G}_m \mathbf{F} \phi_0\|^2 + \Omega(\mathbf{f}).$$

et on retrouve l'approximation linéaire de BORN.

Pour $c_k = 1$, effectivement on retrouve :

$$\mathbf{H}_1(\mathbf{f}) = \mathbf{H}(\mathbf{f}) \longrightarrow J_1(\mathbf{f}) = J(\mathbf{f})$$

A COMPLETER

10.10 Problèmes et questions ouverts

Problème essentiel :

Coût de calcul très élevé du MAP

- potentiels convexes/non convexes $\longrightarrow$ optimisation locale/globale
- Optimisation stochastique ou déterministe

Solution proposée :

Relaxation déterministe: GNC

Non-convexité graduelle (GNC): Technique d'optimisation non-locale

A COMPLETER

Chapitre 11

Choix des hyperparamètres

Dans ce chapitre nous aborderons le problème de l'estimation des hyperparamètres, et en particulier, celui de la détermination du paramètre de la régularisation. Dans un premier temps, nous décrirons brièvement les méthodes classiques du type : adéquation au données, la courbe en L, la validation croisée, et la validation croisée généralisée. Ensuite, nous présenterons le problème général de l'estimation des hyperparamètres dans le cadre bayésien.

11.1 Position du problème

Dans les chapitres précédents, nous avons vu que la résolution d'un problème inverse ne peut se faire d'une manière satisfaisante que qu'avec l'apport d'une information *a priori* appropriée. L'idée de combiner l'information contenue dans les données avec celle de l'*a priori* est à la base de la notion de la régularisation.

Par exemple dans l'approche classique de la régularisation (au sens de Phillips, Towmay et Tikhonov), nous avons vu que la solution est définie par :

$$\hat{\mathbf{f}} = \arg \min_{\mathbf{f}} \left\{ \|\mathbf{g} - \mathbf{H}\mathbf{f}\|^2 + \lambda \|\mathbf{D}\mathbf{f}\|^2 \right\}$$

où λ est le paramètre de régularisation qui a pour rôle de faire un compromis entre la fidélité aux données et l'information *a priori* qui est ici la variation douce de la grandeur inconnue.

D'une manière plus générale on peut définir une solution régularisée par :

$$\hat{\mathbf{f}} = \arg \min_{\mathbf{f}} \left\{ \Delta_1(\mathbf{g}, \mathbf{H}\mathbf{f}) + \lambda \Delta_2(\mathbf{f}, \hat{\mathbf{f}}_\infty) \right\}$$

où Δ_1 et Δ_2 sont deux distances, la première dans l'espace des mesures $\mathbf{g}$ et la deuxième dans l'espace des inconnues $\mathbf{f}$.

Notant, la solution pour $\lambda = 0$ par $\hat{\mathbf{f}}_0$ et la solution pour $\lambda = \infty$ par $\hat{\mathbf{f}}_\infty$ on peut encore généraliser (ou réécrire la définition de la solution régularisée par :

$$\hat{\mathbf{f}} = \arg \min_{\mathbf{f}} \left\{ \Delta_1(\mathbf{f}, \hat{\mathbf{f}}_0) + \lambda \Delta_2(\mathbf{f}, \hat{\mathbf{f}}_\infty) \right\}$$

où Δ_1 et Δ_2 sont deux distances euclidiennes, $\hat{\mathbf{f}}_\infty$ est la solution que l'on obtient lorsque $\lambda \rightarrow \infty$ (la solution *a priori*) et $\hat{\mathbf{f}}_0$ est la solution que l'on obtient lorsque $\lambda \rightarrow 0$ (la solution inverse généralisée par exemple).

Arrivé à ce stade deux questions se sont posées :

- Régulariser de quelle façon ou comment choisir Δ_1 et Δ_2 ?
- Régulariser jusqu'à quel point ou comment choisir le paramètre de la régularisation ?

Les figures de la page suivante montrent, sur un exemple du problème de déconvolution, l'effet du paramètre λ sur la solution.

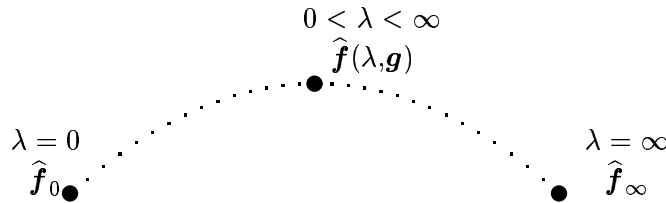
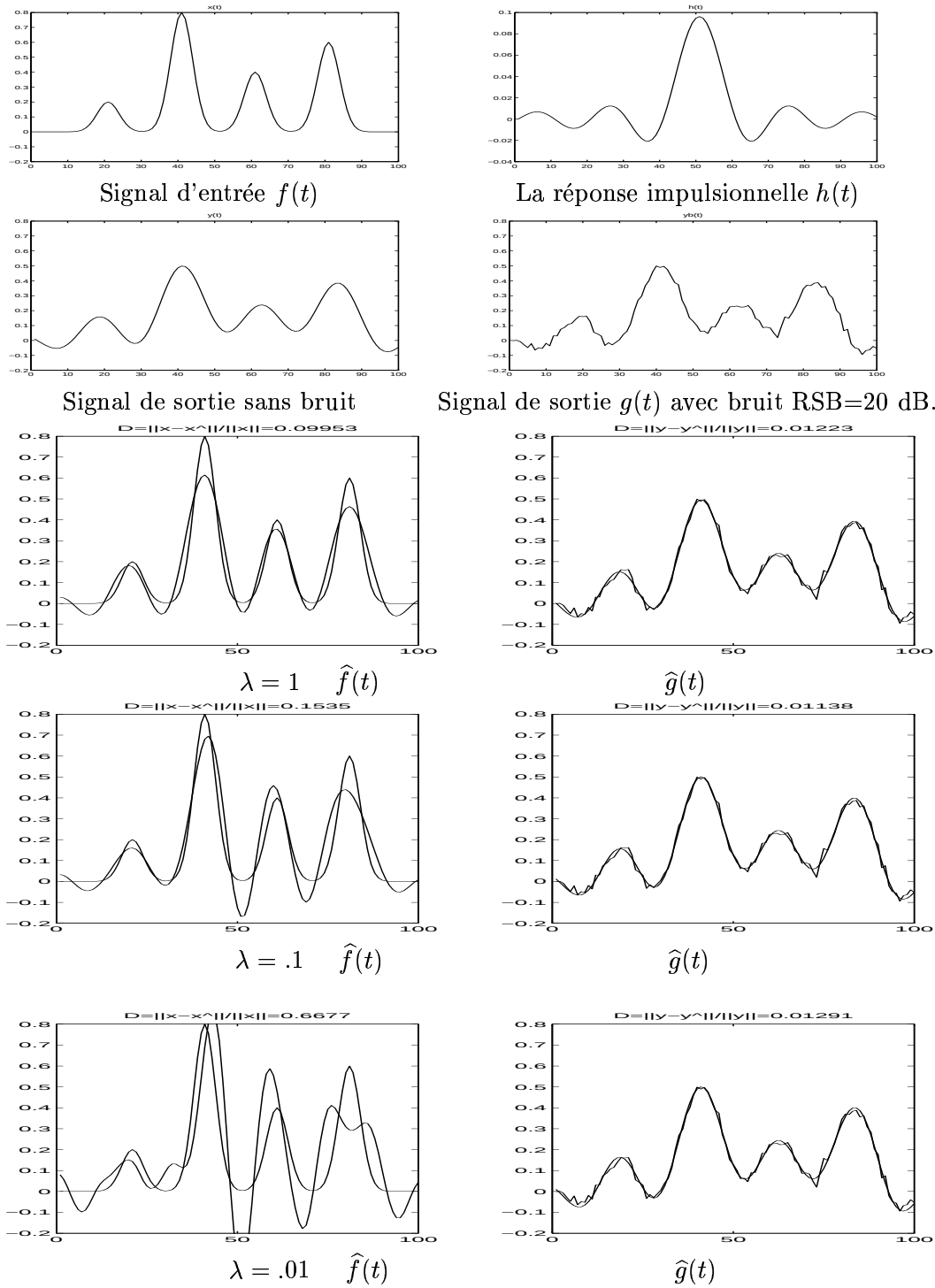


FIG. 11.1 – *Choix du paramètre de régularisation*

FIG. 11.2 – *Effet du coefficients de la régularisation sur le résultat d'une déconvolution.*

Nous avons vu ensuite que les approches probabilistes ont essayé, d'abord, répondre à la première question. Mais, la réponse à la deuxième question est resté à ce jour encore un problème ouvert.

Dans ce chapitre, d'abord, nous allons considérer le problème de la détermination du paramètre de la régularisation, ensuite, nous allons étendre la question dans l'approche bayésienne à la question plus générale de la détermination des hyperparamètres.

11.2 Choix du coefficient de régularisation

Dans cette section, nous allons considérer le problème de la détermination du paramètre de régularisation dans le cas de la régularisation quadratique :

$$\hat{\mathbf{f}} = \arg \min_{\mathbf{f}} \{J(\mathbf{f})\} \quad \text{avec} \quad J(\mathbf{f}) = \|\mathbf{g} - \mathbf{H}\mathbf{f}\|^2 + \lambda \|\mathbf{D}\mathbf{f}\|^2$$

En effet, dans ce cas simple, nous avons une expression analytique pour la solution régularisée :

$$\hat{\mathbf{f}}(\lambda) = [\mathbf{H}^t \mathbf{H} + \lambda \mathbf{D}^t \mathbf{D}]^{-1} \mathbf{H}^t \mathbf{g}$$

ce qui facilite la compréhension des principales méthodes de la détermination du paramètre de régularisation.

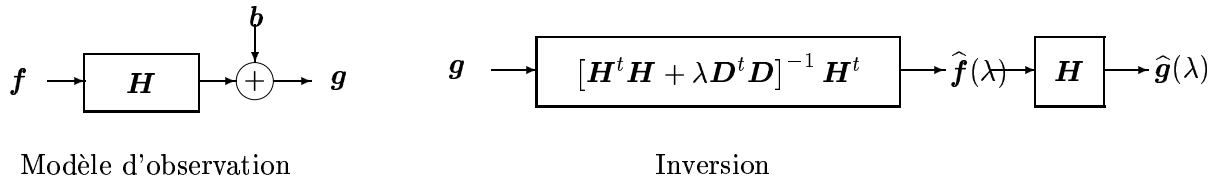


FIG. 11.3 – *Modèle d'observation et d'inversion en régularisation quadratique.*

Notant par :

$$\mathbf{A}(\lambda) = \mathbf{H} [\mathbf{H}^t \mathbf{H} + \lambda \mathbf{D}^t \mathbf{D}]^{-1} \mathbf{H}^t$$

on a les relations suivantes :

$$\begin{aligned} \|\mathbf{g} - \mathbf{H}\hat{\mathbf{f}}(\lambda)\|^2 &= \left\| \mathbf{g} - \mathbf{H} [\mathbf{H}^t \mathbf{H} + \lambda \mathbf{D}^t \mathbf{D}]^{-1} \mathbf{H}^t \mathbf{g} \right\|^2 \\ &= \|\mathbf{[I - A}(\lambda)]\mathbf{g}\|^2 \end{aligned}$$

$$\begin{aligned} \|\mathbf{H}\mathbf{f} - \hat{\mathbf{g}}(\lambda)\|^2 &= \left\| \mathbf{H}\mathbf{f} - \mathbf{H} [\mathbf{H}^t \mathbf{H} + \lambda \mathbf{D}^t \mathbf{D}]^{-1} \mathbf{H}^t \mathbf{g} \right\|^2 \\ &= \|\mathbf{H}\mathbf{f} - \mathbf{A}(\lambda)\mathbf{g}\|^2 \end{aligned}$$

$$\|\mathbf{f} - \hat{\mathbf{f}}(\lambda)\|^2 = \left\| \mathbf{f} - [\mathbf{H}^t \mathbf{H} + \lambda \mathbf{D}^t \mathbf{D}]^{-1} \mathbf{H}^t \mathbf{g} \right\|^2$$

Utilisant ces relations, on peut envisager un certain nombre de méthodes pour la détermination du λ :

- Adéquation aux données $\text{RSS}(\lambda) \stackrel{\text{def}}{=} \|\mathbf{g} - \mathbf{H}\hat{\mathbf{f}}(\lambda)\|^2$
- Risque moyen $\text{E} [\|\mathbf{f} - \hat{\mathbf{f}}(\lambda)\|^2]$
- Validation croisée $\text{E} [\|\mathbf{g} - \hat{\mathbf{g}}(\lambda)\|^2]$

11.2.1 Adéquation aux données

Le problème de la minimisation d'un critère composite

$$\min_{\mathbf{f}} \Delta_1(\mathbf{g}, \mathbf{H}\mathbf{f}) + \lambda \Delta_2(\mathbf{f}, \hat{\mathbf{f}}_\infty)$$

peut être interprété comme un problème d'optimisation sous contrainte

$$\min_{\mathbf{f}} \Delta_2(\mathbf{f}, \hat{\mathbf{f}}_\infty) \quad \text{s.c.} \quad \Delta_1(\mathbf{g}, \mathbf{H}\mathbf{f}) \leq c_1$$

ou encore

$$\min_{\mathbf{f}} \Delta_1(\mathbf{g}, \mathbf{H}\mathbf{f}) \quad \text{s.c.} \quad \Delta_2(\mathbf{f}, \hat{\mathbf{f}}_\infty) \leq c_2$$

où c_1 et c_2 sont deux constantes qu'il faudra déterminer.

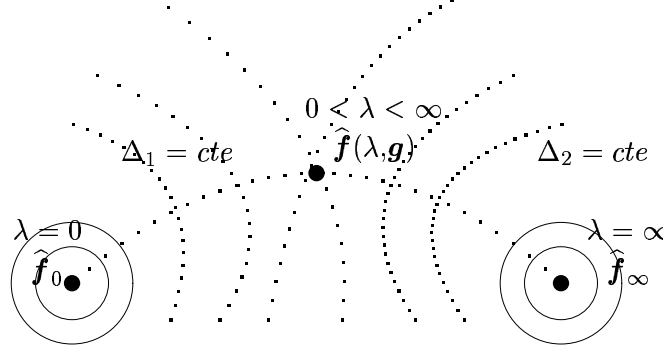


FIG. 11.4 – *Choix du paramètre de régularisation par adéquation aux données*

Pour mieux voir la limitation de cette approche, considérons le deuxième cas. c_2 peut alors être considéré comme une statistique dont sa loi découle de la loi du bruit ou d'une façon plus générale de la loi $p(\mathbf{g}|\mathbf{f})$.

Supposons par exemple que $b_i \sim v.a.i. \mathcal{N}(0, \sigma_b^2)$. Alors Δ_1 devient une distance quadratique. Supposons aussi que Δ_2 est choisi quadratique et égale à $\|\mathbf{D}\mathbf{f}\|^2$. On a alors :

$$\hat{\mathbf{f}}(\lambda) = \arg \min_{\mathbf{f}} \left\{ \|\mathbf{g} - \mathbf{H}\mathbf{f}\|^2 + \lambda \|\mathbf{D}\mathbf{f}\|^2 \right\} = \left[\mathbf{H}^t \mathbf{H} + \lambda \mathbf{D}^t \mathbf{D} \right]^{-1} \mathbf{H}^t \mathbf{g}$$

$$\text{RSS}(\lambda) = \left[\mathbf{g} - \mathbf{H}\hat{\mathbf{f}}(\lambda) \right]^t \left[\mathbf{g} - \mathbf{H}\hat{\mathbf{f}}(\lambda) \right] = \left\| \mathbf{g} - \mathbf{H}\hat{\mathbf{f}}(\lambda) \right\|^2$$

$$c : \chi^2(N) \longrightarrow c = N \text{ ou } N + 2\sqrt{N}$$

$$\text{RSS}(\lambda) = N\sigma_b^2 \longrightarrow \lambda_\chi$$

- Il faut connaître σ_b^2
- $\left[\mathbf{g} - \mathbf{H}\hat{\mathbf{f}}(\lambda) \right] \neq [\mathbf{g} - \mathbf{H}\mathbf{f}]$
- sur-régularisation

11.2.2 Risque moyen

$$\text{RM} \stackrel{\text{def}}{=} \mathbb{E} [\|\mathbf{f} - \hat{\mathbf{f}}(\lambda, \mathbf{g})\|^2] = \int [\mathbf{f} - \hat{\mathbf{f}}(\lambda, \mathbf{g})]^2 p(\mathbf{g}|\mathbf{f}, \lambda) d\mathbf{g}$$

$$\hat{\mathbf{f}}(\lambda) = [\mathbf{H}^t \mathbf{H} + \lambda \mathbf{D}^t \mathbf{D}]^{-1} \mathbf{H}^t \mathbf{g}$$

$$\begin{aligned} \hat{\lambda} &= \arg \min_{\lambda} \left\{ \mathbb{E} [\|\mathbf{f} - \hat{\mathbf{f}}(\lambda)\|^2] \right\} \\ &= \|\mathbf{I} - \mathbf{A}(\lambda)\| \mathbf{H} \mathbf{f}\|^2 + \sigma_b^2 \text{Tr}\{\mathbf{A}(\lambda)\}^2 \end{aligned}$$

où

$$\mathbf{A}(\lambda) = \mathbf{H} [\mathbf{H}^t \mathbf{H} + \lambda \mathbf{D}^t \mathbf{D}]^{-1} \mathbf{H}^t$$

et si $\mathbf{H}$ n'est pas singulière

$$\mathbf{A}(\lambda) = [\mathbf{I} - \lambda \mathbf{Q}]^{-1} \quad \text{avec} \quad \mathbf{Q} = [\mathbf{H}^t]^{-1} \mathbf{D}^t \mathbf{D} \mathbf{H}^{-1}$$

$$\lambda_{\text{RM}} = \arg \min_{\lambda} \{\text{RM}(\lambda, \mathbf{f})\}$$

- Il faut connaître σ_b^2
- $\text{RM}(\lambda, \mathbf{f})$ dépend du vrai $\mathbf{f}$ qui est inconnu

11.2.3 Validation croisée

L'idée de base dans les techniques de validations croisées est de partitionner les données en deux sous ensembles. On utilisera alors les données $\mathbf{g}_1$ dans le première sous ensemble pour calculer la solution $\hat{\mathbf{f}}(\lambda)$ et en déduire (ou prédire) celles du deuxième ; $\hat{\mathbf{g}}_2 = \mathbf{H}\hat{\mathbf{f}}(\lambda)$. On peut alors déterminer la valeur optimale (au sens de la VC) du λ par :

$$\hat{\lambda} = \arg \min_{\lambda} \left\{ \mathbb{E} [\|\mathbf{g}_2 - \hat{\mathbf{g}}_2(\lambda)\|^2] \right\}$$

Une partition la plus simple possible est d'éliminer à chaque étape une seule donnée g_k , c'est à dire :

$$P_1 = \{g_1, \dots, g_{k-1}, g_{k+1}, \dots, g_M\}, \quad P_2 = \{g_k\}$$

Notons par $\mathbf{g}^{(-k)} = [g_1, \dots, g_{k-1}, g_{k+1}, \dots, g_N]$. Le critère de VC peut alors s'écrire :

$$\text{VC}(\lambda) \stackrel{\text{def}}{=} \sum_{k=1}^M \left[g_k - \mathbf{h}_k^t \hat{\mathbf{f}}(\lambda, \mathbf{g}^{(-k)}) \right]^2$$

et

$$\hat{\lambda}_{VC} = \arg \min_{\lambda} \{ \text{VC}(\lambda) \}$$

Pour aller un peu plus dans le détail, considérons le cas de la régularisation quadratique :

$$\hat{\mathbf{f}}(\lambda, \mathbf{g}^{(-k)}) = \left[\mathbf{H}^t \mathbf{H} - \mathbf{h}_k \mathbf{h}_k^t + \lambda \mathbf{P}_0^{-1} \right]^{-1} \left[\mathbf{H}^t \mathbf{g} - \mathbf{h}_k \mathbf{g}^{(-k)} \right]$$

+

lemme d'inversion des matrices

↓

$$\text{VC}(\lambda) = [\mathbf{B}(\lambda) [\mathbf{I} - \mathbf{A}(\lambda)] \mathbf{g}]^2$$

- $\mathbf{B}(\lambda) = \text{diag} \{ (1 - a_{kk})^{-1} \}$
- Le calcul de $\text{VC}(\lambda)$ se ramène à celui des éléments diagonaux de $\mathbf{A}(\lambda)$
- Si $\mathbf{A}(\lambda)$ diagonale, $\text{VC}(\lambda)$ n'a pas de minimum

11.2.4 Validation croisée généralisée

$$\text{VCG}(\lambda) = \frac{\text{VC}(\lambda)}{\text{Tr} \{ \mathbf{I} - \mathbf{A}(\lambda) \}^2}$$

$$\text{VCG}(\lambda) \stackrel{\text{def}}{=} \sum_{k=1}^M w_k(\lambda)^2 \left[g_k - \mathbf{h}_k^t \hat{\mathbf{f}}(\lambda, \mathbf{g}^{(-k)}) \right]^2$$

coeffs de pondération : $w_k(\lambda) = (1 - a_{kk})(1 - \text{Tr} \{ \mathbf{I} - \mathbf{A}(\lambda) \})^{-1}$

$$\hat{\lambda}_{VCG} = \arg \min_{\lambda} \{ \text{VCG}(\lambda) \}$$

Calcul de $\text{VCG}(\lambda)$

Ici aussi, comme dans le cas précédent, en utilisant la lemme d'inversion des matrices, il est possible d'exprimer le critère en fonctions des termes calculable d'une manière récursive :

$$\text{VCG}(\lambda) = \sum_{k=1}^M \left(\frac{u_k}{v_k} \right)^2$$

- $u_k = g_k - \mathbf{h}_k^t \hat{\mathbf{f}}(\lambda, \mathbf{g}^{(-k)})$ résidu
- $v_k = a_{kk} - 1$
- Auxiliaire de calcul :

$$\begin{aligned} \mathbf{P}_{M|M} &= (\mathbf{H}^t \mathbf{H} + \lambda \mathbf{P}_0^t)^{-1} \quad \text{Covariance de } (\hat{\mathbf{f}} - \mathbf{f}) \\ \mathbf{A}(\lambda) &= \mathbf{H} \mathbf{P}_{M|M} \mathbf{H}^t \end{aligned}$$

11.3 Un autre regard

$$\mathbf{f} \sim \mathcal{N}(0, \sigma_x^2 \mathbf{P}_0)$$

$$\mathbf{b} \sim \mathcal{N}(0, \sigma_b^2 \mathbf{I})$$

$$\mathbf{g} | \mathbf{f} \sim \mathcal{N}(\mathbf{H}\mathbf{f}, \sigma_b^2 \mathbf{I})$$

Problème 1 :

Étant donné $\mathbf{H}$, $\sigma_x^2, \mathbf{P}_0, \sigma_b^2$ et $\mathbf{g}$, estimer $\mathbf{f}$

Solution :

$$\mathbf{f} | \mathbf{g} \sim \mathcal{N}(\hat{\mathbf{f}}, \mathbf{P}) \quad \left\{ \begin{array}{l} \hat{\mathbf{f}} = \left(\mathbf{H}^t \mathbf{H} + \lambda \mathbf{P}_0^{-1} \right)^{-1} \mathbf{H}^t \mathbf{g} \\ \mathbf{P} = \sigma_b^2 \left(\mathbf{H}^t \mathbf{H} + \lambda \mathbf{P}_0^{-1} \right)^{-1} \end{array} \right., \quad \lambda = \frac{\sigma_b^2}{\sigma_x^2}$$

Problème 2 :

Étant donné $\mathbf{H}$, $\mathbf{P}_0, \sigma_b^2$ et $\mathbf{g}$, estimer σ_x^2 et $\mathbf{f}$

ou d'une manière équivalente :

Étant donné $\mathbf{H}$, $\mathbf{P}_0, \sigma_b^2$, et $\mathbf{g}$, estimer λ et $\mathbf{f}$

Problème 3 :

Étant donné $\mathbf{H}$, $\mathbf{P}_0$ et $\mathbf{g}$, estimer σ_x^2, σ_b^2 et $\mathbf{f}$

ou d'une manière équivalente :

Étant donné $\mathbf{H}$, $\mathbf{P}_0$ et $\mathbf{g}$, estimer λ, σ_b^2 et $\mathbf{f}$

Notons : $\boldsymbol{\theta} = [\sigma_x^2, \sigma_b^2]$ ou $\boldsymbol{\theta} = [\lambda, \sigma_b^2]$

Adéquation aux données

$$Z = \frac{1}{\sigma_b^2} \|\mathbf{g} - \mathbf{H}\mathbf{f}\|^2 \sim \text{loi de } \chi^2(N)$$

$$\mathbb{E} [\|\mathbf{g} - \mathbf{H}\mathbf{f}\|^2] = N\sigma_b^2$$

– Première approximation :

$$\mathbb{E} [\|\mathbf{g} - \mathbf{H}\mathbf{f}\|^2] \approx \text{RSS}(\lambda) = \|\mathbf{g} - \mathbf{H}\hat{\mathbf{f}}(\lambda)\|^2$$

$$\|\mathbf{g} - \mathbf{H}\hat{\mathbf{f}}(\lambda)\|^2 = N\sigma_b^2 \longrightarrow \hat{\lambda} \quad (11.1)$$

– Deuxième approximation :

$$\mathbb{E} [\|\mathbf{g} - \mathbf{H}\mathbf{f}\|^2] \approx \frac{\text{RSS}(\lambda)}{\text{Tr}\{\mathbf{I} - \mathbf{A}(\lambda)\}^2} = \frac{\|\mathbf{g} - \mathbf{H}\hat{\mathbf{f}}(\lambda)\|^2}{\text{Tr}\{\mathbf{I} - \mathbf{A}(\lambda)\}^2}$$

où

$$\mathbf{A}(\lambda) = \mathbf{H} [\mathbf{H}^t \mathbf{H} + \lambda \mathbf{D}^t \mathbf{D}]^{-1} \mathbf{H}^t$$

$$\frac{\|\mathbf{g} - \mathbf{H}\hat{\mathbf{f}}(\lambda)\|^2}{\text{Tr}\{\mathbf{I} - \mathbf{A}(\lambda)\}^2} = N\sigma_b^2 \longrightarrow \hat{\lambda} \quad (11.2)$$

- Les deux équations (11.1) et (11.2) sont non linéaires
- Dans le cadre des approximations circulantes des matrices, il est possible d'obtenir des relation explicites pour leurs solutions.

Risque moyen

$$\mathbb{E} [\|\mathbf{f} - \hat{\mathbf{f}}(\boldsymbol{\theta}, \mathbf{g})\|^2] = \int [\mathbf{f} - \hat{\mathbf{f}}(\boldsymbol{\theta}, \mathbf{g})]^2 p(\mathbf{g} | \mathbf{f}, \boldsymbol{\theta}) d\mathbf{g}$$

Validation croisée

$$\mathbb{E} [\|\mathbf{g} - \mathbf{H}\hat{\mathbf{f}}(\boldsymbol{\theta}, \mathbf{g})\|^2] = \int [\mathbf{g} - \mathbf{H}\hat{\mathbf{f}}(\boldsymbol{\theta}, \mathbf{g})]^2 p(\mathbf{g} | \mathbf{f}, \boldsymbol{\theta}) d\mathbf{g}$$

Validation croisée généralisée

$$\mathbb{E} [\|\mathbf{H}\mathbf{f} - \hat{\mathbf{g}}(\boldsymbol{\theta}, \mathbf{g})\|^2] = \int [\mathbf{H}\mathbf{f} - \hat{\mathbf{g}}(\boldsymbol{\theta}, \mathbf{g})]^2 p(\mathbf{g} | \mathbf{f}, \boldsymbol{\theta}) d\mathbf{g}$$

Vraisemblance généralisée

$$(\hat{\mathbf{f}}, \hat{\boldsymbol{\theta}}) = \arg \max_{(\mathbf{f}, \boldsymbol{\theta})} \{p(\mathbf{f}, \mathbf{g}; \boldsymbol{\theta})\} = \arg \max_{(\mathbf{f}, \boldsymbol{\theta})} \{p(\mathbf{g}|\mathbf{f}; \boldsymbol{\theta})p(\mathbf{f}; \boldsymbol{\theta})\}$$

Vraisemblance marginale

$$\begin{cases} \hat{\boldsymbol{\theta}} &= \arg \max_{\boldsymbol{\theta}} \{p(\mathbf{g}; \boldsymbol{\theta})\} = \arg \max_{\boldsymbol{\theta}} \left\{ \int p(\mathbf{g}|\mathbf{f}; \boldsymbol{\theta})p(\mathbf{f}; \boldsymbol{\theta}) \, d\mathbf{f} \right\} \\ \hat{\mathbf{f}} &= \arg \max_{\mathbf{f}} \{p(\mathbf{f}|\mathbf{g}; \hat{\boldsymbol{\theta}})\} \end{cases}$$

La principale difficulté réside dans le calcul de $p(\mathbf{g}; \boldsymbol{\theta})$. C'est pourquoi, en général, on utilise l'algorithme EM pour calculer la solution $\hat{\boldsymbol{\theta}}$.

Pseudo-vraisemblance

Lorsque $\boldsymbol{\theta}$ constitue les paramètres de la loi *a priori*, $p(\mathbf{f}; \boldsymbol{\theta})$ et lorsque cette loi est une loi markovienne, on peut envisager d'approximer la vraisemblance par

$$p(\mathbf{g}; \boldsymbol{\theta}) \approx p(\mathbf{g}|\mathbf{f}; \boldsymbol{\theta}) \sum_i p(f_i|f_j; \boldsymbol{\theta}; j \in v_i)$$

Contexte complètement bayésien

$$(\hat{\mathbf{f}}, \hat{\boldsymbol{\theta}}) = \arg \max_{(\mathbf{f}, \boldsymbol{\theta})} \{p(\mathbf{f}, \boldsymbol{\theta}|\mathbf{g})\} = \arg \max_{(\mathbf{f}, \boldsymbol{\theta})} \{p(\mathbf{g}|\mathbf{f}, \boldsymbol{\theta})p(\mathbf{f}|\boldsymbol{\theta})p(\boldsymbol{\theta})\}$$

Ici, la principale difficulté est le choix de la loi $p(\boldsymbol{\theta})$ de telle sorte que $p(\mathbf{f}, \boldsymbol{\theta}|\mathbf{g})$ soit unimodale et que la solution puisse exister et soit unique.

11.4 Algorithme EM

(Expectation-Maximization)

Rechercher la solution à maximum de vraisemblance

Exemple :

$$\hat{\theta} = \arg \max_{\theta} \{p(\mathbf{g}; \theta)\} = \arg \max_{\theta} \{\ln p(\mathbf{g}; \theta)\}$$

- vraisemblance $l(\theta) = p(\mathbf{g}; \theta)$
- log-vraisemblance $L(\theta) = \ln p(\mathbf{g}; \theta)$
- Calcul de $p(\mathbf{g}; \theta)$ n'est souvent pas aisé
- Algorithme EM permet d'atteindre le maximum sans calculer explicitement $p(\mathbf{g}; \theta)$:

Le principe de l'algorithme EM consiste à :

- Introduire une variable auxiliaire $\mathbf{z}$ reliée à $\mathbf{g}$
- $\mathbf{z}$ données complètes
- $\mathbf{g}$ données incomplètes
- $\{\mathbf{z}\} \mapsto \{\mathbf{g}\}$ Projection : $\mathbf{g} = \mathbf{H}\mathbf{z}$

$$p(\mathbf{z}|\mathbf{g}; \theta) = \frac{p(\mathbf{z}, \mathbf{g}; \theta)}{p(\mathbf{g}; \theta)} = \frac{p(\mathbf{z}; \theta)p(\mathbf{g}|\mathbf{z})}{p(\mathbf{g}; \theta)} = \frac{p(\mathbf{z}; \theta)I_{\mathbf{g}}(\mathbf{z})}{p(\mathbf{g}; \theta)}$$

où $I_{\mathbf{g}}(\mathbf{z})$ est la fonction indicatrice : $I_{\mathbf{g}}(\mathbf{z}) = \begin{cases} 1 & \text{si } \mathbf{g} = \mathbf{H}\mathbf{z} \\ 0 & \text{sinon} \end{cases}$

$$p(\mathbf{g}; \theta) = \int_{I_{\mathbf{g}}(\mathbf{z})} p(\mathbf{z}; \theta) d\mathbf{z}$$

- Algorithme itératif:

$$\begin{cases} \text{(E)} & \text{Calculer } Q(\theta; \hat{\theta}^{(k)}) = \mathbb{E}_{\mathbf{z}|\mathbf{Y}; \hat{\theta}^{(k)}} \{\ln p(\mathbf{z}; \theta)\} \\ \text{(M)} & \text{Choisir } \hat{\theta}^{(k+1)} = \arg \max_{\theta} \left\{ Q(\theta; \hat{\theta}^{(k)}) \right\} \end{cases}$$

- $\hat{\theta}^k$ converge vers $\hat{\theta}$

$$\begin{aligned} Q(\theta; \hat{\theta}^{(k)}) &= \mathbb{E} \left[\ln p(\mathbf{z}; \theta) | \mathbf{g}; \hat{\theta}^{(k)} \right] \\ &= \mathbb{E}_{\mathbf{z}|\mathbf{Y}; \hat{\theta}^{(k)}} \{\ln p(\mathbf{z}; \theta)\} \\ &= \int_{I_{\mathbf{g}}(\mathbf{z})} \ln p(\mathbf{z}; \theta) p(\mathbf{z}|\mathbf{g}; \hat{\theta}^{(k)}) d\mathbf{z} \end{aligned}$$

11.5 Exemples et exercices :

Exercice 1 :

$$y = x + b, \quad x \sim \mathcal{N}(0, \sigma_x^2), \quad b \sim \mathcal{N}(0, \sigma_b^2)$$

Vérifier les relations suivantes :

$$\begin{aligned} y|x &\sim \mathcal{N}(x, \sigma_b^2), \\ y &\sim \mathcal{N}(0, \sigma_x^2 + \sigma_b^2) \\ \begin{pmatrix} x \\ y \end{pmatrix} &\sim \mathcal{N}\left(\begin{pmatrix} 0 \\ 0 \end{pmatrix}, \begin{pmatrix} \sigma_x^2 & \sigma_x^2 \\ \sigma_x^2 & \sigma_x^2 + \sigma_b^2 \end{pmatrix}\right) \\ x|y &\sim \mathcal{N}\left(\hat{x}, \frac{\sigma_x^2 \sigma_b^2}{\sigma_x^2 + \sigma_b^2}\right), \quad \text{avec} \quad \hat{x} = \frac{\sigma_x^2}{\sigma_x^2 + \sigma_b^2} y \end{aligned}$$

Exercice 2 :

Étant donné y et $\sigma_b^2 = 1$, estimer x et $\theta = \sigma_x^2$

Vérifier les relations suivantes :

$$\begin{aligned} x \sim \mathcal{N}(0, \theta) &\longrightarrow p(x; \theta) = \frac{1}{\sqrt{2\pi\theta}} \exp\left[-\frac{x^2}{2\theta}\right] \\ b \sim \mathcal{N}(0, 1) &\longrightarrow p(b) = \frac{1}{\sqrt{2\pi}} \exp\left[-\frac{b^2}{2}\right] \\ y|x \sim \mathcal{N}(x, 1) &\longrightarrow p(y|x) = \frac{1}{\sqrt{2\pi}} \exp\left[-\frac{(y-x)^2}{2}\right] \\ y \sim \mathcal{N}(0, \theta + 1) &\longrightarrow p(y; \theta) = \frac{1}{\sqrt{2\pi(\theta + 1)}} \exp\left[-\frac{y^2}{2(\theta + 1)}\right] \\ \Sigma &= \begin{pmatrix} \theta & \theta \\ \theta & \theta + 1 \end{pmatrix}, \quad \det |\Sigma| = \theta \\ \Sigma^{-1} &= \frac{1}{\theta} \begin{pmatrix} \theta + 1 & -\theta \\ -\theta & \theta \end{pmatrix}, \quad \det |\Sigma^{-1}| = 1 \\ \begin{pmatrix} x \\ y \end{pmatrix} \sim \mathcal{N}\left[\begin{pmatrix} 0 \\ 0 \end{pmatrix}, \Sigma\right] &\longrightarrow p(x, y; \theta) = \frac{\det |\Sigma|^{-\frac{1}{2}}}{2\pi} \exp\left[-\frac{1}{2}(x, y)\Sigma^{-1} \begin{pmatrix} x \\ y \end{pmatrix}\right] \\ x|y &\sim \mathcal{N}\left(\hat{x}, \frac{\theta}{\theta + 1}\right), \quad \hat{x} = \frac{\theta}{\theta + 1} y \\ p(x|y; \theta) &= \frac{1}{\sqrt{2\pi \frac{\theta}{\theta + 1}}} \exp\left[-\frac{(\theta + 1)(x - \hat{x})^2}{2\theta}\right] \end{aligned}$$

Vraisemblance généralisée :

$$\begin{aligned}
 \text{VG}(x, \theta|y) &= p(x, y; \theta) = p(y|x)p(x; \theta) \\
 &= \frac{1}{\sqrt{2\pi}} \exp \left[-\frac{(y-x)^2}{2} \right] \frac{1}{\sqrt{2\pi\theta}} \exp \left[-\frac{x^2}{2\theta} \right] \\
 &= \frac{1}{2\pi\sqrt{\theta}} \exp \left[-\frac{1}{2}(y^2 + \frac{\theta+1}{\theta}x^2 - 2xy) \right]
 \end{aligned}$$

Vraisemblance marginale :

$$\begin{aligned}
 \text{VM}(\theta|y) &= \int p(x, y; \theta) dx = p(y; \theta) \\
 &= \frac{1}{\sqrt{2\pi(\theta+1)}} \exp \left[-\frac{y^2}{2(\theta+1)} \right]
 \end{aligned}$$

Estimation au sens du MAP de x si θ connu :

$$\begin{aligned}
 \hat{x} &= \arg \max_x \{p(x|y; \theta)\} = \arg \max_x \{p(y|x; \theta)p(x; \theta)\} \\
 &= \arg \max_x \left\{ \frac{1}{\sqrt{2\pi}} \exp \left[-\frac{(y-x)^2}{2} \right] \frac{1}{\sqrt{2\pi\theta}} \exp \left[-\frac{x^2}{2\theta} \right] \right\} \\
 &= \arg \max_x \left\{ \frac{1}{2\pi\sqrt{\theta}} \exp \left[-\frac{1}{2}(y^2 + \frac{\theta}{\theta+1}x^2 - 2xy) \right] \right\} \\
 &= \frac{\theta}{\theta+1}y \longrightarrow \hat{x} \sim \mathcal{N} \left(\frac{\theta}{\theta+1}y, \frac{\theta}{\theta+1} \right)
 \end{aligned}$$

Estimation au sens du MV de x si θ connu :

$$\begin{aligned}
 \hat{x} &= \arg \max_x \{p(y|x; \theta)\} \\
 &= \arg \max_x \left\{ \frac{1}{\sqrt{2\pi}} \exp \left[-\frac{(y-x)^2}{2} \right] \right\} \\
 &= y \longrightarrow \hat{x} \sim \mathcal{N}(y, 1)
 \end{aligned}$$

Vraisemblance généralisée :

$$(\hat{x}, \hat{\theta}) = \arg \max_{(x, \theta)} \{p(x, y; \theta)\} = \arg \min_{(x, \theta)} \{-\ln p(x, y; \theta)\}$$

$$\begin{aligned} p(x, y; \theta) &= p(y|x)p(x; \theta) = \frac{1}{\sqrt{2\pi}} \exp \left[-\frac{(y-x)^2}{2} \right] \frac{1}{\sqrt{2\pi\theta}} \exp \left[-\frac{x^2}{2\theta} \right] \\ &= \frac{1}{2\pi\sqrt{\theta}} \exp \left[-\frac{1}{2} \left[(y-x)^2 + \frac{x^2}{\theta} \right] \right] \end{aligned}$$

$$L(x, \theta) = -\ln p(x, y; \theta) = \ln(2\pi) + \frac{1}{2} \left[\ln(\theta) + y^2 + \frac{\theta+1}{\theta} x^2 - 2xy \right]$$

$$\begin{cases} \frac{\partial L(x, \theta)}{\partial x} = \left(\frac{\theta+1}{\theta} x - y \right) = 0 & \longrightarrow x = \frac{\theta}{\theta+1} y \\ \frac{\partial L(x, \theta)}{\partial \theta} = \frac{1}{2} \left[\frac{1}{\theta} - \frac{x^2}{\theta^2} \right] = 0 & \longrightarrow \theta = x^2 \end{cases}$$

$$\begin{cases} \hat{x} = \frac{\theta}{\theta+1} y \\ \hat{\theta} = \hat{x}^2 = \left(\frac{\hat{\theta}}{\hat{\theta}+1} \right)^2 y^2 \end{cases}$$

Condition aux limites :

$$\theta = \left(\frac{\hat{\theta}}{\hat{\theta}+1} \right)^2 y^2 \longrightarrow \theta^2 + 2\theta + 1 - \theta y^2 = 0 \longrightarrow \Delta = (2 - y^2)^2 - 4 \geq 0 \longrightarrow 0 > y^2 > 4$$

$$\begin{cases} \hat{x} = \arg \max_x \{p(x|y; \hat{\theta})\} \\ \hat{\theta} = \arg \max_{\theta} \{p(\hat{x}; \theta)\} \end{cases}$$

Vraisemblance marginale :

$$\begin{cases} \hat{\theta} &= \arg \max_{\theta} \{p(y; \theta)\} \\ \hat{x} &= \arg \max_x \{p(x|y; \hat{\theta})\} \end{cases}$$

$$\begin{cases} p(y; \theta) &= \frac{1}{\sqrt{2\pi(\theta+1)}} \exp \left[-\frac{y^2}{2(\theta+1)} \right] \\ p(x|y; \theta) &= \frac{1}{\sqrt{2\pi\frac{\theta}{\theta+1}}} \exp \left[-\frac{(\theta+1)(x-\hat{x})^2}{2\theta} \right] \end{cases}$$

$$\begin{cases} \hat{\theta} &= y^2 - 1 \text{ si } y^2 > 1 \\ \hat{x} &= \frac{\hat{\theta}}{\hat{\theta}+1} y = y - \frac{1}{y} \end{cases}$$

Contexte complètement bayésien :

$$(\hat{x}, \hat{\theta}) = \arg \max_{(x, \theta)} \{p(x, \theta|y)\} = \arg \max_{(x, \theta)} \{p(y|x)p(x|\theta)p(\theta)\} = \arg \min_{(x, \theta)} \{-\ln p(y, x, \theta)\}$$

$$\begin{aligned} L(x, \theta) &= \ln p(x, y; \theta) \\ &= -\ln(2\pi) - \frac{1}{2} \left[\ln(\theta) + y^2 + \frac{\theta+1}{\theta} x^2 - 2xy \right] + \ln p(\theta) \\ \begin{cases} \frac{dL(x, \theta)}{dx} &= -\frac{\theta+1}{\theta} x + y = 0 \quad \longrightarrow x = \frac{\theta}{\theta+1} y \\ \frac{dL(x, \theta)}{d\theta} &= -\frac{1}{2} \left[\frac{3}{\theta} - \frac{x^2}{\theta^2} \right] = 0 \quad \longrightarrow \theta = \frac{1}{3} x^2 \end{cases} \end{aligned}$$

Avec un choix de $p(\theta) \propto \frac{1}{\theta}$ on a :

$$p(\theta) \propto \frac{1}{\theta} \quad \longrightarrow \quad \begin{cases} \hat{x} &= \frac{\theta}{\theta+1} y \\ \hat{\theta} &= \frac{1}{3} \hat{x}^2 \end{cases}$$

$$p(\theta) \propto \frac{1}{\sqrt{\theta}} \quad \longrightarrow \quad \begin{cases} \hat{x} &= \frac{\theta}{\theta+1} y \\ \hat{\theta} &= \frac{1}{2} \hat{x}^2 \end{cases}$$

Algorithme EM :

$$z = \begin{pmatrix} x \\ y \end{pmatrix} \longrightarrow p(z; \theta) = p(x, y; \theta) = \mathcal{N} \left(\begin{pmatrix} 0 \\ 0 \end{pmatrix}, \Sigma \right)$$

$$\Sigma = \begin{pmatrix} \theta & \theta \\ \theta & \theta + 1 \end{pmatrix}$$

$$-\frac{1}{2}(x, y) \Sigma^{-1} \begin{pmatrix} x \\ y \end{pmatrix} = -\frac{1}{2} \left(\frac{\theta + 1}{\theta} x^2 - 2xy + y^2 + \ln \det |\Sigma| + 2 \ln 2\pi \right)$$

$$\ln p(z; \theta) \propto -\ln \theta - \frac{\theta + 1}{\theta} x^2 + 2xy - y^2$$

$$p(z|y; \hat{\theta}) = p(x, y|y; \hat{\theta}) = p(x|y; \hat{\theta})$$

$$\mathbb{E} [\ln p(z; \theta) | y; \hat{\theta}] = \ln \theta - \frac{\theta + 1}{\theta} \mathbb{E} [x^2 | y; \hat{\theta}] + 2y \mathbb{E} [x | y; \hat{\theta}] - y^2$$

$$Q(\theta, \hat{\theta}) = \ln \theta - \frac{\theta + 1}{\theta} \mathbb{E} [x^2 | y; \hat{\theta}] + 2y \mathbb{E} [x | y; \hat{\theta}] + \text{cte}$$

$$\begin{aligned} \frac{\partial Q(\theta, \hat{\theta}^{(k)})}{\partial \theta} = 0 & \longrightarrow \hat{\theta}^{(k+1)} = \mathbb{E} [x^2 | y; \hat{\theta}^{(k)}] \\ & = \frac{\theta^{(k)}}{\theta^{(k)} + 1} + \left(\frac{\theta^{(k)}}{\theta^{(k)} + 1} y \right)^2 \end{aligned}$$

$$\hat{\theta}^{(k+1)} = -\frac{\theta^{(k)}}{\theta^{(k)} + 1} - \left(\frac{\theta^{(k)}}{\theta^{(k)} + 1} y \right)^2$$

Exercice 3 :**Extension au cas :** $\mathbf{g} = \mathbf{H}\mathbf{f} + \mathbf{b}$

$$\mathbf{f} \sim \mathcal{N}(0, \mathbf{R}_X = \sigma_x^2 \mathbf{P}_0), \quad \mathbf{b} \sim \mathcal{N}(0, \mathbf{R}_B = \sigma_b^2 \mathbf{I})$$

Vérifier les relations suivantes :

$$\begin{aligned} \mathbf{g}|\mathbf{f} &\sim \mathcal{N}(\mathbf{H}\mathbf{f}, \mathbf{R}_B = \sigma_b^2 \mathbf{I}) \\ \mathbf{g} &\sim \mathcal{N}(0, \mathbf{R}_Y) \quad \text{avec} \quad \mathbf{R}_Y = \mathbf{H}\mathbf{R}_X\mathbf{H}^t + \mathbf{R}_B \\ \begin{pmatrix} \mathbf{f} \\ \mathbf{g} \end{pmatrix} &\sim \mathcal{N}\left(\begin{pmatrix} 0 \\ 0 \end{pmatrix}, \mathbf{\Sigma}\right) \quad \text{avec} \quad \mathbf{\Sigma} = \begin{pmatrix} \mathbf{R}_X & \mathbf{R}_{XY} \\ \mathbf{R}_{YX} & \mathbf{R}_Y \end{pmatrix} \\ &\quad \text{où} \quad \mathbf{R}_{YX} = \mathbf{R}_{XY}^t = \mathbf{H}\mathbf{R}_X \\ \mathbf{f}|\mathbf{g} &\sim \mathcal{N}(\hat{\mathbf{f}}, \mathbf{P}) \end{aligned}$$

avec

$$\begin{cases} \hat{\mathbf{f}} &= \mathbf{R}_X\mathbf{H}^t(\mathbf{H}\mathbf{R}_X\mathbf{H}^t + \mathbf{R}_B)^{-1}\mathbf{g} = \mathbf{P}\mathbf{H}^t\mathbf{R}_B^{-1}\mathbf{g} \\ \mathbf{P} &= \mathbf{R}_X - \mathbf{R}_X\mathbf{H}^t(\mathbf{H}\mathbf{R}_X\mathbf{H}^t + \mathbf{R}_B)^{-1}\mathbf{H}\mathbf{R}_X = (\mathbf{R}_X^{-1} + \mathbf{H}^t\mathbf{R}_B^{-1}\mathbf{H})^{-1} \end{cases}$$

Avec $\mathbf{R}_B = \sigma_b^2 \mathbf{I}$ et $\mathbf{R}_X = \sigma_x^2 \mathbf{P}_0$ on obtient :

$$\begin{cases} \hat{\mathbf{f}} &= \left(\mathbf{H}^t\mathbf{H} + \lambda\mathbf{P}_0^{-1}\right)^{-1}\mathbf{H}^t\mathbf{g} \quad \text{avec} \quad \lambda = \frac{\sigma_b^2}{\sigma_x^2} = \frac{1}{\theta} \\ \mathbf{P} &= \sigma_b^2 \left(\mathbf{H}^t\mathbf{H} + \lambda\mathbf{P}_0^{-1}\right)^{-1} \end{cases}$$

Exercice 4 :

Étant donnée $\mathbf{H}, \mathbf{g}, \sigma_b^2 = 1$ et $\mathbf{P}_0$ estimer $\mathbf{f}$ et $\theta = \sigma_x^2$.

Vérifier et compléter les relations suivantes :

$$\begin{aligned}\mathbf{f} &\sim \mathcal{N}(0, \theta \mathbf{P}_0) \\ \mathbf{b} &\sim \mathcal{N}(0, \mathbf{I}) \\ \mathbf{g} | \mathbf{f} &\sim \mathcal{N}(\mathbf{H}\mathbf{f}, \mathbf{I}) \\ \mathbf{g} &\sim \mathcal{N}(0, \mathbf{R}_Y)\end{aligned}$$

$$\boldsymbol{\Sigma} = \theta \begin{pmatrix} \mathbf{P}_0 & \mathbf{P}_0^t \mathbf{H}^t \\ \mathbf{H} \mathbf{P}_0 & \mathbf{H} \mathbf{P}_0 \mathbf{H}^t + \mathbf{I} \end{pmatrix},$$

$$\det |\boldsymbol{\Sigma}| = \det |\mathbf{R}_X - \mathbf{R}_{XY} \mathbf{R}_Y^{-1} \mathbf{R}_{YX}| \det |\mathbf{R}_Y|$$

$$\boldsymbol{\Sigma}^{-1} = \begin{pmatrix} \mathbf{P}^{-1} & -\mathbf{P}^{-1} \mathbf{R}_{XY} \mathbf{R}_Y^{-1} \\ -\mathbf{R}_Y^{-1} \mathbf{R}_{YX} \mathbf{P}^{-1} & \mathbf{R}_Y^{-1} + \mathbf{R}_Y^{-1} \mathbf{R}_{YX} \mathbf{P}^{-1} \mathbf{R}_{XY} \mathbf{R}_Y^{-1} \end{pmatrix}$$

$$\det |\boldsymbol{\Sigma}^{-1}| =$$

$$\begin{aligned} \begin{pmatrix} \mathbf{f} \\ \mathbf{g} \end{pmatrix} &\sim \mathcal{N} \left[\begin{pmatrix} 0 \\ 0 \end{pmatrix}, \boldsymbol{\Sigma} \right] \\ &= \frac{\det |\boldsymbol{\Sigma}|^{-\frac{1}{2}}}{(2\pi)^{\frac{N+M}{2}}} \exp \left[-\frac{1}{2} [\mathbf{f}, \mathbf{g}]^t \boldsymbol{\Sigma}^{-1} \begin{pmatrix} \mathbf{f} \\ \mathbf{g} \end{pmatrix} \right] \end{aligned}$$

$$\mathbf{f} | \mathbf{g} \sim \mathcal{N}(\hat{\mathbf{f}}, \mathbf{P}) = (2\pi)^{-\frac{N}{2}} \det |\mathbf{P}|^{-\frac{1}{2}} \exp \left[-\frac{1}{2} [(\mathbf{f} - \hat{\mathbf{f}})]^t \mathbf{P}^{-1} [\mathbf{f} - \hat{\mathbf{f}}] \right]$$

avec

$$\begin{cases} \hat{\mathbf{f}} &= (\mathbf{H}^t \mathbf{H} + \lambda \mathbf{P}_0^{-1})^{-1} \mathbf{H}^t \mathbf{g} \quad \text{avec} \quad \lambda = \frac{1}{\theta} \\ \mathbf{P} &= (\mathbf{H}^t \mathbf{H} + \lambda \mathbf{P}_0^{-1})^{-1} \end{cases}$$

Vraisemblance généralisée :

$$\begin{aligned}
 \text{VG}(\mathbf{f}, \theta | \mathbf{g}) &= p(\mathbf{f}, \mathbf{g}; \theta) = p(\mathbf{g} | \mathbf{f}) p(\mathbf{f}; \theta) \\
 &= (2\pi)^{-\frac{M}{2}} \det |\mathbf{R}_B|^{-\frac{1}{2}} \exp \left[-\frac{1}{2} [\mathbf{g} - \mathbf{H}\mathbf{f}]^t \mathbf{R}_B^{-1} [\mathbf{g} - \mathbf{H}\mathbf{f}] \right] \times \\
 &\quad \times (2\pi)^{-\frac{N}{2}} \det |\mathbf{R}_X|^{-\frac{1}{2}} \exp \left[-\frac{1}{2} [\mathbf{f}^t \mathbf{R}_X^{-1} \mathbf{f}] \right] \\
 &= (2\pi)^{-\frac{N}{2}} (\theta \det |\mathbf{P}_0|)^{-\frac{1}{2}} \exp \left[-\frac{1}{2\theta} [\theta [\mathbf{g} - \mathbf{H}\mathbf{f}]^t [\mathbf{g} - \mathbf{H}\mathbf{f}] + \mathbf{f}^t \mathbf{P}_0^{-1} \mathbf{f}] \right]
 \end{aligned}$$

Vraisemblance marginale :

$$\begin{aligned}
 \text{VM}(\theta | \mathbf{g}) &= \int p(\mathbf{f}, \mathbf{g}; \theta) d\mathbf{f} = p(\mathbf{g}; \theta) \\
 &= (2\pi)^{-\frac{M}{2}} \det |\mathbf{R}_Y|^{-\frac{1}{2}} \exp \left[-\frac{1}{2} [\mathbf{g}^t \mathbf{R}_Y^{-1} \mathbf{g}] \right] \\
 \mathbf{R}_Y &= \theta \left(\mathbf{H} \mathbf{P}_0 \mathbf{H}^t + \frac{1}{\theta} \right)
 \end{aligned}$$

Estimation au sens du MAP de $\mathbf{f}$ si θ connu :

$$\begin{aligned}
 \hat{\mathbf{f}} &= \arg \max_{\mathbf{f}} \{p(\mathbf{f} | \mathbf{g}; \theta)\} = \arg \max_{\mathbf{f}} \{p(\mathbf{g} | \mathbf{f}; \theta) p(\mathbf{f}; \theta)\} \\
 &= \arg \min_{\mathbf{f}} \left\{ [\mathbf{g} - \mathbf{H}\mathbf{f}]^t \mathbf{R}_B^{-1} [\mathbf{g} - \mathbf{H}\mathbf{f}] + \mathbf{f}^t \mathbf{R}_X^{-1} \mathbf{f} \right\} \\
 &= \left(\mathbf{H}^t \mathbf{H} + \lambda \mathbf{P}_0^{-1} \right)^{-1} \mathbf{H}^t \mathbf{g}, \quad \lambda = \frac{1}{\theta}
 \end{aligned}$$

Estimation au sens du MV de $\mathbf{f}$ si θ connu :

$$\begin{aligned}
 \hat{\mathbf{f}} &= \arg \max_{\mathbf{f}} \{p(\mathbf{g} | \mathbf{f}; \theta)\} \\
 &= \arg \max_{\mathbf{f}} \left\{ (2\pi)^{-\frac{M}{2}} \det |\mathbf{R}_B|^{-\frac{1}{2}} \exp \left[-\frac{1}{2} [\mathbf{g} - \mathbf{H}\mathbf{f}]^t \mathbf{R}_B^{-1} [\mathbf{g} - \mathbf{H}\mathbf{f}] \right] \right\} \\
 &= \arg \min_{\mathbf{f}} \left\{ [\mathbf{g} - \mathbf{H}\mathbf{f}]^t \mathbf{R}_B^{-1} [\mathbf{g} - \mathbf{H}\mathbf{f}] \right\} \\
 &= (\mathbf{H}^t \mathbf{H})^{-1} \mathbf{H}^t \mathbf{g}
 \end{aligned}$$

Vraisemblance généralisée

$$(\hat{\mathbf{f}}, \hat{\theta}) = \arg \max_{(\mathbf{f}, \theta)} \{p(\mathbf{f}, \mathbf{g}; \theta)\} = \arg \max_{(\mathbf{f}, \theta)} \{\ln p(\mathbf{f}, \mathbf{g}; \theta)\}$$

$$\begin{aligned} p(\mathbf{f}, \mathbf{g}; \theta) &= p(\mathbf{g}|\mathbf{f})p(\mathbf{f}; \theta) \\ &= (2\pi)^{-\frac{M}{2}} \det |\mathbf{R}_B|^{-\frac{1}{2}} \exp \left[-\frac{1}{2} [\mathbf{g} - \mathbf{H}\mathbf{f}]^t \mathbf{R}_B^{-1} [\mathbf{g} - \mathbf{H}\mathbf{f}] \right] \times \\ &\quad \times (2\pi)^{-\frac{N}{2}} \det |\mathbf{R}_X|^{-\frac{1}{2}} \exp \left[-\frac{1}{2} [\mathbf{f}^t \mathbf{R}_X^{-1} \mathbf{f}] \right] \\ &= K \theta^{-\frac{1}{2}} \exp \left[-\frac{1}{2} [\mathbf{g} - \mathbf{H}\mathbf{f}]^t [\mathbf{g} - \mathbf{H}\mathbf{f}] - \frac{1}{2\theta} \mathbf{f}^t \mathbf{P}_0^{-1} \mathbf{f} \right] \end{aligned}$$

$$L(x, \theta) = \ln p(\mathbf{f}, \mathbf{g}; \theta) \propto \left[\ln \theta + [\mathbf{g} - \mathbf{H}\mathbf{f}]^t [\mathbf{g} - \mathbf{H}\mathbf{f}] + \frac{1}{\theta} \mathbf{f}^t \mathbf{P}_0^{-1} \mathbf{f} \right]$$

$$\begin{cases} \frac{dL(\mathbf{f}, \theta)}{d\mathbf{f}} = 0 & \longrightarrow \mathbf{f} = \left(\mathbf{H}^t \mathbf{H} + \lambda \mathbf{P}_0^{-1} \right)^{-1} \mathbf{H}^t \mathbf{g} \\ \frac{dL(\mathbf{f}, \theta)}{d\theta} = 0 & \longrightarrow \theta = \mathbf{f}^t \mathbf{P}_0^{-1} \mathbf{f} \end{cases}$$

$$\begin{cases} \hat{\mathbf{f}} &= \left(\mathbf{H}^t \mathbf{H} + \lambda \mathbf{P}_0 \right)^{-1} \mathbf{H}^t \mathbf{g} \\ \hat{\theta} &= \hat{\mathbf{f}}^t \mathbf{P}_0^{-1} \hat{\mathbf{f}} \end{cases}$$

$$\begin{cases} \hat{\mathbf{f}} &= \arg \max_{\mathbf{f}} \{p(\mathbf{f}|\mathbf{g}; \hat{\theta})\} \\ \hat{\theta} &= \arg \max_{\theta} \{p(\hat{\mathbf{f}}; \theta)\} \end{cases}$$

Vraisemblance marginale :

$$\begin{cases} \hat{\theta} &= \arg \max_{\theta} \{p(\mathbf{g}; \theta)\} = (2\pi)^{-\frac{M}{2}} \det |\mathbf{R}_Y|^{-\frac{1}{2}} \exp \left[-\frac{1}{2} [\mathbf{g}^t \mathbf{R}_Y^{-1} \mathbf{g}] \right] \\ \hat{\mathbf{f}} &= \arg \max_{\mathbf{f}} \{p(\mathbf{f}|\mathbf{g}; \hat{\theta})\} \end{cases}$$

$$\mathbf{R}_Y = \theta \mathbf{H} \mathbf{P}_0 \mathbf{H}^t + \mathbf{I}$$

$$\begin{aligned} \hat{\theta} &= \arg \min_{\theta} \left\{ \ln \det |\theta \mathbf{H} \mathbf{P}_0 \mathbf{H}^t + \mathbf{I}| - \mathbf{g}^t (\theta \mathbf{H} \mathbf{P}_0 \mathbf{H}^t + \mathbf{I})^{-1} \mathbf{g} \right\} \\ \hat{\mathbf{f}} &= \arg \max_{\mathbf{f}} \{p(\mathbf{f}|\mathbf{g}; \hat{\theta})\} \end{aligned}$$

Contexte complètement bayésien :

$$(\hat{\mathbf{f}}, \hat{\theta}) = \arg \max_{(\mathbf{f}, \theta)} \{p(\mathbf{f}, \theta | \mathbf{g})\} = \arg \max_{(\mathbf{f}, \theta)} \{p(\mathbf{g} | \mathbf{f}) p(\mathbf{f} | \theta) p(\theta)\}$$

$$p(\theta) \propto \frac{1}{\theta}$$

$$L(x, \theta) = \ln p(\mathbf{f}, \mathbf{g}; \theta) \propto \left[\ln \theta + [\mathbf{g} - \mathbf{H} \mathbf{f}]^t [\mathbf{g} - \mathbf{H} \mathbf{f}] + \frac{1}{\theta} \mathbf{f}^t \mathbf{P}_0^{-1} \mathbf{f} \right] - \ln \theta$$

$$\begin{cases} \frac{dL(\mathbf{f}, \theta)}{d\mathbf{f}} = 0 & \longrightarrow \mathbf{f} = (\mathbf{H}^t \mathbf{H} + \lambda \mathbf{P}_0)^{-1} \mathbf{H}^t \mathbf{g} \\ \frac{dL(\mathbf{f}, \theta)}{d\theta} = 0 & \longrightarrow \theta = \mathbf{f}^t \mathbf{P}_0^{-1} \mathbf{f} \end{cases}$$

$$\begin{cases} \hat{\mathbf{f}} &= (\mathbf{H}^t \mathbf{H} + \lambda \mathbf{P}_0)^{-1} \mathbf{H}^t \mathbf{g} \\ \hat{\theta} &= \hat{\mathbf{f}}^t \mathbf{P}_0^{-1} \hat{\mathbf{f}} \end{cases}$$

Algorithme EM

$$\mathbf{z} = \begin{pmatrix} \mathbf{f} \\ \mathbf{g} \end{pmatrix} \longrightarrow p(\mathbf{z}; \theta) = p(\mathbf{f}, \mathbf{g}; \theta) = \mathcal{N} \left(\begin{pmatrix} 0 \\ 0 \end{pmatrix}, \mathbf{\Sigma} \right)$$

$$\ln p(\mathbf{z}; \theta) \propto \ln \det |\mathbf{\Sigma}| - [\mathbf{f}, \mathbf{g}] \mathbf{\Sigma}^{-1} \begin{pmatrix} \mathbf{f} \\ \mathbf{g} \end{pmatrix}$$

$$p(\mathbf{z} | \mathbf{g}; \hat{\theta}) = p(\mathbf{f}, \mathbf{g} | \mathbf{g}; \hat{\theta}) = p(\mathbf{f} | \mathbf{g}; \hat{\theta})$$

$$\mathbb{E} \left[\ln p(\mathbf{z}; \theta) | \mathbf{g}; \hat{\theta} \right] = \ln \det |\mathbf{\Sigma}| - \mathbb{E} \left[[\mathbf{f}, \mathbf{g}] \mathbf{\Sigma}^{-1} \begin{pmatrix} \mathbf{f} \\ \mathbf{g} \end{pmatrix} \right]$$

$$Q(\theta, \hat{\theta}) = \ln \theta - \frac{\theta + 1}{\theta} \mathbb{E} \left[\mathbf{f}^2 | \mathbf{g}; \hat{\theta} \right] + 2\mathbf{g} \mathbb{E} \left[\mathbf{f} | \mathbf{g}; \hat{\theta} \right] + \text{cte}$$

$$\frac{dQ(\theta, \hat{\theta}^{(k)})}{d\theta} = 0 \longrightarrow \hat{\theta}^{(k+1)} = \mathbb{E} \left[\mathbf{f}^2 | \mathbf{g}; \hat{\theta}^{(k)} \right] = \frac{\theta^{(k)}}{\theta^{(k)} + 1} + \left(\frac{\theta^{(k)}}{\theta^{(k)} + 1} \mathbf{g} \right)^2$$

Chapitre 12

Étude de cas et exemples d'applications

Dans ce chapitre nous étudierons un certain nombre de problèmes classiques comme la déconvolution des signaux, la restauration d'image, la reconstruction d'images en tomographie X, la reconstruction d'images en tomographie à ondes diffractées, la reconstruction d'images par courants de FOUCAULT en contrôle non destructif, etc.

12.1 Déconvolution des signaux

12.1.1 Un problème d'instrumentation

Voici un exemple de problème où les mesures effectuées par l'instrument ne sont pas satisfaisantes et il est impossible de changer l'instrument. On voudrait cependant corriger les défauts de l'instrument par une méthode numérique de déconvolution.

Dans cet exemple nous allons étudier successivement les problèmes suivants :

- l'estimation de la réponse impulsionnelle d'un système de mesure à partir de sa réponse indicielle;
- l'estimation de la réponse impulsionnelle d'un système de mesure à partir de sa sortie à une entrée quelconque mais connue;
- l'estimation de l'entrée lorsque la réponse impulsionnelle est connue (déconvolution simple);
- l'estimation de l'entrée lorsque la réponse impulsionnelle n'est pas connue (déconvolution aveugle);

Nous examinerons :

- Méthodes naïves
- Filtre de Wiener
- Régularisation quadratique pour l'inversion
- Régularisation quadratique pour l'estimation de la réponse impulsionnelle
- Une méthode de déconvolution aveugle basée sur la régularisation quadratique

12.1.2 Méthodes naïves

Il s'agit d'examiner les résultats que l'on obtient lorsqu'on utilise une méthode naïve du type filtrage inverse pour inversion :

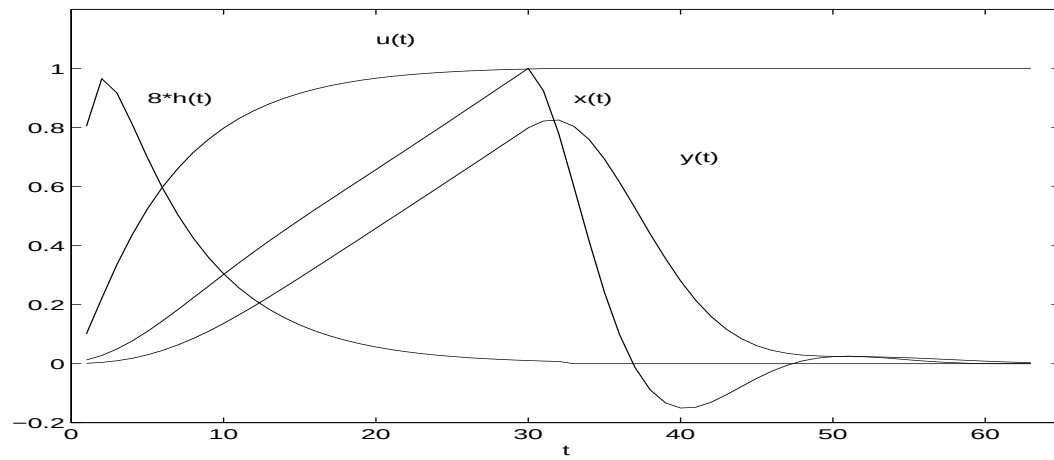
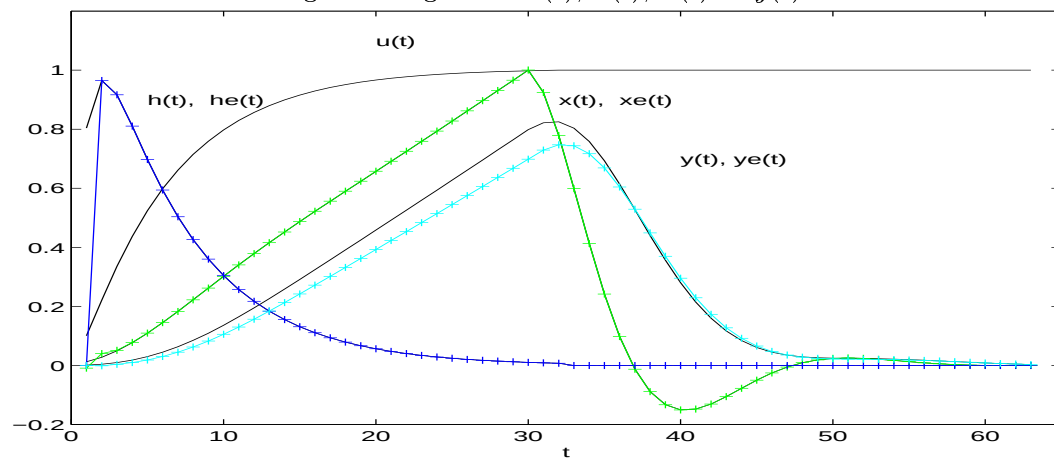
$$y(t) = h(t) * x(t) + b(t) \longrightarrow \text{TF} \longrightarrow Y(\omega) = H(\omega) X(\omega) + B(\omega) \longrightarrow$$

$$\hat{X}(\omega) = \frac{1}{H(\omega)} Y(\omega) \longrightarrow \text{TF inverse} \longrightarrow \hat{x}(t)$$

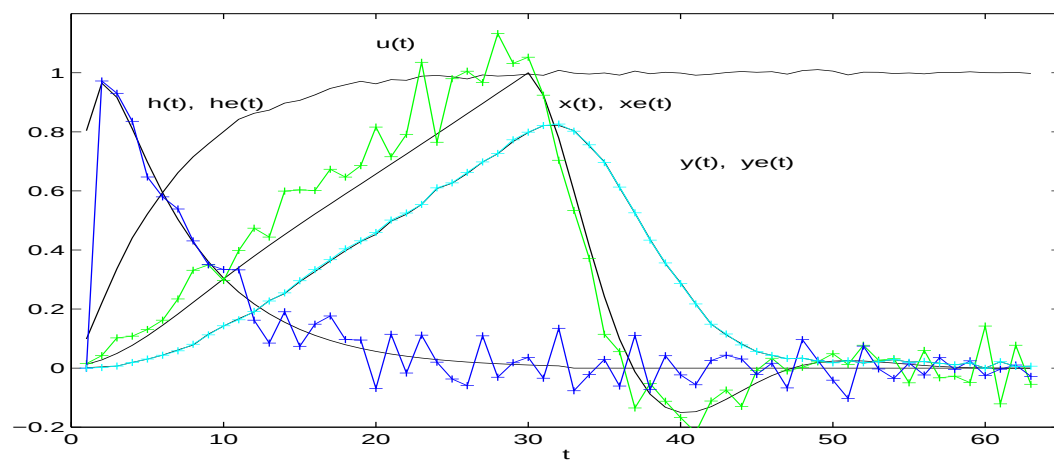
et du type dérivation de la réponse indicielle pour le calcul de la réponse impulsionnelle :

$$u(t) = \int_0^t h(t) dt \longrightarrow h(t) = \frac{du(t)}{dt}$$

Effectivement, comme le montre les figures qui suivent, ces méthodes peuvent marcher lorsqu'il n'y a pas du bruit. Mais, en présence du bruit, même très très faible, ces méthodes échouent très facilement.

Signaux originaux: $u(t)$, $h(t)$, $x(t)$ et $y(t)$ 

Pas de bruit: tout va bien.



Avec bruit: rien ne va plus

FIG. 12.1 – Méthodes naïves:

On calcule $h(t)$ en dérivant $u(t)$ et $x(t)$ par filtrage inverse. Lorsqu'il n'y a pas de bruit, tout va bien, mais lorsqu'on rajoute un bruit sur les données les résultats deviennent très vite inexploitable.

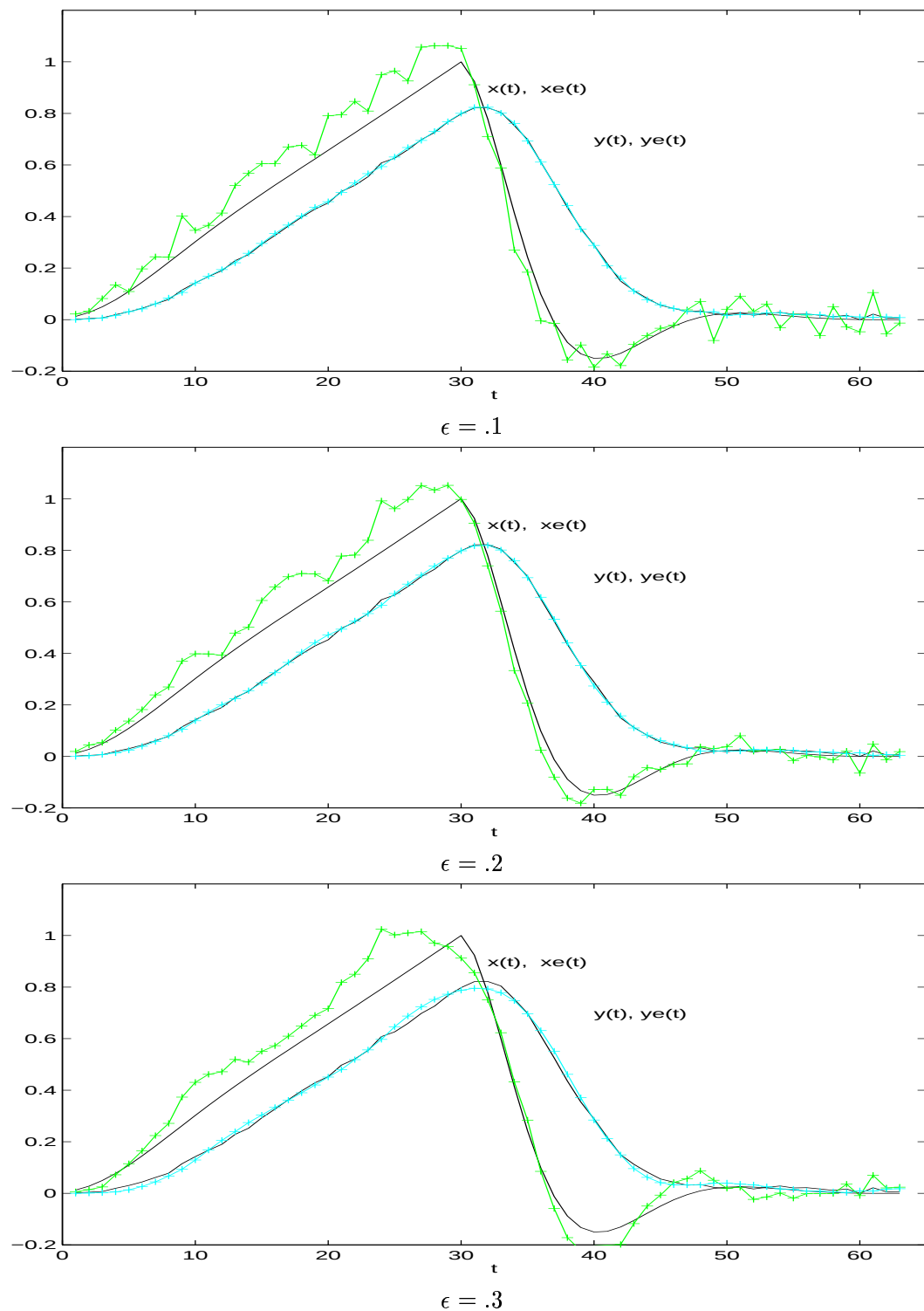


FIG. 12.2 – **Méthodes naïves :** On bricole en remplaçant tous les valeurs de $|H(\omega)| < \epsilon$ par $|H(\omega)| = \epsilon$ pour tenter d'éliminer la source d'ennui, mais sans succès. Trois valeurs de ϵ sont essayées.

12.1.3 Filtre de Wiener pour l'inversion

Dans cette étude, nous considérons seulement le problème de la déconvolution simple, en faisant l'hypothèse que la réponse impulsionnelle est connue parfaitement. La solution est calculée par

$$\hat{X}(\omega) = \frac{H^*(\omega)}{|H(\omega)|^2 + \frac{S_{bb}(\omega)}{S_{xx}(\omega)}} Y(\omega)$$

avec l'hypothèse $\frac{S_{bb}(\omega)}{S_{xx}(\omega)} = r$ et différentes valeurs pour r .

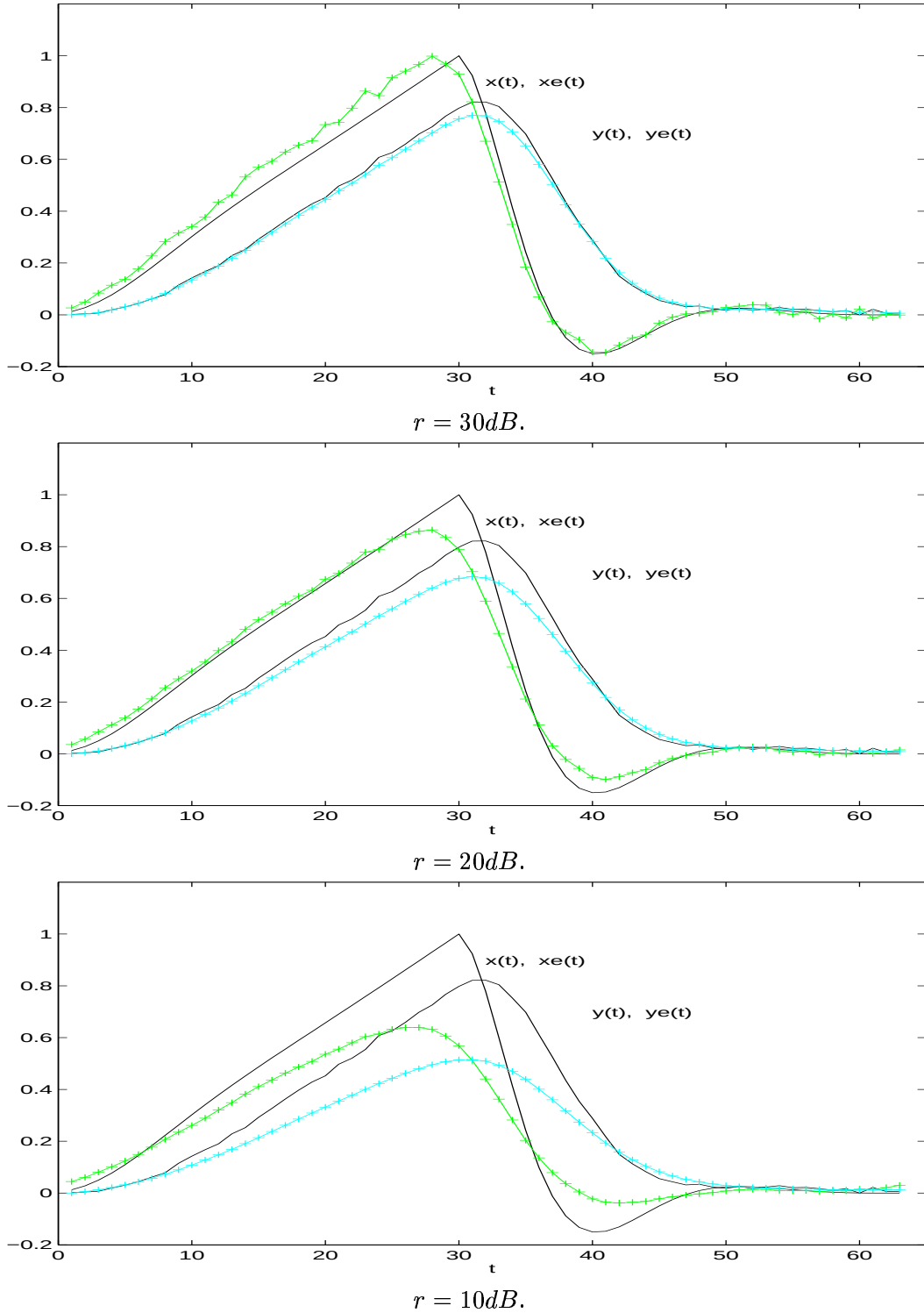


FIG. 12.3 – Filtre de Wiener en déconvolution :

On essay le filtre de Wiener avec l'hypothèse $\frac{S_{xx}(\omega)}{S_{bb}(\omega)} = r$ pour différentes valeurs de r .

12.1.4 Régularisation quadratique pour l'inversion

Dans cette étude, nous considérons seulement le problème de la déconvolution simple, en faisant l'hypothèse que la réponse impulsionnelle est connue parfaitement. La solution est calculée par

$$\hat{\mathbf{x}} = \left(\mathbf{H}^t \mathbf{H} + \lambda \mathbf{D}^t \mathbf{D} \right)^{-1} \mathbf{H}^t \mathbf{y}$$

avec différentes valeurs pour le paramètre de régularisation λ .

Deux cas sont étudiés correspondant à deux choix pour la matrice $\mathbf{D}$:

- régularisation d'ordre zéro : $\mathbf{D} = \mathbf{I}$
- régularisation d'ordre un : $\mathbf{D} = \text{TOEPLITZ}[-1,1]$ est une matrice des différences finies d'ordre un.

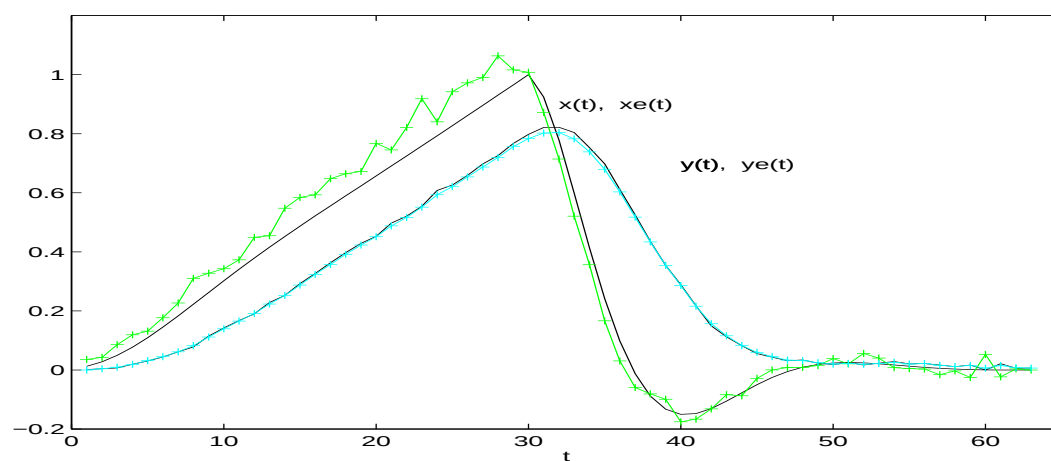
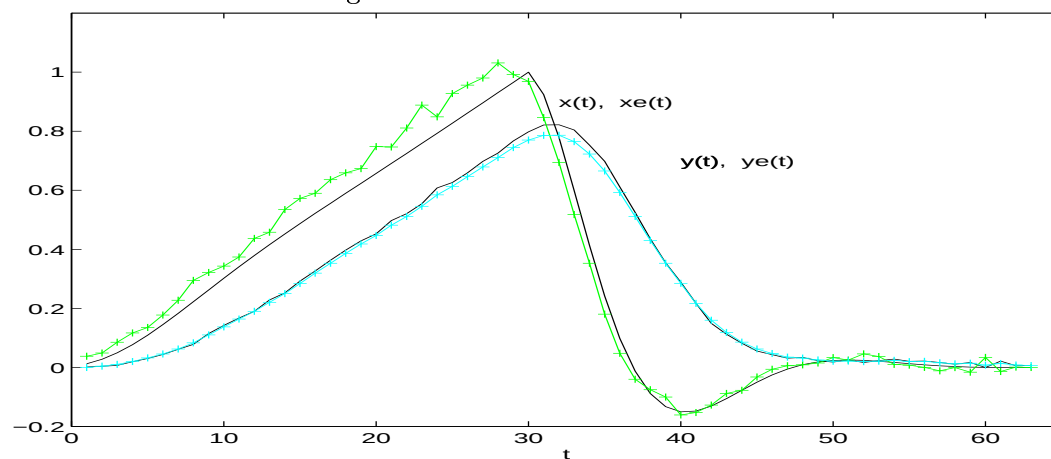
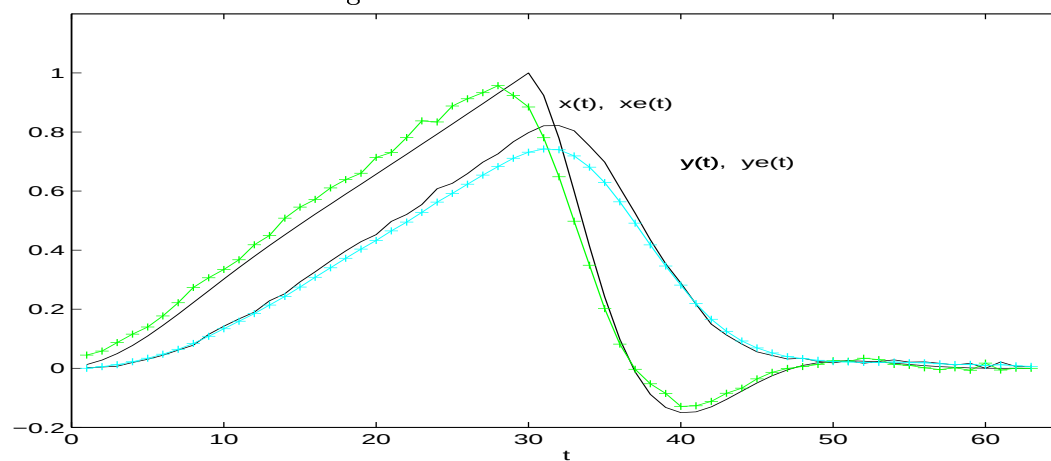
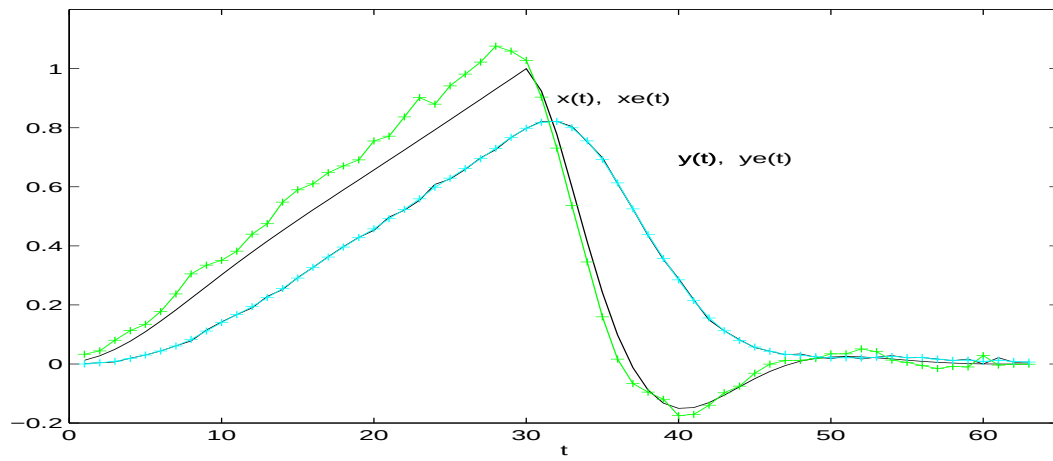
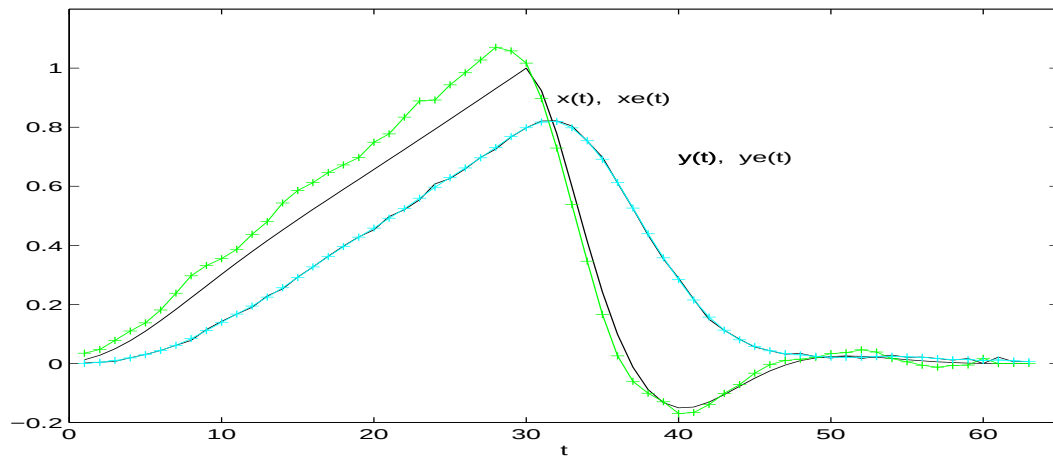
Régularisation d'ordre zéro. $\lambda = .01$ Régularisation d'ordre zéro. $\lambda = .02$ Régularisation d'ordre zéro. $\lambda = .05$

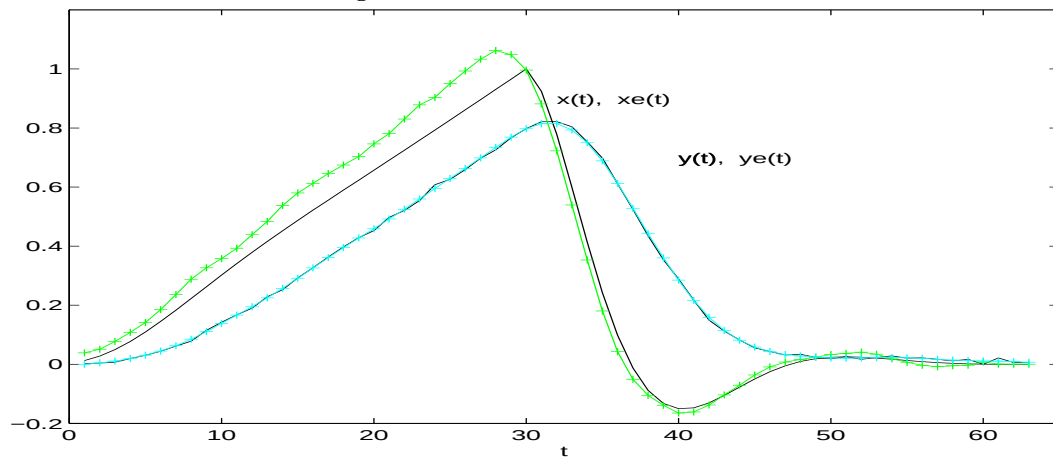
FIG. 12.4 – Déconvolution par régularisation d'ordre zéro.



Régularisation d'ordre un. $\lambda = .01$



Régularisation d'ordre un. $\lambda = .02$



Régularisation d'ordre un. $\lambda = .05$

FIG. 12.5 – Déconvolution par régularisation d'ordre un.

12.1.5 Régularisation quadratique pour l'estimation de la réponse impulsionnelle

Dans cette étude, nous considérons seulement le problème de l'estimation de la réponse impulsionnelle $h(t)$ à partir de la réponse indicielle $u(t)$. Notons que la relation entre ces deux grandeurs s'écrit :

$$u(t) = \int_0^t h(t) dt \longrightarrow h(t) = \frac{du(t)}{dt}$$

Discrétisation :

$$u(m) = \sum_{k=0}^m h(k) \longrightarrow \begin{pmatrix} u(0) \\ \vdots \\ u(M) \end{pmatrix} = \begin{pmatrix} 1 & 0 & \cdots & 0 \\ 1 & 1 & & \vdots \\ \vdots & & \ddots & \\ 1 & \cdots & & 1 \end{pmatrix} \begin{pmatrix} h(0) \\ \vdots \\ h(M) \end{pmatrix}$$

$$\mathbf{u} = \mathbf{A}\mathbf{h}$$

Solution régularisée :

$$\hat{\mathbf{h}} = (\mathbf{A}^t \mathbf{A} + \lambda \mathbf{D}^t \mathbf{D})^{-1} \mathbf{A}^t \mathbf{u}$$

Dans le cas plus général où l'entrée est un signal quelconque, on a :

$$y(t) = [h * f](t) + b(t)$$

Discrétisation :

$$y(m) = \sum_{k=0}^M x(k)h(m-k) + b(m)$$

$$\begin{pmatrix} y(0) \\ y(1) \\ \vdots \\ y(p) \\ \vdots \\ y(M) \end{pmatrix} = \begin{pmatrix} x(0) & 0 & \cdots & \cdots & \cdots & 0 \\ x(1) & x(0) & & & & \vdots \\ \vdots & & & & & \\ x(p) & & & & x(0) & \\ x(p+1) & x(p) & & & x(1) & \\ \vdots & & & & \vdots & \\ x(M) & \cdots & & \cdots & x(M-p) & \end{pmatrix} \begin{pmatrix} h(0) \\ \vdots \\ h(p) \end{pmatrix}$$

$$\mathbf{y} = \mathbf{X}\mathbf{h} + \mathbf{b} \text{ avec } \mathbf{X} = \text{TOEPLITZ}(\mathbf{x}),$$

Solution régularisée :

$$\hat{\mathbf{h}} = (\mathbf{X}^t \mathbf{X} + \lambda \mathbf{D}^t \mathbf{D})^{-1} \mathbf{X}^t \mathbf{y}$$

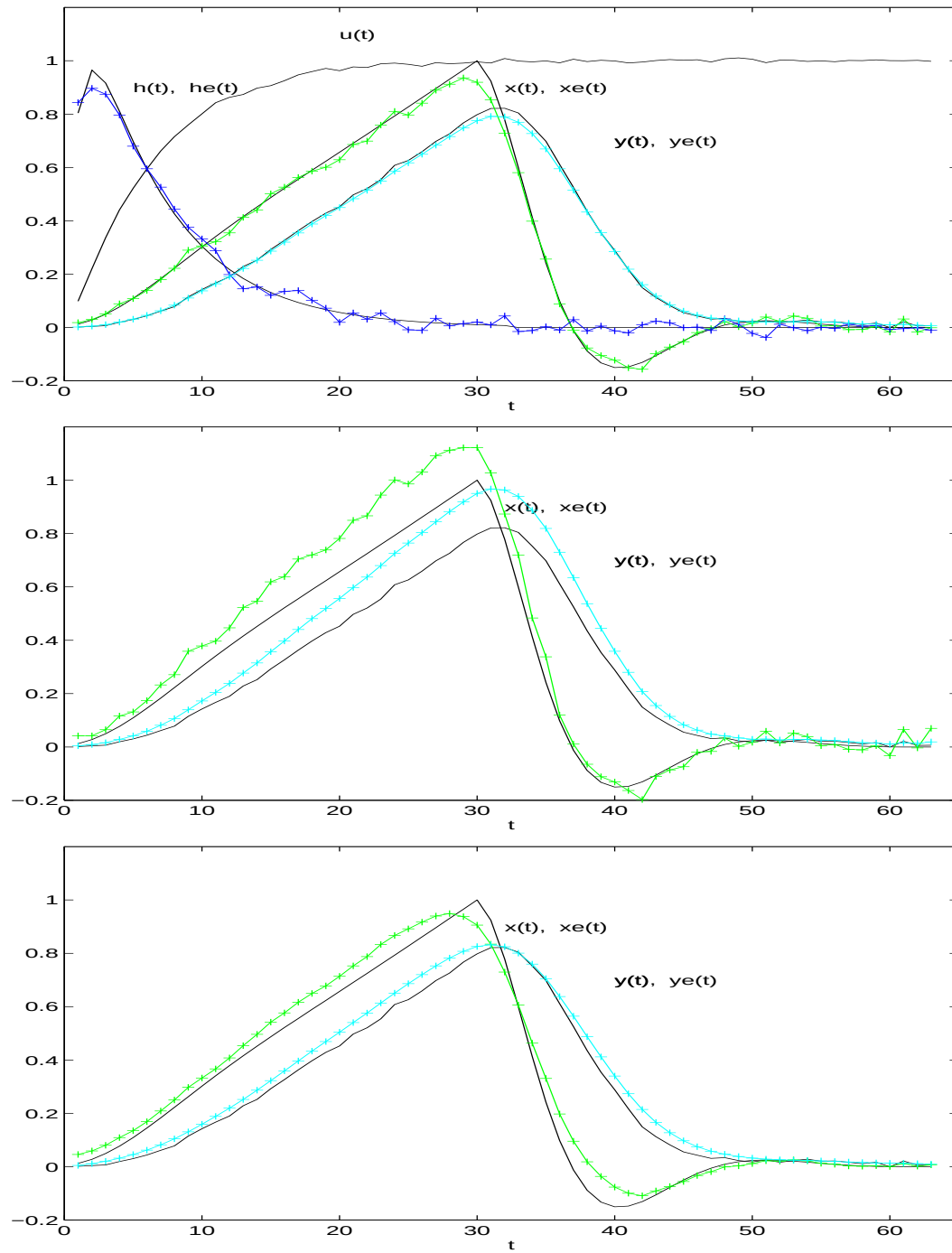


FIG. 12.6 – Exemple de résultat. On estime d'abord $h(t)$ à partir de $u(t)$ et on l'utilise pour l'estimation de $x(t)$ à partir de $y(t)$

12.1.6 Déconvolution aveugle

Estimation simultanée de la réponse impulsionnelle $h(t)$ et l'entrée $x(t)$ à partir de la sortie $y(t)$.

$$y(t) = [h * f](t) + b(t)$$

Discretisation :

$$y(m) = \sum_{k=0}^p h(k)x(m-k) + b(m) \quad \text{ou} \quad y(m) = \sum_{k=0}^M x(k)h(m-k) + b(m)$$

$$\begin{pmatrix} y(0) \\ y(1) \\ \vdots \\ y(p) \\ \vdots \\ y(M) \end{pmatrix} = \begin{pmatrix} h(0) & 0 & \cdots & \cdots & \cdots & 0 \\ h(1) & h(0) & & & & \vdots \\ \vdots & h(1) & & & & \\ h(p) & & & & & \vdots \\ 0 & h(p) & & & & \vdots \\ \vdots & & & & h(0) & 0 \\ 0 & \cdots & 0 & h(p) & \cdots & h(1) & h(0) \end{pmatrix} \begin{pmatrix} x(0) \\ x(1) \\ \vdots \\ x(M-p) \\ \vdots \\ x(M) \end{pmatrix}$$

$$\begin{pmatrix} y(0) \\ y(1) \\ \vdots \\ y(p) \\ \vdots \\ y(M) \end{pmatrix} = \begin{pmatrix} x(0) & 0 & \cdots & \cdots & \cdots & 0 \\ x(1) & x(0) & & & & \vdots \\ \vdots & & & & & \\ x(p) & & & & x(0) & \\ x(p+1) & x(p) & & & x(1) & \\ \vdots & & & & \vdots & \\ x(M) & \cdots & & \cdots & x(M-p) & \end{pmatrix} \begin{pmatrix} h(0) \\ \vdots \\ h(p) \end{pmatrix}$$

$$\mathbf{y} = \mathbf{H}\mathbf{x} + \mathbf{b} \quad \text{ou} \quad \mathbf{y} = \mathbf{X}\mathbf{h} + \mathbf{b} \quad \text{avec} \quad \mathbf{H} = \text{TOEPLITZ}(\mathbf{h}), \quad \mathbf{X} = \text{TOEPLITZ}(\mathbf{x}),$$

$$\begin{aligned} p(\mathbf{y}|\mathbf{x}, \mathbf{h}) &\propto \exp \left[-\frac{1}{2\sigma_b^2} \|\mathbf{y} - \mathbf{H}\mathbf{x}\|^2 \right] \\ p(\mathbf{x}) &\propto \exp \left[-\frac{\alpha}{2} \|\mathbf{D}_x \mathbf{x}\|^2 \right] \\ p(\mathbf{h}) &\propto \exp \left[-\frac{\beta}{2} \|\mathbf{D}_h \mathbf{h}\|^2 \right] \\ p(\mathbf{x}, \mathbf{h}|\mathbf{y}) &\propto \exp \left[-\frac{1}{2\sigma_b^2} J(\mathbf{x}, \mathbf{h}) \right] \\ J(\mathbf{x}, \mathbf{h}) &= \|\mathbf{y} - \mathbf{H}\mathbf{x}\|^2 + \lambda_1 \|\mathbf{D}_h \mathbf{h}\|^2 + \lambda_2 \|\mathbf{D}_x \mathbf{x}\|^2 \end{aligned}$$

Estimation au sens du MAP :

$$(\hat{\mathbf{x}}, \hat{\mathbf{h}}) = \arg \min_{(\mathbf{x}, \mathbf{h})} \{J(\mathbf{x}, \mathbf{h})\}$$

Mise en œuvre :

On constate que J est quadratique en $\mathbf{x}$ à $\mathbf{h}$ fixé et quadratique en $\mathbf{h}$ à $\mathbf{x}$ fixé. On peut utiliser cette constatation pour proposer un algorithme itérative :

$$\hat{\mathbf{x}}^{(k)} = \arg \min_{(\mathbf{x})} \left\{ J(\mathbf{x}, \mathbf{h}^{(k-1)}) \right\}$$

$$\hat{\mathbf{h}}^{(k)} = \arg \min_{(\mathbf{h})} \left\{ J(\mathbf{x}^{(k-1)}, \mathbf{h}) \right\}$$

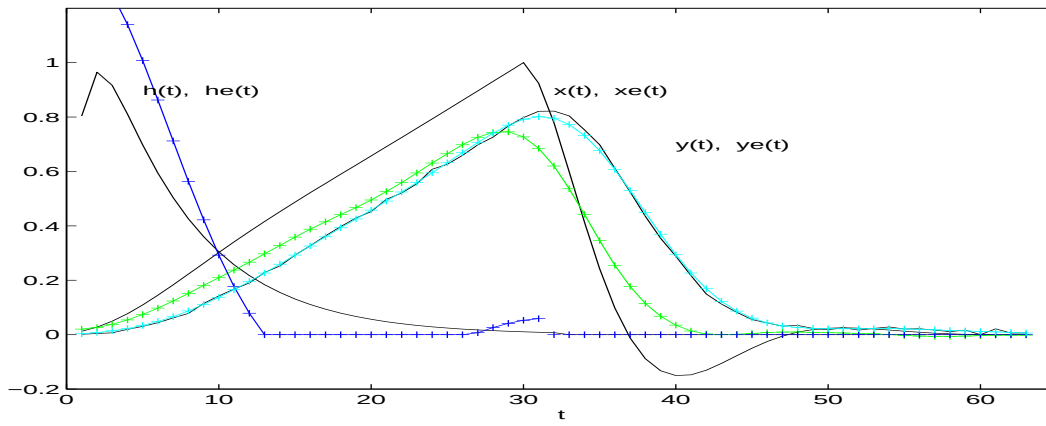
$$\hat{H}(\omega) = \frac{X^*(\omega)}{|X(\omega)|^2 + \lambda_1 D_h(\omega)} Y(\omega)$$

$$\hat{X}(\omega) = \frac{H^*(\omega)}{|H(\omega)|^2 + \lambda_2 D_x(\omega)} Y(\omega)$$

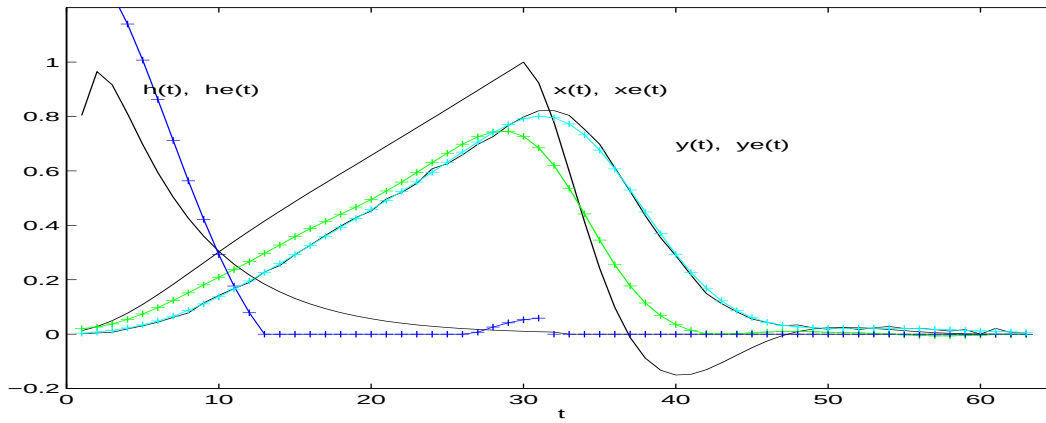
Dans le cas de la restauration d'image :

$$\hat{H}(\omega_x, \omega_y) = \frac{X^*(\omega_x, \omega_y)}{|X(\omega_x, \omega_y)|^2 + \lambda_1 D_h(\omega_x, \omega_y)} Y(\omega_x, \omega_y)$$

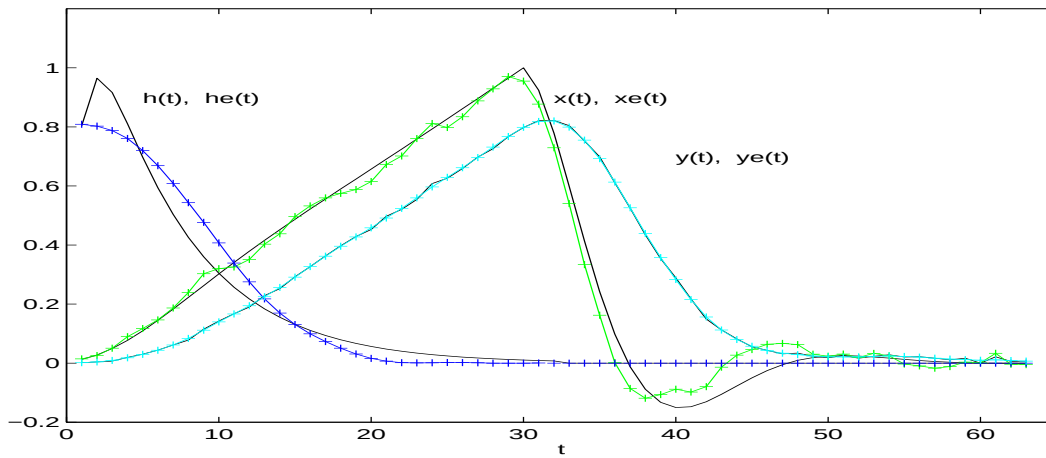
$$\hat{X}(\omega_x, \omega_y) = \frac{H^*(\omega_x, \omega_y)}{|H(\omega_x, \omega_y)|^2 + \lambda_2 D_x(\omega_x, \omega_y)} Y(\omega_x, \omega_y)$$



Résultat à l'itération 1



Résultat à l'itération 2



Résultat à l'itération finale.

FIG. 12.7 – Déconvolution aveugle :

Exemples de résultats de l'estimation simultanée de $h(t)$ et de $x(t)$ à partir de $y(t)$ avec trois jeux de paramètres différents pour $(\lambda_1, \lambda_2) = \{(20,1), (10,1), (20,05)\}$.

Écriture des relations en continu:

Considérons le problème de la déconvolution aveugle des signaux :

$$g(t) = h(t) * f(t) + b(t)$$

et définissons le critère suivant :

$$J(f, h) = \|g - h * f\|^2 + \lambda_1 \|d_1 * f\|^2 + \lambda_2 \|d_2 * h\|^2$$

où $d_1(t)$ et $d_2(t)$ sont deux fonctions connues.

$$\begin{aligned} \frac{\partial J}{\partial f} &= -2h(-t) * [g(t) - h(t) * f(t)] + 2\lambda_1 d_1(-t) * [d_1(t) * f(t)] \\ \frac{\partial J}{\partial h} &= -2f(-t) * [g(t) - h(t) * f(t)] + 2\lambda_2 d_2(-t) * [d_2(t) * h(t)] \end{aligned}$$

$$\begin{aligned} [h(-t) * h(t) + \lambda_1 d_1(-t) * d_1(t)] * f(t) &= h(-t) * g(t) \\ [f(-t) * f(t) + \lambda_2 d_2(-t) * d_2(t)] * h(t) &= f(-t) * g(t) \end{aligned}$$

Passant dans le domaine de FOURIER on obtient :

$$\begin{aligned} [|H(\omega)|^2 + \lambda_1 |D_1(\omega)|^2] F(\omega) &= H^*(\omega) G(\omega) \\ [|F(\omega)|^2 + \lambda_2 |D_2(\omega)|^2] H(\omega) &= F^*(\omega) G(\omega) \end{aligned}$$

$$\begin{aligned} F(\omega) &= \frac{H^*(\omega) G(\omega)}{|H(\omega)|^2 + \lambda_1 |D_1(\omega)|^2} \\ H(\omega) &= \frac{F^*(\omega) G(\omega)}{|F(\omega)|^2 + \lambda_2 |D_2(\omega)|^2} \end{aligned}$$

Déconvolution avec une étape de calibration :

Considérons maintenant le cas où, en plus des données, nous avons eu la possibilité d'effectuer une calibration :

$$\begin{aligned} g_c(t) &= h(t) * f_c(t) + b_1(t) \\ g(t) &= h(t) * f(t) + b_2(t) \end{aligned}$$

et définissons le critère suivant :

$$J(f, h) = \|g_c - h * f_c\|^2 + \|g - h * f\|^2 + \lambda_1 \|d_1 * f\|^2 + \lambda_2 \|d_2 * h\|^2$$

où $d_1(t)$ et $d_2(t)$ sont deux fonctions connues.

$$\begin{aligned} \frac{\partial J}{\partial f} &= -2h(-t) * [g(t) - h(t) * f(t)] + 2\lambda_1 d_1(-t) * [d_1(t) * f(t)] \\ \frac{\partial J}{\partial h} &= -2f_c(-t) * [g_c(t) - h(t) * f_c(t)] - 2f(-t) * [g(t) - h(t) * f(t)] + 2\lambda_2 d_2(-t) * [d_2(t) * h(t)] \end{aligned}$$

$$\begin{aligned} [h(-t) * h(t) + \lambda_1 d_1(-t) * d_1(t)] * f(t) &= h(-t) * g(t) \\ [f(-t) * f(t) + f_c(-t) * f_c(t) + \lambda_2 d_2(-t) * d_2(t)] * h(t) &= f_c(-t) * g_c(t) + f(-t) * g(t) \end{aligned}$$

Passant dans le domaine de FOURIER on obtient :

$$\begin{aligned} [|H(\omega)|^2 + \lambda_1 |D_1(\omega)|^2] F(\omega) &= H^*(\omega) G(\omega) \\ [|F_c(\omega)|^2 + |F(\omega)|^2 + \lambda_2 |D_2(\omega)|^2] H(\omega) &= F_c^*(\omega) G_c(\omega) + F^*(\omega) G(\omega) \end{aligned}$$

$$\begin{aligned} F(\omega) &= \frac{H^*(\omega) G(\omega)}{[|H(\omega)|^2 + \lambda_1 |D_1(\omega)|^2]} \\ H(\omega) &= \frac{F_c^*(\omega) G_c(\omega) + F^*(\omega) G(\omega)}{|F_c(\omega)|^2 + |F(\omega)|^2 + \lambda_2 |D_2(\omega)|^2} \end{aligned}$$

12.2 Restauration d'image

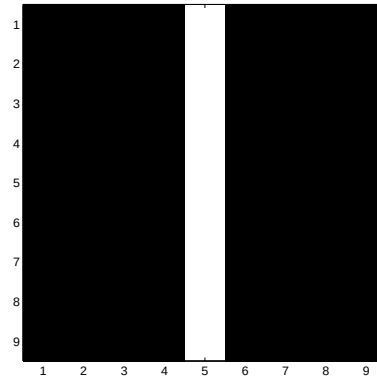
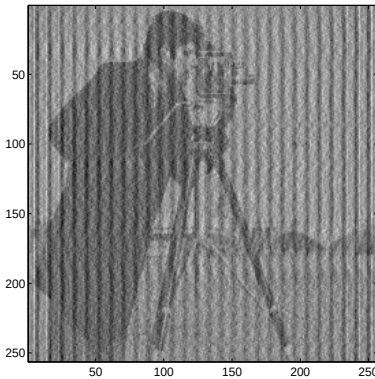
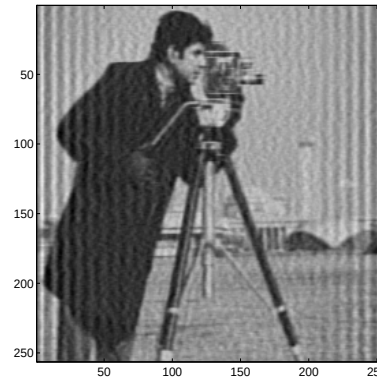
Image d'entrée $f(x,y)$ La réponse impulsionnelle $h(x,y)$ 

Image de sortie sans bruit

Image de sortie $g(x,y)$ avec bruit RSB=20 dB.

Restauration sans régularisation



Restauration avec régularisation

FIG. 12.8 – *Restauration d'image sans et avec régularisation*

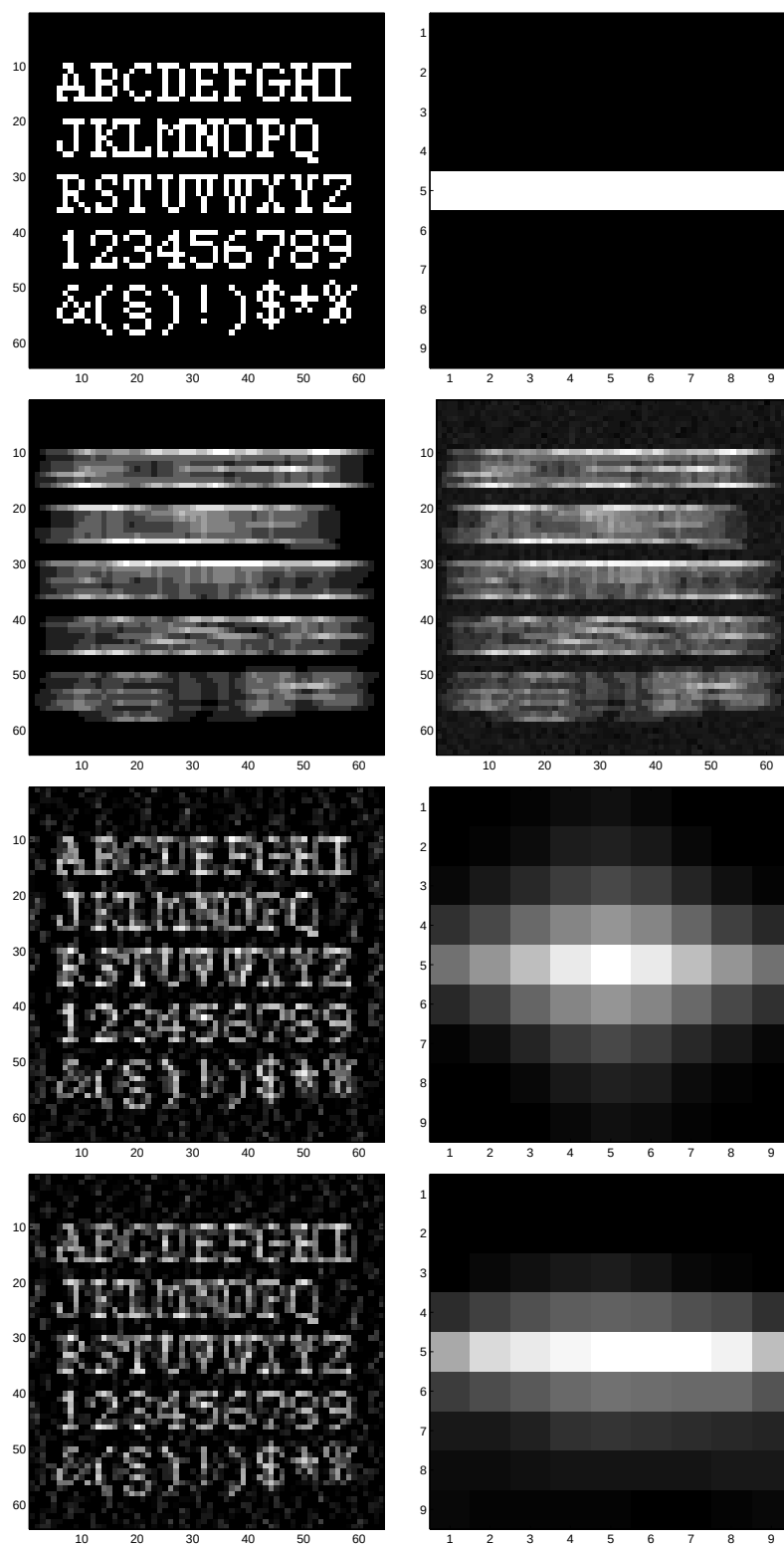
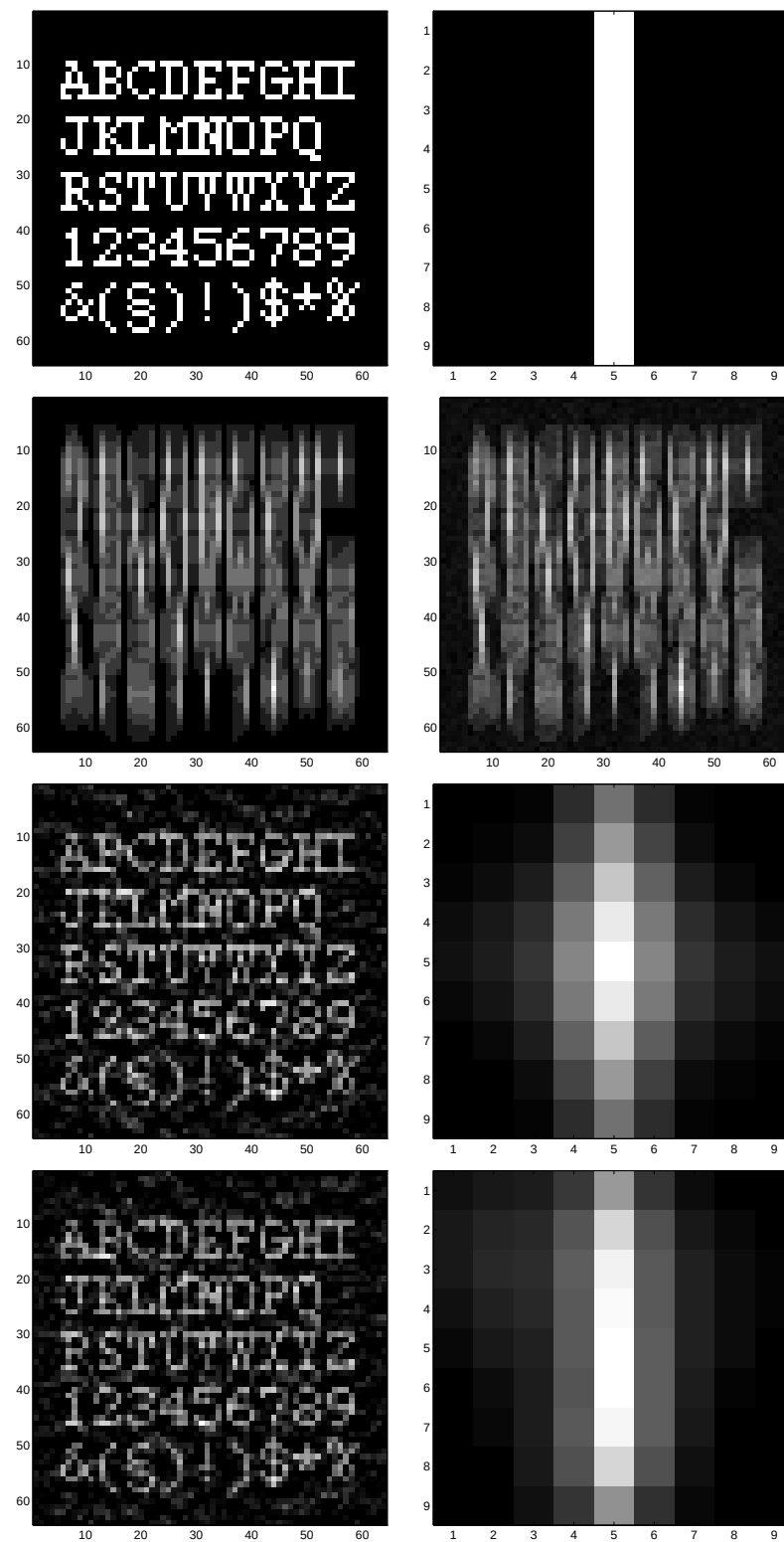


FIG. 12.9 – *Restauration aveugle d'image :*

- a) Image originale, b) réponse impulsionnelle
 c) Image dégradée sans bruit, d) Image dégradée avec bruit (données)
 e) réponse impulsionnelle à initialisation, f) Image estimée à initialisation,
 g) réponse impulsionnelle estimée, h) Image estimée

FIG. 12.10 – *Restauration aveugle d'image :*

- a) Image originale, b) réponse impulsionnelle
 c) Image dégradée sans bruit, d) Image dégradée avec bruit (données)
 e) réponse impulsionnelle à initialisation, f) Image estimée à initialisation,
 g) réponse impulsionnelle estimée, h) Image estimée

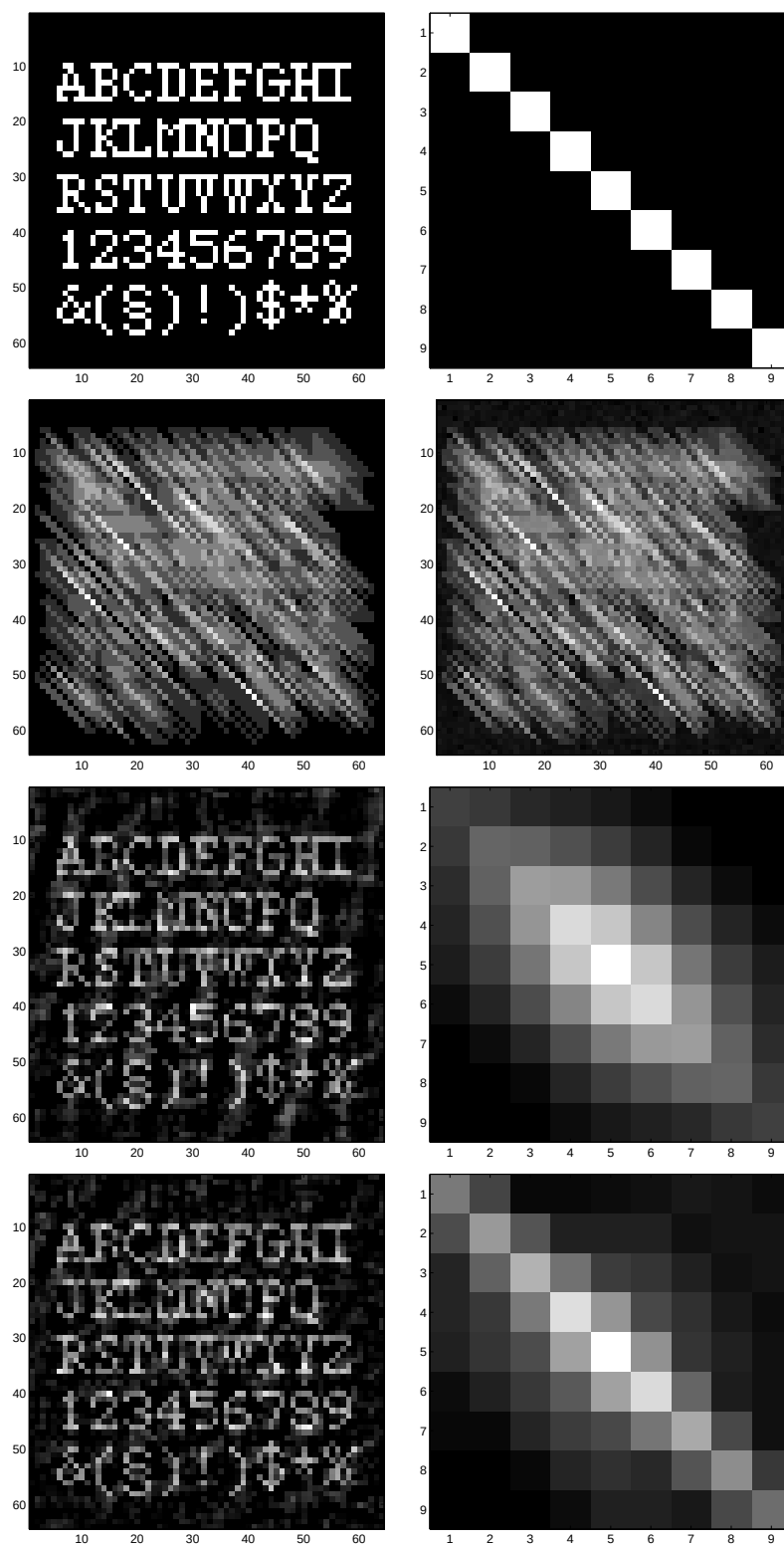
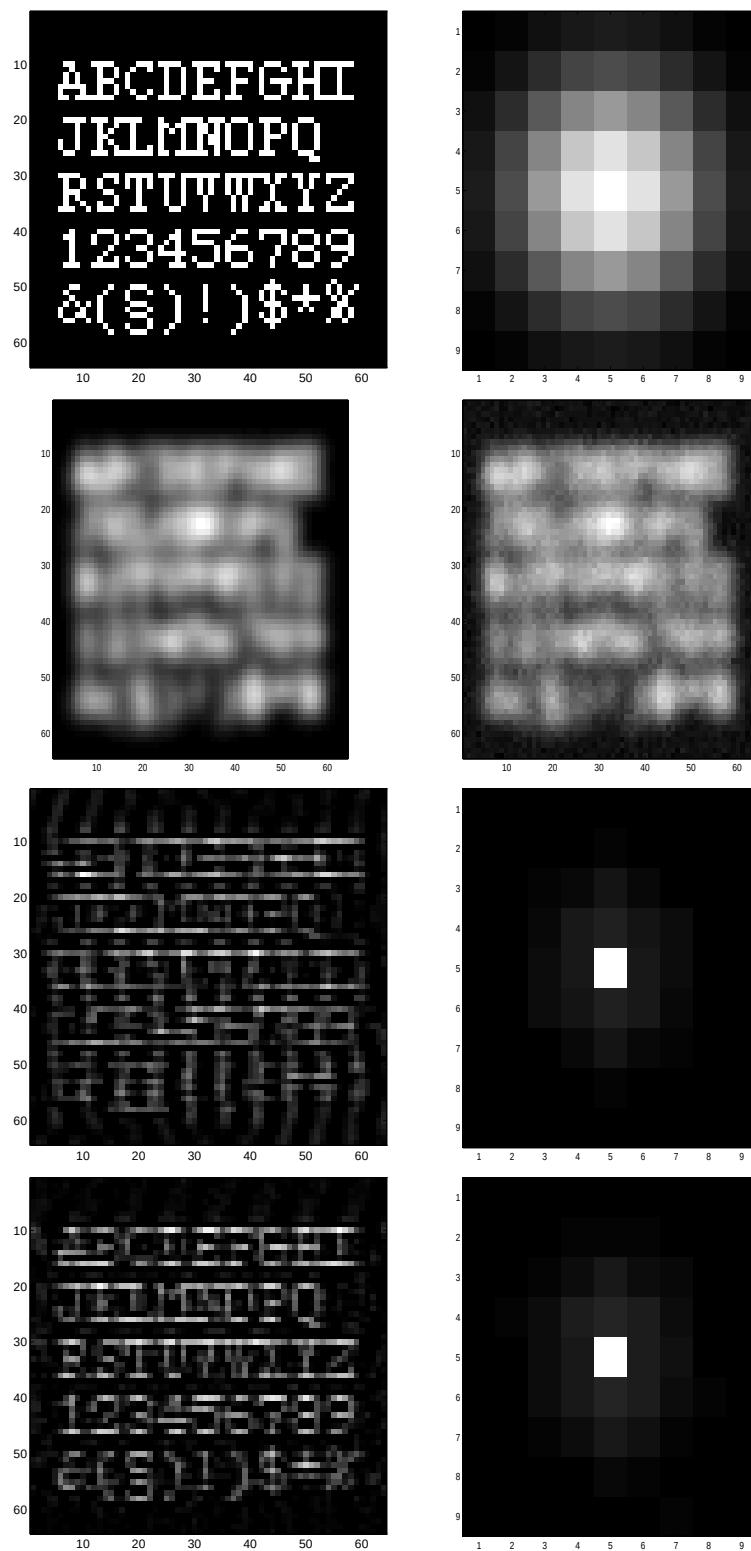


FIG. 12.11 – *Restauration aveugle d'image :*

- a) Image originale, b) réponse impulsionnelle
 c) Image dégradée sans bruit, d) Image dégradée avec bruit (données)
 e) réponse impulsionnelle à initialisation, f) Image estimée à initialisation,
 g) réponse impulsionnelle estimée, h) Image estimée

FIG. 12.12 – *Restauration aveugle d'image :*

- a) Image originale, b) réponse impulsionnelle
 c) Image dégradée sans bruit, d) Image dégradée avec bruit (données)
 e) réponse impulsionnelle à initialisation, f) Image estimée à initialisation,
 g) réponse impulsionnelle estimée, h) Image estimée

12.3 Reconstruction d'image en tomographie X

12.4 Reconstruction d'image en tomographie microondes

12.5 Reconstruction d'image en CND par courant de Foucault

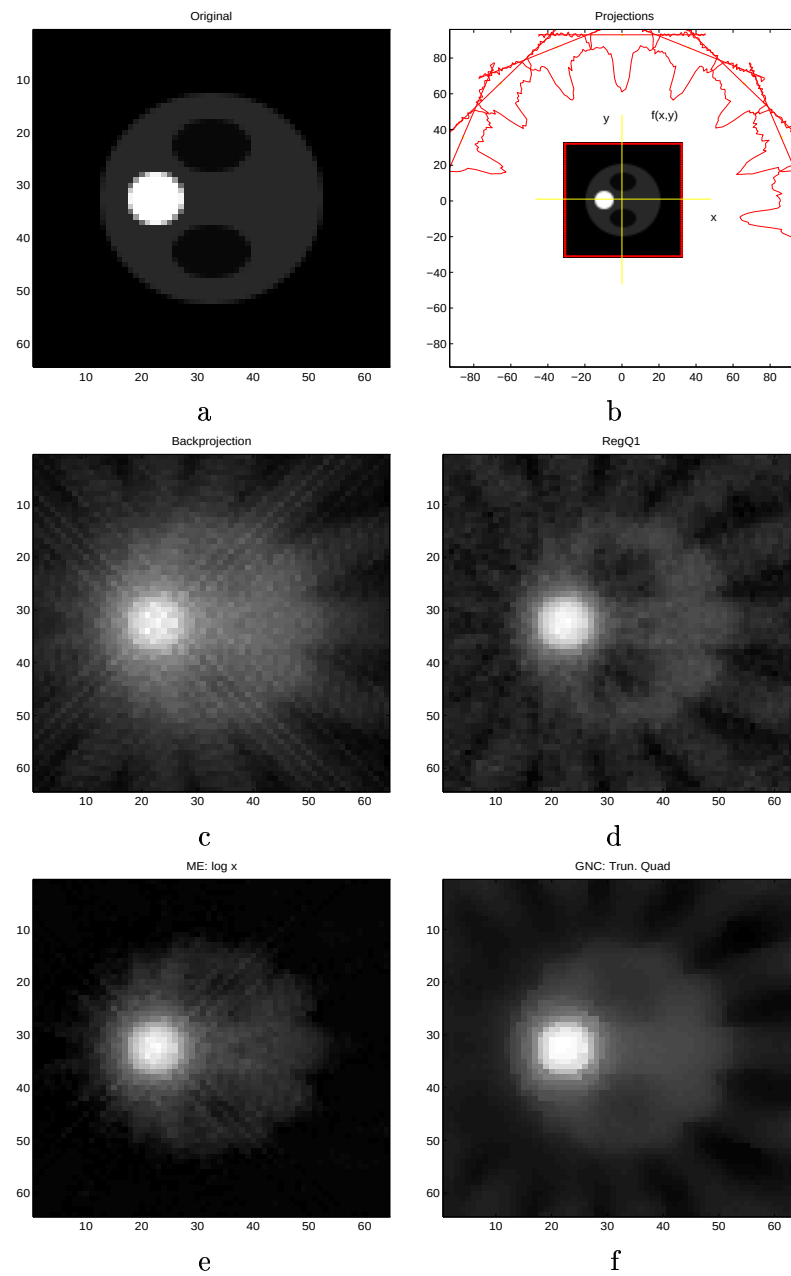


FIG. 12.13 – *Reconstruction d'image en tomographie X :*

- a) objet originale, b) projections (données),
 c) Reconstruction par rétroprojection, d) Reconstruction par régularisation quadratique,
 e) Reconstruction par ME,
 f) Reconstruction par modèle markovien et GNC,*

Annexe A

Algorithmes d'optimisation pour un critère de moindre carrés

Optimisation d'un critère et en particulier un critère de moindre carrés (MC) se trouve au cœur d'un grand nombre de problèmes. L'objet de cette annexe est de fournir une présentation des différents algorithmes d'optimisation que l'on peut utiliser pour cette tâche, surtout dans les cas d'un système non linéaire. Le cas linéaire s'en déduit, bien entendu, comme un cas particulier.

A.1 Introduction

Considérons tout d'abord le cas général de la minimisation d'un critère $J(\mathbf{x})$, c'est à dire le calcul de

$$\hat{\mathbf{x}} = \arg \min_{\mathbf{x}} \{J(\mathbf{x})\}$$

Lorsque le critère est unimodal, on peut alors utiliser les algorithmes d'optimisation classique du type gradient pour calculer la solution. Notant par :

$$\mathbf{g} = \nabla J(\mathbf{x}) = \left[\frac{\partial J}{\partial x_1}, \dots, \frac{\partial J}{\partial x_n} \right]^t \quad \text{le gradient de } J, \quad \mathbf{g}^k = \nabla J(\mathbf{x}^k)$$

et par

$$\mathbf{H} = \nabla^2 J(\mathbf{x}) = \left\{ \frac{\partial^2 J}{\partial x_i \partial x_j} \right\} \quad \text{la matrice Hessienne de } J, \quad \mathbf{H}^k = \nabla^2 J(\mathbf{x}^k)$$

sa matrice Hessienne, la majorité des méthodes d'optimisation locale calcule la solution d'une manière itérative suivante :

$$\begin{aligned} \mathbf{x}^{(k+1)} &= \mathbf{x}^{(k)} - \alpha^{(k)} \mathbf{D}^{(k)} \nabla J(\mathbf{x}^k) \\ &= \mathbf{x}^{(k)} - \alpha^{(k)} \mathbf{D}^{(k)} \mathbf{g}^k \\ &= \mathbf{x}^{(k)} + \alpha^{(k)} \mathbf{d}^{(k)} \end{aligned}$$

Suivant les différents choix pour $\alpha^{(k)}$ et $\mathbf{D}^{(k)}$ on peut distinguer les algorithmes suivants :

1. Méthodes de premier ordre ou de gradient

- gradient à pas fixe : $\mathbf{D}^{(k)} = \mathbf{I}$, $\alpha^{(k)} = \alpha$ et $\mathbf{d}^{(k)} = -\mathbf{g}^{(k)}$
- gradient à pas variable : $\mathbf{D}^{(k)} = \mathbf{I}$ et

$$\alpha^{(k)} = -\frac{\mathbf{g}^{(k)t} \mathbf{g}^{(k)}}{\mathbf{g}^{(k)t} \mathbf{H}^k \mathbf{g}^{(k)}} = \frac{\|\mathbf{g}^k\|^2}{\|\mathbf{g}^k\|_H^2}$$

- gradient conjugué :

$$\begin{aligned} \mathbf{x}^{(k+1)} &= \mathbf{x}^{(k)} + \alpha^{(k)} \mathbf{d}^{(k)} & \alpha^{(k)} &= -\frac{\mathbf{d}^{(k)t} \mathbf{g}^{(k)}}{\mathbf{d}^{(k)t} \mathbf{H}^k \mathbf{d}^{(k)}} \\ \mathbf{d}^{(k+1)} &= \mathbf{d}^{(k)} + \beta^{(k)} \mathbf{g}^{(k)} & \beta^{(k)} &= -\frac{\mathbf{g}^{(k)t} \mathbf{g}^{(k)}}{\mathbf{g}^{(k-1)t} \mathbf{g}^{(k-1)}} \end{aligned}$$

2. Méthodes de second ordre

- Méthode de Newton : $\mathbf{D}^{(k)} = [\mathbf{H}^{(k)}]^{-1}$

$$\mathbf{x}^{(k+1)} = \mathbf{x}^{(k)} - \alpha^{(k)} [\mathbf{H}^{(k)}]^{-1} \mathbf{g}^{(k)}$$

- Méthode de Newton approchée : $\mathbf{D}^{(k)} = \text{diag} \{d_1^{(k)}, \dots, d_n^{(k)}\}$ avec

$$d_i^{(k)} = D_{ii} = \left(\frac{\partial^2 J}{\partial x_i^2} \right)^{-1}$$

les éléments diagonaux de la matrice Hessienne.

– Méthodes de Newton modifiées :

$$\mathbf{D}^{(k)} = \mathbf{D}^{(0)} = [\mathbf{H}^{(0)}]^{-1} = \mathbf{D}$$

ou

$$\mathbf{D}^{(k)} = \mathbf{D}^{(0)} = \text{diag} \left\{ d_1^{(0)}, \dots, d_n^{(0)} \right\} \quad \text{avec} \quad d_i^{(0)} = \left(\frac{\partial^2 J}{\partial x_i^2}(\mathbf{x})^{(0)} \right)^{-1}$$

A.2 Cas des critères moindres carrés

Considérons maintenant le cas des critères de la forme

$$J(\mathbf{x}) = \frac{1}{2} \|\mathbf{f}(\mathbf{x})\|^2 = \frac{1}{2} \sum_{i=1}^M |f_i(\mathbf{x})|^2$$

On a alors

$$\frac{\partial J(\mathbf{x})}{\partial x_n} = \sum_{i=1}^M \frac{\partial f_i}{\partial x_n} f_i$$

et

$$\frac{\partial^2 J(\mathbf{x})}{\partial x_n \partial x_m} = \sum_{i=1}^M \frac{\partial^2 f_i}{\partial x_n \partial x_m} f_i + \frac{\partial f_i}{\partial x_n} \frac{\partial f_i}{\partial x_m}$$

Si on note par

$$\mathbf{F} = \left[\frac{\partial \mathbf{f}}{\partial \mathbf{x}} \right] = \begin{bmatrix} \nabla f_1 & \dots & \nabla f_M \end{bmatrix} = \begin{pmatrix} \frac{\partial f_1}{\partial x_1} & \frac{\partial f_2}{\partial x_1} & \dots & \frac{\partial f_M}{\partial x_1} \\ \frac{\partial f_1}{\partial x_2} & \frac{\partial f_2}{\partial x_2} & & \frac{\partial f_M}{\partial x_2} \\ \vdots & & & \vdots \\ \frac{\partial f_1}{\partial x_N} & \frac{\partial f_2}{\partial x_N} & & \frac{\partial f_M}{\partial x_N} \end{pmatrix}$$

on a

$$\nabla^2 J(\mathbf{x}) = \mathbf{F}^t \mathbf{F} + \sum_{i=1}^M \frac{\partial^2 g_i}{\partial x_n \partial x_m} g_i$$

et on peut alors distinguer les algorithmes suivants :

- Méthode de Newton : $\mathbf{D}^{(k)} = [\nabla^2 J(\mathbf{x}^{(k)})]^{-1}$
- Méthode de Gauss-Newton : $\mathbf{D}^{(k)} = (\mathbf{F}^t \mathbf{F})^{-1}$
- Méthodes de Levenberg-Marquart :

$$\mathbf{D}^{(k)} = (\mathbf{F}^t \mathbf{F} + \lambda^{(k)} \mathbf{I})^{-1}$$

avec $\lambda^{(k)}$ une suite décroissante de nombres positifs.

- pour λ grand : gradient
- pour λ petit : Newton

– Méthodes de Quasi-Newton :

$$\mathbf{D}^{(k+1)} = \mathbf{D}^{(k)} + \frac{\mathbf{p}^{(k)} \mathbf{p}^{t(k)}}{\mathbf{p}^{t(k)} \mathbf{Q}^{(k)}} - \frac{\mathbf{D}^{(k)} \mathbf{Q}^{(k)} \mathbf{Q}^{t(k)} \mathbf{D}^{(k)}}{\mathbf{Q}^{t(k)} \mathbf{D}^{(k)} \mathbf{Q}^{(k)}} + \eta^{(k)} \tau^{(k)} \mathbf{v}^{(k)} \mathbf{v}^{t(k)}$$

$$\begin{aligned} \mathbf{p}^{(k)} &= \mathbf{x}^{(k+1)} - \mathbf{x}^{(k)} \\ \mathbf{Q}^{(k)} &= \nabla J(\mathbf{x}^{(k+1)}) - \nabla J(\mathbf{x}^{(k)}) \\ \mathbf{v}^{(k)} &= \frac{\mathbf{p}^{(k)}}{\mathbf{p}^{t(k)} \mathbf{Q}^{(k)}} - \frac{\mathbf{D}^{(k)} \mathbf{Q}^{(k)}}{\tau^{(k)}} \\ \tau^{(k)} &= \mathbf{Q}^{t(k)} \mathbf{D}^{(k)} \mathbf{Q}^{(k)} \end{aligned}$$

- Méthode de Davidon-Fletcher-Powell (DFP) : $\eta^{(k)} = 0$
- Méthode de Broyden-Fletcher-Golfand-Shayno (BGFGS) : $\eta^{(k)} = 1$

A.3 Cas des problèmes inverses avec un bruit additif

Considérons maintenant le modèle :

$$y_i = h_i(\mathbf{x}) + b_i, \quad i = 1, \dots, M$$

ou sous sa forme vectorielle :

$$\mathbf{y} = \mathbf{h}(\mathbf{x}) + \mathbf{b},$$

où $\mathbf{y}$ sont les données et $\mathbf{x}$ les inconnues.

Le critère des moindres carrés s'écrit :

$$J(\mathbf{x}) = \|\mathbf{y} - \mathbf{h}(\mathbf{x})\|^2 = [\mathbf{y} - \mathbf{h}(\mathbf{x})]^t [\mathbf{y} - \mathbf{h}(\mathbf{x})] = \sum_{i=1}^M (y_i - h_i(\mathbf{x}))^2$$

Lorsque le système est non linéaire, ce critère n'est pas quadratique et, en général, peut être multimodal. Cependant l'usage des méthodes du type gradient est courant pour calculer une solution correspondant à un minimum local de ce critère. L'objet de cette annexe est de fournir une description succincte de ces méthodes.

Notons le gradient du critère par :

$$\nabla J(\mathbf{x}) = \left[\frac{\partial J}{\partial x_j} \right] = \left[\frac{\partial J}{\partial x_1}, \dots, \frac{\partial J}{\partial x_N} \right]^t$$

et la matrice Hessien du critère par :

$$\nabla^2 J(\mathbf{x}) = \left[\frac{\partial^2 J}{\partial x_k \partial x_l} \right] = \begin{pmatrix} \frac{\partial^2 J}{\partial x_1^2} & \frac{\partial^2 J}{\partial x_1 \partial x_2} & \dots & \frac{\partial^2 J}{\partial x_1 \partial x_N} \\ \frac{\partial^2 J}{\partial x_2 \partial x_1} & & & \\ \vdots & & & \\ \frac{\partial^2 J}{\partial x_N \partial x_1} & \frac{\partial^2 J}{\partial x_N \partial x_{N-1}} & \frac{\partial^2 J}{\partial x_N^2} \end{pmatrix}$$

Notons aussi :

$$\nabla h_i(\mathbf{x}) = \left[\frac{\partial h_i}{\partial x_j} \right] = \left[\frac{\partial h_i}{\partial x_1}, \dots, \frac{\partial h_i}{\partial x_N} \right]^t,$$

$$\nabla^2 h_i = \left[\frac{\partial^2 h_i}{\partial x_k \partial x_l} \right] = \begin{pmatrix} \frac{\partial^2 h_i}{\partial x_1^2} & \frac{\partial^2 h_i}{\partial x_1 \partial x_2} & \dots & \frac{\partial^2 h_i}{\partial x_1 \partial x_N} \\ \frac{\partial^2 h_i}{\partial x_2 \partial x_1} & & & \\ \vdots & & & \\ \frac{\partial^2 h_i}{\partial x_N \partial x_1} & \frac{\partial^2 h_i}{\partial x_N \partial x_{N-1}} & \frac{\partial^2 h_i}{\partial x_N^2} \end{pmatrix}$$

et

$$\mathbf{H} = \left[\frac{\partial \mathbf{h}}{\partial \mathbf{x}} \right] = \left[\nabla h_1 : \dots : \nabla h_M \right] = \begin{pmatrix} \frac{\partial h_1}{\partial x_1} & \frac{\partial h_2}{\partial x_1} & \dots & \frac{\partial h_M}{\partial x_1} \\ \frac{\partial h_1}{\partial x_2} & & & \frac{\partial h_M}{\partial x_2} \\ \vdots & & & \\ \frac{\partial h_1}{\partial x_N} & \frac{\partial h_2}{\partial x_N} & & \frac{\partial h_M}{\partial x_N} \end{pmatrix}$$

On a alors

$$\begin{aligned}
 \nabla J(\mathbf{x}) &= -2 \left[\frac{\partial \mathbf{h}}{\partial \mathbf{x}} \right]^t [\mathbf{y} - \mathbf{h}(\mathbf{x})] \\
 &= -2 [\mathbf{H}(\mathbf{x})]^t [\mathbf{y} - \mathbf{h}(\mathbf{x})] \\
 \nabla^2 J(\mathbf{x}) &= 2 \left[\frac{\partial \mathbf{h}}{\partial \mathbf{x}} \right]^t \left[\frac{\partial \mathbf{h}}{\partial \mathbf{x}} \right] + 2 \left[\sum_i [y_i - h_i(\mathbf{x})] \nabla^2 h_i(\mathbf{x}) \right] \\
 &= 2 [\mathbf{H}(\mathbf{x})]^t [\mathbf{H}(\mathbf{x})] + 2 \left[\sum_i [y_i - h_i(\mathbf{x})] \nabla^2 h_i(\mathbf{x}) \right]
 \end{aligned}$$

Notons que pour un système linéaire on a

$$y_i = h_i(\mathbf{x}) = \sum_{j=1}^N h_{ij} x_j \implies \mathbf{y} = \mathbf{A}\mathbf{x} + \mathbf{b},$$

avec

$$\begin{aligned}
 \mathbf{A} &= \begin{pmatrix} h_{11} & h_{12} & \cdots & h_{1N} \\ h_{21} & \cdots & \cdots & \vdots \\ \vdots & & & \vdots \\ h_{M1} & \cdots & \cdots & h_{MN} \end{pmatrix} \\
 J(\mathbf{x}) &= \|\mathbf{y} - \mathbf{A}\mathbf{x}\|^2 = [\mathbf{y} - \mathbf{A}\mathbf{x}]^t [\mathbf{y} - \mathbf{A}\mathbf{x}] \\
 \mathbf{H} &= \mathbf{A} \\
 \nabla^2 h_i &= \left[\frac{\partial^2 h_i}{\partial \mathbf{x}^2} \right] = 0 \\
 \nabla J &= -2\mathbf{A}^t [\mathbf{y} - \mathbf{A}\mathbf{x}] \\
 \nabla^2 J &= 2\mathbf{A}^t \mathbf{A}
 \end{aligned}$$

car $\frac{\partial^2 h_i}{\partial x_k \partial x_l} = 0, \forall k, l$.

Développement à l'ordre deux, en série de Taylor de $J(\mathbf{x})$ autour d'un point $\mathbf{x}^*$ donne

$$J(\mathbf{x}) = J(\mathbf{x}^*) + [\nabla J(\mathbf{x}^*)]^t [\mathbf{x} - \mathbf{x}^*] + \frac{1}{2} [\mathbf{x} - \mathbf{x}^*]^t [\nabla^2 J(\mathbf{x}^*)] [\mathbf{x} - \mathbf{x}^*]$$

Conditions d'existence d'une solution :

$$\begin{cases} \nabla J(\mathbf{x}^*) = \mathbf{0} \\ \nabla^2 J(\mathbf{x}^*) \geq 0 \quad (\text{C.N. dnn}) \\ \nabla^2 J(\mathbf{x}^*) > 0 \quad (\text{C.S. dp}) \end{cases}$$

La première condition d'existence d'une solution donne lieu au système d'équations non linéaires :

$$\nabla J(\mathbf{x}^*) = \mathbf{0} \implies [\mathbf{H}(\mathbf{x}^*)]^t [\mathbf{y} - \mathbf{h}(\mathbf{x}^*)] = \mathbf{0}$$

Méthode de Gauss :

Développement à l'ordre 1 de $\mathbf{y}$ donne :

$$\mathbf{y} = \mathbf{h}(\mathbf{x}^*) + \left[\frac{\partial \mathbf{h}(\mathbf{x}^*)}{\partial \mathbf{x}} \right]^t [\mathbf{x} - \mathbf{x}^*] + \mathbf{b} = \mathbf{h}(\mathbf{x}^*) + \mathbf{H}^t [\mathbf{x} - \mathbf{x}^*] + \mathbf{b}$$

$$\nabla J(\mathbf{x}) = \mathbf{0} \implies \mathbf{x} = \mathbf{x}^* + [\mathbf{H}^t \mathbf{H}]^{-1} \mathbf{H}^t [\mathbf{y} - \mathbf{h}(\mathbf{x}^*)]$$

Algorithme :

$$\mathbf{x}^{(k)} = \mathbf{x}^{(k-1)} + \alpha [\mathbf{H}^t \mathbf{H}]^{-1} \mathbf{H}^t [\mathbf{y} - \mathbf{h}(\mathbf{x}^{(k)})]$$

Méthode de Newton–Raphson :

Développement à l'ordre 2 de $\mathbf{y}$ donne :

$$\nabla J(\mathbf{x}) = \mathbf{0} \implies \mathbf{x} = \mathbf{x}^* - [\nabla^2 J(\mathbf{x}^*)]^{-1} [\nabla J(\mathbf{x}^*)]$$

Algorithme :

$$\mathbf{x}^{(k)} = \mathbf{x}^{(k-1)} + \alpha \left[\mathbf{H}^t \mathbf{H} + \sum_i [y_i - h_i(\mathbf{x}^{(k)})] \frac{\partial^2 h_i}{\partial \mathbf{x}^2}(\mathbf{x}^{(k)}) \right]^{-1} \mathbf{H}^t [\mathbf{y} - \mathbf{h}(\mathbf{x}^{(k)})]$$

A.4 Comparaison entre les méthodes

Notons

$$\mathbf{H}_k = [\mathbf{H}(\mathbf{x}^{(k)})], \quad \text{et} \quad \mathbf{M}_k = [\mathbf{M}(\mathbf{x}^{(k)})].$$

On a alors dans les deux cas :

$$\mathbf{x}^{(k)} = \mathbf{x}^{(k-1)} + \alpha \mathbf{M}_k^{-1} \mathbf{H}_k^t [\mathbf{y} - \mathbf{h}(\mathbf{x}^{(k)})]$$

– Méthode de Newton–Raphson :

$$\mathbf{M}_k = \left[\mathbf{H}_k^t \mathbf{H}_k + \sum_i [y_i - h_i(\mathbf{x}^{(k)})] \nabla^2 h_i(\mathbf{x}^{(k)}) \right]$$

– Méthode de Gauss :

$$\mathbf{M}_k = [\mathbf{H}_k^t \mathbf{H}_k]$$

– Méthode du gradient conjugué :

$$\mathbf{M}_k = [\mathbf{H}_k^t \mathbf{H}_k], \quad \mathbf{H}_k = \left[\frac{\partial \mathbf{h}}{\partial \mathbf{x}}(\mathbf{x}^{(k)}) \right] \quad \text{est supposée Toeplitz}$$

– Méthode du gradient :

$$\mathbf{M}_k = \mathbf{I}$$

Dans le cas d'un système linéaire on a

$$\mathbf{H}_k = \mathbf{H} = \mathbf{A}, \quad \mathbf{M}_k = \mathbf{M} = \mathbf{H}^t \mathbf{H}.$$

Ceci vaut dire qu'il y a l'équivalence entre la méthode de Newton–Raphson et la méthode de Gauss.

Si le système linéaire et invariant par translation on a

$$\mathbf{H}_k = \mathbf{H} = \mathbf{A} \quad (\text{Toeplitz}), \quad \mathbf{M}_k = \mathbf{M} = \mathbf{H}^t \mathbf{H} \quad (\text{Toeplitz symétrique})$$

Ceci vaut dire qu'il y a l'équivalence entre la méthode de Newton–Raphson, la méthode de Gauss et la méthode du gradient conjugué.

Annexe B

Algorithme de Gauss-Seidel pour optimisation d'un critère quadratique

Lorsqu'on cherche à optimiser un critère d'une manière itérative, il existe deux approches:

- les méthodes qui mettent à jour l'ensemble des variables à chaque itérations, comme par exemples des méthodes du type gradient ; et
- les méthodes qui mettent à jour seulement une variable à chaque itérations.

L'objet de cette annexe est étudier la méthode de Gauss-Seidel qui est une méthode de la deuxième approche pour le cas du critère MC ou celui d'un critère régularisé.

B.1 Introduction

Considérons le critère

$$J(\mathbf{x}) = \|\mathbf{y} - \mathbf{Ax}\|^2$$

et son gradient :

$$\nabla J = -2\mathbf{A}^t(\mathbf{y} - \mathbf{Ax})$$

Rappelons le principe d'un algorithme du gradient pour calculer la solution :

$$\begin{aligned} \mathbf{x}^{(0)} &= \mathbf{0} \\ \mathbf{x}^{(k+1)} &= \mathbf{x}^{(k)} + \alpha^{(k)} \nabla J(\mathbf{x}^{(k)}) \\ &= \mathbf{x}^{(k)} + \alpha^{(k)} \mathbf{A}^t(\mathbf{y} - \mathbf{Ax}^{(k)}) = \mathbf{x}^{(k)} + \alpha^{(k)} \mathbf{A}^t \mathbf{y} - \alpha^{(k)} \mathbf{A}^t \mathbf{Ax}^{(k)} \end{aligned}$$

Supposons maintenant, qu'à l'itération k , on voudrait remettre à jour seulement la variable x_k et écrivons alors le critère en fonction de cette seule variable :

$$\begin{aligned} J(x_k) &= \|\mathbf{y} - \mathbf{Ax}\|^2 \\ &= \sum_{i=1}^M \left(y_i - \sum_{j=1}^N a_{ij} x_j \right)^2 \\ &= \sum_{i=1}^M \left(y_i - \sum_{j \neq k}^N a_{ij} x_j - a_{ik} x_k \right)^2 \\ &= \sum_{i=1}^M \left[\left(y_i - \sum_{j \neq k}^N a_{ij} x_j \right)^2 - 2a_{ik} \left(y_i - \sum_{j \neq k}^N a_{ij} x_j \right) x_k + a_{ik}^2 x_k^2 \right] \\ \frac{\partial J}{\partial x_k} &= -2 \sum_{i=1}^M a_{ik} \left(y_i - \sum_{j=1}^N a_{ij} x_j \right) \\ &= -2 \sum_{i=1}^M a_{ik} \left(y_i - \sum_{j \neq k}^N a_{ij} x_j - a_{ik} x_k \right) \\ &= -2 \sum_{i=1}^M \left[a_{ik} \left(y_i - \sum_{j \neq k}^N a_{ij} x_j \right) - a_{ik}^2 x_k \right] \end{aligned}$$

Le critère $J(x_k)$ est quadratique en x_k et sa dérivée est linéaire. On peut alors calculer explicitement la solution en annulant la dérivée :

$$\begin{aligned} \frac{\partial J}{\partial x_k} = 0 \longrightarrow x_k &= \frac{1}{\sum_{i=1}^M a_{ik}^2} \sum_{i=1}^M a_{ik} \left(y_i - \sum_{j \neq k}^N a_{ij} x_j \right) \\ &= \frac{1}{\sum_{i=1}^M a_{ik}^2} \sum_{i=1}^M a_{ik} \left(y_i - \sum_{j=1}^N a_{ij} x_j + a_{ik} x_k \right) \\ &= x_k + \frac{1}{\sum_{i=1}^M a_{ik}^2} \sum_{i=1}^M a_{ik} \left(y_i - \sum_{j=1}^N a_{ij} x_j \right) \end{aligned}$$

$$\begin{aligned}
&= x_k + \frac{1}{\langle \mathbf{a}_{*k}, \mathbf{a}_{*k} \rangle} \sum_{i=1}^M a_{ik} \left(y_i - [\mathbf{A}\mathbf{x}^{(k)}]_i \right) \\
&= x_k + \frac{1}{\langle \mathbf{a}_{*k}, \mathbf{a}_{*k} \rangle} \langle \mathbf{a}_{*k}, (\mathbf{y} - \mathbf{A}\mathbf{x}^{(k)}) \rangle
\end{aligned}$$

Pour simplifier ces écritures, notons par

$$\mathbf{x} = \mathbf{x}_{(-k)} + \mathbf{x}_{(k)}$$

avec

$$\mathbf{x}_{(-k)} = [x_1, \dots, x_{k-1}, 0, x_{k+1}, \dots, x_N]^t, \quad \text{et} \quad \mathbf{x}_{(k)} = [0, \dots, 0, x_k, 0, \dots, 0]^t$$

On a alors

$$J(x_k) = ax_k^2 + bx_k + c$$

avec

$$\begin{aligned}
a &= \sum_i a_{ik}^2 = \|\mathbf{a}_{*k}\|^2 \\
b &= -2 \sum_i a_{ik} \left(y_i - \sum_{j \neq k}^N a_{ij} x_j \right) = -2 \langle \mathbf{a}_{*k}, (\mathbf{y} - \mathbf{A}\mathbf{x}_{(-k)}) \rangle \\
c &= \sum_i \left(y_i - \sum_{j \neq k}^N a_{ij} x_j \right)^2 = \|\mathbf{y} - \mathbf{A}\mathbf{x}_{(-k)}\|^2
\end{aligned}$$

L'algorithme de Gausse-Seidel est fondée sur cette relation et donné par :

$$\begin{aligned}
\mathbf{x}^{(0)} &= \mathbf{0} \\
x_k^{(k+1)} &= x_k^{(k)} + \alpha^{(k)} \frac{\partial J}{\partial x_k} (\mathbf{x}^{(k)}) \\
&= x_k^{(k)} + \alpha^{(k)} [\mathbf{A}^t (\mathbf{y} - \mathbf{A}\mathbf{x}^{(k)})]_k \\
&= x_k^{(k)} + \alpha^{(k)} \sum_{i=1}^M a_{ik} (y_i - [\mathbf{A}\mathbf{x}^{(k)}]_i) \quad \text{avec} \quad \alpha^{(k)} = \frac{1}{\langle \mathbf{a}_{*k}, \mathbf{a}_{*k} \rangle}.
\end{aligned}$$

Rappelons à titre de comparaison la Méthode de Kaczmarz que nous l'avons vue dans le chapitre concernant l'inversion généralisée :

$$\begin{aligned}
\mathbf{x}^{(0)} &= \mathbf{0} \\
\mathbf{x}^{(k+1)} &= \mathbf{x}^{(k)} + \frac{(y_i - [\mathbf{A}\mathbf{x}^{(k)}]_i)}{\langle \mathbf{a}_{i*}, \mathbf{a}_{i*} \rangle} \mathbf{a}_{i*} \\
&= \mathbf{x}^{(k)} + \frac{(y_i - \langle \mathbf{a}_{i*}, \mathbf{x}^{(k)} \rangle)}{\langle \mathbf{a}_{i*}, \mathbf{a}_{i*} \rangle} \mathbf{a}_{i*}, \quad i = 1, \dots, M, 1 \dots, M, 1, \dots
\end{aligned}$$

B.2 Le cas d'un critère régularisé

Le critère et son gradient dans ce cas sont :

$$\begin{aligned} J(\mathbf{x}) &= \|\mathbf{y} - \mathbf{Ax}\|^2 + \lambda \|\mathbf{Cx}\|^2, \\ \nabla J &= -2\mathbf{A}^t(\mathbf{y} - \mathbf{Ax}) + 2\lambda \mathbf{C}^t \mathbf{Cx} = -2 \left[\mathbf{A}^t \mathbf{y} - (\mathbf{A}^t \mathbf{A} + 2\lambda \mathbf{C}^t \mathbf{C}) \mathbf{x} \right] \end{aligned}$$

La même démarche nous conduit :

$$\begin{aligned} J(x_k) &= \|\mathbf{y} - \mathbf{Ax}\|^2 + \lambda \|\mathbf{Cx}\|^2 \\ &= \sum_{i=1}^M \left(y_i - \sum_{j=1}^N a_{ij} x_j \right)^2 + \sum_{j=1}^N \left(\sum_{n=1}^N c_{jn} x_n \right)^2 \\ &= \sum_{i=1}^M \left[\left(y_i - \sum_{j \neq k}^N a_{ij} x_j \right)^2 - 2a_{ik} \left(y_i - \sum_{j \neq k}^N a_{ij} x_j \right) x_k + a_{ik}^2 x_k^2 \right] \\ &\quad + \lambda \sum_{j \neq k}^N \left(\sum_{n \neq k}^N c_{jn} x_n + c_{jk} x_k \right)^2 + \lambda \sum_{n \neq k}^N (c_{kn} x_n + c_{kk} x_k)^2 \\ \frac{\partial J}{\partial x_k} &= -2 \sum_{i=1}^M a_{ik} \left(y_i - \sum_{j=1}^N a_{ij} x_j \right) + 2\lambda \sum_{j=1}^N \sum_{n=1}^N c_{jk} c_{jn} x_n \\ &= -2 \sum_{i=1}^M a_{ik} \left(y_i - \sum_{j \neq k}^N a_{ij} x_j - a_{ik} x_k \right) + 2\lambda \sum_{n=1}^N x_n \sum_{j=1}^N c_{jk} c_{jn} \\ &= -2 \sum_{i=1}^M \left[a_{ik} \left(y_i - \sum_{j \neq k}^N a_{ij} x_j \right) - a_{ik}^2 x_k \right] + 2\lambda \sum_{n \neq k}^N x_n \sum_{j=1}^N c_{jk} c_{jn} + 2\lambda x_k \sum_{j=1}^N c_{jk}^2 \end{aligned}$$

Solution explicite:

$$\nabla J = \mathbf{0} \longrightarrow (\mathbf{A}^t \mathbf{A} + \lambda \mathbf{C}^t \mathbf{C}) \mathbf{x} = \mathbf{A}^t \mathbf{y} \longrightarrow \mathbf{x} = (\mathbf{A}^t \mathbf{A} + \lambda \mathbf{C}^t \mathbf{C})^{-1} \mathbf{A}^t \mathbf{y}$$

Méthode du gradient :

$$\begin{aligned} \mathbf{x}^{(0)} &= \mathbf{0} \\ \mathbf{x}^{(k+1)} &= \mathbf{x}^{(k)} + \alpha^{(k)} \nabla J(\mathbf{x}^{(k)}) \\ &= \mathbf{x}^{(k)} + \alpha^{(k)} \left[\mathbf{A}^t (\mathbf{y} - \mathbf{Ax}^{(k)}) - \lambda \mathbf{C}^t \mathbf{Cx}^{(k)} \right] \\ &= \mathbf{x}^{(k)} + \alpha^{(k)} \mathbf{A}^t \mathbf{y} - \alpha^{(k)} \left(\mathbf{A}^t \mathbf{A} + \frac{\lambda}{\alpha} \mathbf{C}^t \mathbf{C} \right) \mathbf{x}^{(k)} \end{aligned}$$

Algorithme du Gauss-Seidel :

$$\begin{aligned} \frac{\partial J}{\partial x_k} &= 0 \longrightarrow \\ x_k &= \frac{1}{\sum_{i=1}^M a_{ik}^2 + \lambda \sum_{j=1}^N c_{jk}^2} \sum_{i=1}^M a_{ik} \left(y_i - \sum_{j \neq k}^N a_{ij} x_j \right) - \lambda \sum_{n \neq k}^N x_n \sum_{j=1}^N c_{jk} c_{jn} \end{aligned}$$

$$\begin{aligned}
&= \frac{1}{\sum_{i=1}^M a_{ik}^2 + \lambda \sum_{j=1}^N c_{jk}^2} \sum_{i=1}^M a_{ik} \left(y_i - \sum_{j=1}^N a_{ij} x_j + a_{ik} x_k \right) - \lambda \sum_{n=1}^N x_n \sum_{j=1}^N c_{jk} c_{jn} + \lambda x_k \sum_{j=1}^N c_{jk}^2 \\
&= x_k + \frac{1}{\|\mathbf{a}_{*k}\|^2 + \lambda \|\mathbf{c}_{*k}\|^2} \left\{ \sum_{i=1}^M a_{ik} \left(y_i - [\mathbf{A}\mathbf{x}^{(k)}]_i \right) + \lambda \sum_{j=1}^N c_{jk} \sum_{n=1}^N c_{jn} x_n \right\} \\
&= x_k + \frac{1}{\|\mathbf{a}_{*k}\|^2 + \lambda \|\mathbf{c}_{*k}\|^2} \left\{ [\mathbf{A}^t (\mathbf{y} - \mathbf{A}\mathbf{x}^{(k)})]_k + \lambda [\mathbf{C}^t \mathbf{C}\mathbf{x}^{(k)}]_k \right\} \\
&= x_k + \frac{1}{\|\mathbf{a}_{*k}\|^2 + \lambda \|\mathbf{c}_{*k}\|^2} \left\{ [\mathbf{A}^t \mathbf{y}]_k - [(\mathbf{A}^t \mathbf{A} + \lambda \mathbf{C}^t \mathbf{C})\mathbf{x}^{(k)}]_k \right\} \\
&= x_k + \frac{1}{\|\mathbf{a}_{*k}\|^2 + \lambda \|\mathbf{c}_{*k}\|^2} \left\{ \langle \mathbf{a}_{*k}, (\mathbf{y} - \mathbf{A}\mathbf{x}^{(k)}) \rangle + \lambda [\mathbf{C}^t \mathbf{C}\mathbf{x}^{(k)}]_k \right\}
\end{aligned}$$

Algorithme :

$$\begin{aligned}
\mathbf{x}^{(0)} &= \mathbf{0} \\
x_k^{(k+1)} &= x_k^{(k)} + \alpha^{(k)} + \alpha^k \left\{ \langle \mathbf{a}_{*k}, (\mathbf{y} - \mathbf{A}\mathbf{x}^{(k)}) \rangle + \lambda [\mathbf{C}^t \mathbf{C}\mathbf{x}^{(k)}]_k \right\} \\
\text{avec } \alpha^k &= \frac{1}{\|\mathbf{a}_{*k}\|^2 + \lambda \|\mathbf{c}_{*k}\|^2}
\end{aligned}$$

Annexe C

Expressions des différentes fonctions potentiels

Dans cette annexe, nous faisons un inventaires des fonctions potentiels utilisées en modélisation markovienne pour la résolution des problèmes inverses.

Ces fonctions sont classées en deux familles : convexes et non convexes.

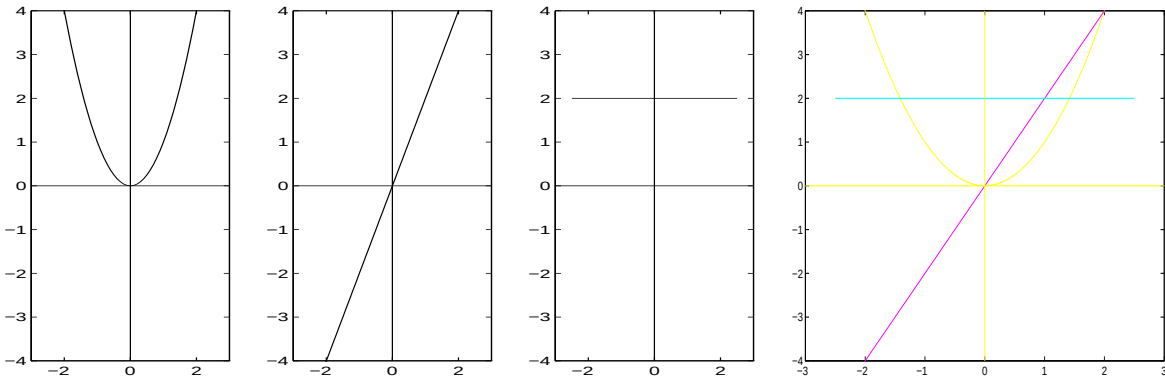
C.1 Potentiels convexes

Quadratique ou L_2 :

$$\phi(t) = t^2$$

$$\phi'(t) = 2t$$

$$\phi''(t) = 2$$

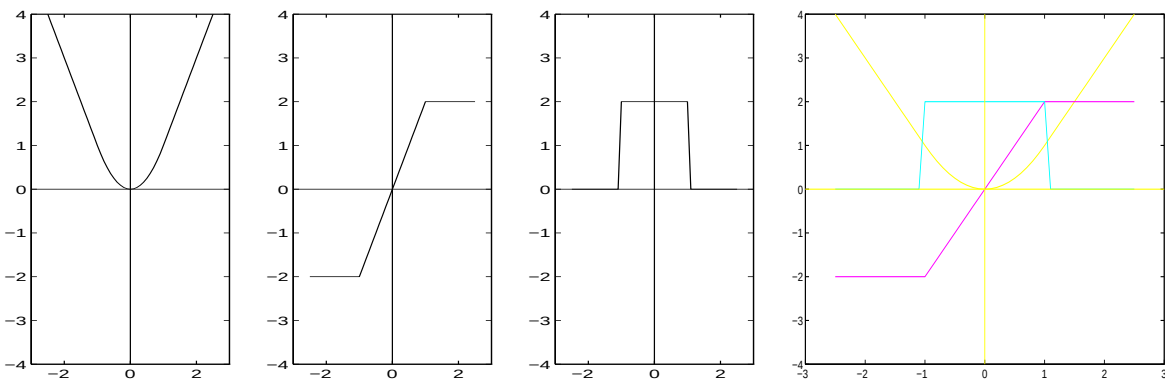


Quadratique prolongée en linéaire :

$$\phi(t) = \begin{cases} t^2 & |t| \leq 1 \\ 2t - 1 & |t| > 1 \end{cases}$$

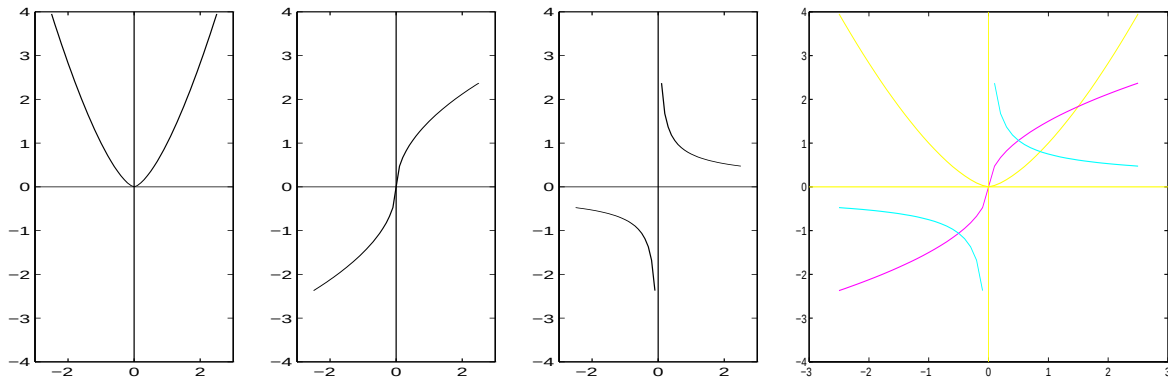
$$\phi'(t) = \begin{cases} 2t & |t| \leq 1 \\ 2 & |t| > 1 \end{cases}$$

$$\phi''(t) = \begin{cases} 2 & |t| \leq 1 \\ 0 & |t| > 1 \end{cases}$$



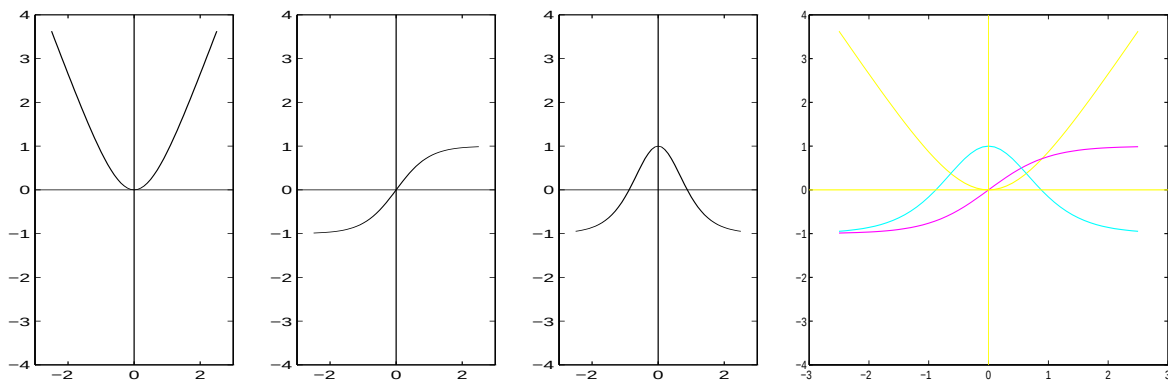
$L_p :$

$$\begin{aligned}\phi(t) &= |t|^p, \quad 1 < p \leq 2 \\ \phi'(t) &= p \operatorname{sign}(t) |t|^{p-1} \\ \phi''(t) &= p(p-1) |t|^{p-2}\end{aligned}$$



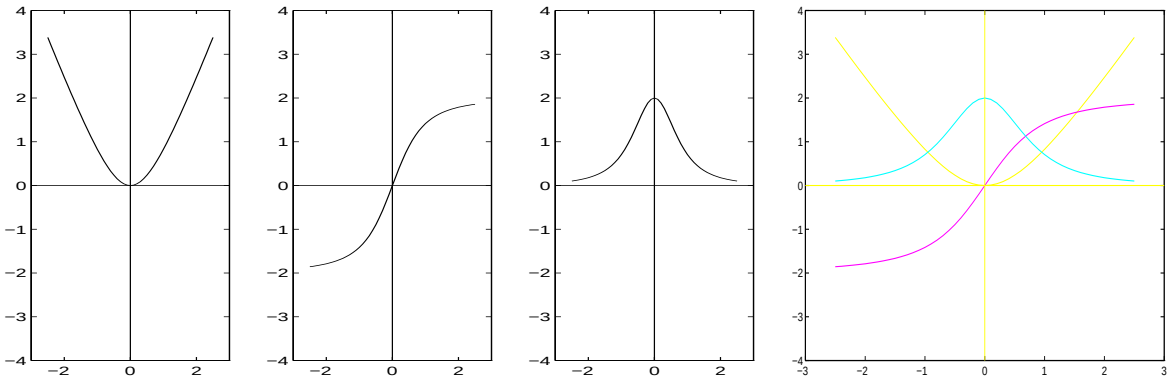
$\log \cosh :$

$$\begin{aligned}\phi(t) &= 2 \log(\cosh(t)) \\ \phi'(t) &= \frac{\sinh(t)}{\cosh(t)} \\ \phi''(t) &= \frac{1 - \sinh^2(t)}{\cosh^2(t)}\end{aligned}$$



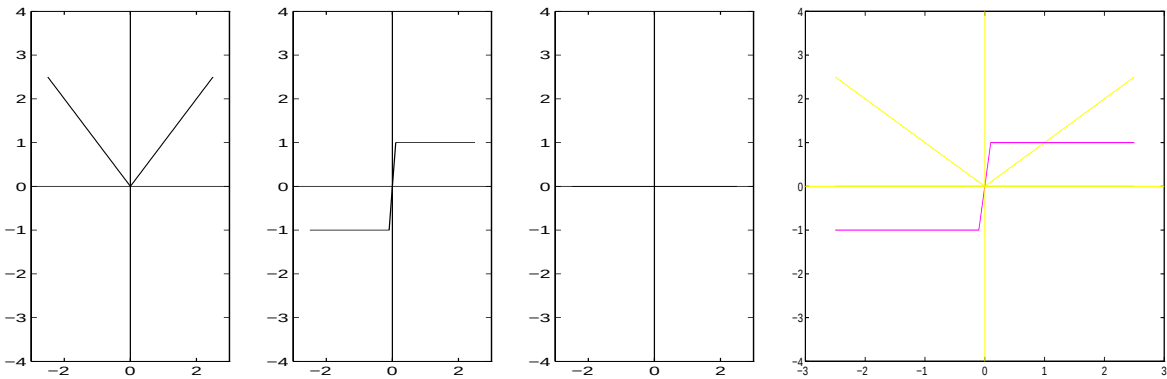
Racine carrée :

$$\begin{aligned}\phi(t) &= 2\sqrt{1+t^2} - 2 \\ \phi'(t) &= \frac{2t}{\sqrt{1+t^2}} \\ \phi''(t) &= \frac{2}{(1+t^2)^{3/2}}\end{aligned}$$



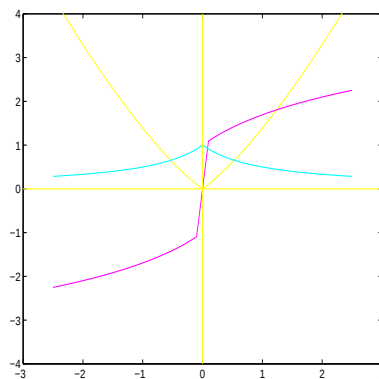
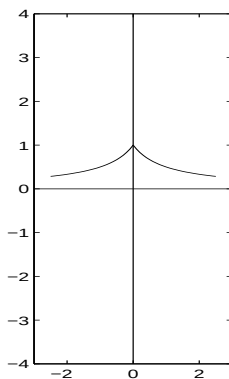
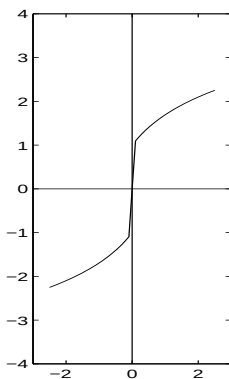
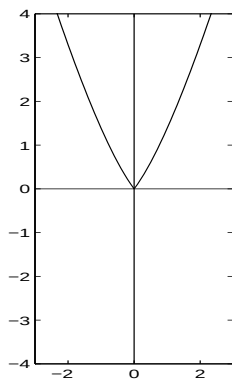
L_1 :

$$\begin{aligned}\phi(t) &= |t| \\ \phi'(t) &= \text{sign}(t) \\ \phi''(t) &= 0\end{aligned}$$



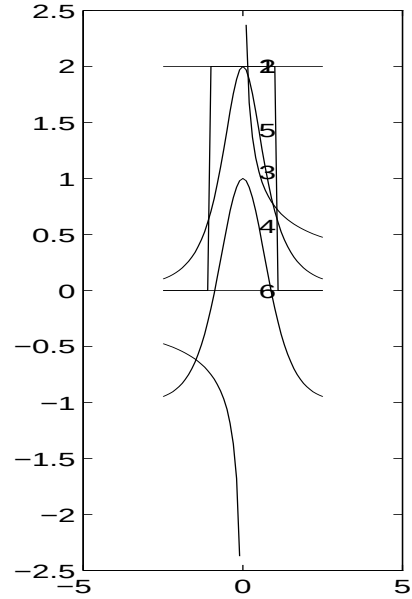
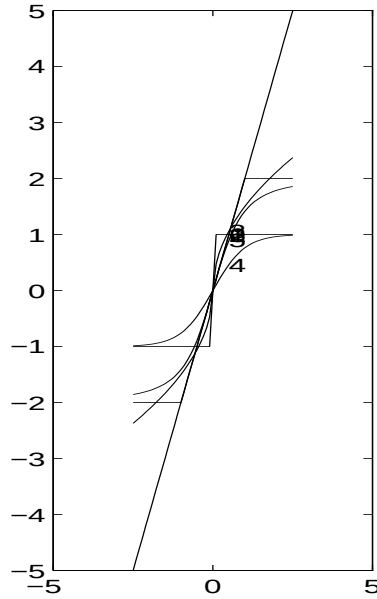
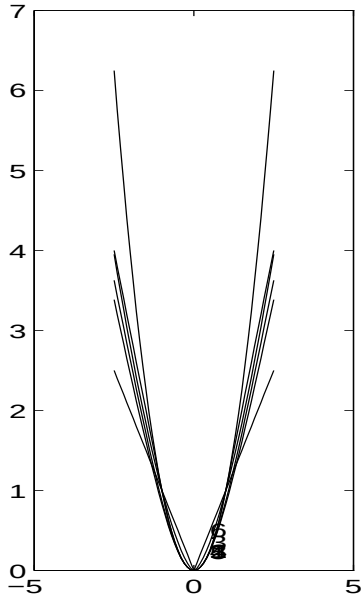
Logarithmique :

$$\begin{aligned}\phi(t) &= (|t| + 1) \log(|t| + 1) \\ \phi'(t) &= \text{sign}(t) (1 + \log(|t| + 1)) \\ \phi''(t) &= \frac{1}{|t| + 1}\end{aligned}$$



Le tableau ci-dessous résume ces fonctions :

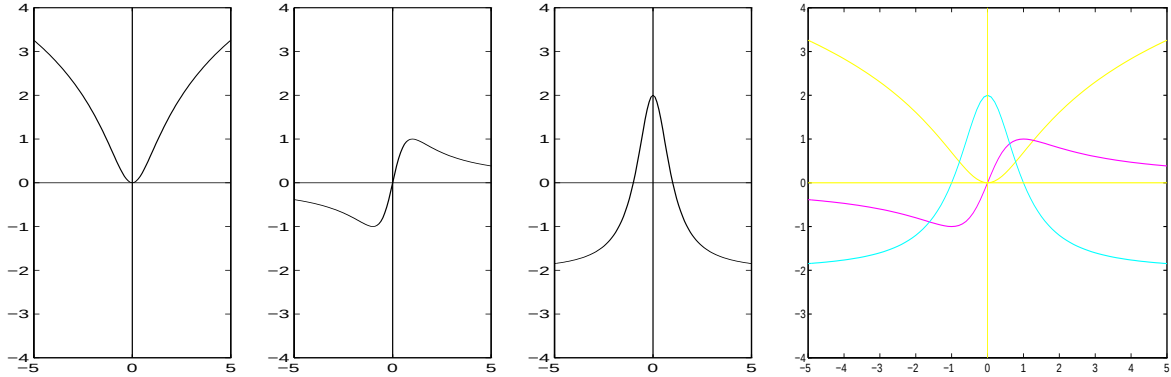
cas	$\phi(t)$	$\phi'(t)$	$\phi''(t)$	références
1	t^2	$2t$	2	Tikhonov
2	$\begin{cases} t^2 & t \leq 1 \\ 2t - 1 & t > 1 \end{cases}$	$\begin{cases} 2t & t \leq 1 \\ 2 & t > 1 \end{cases}$	$\begin{cases} 2 & t \leq 1 \\ 0 & t > 1 \end{cases}$	Hubert 81
3	$ t ^p, \quad 1 < p \leq 2$	$p \operatorname{sign}(t) t ^{p-1}$	$p(p-1) t ^{p-2}$	Bouman & Sauer 93
4	$2 \log(\cosh(t))$	$\frac{\sinh(t)}{\cosh(t)}$	$\frac{1 - \sinh^2(t)}{\cosh^2(t)}$	Green 90
5	$2\sqrt{1+t^2} - 2$	$\frac{2t}{\sqrt{1+t^2}}$	$\frac{2}{(1+t^2)^{3/2}}$	Aubert 94
6	$ t $	$\operatorname{sign}(t)$	0	Rudin Osher 92, Besag 86
7	$(t + 1) \log(t + 1)$	$\operatorname{sign}(t) (1 + \log(t + 1))$	$\frac{1}{ t + 1}$	Zervakis 95
8	$(t - m + t + 1) \log\left(\frac{ t - m + t + 1}{ t + 1} + 1\right)$	$\operatorname{sign}(t - m) \log\left(\frac{ t - m }{ t + 1} + 1\right)$	$\frac{1}{ t - m + t + 1}$	Zervakis 95
9	$\begin{cases} t^2 & t \leq 1 \\ 1 - 2 \log t & t > 1 \end{cases}$	$\begin{cases} 2t & t \leq 1 \\ \frac{2 t }{t^2} & t > 1 \end{cases}$	$\begin{cases} 2 & t \leq 1 \\ \frac{2}{t^2} & t > 1 \end{cases}$	AMD



C.2 Potentiels non convexes

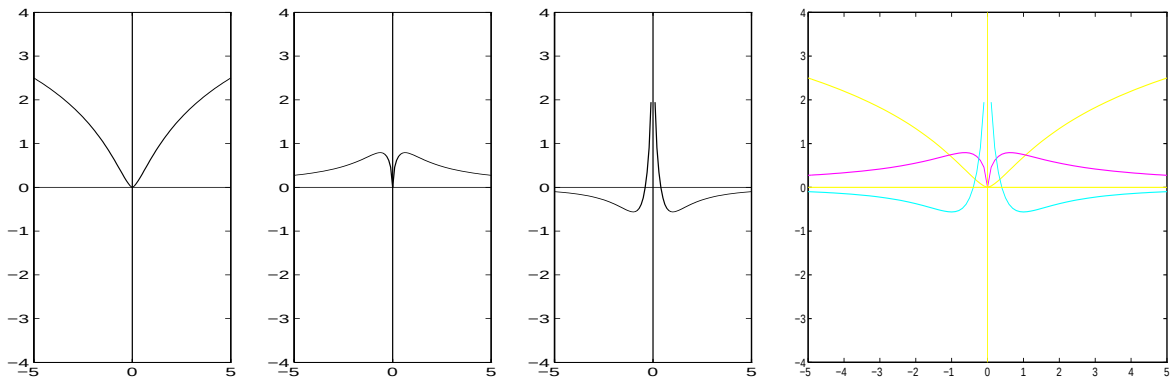
Hebert & Leahy, Perona & Malik :

$$\begin{aligned}\phi(t) &= \log(1 + t^2) \\ \phi'(t) &= \frac{2t}{1 + t^2} \\ \phi''(t) &= \frac{2(1 - t^2)}{(1 + t^2)^2}\end{aligned}$$



Extension :

$$\begin{aligned}\phi(t) &= \log(1 + |t|^p) \\ \phi'(t) &= \frac{p |t|^{p-1}}{1 + t^p} \\ \phi''(t) &= \frac{p(p-1)|t|^{p-2}(1 - t^p) - p^2 t^{2p-2}}{(1 + t^p)^2}\end{aligned}$$

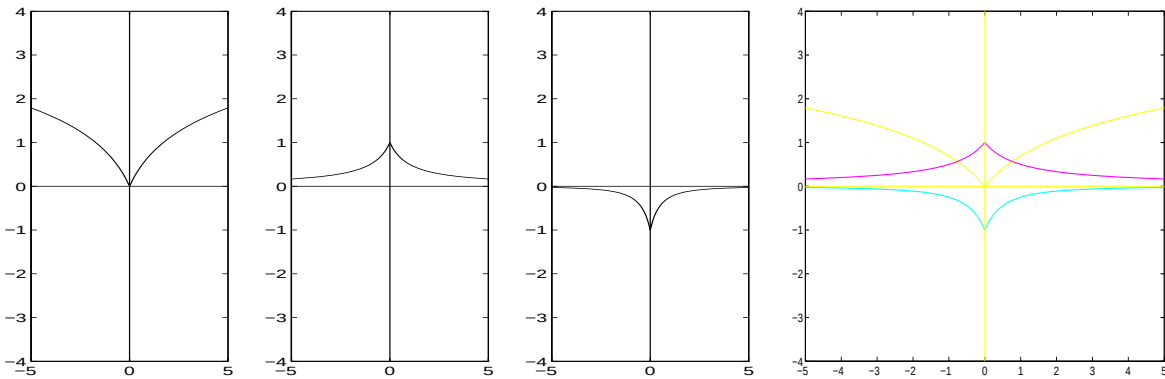


Cas particulier $p = 1$:

$$\phi(t) = \log(1 + |t|)$$

$$\phi'(t) = \frac{\text{sign}(t)}{1 + t}$$

$$\phi''(t) = \frac{-1}{(1 + t)^2}$$

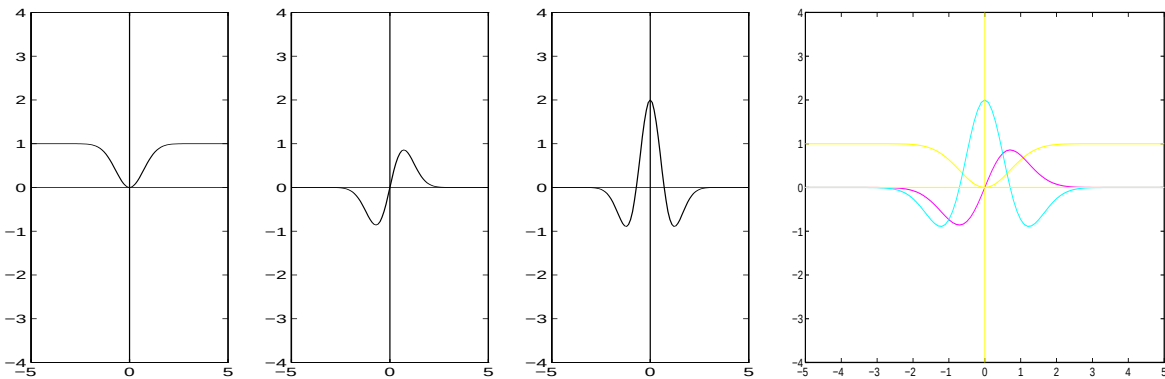


Gaussienne inversée :

$$\phi(t) = 1 - e^{-t^2}$$

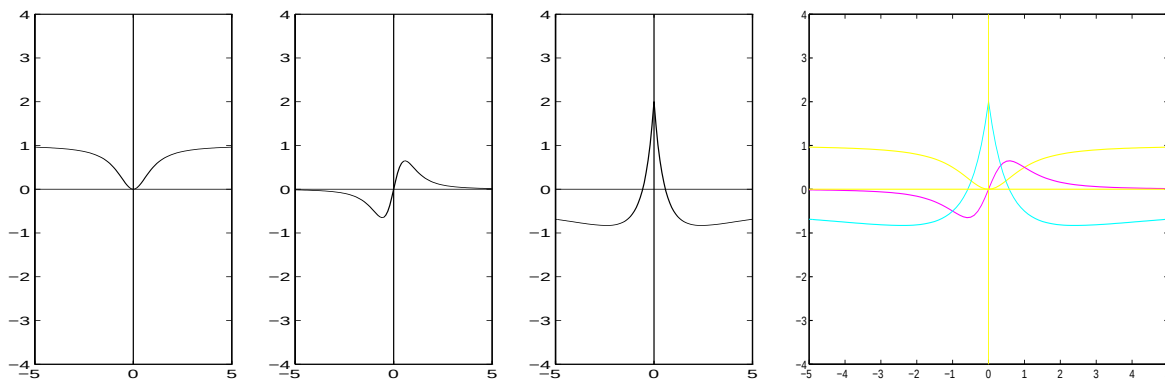
$$\phi'(t) = 2te^{-t^2}$$

$$\phi''(t) = 2e^{-t^2}(1 - 2t^2)$$



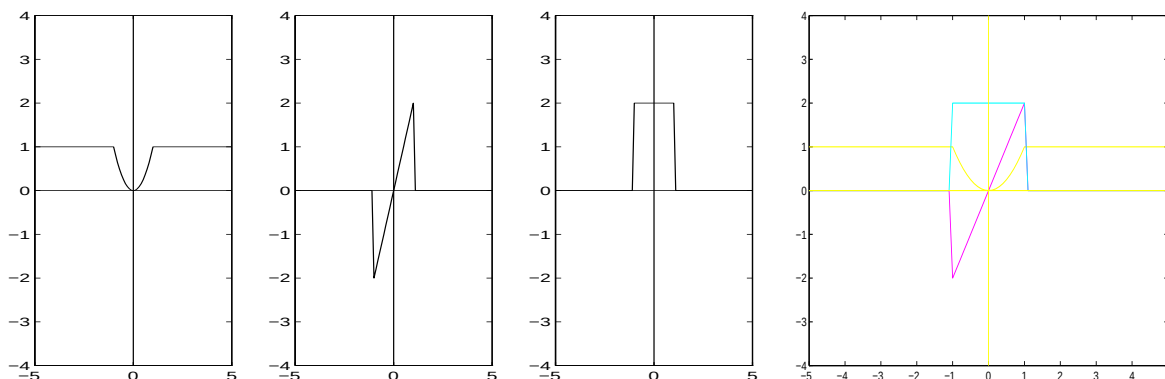
Lorentzienne inversée :

$$\begin{aligned}\phi(t) &= \frac{t^2}{1+t^2} \\ \phi'(t) &= \frac{2t}{(1+t^2)^2} \\ \phi''(t) &= \frac{2(1-3t^2)}{(1+t^2)^3}\end{aligned}$$



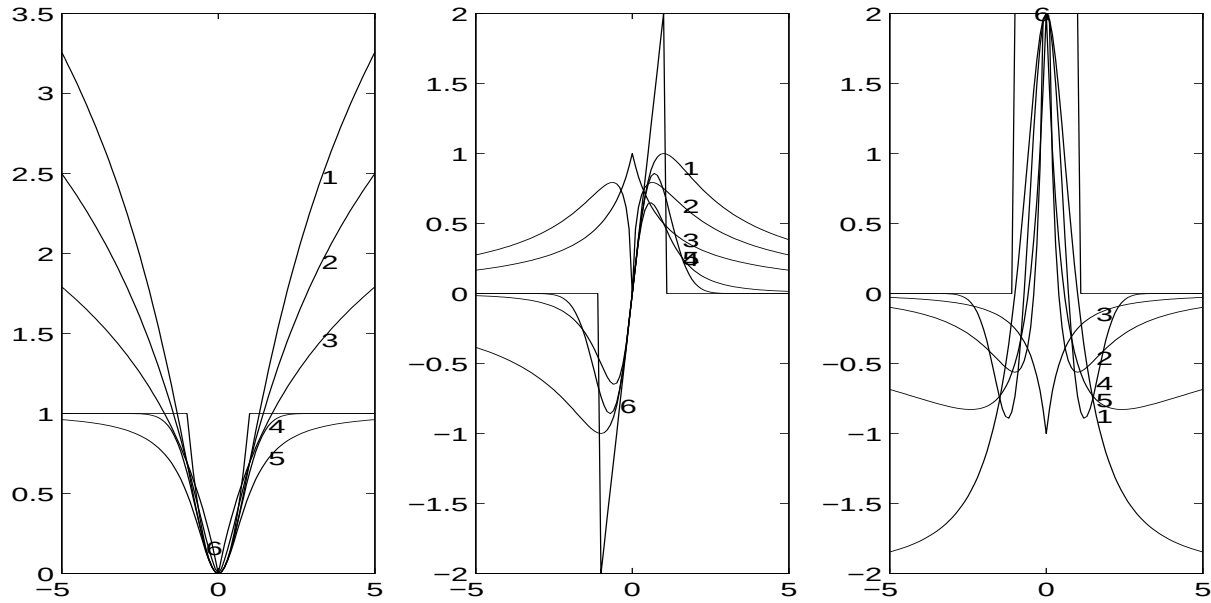
Quadratique tronquée :

$$\begin{aligned}\phi(t) &= \begin{cases} 2t & |t| \leq 1 \\ 0 & |t| > 1 \end{cases} \\ \phi'(t) &= \begin{cases} 2 & |t| \leq 1 \\ 0 & |t| > 1 \end{cases} \\ \phi''(t) &= \begin{cases} 2 & |t| \leq 1 \\ 0 & |t| > 1 \end{cases}\end{aligned}$$



Le tableau ci-dessous résume ces fonctions :

cas	$\phi(t)$	$\phi'(t)$	$\phi''(t)$	références
1	$\log(1 + t^2)$	$\frac{2t}{1+t^2}$	$\frac{2(1-t^2)}{(1+t^2)^2}$	Hebert & Leahy 90 Perona & Malik 90
2	$\log(1 + t ^p)$	$\frac{p t ^{p-1}}{1+t^p}$	$\frac{p(p-1) t ^{p-2}(1-t^p) - p^2 t^{2p-2}}{(1+t^p)^2}$	AMD
3	$\log(1 + t)$	$\text{sign}(t) \frac{1}{1+t}$	$\frac{-1}{(1+t)^2}$	AMD
4	$-e^{-t^2} + 1$	$2te^{-t^2}$	$2e^{-t^2}(1 - 2t^2)$	Perona & Malik 90
5	$\frac{t^2}{1+t^2}$	$\frac{2t}{(1+t^2)^2}$	$\frac{2(1-3t^2)}{(1+t^2)^3}$	Geman & McClure 85 90
6	$\min(1, t^2) = \begin{cases} t^2 & t \leq 1 \\ 1 & t > 1 \end{cases}$	$\begin{cases} 2t & t \leq 1 \\ 0 & t > 1 \end{cases}$	$\begin{cases} 2 & t \leq 1 \\ 0 & t > 1 \end{cases}$	Blake & Zisserman 87 Geman & Geman 84



Annexe D

Quelques exemples simples d'inférence bayésienne

Dans cette annexe, quelques exemples simples de l'inférence bayésienne sont présentés. L'objectif étant plutôt pédagogique.

D.1 Introduction et rappel des notations

Dans ces exemples nous utilisons les notations suivantes :

$$N(m, \lambda) = \sqrt{\frac{\lambda}{2\pi}} \exp \left[-\frac{\lambda}{2} (x - m)^2 \right]$$

$$N(\mathbf{m}, \mathbf{\Lambda}) = \frac{|\mathbf{\Lambda}|^{1/2}}{(2\pi)^{n/2}} \exp \left[-\frac{1}{2} (\mathbf{x} - \mathbf{m})^t \mathbf{\Lambda} (\mathbf{x} - \mathbf{m}) \right]$$

Problème 1 :

$$y = x + b$$

$$b|\beta \sim \mathbf{N}(0, \beta)$$

$$y|x, \beta \sim \mathbf{N}(x, \beta)$$

$$x|x_0, \theta \sim \mathbf{N}(x_0, \theta)$$

$$\begin{pmatrix} y \\ x \end{pmatrix} | \beta, x_0, \theta \sim \mathbf{N} \left(\begin{pmatrix} x_0 \\ x_0 \end{pmatrix}, \begin{pmatrix} \theta + \beta & \theta \\ \theta & \theta \end{pmatrix}^{-1} \right)$$

$$y|\beta, x_0, \theta \sim \mathbf{N} \left(x_0, \frac{\theta\beta}{\theta + \beta} \right)$$

$$\begin{aligned} x|y, \beta, x_0, \theta &\sim \mathbf{N}(x_m, \theta_m), \\ &\text{avec} \\ x_m &= \frac{1}{\theta_m} (\theta x_0 + \beta y) \\ \theta_m &= \theta + \beta \end{aligned}$$

$$y|y, \beta, x_0, \theta \sim \mathbf{N} \left(x_m, \frac{\theta_m \beta}{\theta_m + \beta} \right)$$

$$b|y, \beta, x_0, \theta \sim \mathbf{N} \left(0, \frac{2\theta_m \beta}{2\theta_m + \beta} \right)$$

$$\beta = 0 \longrightarrow x_m = x_0, \quad \theta_m = \theta$$

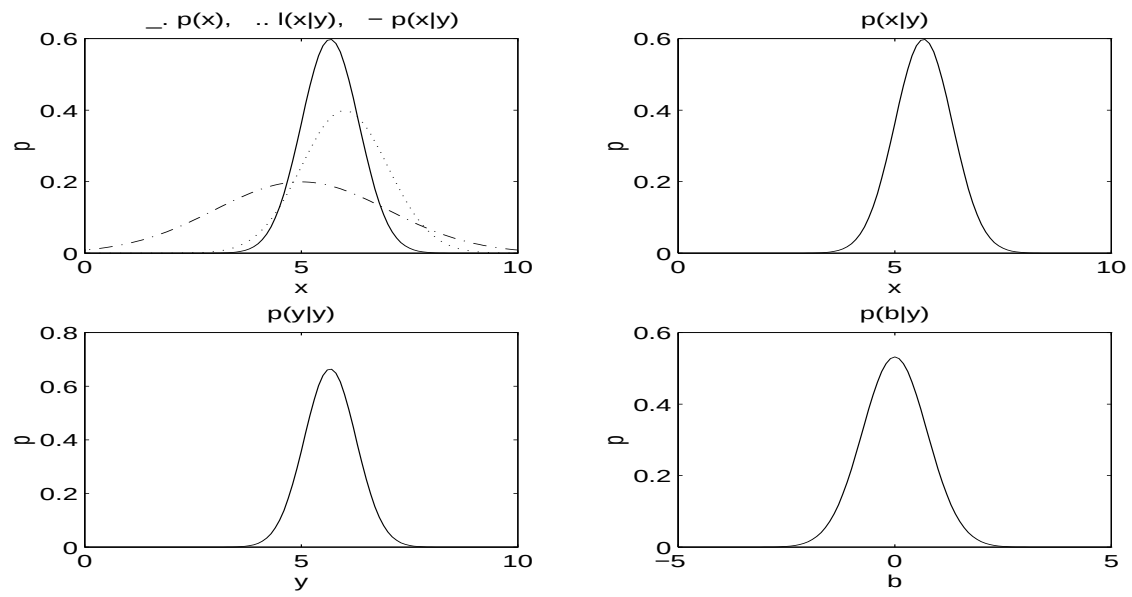
$$\theta = 0 \longrightarrow x_m = y, \quad \theta_m = \beta$$

$$\beta \longrightarrow \infty \longrightarrow x_m = y, \quad \theta_m = \infty$$

$$\theta \longrightarrow \infty \longrightarrow x_m = x_0, \quad \theta_m = \infty$$

$$\theta = r\beta \longrightarrow x_m = \frac{1}{r+1} (rx_0 + y), \quad \theta_m = (r+1)\beta$$

Exemple 1 : $y = 6$, $\beta = 1$, $x_0 = 5$, $\theta = .5 \rightarrow x_m = 5.6666$, $\theta_m = 1.5$



Problème 2 :

$$y_i = x + b_i, \quad i = 1, \dots, m \quad \longrightarrow \quad \mathbf{y} = x\mathbf{1} + \mathbf{b}$$

$$b_i | \beta \sim \mathbf{N}(0, \beta)$$

$$\mathbf{b} | \beta \sim \mathbf{N}(\mathbf{0}, \beta \mathbf{I})$$

$$y_i | x, \beta \sim \mathbf{N}(x, \beta)$$

$$\mathbf{y} | x, \beta \sim \mathbf{N}(x\mathbf{1}, \beta \mathbf{I})$$

$$x | x_0, \theta \sim \mathbf{N}(x_0, \theta)$$

$$\begin{pmatrix} y_i \\ x \end{pmatrix} | \beta, x_0, \theta \sim \mathbf{N} \left(\begin{pmatrix} x_0 \\ x_0 \end{pmatrix}, \begin{pmatrix} \theta + \beta & \theta \\ \theta & \theta \end{pmatrix}^{-1} \right)$$

$$\begin{pmatrix} \mathbf{y} \\ x \end{pmatrix} | \beta, x_0, \theta \sim \mathbf{N} \left(\begin{pmatrix} x_0 \mathbf{1} \\ x_0 \end{pmatrix}, \begin{pmatrix} (\theta + \beta) \mathbf{I} & \theta \mathbf{1} \\ \theta \mathbf{1}^t & \theta \end{pmatrix}^{-1} \right)$$

$$y_i | \beta, x_0, \theta \sim \mathbf{N}(x_0, \frac{\theta\beta}{\theta+\beta})$$

$$\mathbf{y} | \beta, x_0, \theta \sim \mathbf{N}(x_0 \mathbf{1}, \frac{\theta\beta}{\theta+\beta} \mathbf{I})$$

$$x | \mathbf{y}, \beta, x_0, \theta \sim \mathbf{N}(x_m, \theta_m),$$

avec

$$x_m = \frac{1}{\theta_m}(\theta x_0 + m\beta \bar{y}) \longrightarrow x_m = \frac{1}{\theta_m}(\theta_{m-1} x_{m-1} + \beta y_m)$$

$$\theta_m = \theta + m\beta \longrightarrow \theta_m = \theta_{m-1} + \beta$$

$$\bar{y} = \frac{1}{m} \sum_i y_i$$

$$y_k | \mathbf{y}, \beta, x_0, \theta \sim \mathbf{N} \left(x_m, \frac{\theta_m \beta}{\theta_m + \beta} \right)$$

$$b_k | \mathbf{y}, \beta, x_0, \theta \sim \mathbf{N} \left(0, \frac{2\theta_m \beta}{2\theta_m + \beta} \right)$$

$$\beta = 0 \quad \longrightarrow \quad x_m = x_0,$$

$$\theta_m = \theta$$

$$\theta = 0 \quad \longrightarrow \quad x_m = \bar{y},$$

$$\theta_m = m\beta$$

$$\beta \longrightarrow \infty \quad \longrightarrow \quad x_m = \bar{y},$$

$$\theta_m = \infty$$

$$\theta \longrightarrow \infty \quad \longrightarrow \quad x_m = x_0,$$

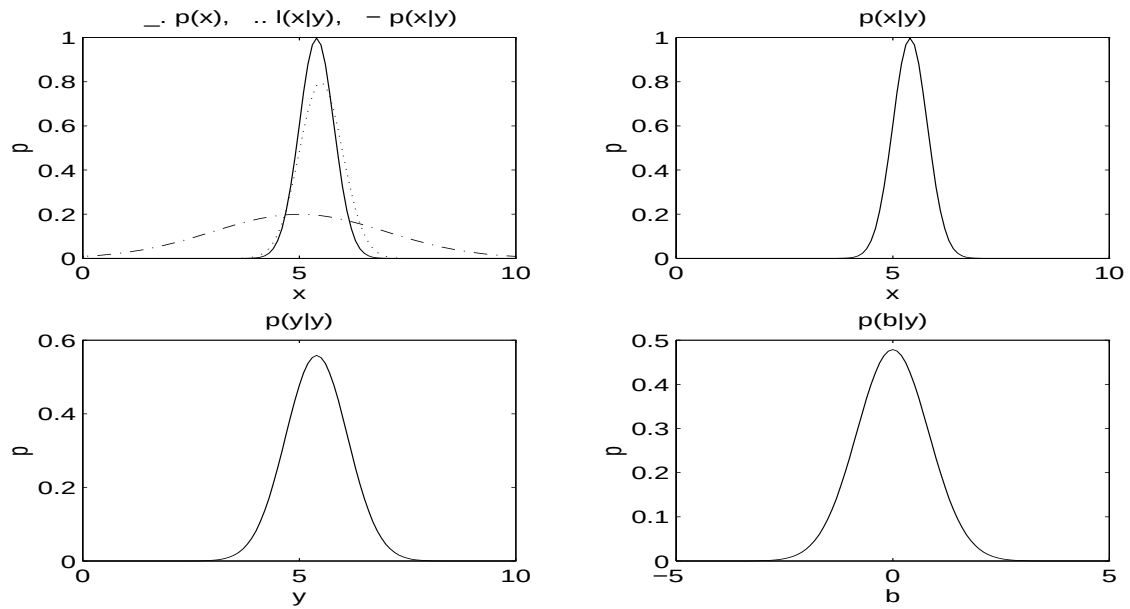
$$\theta_m = \infty$$

$$m \longrightarrow \infty \quad \longrightarrow \quad x_m = \bar{y},$$

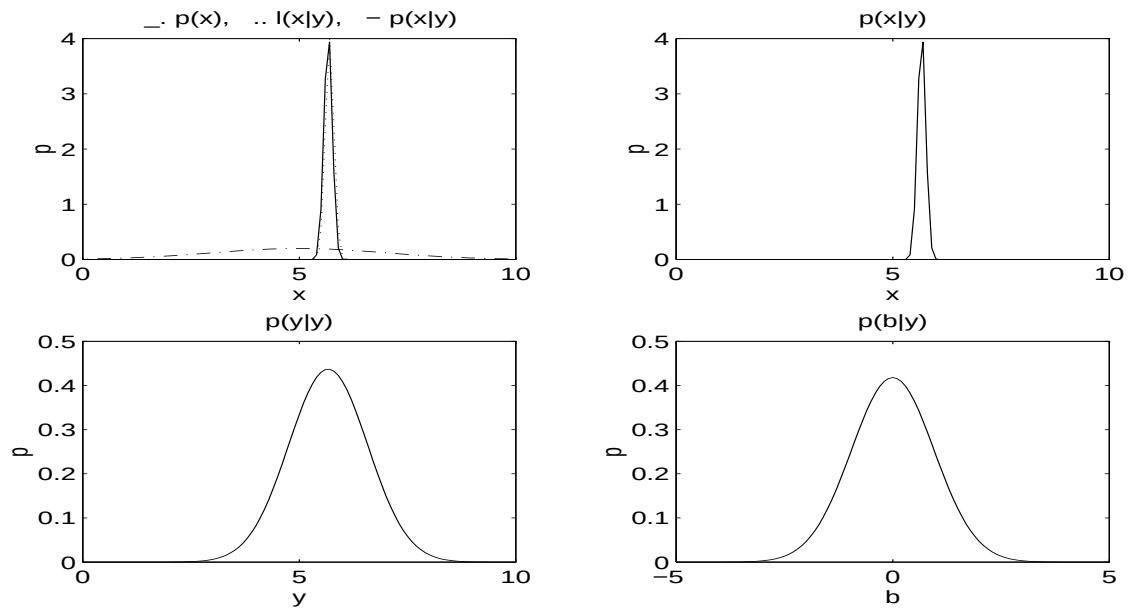
$$\theta_m = \infty$$

$$\theta = r\beta \quad \longrightarrow \quad x_m = \frac{1}{r+m}(rx_0 + m\bar{y}), \quad \theta_m = (r+m)\beta$$

Exemple 2.1: $\mathbf{y} = [6, 5]$, $m = 2$, $\bar{y} = 5.5$, $\beta = 1$, $x_0 = 5$, $\theta = .5$
 $x_m = 5.4$, $\theta_m = 2.5$



Exemple 2.2: $\mathbf{y} = [6, 5, 4, 7, 3, 5, 6, 8, 7, 6]$, $m = 10$, $\bar{y} = 5.7$
 $\beta = 1$, $x_0 = 5$, $\theta = .5$ $x_m = 5.667$, $\theta_m = 10.5$



Problème 3 :

$$y = ax + b$$

$$\begin{aligned} b|\beta &\sim \mathbf{N}(0, \beta) \\ y|x, \beta &\sim \mathbf{N}(ax, \beta) \end{aligned}$$

$$x|x_0, \theta \sim \mathbf{N}(x_0, \theta)$$

$$\begin{pmatrix} y \\ x \end{pmatrix} | \beta, x_0, \theta \sim \mathbf{N} \left(\begin{pmatrix} ax_0 \\ x_0 \end{pmatrix}, \begin{pmatrix} \theta + a^2\beta & a\theta \\ a\theta & \theta \end{pmatrix}^{-1} \right)$$

$$y|\beta, x_0, \theta \sim \mathbf{N} \left(ax_0, \frac{\theta\beta}{\theta + a^2\beta} \right)$$

$$\begin{aligned} x|y, \beta, x_0, \theta &\sim \mathbf{N}(x_m, \theta_m), \\ &\text{avec} \\ x_m &= \frac{1}{\theta_m}(\theta x_0 + a\beta y), \\ \theta_m &= \theta + a^2\beta \end{aligned}$$

$$\begin{aligned} y|y, \beta, x_0, \theta &\sim \mathbf{N} \left(ax_m, \frac{\theta_m\beta}{\theta_m + a^2\beta} \right) \\ b|y, \beta, x_0, \theta &\sim \mathbf{N} \left(0, \frac{2\theta_m\beta}{2\theta_m + a^2\beta} \right) \end{aligned}$$

$$\begin{array}{lll} \beta = 0 & \longrightarrow & x_m = x_0, & \theta_m = \theta \\ \theta = 0 & \longrightarrow & x_m = y/a, & \theta_m = a^2\beta \\ \beta \longrightarrow \infty & \longrightarrow & x_m = y/a, & \theta_m = \infty \\ \theta \longrightarrow \infty & \longrightarrow & x_m = x_0, & \theta_m = \infty \end{array}$$

$$\begin{array}{lll} \theta = r\beta & \longrightarrow & x_m = \frac{1}{r+a^2}(rx_0 + ay), & \theta_m = (r + a^2)\beta \\ r \longrightarrow \infty & \longrightarrow & x_m = x_0, & \theta_m = \infty \\ r = 0 & \longrightarrow & x_m = y/a, & \theta_m = a^2\beta \end{array}$$

Problème 4 :

$$y_i = ax + b_i, \quad i = 1, \dots, m \quad \longrightarrow \quad \mathbf{y} = ax\mathbf{1} + \mathbf{b}$$

$$b_i | \beta \sim \mathbf{N}(0, \beta)$$

$$\mathbf{b} | \beta \sim \mathbf{N}(\mathbf{0}, \beta \mathbf{I})$$

$$y_i | x, \beta \sim \mathbf{N}(ax, \beta)$$

$$\mathbf{y} | x, \beta \sim \mathbf{N}(ax\mathbf{1}, \beta \mathbf{I})$$

$$x | x_0, \theta \sim \mathbf{N}(x_0, \theta)$$

$$\begin{pmatrix} y_i \\ x \end{pmatrix} | \beta, x_0, \theta \sim \mathbf{N} \left(\begin{pmatrix} ax_0 \\ x_0 \end{pmatrix}, \begin{pmatrix} a^2\theta + \beta & a\theta \\ a\theta & \theta \end{pmatrix}^{-1} \right)$$

$$\begin{pmatrix} \mathbf{y} \\ x \end{pmatrix} | \beta, x_0, \theta \sim \mathbf{N} \left(\begin{pmatrix} ax_0\mathbf{1} \\ x_0 \end{pmatrix}, \begin{pmatrix} a^2\theta + \beta\mathbf{I} & a\theta\mathbf{1} \\ a\theta\mathbf{1}^t & \theta \end{pmatrix}^{-1} \right)$$

$$y | \beta, x_0, \theta \sim \mathbf{N} \left(ax_0, \frac{\theta\beta}{\theta + a^2\beta} \right)$$

$$\mathbf{y} | \beta, x_0, \theta \sim \mathbf{N} \left(ax_0\mathbf{1}, \frac{\theta\beta}{\theta + a^2\beta} \mathbf{I} \right)$$

$$x | \mathbf{y}, \beta, x_0, \theta \sim \mathbf{N}(x_m, \theta_m),$$

avec

$$x_m = \frac{1}{\theta_m}(\theta x_0 + ma\beta\bar{y}) \longrightarrow x_m = \frac{1}{\theta_m}(\theta_{m-1}x_0 + a\beta y_m)$$

$$\theta_m = (\theta + ma^2\beta) \longrightarrow \theta_m = \theta_{m-1} + a^2\beta$$

$$\bar{y} = \frac{1}{m} \sum_i y_i$$

$$y_k | \mathbf{y}, \beta, x_0, \theta \sim \mathbf{N} \left(ax_m, \frac{\theta_m\beta}{\theta_m + ma^2\beta} \right)$$

$$b_k | \mathbf{y}, \beta, x_0, \theta \sim \mathbf{N} \left(0, \frac{2\theta_m\beta}{2\theta_m + ma^2\beta} \right)$$

$$\begin{array}{lll} \beta = 0 & \longrightarrow & x_m = x_0, & \theta_m = \theta \\ \theta = 0 & \longrightarrow & x_m = \bar{y}/a, & \theta_m = ma^2\beta \\ \beta \longrightarrow \infty & \longrightarrow & x_m = \bar{y}/a, & \theta_m = \infty \\ \theta \longrightarrow \infty & \longrightarrow & x_m = x_0, & \theta_m = \infty \\ m \longrightarrow \infty & \longrightarrow & x_m = \bar{y}/a, & \theta_m = \infty \end{array}$$

$$\begin{array}{lll} \theta = r\beta & \longrightarrow & x_m = \frac{1}{r+ma^2}(ma\bar{y} + rx_0), & \theta_m = (r + ma^2)\beta \\ r \longrightarrow \infty & \longrightarrow & x_m = x_0, & \theta_m = \infty \\ r = 0 & \longrightarrow & x_m = \bar{y}/a, & \theta_m = ma^2\beta \end{array}$$

Problème 5 :

$$y_i = x + b_i, \quad i = 1, \dots, m \quad \longrightarrow \quad \mathbf{y} = ax\mathbf{1} + \mathbf{b}$$

$$\begin{aligned} b_i | \beta &\sim \mathbf{N}(0, \beta) \\ y_i | x, \beta &\sim \mathbf{N}(x, \beta) \end{aligned}$$

$$\begin{aligned} x | x_0, \theta &\sim \mathbf{N}(x_0, \theta), \quad \theta = r\beta, \quad r = \theta/\beta \\ \theta | a, b &\sim \mathbf{Gam}(a, rb), \quad a = r_0b, \quad r_0 = a/b, \quad \mathbb{E}[\theta] = \frac{a}{rb} = \frac{r_0}{r}, \quad \text{Var}[\theta] = \frac{a}{(rb)^2} \\ (x, \theta) &\sim \mathbf{NGam}(x_0, 1, a, rb) \end{aligned}$$

$$\begin{aligned} x | x_0, r, \beta &\sim \mathbf{N}(x_0, r\beta), \\ \beta | a, b &\sim \mathbf{Gam}(a, b), \quad a = r_0b, \quad r_0 = a/b, \quad \mathbb{E}[\beta] = \frac{a}{b} = r_0, \quad \text{Var}[\beta] = \frac{a}{b^2} = \frac{r_0}{b} \\ (x, \beta) &\sim \mathbf{NGam}(x_0, r, a, b) \end{aligned}$$

$$\begin{aligned} x | x_0, r, r_0, a &\sim \mathbf{St}(x_0, rr_0, 2a), \quad \mathbb{E}[x] = x_0, \quad \text{Var}[x] = \frac{2a}{2a-2} \frac{1}{rr_0} \\ y_i | x_0, r, r_0, a &\sim \mathbf{St}\left(x_0, \frac{r}{r+1}r_0, 2a\right), \quad \mathbb{E}[y] = x_0, \quad \text{Var}[y] = \frac{2a}{2a-2} \frac{r+1}{rr_0} \end{aligned}$$

$$\begin{aligned} x | \mathbf{y}, x_0, r, r_0, a &\sim \mathbf{St}(x_m, (m+r)r_m, 2a_m), \\ &\text{avec} \\ x_m &= \frac{1}{r+m}(rx_0 + m\bar{y}) \longrightarrow x_m = \frac{1}{r+m}[(r+m-1)x_{m-1} + y_m] \\ r_m &= a_m/b_m \\ a_m &= a + m/2, \quad b_m = b + ms^2/2 + \frac{1}{2} \frac{rm}{r+m}(x_0 - \bar{y})^2 \\ \bar{y} &= \frac{1}{m} \sum_{i=1}^m y_i, \quad s^2 = \sum_{i=1}^m (y_i - \bar{y})^2 \\ \mathbb{E}[x | \mathbf{y}] &= x_m, \text{ si } 2a_m > 1 \\ \text{Var}[x | \mathbf{y}] &= \frac{1}{(m+r)r_m} \frac{2a_m}{2a_m-2}, \text{ si } 2a_m > 2 \\ \text{mode}(x | \mathbf{y}) &= x_m \end{aligned}$$

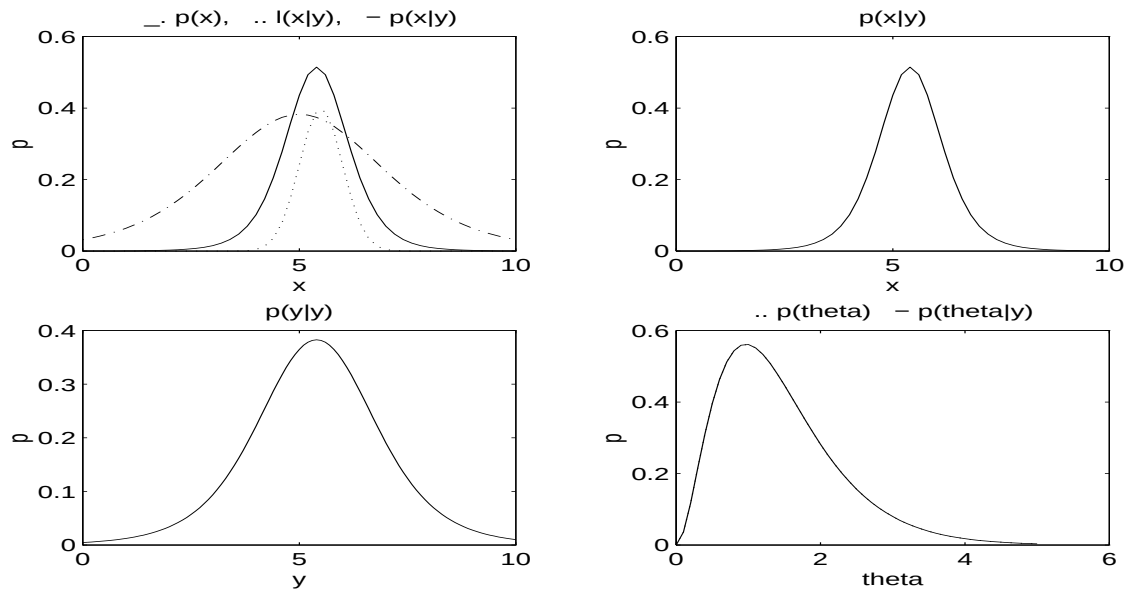
$$\begin{aligned} y_k | \mathbf{y}, x_0, r, r_0, a &\sim \mathbf{St}\left(x_m, \frac{m+r}{m+r+1}r_m, 2a_m\right) \\ \beta | \mathbf{y}, x_0, r, r_0, a &\sim \mathbf{Gam}(a_m, b_m) \\ \theta | \mathbf{y}, x_0, r, r_0, a &\sim \mathbf{Gam}(a_m, rb_m) \\ \mathbb{E}[\beta | \mathbf{y}] &= \frac{a_m}{b_m}, \quad \text{Var}[\beta | \mathbf{y}] = \frac{a_m}{b_m^2}, \quad \text{mode}(\beta | \mathbf{y}) = \frac{a_m-1}{b_m} \\ \mathbb{E}[\theta | \mathbf{y}] &= \frac{a_m}{rb_m}, \quad \text{Var}[\theta | \mathbf{y}] = \frac{a_m}{r^2b_m^2}, \quad \text{mode}(\theta | \mathbf{y}) = \frac{a_m-1}{rb_m} \end{aligned}$$

$$\begin{aligned} \theta = 0 \longrightarrow r = 0 &\longrightarrow x_m = \bar{y}, \quad b_m = b + ms^2/2 \\ \theta \longrightarrow \infty &\longrightarrow x_m = x_0, \quad b_m = b + ms^2/2 + \frac{m}{2}(x_0 - \bar{y})^2 = b + \frac{m}{2}[s^2 + (x_0 - \bar{y})^2] \\ m \longrightarrow \infty &\longrightarrow x_m = \bar{y}, \quad x | \mathbf{y}, x_0, r, r_0, a \sim \mathbf{N}(\bar{y}, 0) \end{aligned}$$

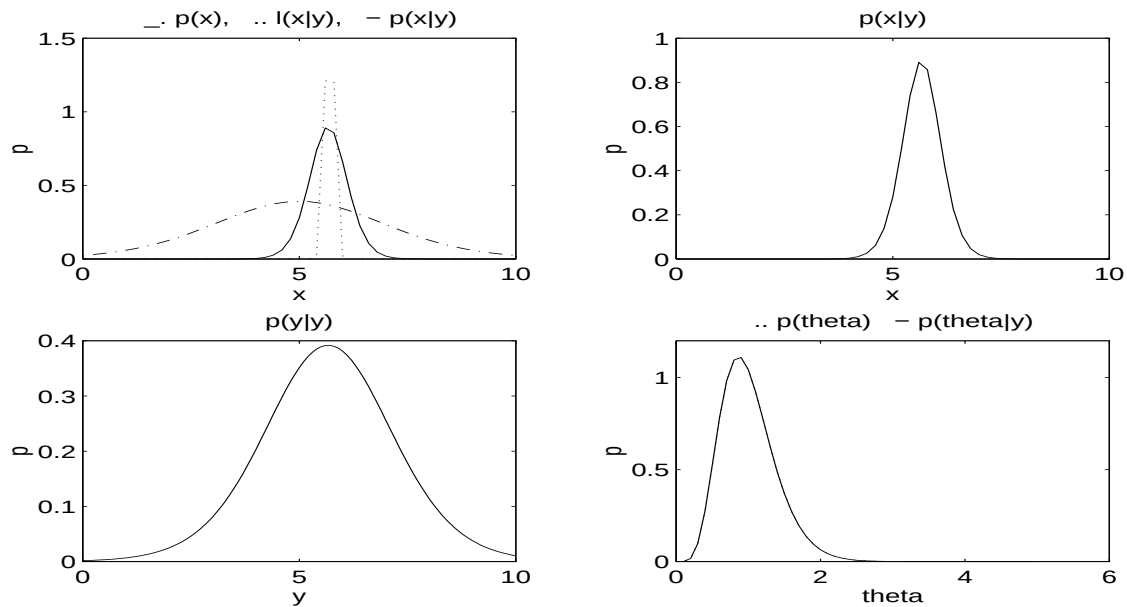
Cas particulier :

$$\begin{aligned} x_0 = \bar{y} &\longrightarrow x_m = \bar{y} & a_m = a + m/2, \quad b_m = b + (m/2)s^2, \\ x_0 = \bar{y}, a = \frac{1}{2}, b = \frac{s^2}{2} &\longrightarrow x_m = \bar{y}, & a_m = \frac{m+1}{2}, b_m = \frac{m+1}{2}s^2, \\ & & r_m = \frac{1}{s^2}, \\ & & \text{Var}[x | \mathbf{y}] = \frac{m+1}{m-1} \frac{1}{m+r} s^2 \end{aligned}$$

Exemple 5.1: $\mathbf{y} = [6, 5]$, $m = 2$, $\bar{y} = 5.5$, $s^2 = 0.25$,
 $\beta = 1, x_0 = 5, \theta = .5, r = .5$, $a = 2, b = 4, r_0 = .5$, $\lambda_0 = 0.25, \text{Var}[x] = 6$
 $a_m = 3, b_m = 4.15, r_m = 0.7229$, $x_m = 5.4, \lambda_m = 1.8072, \text{Var}[x|\mathbf{y}] = .83$



Exemple 5.2: $\mathbf{y} = [6, 5, 4, 7, 3, 5, 6, 8, 7, 6]$, $m = 10$, $\bar{y} = 5.7$, $s^2 = 2.01$,
 $\beta = 1, x_0 = 5, \theta = .5, r = .5$, $a = 2, b = 4, r_0 = .5$, $\lambda_0 = 0.25, \text{Var}[x] = 4.667$
 $a_m = 7, b_m = 13.88, r_m = 0.5042$, $x_m = 5.667, \lambda_m = 5.294, \text{Var}[x|\mathbf{y}] = .2204$



Annexe E

Quelques rappels sur les matrices

Dans cette annexe d'abord nous rappelons un certain nombres de définitions et de relations sur les matrices et le calcul matriciel. Ensuite, les propriétés d'un certain nombres de matrices que l'on trouve souvent en traitement du signal et des images sont étudiées. parmi ces matrices, une place particulière est donnée aux matrices de Toeplitz, de Hankel, circulantes, bloc-Toeplitz et bloc-circulantes.

E.1 Introduction

E.1.1 Définition

Soient $\mathcal{V}$ et $\mathcal{W}$ deux espaces vectoriels sur le même corps (le corps $\mathcal{R}$ des nombres réels ou le corps $\mathcal{C}$ des nombres complexes ou, d'une façon plus générale, le corps $\mathcal{K}$ des scalaires), munis de bases $\{\mathbf{e}_i, i = 1, \dots, m\}$ et $\{\mathbf{f}_j, j = 1, \dots, n\}$ respectivement. Relativement à ces bases, une application linéaire

$$\mathbf{A} : \mathcal{V} \longrightarrow \mathcal{W}$$

est représentée par la matrice $\mathbf{A}$ à m lignes et n colonnes :

$$\mathbf{A} = \begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \vdots & & & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} \end{pmatrix}$$

les éléments a_{ij} de la matrice $\mathbf{A}$ étant définis de façon unique par les relations

$$\mathbf{A}\mathbf{e}_j = \sum_{i=1}^m a_{ij} \mathbf{f}_i, \quad i = 1, \dots, m.$$

Autrement dit, le j -ème vecteur colonne

$$\mathbf{a}_{*j} = [\mathbf{A}]_{*j} = \begin{pmatrix} a_{1j} \\ a_{2j} \\ \vdots \\ a_{mj} \end{pmatrix}$$

de la matrice $\mathbf{A}$ représente le vecteur $\mathbf{A}\mathbf{e}_j$ dans la base $\{\mathbf{f}_i, i = 1, \dots, m\}$. On note

$$\mathbf{a}_{i*} = [\mathbf{A}]_{i*} = (a_{i1} \ a_{i2} \ \cdots \ a_{in})^t$$

le i -ème vecteur ligne de la matrice $\mathbf{A}$.

- Une matrice $\mathbf{A}$ d'éléments a_{ij} est notée $\mathbf{A} = [a_{ij}]$. Une matrice à m lignes et n colonnes est appelée matrice de type $[m, n]$, et on note $\mathcal{A}_{m,n}(\mathcal{K})$ l'espace vectoriel sur le corps $\mathcal{K}$ formé par les matrices de type $[m, n]$ à éléments dans $\mathcal{K}$. Un vecteur colonne est donc une matrice de type $[m, 1]$ et un vecteur ligne une matrice de type $[1, n]$.
- étant donnée une matrice $\mathbf{A} \in \mathcal{A}_{m,n}(\mathcal{C})$, on note $\mathbf{A}^* \in \mathcal{A}_{n,m}(\mathcal{C})$ la matrice adjointe de la matrice $\mathbf{A}$, définie de façon unique par les relations

$$\langle \mathbf{A}\mathbf{u}, \mathbf{v} \rangle_m = \langle \mathbf{u}, \mathbf{A}^*\mathbf{v} \rangle_n, \quad \forall \mathbf{u} \in \mathcal{C}^n, \mathbf{v} \in \mathcal{C}^m$$

qui entraînent $[\mathbf{A}^*]_{ij} = a_{ji}^*$. De la même façon, étant donnée une matrice $\mathbf{A} \in \mathcal{A}_{m,n}(\mathcal{R})$, on note $\mathbf{A}^t \in \mathcal{A}_{n,m}(\mathcal{R})$ la matrice transposée de la matrice $\mathbf{A}$, définie de façon unique par les relations

$$\langle \mathbf{A}\mathbf{u}, \mathbf{v} \rangle_m = \langle \mathbf{u}, \mathbf{A}^t\mathbf{v} \rangle_n \quad \forall \mathbf{u} \in \mathcal{R}^n, \mathbf{v} \in \mathcal{R}^m$$

qui entraînent $[\mathbf{A}^t]_{ij} = a_{ji}$.

E.1.2 Opérations élémentaires

- **Addition :**

$$\mathbf{C} = \mathbf{A} + \mathbf{B}, \quad [\mathbf{C}]_{ij} = [\mathbf{A}]_{ij} + [\mathbf{B}]_{ij}$$

$$[\mathbf{A} + \mathbf{B} + \mathbf{C}] = [\mathbf{C} + \mathbf{B} + \mathbf{A}] = [\mathbf{B} + \mathbf{C} + \mathbf{A}]$$

- **Multiplication par un scalaire :**

$$\mathbf{C} = \lambda \mathbf{A}, \quad [\mathbf{C}]_{ij} = \lambda [\mathbf{A}]_{ij}$$

$$\lambda [\mathbf{A} + \mathbf{B}] = \lambda \mathbf{A} + \lambda \mathbf{B}$$

- **Multiplication de deux matrices :**

À la composition des applications linéaires correspond la multiplication des matrices : si $\mathbf{A} = [a_{ij}]$ est une matrice de type $[I, J]$ et $\mathbf{B} = [b_{jk}]$ est une matrice de type $[J, K]$, leur produit $\mathbf{AB}$ est la matrice $\mathbf{C} = [c_{ik}]$ de type $[I, K]$ définie par

$$\mathbf{C} = \mathbf{AB} \longrightarrow c_{ik} = \sum_{j=1}^J a_{ij} b_{jk}$$

On rappelle que l'on a :

$$[\mathbf{A} + \mathbf{B}]^t = \mathbf{A}^t + \mathbf{B}^t, \quad [\mathbf{A} + \mathbf{B}]^* = \mathbf{A}^* + \mathbf{B}^*$$

$$[\mathbf{AB}]^t = \mathbf{B}^t \mathbf{A}^t, \quad [\mathbf{AB}]^* = \mathbf{B}^* \mathbf{A}^*$$

$$\lambda[\mathbf{AB}] = [\lambda\mathbf{A}]\mathbf{B} = \mathbf{A}[\lambda\mathbf{B}]$$

$$\mathbf{A}[\mathbf{B} + \mathbf{C}] = \mathbf{AB} + \mathbf{AC}$$

$$[\mathbf{A} + \mathbf{B}]\mathbf{C} = \mathbf{AC} + \mathbf{BC}$$

E.1.3 Matrices élémentaires

- Une matrice de type $[n, n]$ est dite matrice carrée d'ordre n . Si $\mathbf{A}$ est une matrice carrée, les éléments a_{ii} sont appelés *éléments diagonaux*, et les éléments $a_{ij}, i \neq j$, sont appelés *éléments hors-diagonaux*.
- Une matrice de type $[n, n]$ est dite *triangulaire supérieure* si $a_{ij} = 0$ pour $i > j$, et *triangulaire inférieure* si $a_{ij} = 0$ pour $i < j$.
- La *matrice unité* est la matrice $\mathbf{I} = [d_{ij}]$.
- Une matrice est *diagonale* si $a_{ij} = 0$, pour $i \neq j$ on la note

$$\mathbf{A} = \text{diag}[a_{ii}] = \text{diag}[a_{11}, a_{22}, \dots, a_{nn}]$$

- Une matrice $\mathbf{A}$ est :

<i>symétrique</i>	si	$\mathbf{A}$ est réelle et $\mathbf{A} = \mathbf{A}^t$
<i>hermitienne</i>	si	$\mathbf{A} = \mathbf{A}^*$
<i>orthogonale</i>	si	$\mathbf{A}$ est réelle et $\mathbf{AA}^t = \mathbf{A}^t \mathbf{A} = \mathbf{I}$
<i>unitaire</i>	si	$\mathbf{AA}^* = \mathbf{A}^* \mathbf{A} = \mathbf{I}$
<i>normale</i>	si	$\mathbf{AA}^* = \mathbf{A}^* \mathbf{A}$

E.1.4 Rang d'une matrice

- Le *rang* d'une matrice $\mathbf{A}$ de type $[M, N]$ est la dimension des sous espaces ligne et colonne engendrés respectivement par les vecteurs lignes $\{\mathbf{a}_{i*}, i = 1, \dots, M\}$ et les vecteurs colonnes $\{\mathbf{a}_{*j}, j = 1, \dots, N\}$ de $\mathbf{A}$.
- Une matrice est dite de *rang plein* lorsque $\text{rang}[\mathbf{A}] = \min(M, N)$.
- $\text{rang}[\mathbf{A}]$ est l'ordre de sa plus grande sous-matrice non singulière.

$$\text{rang}[\mathbf{A}] = \text{rang}[\mathbf{A}^t]$$

- Deux matrices sont dites équivalentes si elle ont le même rang.
- Si $\mathbf{B}$ est une matrice de rang plein régulière, alors les matrices $\mathbf{BA}$ ou $\mathbf{AB}$ ont le même rang que la matrice $\mathbf{A}$:

$$\text{rang}[\mathbf{AB}] = \text{rang}[\mathbf{BA}] = \text{rang}[\mathbf{A}]$$

E.1.5 Matrice inverse

- Une matrice $\mathbf{A}$ est inversible s'il existe une matrice, notée $\mathbf{A}^{-1}$ et appelée matrice inverse de la matrice $\mathbf{A}$, telle que $\mathbf{A}\mathbf{A}^{-1} = \mathbf{A}^{-1}\mathbf{A} = \mathbf{I}$. Dans le cas contraire, on dit que la matrice est singulière.
- Si $\mathbf{A}$ et $\mathbf{B}$ sont des matrices inversibles on a :

$$[\mathbf{A}^t]^{-1} = [\mathbf{A}^{-1}]^t, \quad [\mathbf{A}^*]^{-1} = [\mathbf{A}^{-1}]^* \quad \text{et} \quad [\mathbf{AB}]^{-1} = \mathbf{B}^{-1}\mathbf{A}^{-1}$$

- Une CNS pour qu'une matrice carrée de dimensions $[N, N]$ soit inversible (ou régulière) est qu'elle soit de rang plein, c'est à dire $\text{rang}[\mathbf{A}] = N$.

$$\mathbf{y} = \mathbf{Ax}, \quad \mathbf{x} = \mathbf{A}^{-1}\mathbf{y}$$

$$(\mathbf{ABC})^{-1} = \mathbf{C}^{-1}\mathbf{B}^{-1}\mathbf{A}^{-1}$$

E.1.6 Déterminant d'une matrice

$$\det[\mathbf{A}] = |\mathbf{A}| = \sum_{j=1}^N [\mathbf{A}]_{ij} \mathbf{A}_{ij}^{\dagger}$$

$$\det[\mathbf{A}] = \det[\mathbf{A}^t]$$

$$\det[\mathbf{ABC}] = \det[\mathbf{A}] \det[\mathbf{B}] \det[\mathbf{C}]$$

E.1.7 Matrice singulière

$\mathbf{A}$ est singulière si $\det[\mathbf{A}] = 0$.

$$\mathbf{Ax} = 0, \quad \mathbf{A}^t \mathbf{x} = 0 \quad \forall \quad \mathbf{x} \neq 0.$$

E.1.8 Trace d'une matrice

La trace d'une matrice est définie par

$$\text{tr}[\mathbf{A}] = \sum_{i=1}^N a_{ii}$$

On a les relations suivantes:

$$\text{tr}[\mathbf{A}] = \sum_{i=1}^N a_{ii} = \sum_{i=1}^N \lambda_i$$

$$\text{tr}[\mathbf{ABC}] = \text{tr}[\mathbf{CAB}] = \text{tr}[\mathbf{BCA}]$$

$$\text{tr}[\mathbf{A} + \mathbf{B} + \mathbf{C}] = \text{tr}[\mathbf{A}] + \text{tr}[\mathbf{B}] + \text{tr}[\mathbf{C}]$$

La trace est invariante dans un changement de base.

E.1.9 Vecteurs propres et valeurs propre d'une matrice

- Un vecteur $\mathbf{u}$ est dit un *vecteur propre* de la matrice $\mathbf{A}$ si

$$\mathbf{A}\mathbf{u} = \lambda\mathbf{u};$$

λ est alors une *valeurs propre* associée au vecteur propre $\mathbf{u}$.

- Pour une matrice carrée on définit :
polynôme caractéristique: $f(\lambda) = \det [\mathbf{A} - \lambda\mathbf{I}]$
Matrice caractéristique: $\mathbf{A} - \lambda\mathbf{I}$
Équation caractéristique: $f(\lambda) = \det [\mathbf{A} - \lambda\mathbf{I} = 0]$
- Les *valeurs propres* $\{\lambda_i = \lambda_i(\mathbf{A}), i = 1, \dots, n\}$ d'une matrice $\mathbf{A}$ d'ordre n sont les n racines, réelles ou complexes, distinctes ou confondues, du polynôme caractéristique

$$f(\lambda) = \det [\mathbf{A} - \lambda\mathbf{I}] = (-\lambda)^n + \sum_{i=1}^N a_{ij} (-\lambda)^{n-1} + \dots + \det [\mathbf{A}] = 0$$

- Les v.p. d'une matrice ont une interprétation géométrique simple: Soit $\mathcal{E}^n$ une espace linéaire complexe à n dimensions et $\{\mathbf{e}_1, \dots, \mathbf{e}_n\}$ une base dans cet espace. A la matrice $\mathbf{A}$ et à cette base on peut associer un opérateur linéaire sous la forme de

$$\mathbf{A}\mathbf{e}_j = \sum_{i=1}^N a_{ji} \mathbf{e}_i, \quad j = 1, \dots, n$$

Ainsi les λ_i sont les v.p. de l'opérateur $\mathbf{A}$, c'est à dire qu'il existe un vecteur

$$\mathbf{u} = u_1\mathbf{e}_1 + \dots + u_i\mathbf{e}_i + \dots + u_n\mathbf{e}_n \neq \mathbf{0},$$

appartenant à l'espace $\mathcal{E}^n$ de façon à avoir $\mathbf{A}\mathbf{u} = \lambda\mathbf{u}$. Ce vecteur est appelé le vecteur propre de l'opérateur $\mathbf{A}$ correspondant à la v.p. λ .

- On montre facilement que

$$\det [\mathbf{A}] = \prod_{i=1}^n \lambda_i$$

Une matrice est régulière ssi $\lambda_i \neq 0, i = 1, \dots, n$.

- Le *spectre* de la matrice $\mathbf{A}$ est le sous-ensemble

$$\text{sp} [\mathbf{A}] = \cap_{j=1}^n \{\lambda_i(\mathbf{A})\}$$

du plan complexe.

- Le *rayon spectral* d'une matrice $\mathbf{A}$ est le nombre $\rho(\mathbf{A}) \geq 0$ défini par

$$\rho(\mathbf{A}) = \max_i \{|\lambda_i(\mathbf{A})|, \quad i = 1, \dots, n\}$$

- A tout valeur propre λ d'une matrice $\mathbf{A}$ est associé (au moins) un vecteur propre $\mathbf{v} \neq \mathbf{0}$ tel que: $\mathbf{A}\mathbf{v} = \lambda\mathbf{v}$.
- Si $\lambda \in \text{sp} [\mathbf{A}]$, le sous-espace vectoriel $\{\mathbf{v} \in \mathcal{V}; \mathbf{A}\mathbf{v} = \lambda\mathbf{v}\}$ est appelé sous-espace propre correspondant à la valeur propre λ .

E.1.10 Valeurs singulières d'une matrice

Les valeurs singulières de la matrice $\mathbf{A}$ de dimensions $[M, N]$ sont données par

$$\mu_i^2 \mathbf{A} \mathbf{A}^* \mathbf{u}_i = \mathbf{u}_i, \quad i = 1, \dots, r$$

$$\mu_i^2 \mathbf{A}^* \mathbf{A} \mathbf{v}_i = \mathbf{v}_i, \quad i = 1, \dots, r$$

où $r = \text{rang}[\mathbf{A}]$. Si on définit $\lambda_i = \frac{1}{(\mu_i)^2}$, $i = 1, \dots, r$ en supposant $\lambda_1 \geq \lambda_2 \geq \dots \geq 0$, et si on les complète par des zéros, c'est à dire $\{\lambda_i = 0, \quad i = r+1, \dots, \max(M, N)\}$, alors on peut associer deux systèmes de v.p. $\{\mathbf{u}_i, \quad i = 1, \dots, M\}$ et $\{\mathbf{v}_j, \quad j = 1, \dots, N\}$ aux $\{\lambda_i\}$, tels que

$$\mathbf{A} \mathbf{u}_i = \lambda_i \mathbf{v}_i$$

$$\mathbf{A} \mathbf{v}_i = \lambda_i \mathbf{u}_i$$

$$\mathbf{A} \mathbf{A}^* \mathbf{u}_i = \lambda_i^2 \mathbf{u}_i, \quad i = 1, \dots, M$$

$$\mathbf{A}^* \mathbf{A} \mathbf{v}_i = \lambda_i^2 \mathbf{v}_i, \quad j = 1, \dots, N.$$

$\lambda_i^2 \mathbf{u}_i$ et $\lambda_i^2 \mathbf{v}_i$ constituent les systèmes caractéristiques de $\mathbf{A} \mathbf{A}^*$ et $\mathbf{A}^* \mathbf{A}$, respectivement. La matrice $\mathbf{A}$ peut se factoriser en le produit

$$\mathbf{A} = \mathbf{U} \mathbf{\Lambda} \mathbf{V}^t$$

où $\mathbf{U}$ est la matrice dont les colonnes sont les v.p. $\mathbf{u}_j$, $\mathbf{V}$ est la matrice dont les colonnes sont les v.p. $\mathbf{v}_i$ et $\mathbf{\Lambda}$ est la matrice diagonale formée de v.p. λ_i .

Le rang r de la matrice normale $\mathbf{A}^* \mathbf{A}$ est inférieur ou égale à $\inf(M, N)$. On complète les matrices $\mathbf{U}$ et $\mathbf{V}$ par les $(M - r)$ v.p. $\mathbf{u}_i$ et les $(N - r)$ v.p. $\mathbf{v}_i$ qui correspondent aux v.p. nulles, de façon à obtenir des bases complètes respectivement de Y et de X .

E.1.11 Matrices et les opérations linéaires

Si commençant par une base $\{\mathbf{e}_1, \dots, \mathbf{e}_n\}$, on introduit dans un espace euclidienne $\mathcal{E}^n$ un produit scalaire sous forme de

$$\langle \mathbf{x}, \mathbf{y} \rangle = x_1 y_1^* + \dots + x_i y_i^* + \dots + x_n y_n^*$$

alors une matrice hermitienne $\mathbf{A} : ([\mathbf{A}]_{ji} = [\mathbf{A}]_{ij}^*)$ induit sur cette base un opérateur hermitienne, c'est à dire

$$\langle \mathbf{A} \mathbf{x}, \mathbf{y} \rangle = \langle \mathbf{x}, \mathbf{A} \mathbf{y} \rangle \quad \forall \mathbf{x}, \mathbf{y} \in \mathcal{E}^n$$

– Le produit scalaire a les propriétés suivantes:

$$\forall \mathbf{x}, \mathbf{x}_1, \mathbf{x}_2, \mathbf{y} \quad \text{et} \quad \mathbf{y}_2 \in \mathcal{E}^n$$

$$\langle \mathbf{x}, \mathbf{x} \rangle > 0 \quad \forall \mathbf{x} \neq 0$$

$$\langle (\mathbf{x}_1 + \mathbf{x}_2), \mathbf{x} \rangle = \langle \mathbf{x}_1, \mathbf{x} \rangle + \langle \mathbf{x}_2, \mathbf{x} \rangle$$

$$\langle \alpha \mathbf{x}, \mathbf{y} \rangle = \alpha \langle \mathbf{x}, \mathbf{y} \rangle$$

$$\langle \mathbf{x}, \mathbf{y} \rangle = \langle \mathbf{y}, \mathbf{x} \rangle^*$$

- Si $\mathbf{A}$ est une matrice hermitienne on a

$$\langle \mathbf{A}\mathbf{x}, \mathbf{x} \rangle = \sum_i a_{ii} x_i x_i^*$$

Un telle forme est entièrement déterminée par une matrice hermitienne $\mathbf{A}$. L'ordre N et le rang de la matrice $\mathbf{A}$ sont appelés respectivement l'ordre et le rang de la forme $\langle \mathbf{A}\mathbf{x}, \mathbf{x} \rangle$.

- Une tâche importante de la théorie des formes hermitiennes est la réduction de la forme hermitienne à la somme de carrés suivantes

$$\langle \mathbf{A}\mathbf{x}, \mathbf{x} \rangle = \sum_k \lambda_k |x_k|^2$$

où $x_k = c_{1k}x_1 + c_{2k}x_2 + \dots$

Si r est le rang de $\mathbf{A}$ et x_i sont linéairement indépendantes, alors seul r coefficients λ_k dans cette équation son non-nulle.

- Soient $\mathcal{V}$ un espace vectoriel de dimension finie n , et $\mathcal{A} : \mathcal{V} \rightarrow \mathcal{V}$ une application linéaire, représentée par une matrice carrée $\mathbf{A}$ relativement à une base $\{\mathbf{e}_i, i = 1, \dots, n\}$. Cette même application, relativement à une autre base $\{\mathbf{f}_j, j = 1, \dots, n\}$, est représentée par la matrice

$$\mathbf{B} = \mathbf{P}^{-1} \mathbf{A} \mathbf{P},$$

où $\mathbf{P}$ est la matrice inversible dont le j -ème vecteur colonne est formé des composantes du vecteur $\mathbf{f}_j$ dans la base $\{\mathbf{e}_i, i = 1, \dots, n\}$. La matrice $\mathbf{P}$ est appelée *matrice de passage de la base* $\{\mathbf{e}_j, j = 1, \dots, n\}$ dans la base $\{\mathbf{f}_j, j = 1, \dots, n\}$. Les deux matrices $\mathbf{A}$ et $\mathbf{B}$ sont dites *semblables* s'il existe une matrice inversible $\mathbf{P}$ telle que $\mathbf{B} = \mathbf{P}^{-1} \mathbf{A} \mathbf{P}$.

- Une même application linéaire $\mathbf{A}$ étant représentée par différentes matrices selon la base choisie, le problème de la réduction des matrices consiste alors de trouver une base vis-à-vis de laquelle la matrice représentant l'application soit aussi simple que possible. Le cas le plus favorable est celui où il existe une matrice inversible $\mathbf{P}$ telle que la matrice $\mathbf{P}^{-1} \mathbf{A} \mathbf{P}$ soit diagonale. Dans ce cas on dit que la matrice $\mathbf{A}$ est *diagonalisable*. Les éléments diagonaux de la matrice $\mathbf{P}^{-1} \mathbf{A} \mathbf{P}$ sont alors les valeurs propres $\lambda_1, \lambda_2, \dots, \lambda_n$ de la matrice $\mathbf{A}$, et le j -ème vecteur colonne $\mathbf{p}_j$ de la matrice $\mathbf{P}$ est formé des composantes d'un vecteur propre correspondant à λ_j , et on a :

$$\mathbf{P}^{-1} \mathbf{A} \mathbf{P} = \text{diag}[\lambda_i] \rightarrow \mathbf{A} \mathbf{p}_j = \lambda_j \mathbf{p}_j, \quad j = 1, \dots, n$$

Autrement dit, une matrice est diagonalisable ssi il existe une base de vecteur propres.

- Étant donnée une matrice carrée $\mathbf{A}$, il existe une matrice unitaire $\mathbf{U}$ telle que la matrice $\mathbf{U}^{-1}\mathbf{A}\mathbf{U}$ soit triangulaire.
- Étant donnée une matrice normale $\mathbf{A}$, il existe une matrice unitaire $\mathbf{U}$ telle que la matrice $\mathbf{U}^{-1}\mathbf{A}\mathbf{U}$ soit diagonale.
- Étant donnée une matrice symétrique $\mathbf{A}$, il existe une matrice orthogonale $\mathbf{O}$ telle que la matrice $\mathbf{O}^{-1}\mathbf{A}\mathbf{O}$ soit diagonale.
- On appelle *valeurs singulières* d'une matrice $\mathbf{A}$ carrée les racines carrées positives des valeurs propres de la matrice hermitienne $\mathbf{A}^*\mathbf{A}$ (ou $\mathbf{A}^t\mathbf{A}$ si la matrice $\mathbf{A}$ est réelle). Les valeurs singulières d'une matrice sont toujours ≥ 0 . Elle sont toutes > 0 ssi la matrice $\mathbf{A}$ est inversible.
- Deux matrices $\mathbf{A}$ et $\mathbf{B}$ de type $[m,n]$ sont dites *équivalentes* s'il existe une matrice inversible $\mathbf{Q}$ d'ordre m et une matrice inversible $\mathbf{P}$ d'ordre n telle que $\mathbf{B} = \mathbf{QAP}$.
- Pour une matrice réelle carrée $\mathbf{A}$, il existe deux matrices orthogonales $\mathbf{U}$ et $\mathbf{V}$ telle que

$$\mathbf{U}^t\mathbf{A}\mathbf{U} = \text{diag}[\mu_i].$$

- Pour une matrice complexe carrée $\mathbf{A}$, il existe deux matrices unitaires $\mathbf{U}$ et $\mathbf{V}$ telle que

$$\mathbf{U}^*\mathbf{A}\mathbf{U} = \text{diag}[\mu_i].$$

- Dans les deux cas, les nombres $\mu_i \geq 0$ sont les valeurs singulières de la matrice $\mathbf{A}$.
- Toutes les valeurs propres d'une matrice hermitienne sont réelles. Toute matrice hermitienne est diagonalisable et la matrice de passage sont unitaires.
- Soit $\mathbf{A}$ une matrice carrée représentant une application linéaire d'un espace $\mathcal{V}$ sur le corps $\mathcal{C}$, muni de son produit *canonique*.
- Le *quotient de Rayleigh* de la matrice $\mathbf{A}$ est l'application :

$$R_A : \{\mathbf{V} - \{0\}\} \longrightarrow \mathcal{C}$$

définie par :

$$R_A(\mathbf{v}) = \frac{\langle \mathbf{v}, \mathbf{A}\mathbf{v} \rangle}{\langle \mathbf{v}, \mathbf{v} \rangle} = \frac{\mathbf{v}^* \mathbf{A} \mathbf{v}}{\mathbf{v}^* \mathbf{v}}, \quad \mathbf{v} \neq 0$$

E.1.12 Décomposition tronquée en valeurs singulière

Si les λ_i décroissent rapidement alors on peut approcher la matrice $\mathbf{A}$ par une décomposition tronquée en valeurs singulières ($k < r = \text{rang}[\mathbf{A}]$) :

$$\mathbf{A} = \sum_{i=1}^k \lambda_i \mathbf{u}_i \mathbf{v}_i^* + \mathbf{E}_k$$

où $\mathbf{E}_k$ est la matrice d'erreur de l'approximation. Si on définit J_k comme l'énergie de l'erreur par

$$J_k = \sum_{i=1}^N \sum_{j=1}^N |[\mathbf{E}_k]_{ij}|^2 = \sum_{i=1}^N \sum_{j=1}^N [\mathbf{E}_k]_{ij} [\mathbf{E}_k]_{ij}^*$$

alors on montre facilement que l'on a

$$J_k = \sum_{i=k+1}^r \lambda_i$$

Si on définit α_k par

$$\alpha_k = \frac{J_k}{J_0} = 1 - \frac{\sum_{i=1}^k \lambda_i}{\sum_{i=1}^r \lambda_i}$$

alors $0 \leq \alpha_k \leq 1$, $\alpha_0 = 1$, et $\alpha_r = 0$.

E.1.13 Inverse généralisée d'une matrice

Considérons l'équation

$$\mathbf{y} = \mathbf{A}\mathbf{x}$$

où $\mathbf{A}$ est une matrice de dimension $[m, n]$.

– **Définition de Moore :**

Une matrice $\mathbf{G}$ de dimension $[n, m]$ est dite inverse généralisée de la matrice $\mathbf{A}$ si

$$\mathbf{A}\mathbf{G} = \mathbf{P}_A \quad \text{et} \quad \mathbf{G}\mathbf{A} = \mathbf{P}_G$$

où $\mathbf{P}_A$ est l'opérateur (matrice) de projection des vecteurs de l'espace $\mathcal{E}^m$ sur $\mu(\mathbf{A})$ et $\mathbf{P}_G$ est l'opérateur (matrice) de projection des vecteurs de l'espace $\mathcal{E}^n$ sur $\mu(\mathbf{G})$.

– **Définition de Penrose :**

Une matrice $\mathbf{G}$ de dimension $[n, m]$ est dite inverse généralisée de la matrice $\mathbf{A}$ si

$$\mathbf{A}\mathbf{G}\mathbf{A} = \mathbf{A}, \quad (\mathbf{A}\mathbf{G})^* = \mathbf{A}\mathbf{G}, \quad \mathbf{G}\mathbf{A}\mathbf{G} = \mathbf{G} \quad \text{et} \quad (\mathbf{G}\mathbf{A})^* = \mathbf{G}\mathbf{A}$$

– Ces définitions sont identiques si les normes associées aux espaces $\mathcal{E}^m$ et $\mathcal{E}^n$ sont du type euclidien: $\|\mathbf{x}\| = \langle \mathbf{x}^*, \mathbf{x} \rangle^{1/2}$.

E.2 Quelques matrices célèbres

Matrices identité :

$$[\mathbf{I}]_{ij} = \delta_{ij} = \begin{cases} 1 & i = j \\ 0 & \text{ailleurs} \end{cases}, \quad j = 1, \dots, N$$

$$\mathbf{I}^t \mathbf{I} = \mathbf{I} \mathbf{I}^t = \mathbf{I}, \quad \mathbf{I} \mathbf{A} = \mathbf{A} \mathbf{I} = \mathbf{A}$$

Matrices diagonales :

$$\mathbf{D} = \text{diag}[d_1, \dots, d_N] \longrightarrow [\mathbf{D}]_{ij} = d_i \delta_{ij}$$

$$\mathbf{D}^{-1} = \text{diag} \left[\frac{1}{d_1}, \dots, \frac{1}{d_N} \right]$$

Matrices d'inversion d'ordre

$$\mathbf{J} = \begin{pmatrix} \cdot & \cdot & \cdot & \cdot & 1 \\ \cdot & 0 & \cdot & 1 & \cdot \\ \cdot & \cdot & 1 & \cdot & \cdot \\ \cdot & 1 & \cdot & 0 & \cdot \\ 1 & \cdot & \cdot & \cdot & \cdot \end{pmatrix} \quad [\mathbf{J}]_{ij} = \begin{cases} 1 & i = N - j + 1, \quad j = 1, \dots, N \\ 0 & \text{ailleurs} \end{cases}$$

$$[\mathbf{J}]^{2k} = \mathbf{I}, \quad [\mathbf{J}]^{2k+1} = \mathbf{J}, \quad \mathbf{J}^{-1} = \mathbf{J}$$

- $\mathbf{AJ}$ inverse l'ordre des éléments des lignes de la matrice $\mathbf{A}$.
- $\mathbf{JA}$ inverse l'ordre des éléments des colonnes de la matrice $\mathbf{A}$.
- $\mathbf{JAJ}$ inverse l'ordre des éléments des lignes et des colonnes de la matrice $\mathbf{A}$.
- $\mathbf{Jh}$ inverse l'ordre des éléments du vecteur $\mathbf{h}$.
- Si $\mathbf{h}$ est un vecteur symétrique on a $\mathbf{Jh} = \mathbf{h}$.
- Si $\mathbf{h}$ est un vecteur anti-symétrique on a $\mathbf{Jh} = -\mathbf{h}$.

Matrices symétriques :

$$[\mathbf{S}]_{ij} = [\mathbf{S}]_{ji}, \quad i, j = 1, \dots, N, \quad \mathbf{S}^t = \mathbf{S}$$

- Si $\mathbf{S}$ est une matrice symétrique avec des éléments réels, alors ses v.p. sont réelles et il existe au moins une matrice $\mathbf{R}$ permettant d'avoir: $\mathbf{R}^t \mathbf{S} \mathbf{R} = \mathbf{D}$, $\mathbf{S} = \mathbf{R} \mathbf{D} \mathbf{R}^t$, où $\mathbf{D}$ est une matrice diagonale.
- Étant donnée une matrice symétrique $\mathbf{A}$, il existe une matrice orthogonale $\mathbf{O}$ telle que la matrice $\mathbf{O}^{-1} \mathbf{A} \mathbf{O}$ soit diagonale.

Matrices anti-symétriques :

$$[\mathbf{R}]_{ij} = -[\mathbf{R}]_{ji}, \quad i, j = 1, \dots, N, \quad i \neq j \quad \mathbf{R}^t = -\mathbf{R}$$

Matrices symétriques par rapport à la deuxième diagonale :

$$[\mathbf{S}]_{ij} = [\mathbf{S}]_{(N-j+1)(N-i+1)}, \quad i, j = 1, \dots, N, \quad \mathbf{J} \mathbf{S} \mathbf{J} = \mathbf{S}^t, \quad \mathbf{J} \mathbf{S}^t \mathbf{J} = \mathbf{S}$$

Matrices anti-symétriques par rapport à la deuxième diagonale :

$$[\mathbf{R}]_{ij} = -[\mathbf{R}]_{(N-i+1)(N-j+1)}, \quad i, j = 1, \dots, N, \quad i \neq j, \quad \mathbf{J} \mathbf{S} \mathbf{J} = -\mathbf{S}^t$$

Matrices hermitienne :

$$[\mathbf{H}]_{ij} = [\mathbf{H}^*]_{ji}, \quad i, j = 1, \dots, N, \quad \mathbf{H}^* = \mathbf{H}^*$$

$[\mathbf{H}]_{ii}$ sont réels. Si $\mathbf{H}$ est une matrice hermitienne $\mathbf{H}^{-1}$ (à condition qu'elle existe) est aussi hermitienne, et on a $\mathbf{x}^t \mathbf{H} \mathbf{x} = \alpha \|\mathbf{x}\|^2$ est réel. $\mathbf{A}$ toute matrice hermitienne on peut toujours associer un système de vecteurs propres orthonormés complet. Les v.p. de $\mathbf{H}$ sont réelles et ses vecteurs propres sont orthogonaux.

Matrices anti-hermitiennes :

$$[\mathbf{K}]_{ij} = -[\mathbf{K}^*]_{ji}, \quad i, j = 1, \dots, N, \quad \mathbf{K}^* = -\mathbf{K}^*$$

Si $\mathbf{K}$ est une matrice anti-hermitienne $\mathbf{K}^{-1}$ (à condition qu'elle existe) est aussi anti-hermitienne, et on a $\mathbf{x}^t \mathbf{K} \mathbf{x} = -j\alpha \|\mathbf{x}\|^2$ est imaginaire. Les v.p. de $\mathbf{K}$ sont imaginaires et ses vecteurs propres sont orthogonaux. Une matrice quelconque $\mathbf{A}$ peut toujours être décomposée de la façon suivante :

$$\mathbf{A} = \frac{1}{2}[\mathbf{A} + \mathbf{A}^t] + \frac{1}{2}[\mathbf{A} - \mathbf{A}^t] = \frac{1}{2}[\mathbf{H} + \mathbf{K}].$$

Matrices idempotentes :

$$\mathbf{Q}^t \mathbf{Q} = \mathbf{Q}$$

Si $\mathbf{A}$ est une matrice idempotente et symétrique on a :

$$\mathbf{A}^2 = \mathbf{A}^t \mathbf{A} = \mathbf{A}$$

$$\mathbf{A}^2 \mathbf{x}_i = \mathbf{A} \mathbf{x}_i = \lambda_i \mathbf{x}_i = \mathbf{A}(\lambda_i \mathbf{x}_i) = \lambda_i^2 \mathbf{x}_i \longrightarrow \lambda_i = 0 \text{ ou } 1.$$

Matrices orthogonales :

$$\mathbf{O}^t \mathbf{O} = \mathbf{O} \mathbf{O}^t = \mathbf{O} \mathbf{O}^{-1} = \mathbf{O}^{-1} \mathbf{O} = \mathbf{I}$$

Deux matrices $\mathbf{A}$ et $\mathbf{B}$ sont *similaires* s'il existe une matrice $\mathbf{C}$ telle que

$$\mathbf{B} = \mathbf{C}^{-1} \mathbf{A} \mathbf{C}$$

On a alors : $\mathbf{A} = \mathbf{C}^{-1} \mathbf{B} \mathbf{C}$. Deux matrices $\mathbf{A}$ et $\mathbf{B}$ sont *orthogonalement similaires* ou *congruentes* s'il existe une matrice orthogonale $\mathbf{O}$ telle que $\mathbf{B} = \mathbf{O}^{-1} \mathbf{A} \mathbf{O}$. On a alors $\mathbf{A} = \mathbf{O}^{-1} \mathbf{B} \mathbf{O}$.

Matrices de décalages :

$\mathbf{Z}$ est une matrices de décalage vers le bas (ou vers la droite) si :

$$\mathbf{Z} = \begin{pmatrix} 0 & . & . & . & 0 \\ 1 & 0 & . & . & 0 \\ 0 & 1 & 0 & . & 0 \\ . & 0 & 1 & 0 & 0 \\ . & . & 0 & 1 & 0 \end{pmatrix}, \quad [\mathbf{Z}]_{ij} = \begin{cases} 1 & i = j + 1 \\ 0 & \text{ailleurs} \end{cases}, \quad j = 1, \dots, N - 1$$

$\mathbf{Z}^t$ est une matrices de décalage vers le haut.

Si $\mathbf{Z}$ est une matrice de décalage vers le haut on a :

$$\begin{aligned} \mathbf{B} &= \mathbf{Z}^\alpha \mathbf{A} \mathbf{Z}^\beta & [\mathbf{B}]_{mn} &= [\mathbf{A}]_{m+\alpha n-\beta} \\ \mathbf{B} &= \mathbf{Z}^\alpha \mathbf{A} \mathbf{Z}^{t\beta} & [\mathbf{B}]_{mn} &= [\mathbf{A}]_{m+\alpha n+\beta} \\ \mathbf{B} &= \mathbf{Z}^{t\alpha} \mathbf{A} \mathbf{Z}^\beta & [\mathbf{B}]_{mn} &= [\mathbf{A}]_{m-\alpha n-\beta} \\ \mathbf{B} &= \mathbf{Z}^{t\alpha} \mathbf{A} \mathbf{Z}^{t\beta} & [\mathbf{B}]_{mn} &= [\mathbf{A}]_{m-\alpha n+\beta} \end{aligned}$$

Matrices de décalage circulaire :

$\mathbf{P}_0$ est une matrice de décalage circulaire vers le haut (ou vers la gauche) si :

$$\mathbf{P}_0 = \begin{pmatrix} 0 & 1 & 0 & . & 0 \\ 0 & 0 & 1 & 0 & . \\ 0 & . & 0 & 1 & 0 \\ . & . & . & 0 & 1 \\ 1 & . & . & . & 0 \end{pmatrix}, \quad [\mathbf{P}_0]_{ij} = \begin{cases} 1 & i = 1, \dots, N-1 \\ 1 & j = 1, i = N \\ 0 & \text{ailleurs} \end{cases}$$

$\mathbf{P}_0$ est une matrice circulante. Elle peut donc être factorisée de façon suivante:

$$\mathbf{P}_0 = \mathbf{W} \mathbf{D} \mathbf{W}^{-1}$$

où $\mathbf{D}$ est une matrice diagonale et les éléments des matrices $\mathbf{W}$ et $\mathbf{W}^{-1}$ sont données par:

$$[\mathbf{W}]_{pq} = \exp \left[\frac{2\pi}{N} (p-1)(q-1) \right], \quad p, q = 1, \dots, N$$

$$[\mathbf{W}^{-1}]_{pq} = \frac{1}{N} \exp \left[\frac{-2\pi}{N} (p-1)(q-1) \right], \quad p, q = 1, \dots, N$$

$\mathbf{P}_1 = \mathbf{P}_0^t$ est une matrices de décalage circulaire vers le bas (ou vers la droite).

$$\mathbf{P}_1^{-1} = \mathbf{P}_1^t = \mathbf{P}_0, \quad \mathbf{P}_0^{-1} = \mathbf{P}_0^t = \mathbf{P}_1$$

Matrices triangulaires supérieure :

$$\mathbf{U} = \begin{pmatrix} u_{11} & u_{12} & u_{13} & u_{14} & u_{15} \\ 0 & u_{22} & u_{23} & u_{24} & u_{25} \\ 0 & 0 & u_{33} & u_{34} & u_{35} \\ 0 & 0 & 0 & u_{44} & u_{45} \\ 0 & 0 & 0 & 0 & u_{55} \end{pmatrix}$$

$$\mathbf{U}_1 = \begin{pmatrix} 1 & u_{12} & u_{13} & u_{14} & u_{15} \\ 0 & 1 & u_{23} & u_{24} & u_{25} \\ 0 & 0 & 1 & u_{34} & u_{35} \\ 0 & 0 & 0 & 1 & u_{45} \\ 0 & 0 & 0 & 0 & 1 \end{pmatrix}, \quad \mathbf{U}_0 = \begin{pmatrix} 0 & u_{12} & u_{13} & u_{14} & u_{15} \\ 0 & 0 & u_{23} & u_{24} & u_{25} \\ 0 & 0 & 0 & u_{34} & u_{35} \\ 0 & 0 & 0 & 0 & u_{45} \\ 0 & 0 & 0 & 0 & 0 \end{pmatrix}$$

L'inverse d'une matrice du type $\mathbf{U}_1$ est aussi une matrice du type $\mathbf{U}_1$.

$$\det[\mathbf{U}] = \prod_{i=1}^n u_{ii}, \quad \det[\mathbf{U}_1] = 1, \quad \det[\mathbf{U}_0] = 0$$

Matrices triangulaires inférieure :

$$\mathbf{L} = \begin{pmatrix} l_{11} & 0 & 0 & 0 & 0 \\ l_{21} & l_{22} & 0 & 0 & 0 \\ l_{31} & l_{32} & l_{33} & 0 & 0 \\ l_{41} & l_{42} & l_{43} & l_{44} & 0 \\ l_{51} & l_{52} & l_{53} & l_{54} & l_{55} \end{pmatrix}$$

$$\mathbf{L}_1 = \begin{pmatrix} 1 & 0 & 0 & 0 & 0 \\ l_{21} & 1 & 0 & 0 & 0 \\ l_{31} & l_{32} & 1 & 0 & 0 \\ l_{41} & l_{42} & l_{43} & 1 & 0 \\ l_{51} & l_{52} & l_{53} & l_{54} & 1 \end{pmatrix} \quad \mathbf{L}_0 = \begin{pmatrix} 0 & 0 & 0 & 0 & 0 \\ l_{21} & 0 & 0 & 0 & 0 \\ l_{31} & l_{32} & 0 & 0 & 0 \\ l_{41} & l_{42} & l_{43} & 0 & 0 \\ l_{51} & l_{52} & l_{53} & l_{54} & 0 \end{pmatrix}$$

L'inverse d'une matrice du type $\mathbf{L}_1$ est aussi une matrice du type $\mathbf{L}_1$.

$$\det[\mathbf{L}] = \prod_{i=1}^n l_{ii}, \quad \det[\mathbf{L}_1] = 1, \quad \det[\mathbf{L}_0] = 0$$

Matrices de Hessenberg supérieure :

$$\mathbf{H} = \begin{pmatrix} h_{11} & h_{12} & h_{13} & h_{14} & h_{15} \\ h_{12} & h_{22} & h_{23} & h_{24} & h_{25} \\ 0 & h_{23} & h_{33} & h_{34} & h_{35} \\ 0 & 0 & h_{34} & h_{44} & h_{45} \\ 0 & 0 & 0 & h_{45} & h_{55} \end{pmatrix}$$

Matrices de Hessenberg inférieure :

$$\mathbf{H} = \begin{pmatrix} h_{11} & h_{12} & 0 & 0 & 0 \\ h_{21} & h_{22} & h_{23} & 0 & 0 \\ h_{31} & h_{32} & h_{33} & h_{34} & 0 \\ h_{41} & h_{42} & h_{43} & h_{44} & h_{45} \\ h_{51} & h_{52} & h_{53} & h_{54} & h_{55} \end{pmatrix}$$

Matrices tridiagonale :

$$\mathbf{B} = \begin{pmatrix} b_{11} & b_{12} & 0 & 0 & 0 \\ b_{21} & b_{22} & b_{23} & 0 & 0 \\ 0 & b_{32} & b_{33} & b_{34} & 0 \\ 0 & 0 & b_{43} & b_{44} & b_{45} \\ 0 & 0 & 0 & b_{54} & b_{55} \end{pmatrix}$$

Matrice de transformation de Householder

Définition

Soit $\mathbf{v}$ un vecteur complexe de dimension $[N]$. Alors $\mathbf{v}^{*t}\mathbf{v}$ est un scalaire réel et $\mathbf{v}\mathbf{v}^{*t}$ est une matrice carrée de dimensions $[N, N]$. Soit $\mathbf{I}$ une matrice identité de mêmes dimensions. Alors la matrice

$$\mathbf{H}\mathbf{v} = \mathbf{I} - \frac{2}{\mathbf{v}^*\mathbf{v}}[\mathbf{v}\mathbf{v}^*]$$

est une matrice de réflexion ou matrice de Householder.

Propriétés :

- $\mathbf{H}\mathbf{v}$ est hermitienne et orthonormale.

- La transformation $\mathbf{y} = \mathbf{H}\mathbf{x}$ avec $\mathbf{H}$ une matrice de Householder est appelée la transformation de Householder. Une telle transformation conserve l'énergie du vecteur original $\mathbf{x}$. (Ceci est vrai pour n'importe quelle transformation orthonormale.)
- Soit $\mathbf{e}_j$ un vecteur base ($[\mathbf{e}_j]_i = \delta_{ij}$), alors pour un vecteur quelconque $\mathbf{u}$ il existe une matrice de Householder $\mathbf{H}$ telle que

$$\mathbf{H}\mathbf{u} = \sigma \mathbf{e}_j$$

ssi σ est choisi telle que

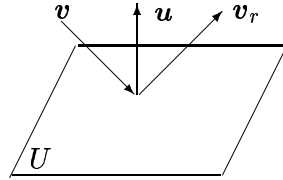
$$\mathbf{u}^* \mathbf{u} = |\sigma|^2 = \sum_{i=1}^N u_i u_i^* \quad \text{ou bien} \quad \sigma = \pm \frac{u_j}{|u_j|} \sqrt{\mathbf{u}^* \mathbf{u}}$$

On a aussi $\mathbf{u}^* \mathbf{H} = \sigma^* \mathbf{e}_j^t$

- Soit $\mathbf{v}$ un vecteur, et $\mathbf{v}_r$ son image à travers un hyperplan U sous-espace orthogonale au vecteur $\mathbf{u}$ (symétrie plane). La matrice de transformation de Householder $\mathbf{H}_u$ est une matrice qui nous permet de faire correspondre $\mathbf{v}_r$ à $\mathbf{v}$ par

$$\mathbf{v}_r = \mathbf{H}_u \mathbf{v}$$

où $\mathbf{H}_u = \mathbf{I} - \beta \mathbf{u} \mathbf{u}^t$ avec $\beta = \frac{2}{(\mathbf{u}^t \mathbf{u})}$



Matrice de transformation de Householder hyperbolique

Soit Φ une matrice diagonale avec des éléments diagonaux de $+1$ et 1 . Soit $\mathbf{v}$ est un vecteur complexe. Alors

$$\mathbf{v}^* \Phi \mathbf{v} = \sum_{i=1}^N |v_i|^2 [\Phi]_{ii}$$

est la norme hyperbolique du vecteur $\mathbf{v}$. Une matrice $\mathbf{Q}$ qui satisfait

$$\mathbf{Q} \Phi \mathbf{Q}^* = \Phi$$

est une matrice hypernormale. Une telle matrice préserve la norme hyperbolique des vecteurs, c'est à dire: si $\mathbf{y}^t = \mathbf{x}^t \mathbf{Q}$ alors $\mathbf{y}^t \mathbf{Q} \mathbf{y} = \mathbf{x}^t \mathbf{Q} \mathbf{x}$. Une matrice définie par

$$\mathbf{H} = \Phi - 2 \frac{\mathbf{Q} \mathbf{Q}^t}{\mathbf{Q}^t \Phi \mathbf{Q}}$$

est appelée une matrice de transformation hyperbolique de Householder. $\mathbf{H}$ est hermitienne et hypernormale par rapport à la matrice Φ , c'est à dire $\mathbf{H} \Phi \mathbf{H}^t = \Phi$.

Propriétés :

- $\det [\mathbf{H}] = 1$; $\mathbf{H}$ est une matrice non-singulière.

- $\mathbf{H}$ n'est, en général, pas symétrique, mais elle a une symétrie hyperbolique, c'est à dire $\mathbf{H}^t \mathbf{A} \mathbf{H} = \mathbf{H} \mathbf{A} \mathbf{H}^t$.
- Les valeurs propres d'une matrice hypernormale forment des paires réciproquement conjuguées, c'est à dire que si λ est une valeur propre de cette matrice $\frac{1}{\lambda^*}$ l'est aussi. De plus si λ est multiple d'ordre n alors $\frac{1}{\lambda^*}$ est aussi de même ordre n .

Décomposition orthogonale par transformation de Householder

Soit $\mathbf{A}$ une matrice de dimensions $[N, M]$ et de rang $N \leq M$. Il existe une matrice orthogonale $\mathbf{H}_u$ de dimensions $[N, M]$ telle que $\mathbf{S} = \mathbf{H}_u \mathbf{A}$ soit triangulaire supérieure. La matrice $\mathbf{A}$ peut être décomposée en le produit d'une matrice orthogonale $\mathbf{H}_u$ et d'une matrice triangulaire supérieure $\mathbf{S}$ de rang plein — $\mathbf{A} = \mathbf{H}_u \mathbf{S}$.

Matrices centro-symétriques

$$[\mathbf{G}]_{ij} = [\mathbf{G}]_{ji} = [\mathbf{G}]_{N+1-i, N+1-j}, \quad i, j = 1, \dots, N$$

Exemple :

$$\mathbf{G} = \begin{pmatrix} g_{11} & g_{12} & g_{13} & g_{14} & g_{15} \\ g_{12} & g_{22} & g_{23} & g_{24} & g_{14} \\ g_{13} & g_{23} & g_{33} & g_{23} & g_{13} \\ g_{14} & g_{24} & g_{23} & g_{22} & g_{12} \\ g_{15} & g_{14} & g_{13} & g_{12} & g_{11} \end{pmatrix}$$

– Si $\mathbf{G}$ est une matrice centro-symétrique on a :

$$\mathbf{J}\mathbf{G} = \mathbf{G}, \quad \mathbf{G}\mathbf{J} = \mathbf{G}, \quad \mathbf{J}\mathbf{G}\mathbf{J} = \mathbf{G}, \quad \text{et} \quad \mathbf{J}\mathbf{G}^{-1}\mathbf{J} = \mathbf{G}^{-1}$$

– Cette dernière relation montre que l'inverse d'une matrice centro-symétrique est aussi une matrice centro-symétrique.

– Si $\mathbf{G}$ est une matrice centro-symétrique de dimension $[N, N]$, alors elle peut être partitionnée de la manière suivante :

– si N est pair ($N = 2r$) on a :

$$\mathbf{G} = \begin{pmatrix} \mathbf{A} & \mathbf{B} \\ \mathbf{J}\mathbf{B}^*\mathbf{J} & \mathbf{J}\mathbf{A}^*\mathbf{J} \end{pmatrix}$$

– si N est impair ($N = 2r + 1$) on a :

$$\mathbf{G} = \begin{pmatrix} \mathbf{A} & \mathbf{x} & \mathbf{B} \\ \mathbf{x}^* & \alpha & \mathbf{x}^*\mathbf{J} \\ \mathbf{J}\mathbf{B}^*\mathbf{J} & \mathbf{J}\mathbf{x} & \mathbf{J}\mathbf{A}^*\mathbf{J} \end{pmatrix}$$

où $\mathbf{A}$ et $\mathbf{B}$ sont deux matrices de dimensions $[r, r]$, $\mathbf{J}$ est une matrice d'inversion d'ordre de dimensions $[r, r]$, $\mathbf{x}$ est un vecteur de dimension r et α est un scalaire.

– Si $\mathbf{G}$ est une matrice centro-symétrique de dimension $[N, N]$, alors

– si N est pair on a

$$\mathbf{G} = \begin{pmatrix} \mathbf{A} & \mathbf{C}^t \\ \mathbf{C} & \mathbf{J}\mathbf{A}\mathbf{J} \end{pmatrix}$$

où $\mathbf{A}, \mathbf{C}$ et $\mathbf{J}$ sont des matrices de dimensions $[N/2, N/2]$ et

$$\mathbf{A}^t = \mathbf{A} \quad \text{et} \quad \mathbf{C}^t = \mathbf{J}\mathbf{C}\mathbf{J}$$

On montre aussi que $\mathbf{G}$ est orthogonalement similaire à et on montre que si $\mathbf{G}$ a des valeurs propres (v.p.) distinctes, elle a $[N/2]$ vecteurs propres symétriques ($\mathbf{h}_i, i = 1, \dots, N/2$) et $[N/2]$ vecteurs propres v.p. anti-symétriques ($\mathbf{g}_i, i = 1, \dots, N/2$).

$$\begin{pmatrix} \mathbf{A} - \mathbf{J}\mathbf{C} & 0 \\ 0 & \mathbf{A} + \mathbf{J}\mathbf{C} \end{pmatrix}$$

Dans ce cas les v.p. symétriques $\mathbf{h}_i$ de $\mathbf{G}$ s'obtiennent par

$$\mathbf{h}_i = \frac{1}{\sqrt{2}}[\mathbf{u}_i, -[\mathbf{J}\mathbf{u}_i]^t]^t$$

De même les v.p. anti-symétriques $\mathbf{g}_i$ de $\mathbf{G}$ s'obtiennent par

$$\mathbf{g}_i = \frac{1}{\sqrt{2}}[\mathbf{v}_i, +[\mathbf{J}\mathbf{v}_i]^t]^t$$

où $\mathbf{u}_i$ et $\mathbf{v}_i$ sont respectivement les v.p. orthonormés des matrices $\mathbf{A} - \mathbf{J}\mathbf{C}$ et $\mathbf{A} + \mathbf{J}\mathbf{C}$.

– si N est impair on a

$$\mathbf{G} = \begin{pmatrix} \mathbf{A} & \mathbf{x} & \mathbf{C}^t \\ \mathbf{x}^t & q & \mathbf{x}^t \mathbf{J} \\ \mathbf{C} & \mathbf{J}\mathbf{x} & \mathbf{J}\mathbf{A}\mathbf{J} \end{pmatrix}$$

où $\mathbf{A}$, $\mathbf{C}$ et $\mathbf{J}$ sont des matrices de dimensions $[N/2 - 1, N/2 - 1]$ et

$$\mathbf{A}^t = \mathbf{A} \quad \text{et} \quad \mathbf{C}^t = \mathbf{J}\mathbf{C}\mathbf{J}$$

et on montre que si $\mathbf{G}$ a des valeurs propres distinctes elle a $[(N+1)/2]$ v.p. symétriques ($\mathbf{h}_i$, $i = 1, \dots, (N+1)/2$) et $[(N-1)/2]$ v.p. anti-symétriques ($\mathbf{g}_i$, $i = 1, \dots, (N-1)/2$)

On montre aussi que $\mathbf{G}$ est orthogonalement similaire à

$$\begin{pmatrix} \mathbf{A} - \mathbf{J}\mathbf{C} & 0 & 0 \\ 0 & q & \sqrt{2}\mathbf{x}^t \\ 0 & \sqrt{2}\mathbf{x} & \mathbf{A} + \mathbf{J}\mathbf{C} \end{pmatrix}$$

Dans ce cas les v.p. symétriques $\mathbf{h}_i$ de $\mathbf{G}$ s'obtiennent par

$$\mathbf{h}_i = \frac{1}{\sqrt{2}}[\mathbf{u}_i^t, 0, -[\mathbf{J}\mathbf{u}_i]^t]^t$$

De même les v.p. anti-symétriques $\mathbf{g}_i$ de $\mathbf{G}$ s'obtiennent par

$$\mathbf{g}_i = \frac{1}{\sqrt{2}}[\mathbf{v}_i^t, \sqrt{2}\alpha_i, [\mathbf{J}\mathbf{v}_i]^t]^t$$

où $\mathbf{u}_i$ sont les v.p. orthonormés de la matrices $\mathbf{A} - \mathbf{J}\mathbf{C}$, et $[\mathbf{v}_i, \alpha_i]^t$ sont les v.p. orthonormés de la matrice

$$\begin{pmatrix} \mathbf{A} + \mathbf{J}\mathbf{C} & \sqrt{2}\mathbf{x} \\ \sqrt{2}\mathbf{x}^t & q \end{pmatrix}$$

Dans tous les cas ces v.p. sont uniques (à une constante multiplicative près) et mutuellement orthogonaux entre eux.

Si $\mathbf{a} = [a_0, a_1, \dots, a_{(N-1)}]^t$ est un v.p. de $\mathbf{G}$, c'est à dire $\mathbf{G}\mathbf{a} = \lambda\mathbf{a}$, le polynôme

$$\mathbf{A}(z) = \sum_{k=0}^{N-1} a_k z^{-k}$$

est un polynôme d'ordre $(N-1)$, appelé polynôme propre de la matrice $\mathbf{G}$. On montre que si $z = z_i$ est une racine de $\mathbf{A}(z)$, alors $z = \frac{1}{z_i}$ en est une aussi.

Matrices d'innovation diagonales (MID)

(En Anglais : Diagonal Innovation matrices DIM)

Considérons les trois vecteurs $\mathbf{x}_{m-1}$ de dimension $[m-1]$, $\mathbf{z}_{m-1}$ de dimension $[m-1]$, et de $\mathbf{y}_m$ de dimension $[m]$. Une matrice MID est une matrice qui est générée par les trois vecteurs génératrices $\mathbf{x}_{m-1}$, $\mathbf{z}_{m-1}$, et $\mathbf{y}_m$ de façon suivante:

$$\mathbf{R} = \begin{pmatrix} y_1 & \vdots & x_1 & \vdots & x_1 & \vdots & x_1 & \vdots & x_1 \\ \dots & & & & & & & & \\ z_1 & & y_2 & \vdots & x_2 & \vdots & x_2 & \vdots & x_2 \\ \dots & \dots & \dots & & & & & & \\ z_1 & & z_2 & & y_3 & \vdots & x_3 & \vdots & x_3 \\ \dots & \dots & \dots & \dots & \dots & & & & \\ z_1 & & z_2 & & z_3 & & y_4 & \vdots & x_4 \\ \dots & \dots & \dots & \dots & \dots & \dots & & & \\ z_1 & & z_2 & & z_3 & & z_4 & & y_5 \end{pmatrix}, \quad \begin{cases} \mathbf{x}_{m+1}^t = [\mathbf{x}_m^t, x_{m+1}] \\ \mathbf{z}_{m+1}^t = [\mathbf{z}_m^t, z_{m+1}] \\ \mathbf{y}_{m+1}^t = [\mathbf{y}_m^t, y_{m+1}] \end{cases}$$

Matrices d'innovation périphériques (MIP)

(En Anglais : Peripheral Innovation matrices PIM)

Une matrice MIP est une matrice qui générée par les trois vecteurs génératrices $\mathbf{x}_{m-1}$, $\mathbf{z}_{m-1}$, et $\mathbf{y}_m$ de façon suivante:

$$\mathbf{R} = \begin{pmatrix} y_1 & \vdots & x_1 & \vdots & x_2 & \vdots & x_3 & \vdots & x_4 \\ \dots & & & & & & & & \\ z_1 & & y_2 & \vdots & x_1 & \vdots & x_2 & \vdots & x_3 \\ \dots & \dots & \dots & & & & & & \\ z_2 & & z_1 & & y_3 & \vdots & x_1 & \vdots & x_2 \\ \dots & \dots & \dots & \dots & \dots & & & & \\ z_3 & & z_2 & & z_1 & & y_4 & \vdots & x_1 \\ \dots & \dots & \dots & \dots & \dots & \dots & & & \\ z_4 & & z_3 & & z_2 & & z_1 & & y_5 \end{pmatrix}, \quad \begin{cases} \mathbf{x}_{m+1}^t = [x_{m+1}, \mathbf{x}_m^t] \\ \mathbf{z}_{m+1}^t = [z_{m+1}, \mathbf{z}_m^t] \\ \mathbf{y}_{m+1}^t = [y_{m+1}, \mathbf{y}_m^t] \end{cases}$$

Les matrices du types $\mathbf{T}$ (Toeplitz), $\mathbf{C}$ (circulante), et $\mathbf{H}$ (Hankel) sont des cas particuliers des matrices MIP. L'inversion des matrices du type MID et MIP demande un coût de calcul de l'ordre de $O(N_2)$.

Matrices de Vandermonde

Matrice de Vandermonde est une matrice dans laquelle les éléments de chaque colonne sont constitués par la puissance n -ième d'un scalaire. Cette matrice est entièrement définie par le vecteur $\mathbf{x} = [x_1, \dots, x_n]$.

$$\mathbf{V} = \begin{pmatrix} 1 & 1 & \dots & 1 \\ x_1 & x_2 & \dots & x_n \\ x_1^2 & x_2^2 & \dots & x_n^2 \\ \dots & \dots & \dots & \dots \\ x_1^{m-1} & x_2^{m-1} & \dots & x_n^{m-1} \end{pmatrix}$$

$$\det[\mathbf{V}] = \prod_{i \neq j} (x_j - x_i)$$

E.3 Matrices blocs

Une matrice quelconque peut toujours être décomposée en blocs. Si ces blocs ont tous les mêmes dimensions on appellera la matrice *matrice-bloc*.

Soit $\mathbf{R} = [\mathbf{R}_{IJ}]$ une matrice-bloc dont chaque bloc est une matrice de dimensions $[P, P]$.

Alors $\mathbf{R}$ est:

<i>bloc-symétrique</i>	si $\mathbf{R}_{IJ} = \mathbf{R}_{JI}$
<i>bloc-antisymétrique</i>	si $\mathbf{R}_{IJ} = -\mathbf{R}_{JI}$
<i>bloc-Hankel</i>	si $\mathbf{R}_{IJ} = \mathbf{R}_{I+J}$
<i>bloc-Toeplitz</i>	si $\mathbf{R}_{IJ} = \mathbf{R}_{I-J}$

Tous ce que l'on peut dire sur les matrices sont valables sur les matrices-bloc si tous les blocs sont inversibles. Par exemple:

Soit $\mathbf{R} = [\mathbf{R}_{IJ}]$, $\mathbf{I}$, $j = 1, \dots, N$ une matrice-bloc avec $\mathbf{R}_{IJ}$ des matrices de dimensions $[P, P]$. Si toutes les $\mathbf{R}_{IJ}$ sont inversibles, alors $\mathbf{R}$ peut être factorisée soit par $\mathbf{R} = \mathbf{L}\mathbf{D}\mathbf{L}^t$, soit par $\mathbf{U}\mathbf{R}\mathbf{U}^t = \mathbf{D}$ où $\mathbf{L}$ est une matrice bloc-triangulaire inférieure avec des matrices d'identités sur son diagonale, $\mathbf{U}$ est une matrice bloc-triangulaire supérieure, $\mathbf{D}$ est une matrice bloc-diagonale, et $\mathbf{U} = \mathbf{L}^{-1}$.

Inversion d'une matrice décomposée en blocs

Soit $\mathbf{A}$ une matrice quelconque telle que

$$\mathbf{A} = \begin{pmatrix} \mathbf{A}_{11} & \mathbf{A}_{12} \\ \mathbf{A}_{21} & \mathbf{A}_{22} \end{pmatrix}$$

Si $\mathbf{A}_{22}^{-1}$ existe, alors $\mathbf{A}$ peut être factorisée de la façon suivante:

$$\mathbf{A} = \begin{pmatrix} \mathbf{I} & \mathbf{A}_{12}\mathbf{A}_{22}^{-1} \\ 0 & \mathbf{I} \end{pmatrix} \begin{pmatrix} \mathbf{A}_{11} - \mathbf{A}_{12}\mathbf{A}_{22}^{-1}\mathbf{A}_{21} & 0 \\ 0 & \mathbf{A}_{22} \end{pmatrix} \begin{pmatrix} \mathbf{I} & 0 \\ \mathbf{A}_{22}^{-1}\mathbf{A}_{21} & \mathbf{I} \end{pmatrix}$$

On a aussi

$$\text{rang}[\mathbf{A}] = \text{rang}[\mathbf{A}_{11} - \mathbf{A}_{12}\mathbf{A}_{22}^{-1}\mathbf{A}_{21}] + \text{rang}[\mathbf{A}_{22}]$$

$\mathbf{A}^{-1}$ existe ssi l'inverse de la matrice $\mathbf{T} = \mathbf{A}_{11} - \mathbf{A}_{12}\mathbf{A}_{22}^{-1}\mathbf{A}_{21}$ existe. Dans ce cas si on note par $\mathbf{B} = \mathbf{A}^{-1}$ on a:

$$\mathbf{B} = \begin{pmatrix} \mathbf{B}_{11} & \mathbf{B}_{12} \\ \mathbf{B}_{21} & \mathbf{B}_{22} \end{pmatrix}$$

$$\mathbf{B}_{11} = \mathbf{T}^{-1} = (\mathbf{A}_{11} - \mathbf{A}_{12}\mathbf{A}_{22}^{-1}\mathbf{A}_{21})^{-1} = \mathbf{A}_{11}^{-1} + \mathbf{A}_{11}^{-1}\mathbf{A}_{12}\mathbf{B}_{22}\mathbf{A}_{21}\mathbf{A}_{11}^{-1}$$

$$\mathbf{B}_{22} = \mathbf{D}^{-1} = (\mathbf{A}_{22} - \mathbf{A}_{21}\mathbf{A}_{11}^{-1}\mathbf{A}_{12})^{-1} = \mathbf{A}_{22}^{-1} + \mathbf{A}_{22}^{-1}\mathbf{A}_{21}\mathbf{B}_{11}\mathbf{A}_{12}\mathbf{A}_{22}^{-1}$$

$$\mathbf{B}_{12} = -\mathbf{A}_{11}^{-1}\mathbf{A}_{12}\mathbf{B}_{22} = -\mathbf{B}_{11}\mathbf{A}_{12}\mathbf{A}_{22}^{-1}$$

$$\mathbf{B}_{21} = -\mathbf{A}_{22}\mathbf{A}_{21}\mathbf{B}_{11} = -\mathbf{B}_{22}\mathbf{A}_{21}\mathbf{A}_{11}^{-1}$$

Les matrices $\mathbf{L}$ et $\mathbf{D}$ sont appelées *le complément Shur* de la matrice $\mathbf{A}$.

Cas particulier 1 :

Si $\mathbf{A}$ est une matrice bloc triangulaire inférieure, c'est à dire si $\mathbf{A}_{21} = 0$, $\mathbf{B}$ est aussi triangulaire inférieure, c'est à dire $\mathbf{B}_{21} = 0$ et on a :

$$\mathbf{B}_{11} = \mathbf{A}_{11}^{-1}, \quad \mathbf{B}_{22} = \mathbf{A}_{22}^{-1}, \quad \mathbf{B}_{12} = -\mathbf{A}_{11}^{-1} \mathbf{A}_{12} \mathbf{A}_{22}^{-1}$$

Cas particulier 2 :

Si $\mathbf{A}$ est une matrice bloc triangulaire supérieure, c'est à dire si $\mathbf{A}_{12} = 0$, $\mathbf{B}$ est aussi triangulaire supérieure, c'est à dire $\mathbf{B}_{12} = 0$ et on a :

$$\mathbf{B}_{11} = \mathbf{A}_{11}^{-1}, \quad \mathbf{B}_{22} = \mathbf{A}_{22}^{-1}, \quad \mathbf{B}_{21} = -\mathbf{A}_{22}^{-1} \mathbf{A}_{21} \mathbf{A}_{11}^{-1}$$

Cas particulier 3 :

Si $\mathbf{A}$ est une matrice bloc-diagonale, c'est à dire si $\mathbf{A}_{12} = \mathbf{A}_{21} = 0$, $\mathbf{B}$ est aussi bloc-diagonale, c'est à dire $\mathbf{B}_{12} = \mathbf{B}_{21} = 0$ et on a :

$$\mathbf{B}_{11} = \mathbf{A}_{11}^{-1}, \quad \mathbf{B}_{22} = \mathbf{A}_{22}^{-1}$$

Cas particulier 4 :

Si $\mathbf{A}_{22}$ est un scalaire et $\mathbf{A}_{21}$ et $\mathbf{A}_{12}$ sont des vecteurs, on a :

$$\mathbf{A} = \begin{pmatrix} \mathbf{A}_{11} & \mathbf{x} \\ \mathbf{z}^t & y \end{pmatrix}, \quad \mathbf{B} = \mathbf{A}^{-1} = \frac{1}{\alpha} \begin{pmatrix} \alpha \mathbf{A}_{11}^{-1} + \mathbf{w} \mathbf{v}^t & \mathbf{w} \\ \mathbf{v}^t & 1 \end{pmatrix}$$

où α , $\mathbf{w}$, et $\mathbf{v}$ s'obtient par :

$$\alpha = \frac{1}{(y - \mathbf{z}^t \mathbf{A}_{11}^{-1} \mathbf{x})} = \frac{|\mathbf{A}|}{|\mathbf{A}_{11}|}, \quad \mathbf{w} = -\mathbf{A}_{11}^{-1} \mathbf{x}, \quad \mathbf{v} = -\mathbf{A}_{11}^{-t} \mathbf{z}$$

E.4 Matrices de Toeplitz et de Hankel

E.4.1 Matrices de Toeplitz

Une matrice de dimensions $[M, N]$ est dite de Toeplitz si

$$[\mathbf{T}]_{ij} = t_{(i-j)}, \quad i = 1, \dots, M, \quad j = 1, \dots, N$$

$$\mathbf{T} = \begin{pmatrix} t_0 & t_{-1} & \cdot & \cdot & t_{-(N-2)} & t_{-(N-1)} \\ t_1 & t_0 & \cdot & \cdot & \cdot & t_{-(N-2)} \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & t_0 & \cdot & \cdot & \cdot \\ t_{(M-2)} & \cdot & \cdot & \cdot & t_0 & t_{-1} \\ t_{(M-1)} & t_{(M-2)} & \cdot & \cdot & t_1 & t_0 \end{pmatrix}$$

Dans le cas où $\mathbf{T}$ est carrée on a $[\mathbf{T}]_{ij} = t_{(i-j)}$, $i, j = 1, \dots, N$. Si de plus $\mathbf{T}$ est symétrique on a $[\mathbf{T}]_{ij} = t_{|i-j|}$, $i, j = 1, \dots, N$.

Par la suite nous considérons seulement les matrices de Toeplitz carrées de dimensions $[N, N]$.

$$\mathbf{T} = \begin{pmatrix} t_0 & t_{-1} & \cdot & \cdot & t_{-(N-2)} & t_{-(N-1)} \\ t_1 & t_0 & \cdot & \cdot & \cdot & t_{-(N-2)} \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & t_0 & \cdot & \cdot & \cdot \\ t_{(N-2)} & \cdot & \cdot & \cdot & t_0 & t_{-1} \\ t_{(N-1)} & t_{(N-2)} & \cdot & \cdot & t_1 & t_0 \end{pmatrix}$$

Une matrice de Toeplitz est entièrement déterminée par sa première ligne

$$[\mathbf{T}]_{1*} = [t_0, t_{-1}, \dots, t_{-(N-2)}, t_{-(N-1)}]^t$$

et sa première colonne

$$[\mathbf{T}]_{*1} = [t_0, t_1, \dots, t_{(N-2)}, t_{(N-1)}]^t$$

c'est à dire par $(2N - 1)$ éléments $\{t_{-(N-1)}, \dots, t_0, \dots, t_{(N-1)}\}$.

Si $\mathbf{T}_N$ est une matrice de Toeplitz de dimensions $[N, N]$ et si on note par $\{\mathbf{T}_n, \quad n = 1, \dots, N\}$ les sous-matrices de $\mathbf{T}_N$ on a les relations suivantes:

$$\mathbf{J}_n \mathbf{T}_n \mathbf{J}_n = \mathbf{T}_n^t, \quad n = 1, \dots, N$$

où $\mathbf{J}_n$ est une matrice d'inversion d'ordre de dimensions $[n, n]$. Une matrice de Toeplitz peut être factorisée de façon suivante:

$$\mathbf{T} = \mathbf{L}\mathbf{U}$$

De plus si elle est hermitienne elle peut être factorisée de façon suivante:

$$\mathbf{T} = \mathbf{L}\mathbf{D}\mathbf{L}^*$$

Les lignes successives de $\mathbf{L}$ sont les coefficients des prédicateurs linéaires d'ordre $0, 1, \dots, N-1$, et les éléments diagonaux de la matrice $\mathbf{D}$ sont les variances d'erreurs de prédictions correspondantes.

Si $\mathbf{T}$ est une matrice de Toeplitz, on a:

$$\mathbf{T}^{-1} = \frac{1}{\lambda_{n-1}} \begin{pmatrix} 1 & \mathbf{b}_{n-1}^t \\ \mathbf{a}_{n-1} & \mathbf{S}_{n-1} \end{pmatrix}$$

Les éléments de $\mathbf{S}_{n-1}$ sont déterminés complètement par λ_{n-1} , $\mathbf{a}_{n-1}$, et $\mathbf{b}_{n-1}$.

Une décomposition moins connue est:

$$\mathbf{T}^{-1} = \frac{1}{\lambda_{n-1}} [\mathbf{A}! \mathbf{B}! - \mathbf{B}\mathbf{A}]$$

où $\mathbf{A}!$ et $\mathbf{B}$ sont des matrices Toeplitz triangulaire supérieur type $\mathbf{L}_1$, et $\mathbf{A}$ et $\mathbf{B}!$ sont des matrices Toeplitz triangulaire inférieur type $\mathbf{U}_1$. Les éléments de ces matrices sont données par:

$$[\mathbf{A}]_{jk} = [\tilde{\mathbf{a}}_{n-1}]_{k-j}, \quad [\mathbf{B}]_{jk} = [\tilde{\mathbf{b}}_{n-1}]_{k-j}$$

$$[\mathbf{A}!]_{jk} = [\mathbf{a}_{n-1}]_{j-k} + \delta_{j-k}, \quad [\mathbf{B}!]_{jk} = [\mathbf{b}_{n-1}]_{k-j} + \delta_{j-k}, \quad 1 \leq j, k \leq n$$

et où pour tout vecteur $\mathbf{g}_m$, $\mathbf{g}_m(k) = 0, \forall k \leq 0$ ou $k > m$.

E.4.2 Matrices de Hankel

$\mathbf{H}$ est une matrice de Hankel si tous ses éléments se trouvant sur des lignes parallèle à sa deuxième diagonale sont identiques. Par exemple:

$$\mathbf{H} = \begin{pmatrix} h_N & h_{-N+1} & \cdot & \cdot & h_{-1} & h_0 \\ h_{-N+1} & \cdot & \cdot & \cdot & h_0 & h_1 \\ \vdots & \cdot & \cdot & \cdot & \cdot & \vdots \\ h_{-1} & h_0 & \cdot & \cdot & \cdot & h_{M-1} \\ h_0 & h_1 & \cdot & \cdot & \cdot & h_M \end{pmatrix}$$

Si $\mathbf{H}$ est une matrice carrée on la note par $\mathbf{H}_N$.

$\mathbf{H}_N$ est une matrice de Hankel d'ordre N si

$$[\mathbf{H}]_{ij} = h_{i+j-2}, \quad i, j = 1, \dots, N$$

$$\mathbf{T} = \begin{pmatrix} h_0 & h_1 & \cdot & \cdot & h_{N-2} & h_{N-1} \\ h_1 & h_2 & \cdot & \cdot & h_{N-1} & h_N \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & h_{N-1} & \cdot & \cdot & h_{2N-3} \\ \cdot & \cdot & \cdot & \cdot & h_{2N-2} & h_{2N-2} \\ h_N & h_{N+1} & \cdot & h_{2N-2} & h_{2N-1} \end{pmatrix}$$

Une matrice de Hankel est toujours symétrique: $\mathbf{H}^t = \mathbf{H}$. Elle est hermitienne ssi tous ses éléments sont réels. Si $\mathbf{T}$ est une matrice de Toeplitz, alors les matrices

$$\mathbf{H}_1 = \mathbf{T}\mathbf{J}, \quad \text{et} \quad \mathbf{H}_2 = \mathbf{J}\mathbf{T}^t$$

sont des matrices de Hankel.

Les trois systèmes d'équations suivantes sont équivalentes:

$$\mathbf{H}\mathbf{a} = -\mathbf{d}, \quad \mathbf{T}[\mathbf{J}\mathbf{a}] = -\mathbf{d}, \quad \mathbf{T}\mathbf{a} = -\mathbf{J}\mathbf{d}$$

E.5 Matrices circulantes

Une matrice circulante est une matrice qui est définie par seule sa première ligne (ou seule par sa première colonne). $\mathbf{C}$ est une matrice *circulante à droite* (*circulante vers le haut*) si :

$$[\mathbf{C}]_{ij} = \begin{cases} c_{j-i} & \text{si } j-i \geq 0 \\ c_{n-j+1} & \text{si } j-i < 0 \end{cases}$$

Exemple :

$$\mathbf{C} = \begin{pmatrix} c_0 & c_1 & c_2 & \cdots & \cdots & c_{n-1} \\ c_{n-1} & c_0 & c_1 & \cdots & \cdots & c_{n-2} \\ c_{n-2} & c_{n-1} & c_0 & \cdots & \cdots & . \\ \vdots & c_{n-2} & c_{n-1} & c_0 & \cdots & \vdots \\ . & \cdots & \cdots & \cdots & \cdots & c_1 \\ c_1 & c_2 & \cdots & \cdots & c_{n-1} & c_0 \end{pmatrix}$$

Une matrice circulante à droite est une matrice de Toeplitz qui est symétrique par rapport à sa deuxième diagonale. $\mathbf{C}$ est une matrice *circulante à gauche* (ou *circulante vers le bas*) si :

$$[\mathbf{C}]_{ij} = \begin{cases} c_{n+1-i-j} & \text{si } j+i \leq n+1 \\ c_{2n+1-i-j} & \text{si } j+i > n+1 \end{cases}$$

Exemple :

$$\mathbf{C} = \begin{pmatrix} c_{n-1} & \cdots & \cdots & c_2 & c_1 & c_0 \\ c_{n-2} & c_{n-1} & \cdots & \cdots & c_0 & c_{n-1} \\ c_{n-3} & \cdots & \cdots & c_0 & c_{n-1} & c_{n-2} \\ \vdots & \cdots & c_0 & c_{n-1} & c_{n-2} & \vdots \\ c_1 & \cdots & \cdots & \cdots & \cdots & c_2 \\ c_0 & c_{n-1} & c_{n-2} & \cdots & c_2 & c_1 \end{pmatrix}$$

Une matrice circulante à gauche est une matrice de Hankel qui est symétrique par rapport à sa diagonale. Si $\mathbf{C}$ est une matrice circulante de dimensions $[N, N]$ avec des éléments réels $[\mathbf{C}]_{ij}$ on a :

$$[\mathbf{C}]_{ij} = c_{(i-j) \bmod N}, \quad \mathbf{C}^t = \mathbf{C}^{-1}$$

$$\mathbf{C} = \begin{pmatrix} c_0 & c_{N-1} & . & . & c_2 & c_1 \\ c_1 & c_0 & . & . & . & c_2 \\ . & . & . & . & . & . \\ . & . & . & . & . & . \\ c_{N-2} & . & . & . & c_0 & c_{N-1} \\ c_{N-1} & c_{N-2} & . & . & c_1 & c_0 \end{pmatrix}$$

Les valeurs propres d'une matrice circulante sont les coefficients de la TFD de sa première ligne et ses vecteurs propres sont les vecteurs de base de la TFD.

Si $\mathbf{C}$ est une matrice circulante, alors

$$\mathbf{C} = \mathbf{F}\mathbf{\Lambda}\mathbf{F}^* = \mathbf{F}^*\mathbf{\Lambda}_c^*\mathbf{F}$$

où

- $\mathbf{F}$ est la matrice de la TFD :

$$\mathbf{F} = \frac{1}{\sqrt{N}}[\mathbf{w}_0, \mathbf{w}_1, \dots, \mathbf{w}_p, \dots, \mathbf{w}_{n-1}]$$

$$[\mathbf{w}_p]_q = \exp\left[-j\frac{2\pi pq}{N}\right], \quad p = 0, \dots, N-1 \quad q = 1, \dots, N$$

- $\mathbf{\Lambda}_c$ est une matrice diagonale contenant les valeurs propres de la matrice $\mathbf{C}$, c'est à dire les coefficients de la TFD de sa première ligne.
- $\mathbf{F}$ est une matrice qui multipliée par un vecteur $\mathbf{x}$ donne un vecteur $\bar{\mathbf{x}}$ qui contient les valeurs de sa TFD.

$$\mathbf{F}\mathbf{x} = \bar{\mathbf{x}} \quad \text{et} \quad \mathbf{x} = \mathbf{F}^*\bar{\mathbf{x}}$$

E.6 Matrices bloc-Toeplitz et bloc-circulantes

E.6.1 Matrices bloc-Toeplitz

La matrice $\mathbf{R}$ est dite bloc-Toeplitz si $\mathbf{R}$ est une matrice-bloc dont la IJ -ème bloc $\mathbf{R}_{IJ}$ est fonction de $I - J$):

$$\mathbf{R}_{IJ} = \mathbf{R}_{I-J}.$$

Si chaque bloc est aussi une matrice de Toeplitz alors $\mathbf{R}$ est une matrice Toeplitz-bloc-Toeplitz.

Matrices bloc-Toeplitz avec des blocs Toeplitz

Soit $\mathbf{R}_{P,M}$ est une matrice bloc d'ordre M avec les blocs $\mathbf{R}_{IJ}$, $I, J = 1, \dots, M$, qui aux même sont des matrices de Toeplitz de dimensions $[P, P]$, c'est à dire:

$$\mathbf{R}_{PM} = \begin{pmatrix} \mathbf{R}_0 & \mathbf{R}_{-1} & \cdot & \cdot & \mathbf{R}_{-(M-2)} & \mathbf{R}_{-(M-1)} \\ \mathbf{R}_1 & \mathbf{R}_0 & \cdot & \cdot & \cdot & \mathbf{R}_{-(M-2)} \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \mathbf{R}_0 & \cdot & \cdot & \cdot \\ \mathbf{R}_{(M-2)} & \cdot & \cdot & \cdot & \mathbf{R}_0 & \mathbf{R}_{-1} \\ \mathbf{R}_{(M-1)} & \mathbf{R}_{(M-2)} & \cdot & \cdot & \mathbf{R}_1 & \mathbf{R}_0 \end{pmatrix}$$

avec

$$\mathbf{R}_J = \begin{pmatrix} r_0 & r_{-1} & \cdot & \cdot & r_{-(P-2)} & r_{-(P-1)} \\ r_1 & r_0 & \cdot & \cdot & \cdot & r_{-(P-2)} \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & r_0 & \cdot & \cdot & \cdot \\ r_{(P-2)} & \cdot & \cdot & \cdot & r_0 & r_{-1} \\ r_{(P-1)} & r_{(P-2)} & \cdot & \cdot & r_1 & r_0 \end{pmatrix}$$

Soit $\mathbf{J}_{P,M}$ une matrice de dimensions $[PM, PM]$ telle que

$$\mathbf{J}_{PM} = \begin{pmatrix} \cdot & \cdot & \cdot & \cdot & \mathbf{J}_P \\ \cdot & 0 & \cdot & \mathbf{J}_P & \cdot \\ \cdot & \cdot & \mathbf{J}_P & \cdot & \cdot \\ \cdot & \mathbf{J}_P & \cdot & 0 & \cdot \\ \mathbf{J}_P & \cdot & \cdot & \cdot & \cdot \end{pmatrix} \quad \text{avec} \quad \mathbf{J}_P = \begin{pmatrix} \cdot & \cdot & \cdot & \cdot & 1 \\ \cdot & 0 & \cdot & 1 & \cdot \\ \cdot & \cdot & 1 & \cdot & \cdot \\ \cdot & 1 & \cdot & 0 & \cdot \\ 1 & \cdot & \cdot & \cdot & \cdot \end{pmatrix}$$

une matrice de dimensions $[P, P]$. Alors on a les propriétés suivantes:

- $\mathbf{R}_{P,M}$ et son inverse sont toutes les deux des matrices centrosymétriques et on a

$$\mathbf{J}_{P,M} \mathbf{R}_{P,M} \mathbf{J}_{P,M} = \mathbf{R}_{P,M}^t \quad \text{et} \quad \mathbf{J}_{P,M} \mathbf{R}_{P,M}^{-1} \mathbf{J}_{P,M} = \mathbf{R}_{P,M}^{-t}$$

- $\mathbf{R}_{P,M}$ qui est une matrice $\mathbf{TT}$ de $[M, M]$ blocs de matrice de Toeplitz de dimensions $[P, P]$ peut être transformée à une matrice $\mathbf{TT}$ de $[P, P]$ blocs de matrices de Toeplitz de dimensions $[M, M]$ avec de simples permutations de ses éléments. Ceci peut être démontré en définissant une matrice de permutation $\mathbf{E}$ de dimensions $[PM, PM]$ telle que

$$\mathbf{E} = (\mathbf{E}_0 \quad . \quad . \quad \mathbf{E}_{P-1})$$

avec $\mathbf{E}_j$ une matrice de dimensions $[M, MP]$ telle que

$$\mathbf{E}_j = \begin{pmatrix} \mathbf{e}_i & . & . & . & . \\ . & \mathbf{e}_i & . & 0 & . \\ . & . & . & . & . \\ . & 0 & . & \mathbf{e}_i & . \\ . & . & . & . & \mathbf{e}_i \end{pmatrix}$$

avec $\mathbf{e}_i$ un vecteur ligne dont toutes ses composantes sont nulles sauf la j -ème qui est égale à 1. On montre alors que

$$\mathbf{E} \mathbf{R}_{P,M} \mathbf{E}^t = \tilde{\mathbf{R}}_{M,P}$$

où $\tilde{\mathbf{R}}_{M,P}$ est une matrice $\mathbf{T}\mathbf{T}$ de $[PxP]$ blocs de matrices de Toeplitz de dimensions $[M, M]$. La ij -ème bloc de $\tilde{\mathbf{R}}_{M,P}$ s'obtient par $\mathbf{E}_j \mathbf{R}_{P,M} \mathbf{E}_j^t$ et le kl -ème élément de ce bloc est $\mathbf{e}_j \mathbf{R}_{l-k} \mathbf{e}_j^t = r_{j-i, l-k}$.

$$\mathbf{E} \mathbf{R}_{P,M} \mathbf{E}^t = \begin{pmatrix} \mathbf{E}_0 \\ . \\ . \\ \mathbf{E}_{P-1} \end{pmatrix} \mathbf{R}_{P,M} [\mathbf{E}^t, \dots, \mathbf{E}_{p-1}^t]$$

E.6.2 Matrices bloc-circulantes

Si $\mathbf{A}$ est une matrice bloc-circulante, composée de $[P, P]$ blocs de matrices circulante de dimensions $[K, K]$, $\mathbf{A}$ peut être factorisée de façon suivante:

$$\mathbf{A} = \mathbf{F}^* \mathbf{D} \mathbf{F}$$

$$\mathbf{F} = \begin{pmatrix} F & . & . & . & . \\ . & F & . & 0 & . \\ . & . & . & . & . \\ 0 & . & F & . & . \\ . & . & . & . & F \end{pmatrix} \quad \mathbf{D} = \begin{pmatrix} D_{11} & . & . & . & D_{1p} \\ . & . & . & . & . \\ . & . & . & . & . \\ . & . & . & . & . \\ D_{p1} & . & . & . & D_{pp} \end{pmatrix}$$

Chaque matrice D_{ij} est une matrice diagonale contenant les valeurs propres de la ij -ème bloc (matrice circulante) de $\mathbf{A}$.

E.7 Inversion récursive des matrices

Dans cette section nous allons présenter une méthode d'inversion récursive des matrices qui est basée sur le partitionnement de matrices en blocs.

Cas général

Si $\mathbf{A}_n$ est une matrice définie positive d'ordre n , alors on peut la partitionner de façon suivante:

$$\mathbf{A}_n = \begin{pmatrix} \mathbf{A}_{n-1} & \mathbf{a}_n \\ \mathbf{b}_n^t & a_{nn} \end{pmatrix}$$

où $\mathbf{a}_n$ et $\mathbf{b}_n$ sont définies par

$$\mathbf{a}_n = [a_{1n}, a_{2n}, \dots, a_{n-1,n}]^t \quad \text{et} \quad \mathbf{b}_n = [b_{1n}, b_{2n}, \dots, b_{n-1,n}]^t$$

Si $\mathbf{Q}_n = \mathbf{A}_n^{-1}$ est aussi partitionnée par

$$\mathbf{Q}_n = \mathbf{A}_n^{-1} = \begin{pmatrix} \mathbf{Q}_{n-1} & \mathbf{Q}_n \\ \mathbf{v}_n^t & q_{nn} \end{pmatrix}$$

on a

$$\mathbf{Q}_{n-1} = \mathbf{A}_{n-1}^{-1} + \alpha_n \mathbf{Q}_n \mathbf{v}_n^t$$

où

$$\alpha_n = a_{nn} - \mathbf{b}_n \mathbf{A}_{n-1}^{-1} \mathbf{a}_n$$

$$q_{nn} = \frac{1}{\alpha_n}, \quad \mathbf{Q}_n = -\frac{1}{\alpha_n} [\mathbf{A}_{n-1}^{-1} \mathbf{a}_n], \quad \mathbf{v}_n = \frac{1}{\alpha_n} [\mathbf{A}_{n-1}^{-1} \mathbf{b}_n].$$

Cas des matrices symétriques

Si $\mathbf{A}$ est une matrice réelle, symétrique et définie positive, alors son inverse $\mathbf{A}^{-1}$ ainsi que toutes ses sous-matrices principales sont réelles, symétriques et définies positives. Si on note alors par $\mathbf{A}_n^{-1} = \mathbf{Q}_n$ et par $[\mathbf{Q}_n]_{ij} = q_{ij}^{(n)}$ on a

$$q_{nn}^{(n)} = \frac{1}{\left(a_{nn} - \sum_{i=1}^{n-1} a_{in} \beta_i^{(n)}\right)}$$

$$q_{ij}^{(n)} = q_{ij}^{(n-1)} - q_{in}^{(n)} \beta_i^{(n)} = q_{ij}^{(n-1)} - \beta_i^{(n)} q_{jn}^{(n)}$$

$$q_{in}^{(n)} = -q_{in}^{(n)} \beta_i^{(n)}, \quad i = 1, \dots, n-1, \quad \beta_i^{(n)} = \sum_{k=1}^{n-1} q_{ik}^{(n-1)} a_{kn}$$

avec initialisation $\beta_1^{(1)} = -1$ et $q_1^{(1)} = \frac{1}{a_{11}}$.

On a aussi la propriété suivante

$$|\mathbf{A}_n| = \prod_{i=1}^n \frac{1}{q_{ii}^{(i)}} = \frac{1}{q_{11}^{(1)} q_{22}^{(2)} \cdots q_{nn}^{(n)}}$$

Cas des matrices de Toeplitz symétriques

$$\beta_n = \sum_{k=1}^{n-1} q_{1k}^{(n-1)} a_{1,n+1-k}, \quad q_{11}^{(n)} = \frac{q_{11}^{(n-1)}}{1 - \beta_n^2}, \quad q_{1n}^{(n)} = -\beta_n q_{11}^{(n)}$$

Si $\mathbf{A}$ est une matrice définie positive, alors on montre que $0 \leq \beta_n < 1$, et que $q_{11}^{(n)}$ croît d'une façon monotone en fonction de n , c'est à dire que $(q_{11}^{(n)}/q_{11}^{(n-1)}) \geq 1$. D'autre part si β_n devient nul à l'ordre n , alors $q_{11}^{(n)}$ devient aussi nul, et les éléments de $n+1$ ème inverse ne changent plus.

Cas des matrices de corrélation d'un signal stationnaire

La matrice de corrélation d'un signal stationnaire exponentiellement est une matrice de Toeplitz symétrique avec la propriété suivante

$$\frac{a_{12}}{a_{11}} = \frac{a_{13}}{a_{12}} = \dots = \frac{a_{1N}}{a_{N-1}}$$

Dans ce cas on montre que l'on peut écrire

$$\beta_2 = \frac{a_{12}}{a_{11}}, \quad \beta_k = 0 \quad \text{pour } k > 2,$$

$$q_{11}^{(n)} = q_{nn}^{(n)} = \frac{1}{a_{11}(1 - \beta_2^2)}, \quad q_{12}^{(n)} = q_{21}^{(n)} = \frac{-\beta_2}{a_{11}(1 - \beta_2^2)}$$

$$q_{jj}^{(n)} = \frac{1 + \beta_2^2}{a_{11}(1 - \beta_2^2)}, \quad j = 2, \dots, n-1, \quad q_{ij}^{(n)} = 0 \quad |i - j| > 2$$

D'une façon plus schématique $\mathbf{A}_n^{-1}$ est sous la forme de

$$\mathbf{A}_n^{-1} = \frac{1}{a_{11}(1 - \beta_2^2)} \begin{pmatrix} 1 & -\beta_2 & \cdot & \cdot & \cdot \\ -\beta_2 & (1 + \beta_2^2) & \cdot & 0 & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & 0 & \cdot & (1 + \beta_2^2) & -\beta_2 \\ \cdot & \cdot & \cdot & -\beta_2 & 1 \end{pmatrix}$$

Cas des matrices de Toeplitz symétriques et tridiagonale

Un cas particulier des matrices de Toeplitz symétrique est une matrice dont les éléments du m -ième diagonal sont tous égaux à zéro pour m supérieure à k . Pour $k = 2$ on a une matrice tridiagonale. Ceci est par exemple le cas de la matrice de covariance d'un processus AR d'ordre 1.

$$\mathbf{A}_n = \begin{pmatrix} a_{11} & a_{12} & 0 & \cdot & 0 \\ a_{12} & a_{11} & \cdot & 0 & \cdot \\ 0 & \cdot & \cdot & \cdot & 0 \\ \cdot & 0 & \cdot & a_{11} & \cdot \\ \cdot & 0 & \cdot & a_{12} & a_{11} \end{pmatrix}$$

Pour cette matrice la valeur de β_n devient

$$\beta_n = \frac{(-1)^n (a_{12}/a_{11})^{n-1}}{\prod_{j=2}^{n-1} (1 - \beta_j^2)^{n-j}} = (-1)^n a_{12}^{n-1} \det [\mathbf{A}_{n-1}^{-1}], \quad n = 2, 3, \dots$$

Une propriété intéressante de ces matrices est le fait que même si la matrice est définie positive à l'ordre n , elle peut ne pas l'être à l'ordre $n+1$. Ceci peut être le cas si par exemple $|\beta_{n+1}| \geq 1$. On peut montrer que si $|\beta_2| \leq 1/2$, alors $|\beta_n| \rightarrow 0$ quand $n \rightarrow \infty$. Ainsi la matrice reste définie positive pour toute valeur de n , et si $|\beta_2| > 1/2$, la matrice deviendra non définie positive.

Cas des matrices triangulaires

Si $\mathbf{A}$ est une matrice réelle, triangulaire et définie positive d'ordre n , alors son inverse $\mathbf{A}_n^{-1}$ ainsi que toutes ses sous-matrices principales sont réelles, triangulaires et définies positives. Ceci est un cas particulier de (1.1) et (1.2) avec $b_n = 0$ et $\mathbf{v}_n = 0$. Ainsi les éléments q_{ij} de la matrice $\mathbf{Q}_n$ sont donnés par

$$q_{ij} = \begin{cases} -\frac{1}{a_{ii}} \sum_{k=1}^{j-1} q_{ik} a_{kj} & \text{pour } j > i \\ \frac{1}{a_{ii}} & \text{pour } j = i \\ 0 & \text{pour } j < i \end{cases}$$

Cas de Toeplitz et triangulaires

Si $\mathbf{A}$ est une matrice de Toeplitz, réelle, triangulaire et définie positive d'ordre n , alors son inverse $\mathbf{A}_n^{-1}$ ainsi que toutes ses sous-matrices principales sont des matrices de Toeplitz, réelles, triangulaires et définies positives.

$$\mathbf{A}_n = \begin{pmatrix} a_0 & a_1 & . & . & a_{n-2} & a_{n-1} \\ 0 & a_0 & a_1 & . & . & a_{n-2} \\ . & 0 & . & . & . & . \\ . & . & . & . & . & . \\ . & . & . & . & . & . \\ . & 0 & . & . & a_0 & a_1 \\ . & . & . & 0 & . & a_0 \end{pmatrix}$$

Dans ce cas on a $a_{ij} = a_{i-j}$, $q_{ij} = q_{i-j} = 0$ pour $i > j$, et on a $q_0 = 1/a_0$, et

$$q_{ij} = -\frac{1}{a_0} \sum_{k=1}^{j-1} q_{k-i} a_{j-k}$$

Si on note par $m = j - i$ et $n = k - i$, alors on a

$$q_m = -\frac{1}{a_0} \sum_{n=1}^{m-1} q_n a_{m-n}$$

$$\mathbf{Q}_n = \mathbf{A}_n^{-1} = \begin{pmatrix} q_0 & q_1 & \cdot & \cdot & q_{n-2} & q_{n-1} \\ 0 & q_0 & q_1 & \cdot & \cdot & q_{n-2} \\ \cdot & 0 & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & 0 & q_0 & q_1 \\ 0 & \cdot & \cdot & \cdot & 0 & q_0 \end{pmatrix}$$

E.8 Résolution des systèmes d'équations linéaires

Dans cette section nous allons présenter les deux problèmes essentiels du calcul matriciel qui sont la résolution d'un système linéaire et le calcul des valeurs propres et des vecteurs propres d'une matrice. L'objectif étant de mettre en évidence les difficultés de ces deux problèmes lors que la matrice est mal conditionnée.

Les deux problèmes fondamentaux de l'analyse numérique matricielles sont:

1. Résolution d'un système linéaire $\mathbf{A}\mathbf{u} = \mathbf{b}$.
2. Calcul des valeurs propres et des vecteurs propres d'une matrice $\mathbf{A}$.

Pour résoudre un problème du type (i) ou (ii), on commence naturellement par choisir une méthode déterminée (selon des critères qui seront précisés au cours des chapitres suivants). Or, quelle que soit la méthode, ce résultat n'est pas en général le résultat exacte, car certaines erreurs sont inévitables au cours du calcul ce sont les erreurs d'arrondi.

Alors que les erreurs d'arrondi sont toujours présentes dans un calcul, le deuxième type d'erreurs n'apparaît que dans certaines méthodes. Une méthode directe conduit à la solution exacte du problème en un nombre fini d'opérations élémentaires s'il n'y avait pas d'erreurs d'arrondi.

Dans une méthode récursive la solution $\mathbf{u}$ du problème est la limite d'une suite $\mathbf{u}_k$ de solutions exactes des sous-matrices de $\mathbf{A}$. Dans une méthode itérative la solution $\mathbf{u}$ du problème est la limite d'une suite $\mathbf{u}_k$ de solutions approchées. Comme on est bien obligé d'arrêter les calcul à l'itération k on commet ainsi une erreur de troncature mesurée par le nombre $\|\mathbf{u} - \mathbf{u}_k\|$ (pour une norme vectorielle donnée). Le choix d'une méthode dépend des propriétés de la matrice: matrice pleine quelconque matrice creuse; matrice symétrique; matrice-bandes; etc.

Une des premières caractéristiques de la matrice à considérer est le nombre et, éventuellement la répartition des éléments nuls de la matrice. C'est ainsi que l'on distingue les matrices pleines et les matrices creuses. Les problèmes de résolution de système linéaire les plus faciles à traiter numériquement sont ceux dont les matrices sont symétriques définies positives.

Les problèmes de calcul des valeurs propres et des vecteurs propres d'une matrice les plus faciles à traiter numériquement sont ceux des matrices symétriques. Autant il existe les trois types de méthodes directe, récursive et itérative pour le problème (1), les méthodes de calcul des valeurs propres ne sauraient qu'être itératives. En effet, les valeurs propres d'une matrice sont les racines d'un polynôme de degré n . Or l'existence d'une méthode directe pour le calcul des valeurs propres équivaldrait à affirmer qu'on peut calculer les racines d'un polynôme arbitraire en un nombre fini d'opérations élémentaires.

E.8.1 Conditionnement d'un problème de résolution des systèmes linéaires

Considérons la matrice inversible

$$\mathbf{A} = \begin{pmatrix} 10 & 7 & 8 & 7 \\ 7 & 5 & 6 & 5 \\ 8 & 6 & 10 & 9 \\ 7 & 5 & 9 & 10 \end{pmatrix}$$

et son inverse

$$\mathbf{A}^{-1} = \begin{pmatrix} 25 & -41 & 10 & -6 \\ -41 & 68 & -17 & 10 \\ 10 & -17 & 5 & -3 \\ -6 & 10 & -3 & 2 \end{pmatrix}$$

et les deux vecteurs $\mathbf{x}^t = [x_1, \dots, x_4]$ et $\mathbf{y}^t = [32, 23, 33, 31]$. La solution exacte de l'équation $\mathbf{Ax} = \mathbf{y}$ est $\mathbf{x}^t = [1, 1, 1, 1]$. Considérons maintenant le vecteur perturbé $[\mathbf{y} + \delta\mathbf{y}]^t = [32.1, 22.9, 33.1, 30.9]$ très légèrement différent à $\mathbf{y}$. La solution de l'équation $\mathbf{Ax} = [\mathbf{y} + \delta\mathbf{y}]$ est

$$\mathbf{x}^t = [9.2, -12.6, 4.5, -1.1]$$

Autrement dit une erreur relative de 1/200 sur les données $\mathbf{y}$ entraîne une erreur relative de 10/1 sur le résultat (soit un rapport d'amplification des erreurs relatives de l'ordre de 2000) de la solution du système linéaire. Considérons également une matrice $\mathbf{A}$ légèrement modifiée :

$$\mathbf{A} = \begin{pmatrix} 10 & 7 & 8.1 & 7.2 \\ 7.08 & 5.04 & 6 & 5 \\ 8 & 5.98 & 9.89 & 9 \\ 6.99 & 4.99 & 9 & 9.98 \end{pmatrix}$$

cette fois ci encore la solution de l'équation $\mathbf{Ax} = \mathbf{y}$ est $\mathbf{x}^t = [-81, 137, -34, 22]$. Là encore, de très petites variations des éléments de la matrice modifient complètement le résultat. Pourtant, la matrice $\mathbf{A}$ du système a un bon aspect: elle est symétrique, son déterminant vaut 1, et la matrice inverse $\mathbf{A}^{-1}$ est tout aussi sympathique! Ces deux exemples montrent qu'il est utile d'étudier le conditionnement d'une matrice quand on a à résoudre un système d'équations linéaires même si la solution est calculée par une méthode directe. Analyse des erreurs permet d'établir facilement les relations suivantes: soient $\mathbf{Ax} = \mathbf{y}$, $\mathbf{A}(\mathbf{x} + \delta\mathbf{x}) = \mathbf{y} + \delta\mathbf{y}$. Alors on a:

$$\frac{\|\delta\mathbf{x}\|}{\|\mathbf{x}\|} \leq \|\mathbf{A}\| \|\mathbf{A}^{-1}\| \frac{\|\delta\mathbf{y}\|}{\|\mathbf{y}\|}$$

soient

$$\mathbf{Ax} = \mathbf{y}, \quad (\mathbf{A} + \Delta\mathbf{A})(\mathbf{x} + \delta\mathbf{x}) = \mathbf{y}$$

Alors

$$\frac{\|\delta\mathbf{x}\|}{\|\mathbf{x} + \delta\mathbf{x}\|} \leq \|\mathbf{A}\| \|\mathbf{A}^{-1}\| \frac{\|\Delta\mathbf{A}\|}{\|\mathbf{A}\|}$$

On définit ainsi le nombre de conditionnement de la matrice $\mathbf{A}$ par

$$\text{cond}\{\mathbf{A}\} = \|\mathbf{A}\| \|\mathbf{A}^{-1}\|$$

Ces deux inégalités montrent que le nombre $\text{cond}\{\mathbf{A}\}$ mesure la sensibilité de la solution $\mathbf{x}$ du système linéaire $\mathbf{Ax} = \mathbf{y}$ vis-à-vis des variations sur les données $\mathbf{A}$ et $\mathbf{y}$, qualité que l'on appelle le conditionnement du système linéaire considéré. Un système linéaire est bien-ou-mal-conditionné selon que $\text{cond}\{\mathbf{A}\}$ est petit ou grand. Par la suite nous présentons un certain nombre de théorèmes et résultats (sans démonstrations) concernant $\text{cond}\{\mathbf{A}\}$.

Théorème 3 Soit $\mathbf{A}$ une matrice inversible, et soit $\mathbf{x}$ et $(\mathbf{x} + \delta\mathbf{x})$ les solutions des systèmes linéaires $\mathbf{Ax} = \mathbf{y}$ et $\mathbf{A}(\mathbf{x} + \delta\mathbf{x}) = \mathbf{y} + \delta\mathbf{y}$. Alors l'inégalité

$$\frac{\|\delta\mathbf{x}\|}{\|\mathbf{x}\|} \leq \text{cond}\{\mathbf{A}\} \frac{\|\delta\mathbf{y}\|}{\|\mathbf{y}\|}$$

est satisfaite, et c'est la meilleure possible — pour une matrice $\mathbf{A}$ donnée, on peut trouver des vecteurs $\mathbf{y} \neq 0$ et $\delta\mathbf{y} \neq 0$ tels qu'elle devienne une égalité.

Théorème 4 Soit $\mathbf{A}$ une matrice inversible, et soit $\mathbf{x}$ et $(\mathbf{x} + \delta\mathbf{x})$ les solutions des systèmes linéaires $\mathbf{Ax} = \mathbf{y}$ et $(\mathbf{A} + \Delta\mathbf{A})(\mathbf{x} + \delta\mathbf{x}) = \mathbf{y}$. Alors l'inégalité

$$\frac{\|\delta\mathbf{x}\|}{\|\mathbf{x} + \delta\mathbf{x}\|} \leq \text{cond}\{\mathbf{A}\} \frac{\|\Delta\mathbf{A}\|}{\|\mathbf{A}\|}$$

est satisfaite, et c'est la meilleure possible — pour une matrice $\mathbf{A}$ donnée, on peut trouver des vecteurs $\mathbf{y} \neq 0$ et une matrice $\Delta\mathbf{A} \neq 0$ tels qu'elle devienne une égalité. On a par ailleurs l'inégalité

$$\frac{\|\delta\mathbf{x}\|}{\|\mathbf{x}\|} \leq \text{cond}\{\mathbf{A}\} \frac{\|\Delta\mathbf{A}\|}{\|\mathbf{A}\|} \frac{1}{1 - \|\mathbf{A}^{-1}\| \|\Delta\mathbf{A}\|}$$

Théorème 5 Pour toute matrice $\mathbf{A}$

1. $\text{cond}\{\mathbf{A}\} \geq 1$, $\text{cond}\{\mathbf{A}\} = \text{cond}\{\mathbf{A}^{-1}\}$, $\text{cond}\{\alpha\mathbf{A}\} = \text{cond}\{\mathbf{A}\}$ si $\alpha \neq 0$.
2. $\text{cond}_2(\mathbf{A}) = \frac{\mu_n(\mathbf{A})}{\mu_1(\mathbf{A})}$ où $\mu_1(\mathbf{A})$ et $\mu_n(\mathbf{A})$ sont respectivement la plus petite et la plus grande valeur singulière de la matrice $\mathbf{A}$.
3. Si $\mathbf{A}$ est une matrice normale; alors

$$\text{cond}_2(\mathbf{A}) = \frac{\max_i |\lambda_i(\mathbf{A})|}{\min_i |\lambda_i(\mathbf{A})|}$$

où $\lambda_i(\mathbf{A})$ sont les valeurs propres de $\mathbf{A}$.

4. Si $\mathbf{A}$ est une matrice unitaire ou orthogonale $\text{cond}_2(\mathbf{A}) = 1$.
5. Le conditionnement $\text{cond}_2(\mathbf{A})$ est invariante par transformation unitaire:

$$\mathbf{U}\mathbf{U}^* = \mathbf{I} \longrightarrow \text{cond}_2(\mathbf{AU}) = \text{cond}_2(\mathbf{UA}) = \text{cond}_2(\mathbf{U}^*\mathbf{AU}) = \text{cond}_2(\mathbf{A})$$

Dans l'exemple (1) on a

$$\lambda_1 = 0.01015 < \lambda_2 = 0.8431 < \lambda_3 = 3.858 < \lambda_4 = 30.2887$$

$$\text{cond}_2(\mathbf{A}) = \frac{\lambda_4}{\lambda_1} = 2984$$

On note que le calcul du $\text{cond}_2(\mathbf{A})$ repose sur la connaissance des valeurs propres extrêmes de la matrice $\mathbf{A}$. En l'absence de cette information on peut utiliser la relation suivante

$$\|\mathbf{A}\|_2 \leq \|\mathbf{A}\|_E \leq \sqrt{n}\|\mathbf{A}\|_2$$

à condition toutefois qu'on connaisse la matrice $\mathbf{A}^{-1}$. Un système linéaire est autant mieux conditionné que le nombre $\text{cond}\{\mathbf{A}\}$ est voisin de 1 (sous-entendu relativement à une norme matricielle subordonnée particulière). Un système linéaire $\mathbf{Ax} = \mathbf{b}$ à matrice orthogonale (ou unitaire) est très bien conditionné puisque $\text{cond}_2(\mathbf{A}) = 1$. Les transformations orthogonales ou unitaires préservent le conditionnement $\text{cond}_2(\mathbf{A})$.

E.8.2 Conditionnement d'un problème de calcul des valeurs singulières d'une matrice

Considérons la matrice carrée d'ordre n

$$\mathbf{A}(\epsilon) = \begin{pmatrix} 0 & & & \epsilon \\ 1 & 0 & & \\ & 1 & & \\ & & \ddots & \ddots \\ & & & 1 & 0 \end{pmatrix}$$

Pour $n = 40$ les valeurs propres de $\mathbf{A}(0)$ sont égales à 0 : par contre les valeurs propres de $\mathbf{A}(10^{-4})$ ont toutes les mêmes module 10^{-1} (les racine n -ème de ϵ). On constate ainsi que l'on a besoin dans ce problème de calcul des valeurs propres; de mesurer le conditionnement de la matrice.

Théorème 6 (de Bauer-Fike) *Soit $\mathbf{A}$ une matrice diagonalisable; $\mathbf{P}$ une matrice telle que*

$$\mathbf{P}^{-1}\mathbf{AP} = \text{diag}[\lambda_i]$$

et $\|\cdot\|$ une norme matricielle telle que

$$\|\text{diag}[d_i]\| = \max_i \|d_i\|$$

pour toute matrice diagonale. Alors pour toute matrice $\Delta\mathbf{A}$

$$\text{sp}[\mathbf{A} + \Delta\mathbf{A}] \in \prod_{i=1}^N x_i \quad \text{où} \quad x_i = \{z \in \mathbb{C} \mid |z - \lambda_i| \leq \text{cond}\{\mathbf{P}\} \|\Delta\mathbf{A}\|\}$$

E.9 Méthodes récursives de résolution des systèmes d'équations linéaires

Principe

On considère le système d'équations linéaires

$$\mathbf{R}\mathbf{a} = -\mathbf{d}$$

où $\mathbf{R}$ est une matrice de dimensions $[N, N]$; $\mathbf{a}$ un vecteur de dimension $[N]$ contenant les paramètres inconnus; et $\mathbf{d}$ un vecteur de données de dimension $[N]$. Supposons que $\mathbf{R}$ soit partitionnée de façon suivante:

$$\mathbf{R}_{m+1} = \begin{pmatrix} \mathbf{R}_m & \mathbf{x}_m \\ \mathbf{z}_m^t & y_m \end{pmatrix}, \quad 1 \leq m \leq N-1$$

$$\mathbf{R}_m \mathbf{a}_m = -\mathbf{d}_m, \quad \mathbf{R}_{m+1} \mathbf{a}_{m+1} = -\mathbf{d}_{m+1}, \quad \mathbf{d}_{m+1}^t = [\mathbf{d}_m^t, d_{m+1}]$$

Appliquant le lemme d'inversion de matrices en bloc à la matrice ainsi partitionnée; on obtient:

$$\alpha_m \mathbf{R}_{m+1}^{-1} = \begin{pmatrix} \alpha_m \mathbf{R}_m^{-1} + \mathbf{w}_m \mathbf{v}_m^t & \mathbf{w}_m \\ \mathbf{v}_m^t & 1 \end{pmatrix}, \quad 1 \leq m \leq N-1$$

où

$$\alpha_m = y_m + \mathbf{z}_m^t \mathbf{w}_m = y_m + \mathbf{v}_m^t \mathbf{x}_m = \frac{\det[\mathbf{R}_{m+1}]}{\det[\mathbf{R}_m]}$$

$$\mathbf{w}_m = -\mathbf{R}_m^{-1} \mathbf{x}_m, \quad \mathbf{v}_m = -\mathbf{R}_m^{-t} \mathbf{z}_m$$

Si on s'intéresse à la résolution de l'équation (1.1); plutôt qu'à l'inversion de la matrice (1.2); ces équations peuvent être écrites sous une autre forme:

$$\mathbf{a}_{m+1} = \begin{pmatrix} \mathbf{a}_m \\ 0 \end{pmatrix} + k_{m+1} \begin{pmatrix} \mathbf{w}_m \\ 1 \end{pmatrix}$$

$$k_{m+1} = -\frac{\beta_m}{\alpha_m}, \quad \beta_m = \mathbf{v}_m^t \mathbf{d}_m + d_{m+1} = \mathbf{z}_m^t \mathbf{a}_m + d_{m+1}$$

Le coût de calcul de cet algorithme est $O(N^3)$. Son intérêt particulier est dans le fait qu'on dispose de la valeur du déterminant de la matrice à chaque itération; ce qui est utile dans le contrôle de stabilité de l'algorithme pour l'inversion des matrices mal-conditionnées. D'autre part; parmi tous les paramètres utilisés dans cet algorithme, c'est à dire $\mathbf{w}_m, \mathbf{v}_m, \alpha_m$ et β_m , seul ce dernier dépend des données.

Cas des matrices Symétriques

Dans ce cas on a

$$\mathbf{x}_m = \mathbf{z}_m = \mathbf{r}_m, \quad \text{et} \quad \mathbf{w}_m = \mathbf{v}_m = \mathbf{R}_m^{-1} \mathbf{r}_m$$

$$\alpha_m = y_m + \mathbf{r}_m^t \mathbf{w}_m$$

$$\alpha_m \mathbf{R}_{m+1}^{-1} = \begin{pmatrix} \alpha_m \mathbf{R}_m^{-1} + \mathbf{v}_m \mathbf{v}_m^t & \mathbf{v}_m \\ \mathbf{v}_m^t & 1 \end{pmatrix}, \quad 1 \leq m \leq N-1$$

Cas des matrices MID

$$\mathbf{w}_m = -\mathbf{R}_m^{-1} \mathbf{x}_m, \quad \mathbf{w}_{m+1} = -\mathbf{R}_{m+1}^{-1} \mathbf{x}_{m+1}$$

En utilisant les Matrices MID; c'est à dire:

$$\mathbf{x}_{m+1}^t = [\mathbf{x}_m^t; x_{m+1}], \quad \mathbf{z}_{m+1}^t = [\mathbf{z}_m^t; z_{m+1}]; \quad \mathbf{y}_{m+1}^t = [\mathbf{y}_m^t; y_{m+1}]$$

on obtient une équation de récurrence pour $\mathbf{w}_{m+1}$ et $\mathbf{v}_{m+1}$ en fonction de $\mathbf{w}_m$ et $\mathbf{v}_m$.

$$\mathbf{w}_{m+1} = \begin{pmatrix} \mathbf{w}_m \\ 0 \end{pmatrix} + k_{m+1}^+ \begin{pmatrix} \mathbf{w}_m \\ 1 \end{pmatrix}, \quad \mathbf{v}_{m+1} = \begin{pmatrix} \mathbf{v}_m \\ 0 \end{pmatrix} + k_{m+1}^- \begin{pmatrix} \mathbf{v}_m \\ 1 \end{pmatrix}$$

$$k_{m+1}^+ = -\frac{\beta_m^+}{\alpha_m} = \frac{\mathbf{v}_m^t \mathbf{x}_m + x_{m+1}}{\alpha_m}, \quad k_{m+1}^- = -\frac{\beta_m^-}{\alpha_m} = \frac{\mathbf{z}_m^t \mathbf{w}_m + z_{m+1}}{\alpha_m}$$

$$\alpha_m = \mathbf{z}_m^t \mathbf{w}_m + y_{m+1} = \mathbf{v}_m^t \mathbf{x}_m + y_{m+1}$$

On peut encore simplifier ces équations en trouvant une relation de récurrence pour $\alpha_m, \beta_m^+, \beta_m^-$ et β_m :

$$\alpha_m = \alpha_{m-1} + k_m^+ \beta_{m-1}^- + (y_{m+1} - y_m)$$

$$\beta_m^+ = (x_{m+1} - y_{m+1}) + \alpha_m, \quad \beta_m^- = (z_{m+1} - y_{m+1}) + \alpha_m$$

$$\beta_m = \beta_{m-1}(1 - k_m^-) + (d_{m+1} - d_m)$$

Ces équations simplifient plus si $y_m = x_m$ et / ou $y_m = z_m$. L'inversion de ces matrice nécessitent $O(p^2)$ opérations élémentaires.

Cas des matrices MIP

Il n'existe pas encore des méthodes rapides directes pour inversion des matrices MIP dans le cas général car ces matrices peuvent avoir un rang de déplacement entre $0, \dots, N$; $[N; N]$ étant les dimensions de la matrice. Par contre un cas particulier de ces matrices est la matrice de Toeplitz ($y_i = y_0, i = 1, \dots, N$) non-symétrique ($x_i \neq z_i$) ou symétrique ($x_i = z_i$).

Cas des matrices de Toeplitz

Une matrice de Toeplitz satisfait la propriété suivant:

$$\mathbf{J} \mathbf{R}_m \mathbf{J} = \mathbf{R}_m^t, \quad m = 1, \dots, N$$

où $\mathbf{J}$ est une matrice d'inversion d'ordre de dimensions $[M, M]$; ce qui vaut dire qu'une matrice de Toeplitz est une matrice symétrique par rapport à sa deuxième diagonale. Si on note par

$$\tilde{\mathbf{x}}_m = [x_1; x_2, \dots, x_m]^t, \quad \tilde{\mathbf{z}}_m = [x_1; x_2, \dots, x_m]^t$$

$$\mathbf{w}_{m+1} = -\mathbf{R}_{m+1}^{-1} \mathbf{J} \tilde{\mathbf{x}}_{m+1}, \quad \mathbf{w}_{m+1} = -\mathbf{J} \mathbf{R}_{m+1}^{-t} \tilde{\mathbf{x}}_{m+1}$$

$$\mathbf{v}_{m+1}^t = -\tilde{\mathbf{z}}_{m+1}^t \mathbf{J} \mathbf{R}_{m+1}^{-1} = -\tilde{\mathbf{z}}_{m+1}^t \mathbf{R}_{m+1}^{-t} \mathbf{J}$$

$$\mathbf{J} \mathbf{w}_{m+1} = \begin{pmatrix} \mathbf{J} \mathbf{w}_m \\ 0 \end{pmatrix} + k_{m+1}^+ \begin{pmatrix} \mathbf{v}_m \\ 1 \end{pmatrix}, \quad \mathbf{J} \mathbf{v}_{m+1} = \begin{pmatrix} \mathbf{J} \mathbf{v}_m \\ 0 \end{pmatrix} + k_{m+1}^- \begin{pmatrix} \mathbf{w}_m \\ 1 \end{pmatrix}$$

$$k_{m+1}^+ = -\frac{\beta_m^+}{\alpha_m} = -\frac{\mathbf{w}_m^t \tilde{\mathbf{x}}_m + x_{m+1}}{\alpha_m}, \quad k_{m+1}^- = -\frac{\beta_m^-}{\alpha_m} = -\frac{\tilde{\mathbf{z}}_m^t \mathbf{w}_m + z_{m+1}}{\alpha_m}$$

$$\alpha_m = \mathbf{z}_m^t \mathbf{w}_m + y_0 = \mathbf{x}_m^t \mathbf{v}_m + y_0$$

Notons que les récursions pour $\mathbf{w}$ et $\mathbf{v}$ sont couplées (ce qui n'était pas le cas pour les matrices MID). Le terme α_m peut être calculé d'une façon plus efficace par

$$\alpha_m = \alpha_{m-1} (1 - k_m^+ k_m^-)$$

On peut obtenir ces mêmes relations pour les matrices de Hankel en remarquant qu'une matrice de Hankel $\mathbf{H}$ et une matrice de Toeplitz $\mathbf{T}$ sont reliées par $\mathbf{H} = \mathbf{T} \mathbf{J}$ ou $\mathbf{H} = \mathbf{J} \mathbf{T}$. Ainsi les trois relations suivantes décrivent un même système d'équations linéaires:

$$\mathbf{H} \mathbf{a} = -\mathbf{d}, \quad \mathbf{T}(\mathbf{J} \mathbf{a}) = -\mathbf{d}, \quad \mathbf{T} \mathbf{a} = -\mathbf{J} \mathbf{d}$$

il suffit alors de remplacer $\mathbf{d}$ par $\mathbf{J} \mathbf{d}$ pour obtenir les relations nécessaires pour des matrices de Hankel.

Cas des matrices de Toeplitz Symétriques

Algorithme de Levinson :

Dans ce cas on a la relation $\mathbf{w} = \mathbf{v}$ ce qui va permettre de simplifier encore plus les relations précédentes. On a

$$k^+ = k^- = k$$

$$\mathbf{J} \mathbf{w}_{m+1} = \begin{pmatrix} \mathbf{J} \mathbf{w}_m \\ 0 \end{pmatrix} + k_{m+1} \begin{pmatrix} \mathbf{v}_m \\ 1 \end{pmatrix}$$

$$\beta = \mathbf{v}_{m+1}^t \tilde{\mathbf{x}}_m + x_{m+1}, \quad \alpha_m = \alpha_{m-1} (1 - k_m^2)$$

Algorithme de Durbin :

Dans le cas particulier où le vecteur $\mathbf{d}$ peut être écrit sous la forme

$$\mathbf{d}_N^t = [\mathbf{d}_{N-1}^t + d_N]$$

comme c'est le cas du problème de la modélisation AR; on a la relation

$$\mathbf{w}_m = \mathbf{J} \mathbf{a}_m, \quad m = 1, 2, \dots, N-1$$

Cas des matrices circulantes

Dans le cas des matrices circulantes on a

$$\mathbf{z}_{N-1} = \mathbf{J}\mathbf{x}_{N-1}, \quad \mathbf{v}_{N-1} = \mathbf{J}\mathbf{w}_{N-1}$$

Il faut noter que l'on peut diagonaliser les matrices circulantes par TFD et obtenir des algorithmes non récurrents.

Cas des matrices de Toeplitz en bande

E.10 Diverses propriétés

Lemme de factorisation d'une matrice

Soit $\mathbf{A}$ une matrice de dimension $[N, P]$. On peut factoriser la matrice $[\mathbf{I}_N \pm \mathbf{A}\mathbf{A}^t]$ de façon suivante:

$$\mathbf{I}_N \pm \mathbf{A}\mathbf{A}^t = [\mathbf{I}_N \pm \mathbf{A}\mathbf{R}^{-1}\mathbf{A}^t][\mathbf{I}_N \pm \mathbf{A}\mathbf{R}^{-1}\mathbf{A}^t]^t$$

où $\mathbf{R} = \mathbf{I}_P + [\mathbf{I}_P \pm \mathbf{A}^t\mathbf{A}]^{1/2}$.

Racines carrées d'une matrice définie positive

Toute matrice définie positive admet une racine carrée. On appelle racine carrée d'une matrice $\mathbf{A}$ toute matrice carrée $\mathbf{S}$, telle que

$$\mathbf{A} = \mathbf{S}\mathbf{S}^t, \quad \mathbf{S} = \mathbf{A}^{1/2}, \quad \mathbf{S}^t = \mathbf{A}^{t/2}.$$

Une telle factorisation n'est pas unique. Si $\mathbf{T}$ est orthogonale ($\mathbf{T}\mathbf{T}^t = \mathbf{I}$), alors $\mathbf{S}\mathbf{T}$ est aussi racine carrée de $\mathbf{A}$ puisque $\mathbf{A} = \mathbf{S}\mathbf{T}\mathbf{T}^t\mathbf{S}^t = \mathbf{S}\mathbf{S}^t$.

Décomposition de Cholesky d'une matrice définie positive

Soit $\mathbf{A}$ une matrice symétrique, réelle, et définie positive d'ordre N . Il existe alors une matrice $\mathbf{O}$ orthogonale réelle d'ordre N telle que $\mathbf{D} = \mathbf{O}^{-1}\mathbf{A}\mathbf{O} = \mathbf{O}^t\mathbf{A}\mathbf{O}$ soit diagonale. De plus, si $\mathbf{A}$ est de rang $K < N$, elle admet une décomposition de la forme

$$\mathbf{A} = \mathbf{S}_1\mathbf{S}_1^t - \mathbf{S}_2\mathbf{S}_2^t$$

où $\mathbf{S}_1$ et $\mathbf{S}_2$ sont de dimensions respectives $[N, \alpha]$ et $[N, K - \alpha]$, α étant le nombre de valeurs propres positives non nulles de $\mathbf{A}$. La décomposition de Cholesky consiste à imposer aux racines carrées d'être des matrices triangulaires. La décomposition est alors unique.

Le rang de déplacement d'une matrice

Le rang de déplacement d'une matrice est le rang de la différence entre la matrice initiale et la matrice décalée d'un cran le long de sa diagonale principale. On définit deux rangs de déplacement, inférieur δ_- et supérieur δ_+ , à partir des deux déplacements possibles (vers le haut-gauche, et vers le bas-droite):

$$\delta_+(\mathbf{R}) = \text{rang} [\mathbf{R} - \mathbf{Z}\mathbf{R}\mathbf{Z}^t], \quad \delta_-(\mathbf{R}) = \text{rang} [\mathbf{R} - \mathbf{Z}^t\mathbf{R}\mathbf{Z}]$$

où $\mathbf{Z}$ est une matrice de décalage vers le bas donnée par:

$$\mathbf{Z} = \begin{pmatrix} 0 & . & . & . & 0 \\ 1 & 0 & . & . & 0 \\ 0 & 1 & 0 & . & 0 \\ . & 0 & 1 & 0 & 0 \\ . & . & 0 & 1 & 0 \end{pmatrix}$$

Pour une matrice quelconque $\mathbf{R}$ on peut toujours construire deux partitions:

$$\mathbf{R} = \begin{pmatrix} r_0 & \mathbf{x}^t \\ \mathbf{y} & \mathbf{R}_i \end{pmatrix} = \begin{pmatrix} \mathbf{R}_s & \mathbf{u} \\ \mathbf{v}^t & s_0 \end{pmatrix}$$

Ainsi, pour cette matrice on a :

$$\mathbf{R} - \mathbf{Z}\mathbf{R}\mathbf{Z}^t = \begin{pmatrix} \mathbf{R}_i - \mathbf{R}_s & \mathbf{u} \\ \mathbf{v}^t & s_0 \end{pmatrix}, \quad \text{et} \quad \mathbf{R} - \mathbf{Z}^t\mathbf{R}\mathbf{Z} = \begin{pmatrix} r_0 & \mathbf{x}^t \\ \mathbf{y} & \mathbf{R}_i - \mathbf{R}_s \end{pmatrix}$$

Dans le cas des matrices symétriques on a $\mathbf{u} = \mathbf{v} = \mathbf{s}$, $\mathbf{x} = \mathbf{y} = \mathbf{r}$. On peut alors construire deux partitions de la matrice initiale $\mathbf{R}$:

$$\mathbf{R} = \begin{pmatrix} r_0 & \mathbf{r}^t \\ \mathbf{r} & \mathbf{R}_i \end{pmatrix} = \begin{pmatrix} \mathbf{R}_s & \mathbf{s} \\ \mathbf{s}^t & s_0 \end{pmatrix}$$

Dans ce cas on a

$$\mathbf{R} - \mathbf{Z}\mathbf{R}\mathbf{Z}^t = \begin{pmatrix} \mathbf{R}_i - \mathbf{R}_s & \mathbf{s} \\ \mathbf{s}^t & s_0 \end{pmatrix} \quad \text{et} \quad \mathbf{R} - \mathbf{Z}^t\mathbf{R}\mathbf{Z} = \begin{pmatrix} r_0 & \mathbf{r}^t \\ \mathbf{r} & \mathbf{R}_i - \mathbf{R}_s \end{pmatrix}$$

Pour une matrice de Toeplitz, on vérifie facilement, qu'on a: $\delta_+(\mathbf{T}) = \delta_-(\mathbf{T}) = 2$ On démontre aussi facilement [] que :

$$\delta_+(\mathbf{R}) = \delta_-\mathbf{R}^{-1} \quad \text{et} \quad \delta_-(\mathbf{R}) = \delta_+\mathbf{R}^{-1} \det[\delta_+(\mathbf{R}) - \delta_-(\mathbf{R})] \leq 2$$

Distance à Toeplitz

Le rang de déplacement $\delta_-(\mathbf{R})$ ou $\delta_+(\mathbf{R})$ ne constitue pas à proprement parler une mesure de la distance à Toeplitz puisque $\delta_-(\mathbf{T}) = \delta_-(\mathbf{T}) = 2$. C'est la raison pour laquelle on utilise aussi une définition différente du rang de déplacement qui fournit directement une *distance à Toeplitz*.

Soit $\mathbf{R}$ une matrice-bloc quelconque de dimensions $[NP, NP]$ $\mathbf{R} = [\mathbf{R}_{ji}]$, $i, j = 1, \dots, N$ où $\mathbf{R}_{ij}$ sont des matrices de dimensions $[P, P]$. Le déplacement est alors défini par

$$\delta\mathbf{R} = [\mathbf{R}_{(i+1),(j+1)} - \mathbf{R}_{ij}], \quad i, j = 1, \dots, N-1$$

et le rang de déplacement devient $\alpha = \text{partie entière de } \left[\frac{\text{rang}[\delta\mathbf{R}]}{P} \right]$.

- Quand les éléments de $\mathbf{R}$ sont scalaires, $P = 1$ et $\alpha = \text{rang}[\delta\mathbf{R}]$.
- Si $\mathbf{R}$ est une matrice quelconque, $\delta\mathbf{R}$ peut être de rang plein et alors $\alpha = N - 1$.
- Si $\mathbf{R}$ est Toeplitz, alors $\alpha = 0$. On a en général $0 \leq \alpha \leq N - 1$, et α peut servir directement de mesure de la distance de $\mathbf{R}$ à Toeplitz.

Annexe F

Quelques transformations utiles

L'objet de cette annexe est de rappeler un certain nombre de transformation linéaires utiles en 1-D, 2-D et n -D. Nous donnons aussi l'expression de la transformée inverse de ces transformations lorsqu'elle existe.

F.1 Transformations 1-D

$$F(u) = \int f(x)h(u,x) \, dx$$

Identité: $h(u,x) = \delta(u-x) \quad F(u) = \int f(x)\delta(u-x) \, dx = f(u)$

Convolution: $h(u,x) = h(u-x) \quad F(u) = \int f(x)h(u-x) \, dx = f(u) * h(u)$

Corrélation: $h(u,x) = h(u+x) \quad F(u) = \int f(x)h(u+x) \, dx = f(u) \star h(u)$

Fourier: $h(u,x) = \exp[-jux] \quad F(u) = \int f(x) \exp[-jux] \, dx$

Laplace: $h(u,x) = \exp[-ux] \quad F(u) = \int f(x) \exp[-ux] \, dx$

Transformations unitaires :

$$F(u) = \int f(x)h(u,x) \, dx$$

$$f(x) = \int F(u)h^*(u,x) \, du$$

Fourier: $h(u,x) = \exp[-jux]$

$$F(u) = \frac{1}{\sqrt{2\pi}} \int f(x) \exp[-jux] \, dx$$

$$f(x) = \frac{1}{\sqrt{2\pi}} \int F(u) \exp[+jux] \, du$$

Cosinus: $h(u,x) = \cos(ux)$

$$F(u) = \frac{2}{\sqrt{\pi}} \int_0^\infty f(x) \cos(ux) \, dx$$

$$f(x) = \frac{2}{\sqrt{\pi}} \int_0^\infty F(u) \cos(ux) \, du$$

Sinus: $h(u,x) = \sin(ux)$

$$F(u) = \frac{2}{\sqrt{\pi}} \int_0^\infty f(x) \sin(ux) \, dx$$

$$f(x) = \frac{2}{\sqrt{\pi}} \int_0^\infty F(u) \sin(ux) \, du$$

Hartley: $h(u,x) = \frac{1}{\sqrt{2\pi}}[\cos(ux) + \sin(ux)]$

$$F(u) = \frac{1}{\sqrt{2\pi}} \int_0^\infty f(x)[\cos(ux) + \sin(ux)] \, dx$$

$$f(x) = \frac{1}{\sqrt{2\pi}} \int_0^\infty F(u)[\cos(ux) + \sin(ux)] \, du$$

Hilbert: $h(u,x) = \frac{1}{\pi}(u-x)^{-1}$

$$F(u) = \frac{1}{\pi} \int \frac{f(x)}{(x-u)} \, dx$$

$$f(x) = \frac{1}{\pi} \int \frac{F(u)}{(x-u)} \, du$$

Hankel: $h(u,x) = J_\nu(ux)$

$$F(u) = \int_0^\infty f(x) J_\nu(ux) \, dx$$

$$f(x) = \int_0^\infty F(u) J_\nu(ux) \, du$$

Mellin: $h(u,x) = x^{u-1}$

$$F(u) = \int f(x)x^{u-1} \, dx$$

Abel: $h(u,x) = \frac{1}{\Gamma(\alpha)}(x-u)^{\alpha-1}$

$$F(u) = \frac{1}{\Gamma(\alpha)} \int f(x)(x-u)^{\alpha-1} \, dx$$

Weyl: $h(u,x) = \frac{1}{\Gamma(\alpha)}(u-x)^{\alpha-1}$

$$F(u) = \frac{1}{\Gamma(\alpha)} \int f(x)(u-x)^{\alpha-1} \, dx$$

F.2 Transformations 2-D

$$F(u,v) = \iint f(x,y)h(u,v;x,y) \, dx \, dy$$

Identité :

$$\begin{aligned} h(u,v;x,y) &= \delta(u-x;v-y) \\ F(u,v) &= \iint f(x,y)\delta(u-x;v-y) \, dx \, dy = f(x,y) \end{aligned}$$

Séparable :

$$\begin{aligned} h(u,v;x,y) &= h_1(u,x)h_2(v,y) \\ F(u,v) &= \iint f(x,y)h_1(u,x)h_2(v,y) \, dx \, dy \end{aligned}$$

Séparable & Symétrique :

$$\begin{aligned} h(u,v;x,y) &= h_1(u,x)h_1(v,y) \\ F(u,v) &= \iint f(x,y)h_1(u,x)h_1(v,y) \, dx \, dy \end{aligned}$$

Convolution :

$$\begin{aligned} h(u,v;x,y) &= h(u-x;v-y) \\ F(u,v) &= \iint f(x,y)h(u-x;v-y) \, dx \, dy = f(x,y) * h(x,y) \end{aligned}$$

Corrélation :

$$\begin{aligned} h(u,v;x,y) &= h(u+x;v+y) \\ F(u,v) &= \iint f(x,y)h(u+x;v+y) \, dx \, dy = f(x,y) \star h(x,y) \end{aligned}$$

Laplace :

$$\begin{aligned} h(u,v;x,y) &= \exp [-(ux+vy)] \\ F(u,v) &= \iint f(x,y) \exp [-(ux+vy)] \, dx \, dy \end{aligned}$$

Transformations 2-D unitaires

Fourier: $h(u, v; x, y) = \exp[-j(ux + vy)]$
 $F(u, v) = \iint f(x, y) \exp[-j(ux + vy)] \, dx \, dy$

Cosinus: $h(u, v; x, y) = \cos(ux) \cos(vy)$
 $F(u, v) = \iint f(x, y) \cos(ux) \cos(vy) \, dx \, dy$

Sinus: $h(u, v; x, y) = \sin(ux) \sin(vy)$
 $F(u, v) = \iint f(x, y) \sin(ux) \sin(vy) \, dx \, dy$

Abel: $h(u, v; x, y) = [(x - u)^2 + (y - v)^2]^{(\alpha-1)/2}$
 $F(u, v) = \iint f(x, y) [(u - x)^2 + (v - y)^2]^{-(\alpha-1)/2} \, dx \, dy$

Weyl: $h(u, v; x, y) = [(u - x)^2 + (v - y)^2]^{-(\alpha-1)/2}$
 $F(u, v) = \iint f(x, y) [(u - x)^2 + (v - y)^2]^{-(\alpha-1)/2} \, dx \, dy$

Hilbert: $h(u, v; x, y) = [(u - x)^2 + (v - y)^2]^{-1/2}$
 $F(u, v) = \iint \frac{f(x, y)}{\sqrt{(u - x)^2 + (v - y)^2}} \, dx \, dy$

Radon: $h(r, \phi; x, y) = \delta(r - x \cos \phi - y \sin \phi)$
 $F(r, \phi) = \iint f(x, y) \delta(r - x \cos \phi - y \sin \phi) \, dx \, dy$

F.3 Transformations n -D

$$F(\mathbf{u}) = \int f(\mathbf{x})h(\mathbf{u}, \mathbf{x}) \, d\mathbf{x}$$

Identité: $h(\mathbf{u}, \mathbf{x}) = \delta(\mathbf{u} - \mathbf{x}) \quad F(\mathbf{u}) = \int f(\mathbf{x})\delta(\mathbf{u} - \mathbf{x}) \, d\mathbf{x} = f(\mathbf{x})$

Convolution: $h(\mathbf{u}, \mathbf{x}) = h(\mathbf{u} - \mathbf{x}) \quad F(\mathbf{u}) = \int f(\mathbf{x})h(\mathbf{u} - \mathbf{x}) \, d\mathbf{x} = f(\mathbf{x}) * h(\mathbf{x})$

Corrélation: $h(\mathbf{u}, \mathbf{x}) = h(\mathbf{u} + \mathbf{x}) \quad F(\mathbf{u}) = \int f(\mathbf{x})h(\mathbf{u} + \mathbf{x}) \, d\mathbf{x} = f(\mathbf{x}) \star h(\mathbf{x})$

Laplace: $h(\mathbf{u}, \mathbf{x}) = \exp[-\mathbf{u}^t \mathbf{x}] \quad F(\mathbf{u}) = \int f(\mathbf{x}) \exp[-\mathbf{u}^t \mathbf{x}] \, d\mathbf{x}$

Transformations unitaires

$$F(\mathbf{u}) = \int f(\mathbf{x})h(\mathbf{u}, \mathbf{x}) \, d\mathbf{x}$$

$$f(\mathbf{x}) = \int F(\mathbf{u})h^*(\mathbf{u}, \mathbf{x}) \, d\mathbf{u}$$

Fourier: $h(\mathbf{u}, \mathbf{x}) = \exp[-j\mathbf{u}^t \mathbf{x}] \quad F(\mathbf{u}) = \int f(\mathbf{x}) \exp[-j\mathbf{u}^t \mathbf{x}] \, d\mathbf{x}$
 $f(\mathbf{x}) = \int F(\mathbf{u}) \exp[+j\mathbf{u}^t \mathbf{x}] \, d\mathbf{u}$

Abel: $h(\mathbf{u}, \mathbf{x}) = \|\mathbf{x} - \mathbf{u}\|^{(\alpha-1)/2} \quad F(\mathbf{u}) = \int f(\mathbf{x})\|\mathbf{x} - \mathbf{u}\|^{(\alpha-1)/2} \, d\mathbf{x}$

Weyl: $h(\mathbf{u}, \mathbf{x}) = \|\mathbf{u} - \mathbf{x}\|^{(\alpha-1)/2} \quad F(\mathbf{u}) = \int f(\mathbf{x})\|\mathbf{u} - \mathbf{x}\|^{(\alpha-1)/2} \, d\mathbf{x}$

Hilbert: $h(\mathbf{u}, \mathbf{x}) = \|\mathbf{u} - \mathbf{x}\|^{-1/2} \quad F(\mathbf{u}) = \int f(\mathbf{x})\|\mathbf{u} - \mathbf{x}\|^{-1/2} \, d\mathbf{x}$

Radon: $h(r, \boldsymbol{\xi}; \mathbf{x}) = \delta(r - \boldsymbol{\xi}^t \mathbf{x}) \quad F(r, \boldsymbol{\xi}) = \int f(\mathbf{x})\delta(r - \boldsymbol{\xi}^t \mathbf{x}) \, d\mathbf{x}$

F.4 Transformations linéaires dans le cas discret

$$g(m) = \sum_{n=0}^{N-1} f(n)h(m,n), \quad m = 0, \dots, N-1 \quad \longrightarrow \quad \mathbf{g} = \mathbf{H}\mathbf{f}$$

Identité: $h(m,n) = \delta(m-n), \quad m,n = 0, \dots, N-1$
 $\mathbf{H}$ matrice Identité

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

Convolution: $h(m,n) = h(m-n), \quad m,n = 0, \dots, N-1$
 $\mathbf{H}$ matrice de Toeplitz

$$\mathbf{H} = \begin{pmatrix} h_0 & h_{-1} & h_{-2} & \cdots & \cdots & h_{-N+1} \\ h_1 & h_0 & h_{-1} & & & \vdots \\ h_2 & h_1 & & & & \vdots \\ \vdots & & & & & \\ \vdots & & & & & h_{-2} \\ \vdots & & & & & h_{-1} \\ h_{N-1} & \cdots & & h_2 & h_1 & h_0 \end{pmatrix}$$

Corrélation: $h(m,n) = h(m+n), \quad m,n = 0, \dots, N-1$
 $\mathbf{H}$ matrice de Hankel

$$\mathbf{H} = \begin{pmatrix} h_0 & h_1 & h_2 & h_3 & \cdots & h_{N-1} \\ h_1 & h_2 & h_3 & & h_{N-1} & \vdots \\ h_2 & h_3 & & & & \\ h_3 & & & & & \\ \vdots & & & & & \vdots \\ & & & & h_{2N-4} & \\ \vdots & h_{N-1} & & h_{2N-4} & h_{2N-3} & \\ h_{N-1} & \cdots & h_{2N-4} & h_{2N-3} & h_{2N-2} \end{pmatrix}$$

Transformations unitaires dans le cas discret

Fourier: $h(m,n) = \frac{1}{N} \exp[-j2\pi mn/N]$, $m,n = 0, \dots, N-1$

H matrice circulante de TFD

$$\mathbf{H} = \frac{1}{N} \begin{pmatrix} 1 & 1 & 1 & \dots & \dots & 1 \\ 1 & w^1 & w^2 & & & w^{N-1} \\ 1 & w^2 & w^4 & & & w^{2(N-1)} \\ \vdots & & & & & \\ \vdots & & & & & w^{(N-3)(N-1)} \\ 1 & w^{N-1} & & w^{(N-2)(N-1)} & w^{(N-1)(N-1)} \end{pmatrix}$$

avec $w = \exp[-j2\pi/N]$

Cosinus: $h(m,n) = \begin{cases} \frac{1}{\sqrt{N}} & \text{si } m = 0 \\ \frac{2}{\sqrt{N}} \cos(\pi m(n+1/2)/N) & \text{si } m \neq 0 \end{cases}$, $m,n = 0, \dots, N-1$

H matrice circulante de TCD

$$\mathbf{H} = \begin{pmatrix} 1 & 1 & 1 & \dots & \dots & 1 \\ 1 & \cos(w) & \cos(2w) & & & \cos[(N-1)w] \\ 1 & \cos(2w) & & & & \vdots \\ \vdots & & & & & \vdots \\ \vdots & & & & & \cos[(N-3)(N-1)w] \\ 1 & \cos[(N-1)w] & & & \cos[(N-2)(N-1)w] \\ & & & & \cos[(N-1)(N-1)w] \end{pmatrix}$$

avec $w = -j2\pi/N$

Sinus: $h(m,n) = \frac{2}{\sqrt{N}} \sin(\pi m(n+1/2)/N)$, $m,n = 0, \dots, N-1$

H matrice circulante de TSD

$$\mathbf{H} = \begin{pmatrix} 1 & 1 & 1 & \dots & \dots & 1 \\ 1 & \sin(w) & \sin(2w) & & & \sin[(N-1)w] \\ 1 & \sin(2w) & & & & \vdots \\ \vdots & & & & & \vdots \\ \vdots & & & & & \sin[(N-3)(N-1)w] \\ 1 & \sin[(N-1)w] & & & \sin[(N-2)(N-1)w] \\ & & & & \sin[(N-1)(N-1)w] \end{pmatrix}$$

avec $w = -j2\pi/N$

Walsh : $h(m,n) = \prod_{i=0}^k (-1)^{b_i(m)b_{n-1-i}(n)}$, $N = 2^k$
H matrice symétrique avec des éléments -1 ou +1
 $b_k(z)$ est le k ème bit de la représentation binaire de z .

$$\mathbf{H} = \begin{pmatrix} 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 & -1 & -1 & -1 & -1 \\ 1 & 1 & -1 & -1 & 1 & 1 & -1 & -1 \\ 1 & 1 & -1 & -1 & -1 & -1 & 1 & 1 \\ 1 & -1 & 1 & -1 & 1 & -1 & 1 & -1 \\ 1 & -1 & 1 & -1 & -1 & 1 & -1 & 1 \\ 1 & -1 & -1 & 1 & 1 & -1 & -1 & 1 \\ 1 & -1 & 1 & -1 & 1 & -1 & 1 & -1 \end{pmatrix}$$

Hadamard : $h(m,n) = (-1)^{\sum_{i=0}^k b_i(m)b_i(n)}$, $N = 2^k$
H matrice symétrique avec des éléments -1 ou +1
 $b_k(z)$ est le k ème bit de la représentation binaire de z .
 Exemple: $z = 6 = 110 \rightarrow b_0(z) = 0, b_1(z) = 1, b_2(z) = 1$.

$$\mathbf{H} = \begin{pmatrix} 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 \\ 1 & -1 & 1 & -1 & 1 & -1 & 1 & -1 \\ 1 & 1 & -1 & -1 & 1 & 1 & -1 & -1 \\ 1 & -1 & -1 & 1 & 1 & -1 & -1 & 1 \\ 1 & 1 & 1 & 1 & -1 & -1 & -1 & -1 \\ 1 & -1 & 1 & -1 & -1 & 1 & -1 & 1 \\ 1 & 1 & -1 & -1 & -1 & -1 & 1 & 1 \\ 1 & -1 & -1 & 1 & -1 & 1 & 1 & -1 \end{pmatrix}$$

Table des figures

2.1	Tomographie à rayons X	15
2.2	Notations utilisées pour la définition de la projection $p(r, \phi)$ d'une fonction $f(x, y)$	16
2.3	Reconstruction d'image en tomographie à rayons X	16
2.4	Théorème de projection en tomographie X.	18
2.5	Reconstruction d'image en tomographie X en passant dans le domaine de FOURIER.	19
2.6	Imagerie active et passive.	20
2.7	Différents modes d'imagerie: par réflexion, par transmission et mode mixte.	21
2.8	Géométrie de l'imagerie micro ondes.	23
2.9	Géométrie de la tomographie de diffraction 2D.	25
2.10	Synthèse de FOURIER en tomographie de diffraction 2D.	26
2.11	Géométrie de la tomographie de diffraction 3D.	27
2.12	Synthèse de FOURIER en tomographie de diffraction 3D.	27
2.13	Géométrie de l'imagerie par courants de FOUCAULT.	29
2.14	Synthèse de FOURIER en imagerie par courants de FOUCAULT.	34
2.15	Principe de la tomographie d'émission: TEP et SPECT.	35
2.16	Système de mesures en tomographie d'émission: TEP et SPECT.	36
2.17	Principe de l'imagerie RMN.	37
2.18	Modes de remplissage du domaine de FOURIER en imagerie RMN.	39
2.19	La géométrie de l'imagerie radar.	40
2.20	Remplissage du domaine de FOURIER en imagerie radar.	41
3.1	Déconvolution de signaux.	49
3.2	Débruitage de signaux.	50
3.3	Débruitage d'images.	50
3.4	Restauration d'image.	51
3.5	Discretisation de la transformée de RADON.	52
3.6	Reconstruction d'image en tomographie X.	53
3.7	Radiographie des objets cylindriques.	54
3.8	Reconstruction d'image en tomographie X pour des objets cylindriques.	55
3.9	Synthèse d'ouverture en radioastronomie.	56
3.10	Synthèse de FOURIER en tomographie X.	57
3.11	Synthèse de FOURIER en imagerie RMN.	58
3.12	Synthèse de FOURIER en tomographie par ondes diffractées.	58
3.13	Géométrie de l'imagerie par ondes diffractées.	61
4.1	Analyse fonctionnelle et topologie	64
4.2	Image et noyau d'un opérateur et de son adjoint: $F = G = \mathbf{R}^2$	65

4.3	Ensemble et opérateur compacts	66
4.4	Image et noyau d'un opérateur et de son adjoint : $F = G = \mathbb{R}^3$	66
4.5	Image et noyau d'un opérateur et de son adjoint : $F = \mathbb{R}^3, G = \mathbb{R}^2$	67
4.6	Déconvolution : un problème mal-posé.	72
4.7	Déconvolution : un problème mal posé	73
5.1	Inversion généralisée et espace des solutions.	84
6.1	Déconvolution par la méthode de Hunt : Exemple 1	99
6.2	Déconvolution par filtre de Wiener : Exemple 1	100
6.3	Déconvolution par la méthode de Hunt : Exemple 2	101
6.4	Déconvolution par filtre de Wiener : Exemple 2	102
6.5	Restauration par la méthode de Hunt : Exemple 1	103
6.6	Restauration par filtre de Wiener : Exemple 1	104
6.7	Restauration par la méthode de Hunt : Exemple 2	105
6.8	Restauration par filtre de Wiener : Exemple 2	106
7.1	Méthodes analytiques pour l'inversion de la TR.	110
7.2	Reconstruction par rétroprojection filtrée	111
7.3	Illustration de la méthode de décomposition sur une base.	118
7.4	chantillonnage et fonctions bases.	119
8.1	Filtrage de WIENER.	125
9.1	Différents types de signaux.	157
10.1	Sites, voisinage et champ.	170
10.2	Principe de base de la technique du GNC.	186
10.3	Fonction quadratique tronquée et sa relaxation	187
11.1	Choix du paramètre de régularisation	192
11.2	Effet du coefficients de la régularisation sur le résultat d'une déconvolution.	193
11.3	Modèle d'observation et d'inversion en régularisation quadratique.	195
11.4	Choix du paramètre de régularisation par adéquation aux données	196
12.1	Méthodes naïves en déconvolution.	217
12.2	Méthodes naïves en déconvolution : Filtrage inverse	218
12.3	Filtre de Wiener en déconvolution.	220
12.4	Déconvolution par régularisation d'ordre zéro.	222
12.5	Déconvolution par régularisation d'ordre un.	223
12.6	Identification de la réponse impulsionnelle par régularisation.	225
12.7	Déconvolution aveugle.	228
12.8	Restauration d'image sans et avec régularisation	231
12.9	Restauration aveugle d'image	232
12.10	Restauration aveugle d'image	233
12.11	Restauration aveugle d'image	234
12.12	Restauration aveugle d'image	235
12.13	Reconstruction d'image en tomographie X	239

Annexe G

Classement bibliographique

Problèmes inverses, régularisation :

[TA77, Tit85b, BDMP88, BPT88, Cul71, KV90], [Dem89, DL91, Dem91]

Inverse généralisée :

[Lan51, Bia59, NW74, Nas81, SF87]

Approche bayésienne :

[BT72, KW93], [Cul79, GK92, GHW79, TBKT91, Tit85a, HT87], [FDG93, GMG92]

Modèles Markoviens :

[Din88, Din90, GY95, HL92, JWHC84, You88], [Nik95, NMD96b]

Préservation de contour :

[GL93, LP88, Ter86, Ter88], [Nik95, NMD96b]

Calcul matriciel :

[Cia82, Cia88, GVL89]

Déconvolution :

[Dem85, DR85]

Déconvolution, algorithmes rapides :

[Com84, DR91]

Déconvolution aveugle :

[Hog74, BGR80, Gou92, Hol93, Sch93, LC95]

Déconvolution impulsionnelle :

[CGM85, Lav93, CID93, Cha93, Gou92] [IG90, Idi91, IG93b, IG93a]

Tomographie d'impédance :

[HWWT91, KM90, Web90]

Tomographie ultrasonore :

[Gre83, HH84, Kak79, ?, Lef85, Mos86, SLW84]

Radio astronomie :

[Bra56, Luc74, Hog74, Fie82, PR84, Hol93, Lan89, Lan91] [TMS84]

Cristallographie :

[Nav85, Nav86, Nav90]

Imagerie Radar :

[AA82]

Choix de l'*a priori* :

[O'S94, O'S95, Zel77, MDD88a, MDD88c]

Sismique :

[Sai90, AH78]

POCS :

[Bre67a, Bre67b, SS82, Com93]

Maximum d'entropie :

[SW48, Jay68, Jay82, Jay85, Bal82, JS84], [BS80, BBF⁺83], [MDD86, MDD87a, MDD87b, MDD88b, MDD88c, MD94], [Ber95, LB93]

Bibliographie

- [AA82] M.F. Adams and A.P. Anderson. Synthetic aperture tomographic (sat) imaging for microwave diagnostics. *Proceedings of the IEE*, 129:83–88, 1982.
- [AH78] V.K. Arya and H.D. Holden. Deconvolution of seismic data-an overview. *IEEE Transactions on Geoscience and Remote Sensing*, GE-16:95–98, 1978.
- [Bal82] R. Balian. *Du microscopique au macroscopique, cours de physique statistique de l'École Polytechnique*, volume 1. Ellipses, Paris, 1982.
- [BBF⁺83] R.K. Bryan, M. Bansal, W. Folkhard, C. Nave, and D.A. Marvin. Maximum-entropy calculation of the electron density at 4 a resolution of pfl filamentous bacteriophage. *Proc. Natl. Acad. Sci. USA*, 80:4728–4731, 1983.
- [BDMP88] M. Bertero, C. De Mol, and E.R. Pike. Linear inverse problems with discrete data: II. stability and regularization. *Inverse Problems*, 4:3, 1988.
- [Ber95] Jean-François Bercher. *Développement de critères de nature entropique pour la résolution des problèmes inverses linéaires*. PhD thesis, Université de Paris-Sud, Orsay, février 1995.
- [Bes74] Julian E. Besag. Spatial interaction and the statistical analysis of lattice systems (with discussion). *Journal of the Royal Statistical Society B*, 36(2):192–236, 1974.
- [Bes86] Julian E. Besag. On the statistical analysis of dirty pictures (with discussion). *Journal of the Royal Statistical Society B*, 48(3):259–302, 1986.
- [Bes89] Julian E. Besag. Digital image processing: Towards Bayesian image analysis. *Journal of Applied Statistics*, 16(3):395–407, 1989.
- [BG93] Julian E. Besag and Peter Green. Spatial statistics and Bayesian computation. *J. R. Statist. Soc. B*, 55:25–37, 1993.
- [BGR80] A. Benveniste, M. Goursat, and G. Ruget. Robust identification of nonminimum phase system: Blind adjustment of a linear equalizer in data communications. *IEEE Transactions on Automatic and Control*, AC-25, 1980.
- [Bia59] H. Bialy. Iterative behandlung linearen funktionalgleichungen. *Arch. Ration. Mech. Anal.*, 4:166–176, 1959.
- [BPT88] M. Bertero, T.A. Poggio, and V. Torre. Ill-posed problems in early vision. *Proceedings of the IEEE*, 76:869–889, August 1988.
- [Bra56] R.N. Bracewell. Strip integration in radio astronomy. *Aust. J. Phys.*, 9:198–201, 1956.
- [Bre67a] L. M. Bregman. The relaxation method of finding the common point of convex sets and its application to the solution of problems in convex programming. *Zh. vychisl. Mat. mat. Fiz.*, 7:147–156, 1967.
- [Bre67b] L. M. Bregman. The relaxation method of finding the common point of convex sets and its application to the solution of problems in convex programming.

- U.S.S.R. Computational Mathematics and Mathematical Physics*, 7:200–217, 1967.
- [BS80] R.K. Bryan and J. Skilling. Deconvolution by maximum entropy as illustrated by application to the jet of m87. *Monthly Notices of the Royal Astronomical Society*, 191:69–79, 1980.
- [BT72] G.E.P Box and G.C. Tiao. *Bayesian inference in statistical analysis*. Addison-Wesley publishing, 1972.
- [BZ87] A. Blake and A. Zisserman. *Visual reconstruction*. The MIT Press, Cambridge, 1987.
- [CGM85] C.Y. Chi, J. Goustias, and Jerry M. Mendel. A fast maximum-likelihood estimation and detection algorithm for Bernoulli-Gaussian processes. In *Proceedings of IEEE ICASSP*, pages 1297–1300, 1985.
- [Cha93] Frédéric Champagnat. *Déconvolution impulsionnelle et extensions pour la caractérisation des milieux inhomogènes en échographie*. PhD thesis, Université de Paris-Sud, Orsay, décembre 1993.
- [Cia82] Philippe G. Ciarlet. *Introduction à l'analyse numérique et à l'optimisation*. Masson, Paris, 1982.
- [Cia88] Philippe G. Ciarlet. *Introduction à l'analyse numérique matricielle et à l'optimisation*. Collection mathématiques appliquées pour la maîtrise. Masson, Paris, 1988. JFG.
- [CID93] Frédéric Champagnat, Jérôme Idier, and Guy Demoment. Deconvolution of sparse spike trains accounting for wavelet phase shifts and colored noise. In *Proceedings of IEEE ICASSP*, volume III, pages 452–455, Minneapolis, U.S.A., 1993.
- [Com84] D. Commenges. The deconvolution problem: Fast algorithms including the preconditioned conjugate-gradient to compute a MAP estimator. *IEEE Transactions on Automatic and Control*, AC-29:229–243, 1984.
- [Com93] P. L. Combettes. The foundation of set theoretic estimation. *Proceedings of the IEEE*, 81(2):182–208, February 1993.
- [Cul71] J. Cullum. Numerical differentiation and regularization. *SIAM Journal of Numerical Analysis*, 8:254–265, 1971.
- [Cul79] J. Cullum. The effective choice of the smoothing norm in regularization. *Math. Comp.*, 33:149–170, 1979.
- [Dem85] Guy Demoment. Déconvolution des signaux. Cours de l'ESE 3086, GPI-LSS, 1985.
- [Dem89] Guy Demoment. Image reconstruction and restoration: Overview of common estimation structure and problems. *IEEE Transactions on Acoustics Speech and Signal Processing*, ASSP-37(12):2024–2036, December 1989.
- [Dem91] Guy Demoment. Problèmes inverses. *Flux*, (141):42–43, novembre 1991.
- [Din88] Jean-Marc Dinten. Tomographic reconstruction with a limited number of projections: regularization using a Markov model. Technical report, Université de Paris-Sud, 1988.
- [Din90] Jean-Marc Dinten. *Tomographie à partir d'un nombre limité de projections: régularisation par champs markoviens*. PhD thesis, Université de Paris-Sud, janvier 1990.

- [DL91] Guy Demoment and André Lannes. Les problèmes inverses. *Le Courrier du CNRS*, 77:–, juin 1991. numéro spécial sur le traitement du signal et de l'image.
- [DR85] Guy Demoment and Roger Reynaud. Fast minimum-variance deconvolution. *IEEE Transactions on Acoustics Speech and Signal Processing*, ASSP-33:1324–1326, 1985.
- [DR91] Guy Demoment and Roger Reynaud. Fast RLS algorithms and chandrasekhar equations. In S. Haykin, editor, *Adaptive Signal Processing*, pages 357–367, San Diego, July 1991. SPIE Conf.
- [FDG93] Natalie Fortier, Guy Demoment, and Yves Goussard. GCV and ML methods of determining parameters in image restoration by regularization: Fast computation in the spatial domain and experimental comparison. *Journal of Visual Communication and Image Representation*, 4(2):157–170, June 1993.
- [Fie82] J.R. Fienup. Phase retrieval algorithms: a comparison. *Applied Optics*, 21(15):2758–2769, August 1982.
- [Gem90] Donald Geman. *École d'Été de Probabilités de Saint-Flour XVIII - 1988*, volume 1427, chapter Random fields and inverse problems in imaging, pages 117–193. Springer-Verlag, lecture notes in mathematics edition, 1990.
- [GG84] Stuart Geman and Donald Geman. Stochastic relaxation, Gibbs distributions, and the Bayesian restoration of images. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, PAMI-6(6):721–741, November 1984.
- [GHW79] G. H. Golub, M. Heath, and G. Wahba. Generalized cross-validation as a method for choosing a good ridge parameter. *Technometrics*, 21(2):215–223, mai 1979. JFG.
- [GK92] N.P. Galatsanos and A.K. Katsaggelos. Methods for choosing the regularization parameter and estimating the noise variance in image restoration and thier telation. *IEEE Transactions on Image Processing*, IP-1(3):322–336, July 1992. JFG.
- [GL93] Muhittin Gökmen and Ching-Chung Li. Edge detection and surface reconstruction using refined regularization. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, PAMI-15(5):492–499, May 1993.
- [GMG92] Élisabeth Gassiat, Fabrice Monfront, and Yves Goussard. On simultaneous signal estimation and parameter identification using a generalized likelihood approach. *IEEE Transactions on Information Theory*, IT-38:157–162, January 1992.
- [Gou92] Yves Goussard. Blind deconvolution of sparse spike trains using stochastic optimization. In *Proceedings of IEEE ICASSP*, volume IV, pages 593–596, 1992.
- [GR92] Stuart Geman and Georges Reynolds. Constrained restoration and recovery of discontinuities. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, PAMI-14(3):367–383, March 1992.
- [Gre83] J.K. Greenleaf. Computerized tomography with ultrasound. *Proceedings of the IEEE*, 71:330, 1983.
- [GVL89] G.H. Golub and C.F. Van Loan. *Matrix computations (2nd Edition)*. xxx, xxx, 1989.

- [GY95] Donald Geman and Chengda Yang. Nonlinear image recovery with half-quadratic regularization. *IEEE Transactions on Image Processing*, IP-4(7):932–946, July 1995.
- [HH84] D. Hiller and H. Hermert. System analysis of ultrasound reflection mode computerized tomography. *IEEE Transactions on Sonics and Ultrasonics*, SU-31:240–250, 1984.
- [HL92] Thomas J. Hebert and Richard Leahy. Statistic-based MAP image reconstruction from poisson data using Gibbs prior. *IEEE Transactions on Signal Processing*, SP-40(9):2290–2303, September 1992.
- [Hog74] Hogbom. Aperture synthesis with a non-regular distribution of interferometer baselines. *Astron. Astrophys. Suppl.*, 15:417–426, 1974.
- [Hol93] Holmes. Blind deconvolution of quantum-limited incoherent imagery: Maximum-likelihood approach. *Journal of the Optical Society of America*, 9(7), July 1993.
- [HT87] P. Hall and D. M. Titterton. Common structure of techniques for choosing smoothing parameter in regression problems. *Journal of the Royal Statistical Society B*, 49(2):184–198, 1987. JFG.
- [HWWT91] P. Hua, E.J. Woo, J.G. Webster, and W.J. Tompkins. Iterative reconstruction methods using regularization and optimal current patterns in electrical impedance tomography. *IEEE Transactions on Medical Imaging*, MI-10:621–628, December 1991.
- [Idi91] Jérôme Idier. *Imagerie des milieux stratifiés : modélisation markovienne, application à la déconvolution sismique*. PhD thesis, Université de Paris-Sud, Orsay, septembre 1991.
- [IG90] Jérôme Idier and Yves Guossard. Stack algorithm for recursive deconvolution of Bernoulli-Gaussian processes. *IEEE Transactions on Geoscience and Remote Sensing*, GE-28(5):975–978, September 1990.
- [IG93a] Jérôme Idier and Yves Guossard. Multichannel seismic deconvolution. *IEEE Transactions on Geoscience and Remote Sensing*, GE-31(5):961–979, September 1993.
- [IG93b] Jérôme Idier and Yves Guossard. Markov modeling for Bayesian restoration of two-dimensional layered structures. *IEEE Transactions on Information Theory*, IT-39(4):1356–1373, July 1993.
- [Jay68] E. T. Jaynes. Prior probabilities. *IEEE Transactions on Systems Science and Cybernetics*, SSC-4(3):227–241, September 1968.
- [Jay82] E. T. Jaynes. On the rationale of maximum-entropy methods. *Proceedings of the IEEE*, 70(9):939–952, September 1982.
- [Jay85] E. T. Jaynes. Where do we go from here? In Jr. C.R. Smith & T. Grandy, editor, *Maximum-Entropy and Bayesian Methods in Inverse Problems*, pages 21–58. 1985.
- [JS84] R. Johnson and J. Shore. Which is better entropy expression for speech processing: -slogs or logs? *IEEE Transactions on Acoustics Speech and Signal Processing*, ASSP-32:129–137, 1984.
- [JWHC84] Valen Johnson, Wing Wong, Xiaoping Hu, and Chin-Tu Chen. Image restoration using Gibbs priors: Boundary modeling, treatment of blurring, and selection of hyperparameter. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, PAMI-13(5):413–425, 1984.

- [Kak79] Avinash C. Kak. Computerized tomography with X-ray, emission, and ultrasound sources. *Proceedings of the IEEE*, 67(9):1245–1272, September 1979.
- [KM90] R. V. Kohn and A. McKenney. Numerical implementation of a variational method for electrical impedance tomography. *Inverse Problems*, pages 389–414, 1990.
- [KV90] Nicolaos Karayiannis and Anastasios Venetsanopoulos. Regularization theory in image restoration – the stabilizing functional approach. *IEEE Transactions on Acoustics Speech and Signal Processing*, ASSP-38(7):1155–1179, July 1990.
- [KW93] Daniel Keren and Michael Werman. Probabilistic analysis of regularization. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, PAMI-15(10):982–995, 1993.
- [Lan51] L. Landweber. An iteration formula for fredholm integral equations of the first kind. *Amer. J. Math.*, 73:615–624, 1951.
- [Lan89] Lannes. Backprojection mechanisms in phase-closure imaging. bispectral analysis of the phase-restoration process. *Experimental Astronomy*, 1:47–76, 1989.
- [Lan91] A. Lannes. Phase and amplitude calibration in aperture synthesis. algebraic structures. *Inverse Problems*, 7:261–298, 1991.
- [Lav93] M. Lavielle. Bayesian deconvolution of Bernoulli-Gaussian processes. *Signal Processing*, 33:67–79, 1993.
- [LB93] Guy Le Besnerais. *Méthode du maximum d'entropie sur la moyenne, critères de reconstruction d'image et synthèse d'ouverture en radio-astronomie*. PhD thesis, Université de Paris-Sud, Orsay, décembre 1993.
- [LC95] J.S. Liu and R. Chen. Blind deconvolution via sequential imputations. *Biometrika*, 90(430):567–576, 1995. MF.
- [Lef85] J.P. Lefebvre. La tomographie d'impédance acoustique. *Traitement du Signal*, 2:103–110, 1985.
- [LP88] David Lee and Theo Pavlidis. One-dimensional regularization with discontinuities. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, PAMI-10(6):417–438, November 1988.
- [Luc74] L. B. Lucy. An iterative technique for the rectification of observed distributions. *The Astronomical Journal*, 79(6):745–754, 1974.
- [MD94] Ali Mohammad-Djafari. Maximum d'entropie et problèmes inverses en imagerie. *Traitement du Signal*, pages 87–116, 1994.
- [MDD86] Ali Mohammad-Djafari and Guy Demoment. Maximum entropy diffraction tomography. In *Proceedings of IEEE ICASSP*, pages 1–34, Tokyo, Japan, 1986.
- [MDD87a] Ali Mohammad-Djafari and Guy Demoment. Maximum entropy Fourier synthesis with application to diffraction tomography. *Applied Optics*, 26(10):1745–1754, 1987.
- [MDD87b] Ali Mohammad-Djafari and Guy Demoment. Tomographie de diffraction et synthèse de Fourier à maximum d'entropie. *Revue de Physique Appliquée*, 22:153–167, 1987.
- [MDD88a] Ali Mohammad-Djafari and Guy Demoment. *Image restoration and reconstruction using entropy as a regularization functional*, volume 2, pages 341–355. Kluwer Academic Publishers, Dordrecht, The Netherlands, 1988.

- [MDD88b] Ali Mohammad-Djafari and Guy Demoment. Maximum entropy reconstruction in X ray and diffraction tomography. *IEEE Transactions on Medical Imaging*, MI-7(4):345–354, 1988.
- [MDD88c] Ali Mohammad-Djafari and Guy Demoment. Utilisation de l'entropie dans les problèmes de restauration et de reconstruction d'images. *Traitement du Signal*, 5(4):235–248, 1988.
- [Mos86] M. Moshfeghi. Ultrasound reflection-mode tomography using fan-shaped-beam insonification. *IEEE Transactions on Ultrasonics Ferroelectrics and Frequency Control*, UFFC-33:299–314, 1986.
- [Nas81] M.Z. Nashed. Operator-theoretic and computational approaches to ill-posed problems with applications to antenna theory. *IEEE Transactions on Antennas and Propagation*, 29:220–231, 1981.
- [Nav85] Jorge Navaza. On the maximum-entropy estimate of the electron density function. *Acta Crystallographica*, A-41:232–244, 1985.
- [Nav86] Jorge Navaza. The use of non-local constraints in maximum-entropy electron density reconstruction. *Acta Crystallographica*, A-42:212–223, 1986.
- [Nav90] Jorge Navaza. Accurate solutions of the maximum entropy equations. their impact on the foundations of direct methods. In *Int. School Comp.*, Bishenberg, Germany, 1990.
- [Nik95] Mila Nikolova. *Inversion markovienne de problèmes linéaires mal posés. application à l'imagerie tomographique*. PhD thesis, Université de Paris-Sud, Orsay, février 1995.
- [NMD96a] Mila Nikolova and Ali Mohammad-Djafari. Eddy current tomography using a binary Markov model. *Signal Processing*, 49:119–132, 1996. Mila Signal Processing.
- [NMD96b] Mila Nikolova and Ali Mohammad-Djafari. Eddy current tomography using a binary Markov model. *Signal Processing*, 49(5):119–132, May 1996.
- [NW74] M.Z. Nashed and G. Wahba. Generalized inverses in reproducing kernel spaces: An approach to regularization of linear operators equations. *SIAM Journal of Mathematical Analysis*, 5:974–987, 1974.
- [O'S94] Joseph A. O'Sullivan. Divergence penalty for image regularization. In *Proceedings of IEEE ICASSP*, volume V, pages 541–544, Adelaide, Australia, April 1994.
- [O'S95] Joseph A. O'Sullivan. Roughness penalties on finite domains. *IEEE Transactions on Image Processing*, IP-4(9):1258–1268, September 1995.
- [PR84] T. J. Pearson and A. C. S. Readhead. Image formation by self-calibration in radio astronomy. *Ann. Rev. Astron. Astrophys.*, 22:97–130, 1984.
- [Sai90] N. Saito. Superresolution of noisy band-limited data by data adaptive regularization and its application to seismic trace inversion. In *Proceedings of IEEE ICASSP*, pages 1237–1240, Albuquerque, U.S.A., April 1990.
- [Sch93] Schultz. Multiframe blind deconvolution of astronomical images. *Journal of the Optical Society of America*, 10(5), May 1993.
- [SF87] Didier Saint-Félix. *Restauration d'image : régularisation d'un problème mal posé et algorithmique associée*. PhD thesis, Thèse de doctorat d'État, Université de Paris-Sud, septembre 1987.
- [SLW84] C.F. Schueler, H. Lee, and G. Wade. Fundamentals of digital ultrasonic imaging. *IEEE Transactions on Sonics and Ultrasonics*, SU-31:195–217, 1984.

- [SS82] M. I. Sezan and H. Stark. Image restoration by the method of convex projections: Part 2 - applications and numerical results. *IEEE Transactions on Medical Imaging*, MI-1(2):95–101, October 1982.
- [SW48] C.E. Shannon and W. Weaver. The mathematical theory of communication. *Bell System Technical Journal*, 27:379–423, 623–656, 1948.
- [TA77] A. Tikhonov and V. Arsenin. *Solutions of Ill-Posed Problems*. Winston, Washington DC, 1977.
- [TBKT91] A. Thompson, J. C. Brown, J. W. Kay, and D. M. Titterington. A study of methods of choosing the smoothing parameter in image restoration by regularization. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, PAMI-13(4):326–339, April 1991.
- [Ter86] D. Terzopoulos. Regularization of inverse visual problems involving discontinuities. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, PAMI-8(4):413–424, July 1986.
- [Ter88] D. Terzopoulos. The computation of visual surface representations. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, PAMI-10(4):417–438, 1988.
- [Tit85a] D. M. Titterington. Common structure of smoothing techniques in statistics. *International Statistical Review*, 53(2):141–170, 1985. JFG,Mila.
- [Tit85b] D. M. Titterington. General structure of regularization procedures in image reconstruction. *A*, 144:381–387, 1985.
- [TMS84] Thompson, Moran, and Swenson. *Interferometry and Synthesis in Radio-astronomy*. Wiley Interscience, New-York, 1984.
- [Web90] J.G. Webster. *Electrical Impedance Tomography*. Adam Hilger, IOP Publishing Ltd, 1990.
- [You88] Laurent Younès. Estimation and annealing for Gibbsian fields. *Annales de l'institut Henri Poincaré*, 24(2):269–294, février 1988.
- [Zel77] Zellner. Maximal data information prior distributions. In A. Aykac and C. Brumat, editors, *New Developments in the Applications of Bayesian Methods*, pages 211–232. North-Holland, Amsterdam, 1977.