



Algorithmes d'optimisation en grande dimension : applications à la résolution de problèmes inverses

Audrey Repetti

► To cite this version:

Audrey Repetti. Algorithmes d'optimisation en grande dimension : applications à la résolution de problèmes inverses. Traitement du signal et de l'image [eess.SP]. Université Paris-Est Marne la Vallée, 2015. Français. NNT : . tel-01281096v1

HAL Id: tel-01281096

<https://theses.hal.science/tel-01281096v1>

Submitted on 1 Mar 2016 (v1), last revised 6 Jun 2016 (v3)

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

UNIVERSITÉ PARIS-EST

Laboratoire d'Informatique Gaspard Monge

UMR CNRS 8049

Thèse de Doctorat

Présentée le 29 juin 2015

par

Audrey REPETTI

ALGORITHMES D'OPTIMISATION EN GRANDE DIMENSION :
APPLICATIONS À LA RÉOLUTION DE PROBLÈMES INVERSES

Rapporteurs :	Jérôme BOLTE	<i>Université Toulouse 1 Capitole</i>
	Said MOUSSAOUI	<i>École Centrale Nantes</i>
Examineurs :	Patrick Louis COMBETTES	<i>Université Pierre et Marie Curie</i>
	Stéphane MALLAT	<i>École Normale Supérieure</i>
	Gabriele STEIDL	<i>Université Technologique de Kaiserslautern</i>
Directeur :	Jean-Christophe PESQUET	<i>Université Paris-Est Marne-la-Vallée</i>
Co-encadrante :	Émilie CHOUZENOUX	<i>Université Paris-Est Marne-la-Vallée</i>

Remerciements

Je souhaite en premier lieu remercier les personnes qui m'ont encadrée, M. Jean-Christophe Pesquet et Mme Emilie Chouzenoux, qui m'ont fait confiance et qui m'ont permis de mener à bien cette thèse. Je leur suis particulièrement reconnaissante pour le temps qu'ils m'ont accordé, pour leurs précieux conseils, ainsi que pour leur dynamisme.

Mes remerciements vont aussi aux membres du jury, qui ont accepté d'évaluer mon travail de thèse. Merci aux rapporteurs MM. Jérôme Bolte et Saïd Moussaoui pour leur relecture attentive de mon manuscrit, à M. Patrick Louis Combettes et Mme Gabriele Steidl d'avoir accepté de faire partie du jury de cette thèse, ainsi qu'à M. Stéphane Mallat d'en avoir accepté la présidence.

Je remercie également Mme Gabriele Steidl de m'avoir accueillie dans son équipe pendant un mois, et M. Laurent Duval de m'avoir permis de travailler avec lui.

Je tiens de plus à remercier Patrice pour m'avoir apporté son aide à plusieurs reprises, autant pour mes problèmes informatiques que lorsque ma voiture restait enfermée, Corinne, Laurence et Séverine pour leur sympathie et leur aide dans les démarches administratives, et Sylvie pour sa bonne humeur et sa patience.

Je souhaite enfin adresser mes remerciements aux doctorants arrivés avant/en même temps/après moi, et qui ont aussi contribué, par leur bonne humeur et leur soutien, à cette thèse. Je remercie Aurélie, Gia-Thuy, Mohamed-Ali, Sonja, et plus particulièrement Ania, Ferial, Mai, Mireille et Yosra, sans qui l'ambiance de la thèse n'aurait pas été aussi bonne.

Finalement, mes remerciements vont à mes amis et ma famille, qui me soutiennent depuis toujours et qui sont certainement ma plus grande source de motivation.

Je remercie mon cousin Marc de m'avoir transmis son intérêt pour les maths, ainsi que Anne-Catherine, Roxanne, Eloise, et Victor.

Je remercie mon cousin Vincent, Valérie, Thomas et Lucie d'avoir été d'une grande patience lors des soirées raclette/travail, et qui ont toujours été à mes côtés pour m'apporter du bonheur. Merci aussi à mon parrain Patrick pour son soutien malgré la distance qui nous sépare.

Je remercie mon grand-père Nando, Michela, Antonella et Giorgio pour leur soutien lors de mes séjours en Italie.

Je remercie Angélique, Benjamin, Benoit, Bérénice, Dorian, Giovanni, Guillaume, Isa, JS, Margot, Sébastien, Walid, qui m'apportent une grande dose de joie de vivre et de folie.

Et enfin, je remercie du fond du coeur mes parents, qui m'ont épaulée tout au long de cette thèse, et pendant toutes les années qui l'ont précédée, ainsi que Clément, qui m'a permis de m'évader au quotidien par sa bonne humeur, son empathie et sa patience.

Table des matières

1	Introduction générale	1
1.1	Contexte	1
1.2	Collaborations	2
1.3	Organisation du document	3
1.4	Publications	5
2	Généralités sur les problèmes inverses et l'optimisation	9
2.1	Introduction	10
2.2	Problèmes inverses	10
2.2.1	Introduction	10
2.2.2	Estimateur du <i>Maximum A Posteriori</i>	12
2.2.3	Fonction d'attache aux données	13
2.2.4	Fonction de régularisation	14
2.3	Outils d'analyse variationnelle	20
2.3.1	Notations et définitions	20
2.3.2	Analyse convexe	23
2.3.3	Calcul sous-différentiel	25
2.3.4	Opérateurs proximaux	28
2.3.5	Inégalité de Kurdyka-Łojasiewicz	36
2.4	Algorithmes d'optimisation	41
2.4.1	Algorithmes de gradient et de sous-gradient	41
2.4.2	Algorithme explicite-implicite	43
2.4.3	Principe de Majoration-Minimisation	46
2.4.4	Algorithmes primaux-duaux	51
2.5	Conclusion	56
Annexe 2.A	Démonstration de la propriété (v) de perturbation quadratique de l'opérateur proximal	56
3	Algorithme explicite-implicite à métrique variable	59
3.1	Introduction	60
3.2	Problème de Minimisation	61
3.3	Méthode proposée	62
3.3.1	Algorithme explicite-implicite à métrique variable	62
3.3.2	Choix de la métrique	63
3.3.3	Algorithme inexact à métrique variable	64
3.4	Analyse de convergence	65

3.4.1	Propriétés de décroissance	65
3.4.2	Résultat de convergence	68
3.5	Conclusion	70
Annexe 3.A	Démonstration du théorème 3.1	71
4	Algorithme explicite-implicite : application à un bruit gaussien dépendant	75
4.1	Introduction	76
4.2	Bruit gaussien dépendant du signal	76
4.2.1	Estimateur du <i>Maximum A Posteriori</i>	76
4.2.2	Fonction d'attache aux données	80
4.3	Application de l'algorithme explicite-implicite à métrique variable	81
4.3.1	Problème d'optimisation	81
4.3.2	Construction de la majorante	82
4.3.3	Implémentation de l'étape proximale	86
4.4	Résultats de simulation	87
4.4.1	Restauration d'image	87
4.4.2	Reconstruction d'image	89
4.5	Conclusion	89
Annexe 4.A	Démonstration de la proposition 4.4	91
5	Algorithme explicite-implicite à métrique variable alterné par bloc	93
5.1	Introduction	94
5.2	Problème de minimisation	94
5.3	État de l'art	96
5.3.1	Algorithme de minimisation alternée par bloc	96
5.3.2	Algorithme proximal alterné par bloc	97
5.3.3	Algorithme explicite-implicite alterné par bloc	98
5.4	Méthode proposée	99
5.4.1	Algorithme explicite-implicite à métrique variable alterné par bloc	99
5.4.2	Hypothèses	100
5.4.3	Version inexacte de l'algorithme explicite-implicite à métrique variable alterné par bloc	102
5.5	Analyse de convergence	103
5.5.1	Propriétés de descente	103
5.5.2	Théorème de convergence	105
5.5.3	Borne supérieure sur la vitesse de convergence	108
5.6	Conclusion	109
Annexe 5.A	Démonstration du théorème 5.1	110
Annexe 5.B	Démonstration du théorème 5.2	113
6	Quelques applications de l'algorithme explicite-implicite à métrique variable alterné par bloc	117
6.1	Introduction	118
6.2	Problème de démixage spectral	118
6.2.1	Formulation du problème	118
6.2.2	Mise en œuvre de l'algorithme	120

6.2.3	Résultats de simulation	122
6.3	Reconstruction de phase	123
6.3.1	Formulation du problème	123
6.3.2	Mise en œuvre de l'algorithme	127
6.3.3	Résultats de simulation	129
6.4	Déconvolution aveugle de signaux sismiques	132
6.4.1	Formulation du problème	132
6.4.2	Mise en œuvre de l'algorithme	137
6.4.3	Résultats de simulation	139
6.5	Conclusion	141
7	Algorithmes primaux-duaux alternés et distribués pour les problèmes d'in-	
	clusion monotone et d'optimisation convexe	143
7.1	Introduction	144
7.2	Opérateurs monotones et optimisation	145
7.2.1	Notations	145
7.2.2	Algorithme explicite-implicite alterné aléatoirement par bloc	147
7.3	Algorithmes primaux-duaux alternés aléatoirement par bloc	151
7.3.1	Problème	152
7.3.2	Première sous-classe d'algorithmes	153
7.3.3	Seconde sous-classe d'algorithmes	159
7.3.4	Application aux problèmes variationnels	163
7.4	Application de l'algorithme primal-dual alterné aléatoirement par bloc	168
7.4.1	Description du problème	168
7.4.2	Méthodes proposées	170
7.4.3	Résultats numériques	172
7.5	Algorithmes distribués	178
7.5.1	Problème et notations	178
7.5.2	Algorithme distribués asynchrones pour les problèmes d'inclusion mo-	
	notone	181
7.5.3	Application aux problèmes variationnels	187
7.6	Conclusion	192
Annexe 7.A	Démonstration du lemme 7.1	194
Annexe 7.B	Démonstration de la proposition 7.3	195
Annexe 7.C	Démonstration de la proposition 7.5	197
Annexe 7.D	Démonstration de la proposition 7.7	198
Annexe 7.E	Démonstration de la proposition 7.11	199
Annexe 7.F	Démonstration de la proposition 7.12	202
8	Opérateur proximal de fonctions quotient	203
8.1	Introduction	204
8.2	Fonctions quotient	205
8.3	Opérateur proximal de fonctions quotient	206
8.3.1	Opérateur proximal de la somme des fonctions quotient	206
8.3.2	Opérateur proximal du maximum des fonctions quotient	210
8.4	Application à l'estimation de la sélectivité	214

8.4.1	Estimation de la sélectivité en SGBD	214
8.4.2	Formulation du problème	214
8.4.3	Solution du problème en SGBD	216
8.5	Conclusion	219
9	Conclusions et perspectives	221
9.1	Bilan	221
9.2	Perspectives	224
	Bibliographie	229

Table des figures

2.1	(a) Image originale <i>peppers</i> $\bar{\mathbf{x}}$. (b) Image dégradée $\mathbf{z}$ par un flou. (c) Image estimé $\hat{\mathbf{x}}$. (d) Maillage 3D original <i>bunny</i> $\bar{\mathbf{x}}$. (e) Maillage bruité $\mathbf{z}$. (f) Maillage estimé $\hat{\mathbf{x}}$	11
2.2	Signal sismique parcimonieux.	15
2.3	Fonctions ℓ_α pour $\alpha \in \{0, 1/10, 1/5, 1/3, 1/2, 1\}$	16
2.4	(a) Image originale <i>clock</i> de dimension $N_1 = N_2 = 256$, (b) gradients verticaux, et (c) gradients horizontaux.	18
2.5	Détail de l'image <i>clock</i> . (a) Voisins locaux (carrés bleus) verticaux et horizontaux d'un pixel donné (carré rouge). (b) Voisins non locaux (carrés bleus) du même pixel (carré rouge). Dans le deuxième cas, les voisins sont sélectionnés de façon à ce que l'intensité des pixels bleus soit similaire à l'intensité du pixel rouge.	19
2.6	(a) Image originale <i>clock</i> . Décomposition sur une base d'ondelettes sur 1 niveau de résolution (b) et sur 2 niveaux de résolution (c) avec une ondelette de Haar.	20
2.7	Illustration du domaine et de l'ensemble des lignes de niveau $\delta > 0$ d'une fonction $\psi: \mathbb{R} \rightarrow \mathbb{R}$	22
2.8	Illustration de l'épigraphe d'une fonction $\mathbf{h}: \mathbb{R}^N \rightarrow \mathbb{R}$	23
2.9	Cônes normaux de l'épigraphe d'une fonction $\mathbf{f}: \mathbb{R} \rightarrow]-\infty, +\infty]$	26
2.10	Opérateurs de seuillage doux (courbe rouge discontinue) et de seuillage dur (courbe noire continue).	33
2.11	(a) Fonction ψ_α pour $\alpha \in \{1, 2, 3, 4\}$. (b) Opérateur proximal $\text{prox}_{\gamma\psi_\alpha}$ pour $\gamma = 1$ et $\alpha \in \{1, 2, 3, 4\}$	35
2.12	Fonction de Huber généralisée pour $\alpha \in \{1/10, 1/5, 1/2, 1\}$, avec (a) $\delta = 1/5$ et (b) $\delta = 1/20$	36
2.13	Fonctions $\rho_{1,\alpha}$ (haut), $\phi_{1,\alpha}$ (centre), et $\text{prox}_{\phi_{1,\alpha}}$ (bas), pour $\alpha \in \{1/10, 1/5, 1/2, 1\}$	37
2.14	Illustration de l'algorithme de Majoration-Minimisation pour la minimisation d'une fonction $\mathbf{h}: \mathbb{R} \rightarrow]-\infty, +\infty]$. À une itération donnée $k \in \mathbb{N}$, on construit une majorante $\mathbf{q}(\cdot, \mathbf{x}_k)$ de $\mathbf{h}$ au point $\mathbf{x}_k$, puis on définit l'itérée $\mathbf{x}_{k+1}$ comme étant le minimiseur de cette majorante.	47
4.1	(a) Image originale <i>phantom</i> de dimension 256×256 . (b) Bruit gaussien dépendant de l'image originale avec $\alpha^{(m)} \equiv 0.1$ et $\beta^{(m)} \equiv 0.001$. (c) Image bruitée avec le bruit dépendant gaussien. (d) Bruit blanc gaussien de variance $\sigma = 0.1154$. (e) Image bruitée avec le bruit blanc gaussien. Les images bruitées (c) et (e) ont le même RSB de 6.6 dB.	79

4.2	Restauration : (a) Image originale <i>Peppers</i> . (b) Image dégradée par un bruit gaussien dépendant du signal, RSB=21.85 dB. (c) Image reconstruite par la méthode proposée, RSB=27.11 dB.	88
4.3	Restauration : Comparaison de la vitesse de convergence de l'algorithme 3.1 explicite-implicite à métrique variable avec $\gamma_k \equiv 1$ (lignes fines continues) et $\gamma_k \equiv 1.9$ (lignes épaisses continues), de l'algorithme 2.4 explicite-implicite avec $\gamma_k \equiv 1$ (lignes fines discontinues) et $\gamma_k \equiv 1.9$ (lignes épaisses discontinues), et de FISTA (lignes pointillés).	88
4.4	Reconstruction : (a) Image originale <i>Zubal</i> . (b) Image dégradée par un bruit gaussien dépendant du signal. (c) Image reconstruite par rétro-projection filtrée, RSB=7 dB. (d) Image reconstruite par la méthode proposée, RSB=18.9 dB.	90
4.5	Reconstruction : Comparaison de la vitesse de convergence de l'algorithme 3.1 explicite-implicite à métrique variable avec $\gamma_k \equiv 1$ (lignes fines continues) et $\gamma_k \equiv 1.9$ (lignes épaisses continues), de l'algorithme 2.4 explicite-implicite avec $\gamma_k \equiv 1$ (lignes fines discontinues) et $\gamma_k \equiv 1.9$ (lignes épaisses discontinues), et de FISTA (lignes pointillés).	91
6.1	Illustration du problème de démixage spectral. L'objectif est d'estimer $\bar{\mathbf{U}}$ et $\bar{\mathbf{V}}$ à partir d'un ensemble de données hyperspectrales $\mathbf{Y}$	119
6.2	Transposée de la matrice $\bar{\mathbf{T}}$. Chaque ligne $p \in \{1, \dots, P\}$ contient les poids de la combinaison pondérée formant le p -ème spectre de $\bar{\mathbf{U}}$. Les carrés noirs correspondent aux coefficients nuls.	122
6.3	Cartes d'abondance $\bar{\mathbf{V}}$. Image correspondant à la première ligne (a), deuxième ligne (b), troisième ligne (c), quatrième ligne (d) et cinquième ligne (e) de la matrice $\bar{\mathbf{V}}$	123
6.4	Spectres exacts $\bar{\mathbf{U}} = \Omega \bar{\mathbf{T}}$ (lignes continues noires) et reconstruits $\hat{\mathbf{U}} = \Omega \hat{\mathbf{T}}$ (lignes discontinues rouges).	124
6.5	Comparaison de la vitesse de convergence de l'algorithme 6.1 (ligne discontinue rouge) et de l'algorithme PALM (ligne continue noire).	124
6.6	Parties réelle (a) et imaginaire (b) de l'image originale $\bar{\mathbf{y}}$	130
6.7	Seuls les coefficients de trame correspondants à la partie réelle de l'image $\bar{\mathbf{y}}_{\mathcal{R}}$ sont représentés dans les trois figures. (a) Exemple de décomposition de trame de $\bar{\mathbf{y}}_{\mathcal{R}}$. (b) Masque $\mathbb{E}$ permettant de forcer les coefficients de l'arrière plan de l'objet à être égaux à 0. (c) Indices d'un bloc $\mathbf{x}^{(j)}$ pour $Q = 32$	131
6.8	(a) Temps de reconstruction nécessaire pour satisfaire le critère d'arrêt $\ \mathbf{x}_k - \hat{\mathbf{x}}\ /\ \hat{\mathbf{x}}\ \leq 0.025$ avec l'algorithme 5.4 en fonction de Q . (b) Profil de convergence de l'algorithme 5.4 (ligne continue) et de l'algorithme PALM de [Bolte et al., 2013] (ligne discontinue), obtenu en utilisant Matlab 7 sur une machine Intel(R) Core(TM) i7-3520M @ 2.9GHz.	132
6.9	Parties réelle (a) (resp. (c)) et imaginaire (b) (resp. (d)) de l'image reconstruite $\hat{\mathbf{y}}$ en utilisant l'algorithme 5.4, RSB = 21.27 dB (resp. l'algorithme de projections alternées régularisé de [Mukherjee et Seelamantula, 2012], RSB = 14.45 dB).	133
6.10	(Haut) Signal sismique inconnu $\bar{\mathbf{x}}$. (Bas) Observations $\mathbf{z}$	134

6.11	(a) Fonction ℓ_0 . (b) Fonction ℓ_1 . (c) Fonction $\ell_{1/2}$. (d) Fonction $\ell_{1,\alpha}/\ell_{2,\eta}$. (e) Fonction $h_2 = \log((\ell_{1,\alpha} + \beta)/\ell_{2,\eta})$	136
6.12	Temps de reconstruction en fonction du nombre d'itérations internes $J_k \equiv J$ effectuées dans l'algorithme 6.2 (moyennes obtenues pour 30 réalisations de bruit).	140
6.13	(1) Erreur résiduelle $\bar{\mathbf{x}} - \hat{\mathbf{x}}$ pour le signal estimé $\hat{\mathbf{x}}$ obtenu en utilisant [Krishnan <i>et al.</i> , 2011] (a) et l'algorithme 6.2 (b). (2) Noyau original $\bar{\mathbf{k}}$ (ligne fine continue bleue), estimé $\hat{\mathbf{k}}$ par l'algorithme 6.2 (ligne épaisse continue noire) et par [Krishnan <i>et al.</i> , 2011] (ligne discontinue épaisse verte).	141
7.1	(a) Maillage original $\bar{\mathbf{x}}$, (b) maillage bruité $\mathbf{z}$ avec EQM = 5.45×10^{-6} , et maillage reconstruit $\hat{\mathbf{x}}$ en utilisant (c) l'algorithme 7.9 avec EQM = 8.89×10^{-7} , et (d) une méthode de lissage du Laplacien avec EQM = 1.29×10^{-6}	173
7.2	Indicateur C(p) obtenu pour différentes valeurs de la probabilité p (moyenné sur dix exécutions de l'algorithme 7.9).	174
7.3	(a) Maillage original $\bar{\mathbf{x}}$, (b) maillage bruité $\mathbf{z}$ avec EQM = 2.89×10^{-6} , et maillage reconstruit $\hat{\mathbf{x}}$ en utilisant (c) l'algorithme 7.10 avec EQM = 8.09×10^{-8} , et (d) une méthode de lissage du Laplacien avec EQM = 5.23×10^{-7}	175
7.4	Temps de reconstruction (cercles) et mémoire nécessaire (carrés) pour différents nombres de blocs p/r	176
7.5	Diagramme résumant les algorithmes primaux-duaux alternés par bloc développés dans les sections 7.3 et 7.4 pour résoudre les problèmes 7.1 et 7.2.	177
7.6	Exemple d'hypergraphe connecté avec $r = 4$ et $m = 6$. Ici, par exemple, le nœud 3 est à la fois dans les hyper-arêtes $\mathbb{V}_1 = \{i(1, 1), i(1, 2), i(1, 3)\}$ et $\mathbb{V}_2 = \{i(2, 1), i(2, 2)\}$. De plus, les voisins du nœud $3 = i(1, 3) = i(2, 1)$ sont les nœuds $1 = i(1, 1) = i(3, 1)$, $2 = i(1, 2)$ et $4 = i(2, 2) = i(3, 2) = i(4, 1)$, et on a $\mathbb{V}_3^* = \{(1, 3), (2, 1)\}$	180
7.7	Diagramme résumant les algorithmes développés dans la section 7.5.	193
8.1	Fonction $\text{prox}_{\gamma \mathbf{q}(\cdot, 1)}$ pour différents γ	209
8.2	Division de $\mathbb{R}^2$ en quatre sous-espaces pour calculer la projection épigraphique dans $\text{epi } \mathbf{q}(\cdot, \mathbf{b})$ pour $\mathbf{b} = 1$	213

Liste des tableaux

6.1	Moyennes sur 200 réalisations de bruit, pour chaque niveau de bruit, des résultats obtenus pour l'estimation de $\bar{\mathbf{x}}$ et $\bar{\mathbf{k}}$ en utilisant [Krishnan <i>et al.</i> , 2011] et l'algorithme 6.2 (Matlab 8 sur une machine Intel(R) Xeon(R) CPU E5-2609 v2@2.5GHz).	142
6.2	Variances sur 200 réalisations de bruit, pour chaque niveau de bruit, des résultats obtenus pour l'estimation de $\bar{\mathbf{x}}$ et $\bar{\mathbf{k}}$ en utilisant [Krishnan <i>et al.</i> , 2011] et l'algorithme 6.2.	142
8.1	Grammaire pour la FND	215
8.2	Erreur $\mathbf{q}_\infty$ obtenue pour l'exemple (8.47).	219

Liste des algorithmes

2.1	Algorithme du gradient	41
2.2	Algorithme de sous-gradient	42
2.3	Algorithme du point proximal	42
2.4	Algorithme explicite-implicite	44
2.5	Algorithme explicite-implicite inexact (erreurs additives)	45
2.6	Algorithme explicite-implicite inexact (erreurs relatives)	45
2.7	Algorithme de Majoration-Minimisation	46
2.8	Algorithme de Majoration-Minimisation Quadratique	48
2.9	Algorithme ADMM	52
2.10	Algorithme primal-dual [Condat, 2013; Vü, 2013]	53
2.11	Algorithme primal-dual (forme symétrique) [Condat, 2013; Vü, 2013]	54
2.12	Algorithme primal-dual simplifié [Chambolle et Pock, 2010]	54
2.13	Algorithme primal-dual simplifié (forme symétrique) [Chambolle et Pock, 2010]	55
2.14	Algorithme primal-dual simplifié (forme symétrique modifiée) [Chambolle et Pock, 2010]	55
3.1	Algorithme explicite-implicite à métrique variable	62
3.2	Algorithme inexact explicite-implicite à métrique variable	64
5.1	Algorithme de minimisation alternée par bloc	96
5.2	Algorithme proximal alterné par bloc	97
5.3	Algorithme explicite-implicite alterné par bloc	98
5.4	Algorithme explicite-implicite à métrique variable alterné par bloc	99
5.5	Algorithme inexact explicite-implicite à métrique variable alterné par bloc	102
6.1	Algorithme explicite-implicite à métrique variable alterné par bloc appliqué au problème de NMF	121
6.2	Algorithme explicite-implicite à métrique variable alterné par bloc appliqué au problème de déconvolution aveugle.	139
7.1	[Combettes et Pesquet, 2015, Algo. (4.1)]	148
7.2	Algorithme préconditionné explicite-implicite alterné aléatoirement par bloc	149
7.3	Algorithme primal-dual alterné aléatoirement par bloc (version 1)	156
7.4	Algorithme primal-dual alterné aléatoirement par bloc (version 1 symétrique)	159
7.5	Algorithme primal-dual alterné aléatoirement par bloc (version 2)	162
7.6	Algorithme primal-dual (version 1) appliqué dans un cadre variationnel	165
7.7	Algorithme primal-dual (version 1 symétrique) appliqué dans un cadre variationnel	166
7.8	Algorithme primal-dual (version 2) appliqué dans un cadre variationnel	167
7.9	Algorithme primal-dual (version 1 symétrique simplifiée) appliqué dans un cadre variationnel	170
7.10	Algorithme primal-dual (version 2 simplifiée) appliqué dans un cadre variationnel	171

7.11	Algorithme distribué asynchrone (version 1)	182
7.12	Algorithme distribué asynchrone (version 1 simplifiée)	184
7.13	Algorithme distribué asynchrone (version 1 simplifiée dans le cas où $(\forall i \in \{1, \dots, m\})$ $D_i^{-1} = 0$ et $B_i = 0$)	185
7.14	Algorithme distribué asynchrone (version 2 simplifiée)	186
7.15	Algorithme distribué asynchrone (version 1 simplifiée) appliqué dans un cadre va- riationnel	188
7.16	Algorithme distribué (version 1 simplifiée) sur un graphe non orienté	189
7.17	Algorithme distribué (version 1 simplifiée) sur un graphe non orienté lorsque $(\forall i \in$ $\{1, \dots, m\}) \mathbf{g}_i = 0$ et $\mathbf{l}_i = \iota_{\{0\}}$	190
8.1	ADMM pour le calcul de $\text{prox}_{\gamma\theta_\infty}(\mathbf{x})$, où $\mathbf{x} \in \mathbb{R}^N$	210
8.2	Algorithme primal-dual pour résoudre le problème (8.44)	217
8.3	Algorithme primal-dual pour résoudre le problème (8.45)	218

Abréviations

BCD	:	Block Coordinate Descent
BC-FB	:	Block Coordinate Forward-Backward
BC-VMFB	:	Block Coordinate Variable Metric Forward-Backward
FB	:	Forward-Backward
KL	:	Kurdyka-Łojasiewicz
MAP	:	Maximum <i>a posteriori</i>
MM	:	Majoration-Minimisation
PALM	:	Proximal Alternating Linearized Minimization
SDP	:	Semi Définie Positive
VMFB	:	Variable Metric Forward-Backward

Résumé

Une approche efficace pour la résolution de problèmes inverses consiste à définir le signal (ou l'image) recherché(e) par minimisation d'un critère pénalisé. Ce dernier s'écrit souvent sous la forme d'une somme de fonctions composées avec des opérateurs linéaires. En pratique, ces fonctions peuvent n'être ni convexes ni différentiables. De plus, les problèmes auxquels on doit faire face sont souvent de grande dimension. L'objectif de cette thèse est de concevoir de nouvelles méthodes pour résoudre de tels problèmes de minimisation, tout en accordant une attention particulière aux coûts de calculs ainsi qu'aux résultats théoriques de convergence.

Une première idée pour construire des algorithmes rapides d'optimisation est d'employer une stratégie de préconditionnement, la métrique sous-jacente étant adaptée à chaque itération. Nous appliquons cette technique à l'algorithme explicite-implicite et proposons une méthode, fondée sur le principe de majoration-minimisation, afin de choisir automatiquement les matrices de préconditionnement. L'analyse de la convergence de cet algorithme repose sur l'inégalité de Kurdyka-Łojasiewicz.

Une seconde stratégie consiste à découper les données traitées en différents blocs de dimension réduite. Cette approche nous permet de contrôler à la fois le nombre d'opérations s'effectuant à chaque itération de l'algorithme, ainsi que les besoins en mémoire, lors de son implémentation. Nous proposons ainsi des méthodes alternées par bloc dans les contextes de l'optimisation non convexe et convexe. Dans le cadre non convexe, une version alternée par bloc de l'algorithme explicite-implicite préconditionné est proposée. Les blocs sont alors mis à jour suivant une règle déterministe acyclique. Lorsque des hypothèses supplémentaires de convexité peuvent être faites, nous obtenons divers algorithmes proximaux primaux-duaux alternés, permettant l'usage d'une règle aléatoire arbitraire de balayage des blocs. L'analyse théorique de ces algorithmes stochastiques d'optimisation convexe se base sur la théorie des opérateurs monotones.

Un élément clé permettant de résoudre des problèmes d'optimisation de grande dimension réside dans la possibilité de mettre en œuvre en parallèle certaines étapes de calculs. Cette parallélisation est possible pour les algorithmes proximaux primaux-duaux alternés par bloc que nous proposons : les variables primales, ainsi que celles duales, peuvent être mises à jour en parallèle, de manière tout à fait flexible. À partir de ces résultats, nous déduisons de nouvelles méthodes distribuées, où les calculs sont répartis sur différents agents communiquant entre eux suivant une topologie d'hypergraphe.

Finalement, nos contributions méthodologiques sont validées sur différentes applications en traitement du signal et des images. Nous nous intéressons dans un premier temps à divers problèmes d'optimisation faisant intervenir des critères non convexes, en particulier en restauration d'images lorsque l'image originale est dégradée par un bruit gaussien dépendant du signal, en mélange spectral, en reconstruction de phase en tomographie, et en déconvolution aveugle pour la reconstruction de signaux sismiques parcimonieux. Puis, dans un second temps, nous abordons des problèmes convexes intervenant dans la reconstruction de maillages 3D et dans l'optimisation de requêtes pour la gestion de bases de données.

Abstract

An efficient approach for solving an inverse problem is to define the recovered signal/image as a minimizer of a penalized criterion which is often split in a sum of simpler functions composed with linear operators. In the situations of practical interest, these functions may be neither convex nor smooth. In addition, large scale optimization problems often have to be faced. This thesis is devoted to the design of new methods to solve such difficult minimization problems, while paying attention to computational issues and theoretical convergence properties.

A first idea to build fast minimization algorithms is to make use of a preconditioning strategy by adapting, at each iteration, the underlying metric. We incorporate this technique in the forward-backward algorithm and provide an automatic method for choosing the preconditioning matrices, based on a majorization-minimization principle. The convergence proofs rely on the Kurdyka-Łojasiewicz inequality.

A second strategy consists of splitting the involved data in different blocks of reduced dimension. This approach allows us to control the number of operations performed at each iteration of the algorithms, as well as the required memory. For this purpose, block alternating methods are developed in the context of both non-convex and convex optimization problems. In the non-convex case, a block alternating version of the preconditioned forward-backward algorithm is proposed, where the blocks are updated according to an acyclic deterministic rule. When additional convexity assumptions can be made, various alternating proximal primal-dual algorithms are obtained by using an arbitrary random sweeping rule. The theoretical analysis of these stochastic convex optimization algorithms is grounded on the theory of monotone operators.

A key ingredient in the solution of high dimensional optimization problems lies in the possibility of performing some of the computation steps in a parallel manner. This parallelization is made possible in the proposed block alternating primal-dual methods where the primal variables, as well as the dual ones, can be updated in a quite flexible way. As an offspring of these results, new distributed algorithms are derived, where the computations are spread over a set of agents connected through a general hypergraph topology.

Finally, our methodological contributions are validated on a number of applications in signal and image processing. First, we focus on optimization problems involving non-convex criteria, in particular image restoration when the original image is corrupted with a signal dependent Gaussian noise, spectral unmixing, phase reconstruction in tomography, and blind deconvolution in seismic sparse signal reconstruction. Then, we address convex minimization problems arising in the context of 3D mesh denoising and in query optimization for database management.

CHAPITRE 1

Introduction générale

1.1 CONTEXTE

Au cours de cette dernière décennie, de nombreuses applications, parmi lesquelles l'imagerie médicale et satellitaire, le traitement de maillages 3D de grande tailles, ou la gestion de bases de données, ont conduit à traiter des données dont la dimension est de plus en plus grande. Il devient alors nécessaire de développer des méthodes permettant un traitement rapide de ces grandes masses de données.

La plupart de ces applications font intervenir un problème inverse visant à estimer des données inconnues à partir d'une version dégradée de ces dernières. Une méthode efficace consiste à définir l'estimée comme étant une solution d'un problème d'optimisation, où l'objectif est de minimiser un critère composé de fonctions qui ne sont ni nécessairement différentiables, ni nécessairement convexes. Dans cette thèse, nous élaborerons plusieurs stratégies qui vont nous permettre d'obtenir des méthodes de minimisation rapides, tout en réduisant les coûts de calculs et les besoins en mémoire.

La première stratégie, permettant d'accélérer la convergence d'algorithmes d'optimisation, consiste à utiliser une technique de préconditionnement se basant sur une modification de la métrique sous-jacente au cours des itérations. Un aspect intéressant de notre étude concerne le choix des matrices de préconditionnement. Nous proposerons une méthode, basée sur une stratégie de majoration-minimisation (MM), permettant de les choisir efficacement. Cette méthode fera l'objet des chapitres 3 et 5. Nos résultats théoriques seront validés par des exemples numériques dans les chapitres 4 et 6.

La seconde stratégie, permettant d'optimiser les besoins en mémoire lors de l'implémentation d'algorithmes d'optimisation est de diviser les données traitées en différents blocs de dimension plus petite. Ainsi, chaque bloc peut être mis à jour séparément. Dans certains domaines du traitement de données de grande taille, par exemple en astronomie, l'image est partitionnée en plusieurs blocs, qui sont traités indépendamment les uns des autres. L'inconvénient de cette approche est qu'elle introduit des artefacts au niveau des frontières des blocs. Afin d'éviter ces effets indésirables, on peut introduire une étape empirique de lissage, et/ou considérer des blocs pouvant se chevaucher. Dans cette thèse, des méthodes rigoureuses opérant par blocs seront proposées dans le cadre de l'optimisation non convexe dans le chapitre 5, et convexe pour les méthodes primales-duales élaborées dans le chapitre 7. Nous montrerons, de plus, dans le chapitre 6 que

les stratégies par bloc peuvent également permettre, en plus d'une réduction de l'occupation mémoire, d'accélérer la vitesse de convergence des algorithmes.

Enfin, une dernière stratégie permettant de résoudre efficacement les problèmes d'optimisation de grande taille consiste à implémenter en parallèle certaines tâches des algorithmes. En raison des progrès technologiques actuels, la conception d'algorithmes parallèles permettant de bénéficier de systèmes multicœurs est souvent d'une importance primordiale pour obtenir des mises en œuvre rapides. Les algorithmes proximaux, basés sur des méthodes primales-duales développés ces dernières années en sont d'excellents exemples. Dans ces méthodes, les calculs des opérateurs proximaux et les étapes de gradient peuvent être effectués en parallèle. Cependant, lorsque le nombre de processeurs disponibles est inférieur au nombre d'opérations à effectuer en parallèle, la plupart des algorithmes proximaux parallèles existants deviennent alors d'un intérêt limité. Dans le chapitre 7 nous proposerons une méthode par bloc permettant de surmonter cette limitation en réduisant de manière flexible le nombre d'opérations effectuées à chaque itération. D'autre part, dans le chapitre 8 nous traiterons un exemple applicatif en bases de données à l'aide d'algorithmes parallèles.

1.2 COLLABORATIONS

Certains travaux présentés dans cette thèse ont été réalisés dans le cadre de collaborations :

- Une première collaboration avec Gabriele Steidl de l'Université de Kaiserslautern (Allemagne) dans le cadre du Labex Bezout a eu pour objectif de développer des méthodes basées sur les opérateurs proximaux et les projections épigraphiques pour traiter des fonctions quotient dans le cadre de l'optimisation de requêtes dans les systèmes de gestion de bases de données. Cette collaboration m'a amenée à passer un mois dans le laboratoire de mathématiques de l'Université de Kaiserslautern. Les résultats obtenus sont présentés dans le chapitre 8.
- Une seconde collaboration avec Laurent Duval de L'IFP Énergies Nouvelles a permis de développer une méthode pour traiter des problèmes de déconvolution aveugle de signaux sismiques parcimonieux. Cette collaboration m'a amené à travailler avec Mai Quyen Pham, doctorante à l'université de Paris-Est et à l'IFP Énergies Nouvelles. Les résultats que nous avons obtenus sont présentés dans le chapitre 6.
- Enfin une dernière collaboration s'est déroulée dans le cadre du projet ANR GRAPHSSIP (Graph Signal Processing) démarré fin 2014, au cours de laquelle nous avons développé des algorithmes primaux-duaux permettant de résoudre des problèmes de débruitage de maillages 3D. Les résultats obtenus sont présentés dans le chapitre 7.

1.3 ORGANISATION DU DOCUMENT

Chapitre 2 : Généralités sur les problèmes inverses et l'optimisation

Le deuxième chapitre sera divisé en trois grandes parties. Tout d'abord, nous expliquerons brièvement comment, partant d'un problème inverse, on peut être conduit à un problème d'optimisation, en utilisant une interprétation de type Maximum a Posteriori (MAP). Dans la seconde partie, nous donnerons les principales notations et outils mathématiques qui seront utilisés tout au long de cette thèse. En particulier, nous nous intéresserons aux opérateurs proximaux de fonctions non nécessairement convexes et à l'inégalité de Kurdyka-Łojasiewicz. Pour finir, dans la dernière partie de ce chapitre, un état de l'art des algorithmes usuels existants permettant de résoudre des problèmes de minimisation sera présenté. En particulier, nous y évoquerons des méthodes telles que les algorithmes de gradient, les algorithmes proximaux, les algorithmes de Majoration-Minimisation, et les méthodes primales-duales.

Chapitre 3 : Algorithme explicite-implicite à métrique variable

Une méthode usuelle permettant de minimiser une somme de deux fonctions, l'une étant différentiable et l'autre convexe, est l'algorithme explicite-implicite, qui alterne, à chaque itération, une étape de gradient sur la fonction différentiable et une étape proximale sur la fonction convexe. Un inconvénient majeur de cet algorithme est sa vitesse de convergence, qui peut se révéler très lente en pratique. Dans ce chapitre nous proposerons donc d'en accélérer la convergence, en changeant, à chaque itération, la métrique sous-jacente, par l'introduction d'une métrique variable. Cette métrique sera choisie en se basant sur une stratégie de Majoration-Minimisation. Nous fournirons une analyse de la convergence de cet algorithme basée sur l'inégalité de Kurdyka-Łojasiewicz. De plus, une version inexacte de l'algorithme sera proposée, permettant de prendre en compte des erreurs de calcul pouvant apparaître au cours des itérations.

Chapitre 4 : Application de l'algorithme explicite-implicite à métrique variable à un bruit gaussien dépendant du signal

Dans ce chapitre nous allons nous intéresser au problème de restauration d'un signal dégradé par un bruit gaussien dépendant. Un tel bruit peut s'observer dans divers systèmes d'imagerie, dans lesquels les acquisitions sont dégradées à la fois par un bruit gaussien dépendant lié au comptage des photons, et par un bruit électronique indépendant du signal d'origine. Nous montrerons que l'interprétation MAP de ce modèle de bruit conduit à une fonction d'attache aux données non convexe. Nous proposerons donc de résoudre le problème de minimisation associé à ce modèle grâce à l'algorithme explicite-implicite à métrique variable, développé au chapitre précédent. L'intérêt d'utiliser une métrique variable sera illustrée sur un exemple pratique en restauration d'images.

Chapitre 5 : Algorithme explicite-implicite à métrique variable alterné par bloc

Dans de nombreuses applications, telles que la résolution de problèmes de déconvolution aveugle, ou lorsque les données considérées sont de très grande dimension, il est intéressant de considérer des méthodes de minimisation alternée. Dans ce chapitre, nous proposerons donc une version alternée par bloc de l'algorithme explicite-implicite à métrique variable développé dans le troisième chapitre. Nous étudierons la convergence et la vitesse de convergence de l'algorithme proposé, dans des versions exacte et inexacte, en considérant une règle de mise à jour acyclique des blocs.

Chapitre 6 : Quelques applications de l'algorithme explicite-implicite à métrique variable alterné par bloc

Afin d'illustrer les bonnes performances pratiques de l'algorithme explicite-implicite à métrique variable alterné par bloc, nous donnerons dans ce chapitre trois exemples de simulation. Dans le premier exemple, nous considérerons un problème de démélange spectral, une application en reconstruction de phase pour un problème de tomographie, et un exemple de déconvolution aveugle pour la restauration de signaux sismiques parcimonieux. Ces applications nous permettront d'illustrer l'intérêt pratique de l'algorithme proposé au chapitre précédent. En particulier, nous montrerons que l'utilisation d'une métrique variable et d'une règle acyclique de mise à jour des blocs conduit à des méthodes de résolution rapides pour les trois problèmes considérés.

Chapitre 7 : Algorithmes primaux-duaux alternés et distribués pour les opérateurs monotones

Dans les chapitre précédents nous avons présenté des méthodes d'optimisation efficaces pour la résolution de problèmes non convexes. Cependant, lorsque le critère fait intervenir des fonctions convexes non lisses composées avec des opérateurs linéaires, ces méthodes requièrent des sous-itérations qui peuvent s'avérer coûteuses.

Dans ce chapitre, nous nous plaçons dans le cadre de l'optimisation convexe, et nous nous intéressons à la classe des algorithmes primaux-duaux. Ces méthodes permettent de traiter des opérateurs linéaires apparaissant dans le problème, sans passer par leur inversion ni par des sous-itérations. Nous nous intéresserons plus généralement aux algorithmes primaux-duaux permettant de résoudre des problèmes d'inclusion monotone. Nous proposerons deux classes de nouveaux algorithmes stochastiques primaux-duaux, permettant de ne mettre à jour, à chaque itération, qu'une sous-partie des données globales sélectionnées de façon aléatoire. Nous analyserons la convergence presque sûre des variables aléatoires générées par les algorithmes proposés. Un exemple applicatif de débruitage de maillage 3D sera considéré. De plus, dans une seconde partie du chapitre, nous déduirons des méthodes précédentes des algorithmes stochastiques distribués, dans lesquels les étapes de calculs sont réparties sur différents agents communiquant entre eux suivant un schéma d'hypergraphe.

Chapitre 8 : Opérateur proximal de fonctions quotient

Dans ce chapitre nous nous intéresserons à l'application des algorithmes proximaux primaux-duaux à l'estimation de requêtes dans les systèmes de gestion de bases de données. Les problèmes sous-jacents dans ce type d'application requièrent la manipulation de fonctions $\mathbf{q}_1$ et $\mathbf{q}_\infty$, respectivement définies comme étant les normes ℓ_1 et ℓ_∞ de fonctions quotients. Nous montrerons d'une part que l'opérateur proximal de la fonction $\mathbf{q}_1$ correspond à un opérateur de seuillage. D'autre part, nous formulerons les problèmes faisant intervenir la fonction $\mathbf{q}_\infty$ comme des problèmes avec contraintes épigraphiques. Pour cela, nous étudierons l'opérateur de projection sur l'épigraphe de $\mathbf{q}_\infty$. Pour finir, nous illustrerons les performances des méthodes proposées par un exemple numérique.

1.4 PUBLICATIONS

*Pour les articles mentionnés par *, les noms des auteurs suivent l'ordre alphabétique, suivant les usages dans les journaux de mathématiques appliquées.*

Articles parus ou acceptés pour publication dans des journaux internationaux :

- ▶ E. Chouzenoux, J.-C. Pesquet, et A. Repetti. *
Variable metric forward-backward algorithm for minimizing the sum of a differentiable function and a convex function,
Journal of Optimization Theory and Applications, vol. 162, no. 1, pp 107–132, Juillet 2014.
- ▶ A. Repetti, M. Q. Pham, L. Duval, E. Chouzenoux et J.-C. Pesquet.
Euclid in a taxicab : sparse blind deconvolution with smoothed ℓ_1/ℓ_2 regularization,
Signal Processing Letters, vol. 22, no. 5, pp. 539–543, Mai 2015.
- ▶ G. Moerkotte, M. Montag, A. Repetti et G. Steidl. *
Proximal operator of quotient functions with application to a feasibility problem in query optimization,
Journal of Computational and Applied Mathematics, vol. 285, pp 243–255, Sept. 2015.
- ▶ J.-C. Pesquet et A. Repetti. *
A class of randomized primal-dual algorithms for distributed optimization,
 Accepté pour publication dans *Journal of Nonlinear and Convex Analysis*, 2014.

Article soumis dans un journal international :

- ▶ E. Chouzenoux, J.-C. Pesquet, et A. Repetti. *
A block coordinate variable metric forward-backward algorithm,
 Soumis à *Global Optimization*, 2013.
http://www.optimization-online.org/DB_HTML/2013/12/4178.html.

Articles de conférences acceptés :

- ▶ A. Repetti, E. Chouzenoux et J.-C. Pesquet.
A penalized weighted least squares approach for restoring data corrupted with signal-dependent noise.
Dans *Proceedings of the 20th European Signal Processing Conference (EUSIPCO 2012)*, pages 1553–1557, Bucarest, Roumanie, 27-31 août 2012.
- ▶ A. Repetti, E. Chouzenoux et J.-C. Pesquet.
Reconstruction d'image en présence de bruit gaussien dépendant par un algorithme explicite-implicite à métrique variable.
Dans *Actes du 24e colloque GRETSI*, Brest, France, 3-6 septembre 2013.
- ▶ A. Repetti, E. Chouzenoux et J.-C. Pesquet.
A nonconvex regularized approach for phase retrieval.
Dans *Proceedings of IEEE International Conference on Image Processing (ICIP 2014)*, pp. 1753–1757, Paris, France, 27-30 octobre 2014.
- ▶ A. Repetti, E. Chouzenoux et J.-C. Pesquet.
A random block-coordinate primal-dual proximal algorithm with application to 3D mesh denoising.
Accepté à *IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2015)*, Brisbane, Australie, 21-25 avril 2015.

Articles de conférences invités :

- ▶ A. Repetti, E. Chouzenoux et J.-C. Pesquet.
A preconditioned forward-backward approach with application to large-scale nonconvex spectral unmixing problems.
Dans *Proceedings of IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2014)*, pp. 1512–1516, Florence, Italie, 4-9 mai 2014.
- ▶ A. Repetti, E. Chouzenoux et J.-C. Pesquet.
Proximal primal-dual optimization methods.
Dans *Proceedings of international Biomedical and Astronomical Signal Processing (BASP) Frontiers workshop*, pp. 25, Villars-sur-Ollon, Switzerland, 25-30 janvier 2015.

Articles de conférences soumis :

- ▶ A. Repetti, E. Chouzenoux et J.-C. Pesquet.
A parallel block-coordinate approach for primal-dual splitting with arbitrary random block selection.
Soumis à *European Signal Processing Conference (EUSIPCO 2015)*.
- ▶ A. Repetti, E. Chouzenoux et J.-C. Pesquet.
Un petit tutoriel sur les méthodes primales-duales pour l'optimisation convexe.
Soumis à *GRETSI 2015*.

Autres présentations :

- ▶ A. Repetti, E. Chouzenoux et J.-C. Pesquet.
Variable metric forward-backward algorithm for minimizing the sum of a differentiable function and a convex function.
Advances in Mathematical Image Processing (AIP), Annweiler, Allemagne, 30 septembre - 2 octobre 2013.
- ▶ A. Repetti, E. Chouzenoux et J.-C. Pesquet.
Une nouvelle approche régularisée pour la reconstruction de phase.
Journées Imagerie Optique Non Conventionnelle (JIONC), Paris, France, 19-20 mai 2014.
- ▶ A. Repetti, E. Chouzenoux et J.-C. Pesquet.
Proximal methods : tools for solving inverse problems on a large scale.
BioImage Informatics and Modeling at the Cellular Scale, Toulouse, France, 30 juin - 2 juillet 2014.
- ▶ A. Repetti, P. L. Combettes et J.-C. Pesquet.
Randomized primal-dual methods based on stochastic quasi-Fejér monotonicity property.
European Conference on Stochastic Programming and Energy Applications (ECSP 2014), Paris, IHP, France, 24 - 26 septembre 2014.
- ▶ A. Repetti, E. Chouzenoux, M. Q. Pham, L. Duval et J.-C. Pesquet.
Algorithme préconditionné explicite-implicite et application à la résolution de problèmes inverses en traitement du signal.
GdR ISIS - Optimisation non-convexe, Paris, France, 16 octobre 2014.
- ▶ A. Repetti, E. Chouzenoux et J.-C. Pesquet.
A random block-coordinate proximal primal-dual optimization method with application to 3D mesh denoising.
EPFL - Laboratoire de systèmes d'information et d'inférence, Lausanne, Suisse, 27 mars 2015.

CHAPITRE 2

Généralités sur les problèmes inverses et l'optimisation

Sommaire

2.1	Introduction	10
2.2	Problèmes inverses	10
2.2.1	Introduction	10
2.2.2	Estimateur du <i>Maximum A Posteriori</i>	12
2.2.3	Fonction d'attache aux données	13
2.2.4	Fonction de régularisation	14
2.3	Outils d'analyse variationnelle	20
2.3.1	Notations et définitions	20
2.3.2	Analyse convexe	23
2.3.3	Calcul sous-différentiel	25
2.3.4	Opérateurs proximaux	28
2.3.5	Inégalité de Kurdyka-Łojasiewicz	36
2.4	Algorithmes d'optimisation	41
2.4.1	Algorithmes de gradient et de sous-gradient	41
2.4.2	Algorithme explicite-implicite	43
2.4.3	Principe de Majoration-Minimisation	46
2.4.4	Algorithmes primaux-duaux	51
2.5	Conclusion	56
Annexe 2.A	Démonstration de la propriété (v) de perturbation quadra- tique de l'opérateur proximal	56

2.1 INTRODUCTION

Ce chapitre a pour but, dans un premier temps, d'introduire les problèmes inverses en traitement du signal et des images, puis, dans un second temps, de définir les outils fondamentaux d'analyse et d'optimisation qui seront utilisés tout au long de cette thèse pour résoudre ces problèmes.

Les problèmes inverses et leurs liens avec les problèmes variationnels seront présentés dans la section 2.2. Puis dans la section 2.3 nous introduirons les outils fondamentaux d'analyse et d'optimisation. Enfin, dans la section 2.4 nous donnerons des exemples d'algorithmes utilisant ces outils pour résoudre des problèmes variationnels.

2.2 PROBLÈMES INVERSES

2.2.1 Introduction

De nombreuses applications en traitement du signal et des images se ramènent à la résolution de problèmes inverses. Cette résolution consiste en l'estimation d'un signal original inconnu, à partir d'une version dégradée (ou incomplète) de celui-ci. Mathématiquement, le signal original peut être représenté par un vecteur $\bar{\mathbf{x}} \in \mathbb{R}^N$, et l'observation dégradée que l'on a de ce signal par un vecteur $\mathbf{z} \in \mathbb{R}^M$ lié à $\bar{\mathbf{x}}$ par la relation

$$\mathbf{z} = \mathbf{H}\bar{\mathbf{x}} + \mathbf{w}, \quad (2.1)$$

où $\mathbf{H} \in \mathbb{R}^{M \times N}$ est une matrice modélisant une dégradation (en traitement des images, cet opérateur peut, par exemple, représenter un flou apparaissant lors de la capture de l'image), et $\mathbf{w} \in \mathbb{R}^M$ est une réalisation d'une variable aléatoire $\mathbf{w}$, représentant un bruit d'acquisition et/ou une erreur de modèle. L'objectif est alors de produire une estimée $\hat{\mathbf{x}} \in \mathbb{R}^N$ du signal original $\bar{\mathbf{x}}$ à partir de $\mathbf{z}$ et $\mathbf{H}$. La figure 2.1 donne deux exemples de problèmes inverses. L'exemple du haut correspond à un problème de restauration d'image, et celui du bas correspond à un problème de restauration de maillage 3D.

Une méthode usuelle pour trouver $\hat{\mathbf{x}}$ est de le définir comme étant un minimiseur d'une somme de deux fonctions [Demoment, 1989] :

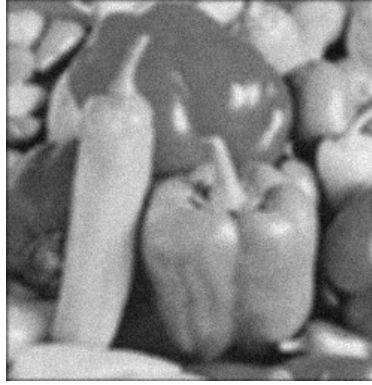
$$\hat{\mathbf{x}} \in \underset{\mathbf{x} \in \mathbb{R}^N}{\text{Argmin}} \ h(\mathbf{x}) + g(\mathbf{x}), \quad (2.2)$$

où $h: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ est le terme d'attache aux données qui contient les informations directement liées au problème traité (par exemple, un estimateur statistique lié à la loi de $\mathbf{w}$), et $g: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ est le terme de régularisation (ou de pénalisation) permettant d'introduire des informations que l'on connaît *a priori* sur le signal recherché (par exemple, parcimonie de la solution, contraintes variées).

Comme nous le verrons dans le chapitre 4, dans certains cas, la matrice $\mathbf{H}$ est elle aussi inconnue. Le problème inverse considéré est alors appelé *problème aveugle*, et l'objectif est de



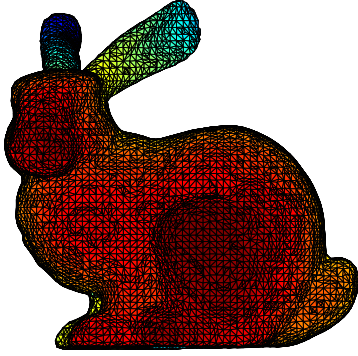
(a)



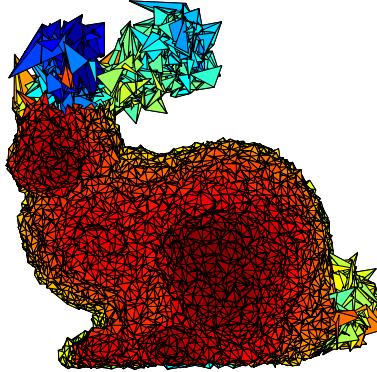
(b)



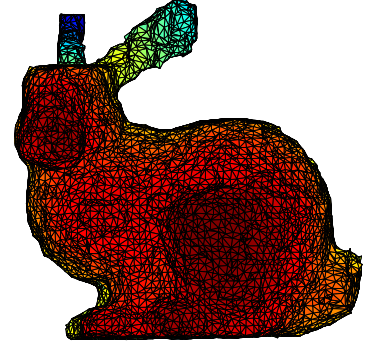
(c)



(d)



(e)



(f)

Figure 2.1 – (a) *Image originale* peppers $\bar{\mathbf{x}}$. (b) *Image dégradée* $\mathbf{z}$ par un flou. (c) *Image estimée* $\hat{\mathbf{x}}$. (d) *Maillage 3D original* bunny $\bar{\mathbf{x}}$. (e) *Maillage bruité* $\mathbf{z}$. (f) *Maillage estimé* $\hat{\mathbf{x}}$.

reconstruire à la fois $\bar{\mathbf{x}}$ et $\mathbf{H}$ à partir de $\mathbf{z}$. Pour cela, on définit les estimées $(\hat{\mathbf{x}}, \hat{\mathbf{H}}) \in \mathbb{R}^N \times \mathbb{R}^{M \times N}$ de $(\bar{\mathbf{x}}, \mathbf{H})$ comme des solutions du problème variationnel :

$$(\hat{\mathbf{x}}, \hat{\mathbf{H}}) \in \underset{(\mathbf{x}, \mathbf{H}) \in \mathbb{R}^N \times \mathbb{R}^{M \times N}}{\text{Argmin}} \quad h(\mathbf{x}, \mathbf{H}) + g_1(\mathbf{x}) + g_2(\mathbf{H}), \quad (2.3)$$

où $h: \mathbb{R}^N \times \mathbb{R}^{M \times N} \rightarrow]-\infty, +\infty]$ est le terme de fidélité, et $g_1: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ (resp. $g_2: \mathbb{R}^{M \times N} \rightarrow]-\infty, +\infty]$) est le terme de régularisation associé au signal $\mathbf{x}$ (resp. à la matrice $\mathbf{H}$). Notons cependant que dans la suite de cette section nous ne nous intéresserons qu'au cas où $\mathbf{H}$ est connue.

Nous allons tout d'abord donner dans la section 2.2.2 une interprétation bayésienne du problème d'estimation (2.1) conduisant à un choix du terme d'attache aux données dans (2.2), et dans la section 2.2.4 nous donnerons quelques exemples de fonctions de régularisation.

Remarque 2.1. Dans ce manuscrit, nous utiliserons des caractères droits pour représenter les

variables déterministes et des caractères italiques pour les variables aléatoires. De plus, les lettres en gras représenteront des vecteurs de variables tandis que les lettres fines représenteront généralement des variables scalaires.

Par exemple, le vecteur $\mathbf{x} = (x^{(n)})_{1 \leq n \leq N} \in \mathbb{R}^N$ sera une réalisation du vecteur aléatoire $\mathbf{x} = (x^{(n)})_{1 \leq n \leq N}$.

2.2.2 Estimateur du *Maximum A Posteriori*

Dans cette section, nous allons nous intéresser au choix de la fonction d'attache aux données $\mathbf{h}$ pour résoudre le problème inverse modélisé par l'équation (2.1). Une approche pour choisir cette fonction est de se baser sur une formulation bayésienne du problème. Supposons que $\mathbf{z}$ et $\overline{\mathbf{x}}$ soient des réalisations de vecteurs aléatoires $\mathbf{z}$ et $\overline{\mathbf{x}}$ absolument continus. En utilisant une approche de type *Maximum A Posteriori* (MAP), nous pouvons définir une estimée $\hat{\mathbf{x}}$ de $\overline{\mathbf{x}}$ par

$$\hat{\mathbf{x}} \in \underset{\mathbf{x} \in \mathbb{R}^N}{\operatorname{Argmax}} f_{\overline{\mathbf{x}}|\mathbf{z}=\mathbf{z}}(\mathbf{x}), \quad (2.4)$$

où $f_{\overline{\mathbf{x}}|\mathbf{z}=\mathbf{z}}$ est la densité de probabilité *a posteriori* de $\overline{\mathbf{x}}$ sachant que $\mathbf{z} = \mathbf{z}$.

Propriété 2.1. Soit $\mathbf{z} \in \mathbb{R}^M$ défini par l'équation (2.1). Supposons que $\overline{\mathbf{x}}$ est indépendant de $\mathbf{w}$. Une solution $\hat{\mathbf{x}}$ du problème (2.4) est donnée par

$$\hat{\mathbf{x}} \in \underset{\mathbf{x} \in \mathbb{R}^N}{\operatorname{Argmin}} h(\mathbf{x}) + g(\mathbf{x}), \quad (2.5)$$

où $h: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ et $g: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ sont respectivement les termes d'attache aux données et de régularisation définis par

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \begin{cases} h(\mathbf{x}) = -\log f_{\mathbf{w}}(\mathbf{z} - \mathbf{H}\mathbf{x}) \\ g(\mathbf{x}) = -\log f_{\overline{\mathbf{x}}}(\mathbf{x}), \end{cases} \quad (2.6)$$

où $f_{\overline{\mathbf{x}}}$ est la densité de probabilité de $\overline{\mathbf{x}}$.

Démonstration. Soit $\mathbf{z} \in \mathbb{R}^M$ tel que $f_{\mathbf{z}}(\mathbf{z}) \neq 0$. D'après la règle de Bayes, on a

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad f_{\overline{\mathbf{x}}|\mathbf{z}=\mathbf{z}}(\mathbf{x}) = \frac{f_{\mathbf{z}|\overline{\mathbf{x}}=\mathbf{x}}(\mathbf{z})f_{\overline{\mathbf{x}}}(\mathbf{x})}{f_{\mathbf{z}}(\mathbf{z})}. \quad (2.7)$$

Ainsi, (2.4) se réécrit

$$\hat{\mathbf{x}} \in \underset{\mathbf{x} \in \mathbb{R}^N}{\operatorname{Argmax}} f_{\mathbf{z}|\overline{\mathbf{x}}=\mathbf{x}}(\mathbf{z})f_{\overline{\mathbf{x}}}(\mathbf{x}). \quad (2.8)$$

Par monotonie de la fonction logarithme, on obtient alors

$$\hat{\mathbf{x}} \in \underset{\mathbf{x} \in \mathbb{R}^N}{\operatorname{Argmin}} \left\{ -\log f_{\mathbf{z}|\overline{\mathbf{x}}=\mathbf{x}}(\mathbf{z}) - \log f_{\overline{\mathbf{x}}}(\mathbf{x}) \right\}. \quad (2.9)$$

En utilisant le modèle (2.1), on a $\mathbf{w} = \mathbf{z} - \mathbf{H}\bar{\mathbf{x}}$. Ainsi, pour tout $\mathbf{x} \in \mathbb{R}^N$,

$$(\forall \mathbf{z} \in \mathbb{R}^M) \quad f_{\mathbf{z}|\bar{\mathbf{x}}=\mathbf{x}}(\mathbf{z}) = f_{\mathbf{w}|\bar{\mathbf{x}}=\mathbf{x}}(\mathbf{w}), \quad (2.10)$$

où $\mathbf{w} = \mathbf{z} - \mathbf{H}\mathbf{x}$. Or, $\bar{\mathbf{x}}$ et $\mathbf{w}$ étant indépendants, on a

$$(\forall \mathbf{w} \in \mathbb{R}^M) \quad f_{\mathbf{w}|\bar{\mathbf{x}}=\mathbf{x}}(\mathbf{w}) = f_{\mathbf{w}}(\mathbf{w}). \quad (2.11)$$

Le résultat s'obtient en combinant les deux dernières inégalités et en les injectant dans (2.9). $\square$

Remarque 2.2. *L'hypothèse d'indépendance entre $\bar{\mathbf{x}}$ et $\mathbf{w}$ revient à supposer que le signal et le bruit additif sont indépendants entre eux. Cependant, en pratique cette hypothèse n'est pas toujours vérifiée. En particulier, nous verrons au chapitre 4 une application en restauration d'image où le bruit de mesure dépend du signal.*

La propriété 2.1 nous indique qu'en utilisant une interprétation de type MAP, le choix de la fonction d'attache aux données $\mathbf{h}$ dépend du type de bruit considéré, tandis que le choix de la fonction de régularisation dépend des caractéristiques des données à estimer.

2.2.3 Fonction d'attache aux données

Bruit gaussien corrélé : Soit $\mathbf{w}$ la variable aléatoire modélisant le bruit additif dans (2.1). Plaçons-nous dans le cas où $\mathbf{w}$ suit une loi normale d'espérance nulle et de matrice de covariance définie positive $\mathbf{\Gamma} \in \mathbb{R}^{M \times M}$:

$$\mathbf{w} \sim \mathcal{N}(\mathbf{0}_M, \mathbf{\Gamma}). \quad (2.12)$$

Sa densité de probabilité est alors donnée par

$$(\forall \mathbf{w} \in \mathbb{R}^M) \quad f_{\mathbf{w}}(\mathbf{w}) = \left(2^M \pi^M \det(\mathbf{\Gamma})\right)^{-1/2} \exp\left(-\frac{1}{2} \mathbf{w}^\top \mathbf{\Gamma}^{-1} \mathbf{w}\right), \quad (2.13)$$

où $\det(\mathbf{\Gamma})$ désigne le déterminant de la matrice $\mathbf{\Gamma}$, et $(\cdot)^\top$ l'opérateur de transposition. D'après la propriété 2.1, la fonction d'attache aux données $\mathbf{h}$ correspond à

$$\begin{aligned} (\forall \mathbf{x} \in \mathbb{R}^N) \quad \mathbf{h}(\mathbf{x}) &= -\log f_{\mathbf{w}}(\mathbf{z} - \mathbf{H}\mathbf{x}) \\ &= \frac{1}{2} (\mathbf{y} - \mathbf{H}\mathbf{x})^\top \mathbf{\Gamma}^{-1} (\mathbf{y} - \mathbf{H}\mathbf{x}) + \kappa, \end{aligned} \quad (2.14)$$

où $\kappa = \frac{1}{2} \log \left(2^M \pi^M \det(\mathbf{\Gamma})\right) \in \mathbb{R}$ est une constante ne dépendant pas de $\mathbf{x}$. On peut donc choisir comme terme d'attache aux données dans (2.5) la fonction

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \mathbf{h}(\mathbf{x}) = \frac{1}{2} (\mathbf{y} - \mathbf{H}\mathbf{x})^\top \mathbf{\Gamma}^{-1} (\mathbf{y} - \mathbf{H}\mathbf{x}). \quad (2.15)$$

Bruit blanc gaussien : On dit que le vecteur aléatoire $\mathbf{w}$ modélise un *bruit blanc gaussien* s'il suit une loi normale d'espérance nulle, et ses composantes sont indépendantes, identiquement distribuées, de variance égale à σ^2 . On se place donc dans un cas particulier de (2.12) où $\mathbf{\Gamma} = \sigma^2 \mathbf{I}_M$, $\mathbf{I}_M$ étant la matrice identité de $\mathbb{R}^{M \times M}$. Ainsi, la fonction d'attache aux données associée se réduit à

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \mathbf{h}(\mathbf{x}) = \frac{1}{2\sigma^2} \|\mathbf{y} - \mathbf{H}\mathbf{x}\|^2. \quad (2.16)$$

2.2.4 Fonction de régularisation

Comme nous l'avons mentionné dans la section 2.2.2, le terme de régularisation g permet de prendre en compte des caractéristiques connues des données à estimer. De manière plus générale, la fonction g peut s'écrire comme une somme de plusieurs fonctions de régularisation :

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad g(\mathbf{x}) = \sum_{j=1}^J \lambda_j g_j(\mathbf{x}), \quad (2.17)$$

où, pour tout $j \in \{1, \dots, J\}$, $\lambda_j \in [0, +\infty]$ est appelé *paramètre de régularisation* et $g_j: \mathbb{R}^N \rightarrow]-\infty, +\infty]$. L'utilisation de plusieurs fonctions de régularisation permet d'espérer de meilleurs résultats de reconstruction, puisque chacune d'elles permet de contraindre l'estimée à vérifier une caractéristique particulière.

Les paramètres de régularisation $(\lambda_j)_{1 \leq j \leq J}$ sont à fixer de façon à obtenir la meilleure qualité de reconstruction possible. Notons que fixer $\lambda_j = 0$, pour tout $j \in \{1, \dots, J\}$, revient à trouver une solution $\hat{\mathbf{x}}$ non régularisée :

$$\text{trouver } \hat{\mathbf{x}} \in \underset{\mathbf{x} \in \mathbb{R}^N}{\text{Argmin}} \quad h(\mathbf{x}). \quad (2.18)$$

Au contraire, faire tendre λ_j vers $+\infty$, pour tout $j \in \{1, \dots, J\}$, revient généralement à trouver une solution $\hat{\mathbf{x}}$ ne prenant en compte que le terme de régularisation :

$$\text{trouver } \hat{\mathbf{x}} \in \underset{\mathbf{x} \in \mathbb{R}^N}{\text{Argmin}} \quad g(\mathbf{x}). \quad (2.19)$$

Il faut donc régler les paramètres de régularisation de manière à obtenir un compromis entre le terme d'attache aux données et le terme de régularisation. En pratique, ces paramètres peuvent être réglés de manière à minimiser l'erreur entre les données originales $\bar{\mathbf{x}}$ et leurs estimations $\hat{\mathbf{x}}$. Par exemple, on peut utiliser l'*erreur quadratique moyenne* (EQM) ou le *rapport signal sur bruit* (RSB) définis par

$$\text{EQM}(\hat{\mathbf{x}}, \bar{\mathbf{x}}) = \frac{1}{N} \|\hat{\mathbf{x}} - \bar{\mathbf{x}}\|^2 \quad \text{et} \quad \text{RSB}(\hat{\mathbf{x}}, \bar{\mathbf{x}}) = 20 \log_{10} \left(\frac{\|\bar{\mathbf{x}}\|}{\|\hat{\mathbf{x}} - \bar{\mathbf{x}}\|} \right), \quad (2.20)$$

où $\|\cdot\|$ désigne la norme euclidienne usuelle. Dans le cas où les données originales ne sont pas disponibles, on peut avoir recours à des estimateurs d'erreur [Chaux *et al.*, 2008; Deledalle *et al.*, 2013, 2014].

Dans cette section, nous allons donner quelques exemples de fonctions de régularisation souvent rencontrés dans le domaine des problèmes inverses.

2.2.4.1 Contraintes

Le problème de minimisation variationnel (2.2) permet de prendre en compte des contraintes lorsque la fonction de régularisation g correspond à la fonction indicatrice (donnée dans la définition 2.3(viii)) d'un sous-ensemble C de $\mathbb{R}^N$:

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad g(\mathbf{x}) = \lambda \iota_C(\mathbf{x}) = \iota_C(\mathbf{x}), \quad (2.21)$$

(le paramètre de régularisation $\lambda > 0$ est “absorbé” par la fonction indicatrice). Ici, C représente un ensemble de contraintes, qui peut s’écrire comme une intersection de contraintes distinctes. Par exemple, on peut considérer des contraintes de type

$$C = \{\mathbf{x} \in \mathbb{R}^N \mid \mathbf{a} \leq \mathbf{L}(\mathbf{x}) \leq \mathbf{b}\}, \quad (2.22)$$

où $\mathbf{L}: \mathbb{R}^N \rightarrow \mathbb{R}^M$ et $(\mathbf{a}, \mathbf{b}) \in [-\infty, +\infty[^M \times]-\infty, +\infty]^M$ ¹. Remarquons qu’en choisissant $\mathbf{a} = \mathbf{b}$, on peut considérer des contraintes d’égalité.

Exemple 2.1. Pour une image en niveaux de gris codée sur 8 bits, l’intensité de chaque pixel est comprise entre 0 et 255. Ainsi, si la variable $\mathbf{x} \in \mathbb{R}^N$ représente une image réorganisée sous forme de vecteur, on peut utiliser la fonction de régularisation $\mathbf{g} = \iota_{[0,255]^N}$. Cela correspond à la contrainte présentée ci-dessus où $\mathbf{L}$ est la fonction identité et, pour tout $n \in \{1, \dots, N\}$, $(\mathbf{a}^{(n)}, \mathbf{b}^{(n)}) = (0, 255)$.

En utilisant des contraintes pour régulariser l’estimée, nous n’avons plus besoin d’optimiser le paramètre de régularisation, ce qui peut être utile dans les cas où nous n’avons pas accès aux estimateurs de type (2.20). Cependant, il faut connaître une estimation des paramètres $(\mathbf{a}, \mathbf{b})$.

2.2.4.2 Parcimonie

Un vecteur $\mathbf{x} \in \mathbb{R}^N$ est dit *parcimonieux* si la plupart de ses composantes sont (approximativement) nulles.

Certains signaux sont parcimonieux dans leur domaine d’observation, par exemple les signaux sismiques sont parcimonieux dans le domaine temporel (c.f. figure 2.2). D’autres signaux

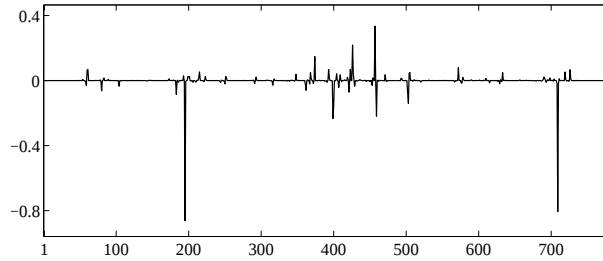


Figure 2.2 — Signal sismique parcimonieux.

sont parcimonieux après une représentation linéaire appropriée, c’est-à-dire après avoir subi une transformation par un opérateur. Dans ce cas, on considèrera une fonction de régularisation de la forme

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \mathbf{g}(\mathbf{x}) = \sum_{s=1}^S \mathbf{g}_s(\mathbf{F}_s \mathbf{x}), \quad (2.23)$$

où, pour tout $s \in \{1, \dots, S\}$, $\mathbf{F}_s \in \mathbb{R}^{M_s \times N}$ et $\mathbf{g}_s: \mathbb{R}^{M_s} \rightarrow]-\infty, +\infty]$. Ici, l’objet $\mathbf{x} \in \mathbb{R}^N$ est supposé parcimonieux après application des matrices $(\mathbf{F}_s)_{1 \leq s \leq S}$, c’est-à-dire que, pour tout $s \in \{1, \dots, S\}$, $\mathbf{F}_s \mathbf{x}$ a un nombre important de coefficients égaux à zéros.

1. La notation $\mathbf{a} \leq \mathbf{L}(\mathbf{x}) \leq \mathbf{b}$ signifie que $\mathbf{L}(\mathbf{x}) - \mathbf{a} \in [0, +\infty[^M$ et $\mathbf{b} - \mathbf{L}(\mathbf{x}) \in [0, +\infty[^M$.

Exemples de pénalisations renforçant la parcimonie : Il existe plusieurs choix de fonctions g_s permettant de promouvoir la parcimonie d'un objet estimé. Nous en donnerons ici quelques exemples. Intuitivement, la fonction permettant de favoriser au mieux la parcimonie d'un vecteur $\mathbf{x} \in \mathbb{R}^N$ est la "pseudo-norme" ℓ_0 [Donoho *et al.*, 1995] qui compte le nombre de coefficients non nuls de $\mathbf{x}$:

$$(\forall \mathbf{x} = (x^{(n)})_{1 \leq n \leq N} \in \mathbb{R}^N) \quad \ell_0(\mathbf{x}) = \|\mathbf{x}\|_0 = \sum_{n=1}^N \chi(x^{(n)}), \quad (2.24)$$

où la fonction $\chi: \mathbb{R} \rightarrow \{0, 1\}$ est définie par

$$(\forall x \in \mathbb{R}) \quad \chi(x) = \begin{cases} 0 & \text{si } x = 0, \\ 1 & \text{sinon.} \end{cases} \quad (2.25)$$

De façon plus générale, les fonctions ℓ_α [Bouman et Sauer, 1996], pour $0 \leq \alpha < 1$, sont des fonctions de régularisation permettant de promouvoir la parcimonie d'un objet :

$$(\forall \mathbf{x} = (x^{(n)})_{1 \leq n \leq N} \in \mathbb{R}^N) \quad \ell_\alpha(\mathbf{x}) = \|\mathbf{x}\|_\alpha = \left(\sum_{n=1}^N |x^{(n)}|^\alpha \right)^{1/\alpha}, \quad (2.26)$$

quand $\alpha \neq 0$. Cependant, ces fonctions ne sont ni différentiables, ni convexes. Elles peuvent donc se révéler difficiles à optimiser en pratique. Une façon usuelle de réduire la difficulté du problème est d'utiliser la norme ℓ_1 [Bect *et al.*, 2004; Donoho, 2006; Figueiredo *et al.*, 2007] comme approximation convexe des pénalisations ci-dessus :

$$(\forall \mathbf{x} = (x^{(n)})_{1 \leq n \leq N} \in \mathbb{R}^N) \quad \ell_1(\mathbf{x}) = \|\mathbf{x}\|_1 = \sum_{n=1}^N |x^{(n)}|. \quad (2.27)$$

La figure 2.3 représente différentes fonctions ℓ_α , pour $0 \leq \alpha \leq 1$. Nous pouvons voir que plus α tend vers 0, moins les coefficient proches de 0 seront pénalisés.

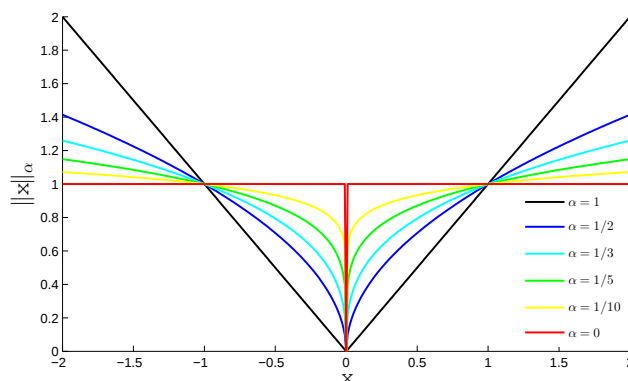


Figure 2.3 — Fonctions ℓ_α pour $\alpha \in \{0, 1/10, 1/5, 1/3, 1/2, 1\}$.

Bien que l'utilisation de la norme ℓ_1 permette de s'affranchir de la non-convexité de la régularisation, cette fonction est non différentiable. Dans certains cas pratiques, il peut être plus

avantageux d'utiliser des approximations lisses de la fonction ℓ_1 , comme par exemple la fonction de Huber [Huber, 1981] définie par $\mathbf{x} = (\mathbf{x}^{(n)})_{1 \leq n \leq N} \in \mathbb{R}^N \mapsto \mathbf{g}(\mathbf{x}) = \sum_{n=1}^N \varphi_\delta(\mathbf{x}^{(n)})$, où $\delta > 0$ et

$$(\forall \mathbf{x} \in \mathbb{R}) \quad \varphi_\delta(\mathbf{x}) = \begin{cases} \frac{\mathbf{x}^2}{2\delta} & \text{si } |\mathbf{x}| \leq \delta, \\ |\mathbf{x}| - \frac{\delta}{2} & \text{sinon.} \end{cases} \quad (2.28)$$

Enfin, des fonctions de type quotient ou différence de normes peuvent aussi être utilisées pour favoriser la parcimonie de données. Ces fonction ne sont, généralement, ni convexes ni différentiables. Dans le chapitre 6, nous définirons une approximation lisse mais non convexe de la pénalisation ℓ_1/ℓ_2 , et nous l'utiliserons dans une application de déconvolution aveugle de signaux sismiques.

Variation totale : Une des fonctions de régularisation les plus utilisées en traitement des images est la *Variation Totale* introduite dans [Rudin et al., 1992] pour des données continues. La version discrète proposée notamment dans [Chambolle, 2004] est définie par :

Définition 2.1. Soit $\mathbf{x} \in \mathbb{R}^{N_1 \times N_2}$ une matrice modélisant une image de dimension $N = N_1 \times N_2$. La variation totale (TV) de $\mathbf{x}$ est donnée par

$$\mathbf{g}(\mathbf{x}) = \text{tv}(\mathbf{x}) = \sum_{i=1}^{N_1} \sum_{j=1}^{N_2} \sqrt{([\nabla_v \mathbf{x}]^{(i,j)})^2 + ([\nabla_h \mathbf{x}]^{(i,j)})^2}, \quad (2.29)$$

où $\nabla_v \mathbf{x} \in \mathbb{R}^{N_1 \times N_2}$ et $\nabla_h \mathbf{x} \in \mathbb{R}^{N_1 \times N_2}$ sont les gradients verticaux et horizontaux de $\mathbf{x} = (\mathbf{x}^{(i,j)})_{1 \leq i \leq N_1, 1 \leq j \leq N_2}$ définis par

$$(\forall (i,j) \in \{1, \dots, N_1\} \times \{1, \dots, N_2\}) \quad \begin{aligned} [\nabla_v \mathbf{x}]^{(i,j)} &= \begin{cases} \mathbf{x}^{(i+1,j)} - \mathbf{x}^{(i,j)} & \text{si } i < N_1, \\ 0 & \text{sinon,} \end{cases} \\ [\nabla_h \mathbf{x}]^{(i,j)} &= \begin{cases} \mathbf{x}^{(i,j+1)} - \mathbf{x}^{(i,j)} & \text{si } j < N_2, \\ 0 & \text{sinon.} \end{cases} \end{aligned}$$

Remarque 2.3. Dans la définition 2.1 nous avons fait l'hypothèse que les coefficients se trouvant sur les bords droit et inférieur de l'image sont égaux à 0. Cependant, d'autres choix sont possibles pour la valeur de ces coefficients (par exemple, l'hypothèse de périodicité ou de symétrie miroir).

Un exemple de gradients verticaux et horizontaux d'une image en niveaux de gris est donné dans la figure 2.4.

Notons que l'équation (2.29) peut se réécrire de la façon suivante :

$$(\forall \mathbf{x} \in \mathbb{R}^{N_1 \times N_2}) \quad \mathbf{g}(\mathbf{x}) = \sum_{n=1}^N \mathbf{g}_n(\mathbf{F}_n \mathbf{x}), \quad (2.30)$$

où $N = N_1 \times N_2$, et, pour tout $n \in \{1, \dots, N\}$, $\mathbf{F}_n = [\nabla_v, \nabla_h]^\top$. Comme il l'est souligné dans [Chouzenoux et al., 2013], diverses fonctions $(\mathbf{g}_n)_{1 \leq n \leq N}$ peuvent être considérées pour pénaliser plus ou moins la parcimonie des gradients de l'image.

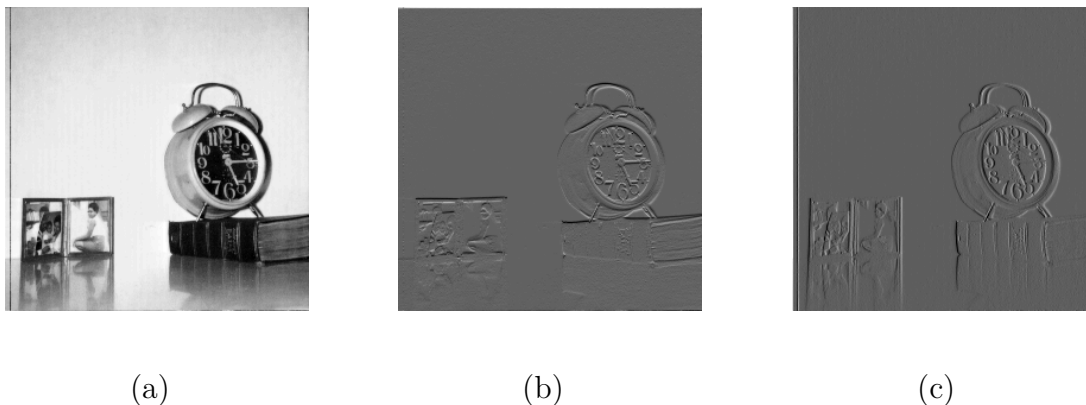


Figure 2.4 – (a) *Image originale* clock de dimension $N_1 = N_2 = 256$, (b) *gradients verticaux*, et (c) *gradients horizontaux*.

La variation totale a été généralisée à l'emploi de *gradients non locaux*. Cette méthode a été introduite dans [Gilboa et Osher, 2008], et dans ce cas, on parle de *Variation Totale Non Locale* (NLTV). La principale différence entre la TV et la NLTV réside dans la définition des voisins d'un pixel. Pour la TV, les voisins d'un pixel d'indice $(i, j) \in \{2, \dots, N_1 - 1\} \times \{2, \dots, N_2 - 1\}$ sont les pixels d'indices

$$\begin{array}{cc} (i-1, j) & (i+1, j) \\ (i, j-1) & (i, j+1). \end{array}$$

Tandis que pour la NLTV, les voisins d'un pixel sont définis selon leur degré de similarité. Une illustration est donnée dans la figure 2.5. Ainsi, on peut utiliser la fonction de régularisation (2.30), avec, pour tout $n \in \{1, \dots, N\}$, $\mathbf{F}_n = [\omega_n^1 \mathbf{F}_n^1, \dots, \omega_n^{M_n} \mathbf{F}_n^{M_n}]^\top$ où M_n détermine le nombre d'orientations considérées au n -ème pixel et les $(\omega_n^m)_{1 \leq m \leq M_n}$ sont des poids à fixer dans une phase de prétraitement. Il a été montré expérimentalement que cette régularisation permettait d'améliorer les résultats de reconstruction dans certains cas de figure [Chierchia et al., 2013a; Peyré, 2011; Werlberger et al., 2010].

Opérateurs de trames : On considère le cas où la fonction de régularisation s'écrit sous la forme (2.23), où pour tout $s \in \{1, \dots, S\}$, $\mathbf{F}_s = \mathbf{F}$ est un opérateur de trame défini par

$$\mathbf{F}: \mathbb{R}^N \rightarrow \mathbb{R}^K: \mathbf{x} \mapsto \mathbf{F}\mathbf{x} = (\langle \mathbf{x}, \mathbf{e}_k \rangle)_{1 \leq k \leq K}, \quad (2.31)$$

où $(\mathbf{e}_k)_{1 \leq k \leq K}$ est un dictionnaire de signaux de $\mathbb{R}^N$ constituant une trame [Mallat, 2009, Chap. 5], i.e., vérifiant la condition suivante :

$$(\exists(\underline{\mu}, \bar{\mu}) \in]0, +\infty[^2) \text{ tel que } (\forall \mathbf{x} \in \mathbb{R}^N) \quad \underline{\mu} \|\mathbf{x}\|^2 \leq \sum_{k=1}^K |\langle \mathbf{x}, \mathbf{e}_k \rangle|^2 \leq \bar{\mu} \|\mathbf{x}\|^2. \quad (2.32)$$

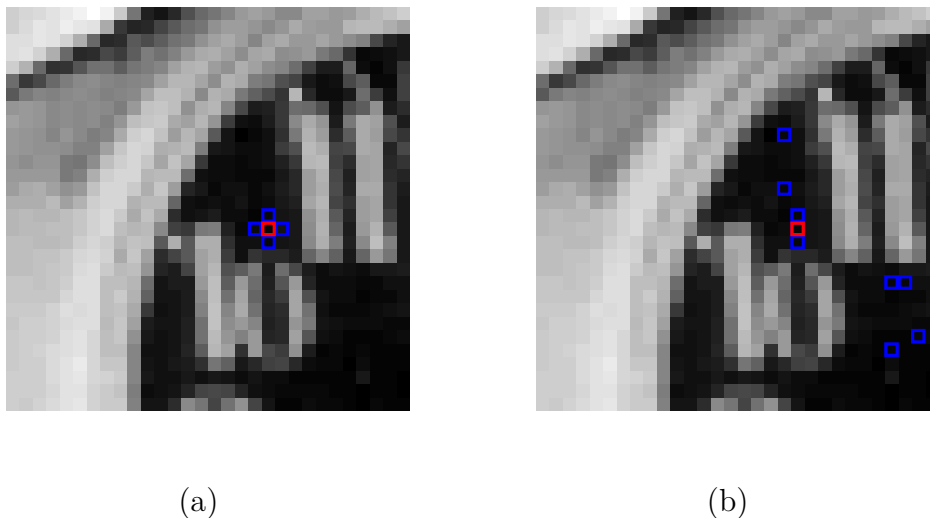


Figure 2.5 – *Détail de l'image clock. (a) Voisins locaux (carrés bleus) verticaux et horizontaux d'un pixel donné (carré rouge). (b) Voisins non locaux (carrés bleus) du même pixel (carré rouge). Dans le deuxième cas, les voisins sont sélectionnés de façon à ce que l'intensité des pixels bleus soit similaire à l'intensité du pixel rouge.*

Notons que dans $\mathbb{R}^N$, l'existence de la borne supérieure est toujours garantie, ce qui n'est pas toujours le cas en dimension infinie. Lorsque $\underline{\mu} = \overline{\mu}$, on dit que la trame est *ajustée*. L'adjoint de l'opérateur de trame est donné par

$$\mathbf{F}^*: \mathbb{R}^K \rightarrow \mathbb{R}^N: \mathbf{y} = (y^{(k)})_{1 \leq k \leq K} \mapsto \mathbf{F}^* \mathbf{y} = \sum_{k=1}^K y^{(k)} \mathbf{e}_k. \quad (2.33)$$

Soit $\mathbf{x} \in \mathbb{R}^N$ un signal donné, on notera alors $\mathbf{y} = \mathbf{F}\mathbf{x}$ ses coefficients de trames.

Des cas particuliers des opérateurs de trames sont les *opérateurs d'ondelettes*. Un exemple de transformée en ondelettes, utilisant une ondelette de Haar est donné dans la figure 2.6. Plus précisément, la figure 2.6(b) représente une décomposition sur une base d'ondelettes de l'image *clock* (figure 2.6(a)) sur un niveau de résolution. Les coefficients de cette décomposition se trouvant dans la partie supérieure gauche sont appelés *coefficients d'approximations* de la figure 2.6(a), tandis que le reste de l'image est composée de *coefficients de détails*. Ce sont les coefficients de détails qui sont les plus parcimonieux. La figure 2.6(c) représente aussi une décomposition en ondelettes de l'image mais sur deux niveaux de résolution, i.e. la décomposition est itérée une seconde fois sur les coefficients d'approximations.

Notons que l'approche que nous avons décrite ici s'appelle *approche à l'analyse*, c'est à dire que l'on cherche à estimer directement l'image via le problème variationnel :

$$\text{trouver } \hat{\mathbf{x}} \in \underset{\mathbf{x} \in \mathbb{R}^N}{\text{Argmin}} \ h(\mathbf{x}) + g(\mathbf{F}\mathbf{x}). \quad (2.34)$$

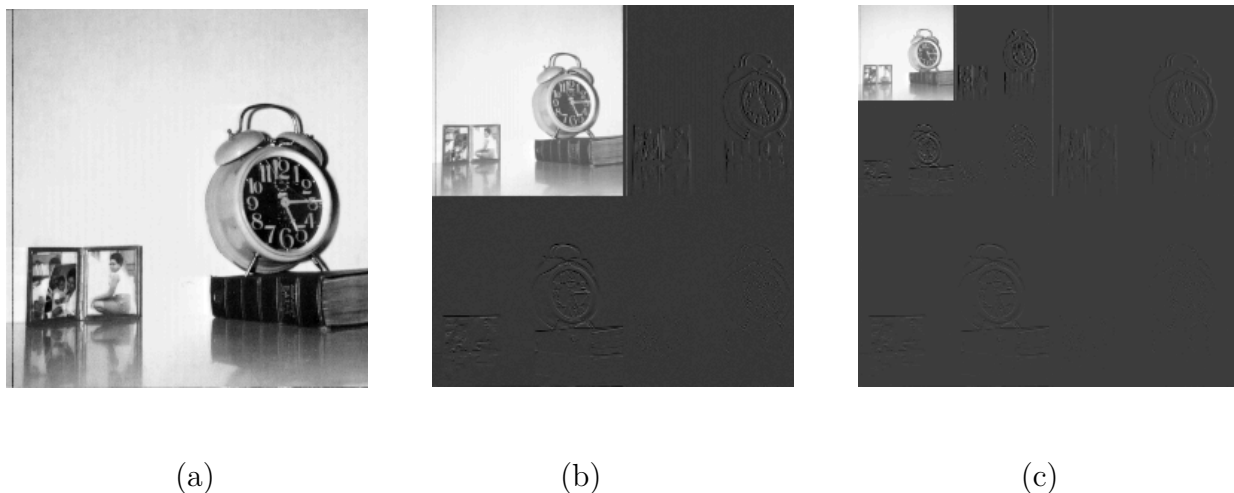


Figure 2.6 – (a) Image originale clock. Décomposition sur une base d'ondelettes sur 1 niveau de résolution (b) et sur 2 niveaux de résolution (c) avec une ondelette de Haar.

Nous verrons un exemple de cette approche dans le chapitre 4. Une autre façon d'utiliser des trames est de considérer une *approche à la synthèse*. Dans ce cas, on cherche à estimer les coefficients de trame en résolvant le problème variationnel :

$$\text{trouver } \hat{\mathbf{y}} \in \underset{\mathbf{y} \in \mathbb{R}^N}{\text{Argmin}} \quad h(\mathbf{F}^* \mathbf{y}) + g(\mathbf{y}), \quad (2.35)$$

où $\hat{\mathbf{y}} = \mathbf{F} \hat{\mathbf{x}}$ sont les coefficients de trame de l'image estimée. L'image est ensuite déduite en utilisant l'adjoint de l'opérateur de trame $\hat{\mathbf{x}} = \mathbf{F}^* \hat{\mathbf{y}}$. Un exemple de cette approche sera présenté dans le chapitre 6.

2.3 OUTILS D'ANALYSE VARIATIONNELLE

2.3.1 Notations et définitions

Nous allons, dans un premier temps, donner nos notations et rappeler des définitions d'analyse qui nous seront utiles dans la suite de ce manuscrit.

Soit $\mathcal{H}$ un espace de Hilbert réel, muni d'un produit scalaire noté $\langle \cdot, \cdot \rangle$ et d'une norme notée $\|\cdot\| = \sqrt{\langle \cdot, \cdot \rangle}$. Nous noterons $\mathcal{B}(\mathcal{H}, \mathcal{G})$ l'ensemble des opérateurs linéaires bornés de $\mathcal{H}$ dans $\mathcal{G}$, où $\mathcal{G}$ est un espace de Hilbert réel, $\mathcal{S}(\mathcal{H})$ l'ensemble des opérateurs linéaires auto-adjoints² bornés de $\mathcal{H}$ dans $\mathcal{H}$, et $\mathcal{S}^+(\mathcal{H})$ l'ensemble des opérateurs de $\mathcal{S}(\mathcal{H})$ fortement positifs³. Remarquons que si $\mathcal{H} = \mathbb{R}^N$, alors $\mathcal{S}_N := \mathcal{S}(\mathbb{R}^N)$ désigne l'ensemble des matrices symétriques de $\mathbb{R}^{N \times N}$, et

2. $\mathbf{U} \in \mathcal{B}(\mathcal{H}) = \mathcal{B}(\mathcal{H}, \mathcal{H})$ est auto-adjoint si $\mathbf{U}^* = \mathbf{U}$.

3. $\mathbf{U} \in \mathcal{S}(\mathcal{H})$ est fortement positif s'il existe $\alpha \in]0, +\infty[$ tel que $(\forall \mathbf{x} \in \mathcal{H}) \quad \langle \mathbf{x} | \mathbf{U} \mathbf{x} \rangle \geq \alpha \|\mathbf{x}\|^2$.

$\mathcal{S}_N^+ := \mathcal{S}^+(\mathbb{R}^N)$ correspond à l'ensemble des matrices symétriques définies positives (notées SDP) de $\mathbb{R}^{N \times N}$. De plus, notons $\mathbf{Id}$ l'opérateur identité de $\mathcal{H}$, et $\mathbf{I}_N$ la matrice identité de $\mathbb{R}^{N \times N}$. On définira l'ordre partiel de Loewner sur $\mathcal{H}$ comme suit :

Définition 2.2. Soient $\mathbf{U}_1$ et $\mathbf{U}_2$ deux opérateurs de $\mathcal{S}(\mathcal{H})$. L'ordre partiel de Loewner sur $\mathcal{H}$ vérifie

$$\mathbf{U}_1 \succcurlyeq \mathbf{U}_2 \quad \Leftrightarrow \quad (\forall \mathbf{x} \in \mathcal{H}) \quad \langle \mathbf{x} \mid \mathbf{U}_1 \mathbf{x} \rangle \geq \langle \mathbf{x} \mid \mathbf{U}_2 \mathbf{x} \rangle.$$

Nous noterons le produit scalaire et la norme pondérés, relativement à $\mathbf{U} \in \mathcal{S}^+(\mathcal{H})$, de la façon suivante :

$$(\forall \mathbf{x} \in \mathcal{H}) (\forall \mathbf{y} \in \mathcal{H}) \quad \langle \mathbf{x} \mid \mathbf{y} \rangle_{\mathbf{U}} := \langle \mathbf{x} \mid \mathbf{U} \mathbf{y} \rangle \quad \text{et} \quad \|\mathbf{x}\|_{\mathbf{U}} := \langle \mathbf{x} \mid \mathbf{U} \mathbf{x} \rangle^{1/2}. \quad (2.36)$$

Remarquons que lorsque $\mathbf{U}$ est égal à $\mathbf{Id}$, alors $\|\cdot\|_{\mathbf{Id}} = \|\cdot\|$.

Définition 2.3. Soit $h: \mathcal{H} \rightarrow]-\infty, +\infty]$.

(i) Le domaine de la fonction h est défini par

$$\text{dom } h := \{\mathbf{x} \in \mathcal{H} \mid h(\mathbf{x}) < +\infty\}.$$

(ii) La fonction h est dite propre si $\text{dom } h$ est non vide.

(iii) Soit $\delta \in \mathbb{R}$. L'ensemble des lignes de niveau δ de h est donné par

$$\text{lev}_{\leq \delta} h := \{\mathbf{x} \in \mathcal{H} \mid h(\mathbf{x}) \leq \delta\}.$$

(iv) La fonction h est dite coercive si

$$\lim_{\|\mathbf{x}\| \rightarrow +\infty} h(\mathbf{x}) = +\infty.$$

(v) Le graphe de h est défini par

$$\text{Graph } h := \{(\mathbf{x}, \xi) \in \mathcal{H} \times \mathbb{R} \mid \xi = h(\mathbf{x})\}.$$

(vi) L'épigraphe de h est défini par

$$\text{epi } h := \{(\mathbf{x}, \xi) \in \mathcal{H} \times \mathbb{R} \mid h(\mathbf{x}) \leq \xi\}.$$

(vii) La fonction h est semi-continue inférieurement (s.c.i) en $\mathbf{x}_0 \in \mathcal{H}$ si

$$\liminf_{\mathbf{x} \rightarrow \mathbf{x}_0} h(\mathbf{x}) \geq h(\mathbf{x}_0).$$

La fonction h est semi-continue inférieurement si elle est s.c.i en tout point $\mathbf{x}_0 \in \mathcal{H}$.

(viii) Soit E un sous-ensemble de $\mathcal{H}$. La fonction indicatrice $\iota_E: \mathcal{H} \rightarrow]-\infty, +\infty]$ de cet ensemble est définie par

$$(\forall \mathbf{x} \in \mathcal{H}) \quad \iota_E(\mathbf{x}) = \begin{cases} 0 & \text{si } \mathbf{x} \in E, \\ +\infty & \text{sinon.} \end{cases}$$

(ix) Soit $g: \mathcal{H} \rightarrow]-\infty, +\infty]$. L'infimale convolution de h et g est définie par

$$h \square g: \mathcal{H} \rightarrow [-\infty, +\infty]: \mathbf{x} \mapsto \inf_{\mathbf{y} \in \mathcal{H}} (h(\mathbf{y}) + g(\mathbf{x} - \mathbf{y})).$$

L'élément neutre de l'infimale convolution est la fonction indicatrice de l'ensemble $\{\mathbf{0}\}$, où $\mathbf{0}$ représente le vecteur nul de $\mathcal{H}$:

$$(\forall \mathbf{x} \in \mathcal{H}) \quad h \square \iota_{\{\mathbf{0}\}}(\mathbf{x}) = h(\mathbf{x}).$$

Une représentation graphique des définitions 2.3(i) et (iii) est donnée dans la figure 2.7. De même, une représentation graphique de la définition 2.3(vi) est donnée dans la figure 2.8.

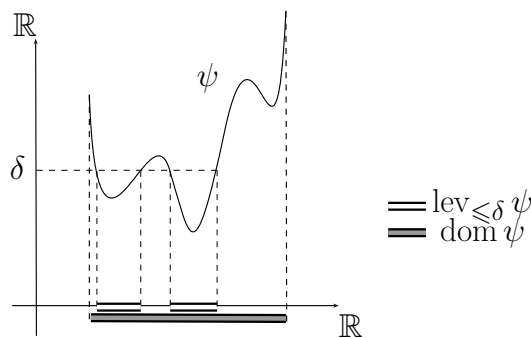


Figure 2.7 – Illustration du domaine et de l'ensemble des lignes de niveau $\delta > 0$ d'une fonction $\psi: \mathbb{R} \rightarrow \mathbb{R}$.

Lorsqu'une fonction associe à un point un ensemble, elle est dite *multivaluée*. La définition suivante généralise la notion de graphe pour ce type de fonctions.

Définition 2.4. Soit $L: \mathcal{H} \rightarrow 2^{\mathcal{G}}$ une fonction multivaluée où $\mathcal{G}$ est un ensemble de Hilbert séparable réel. Le graphe de L est défini par

$$\text{Graph } L := \{(\mathbf{x}, \mathbf{y}) \in \mathcal{H} \times \mathcal{G} \mid \mathbf{y} \in L(\mathbf{x})\}.$$

Les deux théorèmes suivants mettent en avant des relations existantes entre certaines des notions données dans la définition 2.3.

Théorème 2.1. [Rockafellar, 1970, Thm. 1.6][Bauschke et Combettes, 2011, Lem. 1.24] Soit $h: \mathcal{H} \rightarrow]-\infty, +\infty]$. La fonction h est semi-continue inférieurement si et seulement si, pour tout $\delta \in \mathbb{R}$, les ensembles $\text{lev}_{\leq \delta} h$ sont fermés dans $\mathcal{H}$. De plus, si h est coercive, alors $\text{lev}_{\leq \delta} h$ est un ensemble borné.

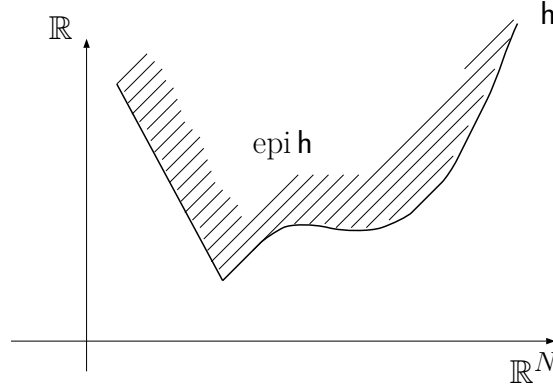


Figure 2.8 – Illustration de l'épigraphe d'une fonction $h: \mathbb{R}^N \rightarrow \mathbb{R}$.

Théorème 2.2. [Rockafellar, 1970, Thm. 1.9] Soit $h: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ une fonction propre, semi-continue inférieurement et coercive. L'ensemble des minimiseurs de h est non vide et compact.

La définition suivante introduit la notion de conjuguée d'une fonction. En particulier, cette notion peut se révéler utile lors de la résolution de certains problèmes inverses.

Définition 2.5. [Bauschke et Combettes, 2011; Rockafellar, 1974] Soit $h: \mathcal{H} \rightarrow]-\infty, +\infty]$ une fonction propre. La conjuguée (de Legendre-Fenchel) $h^*: \mathcal{H} \rightarrow]-\infty, +\infty]$ de la fonction h est définie par

$$(\forall \mathbf{x} \in \mathcal{H}) \quad h^*(\mathbf{x}) := \sup_{\mathbf{y} \in \mathcal{H}} \langle \mathbf{y} | \mathbf{x} \rangle - h(\mathbf{y}).$$

Remarque 2.4. [Bauschke et Combettes, 2011; Komodakis et Pesquet, 2014] Soient $h: \mathcal{H} \rightarrow]-\infty, +\infty]$, $g: \mathcal{G} \rightarrow]-\infty, +\infty]$, où $\mathcal{G}$ est un ensemble de Hilbert réel, et $\mathbf{L} \in \mathcal{B}(\mathcal{H}, \mathcal{G})$. On appellera problème primal le problème de minimisation

$$\text{trouver } \hat{\mathbf{x}} \in \underset{\mathbf{x} \in \mathcal{H}}{\text{Argmin}} \quad h(\mathbf{x}) + g(\mathbf{L}\mathbf{x}), \quad (2.37)$$

et problème dual

$$\text{trouver } \hat{\mathbf{v}} \in \underset{\mathbf{v} \in \mathcal{G}}{\text{Argmin}} \quad h^*(-\mathbf{L}^*\mathbf{v}) + g^*(\mathbf{v}). \quad (2.38)$$

Les solutions du problème (2.37) sont alors appelées solutions primales, et les solutions du problème (2.4) sont les solutions duales.

2.3.2 Analyse convexe

Dans cette section, nous nous intéresserons plus particulièrement à une sous-discipline de l'analyse variationnelle, fondamentale en optimisation, qui est l'analyse convexe.

Définition 2.6. [Bauschke et Combettes, 2011, Déf. 3.1] Soit C un sous-ensemble de $\mathcal{H}$. L'ensemble C est convexe si

$$(\forall \lambda \in]0, 1[) (\forall (\mathbf{x}, \mathbf{y}) \in C^2) \quad \lambda \mathbf{x} + (1 - \lambda) \mathbf{y} \in C.$$

Définition 2.7. Une fonction $f: \mathcal{H} \rightarrow]-\infty, +\infty]$ est dite

(i) [Bauschke et Combettes, 2011, Prop. 8.4] convexe si elle vérifie

$$(\forall \lambda \in]0, 1[) (\forall (\mathbf{x}, \mathbf{y}) \in (\text{dom } f)^2) \quad f(\lambda \mathbf{x} + (1 - \lambda) \mathbf{y}) \leq \lambda f(\mathbf{x}) + (1 - \lambda) f(\mathbf{y}).$$

(ii) [Bauschke et Combettes, 2011, Déf. 8.6] strictement convexe si elle vérifie

$$(\forall \lambda \in]0, 1[) (\forall (\mathbf{x}, \mathbf{y}) \in (\text{dom } f)^2) \\ \mathbf{x} \neq \mathbf{y} \quad \Rightarrow \quad f(\lambda \mathbf{x} + (1 - \lambda) \mathbf{y}) < \lambda f(\mathbf{x}) + (1 - \lambda) f(\mathbf{y}).$$

(iii) [Bauschke et Combettes, 2011, Déf. 10.5] fortement convexe de constante β si $\|f - (\beta/2)\| \cdot \|\cdot\|^2$ est convexe.

(iv) [Bauschke et Combettes, 2011, Déf. 10.18] quasi-convexe si ses lignes de niveaux $(\text{lev}_{\leq \xi} f)_{\xi \in \mathbb{R}}$ sont convexes.

Remarque 2.5. Si $f: \mathcal{H} \rightarrow]-\infty, +\infty]$ est une fonction convexe, alors son domaine est convexe.

Dans la suite nous noterons $\Gamma_0(\mathcal{H})$ l'ensemble des fonctions de $\mathcal{H}$ dans $]-\infty, +\infty]$ convexes, propres et semi-continues inférieurement.

Proposition 2.1. [Bauschke et Combettes, 2011] Si $\varphi \in \Gamma_0(\mathcal{H})$, alors les assertions suivantes sont vérifiées :

- (i) $\text{epi } \varphi$ est un ensemble non vide, convexe et fermé.
- (ii) $\varphi^* \in \Gamma_0(\mathcal{H})$ et $\varphi^{**} = (\varphi^*)^* = \varphi$.

La proposition suivante nous permet d'assurer l'existence de minimiseurs d'une somme de fonctions de $\Gamma_0(\mathcal{H})$.

Proposition 2.2. [Bauschke et Combettes, 2011, Cor. 11.15] Soient f et g des fonctions de $\Gamma_0(\mathcal{H})$. Supposons que $\text{dom } f \cap \text{dom } g \neq \emptyset$, que f est coercive et que g est bornée inférieurement. La somme $f + g$ est alors coercive et admet un minimiseur sur $\mathcal{H}$.

Notons qu'il existe d'autres conditions d'existence de minimiseurs de fonctions convexe, données par exemple dans [Bauschke et Combettes, 2011, Sec. 11.4] et dans [Combettes, 2013, Prop. 5.3].

Définition 2.8. [Bauschke et Combettes, 2011, Déf. 16.1] [Rockafellar, 1970, Sec. 23] Soit $f \in \Gamma_0(\mathcal{H})$. Le sous-différentiel de Moreau de f en $\mathbf{x} \in \text{dom } f$ est défini par

$$\partial f(\mathbf{x}) := \left\{ \mathbf{t} \in \mathcal{H} \mid (\forall \mathbf{y} \in \mathcal{H}) \quad f(\mathbf{y}) - \langle \mathbf{y} - \mathbf{x} \mid \mathbf{t} \rangle \geq f(\mathbf{x}) \right\}$$

Ainsi, les minimiseurs globaux d'une fonction de $\Gamma_0(\mathcal{H})$ peuvent être caractérisés par la *règle de Fermat* :

Théorème 2.3. [Bauschke et Combettes, 2011, Thm. 16.2] *Soit $f \in \Gamma_0(\mathcal{H})$. L'ensemble des minimiseurs globaux de f vérifie alors*

$$\text{Argmin } f = \{\mathbf{x} \in \mathcal{H} \mid \mathbf{0} \in \partial f(\mathbf{x})\}.$$

Comme nous l'avons vu dans la section 2.2, la résolution de problèmes inverses passe souvent par la minimisation d'une fonction f . Le théorème ci-dessus nous permet donc d'établir un lien entre la solution de ce problème et le sous-différentiel de f . Nous allons cependant voir que d'autres définitions du sous-différentiel sont possibles.

2.3.3 Calcul sous-différentiel

Dans la section 2.2 nous avons vu que les fonctions apparaissant dans les problèmes variationnels associés aux problèmes inverses peuvent être non différentiables et/ou non convexes. Dans cette section, nous donnerons donc les notions de base de sous-différentiabilité pour les fonctions non nécessairement convexes. Ainsi nous pourrons faire le lien entre les points critiques d'une fonction et son sous-différentiel.

2.3.3.1 Sous-différentiel

Nous considérerons dans cette section que $\mathcal{H} = \mathbb{R}^N$.

Définition 2.9. [Rockafellar et Wets, 1997, Def. 8.3], [Mordukhovich, 2006, Sec.1.3] *Soit $f: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ une fonction propre et soit $\mathbf{x} \in \text{dom } f$.*

- *Le sous-différentiel de Fréchet (ou sous-différentiel régulier) de f en $\mathbf{x}$ correspond à l'ensemble suivant*

$$\widehat{\partial}f(\mathbf{x}) := \left\{ \hat{\mathbf{t}} \in \mathbb{R}^N \mid \liminf_{\substack{\mathbf{y} \rightarrow \mathbf{x} \\ \mathbf{y} \neq \mathbf{x}}} \frac{1}{\|\mathbf{x} - \mathbf{y}\|} (f(\mathbf{y}) - f(\mathbf{x}) - \langle \mathbf{y} - \mathbf{x}, \hat{\mathbf{t}} \rangle) \geq 0 \right\}.$$

Si $\mathbf{x} \notin \text{dom } f$, alors $\widehat{\partial}f(\mathbf{x}) = \emptyset$.

- *Le sous-différentiel (limite) de f en $\mathbf{x}$ correspond à l'ensemble suivant*

$$\partial f(\mathbf{x}) := \left\{ \mathbf{t} \in \mathbb{R}^N \mid \exists \mathbf{y}_k \rightarrow \mathbf{x}, f(\mathbf{y}_k) \rightarrow f(\mathbf{x}), \hat{\mathbf{t}}_k \in \widehat{\partial}f(\mathbf{y}_k) \rightarrow \mathbf{t} \right\}.$$

Dans le cas où la fonction $f: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ est convexe et propre, alors pour tout $\mathbf{x} \in \mathbb{R}^N$ les sous-différentiels ci-dessus coïncident avec le sous-différentiel utilisé en analyse convexe [Rockafellar et Wets, 1997, Prop. 8.12].

Définition 2.10. *Soit $f: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ une fonction propre. L'ensemble des points critiques de f correspond aux zéros de ∂f .*

Remarque 2.6. Lorsque $f \in \Gamma_0(\mathbb{R}^N)$, les points critiques de f correspondent à ses minimiseurs globaux.

Lorsque $f: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ n'est pas nécessairement convexe, la règle de Fermat donnée par le théorème 2.3 n'est pas valable. Cependant, une des implications reste vraie :

Proposition 2.3. Si $f: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ est une fonction propre, alors on a

$$\mathbf{x} \in \text{Argmin } f \quad \Rightarrow \quad \mathbf{0} \in \partial f(\mathbf{x}).$$

Le théorème suivant permet de faire le lien entre la notion de sous-différentiel d'une fonction propre et la notion de cône normal :

Théorème 2.4. [Rockafellar et Wets, 1997, Thm. 8.9] Soit $f: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ une fonction propre. Pour tout $\mathbf{x} \in \text{dom } f$, on a

$$\partial f(\mathbf{x}) = \{\mathbf{t} \in \mathbb{R}^N \mid (\mathbf{t}, -1) \in N_{\text{epi } f}(\mathbf{x}, f(\mathbf{x}))\},$$

où $N_{\text{epi } f}$ est le cône normal de l'ensemble $\text{epi } f$.

Remarque 2.7. Ce théorème permet de se représenter graphiquement le sous différentiel d'une fonction propre non nécessairement convexe. Un exemple est donné dans la figure 2.9 dans le cas d'une fonction $f: \mathbb{R} \rightarrow]-\infty, +\infty]$.

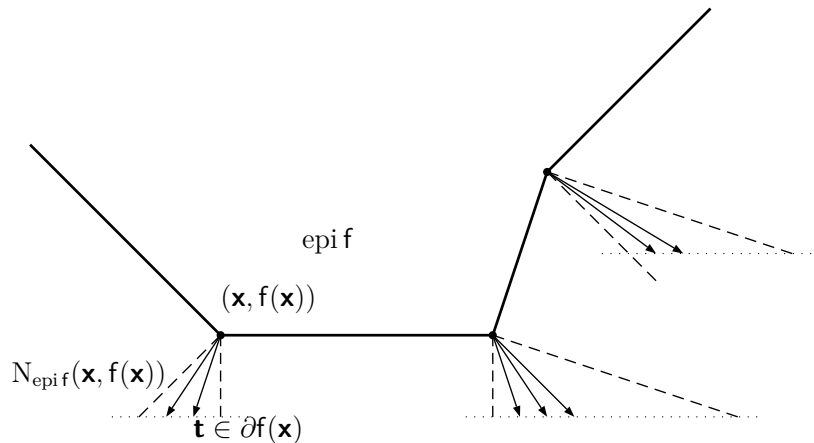


Figure 2.9 – Cônes normaux de l'épigraph d'une fonction $f: \mathbb{R} \rightarrow]-\infty, +\infty]$.

La propriété suivante, de fermeture du graphe de ∂f , est une conséquence directe de la proposition 6.6 de [Rockafellar et Wets, 1997] et du théorème 2.4 précédent :

Proposition 2.4. Soit $(\mathbf{x}_k, \mathbf{t}_k)_{k \in \mathbb{N}}$ une suite de l'ensemble $\text{Graph } \partial f$. Si $(\mathbf{x}_k, \mathbf{t}_k)$ converge vers $(\mathbf{x}, \mathbf{t})$ et $f(\mathbf{x}_k)$ converge vers $f(\mathbf{x})$, alors $(\mathbf{x}, \mathbf{t}) \in \text{Graph } \partial f$.

Notons que cette propriété est utile pour démontrer la convergence de suites générées par des algorithmes minimisant des critères non convexes [Attouch *et al.*, 2010a, 2011; Bolte *et al.*, 2013]. Cependant, comme il est indiqué dans [Rockafellar et Wets, 1997, Chap. 6], elle n'est vérifiée que pour le sous-différentiel limite, mais pas pour le sous-différentiel de Fréchet.

2.3.3.2 Sous-différentiel partiel

Dans le chapitre 5 nous nous intéresserons à la minimisation de fonctions séparables par blocs, définies en dimension finie. Nous allons donc donner ici les règles de sous-différentiation de ces fonctions.

Soit $f: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ une fonction séparable par bloc. Plus précisément, soit $J \in \{1, \dots, N\}$, on peut alors décomposer tout vecteur $\mathbf{x} \in \mathbb{R}^N$ de la façon suivante

$$\mathbf{x} = (\mathbf{x}^{(j)})_{1 \leq j \leq J} \in \mathbb{R}^{N_1} \times \dots \times \mathbb{R}^{N_J}, \quad (2.39)$$

où $\sum_{j=1}^J N_j = N$, et on suppose que la fonction f peut s'écrire

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad f(\mathbf{x}) = \sum_{j=1}^J f_j(\mathbf{x}^{(j)}), \quad (2.40)$$

avec, pour tout $j \in \{1, \dots, J\}$, $f_j: \mathbb{R}^{N_j} \rightarrow]-\infty, +\infty]$.

La proposition suivante nous donne les règles de base de sous-différentiation d'une fonction séparable.

Proposition 2.5. [Rockafellar et Wets, 1997, Prop. 10.5] *Soit $f: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ une fonction propre séparable vérifiant (2.40). On a alors*

$$(\forall \mathbf{x} = (\mathbf{x}^{(j)})_{1 \leq j \leq J} \in \text{dom } f) \quad \partial f(\mathbf{x}) = \partial f_1(\mathbf{x}^{(1)}) \times \dots \times \partial f_J(\mathbf{x}^{(J)}).$$

2.3.3.3 Gradient

Lorsque la fonction $f: \mathcal{H} \rightarrow]-\infty, +\infty]$ (avec $\mathcal{H} = \mathbb{R}^N$ si f est non convexe) est différentiable au point $\mathbf{x} \in \mathcal{H}$, on a

$$\partial f(\mathbf{x}) = \{\nabla f(\mathbf{x})\}.$$

La notion suivante sera utile pour analyser la convergence des algorithmes d'optimisation que nous étudierons par la suite :

Définition 2.11. *Une fonction différentiable $f: \mathcal{H} \rightarrow]-\infty, +\infty]$ est dite de gradient β -Lipschitz sur un sous-ensemble E de $\mathcal{H}$, pour $\beta > 0$, si*

$$(\forall (\mathbf{x}, \mathbf{y}) \in E^2) \quad \|\nabla f(\mathbf{x}) - \nabla f(\mathbf{y})\| \leq \beta \|\mathbf{x} - \mathbf{y}\|.$$

Lorsqu'une fonction $f: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ est la somme d'une fonction différentiable et d'une fonction non différentiable, la règle de différentiation est alors donnée par la proposition suivante :

Proposition 2.6. [Rockafellar et Wets, 1997, Ex. 8.8] Soit $f: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ une fonction propre telle que, pour tout $\mathbf{x} \in \mathbb{R}^N$, $f(\mathbf{x}) = f_1(\mathbf{x}) + f_2(\mathbf{x})$, où f_1 est une fonction propre différentiable sur $\text{dom } f$ et f_2 est une fonction propre. On a alors

$$(\forall \mathbf{x} \in \text{dom } f) \quad \partial f(\mathbf{x}) = \nabla f_1(\mathbf{x}) + \partial f_2(\mathbf{x}).$$

2.3.4 Opérateurs proximaux

2.3.4.1 Définition

L'opérateur proximal en un point $\mathbf{x} \in \mathcal{H}$ a été défini dans [Moreau, 1965] pour une fonction $f \in \Gamma_0(\mathcal{H})$ comme étant l'unique minimiseur de la fonction $f + \frac{1}{2}\|\cdot - \mathbf{x}\|^2$. On a ainsi, pour $\gamma > 0$,

$$\text{prox}_{\gamma f}: \mathcal{H} \rightarrow \mathcal{H}: \mathbf{x} \mapsto \text{prox}_{\gamma f}(\mathbf{x}) := \underset{\mathbf{y} \in \mathcal{H}}{\text{argmin}} \quad f(\mathbf{y}) + \frac{1}{2\gamma}\|\mathbf{y} - \mathbf{x}\|^2. \quad (2.41)$$

Cet opérateur a été largement étudié pour des fonctions convexes usuelles du traitement du signal, par exemple dans [Chaux et al., 2007; Combettes et Pesquet, 2007b, 2010; Combettes et Wajs, 2005; Parikh et Boyd, 2013]. Cette définition peut être généralisée pour des fonctions $f: \mathcal{H} \rightarrow]-\infty, +\infty]$ non nécessairement convexes dans le cas de normes pondérées (c.f. [Hiriart-Urruty et Lemaréchal, 1993, Sec. XV.4] et [Attouch et al., 2011; Chouzenoux et al., 2014b; Combettes et Vũ, 2013]) :

Définition 2.12. Soient $f: \mathcal{H} \rightarrow]-\infty, +\infty]$ une fonction propre et semi-continue inférieurement, $\mathbf{U} \in \mathcal{S}^+(\mathcal{H})$, et $\mathbf{x} \in \mathcal{H}$. L'opérateur proximal de f en $\mathbf{x}$ relativement à la métrique induite par $\mathbf{U}$ est l'ensemble défini par

$$\text{prox}_{\mathbf{U}, f}(\mathbf{x}) = \underset{\mathbf{y} \in \mathcal{H}}{\text{Argmin}} \quad f(\mathbf{y}) + \frac{1}{2}\|\mathbf{y} - \mathbf{x}\|_{\mathbf{U}}^2. \quad (2.42)$$

Par abus de notation, lorsque cet ensemble est égal à un singleton, on note $\text{prox}_{\mathbf{U}, f}(\mathbf{x})$ son unique élément.

Remarque 2.8.

- (i) Si $\mathbf{U} = \gamma^{-1}\mathbf{Id}$, pour $\gamma > 0$, et $f \in \Gamma_0(\mathcal{H})$, alors $\text{prox}_{\gamma^{-1}\mathbf{Id},f} \equiv \text{prox}_{\gamma f}$ correspond à l'opérateur proximal (2.41) de [Moreau, 1965]. De plus, pour tout $\mathbf{x} \in \mathcal{H}$, $\text{prox}_{\gamma f}(\mathbf{x})$ est unique.
- (ii) Si $\mathcal{H} = \mathbb{R}^N$ et $\mathbf{U} = \gamma^{-1}\mathbf{I}$, pour $\gamma > 0$, alors $\text{prox}_{\gamma f}$ est l'opérateur proximal défini dans [Rockafellar et Wets, 1997, Déf. 1.22] et utilisé dans [Attouch et al., 2011; Bolte et al., 2013].
- (iii) Notons que, pour $\mathcal{H} = \mathbb{R}^N$ et $\mathbf{U} \in \mathcal{S}_N^+$,

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \text{prox}_{\mathbf{U},f}(\mathbf{x}) = \mathbf{U}^{-1/2} \text{prox}_{f \circ \mathbf{U}^{-1/2}}(\mathbf{U}^{1/2}\mathbf{x}). \quad (2.43)$$

En dimension finie, lorsque $\mathcal{H} = \mathbb{R}^N$ et $f: \mathcal{H} \rightarrow]-\infty, +\infty]$ est une fonction non nécessairement convexe, la question d'existence de l'opérateur proximal de f se pose. Nous allons donc énoncer un ensemble de conditions permettant d'assurer que, pour tout $\mathbf{x} \in \mathbb{R}^N$, l'ensemble $\text{prox}_{\gamma f}(\mathbf{x})$ soit non vide.

Définition 2.13. [Rockafellar et Wets, 1997, Déf. 1.23] Soit $f: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ une fonction propre et semi-continue inférieurement. On dit que f est *prox-bornée* s'il existe $\gamma > 0$ tel que

$$(\exists \mathbf{x} \in \mathbb{R}^N) \quad \inf_{\mathbf{y} \in \mathbb{R}^N} f(\mathbf{y}) + \frac{1}{2\gamma} \|\mathbf{y} - \mathbf{x}\|^2 > -\infty. \quad (2.44)$$

De plus, on note $\gamma_f = \sup \{ \gamma > 0 \mid (2.44) \text{ est satisfaite} \}$ et f est dite γ_f -prox-bornée.

Théorème 2.5. [Rockafellar et Wets, 1997, Thm. 1.25] Si $f: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ est une fonction propre, semi-continue inférieurement et γ_f -prox-bornée, de seuil $\gamma_f > 0$, alors, pour tout $\gamma \in]0, \gamma_f[$ et tout $\mathbf{x} \in \mathbb{R}^N$, l'ensemble $\text{prox}_{\gamma f}(\mathbf{x})$ est non vide et compact.

En particulier, si $\gamma_f = +\infty$, alors pour tout $\gamma > 0$ l'ensemble $\text{prox}_{\gamma f}(\mathbf{x})$ est non vide et compact.

Proposition 2.7. [Rockafellar et Wets, 1997, Ex. 1.24] Soit $f: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ une fonction propre et semi-continue inférieurement. Les assertions suivantes sont équivalentes :

- (i) f est prox-bornée.
- (ii) Il existe une fonction polynomiale q de degré inférieur ou égal à 2 tel que $f \geq q$.
- (iii) Il existe $\lambda \in \mathbb{R}$ tel que $\inf f + \frac{\lambda}{2} \|\cdot\|^2 > -\infty$.
- (iv) $\liminf_{\|\mathbf{x}\| \rightarrow +\infty} \frac{f(\mathbf{x})}{\|\mathbf{x}\|^2} > -\infty$.

Remarque 2.9. Remarquons que si λ_f est l'infimum des λ vérifiant (iii) alors le seuil de f est $\gamma_f = 1/\max\{0, \lambda_f\}$ (avec $1/0 = \infty$). En particulier, si f est bornée inférieurement, i.e. $\inf f > -\infty$ et donc $\lambda_f = 0$, alors f est γ_f -prox-bornée, avec $\gamma_f = +\infty$, et donc, pour tout $\gamma > 0$ et $\mathbf{x} \in \mathbb{R}^N$, $\text{prox}_{\gamma f}(\mathbf{x}) \neq \emptyset$.

2.3.4.2 Propriétés

L'opérateur proximal étant très utilisé dans la résolution de problèmes variationnels par le biais d'algorithmes proximaux, il a largement été étudié ces dernières années, en particulier pour des fonctions $f \in \Gamma_0(\mathcal{H})$ dans le cadre de la métrique usuelle euclidienne [Chaux *et al.*, 2007; Combettes et Pesquet, 2007b, 2010; Combettes et Wajs, 2005] et de métriques pondérées [Combettes et Vũ, 2013, 2014] (voir aussi [Becker et Fadili, 2012] en dimension finie). Nous listons dans cette section quelques propriétés utiles des opérateurs proximaux de fonctions de $\Gamma_0(\mathcal{H})$, et donnons leur équivalent, lorsque cela est possible, pour des fonctions non convexes quand $\mathcal{H} = \mathbb{R}^N$.

Dans cette section, $\mathcal{H}$ désigne soit un espace de Hilbert réel lorsque les fonctions étudiées sont convexes, soit $\mathbb{R}^N$ lorsque les fonction ne sont pas nécessairement convexes.

(i) Caractérisation :

Soient $\mathbf{U} \in \mathcal{S}^+(\mathcal{H})$, $\mathbf{x} \in \mathcal{H}$ et $f: \mathcal{H} \rightarrow]-\infty, +\infty]$ une fonction propre, semi-continue inférieurement, telle que $\text{prox}_{\mathbf{U},f}(\mathbf{x}) \neq \emptyset$. On a alors

- Si $\mathbf{p} \in \text{prox}_{\mathbf{U},f}(\mathbf{x})$, alors $\mathbf{x} - \mathbf{p} \in \mathbf{U}^{-1}\partial f(\mathbf{p})$.
- Si f est convexe, alors

$$\mathbf{p} = \text{prox}_{\mathbf{U},f}(\mathbf{x}) \Leftrightarrow \mathbf{x} - \mathbf{p} \in \mathbf{U}^{-1}\partial f(\mathbf{p}). \quad (2.45)$$

(ii) Point fixe :

Soient $\mathbf{U} \in \mathcal{S}^+(\mathcal{H})$ et $f: \mathcal{H} \rightarrow]-\infty, +\infty]$ une fonction propre, semi-continue inférieurement, telle que pour tout $\mathbf{x} \in \mathbb{R}^N$ $\text{prox}_{\mathbf{U},f}(\mathbf{x}) \neq \emptyset$. On a alors

- Si $\mathbf{x} \in \text{prox}_{\mathbf{U},f}(\mathbf{x})$, alors $\mathbf{x}$ est un point critique de f .
- Si f est convexe, alors

$$\mathbf{x} = \text{prox}_{\mathbf{U},f}(\mathbf{x}) \Leftrightarrow \mathbf{x} \in \text{Argmin } f.$$

(iii) Translation :

Soient $\mathbf{U} \in \mathcal{S}^+(\mathcal{H})$, $\mathbf{x}_0 \in \mathcal{H}$ et $g: \mathcal{H} \rightarrow]-\infty, +\infty]$ une fonction propre, semi-continue inférieurement, telle que pour tout $\mathbf{x} \in \mathcal{H}$ $\text{prox}_{\mathbf{U},g}(\mathbf{x}) \neq \emptyset$. Soit $f = f(\cdot - \mathbf{x}_0)$. Pour tout $\mathbf{x} \in \mathcal{H}$, $\text{prox}_{\mathbf{U},f}(\mathbf{x}) \neq \emptyset$ et on a alors

$$\text{prox}_{\mathbf{U},f}(\mathbf{x}) = \{\mathbf{x}_0 + \mathbf{p} \mid \mathbf{p} \in \text{prox}_{\mathbf{U},g}(\mathbf{x} - \mathbf{x}_0)\} = \mathbf{x}_0 + \text{prox}_{\mathbf{U},g}(\mathbf{x} - \mathbf{x}_0).$$

(iv) **Changement d'échelle :**

Soient $\mathbf{U} \in \mathcal{S}^+(\mathcal{H})$, $\alpha \in \mathbb{R}^*$ et $g: \mathcal{H} \rightarrow]-\infty, +\infty]$ une fonction propre, semi-continue inférieurement, telle que pour tout $\mathbf{x} \in \mathcal{H}$ $\text{prox}_{\mathbf{U},g}(\mathbf{x}) \neq \emptyset$. Soit $f = g(\cdot/\alpha)$. Pour tout $\mathbf{x} \in \mathcal{H}$, $\text{prox}_{\mathbf{U},f}(\mathbf{x}) \neq \emptyset$ et on a alors

$$\text{prox}_{\mathbf{U},f}(\mathbf{x}) = \alpha \text{prox}_{\mathbf{U},g/\alpha^2}(\mathbf{x}/\alpha).$$

(v) **Perturbation quadratique :**

Soient $\mathbf{U} \in \mathcal{S}^+(\mathcal{H})$, $\mathbf{A} \in \mathcal{S}^+(\mathcal{H})$ et $g: \mathcal{H} \rightarrow]-\infty, +\infty]$ une fonction propre, semi-continue inférieurement, telle que pour tout $\mathbf{x} \in \mathcal{H}$ $\text{prox}_{\mathbf{U},g}(\mathbf{x}) \neq \emptyset$. Définissons $f = g + \frac{1}{2}\|\cdot\|_{\mathbf{A}}^2 + \beta\langle \cdot | \mathbf{x}_0 \rangle + \gamma$, où $(\beta, \gamma) \in \mathbb{R}^2$ et $\mathbf{x}_0 \in \mathcal{H}$. On a alors

$$(\forall \mathbf{x} \in \mathcal{H}) \quad \text{prox}_{\mathbf{U},f}(\mathbf{x}) = \text{prox}_{\mathbf{A}+\mathbf{U},g} \left((\mathbf{A} + \mathbf{U})^{-1}(\mathbf{U}\mathbf{x} - \beta\mathbf{x}_0) \right). \quad (2.46)$$

Une démonstration de cette propriété est donnée dans l'annexe 2.A.

Remarquons que si $\mathbf{U} = \mathbf{Id}$, $\mathbf{A} = \alpha\mathbf{Id}$, où $\alpha \in]0, +\infty[$, et si $g \in \Gamma_0(\mathcal{H})$ dans (2.46), on retrouve l'égalité énoncée dans [Combettes et Wajs, 2005] :

$$(\forall \mathbf{x} \in \mathcal{H}) \quad \text{prox}_f(\mathbf{x}) = \text{prox}_{g/(\alpha+1)} \left(\frac{\mathbf{x} - \beta\mathbf{x}_0}{\alpha+1} \right).$$

(vi) **Décomposition de Moreau :**

Soient $f \in \Gamma_0(\mathcal{H})$ et $\mathbf{U} \in \mathcal{S}^+(\mathcal{H})$. On a alors [Combettes et Vũ, 2014, Ex. 3.9]

$$(\forall \mathbf{x} \in \mathcal{H}) \quad \mathbf{x} = \text{prox}_{\mathbf{U},f^*}(\mathbf{x}) + \mathbf{U}^{-1} \text{prox}_{\mathbf{U}^{-1},f}(\mathbf{U}\mathbf{x}). \quad (2.47)$$

Remarquons que lorsque $\mathbf{U} = \alpha\mathbf{Id}$, pour $\alpha \in]0, +\infty[$, dans l'égalité (2.47), on retrouve la décomposition de Moreau usuelle [Bauschke et Combettes, 2011, Thm. 14.3] :

$$(\forall \mathbf{x} \in \mathcal{H}) \quad \mathbf{x} = \text{prox}_{\alpha f}(\mathbf{x}) + \alpha \text{prox}_{\alpha^{-1}f^*}(\alpha^{-1}\mathbf{x}).$$

(vii) **Fonction additivement séparables :**

Soient $(\mathcal{H}_n)_{1 \leq n \leq N}$ des espaces de Hilbert et $\mathcal{H}$ leur somme directe. Soient, pour tout $n \in \{1, \dots, N\}$, $\mathbf{U}^{(n)} \in \mathcal{S}^+(\mathcal{H}_n)$, $\mathbf{x} = (\mathbf{x}^{(n)})_{1 \leq n \leq N} \mapsto \mathbf{U}\mathbf{x} = (\mathbf{U}_n \mathbf{x}^{(n)})$, et $f: \mathcal{H} \rightarrow]-\infty, +\infty]$ une fonction additivement séparable, i.e. vérifiant

$$(\forall \mathbf{x} = (\mathbf{x}^{(n)})_{1 \leq n \leq N} \in \mathcal{H}) \quad f(\mathbf{x}) = \sum_{n=1}^N f_n(\mathbf{x}^{(n)}). \quad (2.48)$$

Supposons que, pour tout $n \in \{1, \dots, N\}$, $f_n: \mathcal{H}_n \rightarrow]-\infty, +\infty]$ soit une fonction semi-continue inférieurement, propre, et telle que, pour tout $\mathbf{x} \in \mathcal{H}_n$, $\text{prox}_{f_n}(\mathbf{x}) \neq \emptyset$. On a alors

$$(\forall \mathbf{x} \in \mathcal{H}) \quad \text{prox}_{\mathbf{U},f}(\mathbf{x}) = (\mathbf{p}^{(n)})_{1 \leq n \leq N}, \quad (2.49)$$

où, pour tout $n \in \{1, \dots, N\}$, $\mathbf{p}^{(n)} = \text{prox}_{\mathbf{U}_n, f_n}(\mathbf{x}^{(n)})$.

Lorsque $\mathcal{H} = \mathbb{R}^N$ et, pour tout $n \in \{1, \dots, N\}$, $f_n \in \Gamma_0(\mathbb{R})$, un résultat plus général a été énoncé dans [Becker et Fadili, 2012] :

Proposition 2.8. [Becker et Fadili, 2012, Cor. 9] *Soient $f \in \Gamma_0(\mathbb{R}^N)$ une fonction additivement séparable et $\mathbf{U} = \text{Diag}(\mathbf{u}^{(1)}, \dots, \mathbf{u}^{(N)}) + \mathbf{v}\mathbf{v}^\top$ où $\mathbf{v} \in \mathbb{R}^N$ et, pour tout $n \in \{1, \dots, N\}$, $\mathbf{u}^{(n)} > 0$. On a alors*

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \text{prox}_{\mathbf{U},f}(\mathbf{x}) = (\mathbf{p}^{(n)})_{1 \leq n \leq N}, \quad (2.50)$$

où

$$(\forall n \in \{1, \dots, N\}) \quad \mathbf{p}^{(n)} = \text{prox}_{f_n/\mathbf{u}^{(n)}} \left(\mathbf{x}^{(n)} - \frac{\alpha^* \mathbf{v}^{(n)}}{\mathbf{u}^{(n)}} \right), \quad (2.51)$$

et α^* est l'unique zéro de la fonction

$$\alpha \in \mathbb{R} \mapsto \left\langle \mathbf{v}, \mathbf{x} - \left(\text{prox}_{f_n/\mathbf{u}^{(n)}}(\mathbf{x}^{(n)} - \alpha \mathbf{v}^{(n)}/\mathbf{u}^{(n)}) \right) \right\rangle + \alpha.$$

2.3.4.3 Cas particuliers

Dans cette section nous donnerons quelques exemples utiles d'opérateurs proximaux de fonctions usuelles pour la résolution de problèmes inverses.

(i) **Opérateur de projection :**

Définition 2.14. *Soit E un sous-ensemble de $\mathcal{H}$. La projection de $\mathbf{x} \in \mathcal{H}$ sur E est définie par*

$$\Pi_E(\mathbf{x}) \in \underset{\mathbf{y} \in E}{\text{Argmin}} \quad \|\mathbf{y} - \mathbf{x}\|. \quad (2.52)$$

En d'autres termes, un élément de $\Pi_E(\mathbf{x})$ est un point de E tel que sa distance à $\mathbf{x}$ soit minimale.

Remarquons que l'opérateur proximal peut être vu comme une généralisation de l'opérateur de projection au sens où, pour tout sous-ensemble E de $\mathcal{H}$,

$$\Pi_E = \text{prox}_{\iota_E}, \quad (2.53)$$

où ι_E désigne la fonction indicatrice de E .

Proposition 2.9. [Bauschke et Combettes, 2011, Thm 3.14] *Si C est un sous-ensemble convexe fermé non vide de $\mathcal{H}$, alors la projection $\Pi_C(\mathbf{x})$ de $\mathbf{x} \in \mathcal{H}$ sur C existe et est unique.*

(ii) **Opérateurs proximaux de fonctions puissances :**

Dans ce paragraphe, nous allons donner les opérateurs proximaux des fonctions

$$f_\alpha : \mathbb{R}^N \rightarrow]-\infty, +\infty] : \mathbf{x} \mapsto \sum_{n=1}^N |\mathbf{x}^{(n)}|^\alpha, \quad (2.54)$$

pour différents $\alpha \in [0, +\infty[$. Notons que f étant additivement séparable, dans la suite nous n'étudierons que les opérateurs proximaux des fonctions $\psi_\alpha : x \in \mathbb{R} \mapsto |x|^\alpha$.

- Lorsque $\alpha \in \{0, 1\}$, l'opérateur proximal de ψ_α est appelé *opérateur de seuillage*. En particulier, lorsque $\alpha = 1$ nous parlerons de seuillage doux [Combettes et Pesquet, 2010], et lorsque $\alpha = 0$ (en utilisant la convention $0^0 = 0$), nous parlerons de seuillage dur [Blumensath et Davies, 2008] (voir figure 2.10).

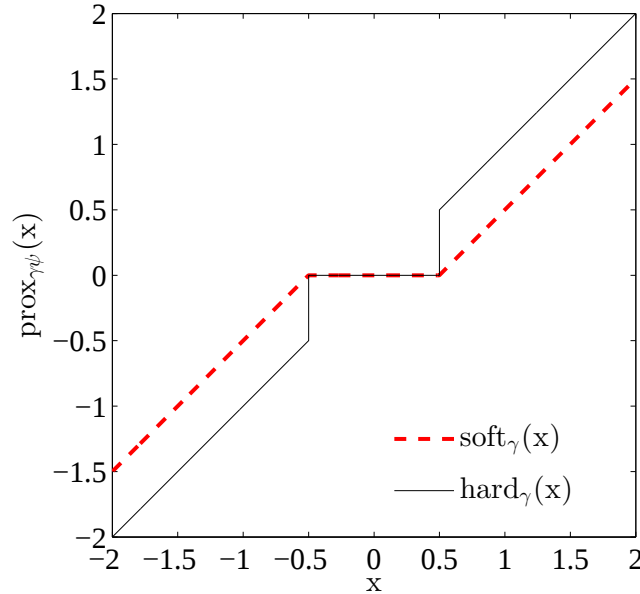


Figure 2.10 – Opérateurs de seuillage doux (courbe rouge discontinue) et de seuillage dur (courbe noire continue).

Définition 2.15.

(a) Soit $\gamma > 0$. On définit l'opérateur de seuillage doux de seuil γ comme suit :

$$(\forall x \in \mathbb{R}) \quad \text{soft}_\gamma(x) = \text{sign}(x) \max\{|x| - \gamma, 0\} \quad (2.55)$$

$$= \begin{cases} x - \gamma & \text{si } x > \gamma, \\ 0 & \text{si } x \in [-\gamma, \gamma], \\ x + \gamma & \text{si } x < -\gamma. \end{cases} \quad (2.56)$$

(b) Soit $\gamma > 0$. On définit l'opérateur de seuillage dur de seuil γ comme suit :

$$(\forall x \in \mathbb{R}) \quad \text{hard}_\gamma(x) = \begin{cases} 0 & \text{si } |x| \leq \gamma, \\ x & \text{sinon.} \end{cases} \quad (2.57)$$

Le résultat suivant permet de construire d'autres opérateurs de seuillage à partir de l'opérateur de seuillage doux :

Proposition 2.10. [Combettes et Pesquet, 2007b, Prop. 3.6] Soit $\varphi \in \Gamma_0(\mathbb{R})$ une fonction différentiable en 0 telle que $\varphi'(0) = 0$. Soit

$$\phi: \mathbb{R} \rightarrow]-\infty, +\infty]: x \mapsto \varphi(x) + \gamma|x|,$$

pour $\gamma > 0$. On a alors

$$\text{prox}_\phi = \text{prox}_\varphi \circ \text{soft}_\gamma. \quad (2.58)$$

- Lorsque $\alpha \in [1, +\infty[$, l'opérateur proximal de ψ_α est donné par la proposition suivante.

Proposition 2.11. [Chaux et al., 2007; Combettes et Pesquet, 2010] Soient $\alpha \geq 1$ et $\gamma > 0$. L'opérateur proximal de ψ_α est donné par

$$(\forall x \in \mathbb{R}) \quad \text{prox}_{\gamma\psi_\alpha}(x) = \text{sign}(x)p, \quad (2.59)$$

où $p \geq 0$ et satisfait $p + \alpha\gamma p^{\alpha-1} = |x|$.

Cette proposition permet d'obtenir une caractérisation de l'opérateur proximal de la fonction ψ_α , pour $\alpha \geq 1$. Cependant, seules certaines valeurs de α conduisent à une forme explicite de cet opérateur. En voici quelques exemples [Chaux et al., 2007, Ex. 4.2 et 4.3] :

$$(\forall x \in \mathbb{R}) \quad \text{prox}_{\gamma\psi_\alpha}(x) = \begin{cases} \text{soft}_\gamma(x) & \text{si } \alpha = 1, \\ \frac{x}{1 + 2\gamma} & \text{si } \alpha = 2, \\ \text{sign}(x) \frac{\sqrt{1 + 12\gamma|x|} - 1}{6\gamma} & \text{si } \alpha = 3, \\ \left(\frac{t+x}{8\gamma}\right)^{1/3} - \left(\frac{t-x}{8\gamma}\right)^{1/3} & \text{si } \alpha = 4. \\ \text{où } t = \sqrt{x^2 + 1/(27\gamma)} \end{cases}$$

Remarquons en particulier que l'on retrouve l'opérateur de seuillage doux lorsque $\alpha = 1$. Une illustration de ces opérateurs proximaux ainsi que des fonctions associées est donnée dans la figure 2.11.

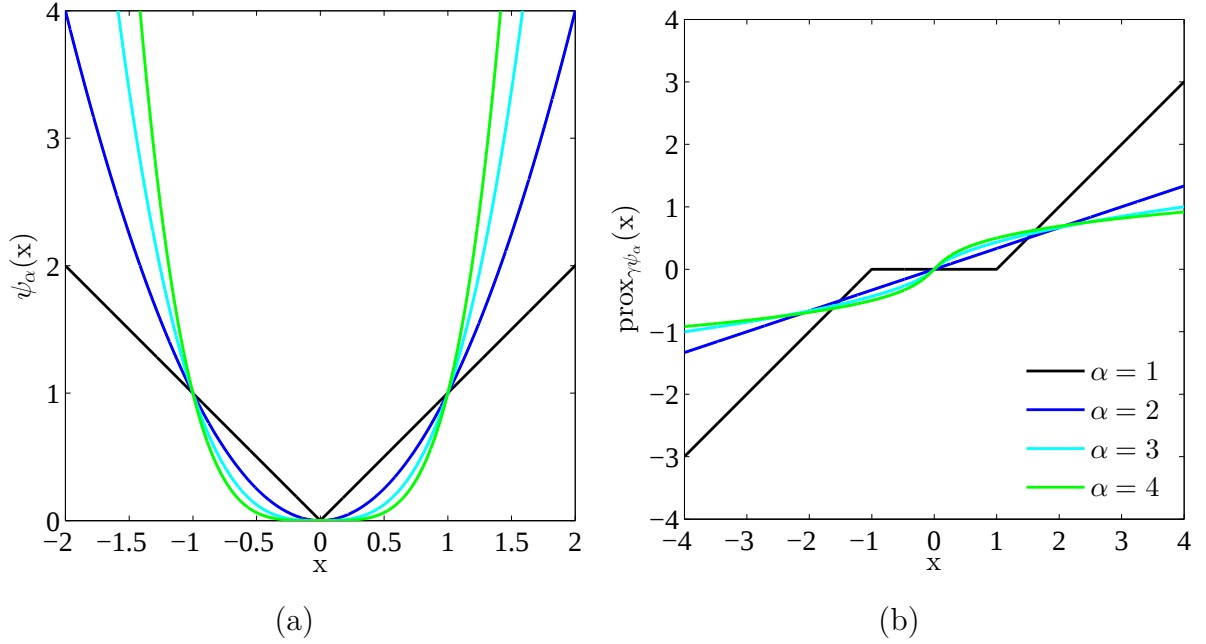


Figure 2.11 – (a) Fonction ψ_α pour $\alpha \in \{1, 2, 3, 4\}$. (b) Opérateur proximal $\text{prox}_{\gamma\psi_\alpha}$ pour $\gamma = 1$ et $\alpha \in \{1, 2, 3, 4\}$.

- Lorsque $\alpha \in]0, 1[$, l'opérateur proximal de la fonction $|\cdot|^\alpha$ correspond à un seuillage et doit être calculé de façon itérative [Antoniadis *et al.*, 2002; Marjanovic et Solo, 2012].
- Nous allons maintenant introduire la fonction de Huber généralisée $\varphi_{\delta,\alpha} : \mathbb{R} \rightarrow \mathbb{R}^+$, pour $\alpha \in]0, 1]$ et $\delta > 0$, définie par

$$\varphi_{\delta,\alpha}(x) = \begin{cases} \frac{x^2}{2\delta} & \text{si } |x| \leq \delta^{1/(2-\alpha)}, \\ \frac{|x|^\alpha}{\alpha} - \left(\frac{1}{\alpha} - \frac{1}{2}\right) \delta^{\alpha/(2-\alpha)} & \text{sinon.} \end{cases} \quad (2.60)$$

Remarquons que, si $\alpha = 1$, $\varphi_{\delta,1} = \varphi_\delta$ est la fonction de Huber définie par (2.28). De plus, la fonction de Huber généralisée peut être vue comme une approximation lisse de la fonction $\psi_\alpha = |\cdot|^\alpha$, pour $\alpha \in]0, 1[$, et, en particulier, elle n'est pas convexe. Une illustration de cette fonction est donnée dans la figure 2.12 pour différentes valeurs de α . Dans [Chartrand, 2012], cette fonction est utilisée pour définir un opérateur de seuillage, que nous utiliserons dans le chapitre 6 dans une application en reconstruction de données hyperspectrales. Ce seuillage est donné dans la proposition suivante.

Proposition 2.12. [Chartrand, 2012] Soient $\alpha \in]0, 1]$, $\delta > 0$, et $\varrho_{\delta,\alpha} \in \Gamma_0(\mathbb{R})$ définie par

$$(\forall x \in \mathbb{R}) \quad \varrho_{\delta,\alpha}(x) = \frac{x^2}{2} - \delta \varphi_{\delta,\alpha}(x). \quad (2.61)$$

L'opérateur proximal de la fonction $\phi_{\delta,\alpha} : \mathbb{R} \rightarrow]-\infty, +\infty]$ définie par

$$(\forall x \in \mathbb{R}) \quad \phi_{\delta,\alpha}(x) = \varrho_{\delta,\alpha}^*(x) - \frac{x^2}{2}, \quad (2.62)$$

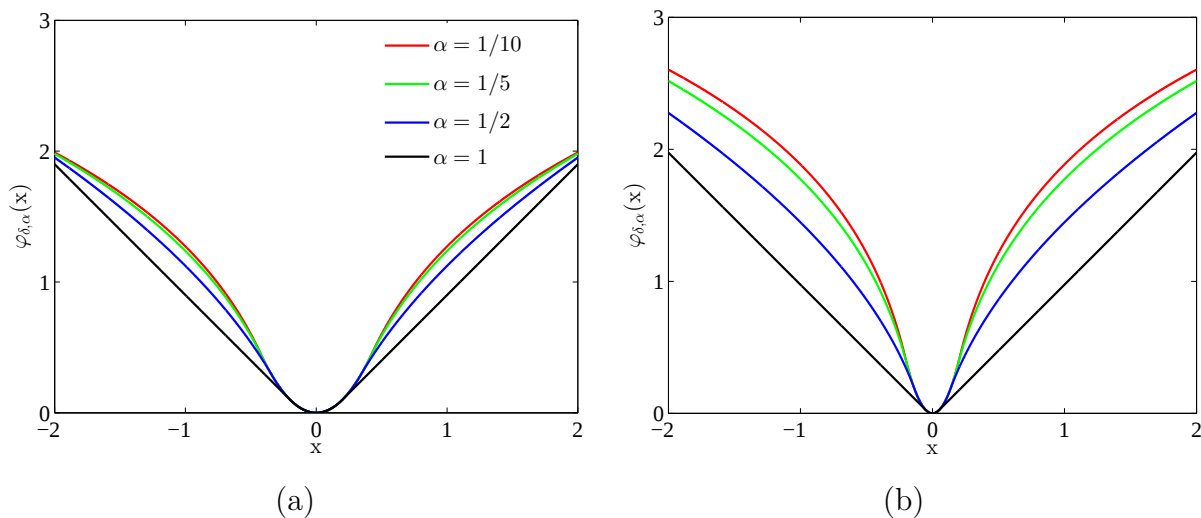


Figure 2.12 – Fonction de Huber généralisée pour $\alpha \in \{1/10, 1/5, 1/2, 1\}$, avec (a) $\delta = 1/5$ et (b) $\delta = 1/20$.

est un seuillage donné par

$$(\forall x \in \mathbb{R}) \quad \text{prox}_{\phi_{\delta,\alpha}}(x) = \begin{cases} \max\{0, |x| - \delta|x|^{\alpha-1}\} \frac{x}{|x|} & \text{si } x \neq 0, \\ 0 & \text{si } x = 0. \end{cases} \quad (2.63)$$

Les fonctions $\varrho_{\delta,\alpha}$, $\phi_{\delta,\alpha}$ et $\text{prox}_{\phi_{\delta,\alpha}}$, définies dans la proposition 2.12, sont tracées dans la figure 2.13.

2.3.5 Inégalité de Kurdyka-Łojasiewicz

2.3.5.1 Motivation

Soit $f: \mathbb{R}^N \rightarrow \mathbb{R}$ une fonction différentiable. On considère la dynamique du gradient donnée par

$$\begin{cases} (\forall t \in]0, +\infty[) & \dot{\mathbf{x}}(t) = -\nabla f(\mathbf{x}(t)), \\ \mathbf{x}(0) & = \mathbf{x}_0. \end{cases}$$

La résolution de cette équation différentielle génère une trajectoire de gradient $t \in]0, +\infty[\mapsto \mathbf{x}(t)$ partant de $\mathbf{x}_0$ et suivant une trajectoire changeant à chaque intersection avec une ligne de niveau de f de façon à prendre la direction orthogonale à la ligne de niveau croisée. La solution $\mathbf{x}(t)$ générée assure que $t \in]0, +\infty[\mapsto f(\mathbf{x}(t))$ soit décroissante. Cette dynamique peut être vue comme une version continue de l'algorithme de descente de gradient donné par

$$(\forall k \in \mathbb{N}) \quad \mathbf{x}_{k+1} = \mathbf{x}_k - \mu_k \nabla f(\mathbf{x}_k),$$

où $\mu_k > 0$.

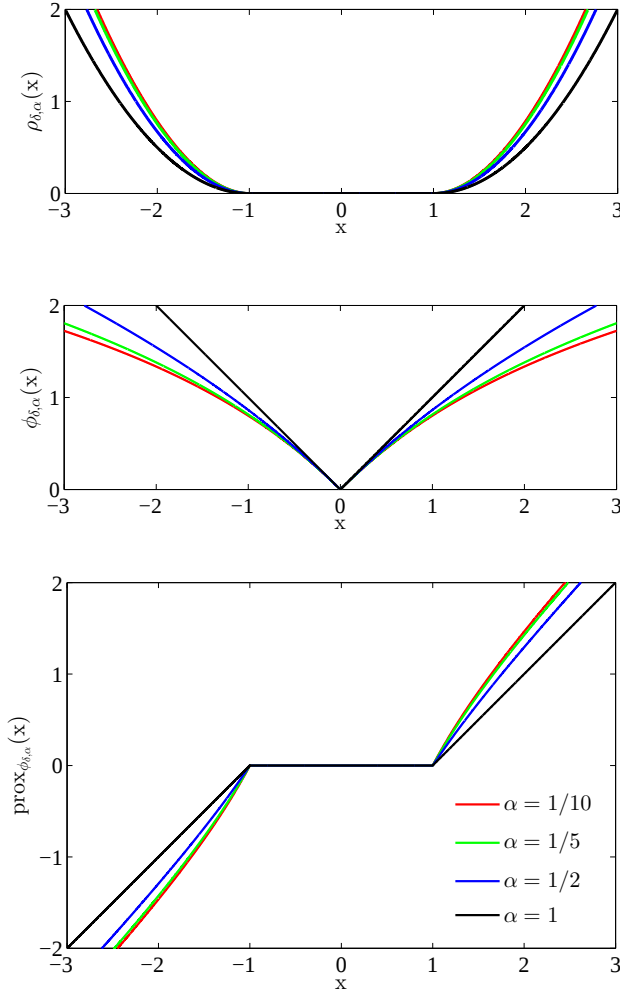


Figure 2.13 – Fonctions $\rho_{1,\alpha}$ (haut), $\phi_{1,\alpha}$ (centre), et $\text{prox}_{\phi_{1,\alpha}}$ (bas), pour $\alpha \in \{1/10, 1/5, 1/2, 1\}$.

Considérons une trajectoire de gradient $\mathbf{x}(t)$ (resp. une suite $(\mathbf{x}_k)_{k \in \mathbb{N}}$) bornée. On voudrait pouvoir assurer la convergence de $\mathbf{x}(t)$ (resp. une suite $(\mathbf{x}_k)_{k \in \mathbb{N}}$) lorsque $t \rightarrow +\infty$ (resp. $k \rightarrow \infty$), vers un point critique de f , i.e. assurer que la trajectoire $\mathbf{x}(t)$ soit de longueur finie.

Dans [Absil *et al.*, 2005; Palis et De Melo, 1982] (voir aussi [Bolte *et al.*, 2007, Ex. 1]), des contre-exemples ont été exhibés proposant des fonctions infiniment dérivables de $\mathbb{R}$ ou de $\mathbb{R}^2$ dont la courbe de gradient $\mathbf{x}(t)$ est bornée, mais de longueur infinie (elle tourne indéfiniment autour du cercle unité). Il a été démontré dans ces mêmes références que les fonctions vérifiant l'inégalité de Kurdyka-Łojasiewicz ne présentaient pas ce type de problème.

2.3.5.2 Inégalité de Łojasiewicz

L'inégalité de Kurdyka-Łojasiewicz (KL) apparaît ainsi comme un outil fondamental en optimisation non convexe qui permet de démontrer la convergence de suites minimisant des fonctions non nécessairement convexes. Cette inégalité a été proposée en premier dans [Łojasiewicz, 1963], et il a été démontré qu'elle est satisfaite par toute fonction analytique réelle :

Propriété 2.2. (Inégalité de Łojasiewicz) *Soient $f: \mathbb{R}^N \rightarrow \mathbb{R}$ une fonction analytique réelle, et $\hat{\mathbf{x}}$ un point critique de cette fonction. Il existe $\theta \in [1/2, 1[$ et $\kappa \in]0, +\infty[$ tels que f satisfait l'inégalité de Łojasiewicz dans un voisinage $\mathcal{V}(\hat{\mathbf{x}})$ de $\hat{\mathbf{x}}$, i.e. :*

$$(\forall \mathbf{x} \in \mathcal{V}(\hat{\mathbf{x}})) \quad |f(\mathbf{x}) - f(\hat{\mathbf{x}})|^\theta \leq \kappa \|\nabla f(\mathbf{x})\|. \quad (2.64)$$

Cette inégalité a été utilisée pour démontrer la convergence d'algorithmes de descente de gradient pour la minimisation de fonctions différentiables non nécessairement convexe [Absil et al., 2005], elle permet d'assurer que la courbe de gradient générée par ce type d'algorithme est de longueur finie [Absil et al., 2005] (voir aussi [Bolte et al., 2010]).

2.3.5.3 Inégalité de Kurdyka-Łojasiewicz

L'inégalité de Łojasiewicz a été généralisée dans [Kurdyka, 1998] pour les fonctions différentiables définissables dans une structure o-minimale (pour les structures o-minimales, voir par exemple [Coste, 1999; Van den Dries et Miller, 1996]). On parlera alors d'inégalité de KL. Plus récemment, dans [Bolte et al., 2006a,b, 2007] il a été démontré que l'inégalité de KL peut être généralisée pour des fonctions non nécessairement différentiables en utilisant les sous-différentiel donné dans la définition 2.9. Cette inégalité, plus générale que (2.64), s'écrit de la façon suivante :

Définition 2.16. [Attouch et al., 2010a, Déf. 7] *La fonction $f: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ satisfait l'inégalité de Kurdyka-Łojasiewicz en $\hat{\mathbf{x}} \in \text{dom } \partial f$ s'il existe $\xi \in]0, +\infty]$, un voisinage $\mathcal{V}(\hat{\mathbf{x}})$ de $\hat{\mathbf{x}}$, et une fonction $\varphi: [0, \xi[\rightarrow [0, +\infty[$ concave tels que*

- (i) $\varphi(0) = 0$,
- (ii) φ est continue en 0 et continument différentiable sur $]0, \xi[$,
- (iii) pour tout $u \in]0, \xi[$, $\varphi'(u) > 0$,
- (iv) $\varphi'(f(\mathbf{x}) - f(\hat{\mathbf{x}})) \text{dist}(\mathbf{0}, \partial f(\mathbf{x})) \geq 1$, pour tout $\mathbf{x} \in E$ tel que $f(\hat{\mathbf{x}}) \leq f(\mathbf{x}) \leq f(\hat{\mathbf{x}}) + \xi$, où $\text{dist}(\mathbf{0}, \partial f(\mathbf{x})) = \inf_{\mathbf{u} \in \partial f(\mathbf{x})} \|\mathbf{u}\|$.

Cette propriété a permis, entre autres, de démontrer la convergence d'algorithmes proximaux [Attouch et Bolte, 2009], et d'algorithmes proximaux alternés [Attouch et al., 2010a] pour des fonctions non nécessairement convexes.

2.3.5.4 Fonctions satisfaisant l'inégalité de Kurdyka-Łojasiewicz

Il a été démontré que l'inégalité de Kurdyka-Łojasiewicz est satisfaite par les fonctions définissables dans une structure o-minimale, et plus particulièrement par les fonctions semi-algébriques et analytiques réelles [Attouch et al., 2010a; Bolte et al., 2006a, 2007, 2010].

Les structures o-minimales réelles, introduites dans [Van den Dries, 1998], sont définies comme suit :

Définition 2.17. [Attouch et al., 2010a, Déf. 13] Soit $\mathcal{O} = \{\mathcal{O}_N\}_{N \in \mathbb{N}}$ tel que, pour tout $N \in \mathbb{N}$, $\mathcal{O}_N$ est une collection de sous-ensembles de $\mathbb{R}^N$. On dit que la famille $\mathcal{O}$ est une structure o-minimale sur $\mathbb{R}$ si les axiomes suivants sont vérifiés :

- (i) Chaque ensemble $\mathcal{O}_N$ est une algèbre booléenne, i.e. $\emptyset \in \mathcal{O}_N$ et, pour tout A, B dans $\mathcal{O}_N$, les ensembles $A \cup B$, $A \cap B$, et $\mathbb{R}^N \setminus A$ sont dans $\mathcal{O}_N$.
- (ii) Pour tout A dans $\mathcal{O}_N$, $A \times \mathbb{R}$ et $\mathbb{R} \times A$ sont dans $\mathcal{O}_{N+1}$.
- (iii) Pour tout A dans $\mathcal{O}_{N+1}$, $\{(x_1, \dots, x_N) \in \mathbb{R}^N \mid (x_1, \dots, x_N, x_{N+1}) \in A\} \in \mathcal{O}_N$.
- (iv) Pour tout $(m, n) \in \{1, \dots, N\}^2$ tel que $m \neq n$, $\{(x_1, \dots, x_N) \in \mathbb{R}^N \mid x_m = x_n\} \in \mathcal{O}_N$.
- (v) L'ensemble $\{(x_1, x_2) \in \mathbb{R}^2 \mid x_1 < x_2\}$ est dans $\mathcal{O}_2$.
- (vi) $\mathcal{O}_1$ coïncide exactement avec l'ensemble des unions finies d'intervalles.

Nous pouvons en déduire la définition d'ensembles et de fonctions définissables :

Définition 2.18. Soit $\mathcal{O} = \{\mathcal{O}_N\}_{N \in \mathbb{N}}$ une structure o-minimale.

- (i) Un ensemble E est définissable dans $\mathcal{O}$ si $E \subset \mathcal{O}_N$.
- (ii) Une fonction $f: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ est définissable dans $\mathcal{O}$ si son graphe est définissable.

Pour certains cas particuliers de structures o-minimales, la forme de la fonction φ de l'inégalité de KL est connue [Attouch et al., 2010a, Sec. 4.3]. Plus particulièrement, dans la suite de cette thèse nous nous intéresserons aux fonctions semi-algébriques réelles (dont la définition est donnée ci-dessous) et analytiques réelles.

Définition 2.19. [Attouch et al., 2011; Bochnak et al., 1998; Shiota, 1997]

- (i) Un sous-ensemble S de $\mathbb{R}^N$ est un ensemble semi-algébrique réel s'il existe un nombre fini de fonctions polynomiales $\mathbf{p}_{i,j}, \mathbf{q}_{i,j}: \mathbb{R}^N \rightarrow \mathbb{R}$, pour $i \in \{1, \dots, p\}$ et $j \in \{1, \dots, q\}$, telles que

$$S = \bigcup_{j=1}^p \bigcap_{i=1}^q \{ \mathbf{x} \in \mathbb{R}^N \mid \mathbf{p}_{i,j}(\mathbf{x}) = 0, \mathbf{q}_{i,j}(\mathbf{x}) < 0 \}.$$

- (ii) Une fonction $\mathbf{f}: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ est semi-algébrique si son graphe est un sous-ensemble semi-algébrique de $\mathbb{R}^{N+1}$.

Remarque 2.10. [Bolte et al., 2006b] La classe des ensembles semi-algébriques est stable pour les opérations de réunion finie, d'intersection finie, de produit cartésien, par passage au complémentaire et par projection (principe de Tarski-Seidenberg [Bochnak et al., 1998]).

Pour les fonctions semi-algébriques réelles et analytiques réelles, la fonction φ de l'inégalité de KL est de la forme $\varphi: u \in [0, \xi[\mapsto u^{1-\theta}$, pour $\theta \in [0, 1[$:

Définition 2.20. [Attouch et Bolte, 2009] La fonction $\mathbf{f}: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ satisfait l'inégalité de Kurdyka-Łojasiewicz si, pour tout $\xi \in \mathbb{R}$, et, pour tout sous-ensemble borné E de $\mathbb{R}^N$, il existe trois constantes $\kappa > 0$, $\zeta > 0$ et $\theta \in [0, 1[$ telles que

$$(\forall \mathbf{t} \in \partial \mathbf{f}(\mathbf{x})) \quad \|\mathbf{t}\| \geq \kappa |\mathbf{f}(\mathbf{x}) - \xi|^\theta,$$

pour tout $\mathbf{x} \in E$ tel que $|\mathbf{f}(\mathbf{x}) - \xi| \leq \zeta$ (avec la convention $0^0 = 0$).

Cette forme de l'inégalité de KL est utilisée pour démontrer des résultats de convergence dans [Attouch et Bolte, 2009], et permet d'obtenir de plus des résultats de vitesse de convergence [Attouch et Bolte, 2009; Attouch et al., 2010a; Frankel et al., 2014].

Remarque 2.11. Comme souligné dans [Attouch et al., 2010a], les fonctions logarithme et exponentielle sont définissables. Il existe donc une fonction φ telle que l'inégalité de KL dans sa forme la plus générale (définition 2.16) soit satisfaite. Cependant, ces fonctions ne sont ni semi-algébriques réelles, ni analytiques réelles, et elles ne satisfont pas l'inégalité donnée dans la définition 2.20.

2.4 ALGORITHMES D'OPTIMISATION

Cette section a pour but de donner un rapide survol des méthodes d'optimisation de base permettant de résoudre des problèmes variationnels usuels.

2.4.1 Algorithmes de gradient et de sous-gradient

Dans un premier temps, nous allons nous intéresser au problème d'optimisation suivant :

$$\text{trouver } \hat{\mathbf{x}} \in \underset{\mathbf{x} \in \mathcal{H}}{\text{Argmin}} \ h(\mathbf{x}), \quad (2.65)$$

où $h \in \Gamma_0(\mathcal{H})$, et $\mathcal{H}$ est un espace de Hilbert réel. D'après la règle de Fermat, $\hat{\mathbf{x}}$ est une solution de (2.65) si et seulement si $\mathbf{0} \in \partial h(\hat{\mathbf{x}})$.

Algorithme de gradient

Dans le cas particulier où h est une fonction différentiable et de gradient β -Lipschitz ($\beta > 0$), cette condition devient alors

$$\mathbf{0} = \nabla h(\hat{\mathbf{x}}). \quad (2.66)$$

On peut donc trouver une solution explicite du problème de minimisation (2.65) pour des fonctions h assez simples. Cependant, un tel $\hat{\mathbf{x}}$ peut se révéler difficile à trouver dès lors que h a une forme plus compliquée et que l'on ne sait plus résoudre (2.66). Dans ce cas, il faut définir la solution du problème de minimisation comme étant la limite d'une suite construite itérativement, par le biais d'algorithmes. Une méthode permettant de trouver une solution itérative du problème de minimisation (2.65) est connue sous le nom d'*algorithme du gradient* (ou de *descente de gradient*) [Bertsekas, 1999, Chap. 1] et est donnée par l'algorithme 2.1.

Algorithme 2.1 Algorithme du gradient

Initialisation : Soient $\mathbf{x}_0 \in \mathbb{R}^N$ et, pour tout $k \in \mathbb{N}$, $\underline{\gamma} \leq \gamma_k \leq (2 - \bar{\gamma})\beta^{-1}$, avec $(\underline{\gamma}, \bar{\gamma}) \in]0, +\infty[^2$.

Itérations :

Pour $k = 0, 1, \dots$

$$\left| \mathbf{x}_{k+1} = \mathbf{x}_k - \gamma_k \nabla h(\mathbf{x}_k). \right.$$

Algorithme de sous-gradient

Lorsque h n'est pas différentiable, une méthode permettant de résoudre le problème (2.65) consiste à utiliser un *algorithme de sous-gradient* [Bertsekas, 1999, Chap. 6], donné par l'algorithme 2.2.

Algorithme 2.2 Algorithme de sous-gradient

Initialisation : Soient $\mathbf{x}_0 \in \mathbb{R}^N$ et, pour tout $k \in \mathbb{N}$, $\gamma_k \in]0, +\infty[$ tels que $\sum_{k \in \mathbb{N}} \gamma_k = +\infty$ et $\sum_{k \in \mathbb{N}} \gamma_k^2 < +\infty$.

Itérations :

Pour $k = 0, 1, \dots$

$\left[\begin{array}{l} \mathbf{x}_{k+1} = \mathbf{x}_k - \gamma_k \mathbf{t}_k, \quad \text{où } \mathbf{t}_k \in \partial h(\mathbf{x}_k). \end{array} \right.$

La convergence de la suite $(h(\mathbf{x}_k))_{k \in \mathbb{N}}$ générée par cet algorithme, vers une solution du problème 2.65, est assurée lorsque le pas $(\gamma_k)_{k \in \mathbb{N}}$ satisfait une condition de carré sommable [Shor, 1985] (voir aussi [Bertsekas, 1999, Ex. 6.3.14] et [Boyd et al., 2003]). Notons que cette hypothèse peut se révéler contraignante en pratique, puisqu'elle peut freiner la convergence et poser des problèmes numériques.

Algorithme proximal

Un moyen pour pallier le problème énoncé dans le paragraphe précédent, est de définir, à chaque itération $k \in \mathbb{N}$, $\mathbf{t}_k$ comme étant un élément du sous-différentiel de h au point $\mathbf{x}_{k+1}$:

$$(\forall k \in \mathbb{N}) \quad \mathbf{x}_{k+1} = \mathbf{x}_k - \gamma_k \mathbf{t}_k, \quad \text{où } \mathbf{t}_k \in \partial h(\mathbf{x}_{k+1}), \quad (2.67)$$

ce qui peut sembler contre-intuitif puisque $\mathbf{x}_{k+1}$ n'est pas encore construit à l'itération k . De façon équivalente, l'équation (2.67) se réécrit

$$(\forall k \in \mathbb{N}) \quad \mathbf{x}_k - \mathbf{x}_{k+1} \in \gamma_k \partial h(\mathbf{x}_{k+1}). \quad (2.68)$$

Autrement dit, d'après la propriété (2.45), pour trouver l'itérée $\mathbf{x}_{k+1}$ telle que (2.67) soit satisfaite, il suffit de la définir comme étant l'opérateur proximal de $\gamma_k h$ en $\mathbf{x}_k$. Ainsi on obtient l'*algorithme du point proximal* [Rockafellar, 1976] donné par l'algorithme 2.3.

Algorithme 2.3 Algorithme du point proximal

Initialisation : Soient $\mathbf{x}_0 \in \mathcal{H}$ et, pour tout $k \in \mathbb{N}$, $\gamma_k \in]0, +\infty[$.

Itérations :

Pour $k = 0, 1, \dots$

$\left[\begin{array}{l} \mathbf{x}_{k+1} \in \text{prox}_{\gamma_k h}(\mathbf{x}_k). \end{array} \right.$

Lorsque h est convexe, rappelons que l'opérateur proximal de h est univalué. Un des résultats de convergence les plus généraux est alors donné par le théorème suivant :

Théorème 2.6. [Bauschke et Combettes, 2011, Thm. 27.1] *Soit $h \in \Gamma_0(\mathcal{H})$ telle que $\text{Argmin } h \neq \emptyset$. Supposons que $(\mathbf{x}_k)_{k \in \mathbb{N}}$ est une suite générée par l'algorithme 2.3, et que $(\gamma_k)_{k \in \mathbb{N}}$ soit telle que $\sum_{k \in \mathbb{N}} \gamma_k = +\infty$. La suite $(\mathbf{x}_k)_{k \in \mathbb{N}}$ converge alors vers une solution du problème 2.65.*

Remarquons que ce résultat de convergence autorise l'emploi de pas $(\gamma_k)_{k \in \mathbb{N}}$ ne convergeant pas vers 0.

Il a été démontré récemment dans [Attouch et al., 2011] que, lorsque $\mathcal{H} = \mathbb{R}^N$, la suite $(\mathbf{x}_k)_{k \in \mathbb{N}}$ générée par l'algorithme 2.3 converge lorsque l'opérateur proximal est calculé de façon approchée et la fonction h n'est pas nécessairement convexe mais satisfait l'inégalité de KL :

Théorème 2.7. [Attouch et al., 2011, Thm. 4.2] *Soit $h: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ une fonction propre, s.c.i., et bornée inférieurement, satisfaisant l'inégalité de KL. On suppose de plus que h est continue sur son domaine. Si $(\mathbf{x}_k)_{k \in \mathbb{N}}$ est une suite bornée générée par l'algorithme 2.3, et $(\gamma_k)_{k \in \mathbb{N}}$ est tel que $\underline{\gamma} \leq \gamma_k \leq \overline{\gamma}$, pour $(\underline{\gamma}, \overline{\gamma}) \in]0, +\infty[$, alors*

- (i) *la suite $(\mathbf{x}_k)_{k \in \mathbb{N}}$ converge vers un point critique $\hat{\mathbf{x}}$ de h ;*
- (ii) *la suite $(\mathbf{x}_k)_{k \in \mathbb{N}}$ est de longueur finie, i.e.*

$$\sum_{k \in \mathbb{N}} \|\mathbf{x}_{k+1} - \mathbf{x}_k\| < +\infty. \quad (2.69)$$

2.4.2 Algorithme explicite-implicite

Algorithme

Dans la section précédente, nous avons présenté des algorithmes permettant de minimiser une fonction convexe, différentiable ou non. Nous allons maintenant nous intéresser au cas où la fonction à minimiser est éclatée en une somme de deux fonctions, l'une étant différentiable tandis que l'autre ne l'est pas nécessairement :

$$\text{trouver } \hat{\mathbf{x}} \in \underset{\mathbf{x} \in \mathcal{H}}{\text{Argmin}} \quad (f(\mathbf{x}) = h(\mathbf{x}) + g(\mathbf{x})), \quad (2.70)$$

où $h \in \Gamma_0(\mathcal{H})$ est différentiable de gradient β -Lipschitz, pour $\beta > 0$, $g \in \Gamma_0(\mathcal{H})$, et $\mathcal{H}$ est un espace de Hilbert réel.

Proposition 2.13. [Combettes et Wajs, 2005, Prop. 3.1] *Soient $\hat{\mathbf{x}} \in \mathcal{H}$ et $\gamma \in]0, +\infty[$. Les assertions suivantes sont équivalentes :*

- (i) *$\hat{\mathbf{x}}$ est une solution du problème (2.70).*
- (ii) *$\hat{\mathbf{x}} = \text{prox}_{\gamma g}(\hat{\mathbf{x}} - \gamma \nabla h(\hat{\mathbf{x}}))$.*

D'après cette caractérisation de point fixe, une idée intuitive pour définir un algorithme itératif permettant de trouver une solution du problème (2.70) est :

$$(\forall k \in \mathbb{N}) \quad \mathbf{x}_{k+1} = \text{prox}_{\gamma g}(\mathbf{x}_k - \gamma \nabla h(\mathbf{x}_k)), \quad (2.71)$$

pour un certain $\gamma > 0$ bien choisi. À chaque itération $k \in \mathbb{N}^*$, la suite $(\mathbf{x}_k)_{k \in \mathbb{N}}$ est donc définie comme étant la composition d'une étape de gradient sur la fonction différentiable h (appelée étape *explicite*), et d'une étape proximale sur la fonction non lisse g (appelée étape *implicite*).

De manière plus rigoureuse, un algorithme permettant de résoudre le problème (2.70) est l'algorithme *explicite-implicite* (ou *forward-backward* en anglais) générant une suite de $\mathcal{H}$ notée $(\mathbf{x}_k)_{k \in \mathbb{N}}$, grâce aux itérations suivantes [Chen et Rockafellar, 1997; Tseng, 1998] :

Algorithme 2.4 Algorithme explicite-implicite

Initialisation : Soient $\mathbf{x}_0 \in \mathcal{H}$ et, pour tout $k \in \mathbb{N}$, $(\gamma_k, \lambda_k) \in]0, +\infty[^2$.

Itérations :

Pour $k = 0, 1, \dots$

$$\left\{ \begin{array}{l} \mathbf{y}_k \in \text{prox}_{\gamma_k \mathbf{g}}(\mathbf{x}_k - \gamma_k \nabla \mathbf{h}(\mathbf{x}_k)), \\ \mathbf{x}_{k+1} = \mathbf{x}_k + \lambda_k (\mathbf{y}_k - \mathbf{x}_k). \end{array} \right.$$

La suite des itérées $(\mathbf{x}_k)_{k \in \mathbb{N}}$ générée par l'algorithme 2.4 converge vers une solution du problème (2.70) :

Théorème 2.8. [Combettes et Wajs, 2005, Thm. 3.4] *Supposons que $\mathbf{h} \in \Gamma_0(\mathcal{H})$ est différentiable et de gradient β -Lipschitz, pour $\beta > 0$, et $\mathbf{g} \in \Gamma_0(\mathcal{H})$. Supposons de plus que $\text{Argmin}(\mathbf{h} + \mathbf{g}) \neq \emptyset$. Soit $(\mathbf{x}_k)_{k \in \mathbb{N}}$ une suite générée par l'algorithme 2.4. Si*

- $0 < \inf_{l \in \mathbb{N}} \gamma_l \leq \sup_{l \in \mathbb{N}} \gamma_l < 2\beta^{-1}$,
- $(\forall k \in \mathbb{N}) \ 0 < \inf_{l \in \mathbb{N}} \lambda_l \leq \lambda_k \leq 1$,

alors, $(\mathbf{x}_k)_{k \in \mathbb{N}}$ converge vers une solution du problème (2.70).

Récemment, dans le cas où $\mathcal{H} = \mathbb{R}^N$, les propriétés de convergence de l'algorithme explicite-implicite ont été étendues au cas où les fonctions $\mathbf{h}$ et $\mathbf{g}$ sont non nécessairement convexes dans [Attouch et Bolte, 2009; Attouch et al., 2011], lorsque $\lambda_k \equiv 1$. Ces résultats de convergence reposent principalement sur l'hypothèse que la fonction $\mathbf{f}$ satisfait l'inégalité de KL :

Théorème 2.9. [Attouch et al., 2011, Thm. 5.] *Supposons que $\mathbf{f}$ est une fonction propre, semi-continue inférieurement, bornée inférieurement, satisfaisant l'inégalité de KL. Supposons de plus que $\mathbf{h} : \mathbb{R}^N \rightarrow \mathbb{R}$ est différentiable, de gradient β -Lipschitz, pour $\beta > 0$, et que $\mathbf{g} : \mathbb{R}^N \rightarrow]-\infty, +\infty]$ est continue sur son domaine. Si $(\mathbf{x}_k)_{k \in \mathbb{N}}$ est une suite bornée générée par l'algorithme 2.4, $\lambda_k \equiv 0$, et $(\gamma_k)_{k \in \mathbb{N}}$ vérifie $0 < \inf_{\ell \in \mathbb{N}} \gamma_\ell \leq \sum_{\ell \in \mathbb{N}} \gamma_\ell < \beta^{-1}$, alors*

- (i) *la suite $(\mathbf{x}_k)_{k \in \mathbb{N}}$ converge vers un point critique $\hat{\mathbf{x}}$ de $\mathbf{f}$;*
- (ii) *la suite $(\mathbf{x}_k)_{k \in \mathbb{N}}$ est de longueur finie, i.e.*

$$\sum_{k \in \mathbb{N}} \|\mathbf{x}_{k+1} - \mathbf{x}_k\| < +\infty. \quad (2.72)$$

Versions inexactes de l'algorithme explicite-implicite

En pratique, l'opérateur proximal de la fonction $\mathbf{g}$ n'a pas nécessairement une forme explicite et doit être calculé de façon itérative. Il existe plusieurs méthodes permettant de prendre en compte des termes d'erreurs éventuels.

L'une d'elles, donnée dans l'algorithme 2.5, consiste à considérer des termes d'erreurs sommables.

Algorithme 2.5 Algorithme explicite-implicite inexact (erreurs additives)

Initialisation : Soient $\mathbf{x}_0 \in \mathcal{H}$ et, pour tout $k \in \mathbb{N}$, $(\gamma_k, \lambda_k) \in]0, +\infty[^2$, $\mathbf{a}_k \in \mathcal{H}$ et $\mathbf{b}_k \in \mathcal{H}$.

Itérations :

Pour $k = 0, 1, \dots$

$$\left\{ \begin{array}{l} \mathbf{y}_k = \mathbf{x}_k + \lambda_k \left(\text{prox}_{\lambda_k \mathbf{g}}(\mathbf{x}_k - \gamma_k \nabla \mathbf{h}(\mathbf{x}_k) + \mathbf{b}_k) + \mathbf{a}_k - \mathbf{x}_k \right), \\ \mathbf{x}_{k+1} = \mathbf{x}_k + \lambda_k (\mathbf{y}_k - \mathbf{x}_k). \end{array} \right.$$

Dans cet algorithme, pour tout $k \in \mathbb{N}$, $\mathbf{a}_k \in \mathcal{H}$ est un terme d'erreur pouvant apparaître lors du calcul de l'opérateur proximal, et $\mathbf{b}_k \in \mathcal{H}$ est un terme d'erreur pouvant apparaître lors du calcul du gradient. Il a été démontré dans [Combettes et Wajs, 2005] que, lorsque $\mathbf{h}$ et $\mathbf{g}$ sont des fonctions convexes, si $\sum_{k \in \mathbb{N}} \|\mathbf{a}_k\| < +\infty$ et $\sum_{k \in \mathbb{N}} \|\mathbf{b}_k\| < +\infty$, alors le théorème 2.8 reste valide.

Une autre méthode modélisant un calcul approché de l'opérateur proximal de la fonction $\mathbf{g}$ est donnée dans l'algorithme 2.6.

Algorithme 2.6 Algorithme explicite-implicite inexact (erreurs relatives)

Initialisation : Soient $\mathbf{x}_0 \in \mathbb{R}^N$, $\mathbf{a} \in]\beta, +\infty[$ et $\mathbf{b} \in]0, +\infty[$.

Itérations :

Pour $k = 0, 1, \dots$

$$\left\{ \begin{array}{l} \text{trouver } \mathbf{x}_{k+1} \in \mathbb{R}^N \text{ et } \mathbf{r}_{k+1} \in \partial \mathbf{g}(\mathbf{x}_{k+1}) \text{ tels que} \\ \mathbf{g}(\mathbf{x}_{k+1}) + \langle \mathbf{x}_{k+1} - \mathbf{x}_k, \nabla \mathbf{h}(\mathbf{x}_k) \rangle + \frac{\mathbf{a}}{2} \|\mathbf{x}_{k+1} - \mathbf{x}_k\|^2 \leq \mathbf{g}(\mathbf{x}_k), \\ \|\mathbf{r}_{k+1} + \nabla \mathbf{h}(\mathbf{x}_k)\| \leq \mathbf{b} \|\mathbf{x}_{k+1} - \mathbf{x}_k\|. \end{array} \right.$$

Dans cet algorithme, la première inégalité est appelée *condition de décroissance suffisante*, et la seconde est appelée *condition d'optimalité*. Il a été démontré dans [Attouch et al., 2011] que le théorème 2.9 reste valide pour l'algorithme 2.6, lorsque l'on considère des fonctions $\mathbf{h}$ et $\mathbf{g}$ non nécessairement convexes.

Remarque 2.12. L'algorithme 2.6 peut être interprété comme une version inexacte de l'algorithme 2.4. En effet, d'une part, par définition de l'opérateur proximal, on a

$$(\forall k \in \mathbb{N}) \quad g(\mathbf{x}_{k+1}) + \langle \mathbf{x}_{k+1} - \mathbf{x}_k, \nabla h(\mathbf{x}_k) \rangle + \frac{1}{2\gamma_k} \|\mathbf{x}_{k+1} - \mathbf{x}_k\|^2 \leq g(\mathbf{x}_k), \quad (2.73)$$

ainsi la condition de décroissance suffisante est satisfaite pour $\mathbf{a} = (\sup_{\ell \in \mathbb{N}} \gamma_\ell)^{-1}$ (puisque $\gamma_k^{-1} \geq (\sup_{\ell \in \mathbb{N}} \gamma_\ell)^{-1} > \beta$, où β est la constante de Lipschitz de ∇h). D'autre part, d'après la caractérisation de l'opérateur proximal, il existe $\mathbf{r}_{k+1} \in \partial g(\mathbf{x}_{k+1})$ tel que

$$\mathbf{r}_{k+1} = -\nabla h(\mathbf{x}_k) + \gamma_k^{-1}(\mathbf{x}_k - \mathbf{x}_{k+1}). \quad (2.74)$$

Ainsi, la condition d'optimalité est obtenue pour $\mathbf{b} = (\inf_{\ell \in \mathbb{N}} \gamma_\ell)^{-1}$.

2.4.3 Principe de Majoration-Minimisation

Définitions et algorithme

Revenons au problème d'optimisation (2.65), où h est supposée différentiable. Un autre algorithme permettant d'aborder ce problème est l'*algorithme de Majoration-Minimisation* (MM), qui consiste à trouver une solution du problème (2.65) de façon itérative en minimisant à chaque itération $k \in \mathbb{N}$ une approximation *tangente majorante* de h .

Définition 2.21. Soit E un sous-ensemble non vide de $\mathbb{R}^N$. Soient $h: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ et $\mathbf{y} \in E$. La fonction $q: \mathbb{R}^N \times \mathbb{R}^N \rightarrow]-\infty, +\infty]$ est une approximation tangente majorante de la fonction h en $\mathbf{y}$ sur E si

$$\begin{cases} (\forall \mathbf{x} \in E) & h(\mathbf{x}) \leq q(\mathbf{x}, \mathbf{y}) \\ h(\mathbf{y}) = q(\mathbf{y}, \mathbf{y}). \end{cases} \quad (2.75)$$

L'algorithme général MM proposé dans [Ortega et Rheinboldt, 1970] (voir aussi [Hunter et Lange, 2004]) est donné par l'algorithme 2.7.

Algorithme 2.7 Algorithme de Majoration-Minimisation

Initialisation : Soit $\mathbf{x}_0 \in \mathbb{R}^N$.

Itérations :

Pour $k = 0, 1, \dots$

$$\left| \begin{array}{l} \mathbf{x}_{k+1} \in \underset{\mathbf{x} \in \mathbb{R}^N}{\text{Argmin}} \quad q(\mathbf{x}, \mathbf{x}_k). \end{array} \right.$$

Une illustration graphique de cet algorithme pour $N = 1$ est donnée dans la figure 2.14.

L'avantage de cet algorithme est de pouvoir remplacer un problème d'optimisation qui peut être difficile à résoudre par une suite de problèmes plus simples, à partir du moment où la fonction q est facile à minimiser. Pour s'assurer que ce soit le cas, l'approximation tangente majorante q est souvent choisie quadratique :

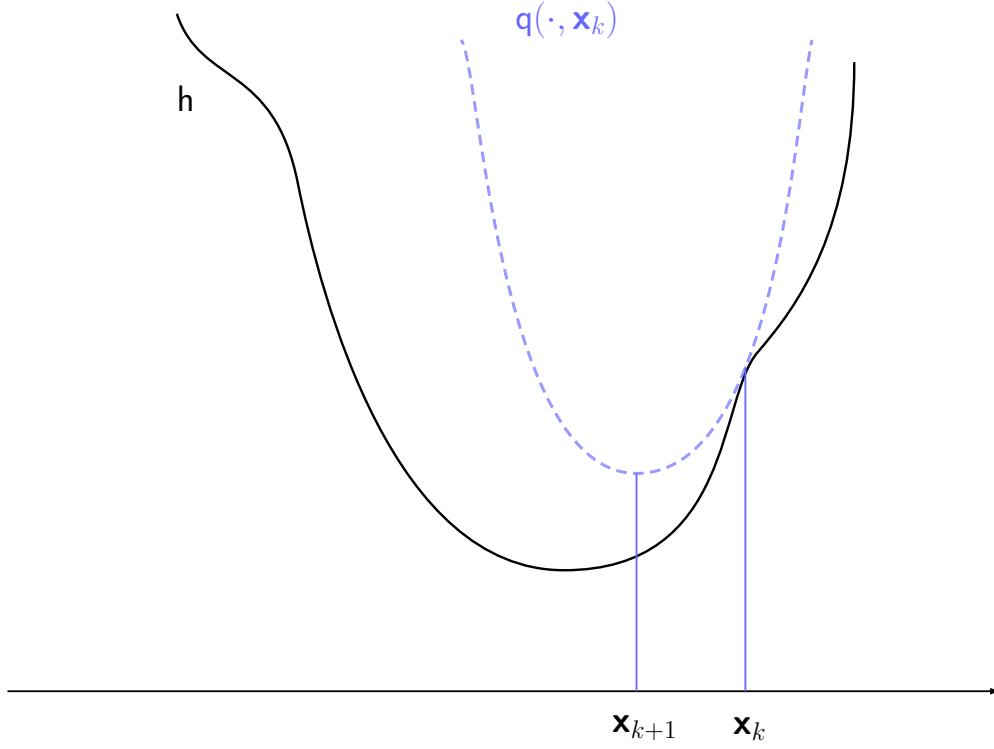


Figure 2.14 – Illustration de l'algorithme de Majoration-Minimisation pour la minimisation d'une fonction $h: \mathbb{R} \rightarrow]-\infty, +\infty]$. À une itération donnée $k \in \mathbb{N}$, on construit une majorante $q(\cdot, \mathbf{x}_k)$ de h au point $\mathbf{x}_k$, puis on définit l'itérée $\mathbf{x}_{k+1}$ comme étant le minimiseur de cette majorante.

Définition 2.22. Soit E un sous-ensemble non vide de $\mathbb{R}^N$. Soient $h: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ et $\mathbf{y} \in E$. La fonction quadratique définie par

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad q(\mathbf{x}, \mathbf{y}) := h(\mathbf{y}) + \langle \mathbf{x} - \mathbf{y}, \nabla h(\mathbf{y}) \rangle + \frac{1}{2} \|\mathbf{x} - \mathbf{y}\|_{\mathbf{A}(\mathbf{y})}^2, \quad (2.76)$$

où $\mathbf{A}(\mathbf{y}) \in \mathcal{S}_N^+$, est une approximation majorante quadratique de h en $\mathbf{y}$ sur E si

$$(\forall \mathbf{x} \in E) \quad h(\mathbf{x}) \leq q(\mathbf{x}, \mathbf{y}). \quad (2.77)$$

Remarquons que l'utilisation d'une majorante quadratique strictement convexe permet d'assurer l'unicité des solutions des sous-problèmes de l'algorithme 2.7. Dans ce cas, la minimisation de $q(\cdot, \mathbf{x}_k)$, pour tout $k \in \mathbb{N}$, dans l'algorithme MM a une forme explicite, donnée dans l'algorithme 2.8.

Algorithme 2.8 Algorithme de Majoration-Minimisation Quadratique

Initialisation : Soit $\mathbf{x}_0 \in \mathbb{R}^N$.

Itérations :

Pour $k = 0, 1, \dots$

$\mathbf{x}_{k+1} = \mathbf{x}_k - \mathbf{A}(\mathbf{x}_k)^{-1} \nabla h(\mathbf{x}_k).$

Lorsque h est de gradient Lipschitz, l'existence de la matrice $\mathbf{A}(\mathbf{y})$ satisfaisant (2.76)-(2.77) est assurée par le lemme suivant :

Lemme 2.1. *Soit E un sous-ensemble convexe non vide de $\mathbb{R}^N$. Soient $h: \mathbb{R}^N \rightarrow \mathbb{R}$ une fonction différentiable de gradient β -Lipschitz sur E , et $\mathbf{A}_k(\mathbf{y}) = \alpha \mathbf{I}_N$, où $\alpha \geq \beta$. La fonction $q: \mathbb{R}^N \times \mathbb{R}^N \rightarrow]-\infty, +\infty]$ définie par (2.76) est une approximation majorante quadratique de h en $\mathbf{y}$ sur E .*

Démonstration. La fonction h étant de gradient β -Lipschitz sur l'ensemble convexe E , on peut appliquer le *Lemme de Descente* [Bertsekas, 1999, Prop.A.24], qui s'exprime par :

$$(\forall (\mathbf{x}, \mathbf{y}) \in E^2) \quad h(\mathbf{x}) \leq h(\mathbf{y}) + \langle \mathbf{x} - \mathbf{y}, \nabla h(\mathbf{y}) \rangle + \frac{\beta}{2} \|\mathbf{x} - \mathbf{y}\|^2.$$

Par conséquent, lorsque $\mathbf{A}_k(\mathbf{y}) = \alpha \mathbf{I}_N$, l'inégalité (2.77) est satisfaite. $\square$

Le lemme précédent permet d'assurer l'existence des matrices $(\mathbf{A}(\mathbf{x}_k))_{k \in \mathbb{N}}$ et de les choisir facilement pour effectuer les itérations de l'algorithme 2.7. Cependant, intuitivement, à chaque itération $k \in \mathbb{N}$, plus la majorante quadratique $q(\cdot, \mathbf{x}_k)$ sera proche de la fonction h dans le voisinage de $\mathbf{x}_k$, meilleure sera l'itérée suivante $\mathbf{x}_{k+1}$. En particulier, lorsque la fonction h est deux fois différentiable, ceci équivaut à choisir $\mathbf{A}(\mathbf{x}_k)$ proche de la matrice hessienne $\nabla^2 h(\mathbf{x}_k)$, pour tout $k \in \mathbb{N}$.

Des techniques utiles de construction de telles matrices sont proposées dans [De Pierro, 1995; Erdogan et Fessler, 1999; Geman et Reynold, 1992; Geman et Yang, 1995; Hunter et Lange, 2004; Idier, 2001; Lange et Fessler, 1995; Sotthivirat et Fessler, 2002; Zheng *et al.*, 2004] pour certaines classes particulières de fonctions h . Nous en proposerons également dans les chapitres 4 et 6.

Par ailleurs, notons qu'il existe des algorithmes de Majoration-Minimisation faisant intervenir des majorantes non nécessairement quadratiques [Bolte et Pauwels, 2014; Fuchs, 2007; Ochs *et al.*, 2015].

Lien avec les algorithmes semi-quadratiques

On s'intéresse au cas où la fonction h est de la forme

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad h(\mathbf{x}) = \frac{1}{2} \|\mathbf{H}\mathbf{x} - \mathbf{z}\|^2 + \sum_{p=1}^P \phi([\mathbf{V}\mathbf{x}]^{(p)}), \quad (2.78)$$

où $\mathbf{z} \in \mathbb{R}^M$, $\mathbf{H} \in \mathbb{R}^{M \times N}$, $\mathbf{V} \in \mathbb{R}^{P \times N}$ et $\phi: \mathbb{R} \rightarrow]-\infty, +\infty]$ est une fonction continument différentiable. On considère les majorantes quadratiques définies dans la proposition suivante :

Proposition 2.14. [Allain et al., 2006; Chan et Mulet, 1999] Soit $\mathbf{y} \in \mathbb{R}^N$. On considère la fonction quadratique

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \mathbf{q}(\mathbf{x}, \mathbf{y}) = h(\mathbf{y}) + \langle \mathbf{x} - \mathbf{y}, \nabla h(\mathbf{y}) \rangle + \frac{1}{2} \|\mathbf{x} - \mathbf{y}\|_{\mathbf{A}(\mathbf{y})}^2.$$

(i) Si ϕ est de gradient Lipschitz, alors $\mathbf{q}(\cdot, \mathbf{y})$ est une fonction majorante de h en $\mathbf{y}$ pour

$$\mathbf{A}(\mathbf{y}) = \mathbf{H}^\top \mathbf{H} + \frac{\mathbf{V}^\top \mathbf{V}}{\alpha},$$

avec $0 < \alpha < \frac{1}{\beta}$.

(ii) Si ϕ est paire, $\phi(\sqrt{\cdot})$ est concave sur $[0, +\infty[$, et, pour tout $t \in \mathbb{R}$, $0 < \dot{\phi}(t)/t < +\infty$, alors $\mathbf{q}(\cdot, \mathbf{y})$ est une fonction majorante de h en $\mathbf{y}$ pour

$$\mathbf{A}(\mathbf{y}) = \mathbf{H}^\top \mathbf{H} + \mathbf{V}^\top \text{Diag}(\omega(\mathbf{V}\mathbf{y})) \mathbf{V},$$

avec, pour tout $t \in \mathbb{R}$, $\omega(t) = \dot{\phi}(t)/t$.

L'algorithme 2.8 peut alors être interprété comme un algorithme semi-quadratique en remplaçant à chaque itération $k \in \mathbb{N}$, la matrice $\mathbf{A}(\mathbf{x}_k)$ par :

- la matrice donnée dans la proposition 2.14(i), on retrouve alors l'algorithme semi-quadratique de Geman & Yang [Geman et Yang, 1995].
- la matrice donnée dans la proposition 2.14(ii), on retrouve alors l'algorithme semi-quadratique de Geman & Reynolds [Geman et Reynold, 1992].

Le résultat de la proposition 2.14(ii) se généralise [Allain et al., 2006; Chouzenoux et al., 2011, 2013] à la minimisation de fonctions de la forme

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad h(\mathbf{x}) = \psi(\mathbf{H}\mathbf{x} - \mathbf{z}) + \|\mathbf{V}_0 \mathbf{x}\|^2 + \sum_{s=1}^S \phi_s(\|\mathbf{V}_s \mathbf{x} - \mathbf{c}_s\|), \quad (2.79)$$

où $\mathbf{z} \in \mathbb{R}^M$, $\mathbf{H} \in \mathbb{R}^{M \times N}$, $\psi: \mathbb{R}^M \rightarrow \mathbb{R}$, $\mathbf{V}_0 \in \mathbb{R}^{P_0 \times N}$, et, pour tout $s \in \{1, \dots, S\}$, $\mathbf{V}_s \in \mathbb{R}^{P_s \times N}$, $\mathbf{c}_s \in \mathbb{R}^{P_s}$, et $\phi_s: \mathbb{R} \rightarrow \mathbb{R}$. On suppose que les assertions suivantes sont satisfaites :

- ψ est différentiable, de gradient β -Lipschitz,
- $(\forall s \in \{1, \dots, S\})$ ϕ_s est différentiable,
- $(\forall s \in \{1, \dots, S\})$ $\phi_s(\sqrt{\cdot})$ est concave sur $[0, +\infty[$,

- $(\forall s \in \{1, \dots, S\})$ il existe $\overline{\omega}_s \in [0, +\infty[$ tel que $(\forall t \in]0, +\infty[) 0 \leq \omega_s(t) \leq \overline{\omega}_s$, où $\omega_s(t) = \phi_s(t)/t$,
- $\lim_{t \rightarrow 0, t \neq 0} \omega_s(t) \in \mathbb{R}$.

Pour tout $\mathbf{y} \in \mathbb{R}^N$, $\mathbf{q}(\cdot, \mathbf{y})$ est alors une fonction majorante quadratique de $\mathbf{h}$ en $\mathbf{y}$ pour

$$\mathbf{A}(\mathbf{y}) = \alpha \mathbf{H}^\top \mathbf{H} + 2\mathbf{V}_0^\top \mathbf{V}_0 + \mathbf{V}^\top \text{Diag}(\mathbf{b}(\mathbf{y})) \mathbf{V}, \quad (2.80)$$

où $\alpha \in [\beta, +\infty[$, $\mathbf{V} = [\mathbf{V}_1 | \dots | \mathbf{V}_S]$ et $\mathbf{b}(\mathbf{y}) = (\mathbf{b}_i(\mathbf{y}))_{1 \leq i \leq SP}$ avec $P = \sum_{s=1}^S P_s$ et $(\forall i \in \{1, \dots, SP\}) \mathbf{b}_i(\mathbf{y}) \in \mathbb{R}$ est défini par

$$(\forall s \in \{1, \dots, S\})(\forall p \in \{1, \dots, P_s\}) \quad \mathbf{b}_{P_1 + \dots + P_{s-1} + p}(\mathbf{y}) = \omega_s(\|\mathbf{V}_s \mathbf{y} - \mathbf{c}_s\|). \quad (2.81)$$

Lien avec l'algorithme explicite-implicite

Comme souligné dans [Beck et Teboulle, 2009; Chaâri *et al.*, 2011], notons qu'il existe un lien étroit entre l'algorithme MM et l'algorithme explicite-implicite.

Considérons l'algorithme 2.4 où les fonctions $\mathbf{h}$ et $\mathbf{g}$ ne sont pas nécessairement convexes, et $\lambda_k \equiv 1$. En utilisant la définition de l'opérateur proximal, on a

$$\begin{aligned} (\forall k \in \mathbb{N}) \quad \mathbf{x}_{k+1} &\in \text{prox}_{\gamma_k \mathbf{g}}(\mathbf{x}_k - \gamma_k \nabla \mathbf{h}(\mathbf{x}_k)) \\ &= \underset{\mathbf{x} \in \mathbb{R}^N}{\text{Argmin}} \quad \mathbf{g}(\mathbf{x}) + \frac{1}{2\gamma_k} \|\mathbf{x} - \mathbf{x}_k + \gamma_k \nabla \mathbf{h}(\mathbf{x}_k)\|^2 \\ &= \underset{\mathbf{x} \in \mathbb{R}^N}{\text{Argmin}} \quad \mathbf{g}(\mathbf{x}) + \frac{1}{2\gamma_k} \left(\|\mathbf{x} - \mathbf{x}_k\|^2 + 2\gamma_k \langle \mathbf{x} - \mathbf{x}_k, \nabla \mathbf{h}(\mathbf{x}_k) \rangle + \|\gamma_k \nabla \mathbf{h}(\mathbf{x}_k)\|^2 \right) \\ &= \underset{\mathbf{x} \in \mathbb{R}^N}{\text{Argmin}} \quad \mathbf{g}(\mathbf{x}) + \underbrace{\mathbf{h}(\mathbf{x}) + \langle \mathbf{x} - \mathbf{x}_k, \nabla \mathbf{h}(\mathbf{x}_k) \rangle + \frac{1}{2\gamma_k} \|\mathbf{x} - \mathbf{x}_k\|^2}_{:= \mathbf{q}_{\gamma_k}(\mathbf{x}, \mathbf{x}_k)}. \end{aligned} \quad (2.82)$$

D'après le théorème 2.9, si

$$(\forall k \in \mathbb{N}) \quad 0 < \inf_{\ell \in \mathbb{N}} \gamma_\ell \leq \gamma_k \leq \sup_{\ell \in \mathbb{N}} \gamma_\ell < \beta^{-1}, \quad (2.83)$$

où β est la constante de Lipschitz de $\mathbf{h}$, alors la suite $(\mathbf{x}_k)_{k \in \mathbb{N}}$ converge vers un point critique de $\mathbf{h} + \mathbf{g}$. Or, sous cette hypothèse, on a, pour tout $k \in \mathbb{N}$, $\gamma_k^{-1} > \beta$, et donc

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \mathbf{q}_{\gamma_k}(\mathbf{x}, \mathbf{x}_k) \geq \mathbf{h}(\mathbf{x}) + \langle \mathbf{x} - \mathbf{x}_k, \nabla \mathbf{h}(\mathbf{x}_k) \rangle + \frac{\beta}{2} \|\mathbf{x} - \mathbf{x}_k\|^2. \quad (2.84)$$

Donc, d'après le lemme 2.1, la fonction $\mathbf{x} \in \mathbb{R}^N \mapsto \mathbf{q}_{\gamma_k}(\mathbf{x}, \mathbf{x}_k)$ majore la fonction $\mathbf{h}$ en $\mathbf{x}_k$:

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \mathbf{q}_{\gamma_k}(\mathbf{x}, \mathbf{x}_k) \geq \mathbf{h}(\mathbf{x}). \quad (2.85)$$

Dans le chapitre 3, nous nous servons de cette interprétation MM de l'algorithme explicite-implicite pour en accélérer la convergence.

2.4.4 Algorithmes praux-duaux

Nous avons vu dans la section 2.2.4 que certaines fonctions de régularisation (par exemple, la variation totale) sont la composition d'une fonction non différentiable et d'un opérateur linéaire. Nous allons donc nous intéresser dans cette partie au problème de minimisation suivant :

$$\text{trouver } \hat{\mathbf{x}} \in \underset{\mathbf{x} \in \mathcal{H}}{\text{Argmin}} \quad f(\mathbf{x}) + h(\mathbf{x}) + g(\mathbf{L}\mathbf{x}), \quad (2.86)$$

où $f \in \Gamma_0(\mathcal{H})$, $g \in \Gamma_0(\mathcal{G})$, $\mathbf{L} \in \mathcal{B}(\mathcal{H}, \mathcal{G})$, $h \in \Gamma_0(\mathcal{H})$ est une fonction différentiable de gradient Lipschitz, $\mathcal{H}$ et $\mathcal{G}$ sont des espaces de Hilbert réels. Un tel problème peut être traité en utilisant, par exemple, l'algorithme explicite-implicite et en calculant l'opérateur proximal de $f + g \circ \mathbf{L}$ de façon approximative (à l'aide de sous-itérations). C'est ce que nous ferons dans la suite de cette thèse lorsque les fonctions ne sont pas nécessairement convexes. Cependant, lorsqu'elles sont convexes, ce problème peut être résolu en traitant à la fois le problème primal (2.86) et le problème dual correspondant [Komodakis et Pesquet, 2014] :

$$\text{trouver } \hat{\mathbf{v}} \in \underset{\mathbf{v} \in \mathcal{G}}{\text{Argmin}} \quad (f^* \square h^*)(-\mathbf{L}^*\mathbf{v}) + g^*(\mathbf{v}). \quad (2.87)$$

Remarquons que les problèmes (2.86)-(2.87) peuvent être reformulés de façon à obtenir un problème de recherche de point selle du Lagrangien [Bauschke et Combettes, 2011, Chap. 19] :

$$\text{trouver } (\hat{\mathbf{x}}, \hat{\mathbf{v}}) \in \underset{\mathbf{x} \in \mathcal{H}}{\text{Argmin}} \left(\underset{\mathbf{v} \in \mathcal{G}}{\text{Argmax}} \quad h(\mathbf{x}) + f(\mathbf{x}) - g^*(\mathbf{v}) + \langle \mathbf{L}\mathbf{x}, \mathbf{v} \rangle \right). \quad (2.88)$$

Ainsi, si le couple $(\hat{\mathbf{x}}, \hat{\mathbf{v}})$ vérifie la condition de Karush-Kuhn-Tucker

$$-\mathbf{L}^*\hat{\mathbf{v}} - \nabla h(\hat{\mathbf{x}}) \in \partial f(\hat{\mathbf{x}}) \quad \text{et} \quad \mathbf{L}\hat{\mathbf{x}} \in \partial g^*(\hat{\mathbf{v}}), \quad (2.89)$$

alors $\hat{\mathbf{x}}$ est une solution du problème primal (2.86), $\hat{\mathbf{v}}$ est une solution du problème dual (2.87), et $(\hat{\mathbf{x}}, \hat{\mathbf{v}})$ est une solution du problème (2.88).

D'autre part, dans la suite de cette section, nous supposerons que

- (i) le problème (2.86) admet au moins une solution ;
- (ii) $\text{ri}(\text{dom } g) \cap \mathbf{L}(\text{dom } f) \neq \emptyset^4$ (ou $\text{dom } g \cap \text{ri}(\mathbf{L}(\text{dom } f)) \neq \emptyset$) sont satisfaites.

Ces conditions permettent d'assurer l'existence d'une solution du problème dual (2.87) [Komodakis et Pesquet, 2014] (notons que d'autres conditions peuvent être considérées, c.f. [Bauschke et Combettes, 2011, Chap. 15]). De plus, l'égalité suivante sera satisfaite

$$f(\hat{\mathbf{x}}) + h(\hat{\mathbf{x}}) + g(\mathbf{L}\hat{\mathbf{x}}) = -\left((f^* \square h^*)(-\mathbf{L}^*\hat{\mathbf{v}}) + g^*(\hat{\mathbf{v}})\right), \quad (2.90)$$

puisqu'il n'y a pas de saut de dualité.

4. $\text{ri}(E)$ désigne l'intérieur relatif de l'ensemble E .

Les algorithmes proximaux primaux-duaux peuvent être décomposés en trois classes :

- des méthodes basées sur l'algorithme explicite-implicite [Chambolle et Pock, 2010; Condat, 2013; Pock et al., 2009; Vü, 2013] : ces méthodes alternent à chaque itération une étape de gradient et une étape proximale.
- des méthodes basées sur l'algorithme explicite-implicite-explicite [Boţ et Hendrich, 2014; Briceños Arias et Combettes, 2011; Combettes, 2013; Combettes et Pesquet, 2011] : en comparaison aux approches précédentes, ces méthodes ajoutent une étape de gradient supplémentaire à chaque itération. Ce sont cependant les premières méthodes primales-duales qui ont été proposées dans la littérature permettant de résoudre (2.86) en exploitant le caractère lisse de h .
- des méthodes basées sur des projections [Alotaibi et al., 2014] : il s'agit de la classe la plus récente, dont le principal avantage, comparativement aux autres méthodes, est de ne pas avoir besoin de connaître la norme de l'opérateur linéaire $\mathbf{L}$.

Dans la suite, nous nous intéresserons à la première classe de ces méthodes.

Algorithme ADMM

Comme indiqué dans [Komodakis et Pesquet, 2014], l'algorithme ADMM (*Alternating Direction Method of Multipliers*) [Boyd et al., 2011; Fortin et Glowinski, 1983; Gabay et Mercier, 1976] peut être vu comme un algorithme primal-dual. Cette méthode permet de résoudre les problèmes (2.86)-(2.87), lorsque $\mathbf{L}^*\mathbf{L}$ est un isomorphisme, grâce aux itérées rappelées dans l'algorithme 2.9.

Algorithme 2.9 Algorithme ADMM

Initialisation : Soient $(\mathbf{y}_0, \mathbf{z}_0) \in \mathcal{G} \times \mathcal{G}$ et $\gamma \in]0, +\infty[$.

Itérations :

Pour $k = 0, 1, \dots$

$$\left[\begin{array}{l} \mathbf{x}_k = \underset{\mathbf{x} \in \mathbb{R}^N}{\operatorname{argmin}} \frac{1}{2} \|\mathbf{L}\mathbf{x} - \mathbf{y}_k + \mathbf{z}_k\|^2 + \frac{1}{\gamma} (f(\mathbf{x}) + h(\mathbf{x})), \\ \mathbf{s}_k = \mathbf{L}\mathbf{x}_k, \\ \mathbf{y}_{k+1} = \operatorname{prox}_{g/\gamma}(\mathbf{z}_k + \mathbf{s}_k), \\ \mathbf{z}_{k+1} = \mathbf{z}_k + \mathbf{s}_k - \mathbf{y}_{k+1}. \end{array} \right.$$

On a alors le résultat de convergence suivant :

Théorème 2.10. [Komodakis et Pesquet, 2014] *Supposons que $\mathbf{L}^*\mathbf{L}$ soit un isomorphisme. La suite $(\mathbf{x}_k, \gamma \mathbf{z}_k)_{k \in \mathbb{N}}$ générée par l'algorithme 2.9 converge alors faiblement vers un couple de solutions $(\hat{\mathbf{x}}, \hat{\mathbf{v}})$ des problèmes (2.86)-(2.87).*

Remarquons que cet algorithme est équivalent à l'algorithme de Douglas-Rachford [Combettes et Pesquet, 2007a; Eckstein et Bertsekas, 1992] lorsque ce dernier est appliqué au problème dual (2.87).

Bien que cet algorithme soit beaucoup utilisé en traitement du signal [Afonso *et al.*, 2011; Figueiredo et Bioucas-Dias, 2010; Figueiredo et Nowak, 2009; Giovannelli et Coulais, 2005; Goldstein et Osher, 2009; Tran-Dinh *et al.*, 2014], en pratique, il peut être difficile à mettre en œuvre. En effet, lorsque la matrice $\mathbf{L}$ est de grande taille et/ou n'a pas une structure simple, le calcul de $\mathbf{x}_k$ à l'itération k peut se révéler coûteux. D'autre part, la différentiabilité de la fonction h n'est pas exploitée de façon explicite dans cette méthode.

Algorithme primal-dual de Condat-Vũ

Une autre méthode permettant de minimiser à la fois le problème primal (2.86) et le problème dual (2.87) est l'algorithme de Condat-Vũ proposé dans [Condat, 2013; Vũ, 2013]. Ses itérées sont données dans l'algorithme 2.10.

Algorithme 2.10 Algorithme primal-dual [Condat, 2013; Vũ, 2013]

Initialisation : Soient $(\mathbf{x}_0, \mathbf{v}_0) \in \mathcal{H} \times \mathcal{G}$, $(\sigma, \tau) \in]0, +\infty[^2$ et, pour tout $k \in \mathbb{N}$, $\lambda_k \in]0, +\infty[$.

Itérations :

Pour $k = 0, 1, \dots$

$$\left\{ \begin{array}{l} \mathbf{y}_k = \text{prox}_{\tau f} \left(\mathbf{x}_k - \tau(\nabla h(\mathbf{x}_k) + \mathbf{b}_k + \mathbf{L}^* \mathbf{v}_k) \right) + \mathbf{a}_k, \\ \mathbf{u}_k = \text{prox}_{\sigma g^*} \left(\mathbf{v}_k + \sigma \mathbf{L}(2\mathbf{y}_k - \mathbf{x}_k) \right) + \mathbf{c}_k, \\ \mathbf{x}_{k+1} = \lambda_k \mathbf{y}_k + (1 - \lambda_k) \mathbf{x}_k, \\ \mathbf{v}_{k+1} = \lambda_k \mathbf{u}_k + (1 - \lambda_k) \mathbf{v}_k. \end{array} \right.$$

Remarquons que, contrairement à l'algorithme 2.9 présenté dans le paragraphe précédent, ici la matrice $\mathbf{L}$ n'a pas besoin de satisfaire de condition particulière.

Dans l'algorithme 2.10, pour tout $k \in \mathbb{N}$, $\mathbf{a}_k$, $\mathbf{b}_k$ et $\mathbf{c}_k$ représentent des termes d'erreurs pouvant apparaître lors des calculs de l'opérateur proximal de τf , du gradient de h et de l'opérateur proximal de σg^* . Remarquons que, lorsque $\mathbf{L} = \mathbf{0}$, et g^* est identiquement nulle, alors on retrouve l'algorithme explicite-implicite vu dans la section 2.4.2.

La convergence de la suite $(\mathbf{x}_k, \mathbf{v}_k)_{k \in \mathbb{N}}$ générée par l'algorithme 2.10 est assurée par le théorème suivant :

Théorème 2.11. [Condat, 2013, Thm. 3.1] *Supposons que les assertions suivantes soient vérifiées :*

- (i) $\tau^{-1} - \sigma \|\mathbf{L}\|^2 > \beta/2$, où β est la constante de Lipschitz de h ,
- (ii) pour tout $k \in \mathbb{N}$, $\lambda_k \in]0, \bar{\lambda}[$, où $\bar{\lambda} = 2 - \beta(\tau^{-1} - \sigma \|\mathbf{L}\|^2)^{-1}/2 \in [1, 2[$,
- (iii) $\sum_{k \in \mathbb{N}} \lambda_k (\bar{\lambda} - \lambda_k) = +\infty$,
- (iv) $\sum_{k \in \mathbb{N}} \lambda_k \|\mathbf{a}_k\| < +\infty$, $\sum_{k \in \mathbb{N}} \lambda_k \|\mathbf{b}_k\| < +\infty$, et $\sum_{k \in \mathbb{N}} \lambda_k \|\mathbf{c}_k\| < +\infty$.

La suite $(\mathbf{x}_k, \mathbf{v}_k)_{k \in \mathbb{N}}$ générée par l'algorithme 2.10 converge alors faiblement vers un couple de solutions $(\tilde{\mathbf{x}}, \tilde{\mathbf{v}})$ des problèmes (2.86)-(2.87).

Notons que le problème primal (2.86) et le problème dual (2.87) ont une structure symétrique. Ainsi, une alternative à l'algorithme 2.10, ayant les mêmes propriétés de convergence que ce dernier, est obtenue en intervertissant les variables primales et duales. Cette méthode est donnée dans l'algorithme 2.11.

Algorithme 2.11 Algorithme primal-dual (forme symétrique) [Condat, 2013; Vũ, 2013]

Initialisation : Soient $(\mathbf{x}_0, \mathbf{v}_0) \in \mathbb{R}^N \times \mathbb{R}^M$, $(\sigma, \tau) \in]0, +\infty[^2$ et, pour tout $k \in \mathbb{N}$, $\lambda_k \in]0, +\infty[$.

Itérations :

Pour $k = 0, 1, \dots$

$$\left\{ \begin{array}{l} \mathbf{u}_k = \text{prox}_{\sigma \mathbf{g}^*}(\mathbf{v}_k + \sigma \mathbf{L} \mathbf{x}_k) + \mathbf{c}_k, \\ \mathbf{y}_k = \text{prox}_{\tau \mathbf{f}}\left(\mathbf{x}_k - \tau(\nabla \mathbf{h}(\mathbf{x}_k) + \mathbf{b}_k + \mathbf{L}^*(2\mathbf{u}_k - \mathbf{v}_k))\right) + \mathbf{a}_k, \\ \mathbf{x}_{k+1} = \lambda_k \mathbf{y}_k + (1 - \lambda_k) \mathbf{x}_k, \\ \mathbf{v}_{k+1} = \lambda_k \mathbf{u}_k + (1 - \lambda_k) \mathbf{v}_k. \end{array} \right.$$

On peut remarquer que l'algorithme présenté dans [Chambolle et Pock, 2010; Pock et al., 2009] est un cas particulier de l'algorithme 2.10. En effet, lorsque $\mathbf{h} \equiv 0$, et les termes d'erreurs $(\mathbf{a}_k)_{k \in \mathbb{N}}$, $(\mathbf{b}_k)_{k \in \mathbb{N}}$ et $(\mathbf{c}_k)_{k \in \mathbb{N}}$ sont tous nuls, l'algorithme 2.10 se réduit à l'algorithme 2.12.

Algorithme 2.12 Algorithme primal-dual simplifié [Chambolle et Pock, 2010]

Initialisation : Soient $(\mathbf{x}_0, \mathbf{v}_0) \in \mathcal{H} \times \mathcal{G}$, $(\sigma, \tau) \in]0, +\infty[^2$ et, pour tout $k \in \mathbb{N}$, $\lambda_k \in]0, +\infty[$.

Itérations :

Pour $k = 0, 1, \dots$

$$\left\{ \begin{array}{l} \mathbf{y}_k = \text{prox}_{\tau \mathbf{f}}(\mathbf{x}_k - \tau \mathbf{L}^* \mathbf{v}_k), \\ \mathbf{u}_k = \text{prox}_{\sigma \mathbf{g}^*}(\mathbf{v}_k + \sigma \mathbf{L}(2\mathbf{y}_k - \mathbf{x}_k)), \\ \mathbf{x}_{k+1} = \lambda_k \mathbf{y}_k + (1 - \lambda_k) \mathbf{x}_k, \\ \mathbf{v}_{k+1} = \lambda_k \mathbf{u}_k + (1 - \lambda_k) \mathbf{v}_k. \end{array} \right.$$

Une version de cet algorithme incluant des termes d'erreur pour le calcul des opérateurs proximaux est donné dans [Condat, 2013]. Notons de plus que, dans ce cas, les conditions de convergences sont moins restrictives que dans le théorème 2.11.

Théorème 2.12. [Condat, 2013, Thm. 3.3] *Supposons que les assertions suivantes soient vérifiées :*

- (i) $\tau\sigma\|\mathbf{L}\|^2 \leq 1$
- (ii) *pour tout $k \in \mathbb{N}$, $\lambda_k \in [\varepsilon, 2 - \varepsilon]$, où $\varepsilon > 0$.*

La suite $(\mathbf{x}_k, \mathbf{v}_k)_{k \in \mathbb{N}}$ générée par l'algorithme 2.12 converge alors faiblement vers un couple de solutions $(\hat{\mathbf{x}}, \hat{\mathbf{v}})$ des problèmes (2.86)-(2.87).

De même, on peut réécrire l'algorithme 2.11 lorsque $\mathbf{h} \equiv 0$, les termes d'erreurs sont nuls et le paramètre de relaxation $\lambda_k \equiv 0$. On obtient alors l'algorithme 2.13, dont la convergence est assurée par le théorème 2.12.

Algorithme 2.13 Algorithme primal-dual simplifié (forme symétrique) [Chambolle et Pock, 2010]

Initialisation : Soient $(\mathbf{x}_0, \mathbf{v}_0) \in \mathcal{H} \times \mathcal{G}$, $(\sigma, \tau) \in]0, +\infty[^2$ et, pour tout $k \in \mathbb{N}$, $\gamma_k \in]0, +\infty[$.

Itérations :

Pour $k = 0, 1, \dots$

$$\begin{cases} \mathbf{v}_{k+1} = \text{prox}_{\sigma \mathbf{g}^*}(\mathbf{v}_k + \sigma \mathbf{L} \mathbf{x}_k), \\ \mathbf{x}_{k+1} = \text{prox}_{\tau \mathbf{f}}(\mathbf{x}_k - \tau \mathbf{L}^*(2\mathbf{v}_{k+1} - \mathbf{v}_k)). \end{cases}$$

En utilisant le théorème de décomposition de Moreau, en effectuant le changement de variable ($\forall k \in \mathbb{N}$) $\mathbf{p}_k = \mathbf{v}_k/\sigma$ puis en réarrangeant les mises à jours des variables, on obtient alors l'algorithme 2.14 [Burger et al., 2014; Esser et al., 2010] lorsque $\theta = 1$ (lorsque $0 < \theta < 1$ on obtient une sous-relaxation de la variable $(\mathbf{z}_k)_{k \in \mathbb{N}}$ [Burger et al., 2014]).

Algorithme 2.14 Algorithme primal-dual simplifié (forme symétrique modifiée) [Chambolle et Pock, 2010]

Initialisation : Soient $(\mathbf{x}_0, \mathbf{p}_0) \in \mathcal{H} \times \mathcal{G}$, $\mathbf{z}_0 = \mathbf{p}_0$, $(\tau, \sigma) \in]0, +\infty[^2$ et $\theta \in]0, 1]$.

Itérations :

Pour $k = 0, 1, \dots$

$$\begin{cases} \mathbf{x}_{k+1} = \text{prox}_{\tau \mathbf{f}}(\mathbf{x}_k - \sigma \tau \mathbf{L}^* \mathbf{z}_k), \\ \mathbf{y}_{k+1} = \text{prox}_{\mathbf{g}/\sigma}(\mathbf{p}_k + \mathbf{L} \mathbf{x}_{k+1}), \\ \mathbf{p}_{k+1} = \mathbf{p}_k + \mathbf{L} \mathbf{x}_{k+1} - \mathbf{y}_{k+1}, \\ \mathbf{z}_{k+1} = \mathbf{p}_{k+1} + \theta(\mathbf{p}_{k+1} - \mathbf{p}_k). \end{cases}$$

On a alors le résultat de convergence suivant :

Corollaire 2.1. [Burger et al., 2014] Supposons que $\tau\sigma\|\mathbf{L}\|^2 < 1$. La suite $(\mathbf{x}_k, \theta\mathbf{p}_k)_{k \in \mathbb{N}}$ générée par l'algorithme 2.14 converge alors faiblement vers un couple de solutions $(\hat{\mathbf{x}}, \hat{\mathbf{v}})$ des problèmes (2.86)-(2.87).

Nous utiliserons ce dernier algorithme dans le chapitre 8.

Par ailleurs, dans le chapitre 7, nous généraliserons ces algorithmes primaux-duaux au cadre stochastique. Nous revisiterons les conditions de convergence dans ce cadre plus général. Cela nous permettra de traiter une classe de problèmes plus large que (2.86)-(2.87).

2.5 CONCLUSION

Une méthode rigoureuse pour trouver une solution d'un problème inverse, est de la définir comme le minimiseur d'un problème variationnel faisant intervenir une fonction d'attache aux données et une fonction de régularisation. Comme nous l'avons vu dans ce chapitre, les fonctions à minimiser peuvent présenter des propriétés mathématiques différentes. Cela nécessite la construction d'algorithmes permettant d'exploiter ces différentes propriétés. En particulier, nous avons donné dans la section précédente des algorithmes usuels permettant de minimiser des fonctions (non) différentiables et/ou (non) convexes.

Deux aspects importants sont à prendre en compte lors de l'utilisation de ces algorithmes pour des applications du traitement du signal et des images : leur vitesse de convergence et leur coût d'implémentation. L'objectif de cette thèse sera de proposer des algorithmes à la fois peu coûteux à implémenter et dont la convergence est assurée et rapide. Ainsi dans un premier temps nous proposerons dans le chapitre suivant une méthode d'accélération de l'algorithme explicite-implicite basée sur l'utilisation d'une métrique variable qui sera choisie grâce au principe MM.

2.A DÉMONSTRATION DE LA PROPRIÉTÉ (V) DE PERTURBATION QUADRATIQUE DE L'OPÉRATEUR PROXIMAL

Soient $\mathbf{x} \in \mathcal{H}$ et $\mathbf{p} = \text{prox}_{\mathbf{A}+\mathbf{U}, \mathbf{g}}((\mathbf{A} + \mathbf{U})^{-1}(\mathbf{U}\mathbf{x} - \beta\mathbf{x}_0))$. D'après la définition 2.12 de l'opérateur proximal, on a

$$\mathbf{p} = \underset{\mathbf{y} \in \mathcal{H}}{\text{Argmin}} \quad \mathbf{g}(\mathbf{y}) + \frac{1}{2}\mathbf{q}(\mathbf{y}),$$

où

$$\begin{aligned} \mathbf{q}(\mathbf{y}) &= \|\mathbf{y} - (\mathbf{A} + \mathbf{U})^{-1}(\mathbf{U}\mathbf{x} - \beta\mathbf{x}_0)\|_{\mathbf{A}+\mathbf{U}}^2 \\ &= \left\langle \mathbf{y} - (\mathbf{A} + \mathbf{U})^{-1}(\mathbf{U}\mathbf{x} - \beta\mathbf{x}_0) \left| (\mathbf{A} + \mathbf{U}) \left(\mathbf{y} - (\mathbf{A} + \mathbf{U})^{-1}(\mathbf{U}\mathbf{x} - \beta\mathbf{x}_0) \right) \right. \right\rangle. \end{aligned}$$

Afin de démontrer l'égalité (2.46), il suffit de montrer que

$$(\forall \mathbf{y} \in \mathcal{H}) \quad \mathbf{g}(\mathbf{y}) + \frac{1}{2}\mathbf{q}(\mathbf{y}) = \mathbf{f}(\mathbf{y}) + \frac{1}{2}\|\mathbf{y} - \mathbf{x}\|_{\mathbf{U}}^2 + \kappa, \quad (2.91)$$

avec $\kappa \in \mathbb{R}$. Pour tout $\mathbf{y} \in \mathcal{H}$, on a

$$\begin{aligned} \mathbf{q}(\mathbf{y}) &= \left\langle \mathbf{y} - (\mathbf{A} + \mathbf{U})^{-1}\mathbf{U}\mathbf{x} + (\mathbf{A} + \mathbf{U})^{-1}\beta\mathbf{x}_0 \mid (\mathbf{A} + \mathbf{U})\mathbf{y} - \mathbf{U}\mathbf{x} + \beta\mathbf{x}_0 \right\rangle \\ &= \left\langle \mathbf{y} - (\mathbf{A} + \mathbf{U})^{-1}\mathbf{U}\mathbf{x} \mid (\mathbf{A} + \mathbf{U})\mathbf{y} - \mathbf{U}\mathbf{x} \right\rangle \\ &\quad + 2\beta \left\langle (\mathbf{A} + \mathbf{U})^{-1}\mathbf{x}_0 \mid (\mathbf{A} + \mathbf{U})\mathbf{y} - \mathbf{U}\mathbf{x} \right\rangle + \beta^2 \left\langle (\mathbf{A} + \mathbf{U})^{-1}\mathbf{x}_0 \mid \mathbf{x}_0 \right\rangle \\ &= \left\langle \mathbf{y} - (\mathbf{A} + \mathbf{U})^{-1}\mathbf{U}\mathbf{x} \mid (\mathbf{A} + \mathbf{U})\mathbf{y} - \mathbf{U}\mathbf{x} \right\rangle + 2\beta \langle \mathbf{y} \mid \mathbf{x}_0 \rangle + \kappa_1, \end{aligned} \quad (2.92)$$

où $\kappa_1 = -2\beta \langle (\mathbf{A} + \mathbf{U})^{-1}\mathbf{x}_0 \mid \mathbf{U}\mathbf{x} \rangle + \beta^2 \langle (\mathbf{A} + \mathbf{U})^{-1}\mathbf{x}_0 \mid \mathbf{x}_0 \rangle$ est une constante par rapport à la variable $\mathbf{y}$. Remarquons que

$$\begin{aligned} &\left\langle \mathbf{y} - (\mathbf{A} + \mathbf{U})^{-1}\mathbf{U}\mathbf{x} \mid (\mathbf{A} + \mathbf{U})\mathbf{y} - \mathbf{U}\mathbf{x} \right\rangle \\ &= \left\langle (\mathbf{y} - \mathbf{x}) + \left(\mathbf{x} - (\mathbf{A} + \mathbf{U})^{-1}\mathbf{U}\mathbf{x} \right) \mid \mathbf{U}(\mathbf{y} - \mathbf{x}) + \mathbf{A}\mathbf{y} \right\rangle \\ &= \|\mathbf{y} - \mathbf{x}\|_{\mathbf{U}}^2 + \|\mathbf{y}\|_{\mathbf{A}}^2 - \langle \mathbf{x} \mid \mathbf{A}\mathbf{y} \rangle + \left\langle \mathbf{x} - (\mathbf{A} + \mathbf{U})^{-1}\mathbf{U}\mathbf{x} \mid (\mathbf{A} + \mathbf{U})\mathbf{y} \right\rangle + \kappa_2, \end{aligned} \quad (2.93)$$

où $\kappa_2 = \langle (\mathbf{A} + \mathbf{U})^{-1}\mathbf{U}\mathbf{x} - \mathbf{x} \mid \mathbf{U}\mathbf{x} \rangle$ est indépendant de $\mathbf{y}$. Remarquons de plus que

$$\begin{aligned} \left\langle \mathbf{x} - (\mathbf{A} + \mathbf{U})^{-1}\mathbf{U}\mathbf{x} \mid (\mathbf{A} + \mathbf{U})\mathbf{y} \right\rangle &= \langle \mathbf{x} \mid \mathbf{A}\mathbf{y} + \mathbf{U}\mathbf{y} \rangle - \langle \mathbf{U}\mathbf{x} \mid \mathbf{y} \rangle \\ &= \langle \mathbf{x} \mid \mathbf{A}\mathbf{y} \rangle. \end{aligned} \quad (2.94)$$

On peut alors déduire (2.91) en combinant les équations (2.92)-(2.94) et en posant $\kappa = \frac{1}{2}(\kappa_1 + \kappa_2) - \gamma$.

CHAPITRE 3

Algorithme explicite-implicite à métrique variable

Sommaire

3.1	Introduction	60
3.2	Problème de Minimisation	61
3.3	Méthode proposée	62
3.3.1	Algorithme explicite-implicite à métrique variable	62
3.3.2	Choix de la métrique	63
3.3.3	Algorithme inexact à métrique variable	64
3.4	Analyse de convergence	65
3.4.1	Propriétés de décroissance	65
3.4.2	Résultat de convergence	68
3.5	Conclusion	70
Annexe 3.A	Démonstration du théorème 3.1	71

3.1 INTRODUCTION

Dans ce chapitre, nous considérerons un problème de minimisation faisant intervenir une somme de deux fonctions dont l'une est différentiable de gradient Lipschitz, tandis que l'autre est convexe mais éventuellement non lisse. Dans ce cadre, une approche usuelle consiste à utiliser l'algorithme proximal explicite-implicite 2.4 [Chen et Rockafellar, 1997; Tseng, 1998]. Cette méthode alterne, à chaque itération, une étape de gradient sur la fonction différentiable et une étape proximale sur la seconde fonction. Récemment, la convergence de cet algorithme a été démontrée lorsque les fonctions à minimiser ne sont pas nécessairement convexes mais satisfont l'inégalité de KL (théorème 2.9).

Dans le contexte de problèmes de grande taille, tels que ceux rencontrés en restauration d'images, un des enjeux principaux est de pouvoir produire de bons résultats numériques en un temps de calcul raisonnable. L'algorithme explicite-implicite est caractérisé par un faible coût de calcul à chaque itération. Cependant, comme la plupart des méthodes de minimisation de premier ordre, il peut converger relativement lentement [Chen et Rockafellar, 1997]. Deux grandes familles de stratégies d'accélération de cet algorithme se distinguent dans la littérature. La première repose sur une accélération de sous-espaces [Beck et Teboulle, 2009; Bioucas-Dias et Figueiredo, 2007; Kowalski, 2010; Nesterov, 2007; Ochs *et al.*, 2014] : c'est-à-dire que la vitesse de convergence va être accélérée en utilisant les informations obtenues lors des itérations précédentes pour construire la nouvelle estimée. La seconde stratégie d'accélération est basée sur l'introduction d'une métrique variant au cours des itérations [Becker et Fadili, 2012; Bonnans *et al.*, 1995; Burke et Qian, 1999; Chen et Rockafellar, 1997; Combettes et Vũ, 2014; Lotito *et al.*, 2009; Parente *et al.*, 2008]. Cette métrique est obtenue grâce à une matrice de préconditionnement. L'algorithme que l'on obtient alors est appelé algorithme explicite-implicite à métrique variable (ou *Variable Metric Forward-Backward* (VMFB)). La convergence de cet algorithme dans le cas où les fonctions à minimiser sont convexes a été démontrée dans [Combettes et Vũ, 2014; Lee *et al.*, 2012; Tran-Dinh *et al.*, 2014].

Dans ce chapitre, nous allons nous poser la question de l'extension de ces résultats de convergence au cadre non convexe. La démonstration de ces résultats reposera d'une part sur l'inégalité de KL, et d'autre part sur un choix original des matrices de préconditionnement basé sur une stratégie de Majoration-Minimisation. Enfin, comme nous l'avons vu dans le chapitre 2, l'opérateur proximal n'est pas toujours calculable de façon explicite. Nous rechercherons un algorithme robuste aux erreurs pouvant apparaître lors du calcul de l'opérateur proximal.

Le plan de ce chapitre est le suivant : nous allons tout d'abord présenter dans la section 3.2 le problème de minimisation étudié. Puis, dans la section 3.3, nous proposerons une version exacte et inexacte de l'algorithme VMFB. Enfin, nous donnerons les résultats de convergence ainsi que leurs démonstrations dans la section 3.4. Dans le chapitre 4, nous présenterons une application de cette méthode pour la résolution d'un problème inverse en traitement d'images.

3.2 PROBLÈME DE MINIMISATION

Nous nous intéressons au problème de minimisation suivant :

$$\text{trouver } \hat{\mathbf{x}} \in \underset{\mathbf{x} \in \mathbb{R}^N}{\text{Argmin}} f(\mathbf{x}), \quad (3.1)$$

où $f : \mathbb{R}^N \rightarrow]-\infty, +\infty]$ est une fonction coercive. De plus, nous supposons que f peut s'écrire de la façon suivante :

$$f = h + g, \quad (3.2)$$

où h est une fonction différentiable et g est une fonction propre et s.c.i. Un algorithme permettant de résoudre ce problème est l'algorithme explicite-implicite donné par l'algorithme 2.4. Ses propriétés de convergence ont été données dans le chapitre 2.

Dans la suite de ce chapitre, nous nous intéresserons au problème de minimisation de la forme (3.1)-(3.2) où les fonctions h et g satisfont les hypothèses suivantes :

Hypothèse 3.1.

- (i) $g : \mathbb{R}^N \rightarrow]-\infty, +\infty]$ est une fonction convexe, propre, semi-continue inférieurement et dont la restriction à son domaine $\text{dom } g$ est continue.
- (ii) $h : \mathbb{R}^N \rightarrow \mathbb{R}$ est une fonction différentiable sur $\mathbb{R}^N$ de gradient β -Lipschitz sur $\text{dom } g$, où $\beta > 0$.
- (iii) f , définie par (3.2), est coercive.

Les remarques suivantes concernant l'hypothèse 3.1 seront utiles par la suite.

Remarque 3.1.

- (i) L'hypothèse 3.1(ii) sur h est moins contraignante que l'hypothèse de gradient Lipschitz adoptée en général pour démontrer la convergence de l'algorithme explicite-implicite [Attouch et al., 2011; Combettes et Wajs, 2005]. En particulier, si le domaine de g est compact et h est deux fois continument différentiable, alors l'hypothèse 3.1(ii) est satisfaite.
- (ii) D'après l'hypothèse 3.1(ii), $\text{dom } g \subset \text{dom } h$. Par conséquent, en utilisant l'hypothèse 3.1(i), $\text{dom } f = \text{dom } g$ est non vide et convexe.
- (iii) D'après l'hypothèse 3.1, f est une fonction propre et semi-continue inférieurement, dont la restriction à son domaine est continue. Ainsi, par coercivité de f , pour tout $\mathbf{x} \in \text{dom } g$, $\text{lev}_{\leq f(\mathbf{x})} f$ est un ensemble compact. De plus, l'ensemble des minimiseurs de f est non-vide et compact.

3.3 MÉTHODE PROPOSÉE

3.3.1 Algorithme explicite-implicite à métrique variable

L'algorithme explicite-implicite à métrique variable a été introduit dans [Becker et Fadili, 2012; Bonnans *et al.*, 1995; Burke et Qian, 1999; Chen et Rockafellar, 1997; Combettes et Vũ, 2014; Lotito *et al.*, 2009; Ochs *et al.*, 2014; Parente *et al.*, 2008] dans le but d'accélérer la convergence de l'algorithme (2.4).

Algorithme 3.1 Algorithme explicite-implicite à métrique variable

Initialisation : Soient $\mathbf{x}_0 \in \text{dom } \mathbf{g}$ et $(\forall k \in \mathbb{N}) (\gamma_k, \lambda_k) \in]0, +\infty[^2$ et $\mathbf{A}_k(\mathbf{x}_k) \in \mathcal{S}_N^+$.

Itérations :

Pour $k = 0, 1, \dots$

$$\left\{ \begin{array}{l} \mathbf{y}_k = \text{prox}_{\gamma_k^{-1} \mathbf{A}_k(\mathbf{x}_k), \mathbf{g}}(\mathbf{x}_k - \gamma_k \mathbf{A}_k(\mathbf{x}_k)^{-1} \nabla \mathbf{h}(\mathbf{x}_k)), \\ \mathbf{x}_{k+1} = \mathbf{x}_k + \lambda_k (\mathbf{y}_k - \mathbf{x}_k). \end{array} \right.$$

Remarque 3.2.

D'après l'hypothèse 3.1(i), $\text{dom } \mathbf{g}$ est convexe. Donc, pour tout $k \in \mathbb{N}$, si $\lambda_k \in [0, 1]$, $\mathbf{x}_k$ et $\mathbf{y}_k$ sont dans $\text{dom } \mathbf{g}$.

Certains liens avec des méthodes existantes peuvent être établies concernant l'algorithme 3.1. D'une part, lorsque $\mathbf{A}_k(\mathbf{x}_k) \equiv \mathbf{I}_N$, l'algorithme 3.1 est équivalent à l'algorithme explicite-implicite usuel donné par l'algorithme 2.4. D'autre part, lorsque $\mathbf{g} \equiv 0$, l'algorithme 3.1 correspond à un algorithme de gradient préconditionné.

La convergence de l'algorithme 3.1 a été étudiée dans divers travaux. Le résultat le plus général est donné dans [Combettes et Vũ, 2014] : les auteurs ont démontré la convergence de la suite des itérés $(\mathbf{x}_k)_{k \in \mathbb{N}}$ vers une solution du problème (3.1)-(3.2) lorsque $\mathbf{h}$ et $\mathbf{g}$ sont convexes et qu'il existe une suite $(\eta_k)_{k \in \mathbb{N}}$ positive bornée telle que, pour tout $k \in \mathbb{N}$,

$$(1 + \eta_k) \mathbf{A}_{k+1}(\mathbf{x}_{k+1}) \succcurlyeq \mathbf{A}_k(\mathbf{x}_k). \quad (3.3)$$

D'autres résultats de convergences existent dans la littérature dans le cas particulier où $\mathbf{g}$ est une fonction indicatrice sur un ensemble convexe [Bertsekas, 1982; Birgin *et al.*, 2000; Bonettini *et al.*, 2009] ou lorsque $\mathbf{h}$ est localement fortement convexe avec un hessien Lipschitz continu et $\mathbf{g}$ est convexe [Lee *et al.*, 2012]. Récemment la convergence de l'algorithme 3.1 a aussi été établie dans [Tran-Dinh *et al.*, 2014] lorsque $\mathbf{g}$ est convexe, $\mathbf{h}$ est une fonction lisse auto-concordante, et, pour tout $k \in \mathbb{N}$, $\mathbf{A}_k(\mathbf{x}_k)$ est soit le hessien de $\mathbf{h}$ en $\mathbf{x}_k$ (ou une version approchée de celui-ci), soit une matrice diagonale. Cependant, dans aucun des travaux précédents, l'étude de la convergence de l'algorithme 3.1 n'est menée dans le cas où la fonction $\mathbf{h}$ est non-convexe. C'est ce que nous nous proposons d'étudier.

3.3.2 Choix de la métrique

L'algorithme 3.1 que nous étudions dans ce chapitre est accéléré grâce à l'introduction d'une métrique variable induite par les matrices $(\mathbf{A}_k(\mathbf{x}_k))_{k \in \mathbb{N}}$. L'une des contributions de ce chapitre est la proposition d'une méthode constructive pour choisir les matrices permettant de définir cette métrique. Pour cela, nous allons nous baser sur une stratégie de type MM. Plus précisément, considérons une suite $(\mathbf{x}_k)_{k \in \mathbb{N}}$ de $\text{dom } \mathbf{g}$. Nous définissons alors la suite de matrices $(\mathbf{A}_k(\mathbf{x}_k))_{k \in \mathbb{N}}$ dans $\mathcal{S}_N^+$ satisfaisant l'hypothèse suivante :

Hypothèse 3.2.

(i) Pour tout $k \in \mathbb{N}$, la fonction quadratique définie par

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \mathbf{q}(\mathbf{x}, \mathbf{x}_k) := h(\mathbf{x}_k) + \langle \mathbf{x} - \mathbf{x}_k, \nabla h(\mathbf{x}_k) \rangle + \frac{1}{2} \|\mathbf{x} - \mathbf{x}_k\|_{\mathbf{A}_k(\mathbf{x}_k)}^2,$$

est une majorante de h en $\mathbf{x}_k$ sur $\text{dom } \mathbf{g}$.

(ii) Il existe $(\underline{\nu}, \overline{\nu}) \in]0, +\infty[^2$ tel que

$$(\forall k \in \mathbb{N}) \quad \underline{\nu} \mathbf{I}_N \preccurlyeq \mathbf{A}_k(\mathbf{x}_k) \preccurlyeq \overline{\nu} \mathbf{I}_N.$$

Cette hypothèse permet de faire le lien entre l'algorithme 3.1 et le principe MM. En effet, remarquons que, d'après l'algorithme 3.1 et l'hypothèse 3.1(i), lorsque $(\forall k \in \mathbb{N}) \gamma_k = 1$, on a

$$\begin{aligned} \mathbf{y}_k &= \text{prox}_{\mathbf{A}_k(\mathbf{x}_k), \mathbf{g}}(\mathbf{x}_k - \mathbf{A}_k(\mathbf{x}_k)^{-1} \nabla h(\mathbf{x}_k)) \\ &= \underset{\mathbf{x} \in \mathbb{R}^N}{\text{argmin}} \quad \mathbf{g}(\mathbf{x}) + \frac{1}{2} \|\mathbf{x} - \mathbf{x}_k + \mathbf{A}_k(\mathbf{x}_k)^{-1} \nabla h(\mathbf{x}_k)\|_{\mathbf{A}_k(\mathbf{x}_k)}^2 \\ &= \underset{\mathbf{x} \in \mathbb{R}^N}{\text{argmin}} \quad \mathbf{g}(\mathbf{x}) + \mathbf{q}(\mathbf{x}, \mathbf{x}_k), \end{aligned}$$

où $\mathbf{q}(\cdot, \mathbf{x}_k)$ est la majorante quadratique de h définie par l'hypothèse 3.2(i).

Remarque 3.3. D'après le lemme 2.1 du chapitre 2, l'existence de telles matrices $(\mathbf{A}_k(\mathbf{x}_k))_{k \in \mathbb{N}}$ est assurée. En effet, d'après l'hypothèse 3.1, $\text{dom } \mathbf{g}$ est convexe et h est de gradient β -Lipschitz sur $\text{dom } \mathbf{g}$. Par conséquent, si, pour tout $k \in \mathbb{N}$, $\mathbf{A}_k(\mathbf{x}_k) = \beta \mathbf{I}_N$, l'hypothèse 3.2 est satisfaite avec $\underline{\nu} = \overline{\nu} = \beta$.

Bien que la remarque précédente permette de choisir facilement les matrices $(\mathbf{A}_k(\mathbf{x}_k))_{k \in \mathbb{N}}$, comme il l'a été précisé dans [Becker et Fadili, 2012; Lee et al., 2012; Tran-Dinh et al., 2014] (voir aussi le chapitre 2), l'obtention de bonnes performances de l'algorithme 3.1 dépend fortement de l'exactitude de l'approximation de second ordre caractérisant la métrique. Plus précisément, il faut construire $(\mathbf{A}_k(\mathbf{x}_k))_{k \in \mathbb{N}}$ de telle façon que, pour tout $k \in \mathbb{N}$, la fonction quadratique $\mathbf{q}(\cdot, \mathbf{x}_k)$ soit aussi proche que possible de h au voisinage de $\mathbf{x}_k$, tout en satisfaisant l'hypothèse 3.2. Comme nous l'avons remarqué dans la section 2.4.3 du chapitre 2, il existe des techniques de construction de telles métriques dans la littérature pour certaines classes particulières de fonctions h [Erdogan et Fessler, 1999; Geman et Reynold, 1992; Geman et Yang, 1995; Hunter et Lange, 2004].

3.3.3 Algorithme inexact à métrique variable

De manière générale, l'opérateur proximal relatif à une métrique quelconque n'a pas d'expression simple explicite. Afin de pallier cette difficulté, nous proposons de résoudre le problème (3.1) grâce à une version inexacte de l'algorithme explicite-implicite à métrique variable, s'inspirant de l'algorithme inexact explicite-implicite donné par l'algorithme 2.6, proposé dans [Attouch *et al.*, 2011, Sec. 5] :

Algorithme 3.2 Algorithme inexact explicite-implicite à métrique variable

Initialisation : Soient $\mathbf{x}_0 \in \text{dom } \mathbf{g}$, $\alpha \in]0, +\infty[$ et $(\forall k \in \mathbb{N}) (\gamma_k, \lambda_k) \in]0, +\infty[^2$ et $\mathbf{A}_k(\mathbf{x}_k) \in \mathcal{S}_N^+$.

Itérations :

Pour $k = 0, 1, \dots$

$\left[\begin{array}{l} \text{trouver } \mathbf{y}_k \in \mathbb{R}^N \text{ et } \mathbf{r}_k \in \partial \mathbf{g}(\mathbf{y}_k) \text{ tels que} \\ \mathbf{g}(\mathbf{y}_k) + \langle \mathbf{y}_k - \mathbf{x}_k, \nabla h(\mathbf{x}_k) \rangle + \gamma_k^{-1} \ \mathbf{y}_k - \mathbf{x}_k\ _{\mathbf{A}_k(\mathbf{x}_k)}^2 \leq \mathbf{g}(\mathbf{x}_k), \\ \ \nabla h(\mathbf{x}_k) + \mathbf{r}_k\ \leq \alpha \ \mathbf{y}_k - \mathbf{x}_k\ _{\mathbf{A}_k(\mathbf{x}_k)}, \\ \mathbf{x}_{k+1} = \mathbf{x}_k + \lambda_k(\mathbf{y}_k - \mathbf{x}_k). \end{array} \right.$

De même que pour l'algorithme 2.6, dans l'algorithme 3.2, la première inégalité correspond à une condition de décroissance suffisante, et la seconde peut être interprétée comme une condition inexacte d'optimalité.

Nous supposons que le pas $(\gamma_k)_{k \in \mathbb{N}}$ et le paramètre de relaxation $(\lambda_k)_{k \in \mathbb{N}}$ présents dans les algorithmes 3.1 et 3.2 satisfont les deux hypothèses suivante :

Hypothèse 3.3.

- (i) Il existe $\underline{\lambda} \in]0, +\infty[$ telle que, pour tout $k \in \mathbb{N}$, $\underline{\lambda} \leq \lambda_k \leq 1$.
- (ii) Il existe $(\underline{\eta}, \overline{\eta}) \in]0, +\infty[^2$ tel que, pour tout $k \in \mathbb{N}$, $\underline{\eta} \leq \lambda_k \gamma_k \leq 2 - \overline{\eta}$.

Hypothèse 3.4.

Il existe $\underline{\tau} \in]0, 1]$ telle que, pour tout $k \in \mathbb{N}$, le paramètre de relaxation λ_k satisfait

$$\mathbf{f}((1 - \lambda_k)\mathbf{x}_k + \lambda_k\mathbf{y}_k) \leq (1 - \underline{\tau})\mathbf{f}(\mathbf{x}_k) + \underline{\tau}\mathbf{f}(\mathbf{y}_k).$$

Remarque 3.4.

- (i) En posant $\underline{\tau} = 1$, on retrouve la condition de décroissance usuelle :

$$(\forall k \in \mathbb{N}) \quad \mathbf{f}((1 - \lambda_k)\mathbf{x}_k + \lambda_k\mathbf{y}_k) \leq \mathbf{f}(\mathbf{y}_k). \quad (3.4)$$

- (ii) Dans la section suivante, nous montrerons que l'hypothèse 3.4 peut être simplifiée dans le cas où $\mathbf{f}$ est une fonction convexe. Cependant, notons que même dans le cas général, il existe au moins une suite $(\lambda_k)_{k \in \mathbb{N}}$ satisfaisant l'hypothèse 3.4 (il suffit de prendre, pour tout $k \in \mathbb{N}$, $\lambda_k = \underline{\tau} = 1$).

Remarque 3.5.

Comme nous l'avons dit précédemment, sous condition que les hypothèses 3.1, 3.2, 3.3 et 3.4 soient satisfaites, l'algorithme 3.2 peut être interprété comme étant une version inexacte de l'algorithme 3.1. Soient $(\mathbf{x}_k)_{k \in \mathbb{N}}$ et $(\mathbf{y}_k)_{k \in \mathbb{N}}$ des suites générées par l'algorithme 3.1. Soit $k \in \mathbb{N}$, en utilisant à la fois la caractérisation variationnelle de l'opérateur proximal et la convexité de $\mathbf{g}$, il existe $\mathbf{r}_k \in \partial \mathbf{g}(\mathbf{y}_k)$ tel que

$$\begin{cases} \mathbf{r}_k = -\nabla \mathbf{h}(\mathbf{x}_k) + \gamma_k^{-1} \mathbf{A}_k(\mathbf{x}_k)(\mathbf{x}_k - \mathbf{y}_k), \\ \langle \mathbf{y}_k - \mathbf{x}_k, \mathbf{r}_k \rangle \geq \mathbf{g}(\mathbf{y}_k) - \mathbf{g}(\mathbf{x}_k), \end{cases} \quad (3.5)$$

ce qui conduit à

$$\mathbf{g}(\mathbf{y}_k) + \langle \mathbf{y}_k - \mathbf{x}_k, \nabla \mathbf{h}(\mathbf{x}_k) \rangle + \gamma_k^{-1} \|\mathbf{y}_k - \mathbf{x}_k\|_{\mathbf{A}_k(\mathbf{x}_k)}^2 \leq \mathbf{g}(\mathbf{x}_k). \quad (3.6)$$

On obtient donc la condition de décroissance de l'algorithme 3.2.

D'autre part, d'après les hypothèses 3.2(ii) et 3.3, on a

$$\|\nabla \mathbf{h}(\mathbf{x}_k) + \mathbf{r}_k\| = \gamma_k^{-1} \|\mathbf{A}_k(\mathbf{x}_k)(\mathbf{x}_k - \mathbf{y}_k)\| \leq \underline{\gamma}^{-1} \overline{\nu}^{1/2} \|\mathbf{x}_k - \mathbf{y}_k\|_{\mathbf{A}_k(\mathbf{x}_k)}. \quad (3.7)$$

La condition d'optimalité de l'algorithme 3.2 est alors obtenue en posant $\alpha = \underline{\gamma}^{-1} \overline{\nu}^{1/2}$.

3.4 ANALYSE DE CONVERGENCE

3.4.1 Propriétés de décroissance

Dans cette section nous allons donner quelques résultats techniques sur le comportement des suites $(\mathbf{f}(\mathbf{x}_k))_{k \in \mathbb{N}}$ et $(\mathbf{f}(\mathbf{y}_k))_{k \in \mathbb{N}}$ générées par l'algorithme 3.2. Ces propriétés nous permettront ensuite d'étudier la convergence de cet algorithme.

Lemme 3.1. *Supposons que les hypothèses 3.1, 3.2 et 3.3 soient satisfaites. Il existe $\mu_1 \in]0, +\infty[$ tel que, pour tout $k \in \mathbb{N}$,*

$$\mathbf{f}(\mathbf{x}_{k+1}) \leq \mathbf{f}(\mathbf{x}_k) - \frac{\mu_1}{2} \|\mathbf{x}_{k+1} - \mathbf{x}_k\|^2 \quad (3.8)$$

$$\leq \mathbf{f}(\mathbf{x}_k) - \underline{\lambda}^2 \frac{\mu_1}{2} \|\mathbf{y}_k - \mathbf{x}_k\|^2, \quad (3.9)$$

où les suites $(\mathbf{x}_k)_{k \in \mathbb{N}}$ et $(\mathbf{y}_k)_{k \in \mathbb{N}}$ sont générées par l'algorithme 3.2.

Démonstration. Pour tout $k \in \mathbb{N}$, la relaxation

$$\mathbf{x}_{k+1} = (1 - \lambda_k) \mathbf{x}_k + \lambda_k \mathbf{y}_k, \quad (3.10)$$

l'hypothèse 3.3(i) et la convexité de $\mathbf{g}$ nous donnent

$$\mathbf{f}(\mathbf{x}_{k+1}) \leq \mathbf{h}(\mathbf{x}_{k+1}) + (1 - \lambda_k)\mathbf{g}(\mathbf{x}_k) + \lambda_k\mathbf{g}(\mathbf{y}_k).$$

De plus, en utilisant l'hypothèse 3.2(i), on obtient

$$\mathbf{f}(\mathbf{x}_{k+1}) \leq \mathbf{h}(\mathbf{x}_k) + \langle \mathbf{x}_{k+1} - \mathbf{x}_k, \nabla \mathbf{h}(\mathbf{x}_k) \rangle + \frac{1}{2} \|\mathbf{x}_{k+1} - \mathbf{x}_k\|_{\mathbf{A}_k(\mathbf{x}_k)}^2 + (1 - \lambda_k)\mathbf{g}(\mathbf{x}_k) + \lambda_k\mathbf{g}(\mathbf{y}_k). \quad (3.11)$$

En reformulant (3.10), nous obtenons

$$\mathbf{x}_{k+1} - \mathbf{x}_k = \lambda_k(\mathbf{y}_k - \mathbf{x}_k). \quad (3.12)$$

En combinant la première inégalité de l'algorithme 3.2 et (3.12), on obtient

$$\langle \mathbf{x}_{k+1} - \mathbf{x}_k, \nabla \mathbf{h}(\mathbf{x}_k) \rangle \leq -\gamma_k^{-1} \lambda_k^{-1} \|\mathbf{x}_{k+1} - \mathbf{x}_k\|_{\mathbf{A}_k(\mathbf{x}_k)}^2 + \lambda_k(\mathbf{g}(\mathbf{x}_k) - \mathbf{g}(\mathbf{y}_k)). \quad (3.13)$$

Les équations (3.11) et (3.13) nous permettent de déduire que

$$\begin{aligned} \mathbf{f}(\mathbf{x}_{k+1}) &\leq \mathbf{f}(\mathbf{x}_k) - (\gamma_k^{-1} \lambda_k^{-1} - \frac{1}{2}) \|\mathbf{x}_{k+1} - \mathbf{x}_k\|_{\mathbf{A}_k(\mathbf{x}_k)}^2 \\ &\leq \mathbf{f}(\mathbf{x}_k) - \frac{1}{2} \frac{\overline{\eta}}{2 - \overline{\eta}} \|\mathbf{x}_{k+1} - \mathbf{x}_k\|_{\mathbf{A}_k(\mathbf{x}_k)}^2, \end{aligned}$$

où la dernière inégalité est obtenue en utilisant l'hypothèse 3.3(ii). Par conséquent, la borne inférieure dans l'hypothèse 3.2(ii) nous permet de déduire l'équation (3.8) en posant $\mu_1 = \frac{\underline{\nu}\overline{\eta}}{2 - \overline{\eta}} > 0$. L'inégalité (3.9) résulte ensuite de (3.12) et de l'hypothèse 3.3(i). $\square$

Le corollaire suivant, permettant de réécrire l'hypothèse 3.4 sous une forme différente, est une conséquence du lemme 3.1.

Corollaire 3.1. *Supposons que les hypothèses 3.1, 3.2 et 3.3 soient satisfaites. Alors l'hypothèse 3.4 est satisfaite si et seulement si il existe $\underline{\tau} \in]0, 1]$ tel que, pour tout $k \in \mathbb{N}$,*

$$(\exists \tau_k \in [\underline{\tau}, 1]) \quad \mathbf{f}(\mathbf{x}_{k+1}) \leq (1 - \tau_k)\mathbf{f}(\mathbf{x}_k) + \tau_k\mathbf{f}(\mathbf{y}_k). \quad (3.14)$$

Démonstration.

- (i) Supposons que l'hypothèse 3.4 soit satisfaite. En utilisant l'étape de relaxation donnée par (3.10), l'inégalité (3.14) est satisfaite en posant $\tau_k = \underline{\tau}$, pour tout $k \in \mathbb{N}$.
- (ii) Montrons maintenant la réciproque du point précédent. Soit $k \in \mathbb{N}$, l'inégalité (3.14) est équivalente à

$$\tau_k(\mathbf{f}(\mathbf{x}_k) - \mathbf{f}(\mathbf{y}_k)) \leq \mathbf{f}(\mathbf{x}_k) - \mathbf{f}(\mathbf{x}_{k+1}). \quad (3.15)$$

Si l'inégalité ci-dessus est satisfaite pour $\tau_k \in [\underline{\tau}, 1]$, alors on peut distinguer deux situations :

- *Cas où $\mathbf{f}(\mathbf{x}_k) \leq \mathbf{f}(\mathbf{y}_k)$.* D'après le lemme 3.1, on a $\mathbf{f}(\mathbf{x}_{k+1}) \leq \mathbf{f}(\mathbf{x}_k)$. D'où

$$\underline{\tau}(\mathbf{f}(\mathbf{x}_k) - \mathbf{f}(\mathbf{y}_k)) \leq 0 \leq \mathbf{f}(\mathbf{x}_k) - \mathbf{f}(\mathbf{x}_{k+1}).$$

- Cas où $f(\mathbf{x}_k) \geq f(\mathbf{y}_k)$. L'équation, (3.15) nous donne alors

$$\mathcal{I}(f(\mathbf{x}_k) - f(\mathbf{y}_k)) \leq \tau_k(f(\mathbf{x}_k) - f(\mathbf{y}_k)) \leq f(\mathbf{x}_k) - f(\mathbf{x}_{k+1}).$$

Ce qui permet de montrer que, si (3.14) est satisfaite, alors, l'hypothèse 3.4 l'est aussi. $\square$

Le lemme qui suit permet de retrouver des hypothèses standards sur le paramètre de relaxation $(\lambda_k)_{k \in \mathbb{N}}$, sous certaines conditions :

Corollaire 3.2. *Supposons que les hypothèses 3.1, 3.2 et 3.3 soient satisfaites. Si f est convexe sur $[\mathbf{x}_k, \mathbf{y}_k]$, pour tout $k \in \mathbb{N}$, alors l'hypothèse 3.4 est satisfaite.*

Démonstration. D'après l'hypothèse 3.3, pour tout $k \in \mathbb{N}$, $\lambda_k \in]0, 1]$. Si f est convexe sur $[\mathbf{x}_k, \mathbf{y}_k]$ pour tout $k \in \mathbb{N}$, alors

$$f((1 - \lambda_k)\mathbf{x}_k + \lambda_k\mathbf{y}_k) \leq (1 - \lambda_k)f(\mathbf{x}_k) + \lambda_k f(\mathbf{y}_k).$$

En utilisant le Corollaire 3.1 et le fait que, pour tout $k \in \mathbb{N}$, λ_k est bornée inférieurement par $\underline{\lambda} > 0$, on peut conclure que l'hypothèse 3.4 est satisfaite. $\square$

Le résultat suivant va nous permettre d'évaluer les variations de f entre les points $\mathbf{x}_k$ et $\mathbf{y}_k$, à chaque itération $k \in \mathbb{N}$ de l'algorithme 3.2.

Lemme 3.2. *Supposons que les hypothèses 3.1, 3.2 et 3.3 soient satisfaites. Il existe $\mu_2 \in \mathbb{R}$ tel que*

$$(\forall k \in \mathbb{N}) \quad f(\mathbf{y}_k) \leq f(\mathbf{x}_k) - \mu_2 \|\mathbf{y}_k - \mathbf{x}_k\|^2.$$

Démonstration. D'après l'hypothèse 3.2(i) et la première inégalité de l'algorithme 3.2, on a

$$h(\mathbf{y}_k) \leq q(\mathbf{y}_k, \mathbf{x}_k) \leq h(\mathbf{x}_k) + g(\mathbf{x}_k) - g(\mathbf{y}_k) - (\gamma_k^{-1} - \frac{1}{2}) \|\mathbf{y}_k - \mathbf{x}_k\|_{\mathbf{A}_k(\mathbf{x}_k)}^2,$$

ce qui, en utilisant l'hypothèse 3.3, nous conduit à

$$f(\mathbf{y}_k) \leq f(\mathbf{x}_k) - \left(\frac{\underline{\lambda}}{2 - \overline{\eta}} - \frac{1}{2} \right) \|\mathbf{y}_k - \mathbf{x}_k\|_{\mathbf{A}_k(\mathbf{x}_k)}^2.$$

Le résultat découle ensuite de l'hypothèse 3.2(ii) en posant

$$\mu_2 = \begin{cases} \underline{\nu} \left(\frac{\underline{\lambda}}{2 - \overline{\eta}} - \frac{1}{2} \right), & \text{si } 2\underline{\lambda} + \overline{\eta} \geq 2, \\ \overline{\nu} \left(\frac{\underline{\lambda}}{2 - \overline{\eta}} - \frac{1}{2} \right), & \text{sinon.} \end{cases}$$

$\square$

3.4.2 Résultat de convergence

Notre démonstration de convergence repose sur le résultat préliminaire suivant :

Lemme 3.3. Soient $(u_k)_{k \in \mathbb{N}}$, $(v_k)_{k \in \mathbb{N}}$, $(v'_k)_{k \in \mathbb{N}}$ et $(\Delta_k)_{k \in \mathbb{N}}$ des suites de réels positifs. Soit $\theta \in]0, 1[$. Supposons que :

- (i) Pour tout $k \in \mathbb{N}$, $u_k^2 \leq v_k^\theta \Delta_k$.
- (ii) $(\Delta_k)_{k \in \mathbb{N}}$ est une suite sommable.
- (iii) Pour tout $k \in \mathbb{N}$, $v_{k+1} \leq (1 - \underline{\alpha})v_k + v'_k$, avec $\underline{\alpha} \in]0, 1[$.
- (iv) Pour tout $k \geq k^*$, $(v'_k)^\theta \leq \varepsilon u_k$, où $\varepsilon > 0$ et $k^* \in \mathbb{N}$.

La suite $(u_k)_{k \in \mathbb{N}}$ est sommable.

Démonstration. En utilisant la sous-additivité de la fonction $t \mapsto t^\theta$, lorsque $\theta \in]0, 1[$, d'après l'assertion (iii), pour tout $k \in \mathbb{N}$, on a

$$v_{k+1}^\theta \leq (1 - \underline{\alpha})^\theta v_k^\theta + (v'_k)^\theta.$$

De plus, en utilisant (iv), on obtient

$$(\forall k \geq k^*) \quad v_{k+1}^\theta \leq (1 - \underline{\alpha})^\theta v_k^\theta + \varepsilon u_k,$$

ce qui implique que, pour tout $K > k^*$,

$$\sum_{k=k^*+1}^K v_k^\theta \leq (1 - \underline{\alpha})^\theta \sum_{k=k^*}^{K-1} v_k^\theta + \varepsilon \sum_{k=k^*}^{K-1} u_k, \quad (3.16)$$

et de manière équivalente, on obtient alors

$$(1 - (1 - \underline{\alpha})^\theta) \sum_{k=k^*}^{K-1} v_k^\theta \leq v_{k^*}^\theta - v_K^\theta + \varepsilon \sum_{k=k^*}^{K-1} u_k. \quad (3.17)$$

D'autre part, l'assertion (i) peut être reformulée de la façon suivante :

$$(\forall k \in \mathbb{N}) \quad u_k^2 \leq \left(\varepsilon^{-1} (1 - (1 - \underline{\alpha})^\theta) v_k^\theta \right) \left(\varepsilon (1 - (1 - \underline{\alpha})^\theta)^{-1} \Delta_k \right).$$

En utilisant l'inégalité $(\forall (t, t') \in [0, +\infty[^2]) \sqrt{tt'} \leq (t + t')/2$ et puisque, pour tout $k \in \mathbb{N}$, $u_k \geq 0$, on obtient

$$(\forall k \in \mathbb{N}) \quad u_k \leq \frac{1}{2} \left(\varepsilon^{-1} (1 - (1 - \underline{\alpha})^\theta) v_k^\theta + \varepsilon (1 - (1 - \underline{\alpha})^\theta)^{-1} \Delta_k \right). \quad (3.18)$$

On peut alors déduire des équations (3.17) et (3.18) que, pour tout $K > k^*$,

$$\sum_{k=k^*}^{K-1} u_k \leq \frac{1}{2} \left(\sum_{k=k^*}^{K-1} u_k + \varepsilon^{-1} (v_{k^*}^\theta - v_K^\theta) + \varepsilon (1 - (1 - \underline{\alpha})^\theta)^{-1} \sum_{k=k^*}^{K-1} \Delta_k \right),$$

et par conséquent,

$$\sum_{k=k^*}^{K-1} u_k \leq \varepsilon^{-1} v_{k^*}^\theta + \varepsilon (1 - (1 - \underline{\alpha})^\theta)^{-1} \sum_{k=k^*}^{K-1} \Delta_k.$$

Il en découle la sommabilité de la suite $(u_k)_{k \in \mathbb{N}}$ en utilisant (ii). □

Nous pouvons maintenant énoncer le résultat principal de ce chapitre concernant la convergence des suites $(\mathbf{x}_k)_{k \in \mathbb{N}}$ et $(\mathbf{y}_k)_{k \in \mathbb{N}}$ générées par l'algorithme 3.1 ou sa version inexacte, i.e. l'algorithme 3.2. Pour cela, en plus des hypothèses utilisées jusqu'ici dans ce chapitre, nous allons supposer que la fonction f satisfait l'inégalité de Kurdyka-Łojasiewicz (c.f. Chapitre 2).

Théorème 3.1. *Supposons que les hypothèses 3.1, 3.2 et 3.3 soient satisfaites. Supposons de plus que la fonction f satisfasse l'inégalité de Kurdyka-Łojasiewicz donnée dans la définition 2.20. Les assertions suivantes sont alors satisfaites :*

- (i) *Les suites $(\mathbf{x}_k)_{k \in \mathbb{N}}$ et $(\mathbf{y}_k)_{k \in \mathbb{N}}$ générées par l'algorithme 3.1 ou l'algorithme 3.2 convergent toutes les deux vers un point critique $\hat{\mathbf{x}}$ de f .*
- (ii) *Ces deux suites satisfont*

$$\sum_{k=0}^{+\infty} \|\mathbf{x}_{k+1} - \mathbf{x}_k\| < +\infty \quad \text{et} \quad \sum_{k=0}^{+\infty} \|\mathbf{y}_{k+1} - \mathbf{y}_k\| < +\infty.$$

- (iii) *Les suites $(f(\mathbf{x}_k))_{k \in \mathbb{N}}$ et $(f(\mathbf{y}_k))_{k \in \mathbb{N}}$ convergent vers $f(\hat{\mathbf{x}})$. De plus, $(f(\mathbf{x}_k))_{k \in \mathbb{N}}$ est décroissante.*

Remarque 3.6. *La démonstration de ce théorème est donnée dans l'annexe 3.A. Elle est constituée de trois étapes principales :*

- (i) *La première étape consiste à se servir des propriétés de décroissances présentées dans la section 3.4.1 pour montrer que les suites $(f(\mathbf{x}_k))_{k \in \mathbb{N}}$ et $(f(\mathbf{y}_k))_{k \in \mathbb{N}}$ convergent vers un réel ξ , et que $(f(\mathbf{x}_k))_{k \in \mathbb{N}}$ est décroissante.*
- (ii) *En utilisant l'inégalité de KL et le lemme 3.3, on montre dans la seconde étape que les suites $(\|\mathbf{x}_{k+1} - \mathbf{x}_k\|)_{k \in \mathbb{N}}$ et $(\|\mathbf{y}_{k+1} - \mathbf{y}_k\|)_{k \in \mathbb{N}}$ sont sommables, et donc que les suites $(\mathbf{x}_k)_{k \in \mathbb{N}}$ et $(\mathbf{y}_k)_{k \in \mathbb{N}}$ sont de Cauchy et convergent vers un point $\hat{\mathbf{x}} \in \mathbb{R}^N$.*
- (iii) *La dernière étape consiste à montrer que $\hat{\mathbf{x}}$ est un point critique de f , en utilisant la proposition 2.4.*

Le théorème de convergence précédent assure que les suites $(\mathbf{x}_k)_{k \in \mathbb{N}}$ et $(\mathbf{y}_k)_{k \in \mathbb{N}}$ convergent vers un point critique de f . De plus, la décroissance de la suite $(f(\mathbf{x}_k))_{k \in \mathbb{N}}$ implique que ce point critique ne sera pas un maximum de f . Cependant, en théorie, les algorithmes 3.1 et 3.2 peuvent converger vers un minimum local de f , ou bien un point selle. Le corollaire qui suit nous permet d'affirmer que, si les algorithmes 3.1 et 3.2 sont initialisés avec un point $\mathbf{x}_0$ suffisamment proche d'un minimum global de f , alors les suites $(\mathbf{x}_k)_{k \in \mathbb{N}}$ et $(\mathbf{y}_k)_{k \in \mathbb{N}}$ convergeront vers cette solution.

Corollaire 3.3. *Supposons que les hypothèses 3.1, 3.2 et 3.3 soient satisfaites. Supposons de plus que la fonction f satisfait l'inégalité de Kurdyka-Łojasiewicz. Soient $(\mathbf{x}_k)_{k \in \mathbb{N}}$ et $(\mathbf{y}_k)_{k \in \mathbb{N}}$ des suites générées par l'algorithme 3.2. Il existe $v > 0$ tel que, si*

$$f(\mathbf{x}_0) \leq \inf_{\mathbf{x} \in \mathbb{R}^N} f(\mathbf{x}) + v, \quad (3.19)$$

alors les suites $(\mathbf{x}_k)_{k \in \mathbb{N}}$ et $(\mathbf{y}_k)_{k \in \mathbb{N}}$ convergent toutes les deux vers une solution du problème (3.1).

Démonstration. D'après la remarque 3.1(iii), on a

$$\xi = \inf_{\mathbf{x} \in \mathbb{R}^N} f(\mathbf{x}) < +\infty.$$

Soit $E = \text{lev}_{\leq \xi + \delta} f$, où $\delta > 0$. Cet ensemble est borné d'après l'hypothèse 3.1(iii). Puisque f satisfait l'inégalité de KL, il existe trois constantes $\kappa > 0$, $\zeta > 0$ and $\theta \in [0, 1[$ telles que

$$(\forall \mathbf{r} \in \partial g(\mathbf{x})) \quad \kappa |f(\mathbf{x}) - \xi|^\theta \leq \|\nabla h(\mathbf{x}) + \mathbf{r}\|, \quad (3.20)$$

pour tout $\mathbf{x} \in E$ tel que $|f(\mathbf{x}) - \xi| \leq \zeta$. En d'autres termes, $f(\mathbf{x}) \leq \xi + \zeta$ puisque, par définition de ξ , $f(x) \geq \xi$ est toujours vraie. Nous posons maintenant $v = \min\{\delta, \zeta\} > 0$ et choisissons $\mathbf{x}_0$ de façon à satisfaire (3.19). Le théorème 3.1(iii) implique que, pour tout $k \in \mathbb{N}$,

$$f(\mathbf{x}_k) \leq \xi + v.$$

La restriction de f à son domaine de définition étant continue, d'après le théorème 3.1(i), les suites $(\mathbf{x}_k)_{k \in \mathbb{N}}$ et $(\mathbf{y}_k)_{k \in \mathbb{N}}$ convergent vers $\hat{\mathbf{x}}$, où $\hat{\mathbf{x}}$ est tel que $f(\hat{\mathbf{x}}) \leq \xi + v$. L'inégalité de KL est donc satisfaite en $\hat{\mathbf{x}}$:

$$(\forall \mathbf{t} \in \partial f(\hat{\mathbf{x}})) \quad \|\mathbf{t}\| \geq \kappa |f(\hat{\mathbf{x}}) - \xi|^\theta.$$

De plus, $\hat{\mathbf{x}}$ étant un point critique de f , on a $\mathbf{0} \in \partial f(\hat{\mathbf{x}})$, et

$$|f(\hat{\mathbf{x}}) - \xi|^\theta \leq 0.$$

Pour conclure, nous obtenons $f(\hat{\mathbf{x}}) = \inf_{\mathbf{x} \in \mathbb{R}^N} f(\mathbf{x})$. □

Remarque 3.7.

- (i) *Remarquons que, bien que la constante de Lipschitz β n'apparaisse pas explicitement dans nos algorithmes, l'hypothèse de Lipschitz différentiability est indispensable dans les démonstrations de nos théorèmes de convergence.*
- (ii) *Nous pouvons remarquer que la convergence de l'algorithme 3.2 a été démontrée dans [Attouch et al., 2011] dans le cas non préconditionné, i.e. lorsque $\mathbf{A}_k(\mathbf{x}_k) \equiv \beta \mathbf{I}_N$ dans le théorème 2.9.*

3.5 CONCLUSION

Dans ce chapitre nous avons proposé de minimiser une fonction $f = h + g$ où h est une fonction différentiable non-nécessairement convexe, g est une fonction convexe non-nécessairement différentiable, et f est une fonction coercive satisfaisant l'inégalité de KL. Pour cela, nous avons donné une version accélérée de l'algorithme explicite-implicite. Cette accélération est basée sur l'introduction d'une métrique variant à chaque itération, choisie grâce au principe de Majoration-Minimisation. Enfin, nous avons démontré la convergence de l'algorithme itératif proposé vers un point critique de la fonction à minimiser. Plus précisément, nous avons démontré :

- un résultat de convergence global, affirmant que les suites générées par l'algorithme explicite-implicite à métrique variable convergent vers un point critique de la fonction f ,
- et un résultat de convergence local, montrant qu'en initialisant judicieusement l'algorithme, alors les suites générées convergent vers un minimum global de h .

Afin d'illustrer l'intérêt de l'utilisation de métriques variables, nous traiterons dans le chapitre 4 une application en reconstruction d'images.

3.A DÉMONSTRATION DU THÉORÈME 3.1

Dans cette annexe, nous allons établir le théorème 3.1 en montrant successivement les étapes (i)-(iii) décrites dans la remarque 3.6.

(i) D'après le lemme 3.1, on a

$$(\forall k \in \mathbb{N}) \quad f(\mathbf{x}_{k+1}) \leq f(\mathbf{x}_k),$$

donc $(f(\mathbf{x}_k))_{k \in \mathbb{N}}$ est une suite décroissante. De plus, les remarques 3.1(iii) et 3.2 assurent que les éléments des suites $(\mathbf{x}_k)_{k \in \mathbb{N}}$ et $(\mathbf{y}_k)_{k \in \mathbb{N}}$ sont dans un sous-ensemble compact E de $\text{lev}_{\leq f(\mathbf{x}_0)} f \subset \text{dom } g$ et que f est bornée inférieurement. Par conséquent, $(f(\mathbf{x}_k))_{k \in \mathbb{N}}$ converge vers un réel ξ , et $(f(\mathbf{x}_k) - \xi)_{k \in \mathbb{N}}$ est une suite positive convergeant vers 0. En outre, le lemme 3.1 nous donne

$$(\forall k \in \mathbb{N}) \quad \lambda^2 \frac{\mu_1}{2} \|\mathbf{y}_k - \mathbf{x}_k\|^2 \leq (f(\mathbf{x}_k) - \xi) - (f(\mathbf{x}_{k+1}) - \xi). \quad (3.21)$$

Ainsi, la suite $(\mathbf{y}_k - \mathbf{x}_k)_{k \in \mathbb{N}}$ converge vers $\mathbf{0}$.

D'autre part, la relaxation

$$(\forall k \in \mathbb{N}) \quad \mathbf{x}_{k+1} = (1 - \lambda_k)\mathbf{x}_k + \lambda_k \mathbf{y}_k, \quad (3.22)$$

et l'hypothèses 3.4 impliquent que, pour tout $k \in \mathbb{N}$,

$$f(\mathbf{x}_{k+1}) - \xi \leq (1 - \underline{\alpha})(f(\mathbf{x}_k) - \xi) + \underline{\alpha}(f(\mathbf{y}_k) - \xi).$$

L'inégalité précédente combinée avec le lemme 3.2, nous donne alors, pour tout $k \in \mathbb{N}$,

$$\underline{\alpha}^{-1}(f(\mathbf{x}_{k+1}) - \xi - (1 - \underline{\alpha})(f(\mathbf{x}_k) - \xi)) \leq f(\mathbf{y}_k) - \xi \leq f(\mathbf{x}_k) - \xi - \mu_2 \|\mathbf{y}_k - \mathbf{x}_k\|^2.$$

Ainsi, puisque $(\|\mathbf{y}_k - \mathbf{x}_k\|)_{k \in \mathbb{N}}$ et $(f(\mathbf{x}_k) - \xi)_{k \in \mathbb{N}}$ convergent vers 0, la suite $(f(\mathbf{y}_k))_{k \in \mathbb{N}}$ converge vers ξ .

- (ii) Revenons à l'équation (3.21). Soit $\psi: [0, +\infty[\rightarrow [0, +\infty[: t \mapsto t^{1/(1-\theta)}$, où $\theta \in [0, 1[$, une fonction convexe à laquelle nous appliquons l'inégalité de gradient suivante

$$(\forall (t, t') \in [0, +\infty[^2) \quad \psi(t) - \psi(t') \leq \dot{\psi}(t)(t - t').$$

Après un changement de variable, cette inégalité peut être réécrite comme suit

$$(\forall (t, t') \in [0, +\infty[^2) \quad t - t' \leq (1 - \theta)^{-1} t^\theta (t^{1-\theta} - t'^{1-\theta}).$$

En prenant $\mathbf{t} = \mathbf{f}(\mathbf{x}_k) - \xi$ et $\mathbf{t}' = \mathbf{f}(\mathbf{x}_{k+1}) - \xi$, on a alors

$$(\forall k \in \mathbb{N}) \quad (\mathbf{f}(\mathbf{x}_k) - \xi) - (\mathbf{f}(\mathbf{x}_{k+1}) - \xi) \leq (1 - \theta)^{-1} (\mathbf{f}(\mathbf{x}_k) - \xi)^\theta \Delta'_k,$$

où

$$(\forall k \in \mathbb{N}) \quad \Delta'_k = (\mathbf{f}(\mathbf{x}_k) - \xi)^{1-\theta} - (\mathbf{f}(\mathbf{x}_{k+1}) - \xi)^{1-\theta}.$$

En combinant l'inégalité ci-dessus avec (3.21) on obtient

$$(\forall k \in \mathbb{N}) \quad \|\mathbf{y}_k - \mathbf{x}_k\|^2 \leq 2\lambda^{-2} \mu_1^{-1} (1 - \theta)^{-1} (\mathbf{f}(\mathbf{x}_k) - \xi)^\theta \Delta'_k. \quad (3.23)$$

Par ailleurs, puisque E est borné, $\mathbf{f}$ satisfait l'inégalité de KL et comme $\mathbf{h}$ est différentiable, il existe trois constantes $\kappa > 0$, $\zeta > 0$ et $\theta \in [0, 1[$ telles que

$$(\forall \mathbf{r} \in \partial \mathbf{g}(\mathbf{x})) \quad \kappa |\mathbf{f}(\mathbf{x}) - \xi|^\theta \leq \|\nabla \mathbf{h}(\mathbf{x}) + \mathbf{r}\|, \quad (3.24)$$

pour tout $\mathbf{x} \in E$ tel que $|\mathbf{f}(\mathbf{x}) - \xi| \leq \zeta$. La suite $(\mathbf{f}(\mathbf{y}_k))_{k \in \mathbb{N}}$ convergeant vers ξ , il existe $k^* \in \mathbb{N}$, tel que, pour tout $k \geq k^*$, $|\mathbf{f}(\mathbf{y}_k) - \xi| < \zeta$. Ainsi, on a, pour tout $\mathbf{r}_k \in \partial \mathbf{g}(\mathbf{y}_k)$,

$$(\forall k \geq k^*) \quad \kappa |\mathbf{f}(\mathbf{y}_k) - \xi|^\theta \leq \|\nabla \mathbf{h}(\mathbf{y}_k) + \mathbf{r}_k\|.$$

Soit $\mathbf{r}_k$ donné dans l'algorithme 3.2. On a alors

$$\begin{aligned} \kappa |\mathbf{f}(\mathbf{y}_k) - \xi|^\theta &\leq \|\nabla \mathbf{h}(\mathbf{y}_k) - \nabla \mathbf{h}(\mathbf{x}_k) + \nabla \mathbf{h}(\mathbf{x}_k) + \mathbf{r}_k\| \\ &\leq \|\nabla \mathbf{h}(\mathbf{y}_k) - \nabla \mathbf{h}(\mathbf{x}_k)\| + \|\nabla \mathbf{h}(\mathbf{x}_k) + \mathbf{r}_k\| \\ &\leq \|\nabla \mathbf{h}(\mathbf{y}_k) - \nabla \mathbf{h}(\mathbf{x}_k)\| + \tau \|\mathbf{x}_k - \mathbf{y}_k\|_{\mathbf{A}_k(\mathbf{x}_k)}. \end{aligned} \quad (3.25)$$

Donc, en utilisant les hypothèses 3.1(ii) et 3.2(ii), on obtient

$$|\mathbf{f}(\mathbf{y}_k) - \xi|^\theta \leq \kappa^{-1} (\beta + \tau \sqrt{\bar{\nu}}) \|\mathbf{x}_k - \mathbf{y}_k\|. \quad (3.26)$$

De plus, d'après l'hypothèse 3.4,

$$\mathbf{f}(\mathbf{x}_{k+1}) - \xi \leq (1 - \underline{\alpha})(\mathbf{f}(\mathbf{x}_k) - \xi) + |\mathbf{f}(\mathbf{y}_k) - \xi|. \quad (3.27)$$

Par ailleurs, nous pouvons remarquer que

$$\begin{aligned} \sum_{k=k^*}^{+\infty} \Delta'_k &= \sum_{k=k^*}^{+\infty} (\mathbf{f}(\mathbf{x}_k) - \xi)^{1-\theta} - (\mathbf{f}(\mathbf{x}_{k+1}) - \xi)^{1-\theta} \\ &= (\mathbf{f}(\mathbf{x}_{k^*}) - \xi)^{1-\theta}, \end{aligned}$$

ce qui montre que $(\Delta'_k)_{k \in \mathbb{N}}$ est une suite sommable. Posons

$$(\forall k \in \mathbb{N}) \quad \begin{cases} \mathbf{u}_k = \|\mathbf{y}_k - \mathbf{x}_k\|, \\ \mathbf{v}_k = \mathbf{f}(\mathbf{x}_k) - \xi \geq 0, \\ \mathbf{v}'_k = |\mathbf{f}(\mathbf{y}_k) - \xi| \geq 0, \\ \varepsilon = \kappa^{-1} (\beta + \tau \sqrt{\bar{\nu}}) > 0, \\ \Delta_k = 2\lambda^{-2} \mu_1^{-1} (1 - \theta)^{-1} \Delta'_k. \end{cases}$$

En utilisant l'équation (3.23), la sommabilité de $(\Delta'_k)_{k \in \mathbb{N}}$ et les équations (3.26) et (3.27), le lemme 3.3 nous permet d'affirmer que $(\|\mathbf{y}_k - \mathbf{x}_k\|)_{k \in \mathbb{N}}$ est une suite sommable dès que $\theta \neq 0$. Dans le cas où $\theta = 0$, puisque $\mathbf{x}_k - \mathbf{y}_k \rightarrow \mathbf{0}$, il existe $k^{**} \geq k^*$ tel que

$$(\forall k \geq k^{**}) \quad \kappa^{-1}(\beta + \tau\sqrt{\nu})\|\mathbf{x}_k - \mathbf{y}_k\| < 1.$$

Ainsi, d'après l'équation (3.26) (rappelons que nous utilisons la convention $0^0 = 0$), on a nécessairement $f(\mathbf{y}_k) = \xi$, pour tout $k \geq k^{**}$. Alors, d'après (3.23), pour tout $k \geq k^{**}$, $\mathbf{x}_k = \mathbf{y}_k$. Ceci prouve que $(\|\mathbf{y}_k - \mathbf{x}_k\|)_{k \in \mathbb{N}}$ est sommable.

D'autre part, en utilisant l'équation (3.22) et l'hypothèse 3.3(i), on a

$$(\forall k \in \mathbb{N}) \quad \|\mathbf{x}_{k+1} - \mathbf{x}_k\| = \lambda_k \|\mathbf{y}_k - \mathbf{x}_k\| \leq \|\mathbf{y}_k - \mathbf{x}_k\|.$$

Par conséquent, la suite $(\mathbf{x}_k)_{k \in \mathbb{N}}$ satisfait

$$\sum_{k=0}^{+\infty} \|\mathbf{x}_{k+1} - \mathbf{x}_k\| < +\infty. \quad (3.28)$$

Cela implique que $(\mathbf{x}_k)_{k \in \mathbb{N}}$ est une suite de Cauchy. Elle converge donc vers un point que l'on notera $\hat{\mathbf{x}}$. Puisque $(\mathbf{x}_k - \mathbf{y}_k)_{k \in \mathbb{N}}$ converge vers $\mathbf{0}$, $(\mathbf{y}_k)_{k \in \mathbb{N}}$ converge de même vers $\hat{\mathbf{x}}$. De plus, on a

$$\begin{aligned} (\forall k \in \mathbb{N}) \quad \|\mathbf{y}_{k+1} - \mathbf{y}_k\| &\leq \|\mathbf{y}_{k+1} - \mathbf{x}_k\| + \|\mathbf{x}_k - \mathbf{y}_k\| \\ &\leq \|\mathbf{y}_{k+1} - \mathbf{x}_{k+1}\| + \|\mathbf{x}_{k+1} - \mathbf{x}_k\| + \|\mathbf{x}_k - \mathbf{y}_k\|, \end{aligned}$$

ainsi, $(\mathbf{y}_k)_{k \in \mathbb{N}}$ satisfait

$$\sum_{k=0}^{+\infty} \|\mathbf{y}_{k+1} - \mathbf{y}_k\| < +\infty. \quad (3.29)$$

- (iii) Il ne nous reste plus qu'à démontrer que la limite $\hat{\mathbf{x}}$ est un point critique de f . Pour cela, posons

$$(\forall k \in \mathbb{N}) \quad \mathbf{t}_k = \nabla h(\mathbf{y}_k) + \mathbf{r}_k,$$

tel que $(\mathbf{y}_k, \mathbf{t}_k) \in \text{Graph } \partial f$, où $\mathbf{r}_k$ est donné dans l'algorithme 3.2. En raisonnant de façon analogue au cheminement suivi pour l'obtention de l'équation (3.26), on a

$$(\forall k \in \mathbb{N}) \quad \|\mathbf{t}_k\| \leq (\beta + \tau\sqrt{\nu})\|\mathbf{x}_k - \mathbf{y}_k\|.$$

Étant donné que les suites $(\mathbf{x}_k)_{k \in \mathbb{N}}$ et $(\mathbf{y}_k)_{k \in \mathbb{N}}$ convergent toutes les deux vers $\hat{\mathbf{x}}$, $(\mathbf{y}_k, \mathbf{t}_k)_{k \in \mathbb{N}}$ converge vers $(\hat{\mathbf{x}}, \mathbf{0})$. Par ailleurs, la remarque 3.1(iii) nous permet d'affirmer que la restriction de f à son domaine est continue. Ainsi, puisque $(\forall k \in \mathbb{N}) \mathbf{y}_k \in \text{dom } f$, la suite $(f(\mathbf{y}_k))_{k \in \mathbb{N}}$ converge vers $f(\hat{\mathbf{x}})$. Au final, ∂f étant fermé (c.f. proposition 2.4), $(\hat{\mathbf{x}}, \mathbf{0}) \in \text{Graph } \partial f$. Autrement dit, $\hat{\mathbf{x}}$ est un point critique de f .

CHAPITRE 4

Algorithme explicite-implicite : application à un bruit gaussien dépendant

Sommaire

4.1	Introduction	76
4.2	Bruit gaussien dépendant du signal	76
4.2.1	Estimateur du <i>Maximum A Posteriori</i>	76
4.2.2	Fonction d'attache aux données	80
4.3	Application de l'algorithme explicite-implicite à métrique variable . .	81
4.3.1	Problème d'optimisation	81
4.3.2	Construction de la majorante	82
4.3.3	Implémentation de l'étape proximale	86
4.4	Résultats de simulation	87
4.4.1	Restauration d'image	87
4.4.2	Reconstruction d'image	89
4.5	Conclusion	89
Annexe 4.A	Démonstration de la proposition 4.4	91

4.1 INTRODUCTION

Dans ce chapitre, nous nous intéressons à un problème de restauration/reconstruction d'image où l'objectif est de trouver une estimée d'une image originale à partir d'une version dégradée de celle-ci. Plus précisément, nous nous intéressons au cas où l'image observée peut s'écrire comme la somme de l'image originale dégradée par un opérateur linéaire et d'un bruit gaussien dépendant. Notons que ce modèle d'observation se rencontre dans de nombreux systèmes d'imagerie [Healey et Kondepudy, 1994; Janesick, 2007; Tian *et al.*, 2001] dans lesquels les acquisitions sont dégradées par un bruit gaussien dépendant lié au comptage de photons, et un bruit électronique indépendant. Le bruit gaussien dépendant peut être également vu comme une approximation au second ordre du bruit Poisson-gaussien fréquemment considéré en astronomie, biologie et imagerie médicale [Foi *et al.*, 2008; Jezierska *et al.*, 2012; Li *et al.*, 2012].

Nous verrons que le problème de minimisation associé à ce modèle, issu du MAP, n'est pas convexe. Nous proposerons donc d'utiliser l'algorithme explicite-implicite à métrique variable, étudié dans le chapitre précédent, pour trouver une solution de ce problème. Cela nous permettra de comparer, sur un exemple pratique, notre algorithme avec sa version non préconditionnée.

Le plan de ce chapitre sera le suivant : nous allons étudier le modèle de bruit gaussien dépendant dans la section 4.2. Puis nous appliquerons l'algorithme développé dans le chapitre précédent à ce modèle dans la section 4.3 et donnerons des résultats de simulations dans la section 4.4. Enfin, nous conclurons dans la section 4.5.

4.2 BRUIT GAUSSIEN DÉPENDANT DU SIGNAL

4.2.1 Estimateur du *Maximum A Posteriori*

On s'intéresse à l'estimation $\hat{\mathbf{x}} \in \mathbb{R}^N$ d'une image originale $\bar{\mathbf{x}} \in \mathbb{R}^N$ à partir d'une image observée $\mathbf{z} \in \mathbb{R}^M$ de la forme

$$\begin{cases} \mathbf{z} = \mathbf{y} + \mathbf{w}, \\ \mathbf{y} = \mathbf{H}\bar{\mathbf{x}}, \end{cases} \quad (4.1)$$

où $\mathbf{H} \in \mathbb{R}^{M \times N}$ est une matrice modélisant la dégradation subie par l'image (par exemple, un opérateur de convolution) et $\mathbf{w} \in \mathbb{R}^M$ est une réalisation d'un vecteur aléatoire absolument continu $\mathbf{w}$ modélisant un bruit additif.

Supposons que $\mathbf{z}$ et $\bar{\mathbf{x}}$ soient des réalisations des vecteurs aléatoires $\mathbf{z}$ et $\bar{\mathbf{x}}$. Comme nous l'avons vu dans le chapitre 2, on peut alors utiliser une approche de type MAP pour définir l'estimée $\hat{\mathbf{x}}$ de $\bar{\mathbf{x}}$:

$$\hat{\mathbf{x}} \in \underset{\mathbf{x} \in \mathbb{R}^N}{\text{Argmax}} \ f_{\bar{\mathbf{x}}|\mathbf{z}=\mathbf{z}}(\mathbf{x}), \quad (4.2)$$

où $f_{\bar{\mathbf{x}}|\mathbf{z}=\mathbf{z}}$ est la densité de probabilité *a posteriori* de $\bar{\mathbf{x}}$ sachant $\mathbf{z} = \mathbf{z}$. Dans la propriété 2.1 nous avons exprimé le problème variationnel associé au modèle (4.1) lorsque le bruit est indépendant de l'image.

Bruit dépendant du signal. Dans ce chapitre, nous nous intéressons au cas où le bruit et l'image originale sont inter-dépendants. Dans ce cas, la proposition suivante est vérifiée :

Proposition 4.1. Soit $\mathbf{z} \in \mathbb{R}^M$ défini par l'équation (4.1). Une solution $\hat{\mathbf{x}}$ du problème (4.2) est donnée par

$$\hat{\mathbf{x}} \in \underset{\mathbf{x} \in \mathbb{R}^N}{\text{Argmin}} \ h(\mathbf{x}) + g(\mathbf{x}), \quad (4.3)$$

où $h: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ et $g: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ sont respectivement les termes d'attache aux données et de régularisation définis par

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \begin{cases} h(\mathbf{x}) = -\log f_{\mathbf{w}|\mathbf{y}=\mathbf{H}\mathbf{x}}(\mathbf{z} - \mathbf{H}\mathbf{x}), \\ g(\mathbf{x}) = -\log f_{\bar{\mathbf{x}}}(\mathbf{x}), \end{cases} \quad (4.4)$$

où $\mathbf{y} = \mathbf{H}\bar{\mathbf{x}}$.

Démonstration. Il suffit de reprendre les équations (2.9) et (2.10) de la démonstration de la propriété 2.1. En les combinant, on obtient

$$\hat{\mathbf{x}} \in \underset{\mathbf{x} \in \mathbb{R}^N}{\text{Argmin}} \quad \left\{ -\log f_{\mathbf{w}|\mathbf{H}\bar{\mathbf{x}}=\mathbf{H}\mathbf{x}}(\mathbf{z} - \mathbf{H}\mathbf{x}) - \log f_{\bar{\mathbf{x}}}(\mathbf{x}) \right\}. \quad (4.5)$$

□

Bruit gaussien dépendant du signal. Supposons maintenant que la variable aléatoire $\mathbf{w}$, modélisant le bruit additif dans le modèle (4.1), suive une loi normale de moyenne nulle et de matrice de covariance $\mathbf{\Gamma}(\mathbf{y}) \in \mathbb{R}^{M \times M}$, telle que $\mathbf{\Gamma}(\mathbf{y}) \succ 0$, dépendant de la réalisation $\mathbf{y} = \mathbf{H}\bar{\mathbf{x}}$ du vecteur aléatoire $\mathbf{y} = \mathbf{H}\bar{\mathbf{x}}$:

$$\mathbf{w}|\mathbf{y} = \mathbf{y} \sim \mathcal{N}(\mathbf{0}_M, \mathbf{\Gamma}(\mathbf{y})). \quad (4.6)$$

Le problème variationnel associé à ce modèle s'écrit alors de la façon suivante :

Lemme 4.1. On considère le modèle (4.1), où $\mathbf{w}$ est une réalisation du vecteur aléatoire $\mathbf{w}$ dont la loi sachant $\mathbf{y} = \mathbf{y}$ est donnée par (4.6) avec $\mathbf{\Gamma}(\mathbf{y}) \in \mathbb{R}^{M \times M}$, telle que $\mathbf{\Gamma}(\mathbf{y}) \succ 0$. Une solution $\hat{\mathbf{x}}$ du problème (4.2) est alors donnée par

$$\hat{\mathbf{x}} \in \underset{\mathbf{x} \in \mathbb{R}^N}{\text{Argmin}} \quad \frac{1}{2}(\mathbf{z} - \mathbf{H}\mathbf{x})^\top \mathbf{\Gamma}(\mathbf{H}\mathbf{x})^{-1}(\mathbf{z} - \mathbf{H}\mathbf{x}) + \frac{1}{2} \log \left(\det(\mathbf{\Gamma}(\mathbf{H}\mathbf{x})) \right) + g(\mathbf{x}), \quad (4.7)$$

où $g: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ est une fonction de régularisation dépendant des caractéristiques physiques de l'image originale.

Démonstration. La fonction de densité de probabilité de $\mathbf{w}$ sachant $\mathbf{y} = \mathbf{y}$ est donnée par

$$(\forall \mathbf{w} \in \mathbb{R}^M) \quad f_{\mathbf{w}|\mathbf{y}=\mathbf{y}}(\mathbf{w}) = \frac{1}{(2\pi)^{N/2}(\det \mathbf{\Gamma}(\mathbf{y}))^{1/2}} \exp \left(-\frac{1}{2} \mathbf{w}^\top \mathbf{\Gamma}(\mathbf{y})^{-1} \mathbf{w} \right). \quad (4.8)$$

Il suffit alors de combiner cette équation à la proposition 4.1. □

Bruit gaussien décorrélé et dépendant du signal. Dans la suite de ce chapitre, nous nous intéresserons au cas particulier où les coefficients de l'image originale $\bar{\mathbf{x}}$ et de la matrice de dégradation $\mathbf{H}$ sont positifs. De plus, nous supposerons que la matrice de covariance $\mathbf{\Gamma}(\mathbf{H}\bar{\mathbf{x}})$, dépendant de l'image $\mathbf{H}\bar{\mathbf{x}} = ([\mathbf{H}\bar{\mathbf{x}}]^{(m)})_{1 \leq m \leq M} \in [0, +\infty[^M$, est une matrice diagonale donnée par

$$\mathbf{\Gamma}(\mathbf{H}\bar{\mathbf{x}}) = \text{Diag} \left(\left(\sigma_m([\mathbf{H}\bar{\mathbf{x}}]^{(m)})^2 \right)_{1 \leq m \leq M} \right), \quad (4.9)$$

où, pour tout $m \in \{1, \dots, M\}$,

$$\begin{aligned} \sigma_m: [0, +\infty[&\rightarrow]0, +\infty[\\ y &\mapsto \left(\alpha^{(m)}y + \beta^{(m)} \right)^{1/2}, \end{aligned} \quad (4.10)$$

avec $(\alpha^{(m)}, \beta^{(m)}) \in [0, +\infty[\times]0, +\infty[$.

Dans ce cadre, le modèle (4.1) peut se réécrire

$$(\forall m \in \{1, \dots, M\}) \quad \mathbf{z}^{(m)} = [\mathbf{H}\bar{\mathbf{x}}]^{(m)} + \sigma_m([\mathbf{H}\bar{\mathbf{x}}]^{(m)}) \mathbf{v}^{(m)}, \quad (4.11)$$

où, pour tout $m \in \{1, \dots, M\}$, $\mathbf{v}^{(m)}$ est une réalisation d'une variable aléatoire $v^{(m)}$ suivant une loi normale de moyenne nulle et de variance égale à 1.

Remarque 4.1.

- (i) Lorsque, pour tout $m \in \{1, \dots, M\}$, $\alpha_m = 0$ et $\beta_m = \sigma^2 \in]0, +\infty[$, le bruit gaussien dépendant se réduit au bruit blanc gaussien (c.f. section 2.2.3 du chapitre 2).
- (ii) En pratique, afin d'assurer la positivité de l'estimée $\hat{\mathbf{x}}$, on peut utiliser comme terme de régularisation la fonction indicatrice sur l'orthant positif :

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \mathbf{g}(\mathbf{x}) = \iota_{[0, +\infty[^N}(\mathbf{x}). \quad (4.12)$$

Dans la figure 4.1, une comparaison est faite entre une image bruitée par un bruit blanc gaussien (figures 4.1(b) et (c)) et une image bruitée par un bruit dépendant gaussien (figures 4.1(d) et (e)). L'image originale (figure 4.1(a)) correspond à l'image *phantom* générée à l'aide de Matlab.

Proposition 4.2. On considère le modèle (4.11), où $\mathbf{H} \in [0, +\infty[^{M \times N}$ et $\bar{\mathbf{x}} \in [0, +\infty[^N$. L'estimée $\hat{\mathbf{x}}$ de $\bar{\mathbf{x}}$ associée au MAP est définie par

$$\hat{\mathbf{x}} \in \underset{\mathbf{x} \in [0, +\infty[^N}{\text{Argmin}} \{ \mathbf{h}_1(\mathbf{x}) + \mathbf{h}_2(\mathbf{x}) + \mathbf{g}(\mathbf{x}) \}, \quad (4.13)$$

où $\mathbf{h}_1$ et $\mathbf{h}_2$ sont des fonctions de $[0, +\infty[^N$ vers $] -\infty, +\infty]$, données par

$$(\forall \mathbf{x} \in [0, +\infty[^N) \quad \mathbf{h}_1(\mathbf{x}) = \sum_{m=1}^M \phi_m([\mathbf{H}\mathbf{x}]^{(m)}) \quad \text{et} \quad \mathbf{h}_2(\mathbf{x}) = \sum_{m=1}^M \psi_m([\mathbf{H}\mathbf{x}]^{(m)}) \quad (4.14)$$

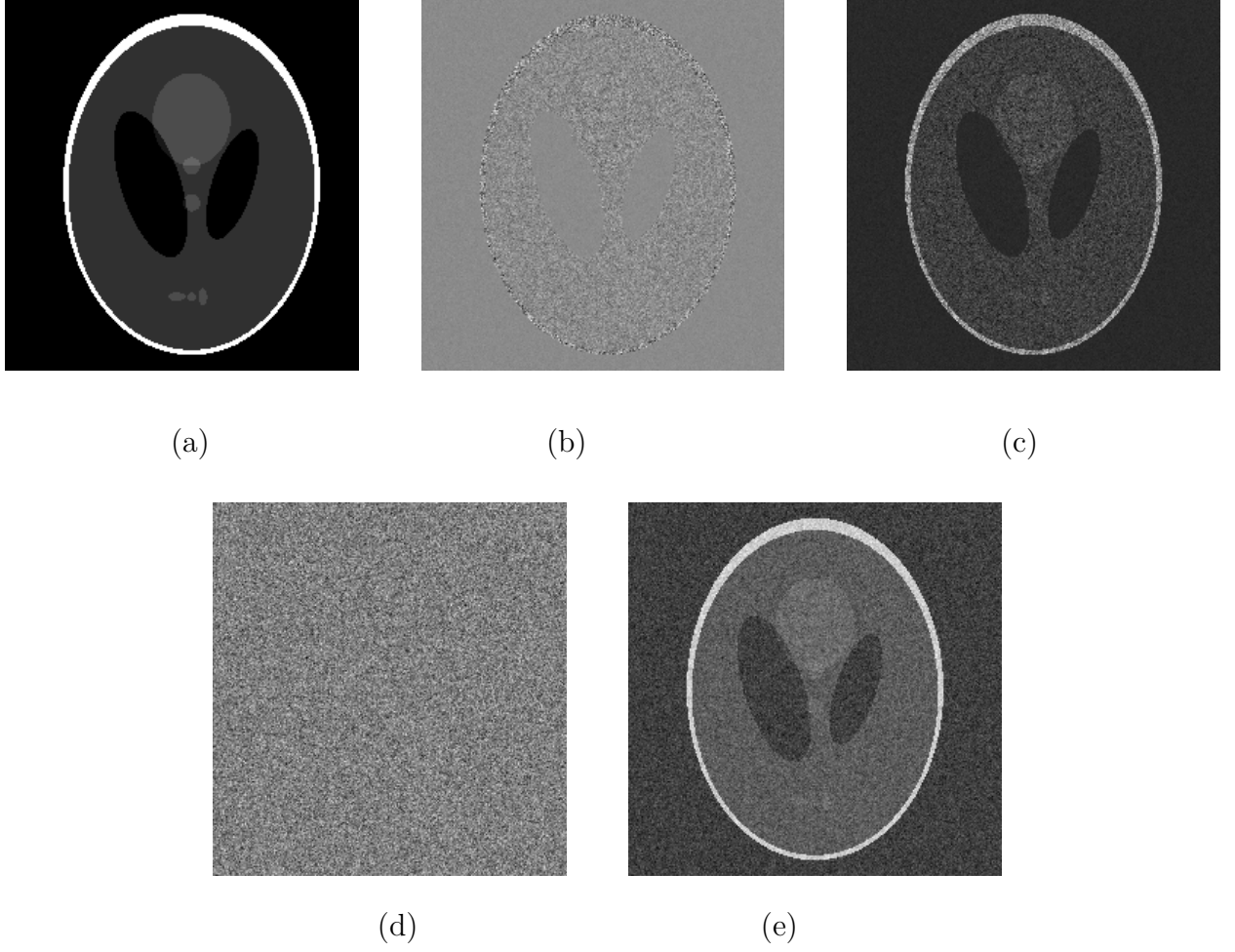


Figure 4.1 – (a) Image originale phantom de dimension 256×256 . (b) Bruit gaussien dépendant de l'image originale avec $\alpha^{(m)} \equiv 0.1$ et $\beta^{(m)} \equiv 0.001$. (c) Image bruitée avec le bruit dépendant gaussien. (d) Bruit blanc gaussien de variance $\sigma = 0.1154$. (e) Image bruitée avec le bruit blanc gaussien. Les images bruitées (c) et (e) ont le même RSB de 6.6 dB.

avec, pour tout $m \in \{1, \dots, M\}$,

$$(\forall y \in [0, +\infty[) \quad \phi_m(y) = \frac{1}{2} \frac{(z^{(m)} - y)^2}{\alpha^{(m)}y + \beta^{(m)}}, \quad (4.15)$$

$$\psi_m(y) = \frac{1}{2} \log \left(\alpha^{(m)}y + \beta^{(m)} \right), \quad (4.16)$$

et $g: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ une fonction de régularisation dépendant des caractéristiques physiques de l'image.

Démonstration. Ce résultat est obtenu en injectant les équations (4.9) et (4.10) dans le lemme 4.1. $\square$

4.2.2 Fonction d'attache aux données

Dans cette section, nous étudions les propriétés des fonctions h_1 et h_2 définies dans la proposition 4.2. Pour cela nous définissons les fonctions

$$(\forall y \in [0, +\infty[) \quad \chi_m(y) = \phi_m(y) + \psi_m(y), \quad (4.17)$$

$$(\forall \mathbf{H} \in [0, +\infty[^{M \times N})(\forall \mathbf{x} \in [0, +\infty[^N) \quad h(\mathbf{x}) = h_1(\mathbf{x}) + h_2(\mathbf{x}) = \sum_{m=1}^M \chi_m([\mathbf{H}\mathbf{x}]^{(m)}). \quad (4.18)$$

Pour tout $m \in \{1, \dots, M\}$, les dérivées premières et secondes des fonctions ϕ_m et ψ_m en $y \in [0, +\infty[$ sont données par

$$\dot{\phi}_m(y) = \frac{(y - z^{(m)})(\alpha^{(m)}y + \alpha^{(m)}z^{(m)} + 2\beta^{(m)})}{2(\alpha^{(m)}y + \beta^{(m)})^2}, \quad \ddot{\phi}_m(y) = \frac{(\alpha^{(m)}z^{(m)} + \beta^{(m)})^2}{(\alpha^{(m)}y + \beta^{(m)})^3}, \quad (4.19)$$

$$\dot{\psi}_m(y) = \frac{\alpha^{(m)}}{2(\alpha^{(m)}y + \beta^{(m)})}, \quad \ddot{\psi}_m(y) = -\frac{(\alpha^{(m)})^2}{2(\alpha^{(m)}y + \beta^{(m)})^2}. \quad (4.20)$$

Ainsi, les dérivées première et seconde de la fonction χ_m sont données par

$$\dot{\chi}_m(y) = \frac{(y - z^{(m)})(\alpha^{(m)}y + \alpha^{(m)}z^{(m)} + 2\beta^{(m)})}{2(\alpha^{(m)}y + \beta^{(m)})^2} + \frac{\alpha^{(m)}}{2(\alpha^{(m)}y + \beta^{(m)})}, \quad (4.21)$$

$$\ddot{\chi}_m(y) = \frac{(\alpha^{(m)}z^{(m)} + \beta^{(m)})^2}{(\alpha^{(m)}y + \beta^{(m)})^3} - \frac{1}{2} \left(\frac{\alpha^{(m)}}{\alpha^{(m)}y + \beta^{(m)}} \right)^2. \quad (4.22)$$

Convexité. D'une part, pour tout $m \in \{1, \dots, M\}$, d'après (4.19), on peut remarquer que la fonction $\ddot{\phi}_m$ est positive sur $[0, +\infty[$, et donc ϕ_m est convexe sur $[0, +\infty[$. Par conséquent, la fonction h_1 est convexe sur $[0, +\infty[^N$.

D'autre part, pour tout $m \in \{1, \dots, M\}$, d'après (4.20), la fonction $\ddot{\psi}_m$ est négative sur $[0, +\infty[$, et donc ψ_m est concave sur $[0, +\infty[$. Par conséquent, la fonction h_1 est concave sur $[0, +\infty[^N$.

Ainsi, la fonction h est une différence de fonctions convexes sur $[0, +\infty[^N$. Dans la proposition suivante, nous allons étudier plus en détails la courbure de la fonction χ_m .

Proposition 4.3. Soit $m \in \{1, \dots, M\}$. Posons

$$y_{\max}^{(m)} = \frac{1}{\alpha^{(m)}} \left(2 \left(\frac{\alpha^{(m)}z^{(m)} + \beta^{(m)}}{\alpha^{(m)}} \right)^2 - \beta^{(m)} \right), \quad (4.23)$$

où $\alpha^{(m)}, \beta^{(m)}$ sont définis dans (4.10) et $z^{(m)}$ correspond au m -ème élément de l'image observée $\mathbf{z}$ définie par (4.11). La fonction χ_m définie par (4.17) est convexe sur $[0, y_{\max}^{(m)}]$ et concave sur $]y_{\max}^{(m)}, +\infty[$.

1. On considèrera la dérivée à droite lorsque $y = 0$.

Démonstration. Soit $m \in \{1, \dots, M\}$. D'après (4.22), la dérivée seconde de χ_m est donnée par

$$(\forall y \in [0, +\infty[) \quad \ddot{\chi}_m(y) = \frac{2(\alpha^{(m)}z^{(m)} + \beta^{(m)})^2 - (\alpha^{(m)})^2(\alpha^{(m)}y + \beta^{(m)})}{2(\alpha^{(m)}y + \beta^{(m)})^3}.$$

La fonction χ_m est donc convexe pour tout $y \in [0, +\infty[$ satisfaisant $\ddot{\chi}_m(y) \geq 0$, ce qui est équivalent à

$$2 \left(\frac{\alpha^{(m)}z^{(m)} + \beta^{(m)}}{\alpha^{(m)}} \right)^2 \geq \alpha^{(m)}y + \beta^{(m)}.$$

Cette inégalité est vérifiée lorsque $y \leq y_{\max}^{(m)}$. □

Constantes de Lipschitz. Comme nous l'avons vu dans les chapitres précédents, une propriété intéressante permettant de construire des algorithmes d'optimisation est la propriété de gradient Lipschitz. C'est ce que nous allons étudier dans ce paragraphe.

Proposition 4.4. *Soit $m \in \{1, \dots, M\}$. La dérivée $\dot{\chi}_m$ donnée par (4.21) de χ_m est une fonction $\mu^{(m)}$ -Lipschitzienne sur $[0, +\infty[$, où $\mu^{(m)}$ est donnée par*

$$\mu^{(m)} = \max \left\{ \frac{1}{(\beta^{(m)})^2} \left(\frac{(\alpha^{(m)}z^{(m)} + \beta^{(m)})^2}{\beta^{(m)}} - \frac{(\alpha^{(m)})^2}{2} \right), \frac{1}{2} \frac{(\alpha^{(m)})^6}{27(\alpha^{(m)}z^{(m)} + \beta^{(m)})^4} \right\}. \quad (4.24)$$

La démonstration de cette proposition est donnée dans l'annexe 4.A. On peut en déduire la constante de Lipschitz du gradient de la fonction d'attache aux données h :

Proposition 4.5. *La fonction h est de gradient Lipschitz, de constante*

$$L = \|\mathbf{H}^\top \text{Diag}(\boldsymbol{\mu})\mathbf{H}\|, \quad (4.25)$$

où $\boldsymbol{\mu} = (\mu^{(m)})_{1 \leq m \leq M}$, avec, pour tout $m \in \{1, \dots, M\}$, $\mu^{(m)}$ la constante donnée par (4.24).

4.3 APPLICATION DE L'ALGORITHME EXPLICITE-IMPLICITE À MÉTRIQUE VARIABLE

4.3.1 Problème d'optimisation

Dans cette section, nous considérons un problème inverse où une version dégradée $\mathbf{z} = (z^{(m)})_{1 \leq m \leq M}$ d'une image originale $\bar{\mathbf{x}} \in [0, +\infty[^N$ est observée suivant le modèle (4.10)-(4.11), où $\mathbf{H} \in [0, +\infty[^M$. Notre objectif est de produire une estimation $\hat{\mathbf{x}} \in [0, +\infty[^N$ de l'image originale $\bar{\mathbf{x}}$ à partir des données observées $\mathbf{z}$.

Comme nous l'avons vu dans la proposition 4.2, $\bar{\mathbf{x}}$ peut être estimée en résolvant le problème de minimisation (4.13). Ce problème faisant intervenir la fonction différentiable non convexe h définie par (4.18), nous proposons une solution utilisant l'algorithme 3.1 explicite-implicite à métrique variable.

Remarque 4.2.

D'après le théorème 3.1, afin d'assurer la convergence de l'algorithme 3.1 (ou de sa version inexacte donnée par l'algorithme 3.2), la fonction

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad f(\mathbf{x}) = h_1(\mathbf{x}) + h_2(\mathbf{x}) + g(\mathbf{x}), \quad (4.26)$$

où h_1, h_2 sont définies par (4.14), doit satisfaire l'inégalité de Kurdyka-Łojasiewicz donnée dans la définition 2.20. Comme indiqué dans la remarque 2.11, ce n'est pas le cas de la fonction logarithme. Nous remplacerons donc, pour tout $m \in \{1, \dots, M\}$, la fonction ψ_m donnée dans (4.16) par la fonction

$$(\forall y \in [0, +\infty[) \quad \widehat{\psi}_m(y) = \frac{1}{2} \widehat{\log}(\alpha^{(m)}y + \beta^{(m)}). \quad (4.27)$$

Dans l'équation (4.27), $\widehat{\log}$ est une approximation semi-algébrique du logarithme définie sur $]0, +\infty[$, qui, comme la fonction originale, est concave et de gradient Lipschitz sur tout intervalle $[\underline{\beta}, +\infty[$, avec $\underline{\beta} \in]0, +\infty[$. Ce type d'approximation est souvent utilisée pour l'implémentation numérique de la fonction logarithme [Dahlquist et Bjorck, 2003, Chap.4]. Ainsi, la fonction h sera la somme de la fonction h_1 définie dans (4.14) et de la fonction

$$(\forall \mathbf{x} \in [0, +\infty[^N) \quad \widehat{h}_2(\mathbf{x}) = \sum_{m=1}^M \widehat{\psi}_m([\mathbf{H}\mathbf{x}]^{(m)}). \quad (4.28)$$

D'autre part, nous considérons une fonction de pénalisation hybride, composée de deux termes $g = g_1 + g_2$. Pour le premier terme, afin de contraindre la dynamique de l'image originale $\bar{\mathbf{x}}$, nous définissons

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad g_1(\mathbf{x}) = \iota_{[\mathbf{x}_{\min}, \mathbf{x}_{\max}]^N}(\mathbf{x}), \quad (4.29)$$

où $\mathbf{x}_{\min} \in [0, +\infty[$ et $\mathbf{x}_{\max} \in]\mathbf{x}_{\min}, +\infty[$ sont, respectivement, les valeurs minimales et maximales des composantes de $\bar{\mathbf{x}}$. Le second terme sera choisi, dans un premier exemple de simulation, comme étant une régularisation de type NLTV [Gilboa et Osher, 2008], puis dans un second exemple nous considérerons une régularisation parcimonieuse dans une approche de type trame à l'analyse [Chaâri et al., 2009; Elad et al., 2007; Pustelnik et al., 2012]. Comme nous l'avons vu dans la section 2.2.4.2, ces deux termes de régularisation s'écrivent

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad g_2(\mathbf{x}) = \sum_{j=1}^J g_{2,j}(\mathbf{F}_j \mathbf{x}), \quad (4.30)$$

où, pour tout $j \in \{1, \dots, J\}$, $g_{2,j}$ est une fonction semi-algébrique et $\mathbf{F}_j$ correspond soit à l'opérateur de gradient non local intervenant dans la NLTV, soit à un opérateur de trame ajusté. Ainsi, dans les deux cas, l'hypothèse 3.1 et l'inégalité de KL sont bien vérifiées.

4.3.2 Construction de la majorante

L'accélération des algorithmes 3.1 et 3.2 repose sur l'utilisation de majorantes quadratique approximant la fonction h autour de l'itérée courante, et dont les matrices de courbure $(\mathbf{A}_k(\mathbf{x}_k))_{k \in \mathbb{N}}$

sont simples à implémenter. Dans cette section, nous proposons une technique de majoration qui combine des stratégies proposées dans [Hunter et Lange, 2004] et [Erdogan et Fessler, 1999], afin de construire des matrices diagonales $(\mathbf{A}_k(\mathbf{x}_k))_{k \in \mathbb{N}}$ satisfaisant l'hypothèse 3.2.

Comme nous l'avons remarqué dans la section 4.2.2, sur son domaine de définition, h est la somme d'une fonction convexe h_1 définie par (4.14) et d'une fonction concave $\hat{h}_2$ définie par (4.28). Nous allons majorer ces deux termes séparément.

Pour tout $k \in \mathbb{N}$, soit $\mathbf{x}_k$ la k -ème itérée générée par l'algorithme 3.2. Une fonction majorant $\hat{h}_2$ sur $[0, +\infty[^N$ en $\mathbf{x}_k$ est

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad q_2(\mathbf{x}, \mathbf{x}_k) = \hat{h}_2(\mathbf{x}_k) + \langle \mathbf{x} - \mathbf{x}_k, \nabla \hat{h}_2(\mathbf{x}_k) \rangle, \quad (4.31)$$

où $\nabla \hat{h}_2(\mathbf{x}_k)$ désigne le gradient de $\hat{h}_2$ en $\mathbf{x}_k$.

Le lemme suivant nous permet de construire une majorante quadratique de la fonction h_1 en $\mathbf{x}_k$. Pour cela, nous avons besoin de définir, dans un premier temps, la fonction multivaluée

$$\Omega: [0, +\infty[^M \rightarrow \mathbb{R}^M: (y^{(m)})_{1 \leq m \leq M} \mapsto (\omega_m(y^{(m)}))_{1 \leq m \leq M},$$

où, pour tout $m \in \{1, \dots, M\}$,

$$(\forall y \in [0, +\infty[) \quad \omega_m(y) = \begin{cases} \ddot{\phi}_m(0), & \text{si } y = 0, \\ 2 \frac{\phi_m(0) - \phi_m(y) + y\dot{\phi}_m(y)}{y^2}, & \text{si } y > 0, \end{cases} \quad (4.32)$$

et ϕ_m est définie par (4.15).

Lemme 4.2. Soit h_1 définie par (4.14) où $\mathbf{H} = (\mathbf{H}^{(m,n)})_{1 \leq m \leq M, 1 \leq n \leq N} \in [0, +\infty[^{M \times N}$. Pour tout $k \in \mathbb{N}$, on pose

$$\mathbf{A}_k(\mathbf{x}_k) = \text{Diag}(\mathbf{P}^\top \Omega(\mathbf{H}\mathbf{x}_k)) + \varepsilon \mathbf{I}_N, \quad (4.33)$$

où $\varepsilon \in [0, +\infty[$ et $\mathbf{P} = (\mathbf{P}^{(m,n)})_{1 \leq m \leq M, 1 \leq n \leq N}$ est la matrice dont les éléments sont donnés par

$$(\forall m \in \{1, \dots, M\})(\forall n \in \{1, \dots, N\}) \quad \mathbf{P}^{(m,n)} = \mathbf{H}^{(m,n)} \sum_{p=1}^N \mathbf{H}^{(m,p)}. \quad (4.34)$$

La fonction q_1 définie par

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad q_1(\mathbf{x}, \mathbf{x}_k) = h_1(\mathbf{x}_k) + \langle \mathbf{x} - \mathbf{x}_k, \nabla h_1(\mathbf{x}_k) \rangle + \frac{1}{2} \|\mathbf{x} - \mathbf{x}_k\|_{\mathbf{A}_k(\mathbf{x}_k)}^2, \quad (4.35)$$

est alors une majorante quadratique de h_1 sur $[0, +\infty[^N$ en $\mathbf{x}_k$.

Démonstration. Pour tout $m \in \{1, \dots, M\}$, ϕ_m est convexe et infiniment dérivable sur $[0, +\infty[$.

On pose

$$(\forall (\mathbf{x}, \mathbf{x}') \in [0, +\infty[^2)$$

$$\rho_m(\mathbf{x}, \mathbf{x}') = \phi_m(\mathbf{x}') + (\mathbf{x} - \mathbf{x}')\dot{\phi}_m(\mathbf{x}') + \frac{1}{2}\omega_m(\mathbf{x}')(\mathbf{x} - \mathbf{x}')^2, \quad (4.36)$$

où les fonctions $\dot{\phi}_m$ et ω_m sont définies, respectivement, par (4.15) et (4.32).

Si $\alpha^{(m)} = 0$, alors ϕ_m est une fonction quadratique et on a

$$(\forall x, x') \in [0, +\infty[^2] \quad \phi_m(x) = \rho_m(x, x').$$

Supposons maintenant que $\alpha^{(m)} \in]0, +\infty[$. La dérivée troisième de ϕ_m est donnée par

$$(\forall x \in [0, +\infty[) \quad \ddot{\phi}_m(x) = -3\alpha^{(m)} \frac{(\alpha^{(m)}z^{(m)} + \beta^{(m)})^2}{(\alpha^{(m)}x + \beta^{(m)})^4},$$

et est strictement négative. Il apparaît alors que $\dot{\phi}_m$ est une fonction strictement concave sur $[0, +\infty[$. En utilisant une version simplifiée de la démonstration donnée dans [Erdogan et Fessler, 1999, Annexe B], nous allons démontrer la positivité de la fonction $\rho_m(\cdot, x') - \phi_m$ sur $[0, +\infty[$, pour tout $x' \in]0, +\infty[$. La dérivée seconde de cette dernière fonction est donnée par

$$(\forall x \in [0, +\infty[) \quad \ddot{\rho}_m(x, x') - \ddot{\phi}_m(x) = \omega_m(x') - \ddot{\phi}_m(x). \quad (4.37)$$

De plus, d'après la formule de Taylor de second ordre, il existe $\tilde{x} \in]0, x'[$ tel que

$$\phi_m(0) = \phi_m(x') - x' \dot{\phi}_m(x') + \frac{1}{2} x'^2 \ddot{\phi}_m(\tilde{x}).$$

Ainsi, en utilisant les équations (4.32) et (4.37), on obtient

$$(\forall x \in [0, +\infty[) \quad \ddot{\rho}_m(x, x') - \ddot{\phi}_m(x) = \ddot{\phi}_m(\tilde{x}) - \ddot{\phi}_m(x). \quad (4.38)$$

Or, puisque $\ddot{\phi}_m$ est strictement négative, $\ddot{\phi}_m$ est strictement décroissante, et (4.38) implique que $\dot{\rho}_m(\cdot, x') - \dot{\phi}_m$ est d'abord strictement décroissante sur $]0, \tilde{u}[$, puis strictement croissante sur $]\tilde{u}, +\infty[$. D'une part, (4.32) et (4.36) nous donnent

$$\begin{cases} \rho_m(x', u') - \phi_m(u') &= 0, \\ \rho_m(0, u') - \phi_m(0) &= 0. \end{cases} \quad (4.39)$$

Ainsi, d'après le théorème des valeurs intermédiaires, il existe $x^* \in]0, x'[$ tel que $\dot{\rho}_m(x^*, x') - \dot{\phi}_m(x^*) = 0$. D'autre part, d'après l'équation (4.36), on a

$$\dot{\rho}_m(x', x') - \dot{\phi}_m(x') = 0.$$

Par conséquent, par monotonie de la fonction $\dot{\rho}_m(\cdot, x') - \dot{\rho}_1^{(m)}$, on peut déduire que x^* est l'unique zéro de cette fonction sur $]0, u'[$, et

$$\begin{cases} \dot{\rho}_m(x, x') - \dot{\phi}_m(x') &> 0, & \forall x \in [0, x^*[, \\ \dot{\rho}_m(x, x') - \dot{\phi}_m(x') &< 0, & \forall x \in]x^*, x'[, \\ \dot{\rho}_m(x, x') - \dot{\phi}_m(x') &> 0, & \forall x \in]x', +\infty[. \end{cases} \quad (4.40)$$

L'équation (4.40) implique que $\rho_m(\cdot, \mathbf{x}') - \phi_m(\cdot)$ est strictement croissante sur $[0, \mathbf{x}^*]$, strictement décroissante sur $]\mathbf{x}^*, \mathbf{x}'[$ et strictement croissante sur $]\mathbf{x}', +\infty[$. Par conséquent, d'après (4.39),

$$(\forall (\mathbf{x}, \mathbf{x}') \in [0, +\infty[\times]0, +\infty[) \quad \phi_m(\mathbf{x}) \leq \rho_m(\mathbf{x}, \mathbf{x}'). \quad (4.41)$$

De plus, d'après l'expression de $\ddot{\phi}_m$ dans (4.19), on peut en déduire qu'elle est maximale en zéro, et les équations (4.32) et (4.36) montrent que

$$(\forall \mathbf{x} \in [0, +\infty[) \quad \phi_m(\mathbf{x}) \leq \rho_m(\mathbf{x}, 0). \quad (4.42)$$

Ainsi, en combinant les équations (4.41) et (4.42), on obtient

$$(\forall (\mathbf{x}, \mathbf{x}') \in [0, +\infty[^2) \quad \phi_m(\mathbf{x}) \leq \rho_m(\mathbf{x}, \mathbf{x}'), \quad (4.43)$$

qui, comme nous l'avons souligné précédemment, reste satisfaite lorsque $\alpha^{(m)} = 0$.

La majoration donnée dans (4.43) implique que, pour tout $k \in \mathbb{N}$,

$$\begin{aligned} (\forall \mathbf{x} \in [0, +\infty[^N) \quad h_1(\mathbf{x}) &\leq h_1(\mathbf{x}_k) + \langle \mathbf{x} - \mathbf{x}_k, \nabla h_1(\mathbf{x}_k) \rangle \\ &\quad + \frac{1}{2} \langle \mathbf{H}\mathbf{x} - \mathbf{H}\mathbf{x}_k, \text{Diag}(\Omega(\mathbf{H}\mathbf{x}_k))(\mathbf{H}\mathbf{x} - \mathbf{H}\mathbf{x}_k) \rangle. \end{aligned}$$

On définit la matrice $(\Lambda^{(m,n)})_{1 \leq m \leq M, 1 \leq n \leq N} \in [0, +\infty[^{M \times N}$ telle que

$$(\forall (m, n) \in \{1, \dots, M\} \times \{1, \dots, N\}) \quad \Lambda^{(m,n)} = \begin{cases} 0, & \text{si } H^{(m,n)} = 0, \\ \frac{H^{(m,n)}}{\sum_{p=1}^N H^{(m,p)}}, & \text{si } H^{(m,n)} > 0. \end{cases}$$

D'après l'inégalité de Jensen et (4.34), pour tout $m \in \{1, \dots, M\}$ et pour tout $\mathbf{x} \in [0, +\infty[^N$,

$$([\Lambda^{(m,1)}, \dots, \Lambda^{(m,N)}](\mathbf{x} - \mathbf{x}_k))^2 \leq \sum_{n=1}^N \Lambda^{(m,n)} (\mathbf{x}^{(n)} - \mathbf{x}_k^{(n)})^2,$$

ce qui est équivalent à

$$([H^{(m,1)}, \dots, H^{(m,N)}](\mathbf{x} - \mathbf{x}_k))^2 \leq \sum_{n=1}^N P^{(m,n)} (\mathbf{x}^{(n)} - \mathbf{x}_k^{(n)})^2.$$

Puisque la convexité de ϕ_m , pour tout $m \in \{1, \dots, M\}$, implique la positivité de $\omega^{(m)}$ sur $[0, +\infty[$, on peut en déduire que

$$\langle \mathbf{H}\mathbf{x} - \mathbf{H}\mathbf{x}_k, \text{Diag}(\Omega(\mathbf{H}\mathbf{x}_k))(\mathbf{H}\mathbf{x} - \mathbf{H}\mathbf{x}_k) \rangle \leq \langle \mathbf{x} - \mathbf{x}_k, \text{Diag}(\mathbf{P}^\top \Omega(\mathbf{H}\mathbf{x}_k))(\mathbf{x} - \mathbf{x}_k) \rangle.$$

La fonction $q_1(\cdot, \mathbf{x}_k)$ définie par (4.35) est donc une majorante quadratique de h_1 sur $[0, +\infty[^N$ en $\mathbf{x}_k$. $\square$

Nous pouvons donc déduire de l'équation (4.31) et du lemme 4.2 que l'hypothèse 3.2(i) est satisfaite pour $\mathbf{q} = \mathbf{q}_1 + \mathbf{q}_2$.

Notons que, pour tout $m \in \{1, \dots, M\}$, la dérivée de ω_m en $\mathbf{x}' \in]0, +\infty[$ est donnée par

$$\begin{aligned}\dot{\omega}_m(\mathbf{x}') &= 2 \frac{\mathbf{x}'^2 \ddot{\phi}_m(\mathbf{x}') - 2(\phi_m(0) - \phi_m(\mathbf{x}') + \mathbf{x}' \dot{\phi}_m(\mathbf{x}'))}{\mathbf{x}'^3} \\ &= 2 \frac{\ddot{\phi}_m(\mathbf{x}') - \ddot{\phi}_m(\tilde{\mathbf{x}})}{\mathbf{x}'},\end{aligned}$$

où $\tilde{\mathbf{x}} \in]0, \mathbf{x}'[$. Or, la dérivée troisième de ϕ_m étant négative, ω_m est une fonction positive strictement décroissante. De plus, d'après l'expression de la fonction $\mathbf{g}$, $\text{dom } \mathbf{g} = [\mathbf{x}_{\min}, \mathbf{x}_{\max}]^N$, donc son image $\{\mathbf{H}\mathbf{x} \mid \mathbf{x} \in \text{dom } \mathbf{g}\}$ est un ensemble compact. Ainsi, pour tout $m \in \{1, \dots, M\}$ et pour tout $\mathbf{x} = (\mathbf{x}^{(m)})_{1 \leq m \leq M}$, la fonction $\mathbf{x}^{(m)} \mapsto \omega_m([\mathbf{H}\mathbf{x}]^{(m)})$ admet un minimum noté $\underline{\omega}_m > 0$ sur $\text{dom } \mathbf{g}$. Ainsi, l'hypothèse 3.2(ii) est satisfaite pour les matrices $(\mathbf{A}_k(\mathbf{x}_k))_{k \in \mathbb{N}}$ définies par (4.33) avec

$$\begin{cases} \underline{\nu} = \varepsilon + \min_{1 \leq n \leq N} \sum_{m=1}^M \mathbf{P}^{(m,n)} \underline{\omega}_m, \\ \overline{\nu} = \varepsilon + \max_{1 \leq n \leq N} \sum_{m=1}^M \mathbf{P}^{(m,n)} \ddot{\phi}_m(0). \end{cases} \quad (4.44)$$

Remarque 4.3.

Si aucune des colonnes de la matrice $\mathbf{H}$ n'est nulle, nous pouvons choisir $\varepsilon = 0$ dans l'équation (4.44). Dans le cas contraire, il faut choisir $\varepsilon > 0$.

4.3.3 Implémentation de l'étape proximale

Dans la section suivante nous donnerons deux exemples de simulations pour lesquelles, comme nous l'avons indiqué, nous utiliserons deux régularisations différentes. La première sera une régularisation NLTV et l'autre une régularisation de type trame à l'analyse. Dans les deux cas, à chaque itération $k \in \mathbb{N}$, il nous faut calculer la variable $\mathbf{y}_k$ de l'algorithme 3.1 correspondant à l'opérateur proximal de la fonction $\mathbf{g}$ au point $\tilde{\mathbf{x}}_k = \mathbf{x}_k - \gamma_k \mathbf{A}_k(\mathbf{x}_k) \nabla \mathbf{h}(\mathbf{x}_k)$, relativement à la métrique induite par $\gamma_k^{-1} \mathbf{A}_k(\mathbf{x}_k)$, i.e.

$$\mathbf{y}_k = \underset{\mathbf{x} \in \mathbb{R}^N}{\text{argmin}} \left\{ \sum_{j=1}^J \mathbf{g}_{2,j}(\mathbf{F}_j \mathbf{x}) + \iota_{[\mathbf{x}_{\min}, \mathbf{x}_{\max}]^N}(\mathbf{x}) + \frac{1}{2} \|\mathbf{x} - \tilde{\mathbf{x}}_k\|_{\gamma_k^{-1} \mathbf{A}_k(\mathbf{x}_k)}^2 \right\}. \quad (4.45)$$

En raison de la présence des opérateurs linéaires $(\mathbf{F}_j)_{1 \leq j \leq J}$, il n'est pas possible d'obtenir une expression explicite de $\mathbf{y}_k$. Cependant, dans le chapitre 3, nous avons montré que la convergence de l'algorithme 3.1 reste assurée lorsque l'opérateur proximal est calculé de façon approchée. Ainsi, nous calculerons (4.45) à l'aide de sous-itérations. Pour cela, remarquons que l'équation (4.45) est équivalente à

$$\begin{aligned}\mathbf{y}_k &= \gamma_k^{1/2} \mathbf{A}_k(\mathbf{x}_k)^{-1/2} \underset{\mathbf{x} \in \mathbb{R}^N}{\text{argmin}} \left\{ \frac{1}{2} \|\mathbf{x} - \gamma_k^{-1/2} \mathbf{A}_k(\mathbf{x}_k)^{1/2} \tilde{\mathbf{x}}_k\|^2 \right. \\ &\quad \left. + \gamma_k^{1/2} \sum_{j=1}^J \mathbf{g}_{2,j}(\mathbf{F}_j \mathbf{A}_k(\mathbf{x}_k)^{-1/2} \mathbf{x}) + \iota_{[\mathbf{x}_{\min}, \mathbf{x}_{\max}]^N}(\gamma_k^{1/2} \mathbf{A}_k(\mathbf{x}_k)^{-1/2} \mathbf{x}) \right\}.\end{aligned}$$

Il existe divers algorithmes permettant de résoudre ce problème d'optimisation. Les résultats numériques présentés dans la section suivante ont été obtenus en utilisant l'algorithme explicite-implicite dual de [Combettes *et al.*, 2011].

4.4 RÉSULTATS DE SIMULATION

Nous allons maintenant évaluer les performances de notre algorithme sur deux exemples de simulations. Dans le premier exemple, nous traiterons une application en restauration d'image et, dans le second exemple, une application en reconstruction d'image en tomographie.

4.4.1 Restauration d'image

Dans le premier scénario considéré, l'image originale $\bar{\mathbf{x}}$ correspond à l'image standard *Peppers* de taille 256×256 (<http://sipi.usc.edu/database/>). Cette image est dégradée par un opérateur de flou $\mathbf{H}$ dont le noyau correspond à un flou gaussien d'écart type 1 et de taille 7×7 . De plus, l'image résultante $\mathbf{H}\bar{\mathbf{x}}$ est bruitée par un bruit gaussien dépendant du signal tel que décrit par le modèle (4.11) où, pour tout $m \in \{1, \dots, M\}$, $\alpha^{(m)} = 0.5$ et $\beta^{(m)} = 1$. Dans ce cas, $\mathbf{x}_{\min} = 0$, $\mathbf{x}_{\max} = 255$, et la constante de Lipschitz de $\nabla \mathbf{h}$ est égale à $L = 18668$.

Dans cette application, on considère comme terme de régularisation la NLTV [Gilboa et Osher, 2008] (voir aussi le paragraphe sur la TV dans la section 2.2.4.2). Ainsi, la fonction $\mathbf{g}_2$ est donnée par (4.30) avec $J = N$ et

$$(\forall n \in \{1, \dots, N\})(\forall \mathbf{x} \in \mathbb{R}^N) \quad \begin{cases} \mathbf{g}_{2,n}(\mathbf{x}) = \theta \ell_2(\mathbf{F}_n \mathbf{x}), \\ \mathbf{F}_n = [\omega_n^{(1)} \mathbf{F}_n^{(1)}, \dots, \omega_n^{(J_n)} \mathbf{F}_n^{(J_n)}]^\top, \end{cases} \quad (4.46)$$

où $\theta \in]0, +\infty[$ est choisi de façon à maximiser le RSB entre $\bar{\mathbf{x}}$ et $\hat{\mathbf{x}}$ défini par (2.20), $(\omega_n^{(j)})_{1 \leq j \leq J_n} \in \mathbb{R}^{J_n}$ sont des poids déterminés par un prétraitement [Bresson, 2009; Jezierska *et al.*, 2013], J_n détermine le nombre de voisins non locaux considérées au n -ème pixel, et ℓ_2 correspond à la norme euclidienne.

La figure 4.2(a) représente l'image originale et la figure 4.2(b) l'image dégradée correspondant au modèle décrit ci-dessus, dont le RSB est de 21.85 dB. Nous présentons en figure 4.2(c) l'image restaurée en utilisant la méthode proposée. Le RSB est égal à 27.11 dB.

Afin de montrer l'intérêt d'utiliser une métrique variable, nous proposons de comparer l'algorithme 3.1 avec sa version non préconditionnée donnée par l'algorithme 2.4 et une version visant à l'accélérer appelée FISTA (*Fast Iterative Shrinkage-Thresholding Algorithm*) développée par [Beck et Teboulle, 2009]. Notons que, bien que peu de résultats théoriques existent concernant la convergence de FISTA dans le cadre non convexe, nous n'avons pas observé de divergence de l'algorithme lors de nos simulations. Pour les algorithmes 3.1 et 2.4 nous avons choisi $\lambda_k \equiv 1$, et deux valeurs du pas $(\gamma_k)_{k \in \mathbb{N}}$ ont été testées : $\gamma_k \equiv 1$ et $\gamma_k \equiv 1.9$. Remarquons que pour ces valeurs de $(\lambda_k)_{k \in \mathbb{N}}$ et $(\gamma_k)_{k \in \mathbb{N}}$, l'hypothèse 3.3 est satisfaite. De plus, d'après la remarque 3.4(ii), l'hypothèse 3.4 est aussi vérifiée.

La figure 4.3 illustre les variations de $(f(\mathbf{x}_k) - f^*)_k$ et $(\|\mathbf{x}_k - \hat{\mathbf{x}}\|)_k$ en fonction du temps de calcul, obtenues par les algorithmes 3.1, 2.4 et FISTA, pour des tests réalisés sur une machine Intel(R)

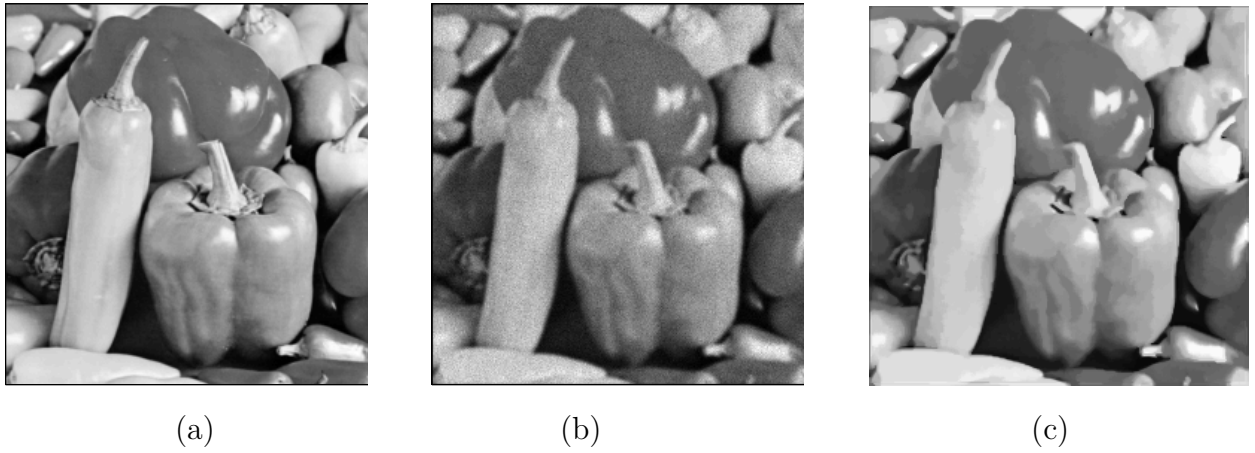


Figure 4.2 – Restoration : (a) Image originale Peppers. (b) Image dégradée par un bruit gaussien dépendant du signal, $RSB=21.85$ dB. (c) Image reconstruite par la méthode proposée, $RSB=27.11$ dB.

Xeon(R) CPU X5570 cadencé à 2.93GHz, implémenté à l’aide de Matlab 7. La solution optimale $\hat{\mathbf{x}}$ a été calculée au préalable, pour chacun des algorithmes, avec un grand nombre d’itérations. D’autre part, f^* représente le minimum des (possibles) différentes valeurs $f(\hat{\mathbf{x}})$ obtenues pour chaque algorithme. Nous pouvons remarquer que notre méthode est la plus rapide dans cet exemple, autant pour la convergence du critère que pour la convergence des itérées.

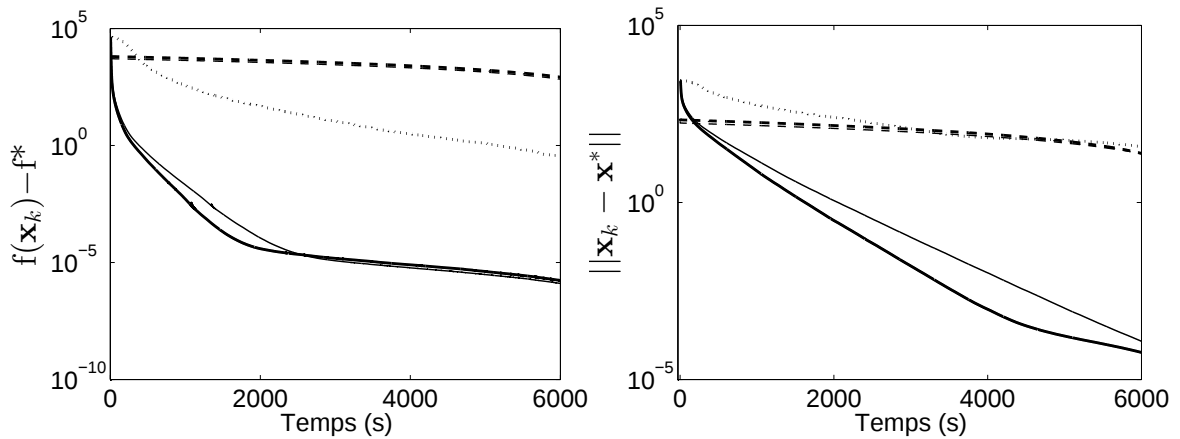


Figure 4.3 – Restoration : Comparaison de la vitesse de convergence de l’algorithme 3.1 explicite-implicite à métrique variable avec $\gamma_k \equiv 1$ (lignes fines continues) et $\gamma_k \equiv 1.9$ (lignes épaisses continues), de l’algorithme 2.4 explicite-implicite avec $\gamma_k \equiv 1$ (lignes fines discontinues) et $\gamma_k \equiv 1.9$ (lignes épaisses discontinues), et de FISTA (lignes pointillés).

4.4.2 Reconstruction d'image

Dans ce second scénario, $\bar{\mathbf{x}}$ correspond à une coupe du fantôme *Zubal* [Zubal *et al.*, 1994] de dimension 128×128 , et $\mathbf{H}$ modélise $M = 16384$ projections parallèles à partir d'acquisitions de 128 lignes et 128 angles. Ce type de modèle se rencontre, par exemple, en tomographie par rayons X [Slaney et Kak, 1988]. Le vecteur $\mathbf{H}\bar{\mathbf{x}}$ est ensuite bruité par un bruit gaussien dépendant du signal d'après le modèle (4.11) où, pour tout $m \in \{1, \dots, M\}$, $\alpha^{(m)} = 0.01$ et $\beta^{(m)} = 0.1$. Dans cette configuration, $x_{\min} = 0$, $x_{\max} = 1$, et la constante de Lipschitz de ∇h est égale à $L = 6.0 \times 10^6$.

Nous utilisons une approche à l'analyse pour le terme de régularisation g_2 donné par (4.30) avec

$$(\forall j \in \{1, \dots, J\}) \quad \begin{cases} g_{2,j} = \theta^{(j)} | \cdot |, \\ \mathbf{F}_j = [\mathbf{W} \cdot]^{(j)}, \end{cases} \quad (4.47)$$

où $(\theta^{(j)})_{1 \leq j \leq J} \in]0, +\infty[^J$ et $\mathbf{W}$ correspond à une transformation redondante en ondelettes de Daubechies d'ordre 8 sur 3 niveaux de résolution. Cette décomposition correspond à une trame à l'analyse ajustée (*tight frame*) de constante égale à 64. Les paramètres $(\theta^{(j)})_{1 \leq j \leq J}$ sont choisis de façon à maximiser le RSB entre $\bar{\mathbf{x}}$ et $\hat{\mathbf{x}}$.

La figure 4.4(a) représente l'image originale, (b) les projections bruitées correspondantes, (c) l'image reconstruite par la méthode standard de rétro-projection filtrée [Slaney et Kak, 1988], et (d) l'image reconstruite par la méthode proposée.

Comme dans l'exemple précédent, nous comparons la vitesse de convergence de la méthode que nous proposons, donnée par l'algorithme 3.1, avec sa version non préconditionnée, donnée dans l'algorithme 2.4, ainsi qu'avec FISTA. Nous choisissons le paramètre de relaxation λ_k égal à 1 pour tout $k \in \mathbb{N}$, et nous testons deux valeurs différentes pour le pas γ_k : $\gamma_k \equiv 1$ et $\gamma_k \equiv 1.9$.

La figure 4.5 illustre les variations de $(f(\mathbf{x}_k) - f^*)_k$ et $(\|\mathbf{x}_k - \hat{\mathbf{x}}\|)_k$ en fonction du temps de calcul, en utilisant les trois algorithmes. Les solutions optimales $\hat{\mathbf{x}}$ et f^* sont calculées de la même façon que dans la section précédente. Dans cet exemple, nous pouvons observer que FISTA et l'algorithme 2.4 ont des comportements similaires. Cependant, contrairement à l'exemple précédent, le choix du pas $(\gamma_k)_{k \in \mathbb{N}}$ influe de façon sensible sur la vitesse de convergence des algorithmes 3.1 et 2.4. Prendre $(\gamma_k)_{k \in \mathbb{N}}$ le plus proche possible de sa borne supérieure permet ici d'accélérer la convergence des critères et des itérées.

4.5 CONCLUSION

Dans ce chapitre nous nous sommes intéressés à la restauration d'une image dégradée par un bruit gaussien dépendant du signal. Pour cela nous avons étudié dans un premier temps les propriétés du terme d'attache aux données issu du MAP associé à ce modèle de bruit. En particulier, nous avons démontré la non-convexité et le caractère gradient Lipschitz de ce critère.

Dans un deuxième temps, nous avons traité deux exemples applicatifs en traitement des images. Le premier était une application en restauration et le second un problème de reconstruction en tomographie. Dans les deux cas, les images étaient dégradées par un opérateur linéaire puis bruitées par un bruit gaussien dépendant. Nous avons utilisé l'algorithme 3.1 explicite-implicite à métrique variable proposé dans le chapitre 3 pour trouver des solutions de ces problèmes inverses.

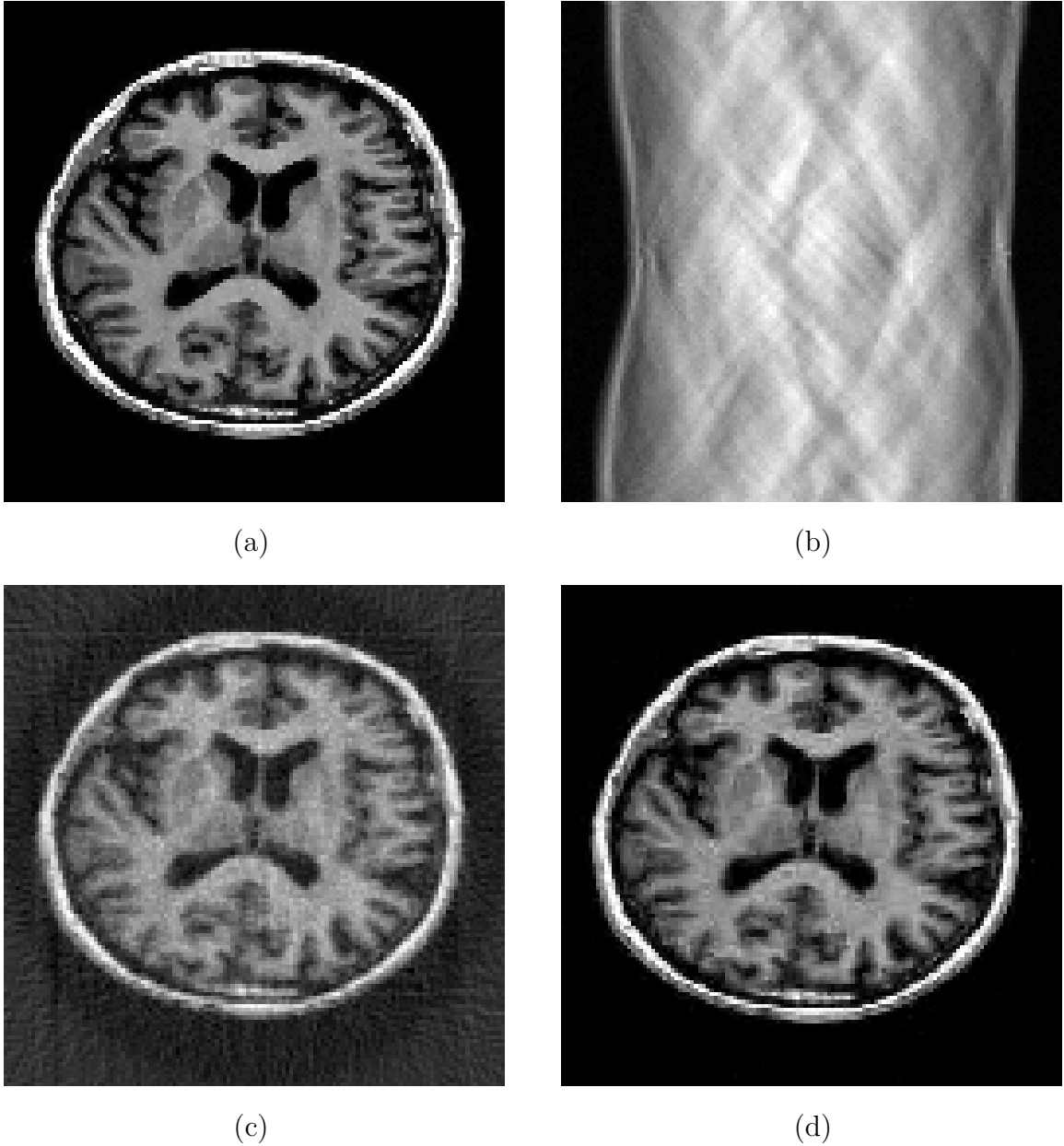


Figure 4.4 – *Reconstruction* : (a) *Image originale Zubal*. (b) *Image dégradée par un bruit gaussien dépendant du signal*. (c) *Image reconstruite par rétro-projection filtrée, $RSB=7$ dB*. (d) *Image reconstruite par la méthode proposée, $RSB=18.9$ dB*.

Des simulations ont montré que notre algorithme se compare très favorablement à l'algorithme explicite-implicite usuel et à ses formes accélérées telles que FISTA. En effet, nous avons vu que l'utilisation d'une métrique variable permet d'accélérer fortement la convergence du critère et des itérées.

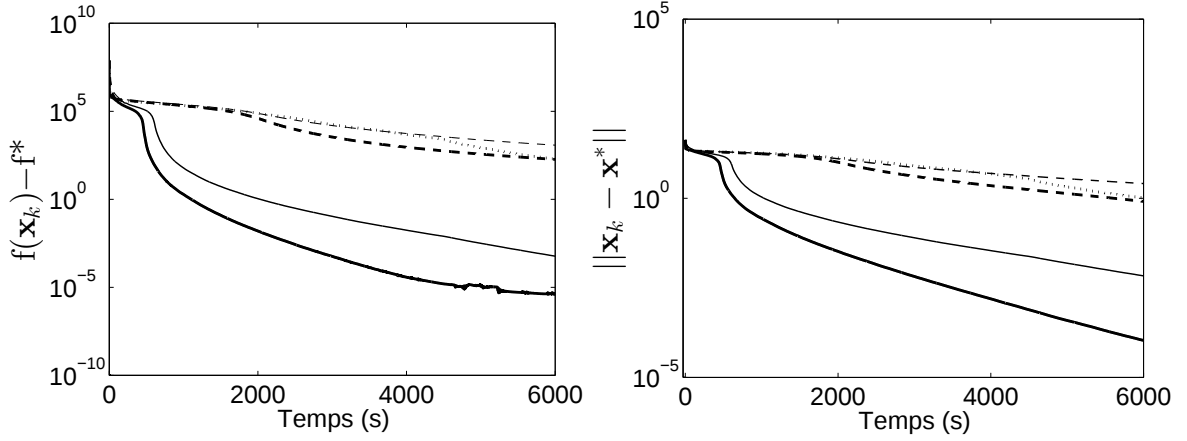


Figure 4.5 – Reconstruction : Comparaison de la vitesse de convergence de l’algorithme 3.1 explicite-implicite à métrique variable avec $\gamma_k \equiv 1$ (lignes fines continues) et $\gamma_k \equiv 1.9$ (lignes épaisses continues), de l’algorithme 2.4 explicite-implicite avec $\gamma_k \equiv 1$ (lignes fines discontinues) et $\gamma_k \equiv 1.9$ (lignes épaisses discontinues), et de FISTA (lignes pointillés).

Dans le chapitre 5, nous proposerons de combiner cette méthode d’accélération à une méthode de minimisation par bloc afin de réduire le coût de calcul par itération lorsque l’on traite des données de grande taille.

4.A DÉMONSTRATION DE LA PROPOSITION 4.4

Soit $m \in \{1, \dots, M\}$. Rappelons que la constante de Lipschitz de $\dot{\chi}_m$ sur $[0, +\infty[$ est définie par

$$\mu^{(m)} = \sup_{y \in [0, +\infty[} |\ddot{\chi}_m(y)|.$$

Or, d’après la proposition 4.3, la fonction χ_m est convexe sur $[0, y_{\max}^{(m)}]$ et concave sur $]y_{\max}^{(m)}, +\infty[$. Par conséquent, sa dérivée seconde vérifie $\ddot{\chi}_m(y) \geq 0$ pour tout $y \in [0, y_{\max}^{(m)}]$, et $\ddot{\chi}_m(y) \leq 0$ pour tout $y \in]y_{\max}^{(m)}, +\infty[$. Ainsi, on a

$$\mu^{(m)} = \max\{\mu_1^{(m)}, \mu_2^{(m)}\}, \quad (4.48)$$

où $\mu_1^{(m)}$ (resp. $\mu_2^{(m)}$) est la constante de Lipschitz de $\dot{\chi}_m$ sur $[0, y_{\max}^{(m)}[$ (resp. sur $]y_{\max}^{(m)}, +\infty[$) :

$$\begin{aligned} \mu_1^{(m)} &= \sup_{y \in [0, y_{\max}^{(m)}[} \ddot{\chi}_m(y) \\ \mu_2^{(m)} &= \sup_{y \in]y_{\max}^{(m)}, +\infty[} -\ddot{\chi}_m(y) = - \inf_{y \in]y_{\max}^{(m)}, +\infty[} \ddot{\chi}_m(y). \end{aligned}$$

Nous allons donc étudier les variations de la fonction $y \in [0, +\infty[\mapsto \ddot{\chi}_m(y)$. La dérivée de $\ddot{\chi}_m$ est donnée par

$$(\forall y \in [0, +\infty[) \quad \ddot{\chi}_m(y) = -3\alpha^{(m)} \frac{(\alpha^{(m)}z^{(m)} + \beta^{(m)})^2}{(\alpha^{(m)}y + \beta^{(m)})^4} + \left(\frac{\alpha^{(m)}}{\alpha^{(m)}y + \beta^{(m)}} \right)^3.$$

On peut donc en déduire que la fonction $y \in [0, +\infty[\mapsto \ddot{\chi}_m(y)$ est décroissante sur $[0, \theta^{(m)}]$, et croissante sur $[\theta^{(m)}, +\infty[$, avec

$$\theta^{(m)} = \frac{1}{\alpha^{(m)}} \left(3 \left(\frac{\alpha^{(m)}z^{(m)} + \beta^{(m)}}{\alpha^{(m)}} \right)^2 - \beta^{(m)} \right). \quad (4.49)$$

(i) Remarquons d'une part que

$$\begin{cases} \ddot{\chi}_m(0) = \frac{1}{(\beta^{(m)})^2} \left(\frac{(\alpha^{(m)}z^{(m)} + \beta^{(m)})^2}{\beta^{(m)}} - \frac{(\alpha^{(m)})^2}{2} \right) \\ \lim_{y \rightarrow +\infty} \ddot{\chi}_m(y) = 0. \end{cases}$$

Donc la borne supérieure de $\ddot{\chi}_m$ sur $[0, +\infty[$ (et donc sur $[0, y_{\max}^{(m)}]$) est atteinte en 0, et on a alors

$$\mu_1^{(m)} = \ddot{\chi}_m(0) = \frac{1}{(\beta^{(m)})^2} \left(\frac{(\alpha^{(m)}z^{(m)} + \beta^{(m)})^2}{\beta^{(m)}} - \frac{(\alpha^{(m)})^2}{2} \right). \quad (4.50)$$

(ii) D'autre part, remarquons que $\theta^{(m)} \geq y_{\max}^{(m)}$. Ainsi, la borne inférieure de $\ddot{\chi}_m$ est atteinte en $\theta^{(m)}$, et, d'après (4.22), on a alors

$$\begin{aligned} \mu_2^{(m)} &= -\ddot{\chi}_m(\theta^{(m)}) \\ &= -\frac{2(\alpha^{(m)}z^{(m)} + \beta^{(m)})^2 - (\alpha^{(m)})^2(\alpha^{(m)}\theta^{(m)} + \beta^{(m)})}{2(\alpha^{(m)}\theta^{(m)} + \beta^{(m)})^3}. \end{aligned}$$

Or, d'après (4.49), on a

$$\alpha^{(m)}\theta^{(m)} + \beta^{(m)} = 3 \left(\frac{\alpha^{(m)}z^{(m)} + \beta^{(m)}}{\alpha^{(m)}} \right)^2.$$

En combinant les deux dernières équations, on obtient

$$\begin{aligned} \mu_2^{(m)} &= \frac{3(\alpha^{(m)}z^{(m)} + \beta^{(m)})^2 - 2(\alpha^{(m)}z^{(m)} + \beta^{(m)})^2}{2 \left(3 \left(\frac{\alpha^{(m)}z^{(m)} + \beta^{(m)}}{\alpha^{(m)}} \right)^2 \right)^3} \\ &= \frac{1}{2} \frac{(\alpha^{(m)})^6}{27(\alpha^{(m)}z^{(m)} + \beta^{(m)})^4}. \end{aligned}$$

CHAPITRE 5

Algorithme explicite-implicite à métrique variable alterné par bloc

Sommaire

5.1	Introduction	94
5.2	Problème de minimisation	94
5.3	État de l'art	96
5.3.1	Algorithme de minimisation alternée par bloc	96
5.3.2	Algorithme proximal alterné par bloc	97
5.3.3	Algorithme explicite-implicite alterné par bloc	98
5.4	Méthode proposée	99
5.4.1	Algorithme explicite-implicite à métrique variable alterné par bloc	99
5.4.2	Hypothèses	100
5.4.3	Version inexacte de l'algorithme explicite-implicite à métrique variable alterné par bloc	102
5.5	Analyse de convergence	103
5.5.1	Propriétés de descente	103
5.5.2	Théorème de convergence	105
5.5.3	Borne supérieure sur la vitesse de convergence	108
5.6	Conclusion	109
Annexe 5.A	Démonstration du théorème 5.1	110
Annexe 5.B	Démonstration du théorème 5.2	113

5.1 INTRODUCTION

Dans ce chapitre, nous considérons un problème d'optimisation consistant à minimiser un critère $f: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ qui peut s'écrire comme la somme d'une fonction h différentiable de gradient Lipschitz et d'une fonction g séparable par bloc et non nécessairement lisse (voir la section 2.3.3.2).

L'hypothèse de séparabilité suggère l'utilisation d'une stratégie d'optimisation alternée, où à chaque itération, un seul bloc est mis à jour tandis que les autres restent inchangés. Ce type de stratégie permet de réduire le coût de calcul par itération. La plus simple d'entre elles est l'algorithme de minimisation alternée par bloc (*Block Coordinate Descent* (BCD) algorithm) qui minimise, à chaque itération, de façon exacte la fonction f par rapport à un bloc donné [Bertsekas, 1999; Luenberger, 1973; Ortega et Rheinboldt, 1970; Zangwill, 1969]. Une autre méthode pouvant être utilisée est la version proximale de l'algorithme BCD, introduite par [Auslender, 1992], qui, à chaque itération, effectue une étape proximale sur la fonction f relativement à un bloc donné. Cependant, aucune de ces deux méthodes ne permet d'exploiter la propriété de gradient Lipschitz de h . Une méthode pouvant en tirer parti est l'algorithme explicite-implicite alterné par bloc (*Block Coordinate Forward-Backward* (BC-FB) algorithm [Luo et Tseng, 1992a,b]). Notons que la convergence dans le cas où f n'est pas nécessairement convexe n'a été démontrée que pour l'algorithme proximal alterné par bloc et pour l'algorithme BC-FB (voir [Attouch et al., 2010a, 2011] et [Bolte et al., 2013; Xu et Yin, 2013] lorsque f satisfait l'inégalité de KL).

Dans ce chapitre nous nous intéressons plus particulièrement au troisième algorithme. Nous avons vu dans le chapitre 4 que l'utilisation d'une métrique variable permet d'accélérer la vitesse de convergence de l'algorithme explicite-implicite. Nous nous proposons d'étendre cette technique d'accélération au cas de l'algorithme BC-FB. Nous analyserons la convergence de cette méthode (ainsi que de sa version inexacte) dans le cas où la fonction f n'est pas nécessairement convexe.

Le reste du chapitre est organisé comme suit : nous formalisons le problème étudié et introduisons nos notations dans la section 5.2. Puis nous dressons, dans la section 5.3, un état de l'art des méthodes alternées par bloc afin de situer notre travail. Dans la section 5.4 nous proposons des versions exacte et inexacte de l'algorithme explicite-implicite à métrique variable alterné par bloc. Nous étudions leurs propriétés de convergence et de vitesse de convergence dans la section 5.5. Une conclusion sera fournie dans la section 5.6. Dans le chapitre 6 nous donnerons des exemples applicatifs de cet algorithme afin d'en illustrer son intérêt pratique.

5.2 PROBLÈME DE MINIMISATION

Dans ce chapitre, nous nous intéressons au problème d'optimisation suivant :

$$\text{trouver } \hat{\mathbf{x}} \in \underset{\mathbf{x} \in \mathbb{R}^N}{\text{Argmin}} (f(\mathbf{x}) = h(\mathbf{x}) + g(\mathbf{x})), \quad (5.1)$$

où $f: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ est une fonction coercive, h est une fonction différentiable et g est une fonction propre, semi-continue inférieurement et séparable par bloc. De façon formelle, on notera $(\mathbb{J}_j)_{1 \leq j \leq J}$ une partition de $\{1, \dots, N\}$ formée de $J \geq 2$ sous-ensembles, et, pour tout

$j \in \{1, \dots, J\}$, $N_j \neq 0$ désignera le cardinal de $\mathbb{J}_j$. Tout vecteur $\mathbf{x} = (\mathbf{x}^{(n)})_{1 \leq n \leq N} \in \mathbb{R}^N$ peut alors se décomposer par blocs :

$$\mathbf{x} = (\mathbf{x}^{(j)})_{1 \leq j \leq J} \in \mathbb{R}^{N_1} \times \dots \times \mathbb{R}^{N_J}, \quad (5.2)$$

où, pour tout $j \in \{1, \dots, J\}$, $\mathbf{x}^{(j)} = (\mathbf{x}^{(n)})_{n \in \mathbb{J}_j} \in \mathbb{R}^{N_j}$ et l'hypothèse de séparabilité sur $\mathbf{g}$ s'écrit

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \mathbf{g}(\mathbf{x}) := \sum_{j=1}^J \mathbf{g}_j(\mathbf{x}^{(j)}), \quad (5.3)$$

où, pour tout $j \in \{1, \dots, J\}$, $\mathbf{g}_j: \mathbb{R}^{N_j} \rightarrow]-\infty, +\infty]$.

Dans le reste de ce chapitre, nous utiliserons les notations suivantes. Pour tout $j \in \{1, \dots, J\}$, on notera $\bar{j}$ l'ensemble complémentaire de $\{j\}$ dans $\{1, \dots, J\}$, i.e. $\bar{j} := \{1, \dots, J\} \setminus \{j\}$. Ainsi, pour tout $\mathbf{x} \in \mathbb{R}^N$, $\mathbf{x}^{(\bar{j})} := (\mathbf{x}^{(1)}, \dots, \mathbf{x}^{(j-1)}, \mathbf{x}^{(j+1)}, \dots, \mathbf{x}^{(J)})$. De plus, pour tout $\mathbf{x}^{(\bar{j})} \in \prod_{i \in \bar{j}} \mathbb{R}^{N_i}$, la fonction $\mathbf{h}_j(\cdot, \mathbf{x}^{(\bar{j})}): \mathbb{R}^{N_j} \rightarrow \mathbb{R}$ désignera la fonction $\mathbf{h}$ restreinte au j -ème bloc, en $\mathbf{x}^{(\bar{j})}$, définie par

$$(\forall \mathbf{y} \in \mathbb{R}^{N_j}) \quad \mathbf{h}_j(\mathbf{y}, \mathbf{x}^{(\bar{j})}) := \mathbf{h}(\mathbf{x}^{(1)}, \dots, \mathbf{x}^{(j-1)}, \mathbf{y}, \mathbf{x}^{(j+1)}, \dots, \mathbf{x}^{(J)}). \quad (5.4)$$

Dans la suite, pour tout $\mathbf{x} \in \mathbb{R}^N$ et $j \in \{1, \dots, J\}$, on notera $\nabla_j \mathbf{h}(\mathbf{x}) \in \mathbb{R}^{N_j}$ le gradient partiel de $\mathbf{h}$ par rapport à $\mathbf{x}^{(j)}$ calculé en $\mathbf{x}$.

Le principe des méthodes alternées par blocs réside en la mise à jour de façon alternée au cours des itérations de chacun des blocs $(\mathbf{x}^{(j)})_{1 \leq j \leq J}$. Plus précisément, à chaque itération $k \in \mathbb{N}$ de l'algorithme, un indice $j_k \in \{1, \dots, J\}$ est sélectionné : le bloc $\mathbf{x}^{(j_k)}$ est alors mis à jour, tandis que tous les autres blocs $\mathbf{x}^{(\bar{j}_k)}$ restent inchangés. Les principales règles de sélection des blocs sont données dans la définition suivante :

Définition 5.1.

Soit $(j_k)_{k \in \mathbb{N}}$ la suite d'indices des blocs mis à jour au cours des itérations $k \in \mathbb{N}$.

(i) Les blocs sont mis à jour selon une règle cyclique si

$$(\forall k \in \mathbb{N}) \quad j_k - 1 = k \pmod{J}. \quad (5.5)$$

(ii) Les blocs sont mis à jour selon une règle essentiellement cyclique s'il existe une constante $K \geq J$ telle que,

$$(\forall k \in \mathbb{N}) \quad \{1, \dots, j\} \subset \{j_k, \dots, j_{k+K-1}\}. \quad (5.6)$$

(iii) Les blocs sont mis à jour selon une règle aléatoire uniforme si, pour tout $k \in \mathbb{N}$, j_k est une réalisation d'une variable aléatoire suivant une loi uniforme sur $\{1, \dots, J\}$.

La règle quasi-cyclique donnée par la définition 5.1(ii) a été introduite par [Luo et Tseng, 1992a] et utilisée sous le nom de règle essentiellement cyclique dans [Tseng, 2001] pour la mise à jour des blocs de l'algorithme BC-FB. Elle peut être interprétée de la façon suivante : les blocs peuvent être mis à jour dans un ordre arbitraire à partir du moment où ils sont tous mis à jour au moins une fois sur un nombre fini d'itérations. En particulier, sous l'hypothèse de mise à jour essentiellement cyclique, les blocs peuvent être sélectionnés suivant un ordre arbitraire

changeant toutes les K itérations. Ainsi, on comprend que la règle cyclique donnée par (i) est un cas particulier de la règle (ii).

En outre, pour toute itération $k \in \mathbb{N}$ et pour tout indice de bloc $j \in \{1, \dots, J\}$, on notera

$$l_{k,j} = \min \{l \in \{0, \dots, K-1\} \mid j_{k+l} = j\}, \quad (5.7)$$

la première fois où le j -ème bloc est mis à jour après la k -ème itération. Par exemple, lorsque la règle cyclique est utilisée, on a $l_{k,j} = j \bmod (J)$. Nous allons de plus introduire la permutation $\sigma_k: \{1, \dots, J\} \rightarrow \{1, \dots, J\}$ qui assure que la suite $(l_{k,\sigma_k(i)})_{1 \leq i \leq J}$ est croissante. Ces deux notations seront particulièrement utiles pour nous permettre de manipuler la règle de mise à jour essentiellement cyclique dans notre analyse de convergence.

5.3 ÉTAT DE L'ART

Nous résumons dans cette section les principales méthodes alternées existants dans la littérature afin de pouvoir situer notre travail et souligner ses avantages par rapport à celles-ci.

5.3.1 Algorithme de minimisation alternée par bloc

Une approche standard pour résoudre le problème (5.1) est d'utiliser l'algorithme de minimisation alternée par bloc donné par l'algorithme 5.1.

Algorithme 5.1 Algorithme de minimisation alternée par bloc

Initialisation : Soit $\mathbf{x}_0 \in \mathbb{R}^N$.

Itérations :

Pour $k = 0, 1, \dots$

$$\left[\begin{array}{l} \text{soit } j_k \in \{1, \dots, J\}, \\ \mathbf{x}_{k+1}^{(j_k)} \in \underset{\mathbf{y} \in \mathbb{R}^{N_{j_k}}}{\text{Argmin}} \left(h_{j_k}(\mathbf{y}, \mathbf{x}_k^{(\bar{j}_k)}) + g_{j_k}(\mathbf{y}) \right), \\ \mathbf{x}_{k+1}^{(\bar{j}_k)} = \mathbf{x}_k^{(\bar{j}_k)}. \end{array} \right.$$

Cette méthode peut aussi être vue comme une généralisation de l'algorithme de Gauss-Seidel permettant de résoudre les systèmes linéaires [Golub et Van Loan, 1996]. Elle est d'ailleurs désignée dans certains ouvrages comme *algorithme non-linéaire de Gauss-Seidel* (*nonlinear Gauss-Seidel method*) [Bertsekas, 1999, Chap.2], [Ortega et Rheinboldt, 1970, Chap.7]. Dans la plupart des ouvrages de référence tels que [Bertsekas, 1999; Luenberger, 1973; Ortega et Rheinboldt, 1970; Zangwill, 1969], la mise à jour des blocs dans l'algorithme 5.1 se fait selon une *règle cyclique* (c.f. définition 5.1(i)). Cependant, l'analyse de convergence effectuée dans [Tseng, 2001] suppose seulement que les blocs sont mis à jour selon une *règle essentiellement cyclique* (c.f. définition 5.1(ii)). Les résultats de convergence présentés dans [Tseng, 2001] semblent de plus être les plus généraux concernant cet algorithme. Ils assurent la convergence de la suite $(\mathbf{x}_k)_{k \in \mathbb{N}}$ dès lors que

- la fonction f est quasi-convexe et hemivariée¹ en chaque bloc,
- soit la fonction f est pseudo-convexe² sur chaque paire de blocs, soit elle admet au moins un minimum par rapport à chacun des blocs.

Cependant, comme il est souligné dans [Tseng, 2001], la deuxième hypothèse peut être contraignante au sens où l'algorithme 5.1 risque de ne pas converger si l'on suppose uniquement la convexité de f par rapport à chaque bloc (c.f. [Powell, 1973] pour une illustration).

5.3.2 Algorithme proximal alterné par bloc

Pour pallier les problèmes de convergence rencontrés avec l'algorithme 5.1, une version proximale de ce dernier a été introduite par [Auslender, 1992]. Elle est définie par l'algorithme 5.2.

Algorithme 5.2 Algorithme proximal alterné par bloc

Initialisation : Soit $\mathbf{x}_0 \in \mathbb{R}^N$. Soient, pour tout $k \in \mathbb{N}$, $\gamma_k \in]0, +\infty[$ et $\mathbf{A}_{j_k}(\mathbf{x}_k) \in \mathcal{S}_{N_{j_k}}^+$.

Itérations :

Pour $k = 0, 1, \dots$

$$\left\{ \begin{array}{l} \text{soit } j_k \in \{1, \dots, J\}, \\ \mathbf{x}_{k+1}^{(j_k)} \in \text{prox}_{\gamma_k^{-1} \mathbf{A}_{j_k}(\mathbf{x}_k), h_{j_k}(\cdot, \mathbf{x}_k^{(\bar{j}_k)}) + g_{j_k}} \left(\mathbf{x}_k^{(j_k)} \right), \\ \mathbf{x}_{k+1}^{(\bar{j}_k)} = \mathbf{x}_k^{(\bar{j}_k)}. \end{array} \right.$$

La convergence de la suite $(\mathbf{x}_k)_{k \in \mathbb{N}}$ générée par l'algorithme 5.2 vers une solution du problème (5.1) a été étudiée dans [Auslender, 1992] dans le cas où :

- la fonction h est une fonction différentiable, de gradient Lipschitz,
- les fonctions $(g_j)_{1 \leq j \leq J}$ sont convexes, propres, semi-continues inférieurement,
- pour tout $\ell \in \mathbb{N}$, $\mathbf{A}_{j_k}(\mathbf{x}_k) = \mathbf{I}_{N_{j_k}}$.

Récemment, la convergence vers un point critique de f , des itérées de l'algorithme 5.2 a été établie dans [Attouch et al., 2010a] dans le cas où les fonctions h et $(g_j)_{1 \leq j \leq J}$ ne sont pas nécessairement convexes et, pour tout $k \in \mathbb{N}$, $\mathbf{A}_{j_k}(\mathbf{x}_k) = \mathbf{I}_{N_{j_k}}$. Ces résultats ont été généralisés dans [Attouch et al., 2011] pour des matrices $(\mathbf{A}_{j_k}(\mathbf{x}_k))_{k \in \mathbb{N}}$ SDP quelconques. Notons que les études de convergences faites dans [Attouch et al., 2010a, 2011] reposent sur l'hypothèse que la fonction f satisfait l'inégalité de Kurdyka-Łojasiewicz. Dans les trois études citées ci-dessus, la mise à jour des indices de blocs $(j_\ell)_{\ell \in \mathbb{N}}$ se fait selon la règle cyclique.

Notons que la convergence de l'algorithme 5.2 a été étudiée dans [Bauschke et al., 2006] dans le cadre des opérateurs de projection de Bregman, dans le cas où $J = 2$, h est la distance de Bregman et g_1, g_2 sont des fonctions convexes. D'autre part, nous pouvons remarquer que

1. Une fonction $f: D \subset \mathbb{R}^N \rightarrow \mathbb{R}$ est *hemivariée* sur un sous-ensemble D_0 de D si elle n'est constante sur aucun segment de D_0 , i.e. il n'existe pas de couple $(\mathbf{x}, \mathbf{y}) \in D_0^2$ tel que $\mathbf{x} \neq \mathbf{y}$, et, pour tout $t \in [0, 1]$, $(1-t)\mathbf{x} + t\mathbf{y} \in D_0$ et $f((1-t)\mathbf{x} + t\mathbf{y}) = f(\mathbf{x})$ [Ortega et Rheinboldt, 1970, Déf. 14.1.1].

2. Une fonction différentiable $f: \mathbb{R}^N \rightarrow \mathbb{R}$ est *pseudo-convexe* si, pour tout $(\mathbf{x}, \mathbf{y}) \in (\mathbb{R}^N)^2$ tel que $\langle \mathbf{x} - \mathbf{y}, \nabla f(\mathbf{x}) \rangle \geq 0$, on a $f(\mathbf{y}) \geq f(\mathbf{x})$ [Mangasarian, 1969, Chap. 9.3]. Cette notion peut-être généralisée au cas où la fonction f n'est pas différentiable en considérant les éléments de son sous-différentiel [Hadjisavvas et Schaible, 2009].

le célèbre algorithme POCS (*Projection Onto Convex Sets* - algorithme de projection sur des ensembles convexes) introduit par [Bregman, 1965] peut être vu comme un cas particulier de l'algorithme 5.2 lorsque $h \equiv 0$ et, pour tout $j \in \{1, \dots, J\}$, g_j est une fonction indicatrice sur un ensemble convexe.

Une des difficultés majeures rencontrées lors de la mise en application de l'algorithme 5.2 est le calcul, à chaque itération $k \in \mathbb{N}$, de l'opérateur proximal de la fonction $h_{j_k}(\cdot, \mathbf{x}_k^{(\bar{j}_k)}) + g_{j_k}$. En effet, même dans le cas où les opérateurs proximaux des fonctions $h_{j_k}(\cdot, \mathbf{x}_k^{(\bar{j}_k)})$ et g_{j_k} ont des formes simples, rien ne permet d'affirmer que l'opérateur proximal de $h_{j_k}(\cdot, \mathbf{x}_k^{(\bar{j}_k)}) + g_{j_k}$ aura une forme explicite. Ainsi, [Attouch et al., 2011] ont proposé une forme inexacte de l'algorithme 5.2 afin de calculer de façon approchée l'opérateur proximal.

5.3.3 Algorithme explicite-implicite alterné par bloc

Une autre stratégie permettant de pallier la possible difficulté rencontrée lors du calcul des opérateurs proximaux dans l'algorithme 5.2 est de remplacer à chaque itération l'étape proximale par une étape explicite-implicite (c.f. algorithme 2.4). Les itérées obtenues sont décrites dans l'algorithme 5.3.

Algorithme 5.3 Algorithme explicite-implicite alterné par bloc

Initialisation : Soient $\mathbf{x}_0 \in \mathbb{R}^N$ et, pour tout $k \in \mathbb{N}$, $\gamma_k \in]0, +\infty[$.

Itérations :

Pour $k = 0, 1, \dots$

$$\left| \begin{array}{l} \text{soit } j_k \in \{1, \dots, J\}, \\ \mathbf{x}_{k+1}^{(j_k)} \in \text{prox}_{\gamma_k g_{j_k}} \left(\mathbf{x}_k^{(j_k)} - \gamma_k \nabla_{j_k} h(\mathbf{x}_k) \right), \\ \mathbf{x}_{k+1}^{(\bar{j}_k)} = \mathbf{x}_k^{(\bar{j}_k)}. \end{array} \right.$$

Cette méthode a tout d'abord été proposée par [Censor et Lent, 1987] afin de minimiser la fonction entropique de Burg sous des contraintes linéaires. Elle a ensuite été généralisée dans [Luo et Tseng, 1992a,b] au cas où la fonction h est convexe différentiable de gradient Lipschitz et les indices $(j_k)_{k \in \mathbb{N}}$ sont mis à jour selon la règle essentiellement cyclique.

Récemment, une étude de convergence de l'algorithme 5.3 a été menée dans [Combettes et Pesquet, 2015] dans le cas où h est convexe de gradient Lipschitz, g est convexe, les blocs sont mis à jour suivant une règle aléatoire, et plusieurs blocs peuvent être mis à jours simultanément à chaque itération. La suite des itérées générée par l'algorithme 5.3 est alors une suite de vecteurs aléatoires, et les auteurs en ont démontré la convergence presque sûre vers un minimiseur de f .

Finalement, les propriétés de convergence de l'algorithme 5.3 ont été étudiées dans [Xu et Yin, 2013] lorsque h est une fonction différentiable de gradient Lipschitz, g est non nécessairement lisse, et les deux fonctions sont convexes par rapport à chaque bloc. Les auteurs ont démontré que les suites $(\mathbf{x}_k)_{k \in \mathbb{N}}$ générées par l'algorithme 5.3 convergent vers un point critique de f . Ces résultats de convergence ont été généralisés dans [Bolte et al., 2013] dans le cas où h et g ne sont pas nécessairement convexes. Dans cet article, l'algorithme 5.3 est dénommé algorithme *Proximal Alternating Linearized Minimization* (PALM). Dans les deux travaux [Bolte et al., 2013; Xu et

Yin, 2013], les démonstrations de convergence reposent sur l'hypothèse que f satisfait l'inégalité de KL, et que les blocs sont mis à jour selon la règle cyclique.

5.4 MÉTHODE PROPOSÉE

5.4.1 Algorithme explicite-implicite à métrique variable alterné par bloc

Comme nous l'avons vu dans le chapitre 3, l'introduction d'une métrique variant au cours des itérations permet d'accélérer la vitesse de convergence de l'algorithme explicite-implicite. Ainsi, nous proposons de faire de même pour l'algorithme 5.3. Les itérées de cette méthode sont données dans l'algorithme 5.4, et est appelé algorithme explicite-implicite à métrique variable alterné par bloc (*Block Coordinate Variable Metric Forward-Backward* (BC-VMFB) algorithm).

Algorithme 5.4 Algorithme explicite-implicite à métrique variable alterné par bloc

Initialisation : Soient $\mathbf{x}_0 \in \mathbb{R}^N$, et, pour tout $k \in \mathbb{N}$, $\gamma_k \in]0, +\infty[$ et $\mathbf{A}_{j_k}(\mathbf{x}_k) \in \mathcal{S}_{N_{j_k}}^+$.

Itérations :

Pour $k = 0, 1, \dots$

$$\left[\begin{array}{l} \text{soit } j_k \in \{1, \dots, J\}, \\ \mathbf{x}_{k+1}^{(j_k)} \in \text{prox}_{\gamma_k^{-1} \mathbf{A}_{j_k}(\mathbf{x}_k), \mathbf{g}_{j_k}} \left(\mathbf{x}_k^{(j_k)} - \gamma_k (\mathbf{A}_{j_k}(\mathbf{x}_k))^{-1} \nabla_{j_k} h(\mathbf{x}_k) \right), \\ \mathbf{x}_{k+1}^{(\bar{j}_k)} = \mathbf{x}_k^{(\bar{j}_k)}, \end{array} \right.$$

L'algorithme 5.4 (ainsi qu'une version inexacte de ce dernier) a été étudié dans [Richtárik et Talác, 2014] (resp. [Tappenden et al., 2013]) lorsque les blocs sont mis à jour suivant une règle aléatoire uniforme (donnée dans la définition 5.1(iii)). Les auteurs ont établi la convergence de la suite $(f(\mathbf{x}_k))_{k \in \mathbb{N}}$ lorsque les fonctions h et $\mathbf{g}_j$ sont convexes au sens où, pour tout $\delta \geq 0$ et $\epsilon \geq 0$, il existe $k_0 \in \mathbb{N}$ tel que la probabilité d'avoir $f(\mathbf{x}_{k_0}) - f(\bar{\mathbf{x}}) \leq \epsilon$ est supérieure à $1 - \delta$.

Par ailleurs, comme il l'est indiqué dans [Razaviyayn et al., 2013], lorsque les matrices de pré-conditionnement $(\mathbf{A}_{j_k}(\mathbf{x}_k))_{k \in \mathbb{N}}$ sont choisies judicieusement, l'algorithme 5.4 peut être vu comme un algorithme MM alterné par bloc. Ce type d'approche a été proposée dans [Fessler, 1997; Saquib et al., 1995; Sotthivirat et Fessler, 2002] dans un contexte de reconstruction d'image. Ainsi, certaines propriétés de convergence de l'algorithme 5.4 peuvent être directement déduites de celles énoncées dans [Jacobson et Fessler, 2007] dans le cas où les fonctions $\mathbf{g}_j$ sont des fonctions indicatrices de sous-ensembles convexes fermés de $\mathbb{R}^{N_j}$, et de celles énoncées dans [Razaviyayn et al., 2013] lorsque les fonctions $\mathbf{g}_j$ sont non lisses et convexes. Cependant, notons que dans [Jacobson et Fessler, 2007; Razaviyayn et al., 2013], la convergence des itérées $(\mathbf{x}_k)_{k \in \mathbb{N}}$ vers une solution du problème (5.1) est démontrée seulement sous des hypothèses spécifiques. En particulier, il faut que la solution du problème (5.1) soit unique, et que chaque fonction $h_j(\cdot, \mathbf{x}^{(\bar{j})}) + \mathbf{g}_j$ n'admette qu'un unique minimiseur, pour $\mathbf{x}^{(\bar{j})}$ fixé.

Remarque 5.1. Lorsque $J = 1$ et $N_1 = N$, si $\mathbf{g} = \mathbf{g}_1$ est convexe, on retrouve l'algorithme 3.1 étudié dans le chapitre 3 dans le cas où il n'y a pas d'étape de relaxation dans ce dernier (i.e., pour tout $k \in \mathbb{N}$ dans l'algorithme 3.1, $\lambda_k = 1$).

5.4.2 Hypothèses

Dans le reste de ce chapitre, nous nous intéresserons aux fonctions $\mathbf{f}$, $\mathbf{h}$ et $\mathbf{g}$ satisfaisant les hypothèses suivantes :

Hypothèse 5.1.

- (i) Pour tout $j \in \{1, \dots, J\}$, $\mathbf{g}_j: \mathbb{R}^{N_j} \rightarrow]-\infty, +\infty]$ est une fonction propre, semi-continue inférieurement, bornée inférieurement par une fonction affine et sa restriction à son domaine de définition est continue.
- (ii) La fonction $\mathbf{h}: \mathbb{R}^N \rightarrow \mathbb{R}$ est différentiable, et $\mathbf{h}$ est de gradient β -Lipschitzien sur $\text{dom } \mathbf{g}$, avec $\beta > 0$.
- (iii) La fonction $\mathbf{f}$ est coercive.

La remarque suivante sera utile dans la suite du chapitre, en particulier pour l'étude de convergence de l'algorithme 5.4.

Remarque 5.2.

- (i) L'hypothèse 5.1(ii) est moins contraignante que l'hypothèse de Lipschitz différentiabilité de $\mathbf{h}$ utilisée en général pour démontrer la convergence de l'algorithme explicite-implicite [Attouch et al., 2011; Combettes et Wajs, 2005]. En particulier, si $\text{dom } \mathbf{g}$ est compact et $\mathbf{h}$ est deux fois continument différentiable, l'hypothèse 5.1(ii) est alors satisfaite.
- (ii) D'après l'hypothèse 5.1(ii), $\text{dom } \mathbf{g} \subset \text{dom } \mathbf{h} = \mathbb{R}^N$. Ainsi, en conséquence de l'hypothèse 5.1(i), $\text{dom } \mathbf{f} = \text{dom } \mathbf{g}$ est non vide.
- (iii) Si l'hypothèse 5.1 est satisfaite, alors $\mathbf{f}$ est une fonction propre, semi-continue inférieurement, et sa restriction à son domaine de définition est continue. Ainsi, la fonction $\mathbf{f}$ étant coercive, pour tout $\mathbf{x} \in \text{dom } \mathbf{g}$, $\text{lev}_{\leq \mathbf{f}(\mathbf{x})} \mathbf{f}$ est un ensemble compact. De plus, l'ensemble des minimiseurs de $\mathbf{f}$ est non vide et compact.
- (iv) Si, pour tout $j \in \{1, \dots, J\}$, la fonction $\mathbf{g}_j$ est propre, semi-continue inférieurement et convexe, alors elle est bornée inférieurement par une fonction affine.

Comme nous l'avons montré dans les chapitres 3 et 4, les matrices de préconditionnement $(\mathbf{A}_{j_k}(\mathbf{x}_k))_{k \in \mathbb{N}}$, induisant une métrique variable, jouent un rôle central dans l'algorithme 5.4. Nous proposons de les choisir suivant une stratégie MM. Plus précisément, soit $k \in \mathbb{N}$ une itération de l'algorithme 5.4, $\mathbf{x}_k \in \text{dom } \mathbf{g}$ l'itérée associée, et $j_k \in \{1, \dots, J\}$ l'indice du bloc sélectionné. La matrice $\mathbf{A}_{j_k}(\mathbf{x}_k) \in \mathcal{S}_{N_{j_k}}^+$ sera alors choisie de façon à satisfaire la condition de majoration suivante :

Hypothèse 5.2.

(i) La fonction quadratique définie par

$$(\forall \mathbf{y} \in \mathbb{R}^{N_{j_k}}) \quad q_{j_k}(\mathbf{y}, \mathbf{x}_k) := h(\mathbf{x}_k) + \langle \mathbf{y} - \mathbf{x}_k^{(j_k)}, \nabla_{j_k} h(\mathbf{x}_k) \rangle + \frac{1}{2} \langle \mathbf{y} - \mathbf{x}_k^{(j_k)}, \mathbf{A}_{j_k}(\mathbf{x}_k)(\mathbf{y} - \mathbf{x}_k^{(j_k)}) \rangle,$$

est une fonction majorante de $h_{j_k}(\cdot, \mathbf{x}_k^{(\bar{j}_k)})$ en $\mathbf{x}_k^{(j_k)}$ sur $\text{dom } \mathbf{g}_{j_\ell}$, i.e.,

$$(\forall \mathbf{y} \in \text{dom } \mathbf{g}_{j_k}) \quad h_{j_k}(\mathbf{y}, \mathbf{x}_k^{(\bar{j}_k)}) \leq q_{j_k}(\mathbf{y}, \mathbf{x}_k).$$

(ii) Il existe $(\underline{\nu}, \bar{\nu}) \in]0, +\infty[^2$ tel que

$$(\forall \ell \in \mathbb{N}) \quad \underline{\nu} \mathbf{I}_{N_{j_\ell}} \preccurlyeq \mathbf{A}_{j_k}(\mathbf{x}_k) \preccurlyeq \bar{\nu} \mathbf{I}_{N_{j_\ell}}.$$

Remarque 5.3.

- (i) Remarquons qu'il n'est pas nécessaire de construire des majorantes quadratiques de $h_j(\cdot, \mathbf{x}^{(\bar{j})})$ sur $\text{dom } \mathbf{g}_j$ pour tout $j \in \{1, \dots, J\}$ et pour tout $\mathbf{x}^{(\bar{j})} \in \bigtimes_{i \in \bar{j}} \text{dom } \mathbf{g}_i$. Il suffit que la majorante soit construite aux points $(\mathbf{x}_k)_{k \in \mathbb{N}}$.
- (ii) Supposons que, pour tout $\mathbf{x}' \in \text{dom } \mathbf{g}$, une fonction majorante quadratique de h sur $\text{dom } \mathbf{g}$ soit donnée par

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad q(\mathbf{x}, \mathbf{x}') := h(\mathbf{x}') + \langle \mathbf{x} - \mathbf{x}', \nabla h(\mathbf{x}') \rangle + \frac{1}{2} \langle \mathbf{x} - \mathbf{x}', \mathbf{B}(\mathbf{x}')(\mathbf{x} - \mathbf{x}') \rangle, \quad (5.8)$$

où $\mathbf{B}(\mathbf{x}') = \left(\mathbf{B}(\mathbf{x}_k)^{(n, n')} \right)_{(n, n') \in \{1, \dots, N\}^2} \in \mathcal{S}_N^+$. L'hypothèse 5.2(i) est alors satisfaite lorsque $\mathbf{A}_{j_k}(\mathbf{x}_k) = \left(\mathbf{B}(\mathbf{x}_k)^{(n, n')} \right)_{(n, n') \in \mathbb{J}_{j_k}^2}$. De plus, s'il existe $(\underline{\nu}, \bar{\nu}) \in]0, +\infty[^2$ tel que, pour tout $\mathbf{x}' \in \text{dom } \mathbf{g}$, $\underline{\nu} \mathbf{I}_N \preccurlyeq \mathbf{B}(\mathbf{x}') \preccurlyeq \bar{\nu} \mathbf{I}_N$, alors l'hypothèse 5.2(ii) est aussi satisfaite.

- (iii) D'après la lemme 2.1, si $\text{dom } \mathbf{g}$ est convexe, l'existence de la fonction majorante (5.8) est assurée lorsque h satisfait l'hypothèse 5.1(ii).

D'autre part, afin d'assurer que chaque bloc soit mis à jour une infinité de fois au cours des itérations de l'algorithme 5.4, nous allons considérer la règle de mise à jour des blocs suivante [Luo et Tseng, 1992b; Tseng, 2001] :

Hypothèse 5.3.

Les indices $(j_k)_{k \in \mathbb{N}}$ des blocs sont mis à jour selon la règle essentiellement cyclique donnée dans la définition 5.1.

Pour finir, nous supposons que, pour tout $k \in \mathbb{N}$, le pas γ_k utilisé dans l'algorithme 5.4 satisfait l'hypothèse suivante :

Hypothèse 5.4.

L'une des assertions suivantes est vérifiée :

- (i) il existe $(\underline{\gamma}, \overline{\gamma}) \in]0, +\infty[^2$ tel que, pour tout $k \in \mathbb{N}$, $\underline{\gamma} \leq \gamma_k \leq 1 - \overline{\gamma}$;
- (ii) pour tout $j \in \{1, \dots, J\}$, $\mathbf{g}_j$ est une fonction convexe et il existe $(\underline{\gamma}, \overline{\gamma}) \in]0, +\infty[^2$ tel que, pour tout $k \in \mathbb{N}$, $\underline{\gamma} \leq \gamma_k \leq 2 - \overline{\gamma}$.

5.4.3 Version inexacte de l'algorithme explicite-implicite à métrique variable alterné par bloc

De la même façon que dans le chapitre 3, nous proposons une version inexacte de l'algorithme 5.4 pour résoudre le problème (5.1) dans le cas où les opérateurs proximaux des fonctions $(\mathbf{g}_j)_{1 \leq j \leq J}$ n'ont pas de forme explicite :

Algorithme 5.5 Algorithme inexact explicite-implicite à métrique variable alterné par bloc

Initialisation : Soient $\mathbf{x}_0 \in \mathbb{R}^N$, $\alpha_1 \in]1/2, +\infty[$, $\alpha_2 \in]0, +\infty[$. Soient, pour tout $k \in \mathbb{N}$, $\gamma_k \in]0, +\infty[$ et $\mathbf{A}_{j_k}(\mathbf{x}_k) \in \mathcal{S}_{N_{j_k}}^+$.

Itérations :

Pour $k = 0, 1, \dots$

$$\left[\begin{array}{l} \text{soit } j_k \in \{1, \dots, J\}, \\ \text{trouver } \mathbf{x}_{k+1}^{(j_k)} \in \mathbb{R}^{N_{j_k}} \text{ et } \mathbf{r}_{k+1}^{(j_k)} \in \partial \mathbf{g}_{j_k}(\mathbf{x}_{k+1}^{(j_k)}) \text{ tels que} \\ \quad \mathbf{g}_{j_k}(\mathbf{x}_{k+1}^{(j_k)}) + \langle \mathbf{x}_{k+1}^{(j_k)} - \mathbf{x}_k^{(j_k)}, \nabla_{j_k} \mathbf{h}(\mathbf{x}_k) \rangle + \alpha_1 \|\mathbf{x}_{k+1}^{(j_k)} - \mathbf{x}_k^{(j_k)}\|_{\mathbf{A}_{j_k}(\mathbf{x}_k)}^2 \leq \mathbf{g}_{j_k}(\mathbf{x}_k^{(j_k)}), \\ \quad \|\nabla_{j_k} \mathbf{h}(\mathbf{x}_k) + \mathbf{r}_{k+1}^{(j_k)}\| \leq \alpha_2 \|\mathbf{x}_{k+1}^{(j_k)} - \mathbf{x}_k^{(j_k)}\|_{\mathbf{A}_{j_k}(\mathbf{x}_k)}, \\ \quad \mathbf{x}_{k+1}^{(\overline{j}_k)} = \mathbf{x}_k^{(\overline{j}_k)}. \end{array} \right.$$

Dans l'algorithme 5.5, la première inégalité est une condition de décroissance suffisante et la seconde inégalité correspond à une condition d'optimalité.

La remarque suivante explique en quoi l'algorithme 5.5 peut être interprété comme une version inexacte de l'algorithme 5.3.

Remarque 5.4.

Soient $(\mathbf{x}_k)_{k \in \mathbb{N}}$ et $(j_k)_{k \in \mathbb{N}}$ des suites générées par l'algorithme (5.3).

- (i) Dans un premier temps, supposons que l'hypothèse 5.4(i) soit satisfaite. D'une part, par définition de l'opérateur proximal, on a, pour tout $k \in \mathbb{N}$,

$$\mathbf{g}_{j_k}(\mathbf{x}_{k+1}^{(j_k)}) + \langle \mathbf{x}_{k+1}^{(j_k)} - \mathbf{x}_k^{(j_k)}, \nabla_{j_k} \mathbf{h}(\mathbf{x}_k) \rangle + \frac{\gamma_k^{-1}}{2} \|\mathbf{x}_{k+1}^{(j_k)} - \mathbf{x}_k^{(j_k)}\|_{\mathbf{A}_{j_k}(\mathbf{x}_k)}^2 \leq \mathbf{g}_{j_k}(\mathbf{x}_k^{(j_k)}),$$

ainsi, la condition de décroissance suffisante de l'algorithme 5.5 est satisfaite pour $\alpha_1 = (1 - \overline{\gamma})^{-1}/2$ (puisque $\gamma_k^{-1} \geq (1 - \overline{\gamma})^{-1} > 1$). D'autre part, en considérant la caractérisation

variationnelle de l'opérateur proximal, il existe $\mathbf{r}_{k+1}^{(j_k)} \in \partial \mathbf{g}_{j_k}(\mathbf{x}_{k+1}^{(j_k)})$ tel que

$$\mathbf{r}_{k+1}^{(j_k)} = -\nabla_{j_k} \mathbf{h}(\mathbf{x}_k) + \gamma_k^{-1} \mathbf{A}_{j_k}(\mathbf{x}_k)(\mathbf{x}_k^{(j_k)} - \mathbf{x}_{k+1}^{(j_k)}).$$

Sous les hypothèses 5.2(ii) et 5.4(i), on obtient alors

$$\|\mathbf{r}_{k+1}^{(j_k)} + \nabla_{j_k} \mathbf{h}(\mathbf{x}_k)\| = \gamma_k^{-1} \|\mathbf{A}_{j_k}(\mathbf{x}_k)(\mathbf{x}_k^{(j_k)} - \mathbf{x}_{k+1}^{(j_k)})\| \leq \underline{\gamma}^{-1} \sqrt{\underline{\nu}} \|\mathbf{x}_k^{(j_k)} - \mathbf{x}_{k+1}^{(j_k)}\|_{\mathbf{A}_{j_k}(\mathbf{x}_k)},$$

qui correspond à la condition d'optimalité donnée par la seconde inégalité de l'algorithme 5.5 pour $\alpha_2 = \underline{\gamma}^{-1} \sqrt{\underline{\nu}}$.

- (ii) Supposons maintenant que l'hypothèse 5.4(ii) soit satisfaite. Soit $k \in \mathbb{N}$. D'après la caractérisation variationnelle de l'opérateur proximal ainsi que la convexité de $\mathbf{g}_{j_k}$, il existe $\mathbf{r}_{k+1}^{(j_k)} \in \partial \mathbf{g}_{j_k}(\mathbf{x}_{k+1}^{(j_k)})$ tel que

$$\begin{cases} \mathbf{r}_{k+1}^{(j_k)} = -\nabla_{j_k} \mathbf{h}(\mathbf{x}_k) + \gamma_k^{-1} \mathbf{A}_{j_k}(\mathbf{x}_k)(\mathbf{x}_k^{(j_k)} - \mathbf{x}_{k+1}^{(j_k)}) \\ \langle \mathbf{x}_{k+1}^{(j_k)} - \mathbf{x}_k^{(j_k)}, \mathbf{r}_{k+1}^{(j_k)} \rangle \geq \mathbf{g}_{j_k}(\mathbf{x}_{k+1}^{(j_k)}) - \mathbf{g}_{j_k}(\mathbf{x}_k^{(j_k)}), \end{cases}$$

ce qui implique

$$\mathbf{g}_{j_k}(\mathbf{x}_{k+1}^{(j_k)}) + \langle \mathbf{x}_{k+1}^{(j_k)} - \mathbf{x}_k^{(j_k)}, \nabla_{j_k} \mathbf{h}(\mathbf{x}_k) \rangle + \gamma_k^{-1} \|\mathbf{x}_{k+1}^{(j_k)} - \mathbf{x}_k^{(j_k)}\|_{\mathbf{A}_{j_k}(\mathbf{x}_k)}^2 \leq \mathbf{g}_{j_k}(\mathbf{x}_k^{(j_k)}),$$

ainsi, la condition suffisante de décroissance de l'algorithme 5.5 est satisfaite pour $\alpha_1 = (2 - \bar{\gamma})^{-1}$ (puisque $\gamma_k^{-1} \geq (2 - \bar{\gamma})^{-1} > 1/2$). La condition d'optimalité inexacte est obtenue par un raisonnement similaire à celui de (i).

5.5 ANALYSE DE CONVERGENCE

5.5.1 Propriétés de descente

Dans cette section, nous donnons des résultats techniques concernant le comportement asymptotique de la suite $(\mathbf{f}(\mathbf{x}_k))_{k \in \mathbb{N}}$ générée par l'algorithme 5.5. Ces résultats préliminaires nous permettront ensuite de démontrer la convergence de la méthode proposée.

Lemme 5.1. Soit $(\mathbf{x}_k)_{k \in \mathbb{N}}$ une suite générée par l'algorithme 5.5. Supposons que les hypothèses 5.1 et 5.2 sont vérifiées. Il existe $\mu \in]0, +\infty[$ tel que, pour tout $k \in \mathbb{N}$,

$$\mathbf{f}(\mathbf{x}_{k+1}) \leq \mathbf{f}(\mathbf{x}_k) - \frac{\mu}{2} \|\mathbf{x}_{k+1}^{(j_k)} - \mathbf{x}_k^{(j_k)}\|^2 = \mathbf{f}(\mathbf{x}_k) - \frac{\mu}{2} \|\mathbf{x}_{k+1} - \mathbf{x}_k\|^2. \quad (5.9)$$

Démonstration. Soit $k \in \mathbb{N}$. On a

$$\mathbf{f}(\mathbf{x}_{k+1}) = \mathbf{h}(\mathbf{x}_{k+1}) + \mathbf{g}(\mathbf{x}_{k+1}).$$

D'une part, d'après l'hypothèse 5.2(i),

$$h(\mathbf{x}_{k+1}) \leq h(\mathbf{x}_k) + \left\langle \mathbf{x}_{k+1}^{(j_k)} - \mathbf{x}_k^{(j_k)}, \nabla_{j_k} h(\mathbf{x}_k) \right\rangle + \frac{1}{2} \|\mathbf{x}_{k+1}^{(j_k)} - \mathbf{x}_k^{(j_k)}\|_{\mathbf{A}_{j_k}(\mathbf{x}_k)}^2. \quad (5.10)$$

D'autre part, d'après l'algorithme 5.5, on a

$$(\forall j \in \bar{j}_k) \quad \mathbf{x}_{k+1}^{(j)} = \mathbf{x}_k^{(j)}, \quad (5.11)$$

ainsi

$$\begin{aligned} \mathbf{g}(\mathbf{x}_{k+1}) &= \mathbf{g}_{j_k}(\mathbf{x}_{k+1}^{(j_k)}) + \sum_{j \in \bar{j}_k} \mathbf{g}_j(\mathbf{x}_{k+1}^{(j)}) \\ &= \mathbf{g}_{j_k}(\mathbf{x}_{k+1}^{(j_k)}) + \sum_{j \in \bar{j}_k} \mathbf{g}_j(\mathbf{x}_k^{(j)}) \\ &= \mathbf{g}(\mathbf{x}_k) + \left(\mathbf{g}_{j_k}(\mathbf{x}_{k+1}^{(j_k)}) - \mathbf{g}_{j_k}(\mathbf{x}_k^{(j_k)}) \right). \end{aligned}$$

Alors, la première inégalité de l'algorithme 5.5 nous donne

$$\mathbf{g}(\mathbf{x}_{k+1}) \leq \mathbf{g}(\mathbf{x}_k) - \left\langle \mathbf{x}_{k+1}^{(j_k)} - \mathbf{x}_k^{(j_k)}, \nabla_{j_k} h(\mathbf{x}_k) \right\rangle - \alpha_1 \|\mathbf{x}_{k+1}^{(j_k)} - \mathbf{x}_k^{(j_k)}\|_{\mathbf{A}_{j_k}(\mathbf{x}_k)}^2. \quad (5.12)$$

Par conséquent, en combinant (5.10) et (5.12), on obtient

$$f(\mathbf{x}_{k+1}) \leq f(\mathbf{x}_k) - \left(\alpha_1 - \frac{1}{2} \right) \|\mathbf{x}_{k+1}^{(j_k)} - \mathbf{x}_k^{(j_k)}\|_{\mathbf{A}_{j_k}(\mathbf{x}_k)}^2. \quad (5.13)$$

Pour conclure, (5.9) découle de l'hypothèse 5.2(ii), en posant $\mu = \underline{\nu}(2\alpha_1 - 1)$ (puisque $\alpha_1 \in]1/2, +\infty[$), et en utilisant l'équation (5.11). $\square$

Soit $(\mathbf{x}_k)_{k \in \mathbb{N}}$ la suite définie par

$$(\forall k \in \mathbb{N}) \quad \mathbf{x}_k = (\mathbf{x}_{k+l+1} - \mathbf{x}_{k+l})_{0 \leq l \leq K-1} \in (\mathbb{R}^N)^K, \quad (5.14)$$

où $(\mathbf{x}_k)_{k \in \mathbb{N}}$ est une suite générée par l'algorithme 5.5 et K est la constante positive désignant le nombre d'itérations maximal au cours desquelles chaque bloc doit être mis à jour (d'après l'hypothèse 5.3). On a alors

$$\|\mathbf{x}_k\|^2 = \sum_{l=0}^{K-1} \|\mathbf{x}_{k+l+1} - \mathbf{x}_{k+l}\|^2,$$

et la propriété suivante est satisfaite.

Lemme 5.2. *Soit $(\mathbf{x}_k)_{k \in \mathbb{N}}$ une suite générée par l'algorithme 5.5. Supposons que les hypothèses 5.1, 5.2 et 5.3 sont satisfaites. Pour tout $k \in \mathbb{N}$,*

$$f(\mathbf{x}_{k+K}) \leq f(\mathbf{x}_k) - \frac{\mu}{2} \|\mathbf{x}_k\|^2,$$

où $\mu \in]0, +\infty[$ est la même constante que celle donnée dans le lemme 5.1.

Démonstration. Soit $k \in \mathbb{N}$. D'après le lemme 5.1, on a

$$\begin{aligned} f(\mathbf{x}_{k+K}) &\leq f(\mathbf{x}_{k+K-1}) - \frac{\mu}{2} \|\mathbf{x}_{k+K} - \mathbf{x}_{k+K-1}\|^2 \\ &\leq f(\mathbf{x}_{k+K-2}) - \frac{\mu}{2} \left(\|\mathbf{x}_{k+K-1} - \mathbf{x}_{k+K-2}\|^2 + \|\mathbf{x}_{k+K} - \mathbf{x}_{k+K-1}\|^2 \right) \\ &\vdots \\ &\leq f(\mathbf{x}_k) - \frac{\mu}{2} \sum_{l=0}^{K-1} \|\mathbf{x}_{k+l+1} - \mathbf{x}_{k+l}\|^2. \end{aligned}$$

□

5.5.2 Théorème de convergence

Dans cette section nous allons énoncer deux théorèmes de convergence :

- Le premier assure la convergence de la suite $(\mathbf{x}_k)_{k \in \mathbb{N}}$ générée par l'algorithme 5.5 vers un point critique de f .
- Le second est un résultat de convergence local assurant que si l'algorithme 5.5 est initialisé de façon judicieuse, alors la suite $(\mathbf{x}_k)_{k \in \mathbb{N}}$ converge vers une solution du problème 5.1.

Pour cela, nous allons, dans un premier temps, donner deux lemmes qui nous permettront de manipuler la règle de mise à jour des bloc essentiellement cyclique.

Lemme 5.3. Soit $(\mathbf{x}_k)_{k \in \mathbb{N}}$ une suite générée par l'algorithme 5.5. Soit $k_0 \in \mathbb{N}$. On définit $\mathcal{J}_{k_0}$ comme étant un sous-ensemble de $\{1, \dots, J\}$ tel que $j_{k_0} \in \mathcal{J}_{k_0}$. On suppose que les hypothèses 5.1 et 5.2 sont satisfaites. Alors, on a

$$\begin{aligned} \sum_{j \in \mathcal{J}_{k_0}} \|\nabla_j \mathbf{h}(\mathbf{x}_{k_0+1}) + \mathbf{r}_{k_0+1}^{(j)}\|^2 &\leq 2(\beta^2 + \alpha_2^2 \bar{\nu}) \|\mathbf{x}_{k_0+1} - \mathbf{x}_{k_0}\|^2 \\ &\quad + 2 \sum_{j \in \mathcal{J}_{k_0} \setminus \{j_{k_0}\}} \|\nabla_j \mathbf{h}(\mathbf{x}_{k_0}) + \mathbf{r}_{k_0}^{(j)}\|^2, \end{aligned} \quad (5.15)$$

où $\mathbf{r}_{k_0+1}^{(j_{k_0})}$ est donné par l'algorithme 5.5 et, pour tout $j \in \mathcal{J}_{k_0} \setminus \{j_{k_0}\}$, $\mathbf{r}_{k_0+1}^{(j)} \in \partial \mathbf{g}_j(\mathbf{x}_{k_0+1})$ et $\mathbf{r}_{k_0}^{(j)} \in \partial \mathbf{g}_j(\mathbf{x}_{k_0})$.

Démonstration. Soit $k_0 \in \mathbb{N}$. D'après l'inégalité de Jensen, on a

$$\begin{aligned} \sum_{j \in \mathcal{J}_{k_0}} \|\nabla_j \mathbf{h}(\mathbf{x}_{k_0+1}) + \mathbf{r}_{k_0+1}^{(j)}\|^2 &\leq 2 \sum_{j \in \mathcal{J}_{k_0}} \|\nabla_j \mathbf{h}(\mathbf{x}_{k_0+1}) - \nabla_j \mathbf{h}(\mathbf{x}_{k_0})\|^2 \\ &\quad + 2 \sum_{j \in \mathcal{J}_{k_0}} \|\nabla_j \mathbf{h}(\mathbf{x}_{k_0}) + \mathbf{r}_{k_0+1}^{(j)}\|^2. \end{aligned} \quad (5.16)$$

D'une part, puisque $\sum_{j=1}^J \|\nabla_j \mathbf{h}(\mathbf{x}_{k_0+1}) - \nabla_j \mathbf{h}(\mathbf{x}_{k_0})\|^2 = \|\nabla \mathbf{h}(\mathbf{x}_{k_0+1}) - \nabla \mathbf{h}(\mathbf{x}_{k_0})\|^2$, d'après l'hypothèse 5.1(ii) on a

$$\sum_{j \in \mathcal{J}_{k_0}} \|\nabla_j \mathbf{h}(\mathbf{x}_{k_0+1}) - \nabla_j \mathbf{h}(\mathbf{x}_{k_0})\|^2 \leq \beta^2 \|\mathbf{x}_{k_0+1} - \mathbf{x}_{k_0}\|^2. \quad (5.17)$$

D'autre part, puisque $j_{k_0} \in \mathcal{J}_{k_0}$,

$$\sum_{j \in \mathcal{J}_{k_0}} \|\nabla_j \mathbf{h}(\mathbf{x}_{k_0}) + \mathbf{r}_{k_0+1}^{(j)}\|^2 = \|\nabla_{j_{k_0}} \mathbf{h}(\mathbf{x}_{k_0}) + \mathbf{r}_{k_0+1}^{(j_{k_0})}\|^2 + \sum_{j \in \mathcal{J}_{k_0} \setminus \{j_{k_0}\}} \|\nabla_j \mathbf{h}(\mathbf{x}_{k_0}) + \mathbf{r}_{k_0+1}^{(j)}\|^2.$$

Par ailleurs, pour tout $j \in \mathcal{J}_{k_0} \setminus \{j_{k_0}\}$, on a $\mathbf{x}_{k_0+1}^{(j)} = \mathbf{x}_{k_0}^{(j)}$. Ainsi, d'après la deuxième inégalité de l'algorithme 5.5 et l'hypothèse 5.2(ii), on a

$$\sum_{j \in \mathcal{J}_{k_0}} \|\nabla_j \mathbf{h}(\mathbf{x}_{k_0}) + \mathbf{r}_{k_0+1}^{(j)}\|^2 \leq \alpha_2^2 \bar{\nu} \|\mathbf{x}_{k_0+1} - \mathbf{x}_{k_0}\|^2 + \sum_{j \in \mathcal{J}_{k_0} \setminus \{j_{k_0}\}} \|\nabla_j \mathbf{h}(\mathbf{x}_{k_0}) + \mathbf{r}_{k_0}^{(j)}\|^2. \quad (5.18)$$

Pour conclure, (5.15) peut être déduite en combinant les équation (5.16), (5.17) et (5.18). $\square$

Lemme 5.4. Soit $(\mathbf{x}_k)_{k \in \mathbb{N}}$ une suite générée par l'algorithme 5.5. Soient $k_0 \in \mathbb{N}$ et $k'_0 \in \mathbb{N}$ deux itérations telles que $k_0 \leq k'_0$. On définit le sous-ensemble $\mathcal{J}_{k_0, k'_0}$ de $\{1, \dots, J\}$ tel que, pour tout $k \in \{k_0, \dots, k'_0\}$, $j_k \in \mathcal{J}_{k_0, k'_0}$. On suppose que les hypothèses 5.1 et 5.2 sont satisfaites. On a alors

$$\begin{aligned} & \sum_{j \in \mathcal{J}_{k_0, k'_0}} \|\nabla_j \mathbf{h}(\mathbf{x}_{k'_0+1}) + \mathbf{r}_{k'_0+1}^{(j)}\|^2 \\ & \leq (\beta^2 + \alpha_2^2 \bar{\nu}) \sum_{k=k_0}^{k'_0} 2^{k'_0+1-k} \|\mathbf{x}_{k+1} - \mathbf{x}_k\|^2 + 2^{k'_0+1-k_0} \sum_{j \in \mathcal{J}_{k_0, k'_0} \setminus \{j_{k_0}\}} \|\nabla_j \mathbf{h}(\mathbf{x}_{k_0}) + \mathbf{r}_{k_0}^{(j)}\|^2, \end{aligned}$$

où $\mathbf{r}_{k'_0+1}^{(j_{k'_0})}$ est donné par l'algorithme 5.5, pour tout $j \in \mathcal{J}_{k_0, k'_0} \setminus \{j_{k'_0}\}$, $\mathbf{r}_{k'_0+1}^{(j)} \in \partial \mathbf{g}_j(\mathbf{x}_{k'_0+1}^{(j)})$ et, pour tout $j \in \mathcal{J}_{k_0, k'_0} \setminus \{j_{k_0}\}$, $\mathbf{r}_{k_0}^{(j)} \in \partial \mathbf{g}_j(\mathbf{x}_{k_0}^{(j)})$.

Démonstration. Soit $(k_0, k'_0) \in \mathbb{N}^2$ tel que $k_0 \leq k'_0$. Sous les hypothèses considérées, on peut

appliquer successivement le lemme 5.3 pour $k'_0, k'_0 - 1, \dots, k_0$. Ainsi, on obtient

$$\begin{aligned}
 & \sum_{j \in \mathcal{J}_{k_0, k'_0}} \|\nabla_j \mathbf{h}(\mathbf{x}_{k'_0+1}) + \mathbf{r}_{k'_0+1}^{(j)}\|^2 \\
 & \leq (\beta^2 + \alpha_2^2 \overline{\nu}) 2 \|\mathbf{x}_{k'_0+1} - \mathbf{x}_{k'_0}\|^2 + 2 \sum_{j \in \mathcal{J}_{k_0, k'_0} \setminus \{j_{k'_0}\}} \|\nabla_j \mathbf{h}(\mathbf{x}_{k'_0}) + \mathbf{r}_{k'_0}^{(j)}\|^2 \\
 & \leq (\beta^2 + \alpha_2^2 \overline{\nu}) 2 \|\mathbf{x}_{k'_0+1} - \mathbf{x}_{k'_0}\|^2 + 2 \sum_{j \in \mathcal{J}_{k_0, k'_0}} \|\nabla_j \mathbf{h}(\mathbf{x}_{k'_0}) + \mathbf{r}_{k'_0}^{(j)}\|^2 \\
 & \leq (\beta^2 + \alpha_2^2 \overline{\nu}) \left(2 \|\mathbf{x}_{k'_0+1} - \mathbf{x}_{k'_0}\|^2 + 2^2 \|\mathbf{x}_{k'_0} - \mathbf{x}_{k'_0-1}\|^2 \right) \\
 & \quad + 2^2 \sum_{j \in \mathcal{J}_{k_0, k'_0} \setminus \{j_{k'_0-1}\}} \|\nabla_j \mathbf{h}(\mathbf{x}_{k'_0-1}) + \mathbf{r}_{k'_0-1}^{(j)}\|^2 \\
 & \leq (\beta^2 + \alpha_2^2 \overline{\nu}) \left(2 \|\mathbf{x}_{k'_0+1} - \mathbf{x}_{k'_0}\|^2 + 2^2 \|\mathbf{x}_{k'_0} - \mathbf{x}_{k'_0-1}\|^2 + 2^3 \|\mathbf{x}_{k'_0-1} - \mathbf{x}_{k'_0-2}\|^2 \right) \\
 & \quad + 2^3 \sum_{j \in \mathcal{J}_{k_0, k'_0} \setminus \{j_{k'_0-2}\}} \|\nabla_j \mathbf{h}(\mathbf{x}_{k'_0-2}) + \mathbf{r}_{k'_0-2}^{(j)}\|^2 \\
 & \quad \vdots \\
 & \leq (\beta^2 + \alpha_2^2 \overline{\nu}) \sum_{k=k_0}^{k'_0} 2^{k'_0+1-k} \|\mathbf{x}_{k+1} - \mathbf{x}_k\|^2 \\
 & \quad + 2^{k'_0+1-k_0} \sum_{j \in \mathcal{J}_{k_0, k'_0} \setminus \{j_{k_0}\}} \|\nabla_j \mathbf{h}(\mathbf{x}_{k_0}) + \mathbf{r}_{k_0}^{(j)}\|^2.
 \end{aligned}$$

□

Nous pouvons maintenant énoncer notre résultat principal concernant la convergence de l'algorithme 5.5 :

Théorème 5.1. *Soit $(\mathbf{x}_k)_{k \in \mathbb{N}}$ une suite générée par l'algorithme 5.5. Supposons que les hypothèses 5.1-5.3 sont satisfaites. De plus, supposons que la fonction $\mathbf{f}$ satisfait l'inégalité de Kurdyka-Łojasiewicz donnée dans la définition 2.20. Alors, les assertions suivantes sont vérifiées :*

- (i) *La suite $(\mathbf{x}_k)_{k \in \mathbb{N}}$ converge vers un point critique $\widehat{\mathbf{x}}$ de $\mathbf{f}$.*
- (ii) *Cette suite vérifie*

$$\sum_{k=0}^{+\infty} \|\mathbf{x}_{k+1} - \mathbf{x}_k\| < +\infty.$$

- (iii) *La suite $(\mathbf{f}(\mathbf{x}_k))_{k \in \mathbb{N}}$ est décroissante et converge vers $\mathbf{f}(\widehat{\mathbf{x}})$.*

La démonstration du théorème 5.1 est donnée dans l'annexe 5.A. On peut en déduire le corollaire suivant, énonçant la convergence locale de l'algorithme 5.5 vers un minimiseur global de $\mathbf{f}$.

Corollaire 5.1. Soit $(\mathbf{x}_k)_{k \in \mathbb{N}}$ une suite générée par l'algorithme (5.5). Supposons que les hypothèses 5.1-5.3 sont satisfaites. De plus, supposons que f satisfait l'inégalité de Kurdyka-Łojasiewicz. Il existe $v \in]0, +\infty[$ tel que, si

$$f(\mathbf{x}_0) \leq \inf_{\mathbf{x} \in \mathbb{R}^N} f(\mathbf{x}) + v,$$

alors, $(\mathbf{x}_k)_{k \in \mathbb{N}}$ converge vers une solution du problème (5.1).

Démonstration. La démonstration de ce corollaire se fait de façon analogue à la démonstration du Corollaire 3.3 du chapitre 3. $\square$

5.5.3 Borne supérieure sur la vitesse de convergence

D'après le théorème 5.1, la limite $\hat{\mathbf{x}}$ de la suite $(\mathbf{x}_k)_{k \in \mathbb{N}}$ générée par l'algorithme (5.5) est un point critique de la fonction f , dès lors que les hypothèses 5.1-5.3 sont vérifiées et que f satisfait l'inégalité de KL. Ainsi, en procédant de la même façon que pour l'obtention de l'équation (5.29), il existe $\zeta \in]0, +\infty[$, $\kappa \in]0, +\infty[$ et $\theta \in [0, 1[$ tels que ((ii)) soit vérifiée pour tout $\mathbf{x} \in \mathbb{R}^N$ satisfaisant $G(\mathbf{x}) \leq G(\hat{\mathbf{x}}) + \zeta$. La valeur θ est alors appelée un *exposant de Łojasiewicz de G en $\hat{\mathbf{x}}$* . Comme d'autres algorithmes basés sur l'inégalité de KL [Attouch et Bolte, 2009; Attouch et al., 2010a], la vitesse de convergence locale de l'algorithme 5.5 dépend de cet exposant.

Le lemme ci-dessous nous sera utile pour déterminer la vitesse de convergence de l'algorithme 5.5 :

Lemme 5.5. [Attouch et Bolte, 2009] Soit $(\Lambda_m)_{m \in \mathbb{N}}$ une suite positive réelle qui décroît vers 0. Supposons qu'il existe $m^* \in \mathbb{N}^*$ et $C \in]0, +\infty[$ tels que, pour tout $m \geq m^*$,

$$\Lambda_m \leq (\Lambda_{m-1} - \Lambda_m) + C(\Lambda_{m-1} - \Lambda_m)^{\frac{1-\theta}{\theta}}, \quad (5.19)$$

où $\theta \in]0, 1[$. On observe deux cas :

(i) Si $\theta \in]\frac{1}{2}, 1[$, alors il existe $\lambda \in]0, +\infty[$ tel que

$$(\forall m \geq 1) \quad \Lambda_m \leq \lambda m^{-\frac{1-\theta}{2\theta-1}}.$$

(ii) Si $\theta \in]0, \frac{1}{2}]$, alors il existe $\lambda \in]0, +\infty[$ et $\tau \in [0, 1[$ tels que

$$(\forall m \in \mathbb{N}) \quad \Lambda_m \leq \lambda \tau^m.$$

Théorème 5.2. Soit $(\mathbf{x}_k)_{k \in \mathbb{N}}$ une suite générée par l'algorithme (5.5). Supposons que les hypothèses 5.1-5.3 sont satisfaites et que f satisfait l'inégalité de KL. Soit θ l'exposant de Łojasiewicz de G au point limite $\hat{\mathbf{x}}$ de la suite $(\mathbf{x}_k)_{k \in \mathbb{N}}$. Les propriétés suivantes sont alors satisfaites :

(i) Si $\theta \in]\frac{1}{2}, 1[$, alors il existe $(\lambda', \lambda'') \in]0, +\infty[^2$ tel que

$$(\forall k > K) \quad \|\mathbf{x}_k - \hat{\mathbf{x}}\| \leq \lambda' \left(\frac{k}{K} - 1 \right)^{-\frac{1-\theta}{2\theta-1}}, \quad (5.20)$$

$$(\forall k > 2K) \quad G(\mathbf{x}_k) - G(\hat{\mathbf{x}}) \leq \lambda'' \left(\frac{k}{K} - 2 \right)^{-\frac{1-\theta}{\theta(2\theta-1)}}. \quad (5.21)$$

(ii) Si $\theta \in]0, \frac{1}{2}]$, alors il existe $(\lambda', \lambda'') \in]0, +\infty[^2$ et $\tau' \in [0, 1[$ tels que

$$(\forall k \in \mathbb{N}) \quad \|\mathbf{x}_k - \hat{\mathbf{x}}\| \leq \lambda'(\tau')^k, \quad (5.22)$$

$$G(\mathbf{x}_k) - G(\hat{\mathbf{x}}) \leq \lambda''(\tau')^{\frac{k}{\theta}}. \quad (5.23)$$

(iii) Si $\theta = 0$, alors la suite $(\mathbf{x}_k)_{k \in \mathbb{N}}$ converge en un nombre fini d'itérations.

La démonstration du théorème 5.2 est donnée dans l'annexe 5.B.

Remarque 5.5.

- (i) Remarquons que, lorsque f est fortement convexe, l'exposant de Łojasiewicz θ de f est égal à $1/2$. Dans ce cas, $\hat{\mathbf{x}}$ est un minimiseur global de f et les suites $(\|\mathbf{x}_k - \hat{\mathbf{x}}\|)_{k \in \mathbb{N}}$ et $(f(\mathbf{x}_k) - f(\hat{\mathbf{x}}))_{k \in \mathbb{N}}$ convergent linéairement.
- (ii) On remarque que, si $\theta \in]0, 1/2]$, alors, pour m suffisamment grand, (5.19) implique

$$\Lambda_m \leq (1 + C)(\Lambda_{m-1} - \Lambda_m),$$

et donc la constante τ' apparaissant dans les équations (5.22)-(5.23) peut être choisie égale à $((1 + C)/(2 + C))^{1/K}$ où C est donné par (5.40).

Remarquons, d'autre part, que la convergence et la vitesse de convergence de l'algorithme 5.4 ont aussi été étudiés par [Frankel et al., 2014]. Les auteurs ont démontré la convergence de l'algorithme pour la résolution d'une classe de problèmes non-convexes incluant celui étudié dans ce chapitre. La principale différence existant entre ces deux travaux réside dans la règle de mise à jour des blocs. En effet, dans [Frankel et al., 2014], les auteurs présentent des propriétés de convergence de l'algorithme 5.4 lorsque les blocs sont mis à jour selon une règle cyclique. Les démonstrations de convergence ne sont donc pas les mêmes que celles présentées dans la section 5.5. De plus, nous verrons dans le chapitre 4 une application pour laquelle il s'est révélé très intéressant numériquement d'utiliser une règle essentiellement cyclique.

5.6 CONCLUSION

Dans ce chapitre nous avons proposé une méthode d'optimisation permettant de minimiser la somme d'une fonction différentiable h et d'une fonction non nécessairement lisse g additivement séparable par bloc. La méthode repose à la fois sur un algorithme explicite-implicite à métrique variable et sur une stratégie alternée par bloc. Nous avons vu dans les chapitres 3 et 4 que la métrique variable permet d'accélérer en pratique la convergence de l'algorithme explicite-implicite. Combiner cette méthode avec un principe de minimisation par bloc nous permet en plus de traiter des données de grande taille en réduisant le coup de calcul ainsi que l'allocation de mémoire par itération.

De plus, nous avons proposé une méthode basée sur le principe MM pour choisir la métrique et nous avons donné une règle générale acyclique de mise à jour des blocs. La convergence des

suites $(\mathbf{x}_k)_{k \in \mathbb{N}}$ générées par l'algorithme proposé, ainsi que par sa version inexacte, et du critère associé a été démontrée en utilisant l'inégalité de KL. Enfin, nous avons étudié la vitesse de convergence de $(\mathbf{x}_k)_{k \in \mathbb{N}}$ et du critère $(f(\mathbf{x}_k))_{k \in \mathbb{N}}$ en fonction de l'exposant de Łojasiewicz de la fonction $f = h + g$.

Afin de montrer l'intérêt pratique de l'algorithme 5.4 (et de sa version inexacte), nous en présenterons plusieurs exemples d'application dans le chapitre 6. En particulier, nous verrons que la mise à jour des blocs de façon acyclique permet dans certains cas d'accélérer de manière significative la convergence de l'algorithme.

5.A DÉMONSTRATION DU THÉORÈME 5.1

La démonstration du théorème 5.1 se fait en trois étapes : une première étape pour montrer que la suite $(f(\mathbf{x}_k))_{k \in \mathbb{N}}$ est convergente, une seconde étape pour démontrer la convergence de la suite $(\mathbf{x}_k)_{k \in \mathbb{N}}$ vers un point $\hat{\mathbf{x}} \in \mathbb{R}^N$, et une dernière étape permettant de s'assurer que $\hat{\mathbf{x}}$ est un point critique de f .

(i) D'après le lemme 5.1, on a

$$(\forall k \in \mathbb{N}) \quad f(\mathbf{x}_{k+1}) \leq f(\mathbf{x}_k),$$

donc, $(f(\mathbf{x}_k))_{k \in \mathbb{N}}$ est une suite décroissante. De plus, puisque $\mathbf{x}_0 \in \text{dom } g$, d'après la remarque 5.2(iii), les éléments de la suite $(\mathbf{x}_k)_{k \in \mathbb{N}}$ sont dans le sous-ensemble compact $E = \text{lev}_{\leq f(\mathbf{x}_0)} f$ de $\text{dom } g$, et f est bornée inférieurement. Ainsi, $(f(\mathbf{x}_k))_{k \in \mathbb{N}}$ converge vers un réel noté ξ , et $(f(\mathbf{x}_k) - \xi)_{k \in \mathbb{N}}$ est une suite positive convergeant vers 0.

(ii) Par ailleurs, d'après le lemme 5.2, on a

$$(\forall k \in \mathbb{N}) \quad \frac{\mu}{2} \|\chi_k\|^2 \leq (f(\mathbf{x}_k) - \xi) - (f(\mathbf{x}_{k+K}) - \xi). \quad (5.24)$$

Soit $\psi: [0, +\infty[\rightarrow [0, +\infty[: u \mapsto u^{\frac{1}{1-\theta}}$ une fonction convexe, où $\theta \in [0, 1[$. L'inégalité de gradient nous donne

$$(\forall (u, v) \in [0, +\infty[^2) \quad \psi(u) - \psi(v) \leq \dot{\psi}(u)(u - v),$$

qui, après un changement de variable, peut être réécrite

$$(\forall (u, v) \in [0, +\infty[^2) \quad u - v \leq (1 - \theta)^{-1} u^\theta (u^{1-\theta} - v^{1-\theta}).$$

En utilisant l'inégalité ci-dessus avec $u = f(\mathbf{x}_k) - \xi$ et $v = f(\mathbf{x}_{k+K}) - \xi$, on obtient

$$(\forall k \in \mathbb{N}) \quad (f(\mathbf{x}_k) - \xi) - (f(\mathbf{x}_{k+K}) - \xi) \leq (1 - \theta)^{-1} (f(\mathbf{x}_k) - \xi)^\theta \Delta_k,$$

où

$$(\forall k \in \mathbb{N}) \quad \Delta_k = (f(\mathbf{x}_k) - \xi)^{1-\theta} - (f(\mathbf{x}_{k+K}) - \xi)^{1-\theta}.$$

Donc, en combinant cette dernière inégalité avec l'équation (5.24), on a

$$(\forall k \in \mathbb{N}) \quad \|\chi_k\|^2 \leq 2\mu^{-1} (1 - \theta)^{-1} (f(\mathbf{x}_k) - \xi)^\theta \Delta_k. \quad (5.25)$$

Posons

$$(\forall k \in \mathbb{N}) \quad \mathbf{t}_k = \left(\nabla_j \mathbf{h}(\mathbf{x}_k) + \mathbf{r}_\ell^{(j)} \right)_{1 \leq j \leq J} \in \mathbb{R}^{N_1} \times \dots \times \mathbb{R}^{N_J}, \quad (5.26)$$

où, pour tout $j \in \{1, \dots, J\}$, $\mathbf{r}_k^{(j)} \in \partial \mathbf{g}_j(\mathbf{x}_k^{(j)})$. La règle de différentiabilité pour les fonctions séparables nous donne

$$\mathbf{r}_k = \left(\mathbf{r}_k^{(j)} \right)_{1 \leq j \leq J} \in \partial \mathbf{g}(\mathbf{x}_k). \quad (5.27)$$

Ainsi, pour tout $k \in \mathbb{N}$,

$$\mathbf{t}_k \in \partial \mathbf{f}(\mathbf{x}_k). \quad (5.28)$$

Puisque E est borné, d'après l'inégalité de KL, il existe trois constantes $\kappa > 0$, $\zeta > 0$ et $\theta \in [0, 1[$ telles que

$$(\forall \mathbf{t} \in \partial \mathbf{f}(\mathbf{x})) \quad \kappa |\mathbf{f}(\mathbf{x}) - \xi|^\theta \leq \|\mathbf{t}\|,$$

pour tout $\mathbf{x} \in E$ tel que $|\mathbf{f}(\mathbf{x}) - \xi| \leq \zeta$. La suite $(\mathbf{f}(\mathbf{x}_k))_{k \in \mathbb{N}}$ convergeant vers ξ , il existe $k^* \in \mathbb{N}$, tel que, pour tout $k \geq k^*$, $|\mathbf{f}(\mathbf{x}_k) - \xi| < \zeta$. Par conséquent, on a

$$(\forall k \geq k^*) \quad \kappa |\mathbf{f}(\mathbf{x}_k) - \xi|^\theta \leq \|\mathbf{t}_k\|, \quad (5.29)$$

où $\mathbf{t}_k$ est donné par (5.26). Soit K , défini par la définition 5.1(ii), le nombre maximal d'itérations au cours desquelles chaque bloc doit être mis à jour au moins une fois. Pour tout $k \in \mathbb{N}$,

$$\|\mathbf{t}_{k+K}\|^2 = \left\| \left(\nabla_j \mathbf{h}(\mathbf{x}_{k+K}) + \mathbf{r}_{k+K}^{(j)} \right)_{1 \leq j \leq J} \right\|^2 = \sum_{j=1}^J \left\| \nabla_j \mathbf{h}(\mathbf{x}_{k+K}) + \mathbf{r}_{k+K}^{(j)} \right\|^2.$$

Soit $l_{k, \sigma_k(J)}$ donné par (5.7). Pour tout $l \in \{k + l_{k, \sigma_k(J)}, \dots, k + K - 1\}$, soit $\mathbf{r}_{l+1}^{(j_l)} \in \partial \mathbf{g}_{j_l}(\mathbf{x}_{l+1}^{(j_l)})$ défini par l'algorithme 5.5. On peut alors appliquer le lemme 5.4 avec $k_0 = k + l_{k, \sigma_k(J)}$, $k'_0 = k + K - 1$ et $\mathcal{J}_{k_0, k'_0} = \{1, \dots, J\}$:

$$\begin{aligned} \|\mathbf{t}_{k+K}\|^2 &\leq (\beta^2 + \alpha_2^2 \overline{\nu}) \sum_{l=k+l_{k, \sigma_k(J)}}^{k+K-1} 2^{k+K-l} \|\mathbf{x}_{l+1} - \mathbf{x}_l\|^2 \\ &\quad + 2^{K-l_{k, \sigma_k(J)}} \sum_{\substack{j=1 \\ j \neq \sigma_k(J)}}^J \left\| \nabla_j \mathbf{h}(\mathbf{x}_{k+l_{k, \sigma_k(J)}}) + \mathbf{r}_{k+l_{k, \sigma_k(J)}}^{(j)} \right\|^2. \end{aligned}$$

En appliquant une nouvelle fois le lemme 5.4 à la somme

$$\sum_{\substack{j=1 \\ j \neq \sigma_k(J)}}^J \left\| \nabla_j \mathbf{h}(\mathbf{x}_{k+l_{k, \sigma_k(J)}}) + \mathbf{r}_{k+l_{k, \sigma_k(J)}}^{(j)} \right\|^2$$

avec $k_0 = k + l_{k, \sigma_k(J-1)}$, $k'_0 = k + l_{k, \sigma_k(J)} - 1$ et $\mathcal{J}_{k_0, k'_0} = \{1, \dots, J\} \setminus \{\sigma_k(J)\}$, on obtient

$$\begin{aligned} \|\mathbf{t}_{k+K}\|^2 &\leq (\beta^2 + \alpha_2^2 \bar{\nu}) \sum_{l=k+l_{k, \sigma_k(J)}}^{k+K-1} 2^{k+K-l} \|\mathbf{x}_{l+1} - \mathbf{x}_l\|^2 \\ &\quad + (\beta^2 + \alpha_2^2 \bar{\nu}) \sum_{l=k+l_{k, \sigma_k(J-1)}}^{k+l_{k, \sigma_k(J)}-1} 2^{k+K-l} \|\mathbf{x}_{l+1} - \mathbf{x}_l\|^2 \\ &\quad + 2^{K-l_{k, \sigma_k(J-1)}} \sum_{\substack{j=1 \\ j \neq \sigma_k(i), i \in \{J-1, J\}}}^J \left\| \nabla_j \mathbf{h}(\mathbf{x}_{k+l_{k, \sigma_k(J-1)}}) + \mathbf{r}_{k+l_{k, \sigma_k(J-1)}}^{(j)} \right\|^2. \end{aligned}$$

En réitérant le procédé de façon similaire pour $i \in \{1, \dots, J-2\}$, on a

$$\begin{aligned} \|\mathbf{t}_{k+K}\|^2 &\leq (\beta^2 + \alpha_2^2 \bar{\nu}) \sum_{l=k+l_{k, \sigma_k(J)}}^{k+K-1} 2^{k+K-l} \|\mathbf{x}_{l+1} - \mathbf{x}_l\|^2 \\ &\quad + (\beta^2 + \alpha_2^2 \bar{\nu}) \sum_{i=1}^{J-1} \sum_{l=k+l_{k, \sigma_k(j)}}^{k+l_{k, \sigma_k(j+1)}-1} 2^{k+K-l} \|\mathbf{x}_{l+1} - \mathbf{x}_l\|^2, \quad (5.30) \end{aligned}$$

où l'on a utilisé le fait que $\{1, \dots, J\} \setminus \{\sigma_k(1), \dots, \sigma_k(J)\} = \emptyset$. Ainsi

$$\sum_{\substack{j=1 \\ j \neq \sigma_k(i), i \in \{1, \dots, J\}}}^J \left\| \nabla_j \mathbf{h}(\mathbf{x}_k) + \mathbf{r}_k^{(j)} \right\|^2 = 0.$$

Comme $l_{k, \sigma_k(1)} = 0$ et, pour tout $l \in \{k, \dots, k+K-1\}$, $2^{k+K-l} \leq 2^K$, il découle de (5.14) et (5.30) que

$$(\forall k \in \mathbb{N}) \quad \|\mathbf{t}_{k+K}\|^2 \leq 2^K (\beta^2 + \alpha_2^2 \bar{\nu}) \sum_{l=k}^{k+K-1} \|\mathbf{x}_{l+1} - \mathbf{x}_l\|^2 = 2^K (\beta^2 + \alpha_2^2 \bar{\nu}) \|\mathbf{x}_k\|^2. \quad (5.31)$$

En combinant les équations (5.25), (5.29) et (5.31) on obtient

$$(\forall k \geq \max\{k^*, K\}) \quad \|\mathbf{x}_k\|^2 \leq 2\mu^{-1}(1-\theta)^{-1}\kappa^{-1}2^{K/2}(\beta^2 + \alpha_2^2 \bar{\nu})^{1/2} \|\mathbf{x}_{k-K}\| \Delta_k.$$

En utilisant l'inégalité

$$(\forall (\mathbf{u}, \mathbf{v}) \in [0, +\infty[^2) \quad (\mathbf{u}\mathbf{v})^{1/2} \leq \frac{1}{2}(\mathbf{u} + \mathbf{v}),$$

et en posant $\mathbf{u} = \|\mathbf{x}_{k-K}\|$ et $\mathbf{v} = 2\mu^{-1}(1-\theta)^{-1}\kappa^{-1}2^{K/2}(\beta^2 + \alpha_2^2 \bar{\nu})^{1/2} \Delta_k$, on obtient

$$(\forall k \geq \max\{k^*, K\}) \quad \|\mathbf{x}_k\| \leq \frac{1}{2} \|\mathbf{x}_{k-K}\| + \mu^{-1}(1-\theta)^{-1}\kappa^{-1}2^{K/2}(\beta^2 + \alpha_2^2 \bar{\nu})^{1/2} \Delta_k. \quad (5.32)$$

Par ailleurs, remarquons que

$$\begin{aligned} \sum_{k=k^*}^{+\infty} \Delta_k &= \sum_{k=k^*}^{+\infty} (f(\mathbf{x}_k) - \xi)^{1-\theta} - (f(\mathbf{x}_{k+K}) - \xi)^{1-\theta} \\ &= \sum_{k=k^*}^{k^*+K-1} (f(\mathbf{x}_k) - \xi)^{1-\theta}. \end{aligned}$$

On peut donc en déduire que $(\Delta_k)_{k \in \mathbb{N}}$ est une suite sommable. De plus, $(\|\mathbf{x}_k\|)_{k \geq \max\{k^*, K\}}$ satisfaisant l'inégalité (5.32), $(\|\mathbf{x}_k\|)_{k \in \mathbb{N}}$ est aussi une suite sommable. Enfin, d'après l'équation (5.14),

$$(\forall k \in \mathbb{N}) \quad \|\mathbf{x}_{k+1} - \mathbf{x}_k\| \leq \|\mathbf{x}_k\|,$$

donc $(\|\mathbf{x}_{k+1} - \mathbf{x}_k\|)_{k \in \mathbb{N}}$ est une suite sommable.

Par conséquent, la suite $(\mathbf{x}_k)_{k \in \mathbb{N}}$ satisfait l'assertion (ii). De plus, cette condition implique que $(\mathbf{x}_k)_{k \in \mathbb{N}}$ est une suite de Cauchy. Ainsi, on peut en déduire qu'elle converge vers un point $\hat{\mathbf{x}}$.

- (iii) Il nous reste à démontrer que la limite $\hat{\mathbf{x}}$ est un point critique de f . D'après (5.28), on a, pour tout $k \in \mathbb{N}$,

$$(\mathbf{x}_k, \mathbf{t}_k) \in \text{Graph } \partial f.$$

Par ailleurs, la suite $(\|\mathbf{x}_k\|)_{k \in \mathbb{N}}$ étant sommable, elle converge vers 0. De plus, l'équation (5.31) implique que

$$(\forall k \in \mathbb{N}) \quad \|\mathbf{t}_k\| \leq 2^{K/2}(\beta^2 + \alpha_2^2 \bar{\nu})^{1/2} \|\mathbf{x}_{k-K}\|,$$

donc $(\mathbf{x}_k, \mathbf{t}_k)_{k \in \mathbb{N}}$ converge vers $(\hat{\mathbf{x}}, \mathbf{0})$. D'autre part, la remarque 5.2(iii) assure que la restriction de f à son domaine de définition est continue. Ainsi, puisque, pour tout $k \in \mathbb{N}$, $\mathbf{x}_k \in \text{dom } f$, la suite $(f(\mathbf{x}_k))_{k \in \mathbb{N}}$ converge vers $f(\hat{\mathbf{x}})$. Finalement, en utilisant la propriété de fermeture de ∂f (c.f. remarque 2.4), on peut conclure que $(\hat{\mathbf{x}}, \mathbf{0}) \in \text{Graph } \partial f$ i.e., $\hat{\mathbf{x}}$ est un point critique de f .

5.B DÉMONSTRATION DU THÉORÈME 5.2

Nous utiliserons dans cette démonstration les mêmes notations que dans la démonstration du théorème 5.1.

Soit K donné par la définition 5.1(ii). Pour tout $k \in \mathbb{N}$, il existe $m \in \mathbb{N}$ et $l \in \{0, \dots, K-1\}$ tels que $k = mK + l$. D'après l'inégalité triangulaire, on a alors

$$\|\mathbf{x}_k - \hat{\mathbf{x}}\| \leq \|\mathbf{x}_{mK} - \hat{\mathbf{x}}\| + \|\mathbf{x}_k - \mathbf{x}_{mK}\|. \quad (5.33)$$

En utilisant une nouvelle fois l'inégalité triangulaire, on obtient

$$\begin{aligned}
 \|\mathbf{x}_{mK} - \widehat{\mathbf{x}}\| &= \left\| \sum_{p=m}^{+\infty} (\mathbf{x}_{(p+1)K} - \mathbf{x}_{pK}) \right\| \\
 &= \left\| \sum_{p=m}^{+\infty} \sum_{l'=0}^{K-1} (\mathbf{x}_{pK+l'+1} - \mathbf{x}_{pK+l'}) \right\| \\
 &\leq \sum_{p=m}^{+\infty} \left\| \sum_{l'=0}^{K-1} (\mathbf{x}_{pK+l'+1} - \mathbf{x}_{pK+l'}) \right\|,
 \end{aligned} \tag{5.34}$$

et selon l'inégalité de Jensen et (5.14),

$$(\forall p \geq m) \quad \left\| \sum_{l'=0}^{K-1} (\mathbf{x}_{pK+l'+1} - \mathbf{x}_{pK+l'}) \right\|^2 \leq K \|\chi_{pK}\|^2. \tag{5.35}$$

Pour tout $m' \in \mathbb{N}$, soit $\Lambda_{m'} = \sum_{p=m'}^{+\infty} \|\chi_{pK}\|$, qui est une somme finie d'après le théorème 5.1. Par conséquent, en combinant les deux inégalités ci-dessus, on obtient

$$\|\mathbf{x}_{mK} - \widehat{\mathbf{x}}\| \leq \sqrt{K} \Lambda_m. \tag{5.36}$$

En appliquant de nouveau l'inégalité de Jensen, on a

$$\begin{aligned}
 \|\mathbf{x}_{mK} - \mathbf{x}_k\|^2 &= \left\| \sum_{l'=0}^{l-1} (\mathbf{x}_{mK+l'+1} - \mathbf{x}_{mK+l'}) \right\|^2 \\
 &\leq l \sum_{l'=0}^{l-1} \|\mathbf{x}_{mK+l'+1} - \mathbf{x}_{mK+l'}\|^2 \leq (K-1) \|\chi_{mK}\|^2.
 \end{aligned} \tag{5.37}$$

En combinant les équations (5.33), (5.36), et (5.37), on obtient alors

$$(\forall k \in \mathbb{N}) \quad \|\mathbf{x}_k - \widehat{\mathbf{x}}\| \leq \sqrt{K} \Lambda_m + \sqrt{K-1} \|\chi_{mK}\| \leq 2\sqrt{K} \Lambda_m. \tag{5.38}$$

D'après l'équation (5.32), on a, pour tout $m \geq \max\{k^*/K, 1\}$,

$$\|\chi_{mK}\| \leq \frac{1}{2} \|\chi_{(m-1)K}\| + \mu^{-1} (1-\theta)^{-1} \kappa^{-1} 2^{K/2} (\beta^2 + \alpha_2^2 \overline{\nu})^{1/2} \Delta_{mK},$$

où $\Delta_{mK} = (f(\mathbf{x}_{mK}) - f(\widehat{\mathbf{x}}))^{1-\theta} - (f(\mathbf{x}_{(m+1)K}) - f(\widehat{\mathbf{x}}))^{1-\theta}$. Ainsi, puisque $(f(\mathbf{x}_k) - f(\widehat{\mathbf{x}}))_{k \in \mathbb{N}}$ est une suite positive qui converge vers 0, on obtient

$$\Lambda_m \leq (\Lambda_{m-1} - \Lambda_m) + 2\mu^{-1} (1-\theta)^{-1} \kappa^{-1} 2^{K/2} (\beta^2 + \alpha_2^2 \overline{\nu})^{1/2} (f(\mathbf{x}_{mK}) - f(\widehat{\mathbf{x}}))^{1-\theta}.$$

Supposons maintenant que $\theta \neq 0$. Selon les équations (5.29) et (5.31), on a

$$\kappa (f(\mathbf{x}_{mK}) - f(\widehat{\mathbf{x}}))^\theta \leq (2^K (\beta^2 + \alpha_2^2 \overline{\nu}))^{1/2} \|\chi_{(m-1)K}\|,$$

ce qui implique

$$(f(\mathbf{x}_{mK}) - f(\widehat{\mathbf{x}}))^{1-\theta} \leq \kappa^{-\frac{1-\theta}{\theta}} (2^K (\beta^2 + \alpha_2^2 \overline{\nu}))^{\frac{1-\theta}{2\theta}} \|\chi_{(m-1)K}\|^{\frac{1-\theta}{\theta}}. \tag{5.39}$$

Donc, en posant

$$C = 2\mu^{-1}(1-\theta)^{-1}\kappa^{-\frac{1}{\theta}}(2^K(\beta^2 + \alpha_2^2\bar{\nu}))^{\frac{1}{2\theta}}, \quad (5.40)$$

nous obtenons, pour tout $m \geq \max\{k^*/K, 1\}$,

$$\Lambda_m \leq (\Lambda_{m-1} - \Lambda_m) + C\|\chi_{(m-1)K}\|^{\frac{1-\theta}{\theta}}.$$

Ainsi, (5.19) est satisfaite.

Par conséquent, d'après le lemme 5.5 et l'équation (5.38), si $\theta \in]\frac{1}{2}, 1[$, il existe $\lambda \in]0, +\infty[$ tel que

$$(\forall k > K) \quad \|\mathbf{x}_k - \widehat{\mathbf{x}}\| \leq 2\sqrt{K}\lambda m^{-\frac{1-\theta}{2\theta-1}} \leq 2\sqrt{K}\lambda \left(\frac{k}{K} - 1\right)^{-\frac{1-\theta}{2\theta-1}},$$

où m correspond à la partie entière inférieure de k/K . L'inégalité (5.20) est alors obtenue en posant $\lambda' = 2\sqrt{K}\lambda$. De façon analogue, si $\theta \in]0, \frac{1}{2}]$, alors il existe $\lambda \in]0, +\infty[$ et $\tau \in [0, 1[$ tels que

$$(\forall k > K) \quad \|\mathbf{x}_k - \widehat{\mathbf{x}}\| \leq 2\sqrt{K}\lambda\tau^m \leq 2\sqrt{K}\lambda\tau^{k/K-1}.$$

Ainsi, si $\tau \neq 0$, (5.22) est satisfaite en posant $\lambda' = 2\sqrt{K}\lambda/\tau$ et $\tau' = \tau^{1/K}$, tandis que si $\tau = 0$, (5.22) est vérifiée de façon immédiate.

De plus, comme $(f(\mathbf{x}_k) - f(\widehat{\mathbf{x}}))_{k \in \mathbb{N}}$ est une suite décroissante, pour tout $k \in \mathbb{N}$,

$$f(\mathbf{x}_k) - f(\widehat{\mathbf{x}}) \leq f(\mathbf{x}_{mK}) - f(\widehat{\mathbf{x}}),$$

où m correspond à la partie entière inférieure de k/K . D'après (5.39), si $m \geq \max\{k^*/K, 1\}$, alors

$$\begin{aligned} f(\mathbf{x}_k) - f(\widehat{\mathbf{x}}) &\leq \kappa^{-1/\theta} \left(2^K(\beta^2 + \alpha_2^2\bar{\nu})\right)^{\frac{1}{2\theta}} \|\chi_{(m-1)K}\|^{1/\theta} \\ &\leq \kappa^{-1/\theta} \left(2^K(\beta^2 + \alpha_2^2\bar{\nu})\right)^{\frac{1}{2\theta}} \Lambda_{m-1}^{1/\theta}. \end{aligned}$$

Ainsi, si $\theta \in]\frac{1}{2}, 1[$, en appliquant de nouveau le lemme 5.5, il existe $\lambda \in]0, +\infty[$ tel que, pour tout $m > 2$,

$$\begin{aligned} f(\mathbf{x}_k) - f(\widehat{\mathbf{x}}) &\leq \kappa^{-1/\theta} \left(2^K(\beta^2 + \alpha_2^2\bar{\nu})\right)^{\frac{1}{2\theta}} \lambda(m-1)^{-\frac{1-\theta}{\theta(2\theta-1)}} \\ &\leq \kappa^{-1/\theta} \left(2^K(\beta^2 + \alpha_2^2\bar{\nu})\right)^{\frac{1}{2\theta}} \lambda \left(\frac{k}{K} - 2\right)^{-\frac{1-\theta}{\theta(2\theta-1)}}. \end{aligned}$$

Par conséquent, on peut trouver $\lambda'' \in]0, +\infty[$ tel que, pour tout $k > 2K$, (5.21) soit vérifiée. D'autre part, d'après le lemme 5.5, si $\theta \in]0, \frac{1}{2}]$, il existe $\lambda \in]0, +\infty[$ et $\tau \in [0, 1[$ tels que

$$\begin{aligned} f(\mathbf{x}_k) - f(\widehat{\mathbf{x}}) &\leq \kappa^{-1/\theta} \left(2^K(\beta^2 + \alpha_2^2\bar{\nu})\right)^{\frac{1}{2\theta}} \lambda\tau^{\frac{m-1}{\theta}} \\ &\leq \kappa^{-1/\theta} \left(2^K(\beta^2 + \alpha_2^2\bar{\nu})\right)^{\frac{1}{2\theta}} \lambda\tau^{\frac{k/K-2}{\theta}}. \end{aligned}$$

On peut donc trouver $\lambda'' \in]0, +\infty[$ tel que, pour tout $k \in \mathbb{N}$, (5.23) soit vérifiée.

Nous allons maintenant démontrer la propriété (iii) en supposant que $\theta = 0$. Posons $\mathcal{L} = \{k \in \mathbb{N} \mid \mathbf{x}_k \neq \widehat{\mathbf{x}}\}$. Soit $k \in \mathcal{L}$ tel que $k \geq \max\{k^*, K\}$. D'après les lemmes 5.1 et 5.2,

$$f(\mathbf{x}_{k+1}) \leq f(\mathbf{x}_k) - \frac{\mu}{2} \|\mathbf{x}_{k+1} - \mathbf{x}_k\|^2 \leq f(\mathbf{x}_{k-K}) - \frac{\mu}{2} \|\mathbf{x}_{k-K}\|^2.$$

En utilisant l'équation (5.31), on obtient alors

$$f(\mathbf{x}_k) - f(\widehat{\mathbf{x}}) - \frac{\mu}{2} \|\mathbf{x}_{k+1} - \mathbf{x}_k\|^2 \leq f(\mathbf{x}_{k-K}) - f(\widehat{\mathbf{x}}) - \frac{\mu'}{2} \|\mathbf{t}_k\|^2,$$

où $\mu' \in]0, +\infty[$. Puisque $\theta = 0$, en combinant l'inégalité ci-dessus avec (5.29), on a

$$f(\mathbf{x}_k) - f(\widehat{\mathbf{x}}) - \frac{\mu}{2} \|\mathbf{x}_{k+1} - \mathbf{x}_k\|^2 \leq f(\mathbf{x}_{k-K}) - f(\widehat{\mathbf{x}}) - \frac{\mu'}{2} \kappa^2 |f(\mathbf{x}_k) - f(\widehat{\mathbf{x}})|^0,$$

c'est-à-dire,

$$f(\mathbf{x}_k) - f(\widehat{\mathbf{x}}) - \frac{\mu}{2} \|\mathbf{x}_{k+1} - \mathbf{x}_k\|^2 \leq f(\mathbf{x}_{k-K}) - f(\widehat{\mathbf{x}}) - \frac{\mu'}{2} \kappa^2.$$

Comme $\lim_{k \rightarrow +\infty} f(\mathbf{x}_k) = f(\widehat{\mathbf{x}})$, cette dernière inégalité implique que $\mathcal{L}$ est un ensemble fini. La propriété (iii) en découle alors de façon immédiate.

CHAPITRE 6

Quelques applications de l'algorithme explicite-implicite à métrique variable alterné par bloc

Sommaire

6.1	Introduction	118
6.2	Problème de démixage spectral	118
6.2.1	Formulation du problème	118
6.2.2	Mise en œuvre de l'algorithme	120
6.2.3	Résultats de simulation	122
6.3	Reconstruction de phase	123
6.3.1	Formulation du problème	123
6.3.2	Mise en œuvre de l'algorithme	127
6.3.3	Résultats de simulation	129
6.4	Déconvolution aveugle de signaux sismiques	132
6.4.1	Formulation du problème	132
6.4.2	Mise en œuvre de l'algorithme	137
6.4.3	Résultats de simulation	139
6.5	Conclusion	141

6.1 INTRODUCTION

Nous avons proposé dans le chapitre précédent un algorithme permettant de minimiser la somme d'une fonction différentiable de gradient Lipschitz et d'une fonction séparable par bloc. Cette méthode combine une technique d'accélération utilisant une métrique variable et un principe de minimisation alternée par bloc. Nous allons maintenant appliquer cet algorithme à la résolution de problèmes inverses faisant intervenir un critère non lisse et non convexe, afin d'étudier son comportement pratique. En particulier nous allons montrer l'utilité de la métrique variable pour accélérer la vitesse de convergence de la méthode. Nous allons également étudier l'influence du choix de la taille des blocs ainsi que de l'ordre de leur mise à jour.

Pour cela, nous allons considérer trois exemples applicatifs différents. La première application présentée dans la section 6.2, est un problème de démélange spectral. Ensuite, la section 6.3 traitera d'une application en reconstruction de phase. Enfin, dans la section 6.4, nous présenterons un exemple de déconvolution aveugle de signaux sismiques. Nous concluons dans la section 6.5.

6.2 PROBLÈME DE DÉMÉLANGE SPECTRAL

6.2.1 Formulation du problème

Considérons un ensemble de données hyperspectrales, noté $\mathbf{Y} \in \mathbb{R}^{S \times M}$, modélisant un ensemble d'images de $\mathbb{R}^M$ (réorganisées en colonnes avec $M = M_1 \times M_2$), acquises en S différentes bandes spectrales. On suppose que l'on observe le modèle de mélange linéaire suivant :

$$\mathbf{Y} = \bar{\mathbf{U}}\bar{\mathbf{V}} + \mathbf{E}, \quad (6.1)$$

où les colonnes de $\bar{\mathbf{U}} \in \mathbb{R}^{S \times P}$ représentent les spectres des P composantes distinctes présentes dans l'image, $\bar{\mathbf{V}} \in \mathbb{R}^{P \times M}$ leurs proportions respectives (abondances) en chaque pixel, et $\mathbf{E} \in \mathbb{R}^{S \times M}$ le bruit d'acquisition. Ce type de modèle se rencontre dans diverses applications telles que la télédétection [Asner et Heidebrecht, 2002], la planétologie [Themelis et al., 2012], le contrôle alimentaire [Gowen et al., 2007], ou la spectromicroscopie [Dobigeon et Brun, 2012]. Une illustration de ce modèle est donnée dans la figure 6.1. L'objectif du problème de démélange spectral est d'estimer $\bar{\mathbf{U}}$ et $\bar{\mathbf{V}}$ à partir des données observées $\mathbf{Y}$ [Bioucas-Dias et al., 2012].

Les méthodes standards de démélange spectral reposent sur des approches supervisées, au sens où les spectres composant $\bar{\mathbf{U}}$ sont supposés être issus d'un dictionnaire connu ou bien avoir été préalablement estimés par un algorithme d'extraction de spectres [Bioucas-Dias et al., 2012; Chouzenoux et al., 2014a; Keshava et Mustard, 2002; Ma et al., 2014; Moussaoui et al., 2012]. Cependant, il existe un intérêt croissant pour des méthodes permettant d'estimer conjointement $\bar{\mathbf{U}}$ et $\bar{\mathbf{V}}$ basées sur la factorisation de matrices non-négative (*Nonnegative Matrix Factorization* (NMF)) [Ma et al., 2014; Tomazeli Duarte et al., 2014].

Dans l'approche standard de NMF, les estimations de $\bar{\mathbf{U}}$ et $\bar{\mathbf{V}}$ sont définies comme étant des solutions du problème de minimisation du critère de moindres carrés associé au modèle (6.1) sous

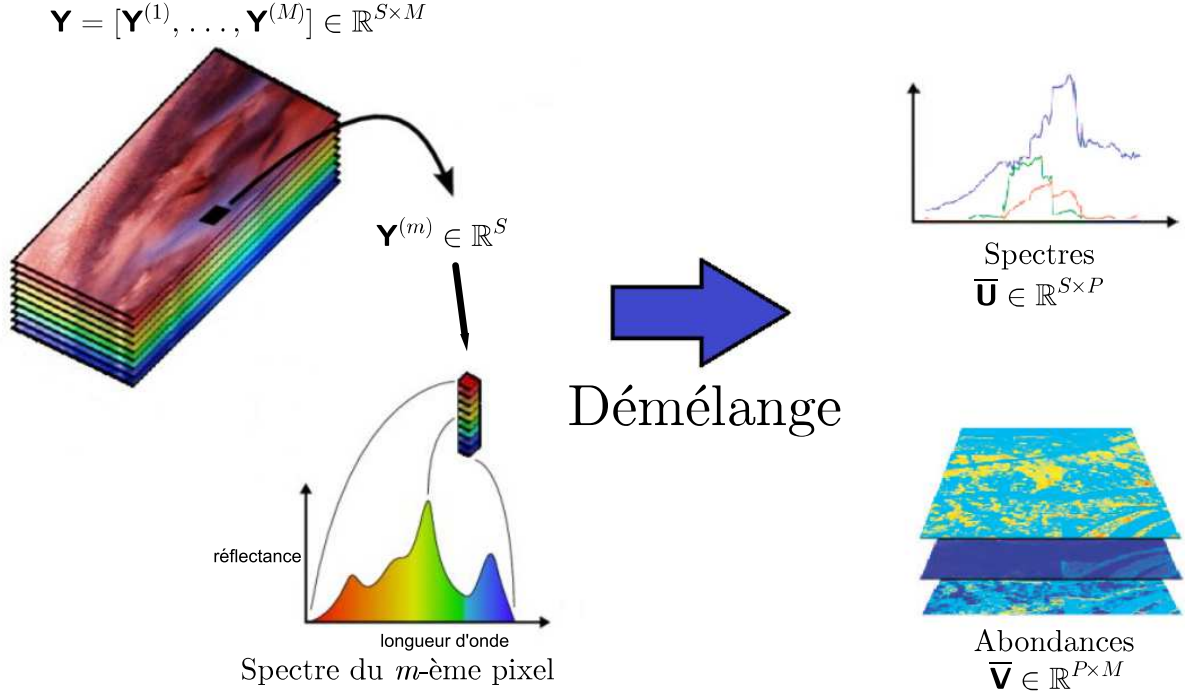


Figure 6.1 – Illustration du problème de démelange spectral. L'objectif est d'estimer $\bar{\mathbf{U}}$ et $\bar{\mathbf{V}}$ à partir d'un ensemble de données hyperspectrales $\mathbf{Y}$.

des contraintes de positivité sur les éléments des deux matrices [Paatero et Tapper, 1994] :

$$\text{trouver } (\hat{\mathbf{U}}, \hat{\mathbf{V}}) \in \underset{(\mathbf{U}, \mathbf{V}) \in \mathbb{R}^{S \times P} \times \mathbb{R}^{P \times M}}{\text{Argmin}} \left\{ \frac{1}{2} \|\mathbf{UV} - \mathbf{Y}\|^2 + \iota_{[0, +\infty[^{S \times P}}(\mathbf{U}) + \iota_{[0, +\infty[^{P \times M}}(\mathbf{V}) \right\}.^1 \quad (6.2)$$

Le critère à minimiser n'est pas convexe et présente des minimas locaux [Lee et Seung, 2001], ainsi, les contraintes de positivité sur les matrices $\bar{\mathbf{U}}$ et $\bar{\mathbf{V}}$ ne sont pas toujours suffisantes pour obtenir des résultats corrects [Chu et al., 2004]. Cependant, comme il est souligné dans [Jia et Qian, 2009; Ma et al., 2014], une nette amélioration de leurs qualités peut être obtenue lorsque des connaissances *a priori* sur les matrices $\bar{\mathbf{U}}$ et $\bar{\mathbf{V}}$ sont introduites dans le problème de minimisation (6.2).

Dans cette section, nous nous focaliserons sur le cas où chaque spectre composant les colonnes de $\bar{\mathbf{U}}$ est une combinaison linéaire de quelques spectres de référence issus d'un grand dictionnaire de taille $Q > P$, que nous noterons $\mathbf{\Omega} \in \mathbb{R}^{S \times Q}$. Ainsi, le modèle d'observation (6.1) est ré-exprimé de la façon suivante :

$$\mathbf{Y} = \mathbf{\Omega} \bar{\mathbf{T}} \bar{\mathbf{V}} + \mathbf{E}, \quad (6.3)$$

où $\bar{\mathbf{T}} \in \mathbb{R}^{Q \times P}$ est une matrice supposée parcimonieuse donnant les coefficients de la combinaison linéaire permettant d'obtenir $\bar{\mathbf{U}} = \mathbf{\Omega} \bar{\mathbf{T}}$.

1. Ici $\|\cdot\|$ désigne la norme de Frobenius.

Nous définissons les estimées $\hat{\mathbf{T}}$ et $\hat{\mathbf{V}}$ de $\bar{\mathbf{T}}$ et $\bar{\mathbf{V}}$ comme étant des solutions du problème

$$\text{trouver } (\hat{\mathbf{T}}, \hat{\mathbf{V}}) \in \underset{(\mathbf{T}, \mathbf{V}) \in \mathbb{R}^{Q \times P} \times \mathbb{R}^{P \times M}}{\text{Argmin}} \quad h(\mathbf{T}, \mathbf{V}) + g_1(\mathbf{T}) + g_2(\mathbf{V}), \quad (6.4)$$

où $h: \mathbb{R}^{Q \times P} \times \mathbb{R}^{P \times M} \rightarrow \mathbb{R}$ correspond au critère de moindre carré associé au modèle (6.3) défini par

$$(\forall (\mathbf{T}, \mathbf{V}) \in \mathbb{R}^{Q \times P} \times \mathbb{R}^{P \times M}) \quad h(\mathbf{T}, \mathbf{V}) = \frac{1}{2} \|\Omega \mathbf{T} \mathbf{V} - \mathbf{Y}\|^2, \quad (6.5)$$

et la fonction $g_1: \mathbb{R}^{Q \times P} \rightarrow]-\infty, +\infty]$ (resp. $g_2: \mathbb{R}^{P \times M} \rightarrow]-\infty, +\infty]$) est une fonction de régularisation propre et semi-continue inférieurement, incorporant des informations *a priori* sur $\hat{\mathbf{T}}$ (resp. $\hat{\mathbf{V}}$).

Afin de promouvoir la parcimonie de $\hat{\mathbf{T}}$, nous proposons de choisir

$$(\forall \mathbf{T} = (\mathbf{T}^{(q,p)})_{1 \leq q \leq Q, 1 \leq p \leq P} \in \mathbb{R}^{Q \times P})$$

$$g_1(\mathbf{T}) = \sum_{q=1}^Q \sum_{p=1}^P \left(\iota_{[\mathbf{T}_{\min}, \mathbf{T}_{\max}]}(\mathbf{T}^{(q,p)}) + \phi_{\delta, \beta}(\mathbf{T}^{(q,p)}) \right), \quad (6.6)$$

où $(\delta, \mathbf{T}_{\min}, \mathbf{T}_{\max}) \in]0, +\infty[^3$ et $\phi_{\delta, \beta}$ est la fonction de régularisation donnée dans la proposition 2.12 [Chartrand, 2012] de paramètre $\beta \in]0, 1]$. Rappelons que cette fonction est convexe si et seulement si $\beta = 1$, i.e. lorsque $\phi_{\delta, \beta}$ correspond à la fonction valeur absolue. D'autre part, on définit

$$g_2 = \iota_{\mathcal{V}}, \quad (6.7)$$

où $\mathcal{V} \subset \mathbb{R}^{P \times M}$ est donné par

$$\mathcal{V} = \left\{ \mathbf{V} = (\mathbf{V}^{(p,m)})_{1 \leq p \leq P, 1 \leq m \leq M} \in \mathbb{R}^{P \times M} \mid (\forall m \in \{1, \dots, M\}) \quad \sum_{p=1}^P \mathbf{V}^{(p,m)} = 1, \right.$$

$$\left. (\forall p \in \{1, \dots, P\}) (\forall m \in \{1, \dots, M\}) \quad \mathbf{V}^{(p,m)} \geq \mathbf{V}_{\min} \right\},$$

avec $\mathbf{V}_{\min} > 0$.

6.2.2 Mise en œuvre de l'algorithme

Nous proposons d'appliquer l'algorithme 5.4, donné dans le chapitre 5, au problème de minimisation (6.4). Pour cela, il nous faut définir, pour tout $(\tilde{\mathbf{T}}, \tilde{\mathbf{V}}) \in \text{dom } g_1 \times \text{dom } g_2$, des fonctions quadratiques $q_1(\cdot | \tilde{\mathbf{T}}, \tilde{\mathbf{V}})$ et $q_2(\cdot | \tilde{\mathbf{T}}, \tilde{\mathbf{V}})$ majorant, respectivement, $h(\cdot, \tilde{\mathbf{V}})$ en $\tilde{\mathbf{T}}$ et $h(\tilde{\mathbf{T}}, \cdot)$ en $\tilde{\mathbf{V}}$. D'après (6.6) et (6.7), les éléments des matrices $\tilde{\mathbf{T}}$ et $\tilde{\mathbf{V}}$ sont strictement positifs. Ainsi, on peut déduire des fonctions quadratiques majorantes proposées dans [Lee et Seung, 2001] dans le contexte de la NMF, les fonctions quadratiques majorantes données par la proposition suivante :

Proposition 6.1. [Lee et Seung, 2001] *Soit $(\tilde{\mathbf{T}}, \tilde{\mathbf{V}}) \in \text{dom } g_1 \times \text{dom } g_2$. Une fonction quadratique majorant $h(\cdot, \tilde{\mathbf{V}})$ en $\tilde{\mathbf{T}}$ est donnée par*

$$(\forall \mathbf{T} \in \mathbb{R}^{Q \times P}) \quad q_1(\mathbf{T} | \tilde{\mathbf{T}}, \tilde{\mathbf{V}}) = h(\tilde{\mathbf{T}}, \tilde{\mathbf{V}}) + \text{Tr} \left((\mathbf{T} - \tilde{\mathbf{T}}) \nabla_1 h(\tilde{\mathbf{T}}, \tilde{\mathbf{V}})^\top \right)$$

$$+ \frac{1}{2} \text{Tr} \left(((\mathbf{T} - \tilde{\mathbf{T}}) \odot \mathbf{A}_1(\tilde{\mathbf{T}}, \tilde{\mathbf{V}})) (\mathbf{T} - \tilde{\mathbf{T}})^\top \right),$$

et une fonction quadratique majorant $h(\tilde{\mathbf{T}}, \cdot)$ en $\tilde{\mathbf{V}}$ est donnée par

$$(\forall \mathbf{V} \in \mathbb{R}^{P \times M}) \quad q_2(\mathbf{V} | \tilde{\mathbf{T}}, \tilde{\mathbf{V}}) = h(\tilde{\mathbf{T}}, \tilde{\mathbf{V}}) + \text{Tr} \left((\mathbf{V} - \tilde{\mathbf{V}}) \nabla_2 h(\tilde{\mathbf{T}}, \tilde{\mathbf{V}})^\top \right) + \frac{1}{2} \text{Tr} \left(((\mathbf{V} - \tilde{\mathbf{V}}) \odot \mathbf{A}_2(\tilde{\mathbf{T}}, \tilde{\mathbf{V}})) (\mathbf{V} - \tilde{\mathbf{V}})^\top \right),$$

avec

$$\mathbf{A}_1(\tilde{\mathbf{T}}, \tilde{\mathbf{V}}) = \left((\Omega^\top \Omega) \tilde{\mathbf{T}} (\tilde{\mathbf{V}} \tilde{\mathbf{V}}^\top) \right) \oslash \tilde{\mathbf{T}}, \quad (6.8)$$

$$\mathbf{A}_2(\tilde{\mathbf{T}}, \tilde{\mathbf{V}}) = \left((\Omega \tilde{\mathbf{T}})^\top \Omega \tilde{\mathbf{T}} \tilde{\mathbf{V}} \right) \oslash \tilde{\mathbf{V}}, \quad (6.9)$$

où $\text{Tr}(\cdot)$ désigne l'opérateur trace, et $\odot$ (resp. $\oslash$) le produit (resp. division) de Hadamard entre deux matrices de même taille.

Dans ce cas particulier, l'algorithme 5.4 se réduit à l'algorithme 6.1.

Algorithme 6.1 Algorithme explicite-implicite à métrique variable alterné par bloc appliqué au problème de NMF

Initialisation : Soient $\mathbf{T}_0 \in \text{dom } \mathbf{g}_1$, $\mathbf{V}_0 \in \text{dom } \mathbf{g}_2$, et, pour tout $k \in \mathbb{N}$, $\gamma_k \in]0, +\infty[$, $\mathbf{A}_1(\mathbf{T}_k, \mathbf{V}_k)$ et $\mathbf{A}_2(\mathbf{T}_{k+1}, \mathbf{V}_k)$ définies par (6.8) et (6.9).

Itérations :

Pour $k = 0, 1, \dots$

$$\left\{ \begin{array}{l} \tilde{\mathbf{T}}_k = \mathbf{T}_k - \gamma_k \nabla_1 h(\mathbf{T}_k, \mathbf{V}_k) \oslash \mathbf{A}_1(\mathbf{T}_k, \mathbf{V}_k), \\ \mathbf{T}_{k+1} = \text{prox}_{\gamma_k^{-1} \mathbf{A}_1(\mathbf{T}_k, \mathbf{V}_k), \mathbf{g}_1} \left(\tilde{\mathbf{T}}_k \right), \\ \tilde{\mathbf{V}}_k = \mathbf{V}_k - \gamma_k \nabla_2 h(\mathbf{T}_k, \mathbf{V}_k) \oslash \mathbf{A}_2(\mathbf{T}_{k+1}, \mathbf{V}_k), \\ \mathbf{V}_{k+1} = \text{prox}_{\gamma_k^{-1} \mathbf{A}_2(\mathbf{T}_{k+1}, \mathbf{V}_k), \mathbf{g}_2} \left(\tilde{\mathbf{V}}_k \right). \end{array} \right.$$

Dans l'algorithme 6.1, pour tout $k \in \mathbb{N}$, les matrices de préconditionnement $\mathbf{A}_1(\mathbf{T}_k, \mathbf{V}_k)$ et $\mathbf{A}_2(\mathbf{T}_{k+1}, \mathbf{V}_k)$ sont appliquées terme à terme. Ainsi, si les matrices $(\mathbf{T}_k)_{k \in \mathbb{N}}$ et $(\mathbf{V}_k)_{k \in \mathbb{N}}$ étaient réorganisées en vecteurs colonnes, les matrices $\mathbf{A}_1(\mathbf{T}_k, \mathbf{V}_k)$ et $\mathbf{A}_2(\mathbf{T}_{k+1}, \mathbf{V}_k)$ seraient des matrices diagonales. La fonction $\mathbf{g}_1$ étant additivement séparable, son opérateur proximal (voir section 2.3.4.2 propriété (vii)) est donné par

$$(\forall k \in \mathbb{N}) \quad \mathbf{T}_{k+1} = (\mathbf{T}_{k+1}^{(q,p)})_{1 \leq q \leq Q, 1 \leq p \leq P} \quad (6.10)$$

avec

$$(\forall (q,p) \in \{1, \dots, P\} \times \{1, \dots, Q\}) \quad \mathbf{T}_{k+1}^{(q,p)} = \Pi_{[\mathbf{T}_{\min}, \mathbf{T}_{\max}]} \left(\text{prox}_{\gamma_k^{-1} \mathbf{A}_1(\mathbf{T}_k, \mathbf{V}_k)^{(q,p), \phi_{\delta, \beta}}} \left(\tilde{\mathbf{T}}_k^{(q,p)} \right) \right), \quad (6.11)$$

où $\Pi_{[\mathbf{T}_{\min}, \mathbf{T}_{\max}]}$ désigne l'opérateur de projection dans l'ensemble $[\mathbf{T}_{\min}, \mathbf{T}_{\max}]$, et l'opérateur proximal de la fonction $\phi_{\delta, \beta}$ est donné dans la proposition 2.12 (chapitre 2). D'autre part, pour tout $k \in \mathbb{N}$, l'opérateur proximal donnant $\mathbf{V}_{k+1}$ dans l'algorithme 6.1 correspond à la projection dans l'ensemble $\mathcal{V}$ relativement à la métrique $\gamma_k^{-1} \mathbf{A}_2(\mathbf{T}_{k+1}, \mathbf{V}_k)$. Cette projection peut être calculée de façon exacte grâce à l'algorithme itératif donné dans [Michelot, 1986].

6.2.3 Résultats de simulation

Afin de simuler des données hyperspectrales réalistes, nous composons le dictionnaire Ω d'une sélection non redondante de $Q = 62$ spectres de $S = 224$ bandes spectrales, allant de 383 nm à 2508 nm, à partir de la bibliothèque disponible sur <http://www.lx.it.pt/~bioucas> (U.S. Geological Survey library [Clark *et al.*, 1993]). La matrice $\bar{\mathbf{U}}$ est composée de $P = 5$ spectres distincts résultant d'une combinaison pondérée de quelques (typiquement 3) spectres purs sélectionnés de façon aléatoire parmi les colonnes de Ω , les poids étant stockés dans la matrice $\bar{\mathbf{T}}$. Chaque

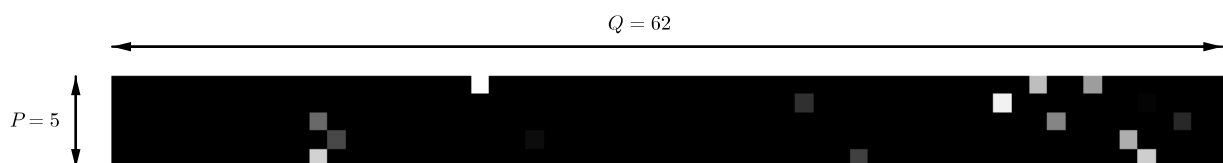


Figure 6.2 – Transposée de la matrice $\bar{\mathbf{T}}$. Chaque ligne $p \in \{1, \dots, P\}$ contient les poids de la combinaison pondérée formant le p -ème spectre de $\bar{\mathbf{U}}$. Les carrés noirs correspondent aux coefficients nuls.

ligne de la matrice d'abondance $\bar{\mathbf{V}}$ est construite comme étant la superposition de gaussiennes 2D composées de $M = 128 \times 128$ pixels, d'espérances et de variances aléatoires, normalisées afin d'assurer la contrainte de somme à un. Finalement, on définit la matrice $\mathbf{E}$ comme une réalisation d'un vecteur aléatoire gaussien centré et dont la variance est choisie de manière à ce que le RSB (c.f. équation (2.20)) soit égal à 20 dB. La figure 6.2 représente la matrice de coefficients $\bar{\mathbf{T}}$, et les abondances $\bar{\mathbf{V}}$ sont représentées dans la figure 6.3.

Les paramètres (δ, β) sont choisis de façon à maximiser le RSB des estimations $\hat{\mathbf{T}}$ et $\hat{\mathbf{V}}$. En particulier, nous avons pu observer que $\beta = 0.1$ donne les meilleurs résultats de reconstruction.

Dans la figure 6.4 nous donnons les spectres exacts $\bar{\mathbf{U}} = \Omega \bar{\mathbf{T}}$ et reconstruits $\hat{\mathbf{U}} = \Omega \hat{\mathbf{T}}$.

Dans cet exemple, la norme résiduelle $r(\hat{\mathbf{T}})$ entre $\hat{\mathbf{T}}$ et $\bar{\mathbf{T}}$ (resp. $r(\hat{\mathbf{V}})$ entre $\hat{\mathbf{V}}$ et $\bar{\mathbf{V}}$) est égal à 1.1×10^{-3} (resp. 7.5×10^{-5}). En comparaison, en appliquant la méthode de NMF proposée dans [Lee et Seung, 2001] au modèle (6.1) on obtient $r(\hat{\mathbf{V}}) = 5 \times 10^{-4}$ (après renormalisation des colonnes de $\hat{\mathbf{V}}$). La figure 6.5 illustre les variations de $(r(\mathbf{T}_k))_k$ en fonction du temps de calcul obtenues avec l'algorithme 6.1 pour $\gamma_k \equiv 0.99$, ou avec l'algorithme PALM proposé dans [Bolte *et al.*, 2013]. Pour obtenir l'algorithme PALM, on remplace, pour tout $k \in \mathbb{N}$, dans l'algorithme 6.1 les matrices $\mathbf{A}_1(\mathbf{T}_k, \mathbf{V}_k)$ et $\mathbf{A}_2(\mathbf{T}_{k+1}, \mathbf{V}_k)$ par

$$\mathbf{A}_1(\mathbf{T}_k, \mathbf{V}_k) = \|\Omega\|^2 \|\mathbf{V}_k\|^2 \mathbf{1}_{Q \times P}, \quad (6.12)$$

$$\mathbf{A}_2(\mathbf{T}_{k+1}, \mathbf{V}_k) = \|\Omega \mathbf{T}_{k+1}\|^2 \mathbf{1}_{P \times M}, \quad (6.13)$$

où, pour tout $(N, M) \in \mathbb{N}^* \times \mathbb{N}^*$, $\mathbf{1}_{N \times M}$ désigne la matrice de dimension $N \times M$ dont tous les coefficients sont égaux à 1. Nous pouvons observer que, dans cet exemple, la métrique variable permet d'accélérer la vitesse de convergence de l'algorithme 6.1 en terme de décroissance de l'erreur résiduelle.

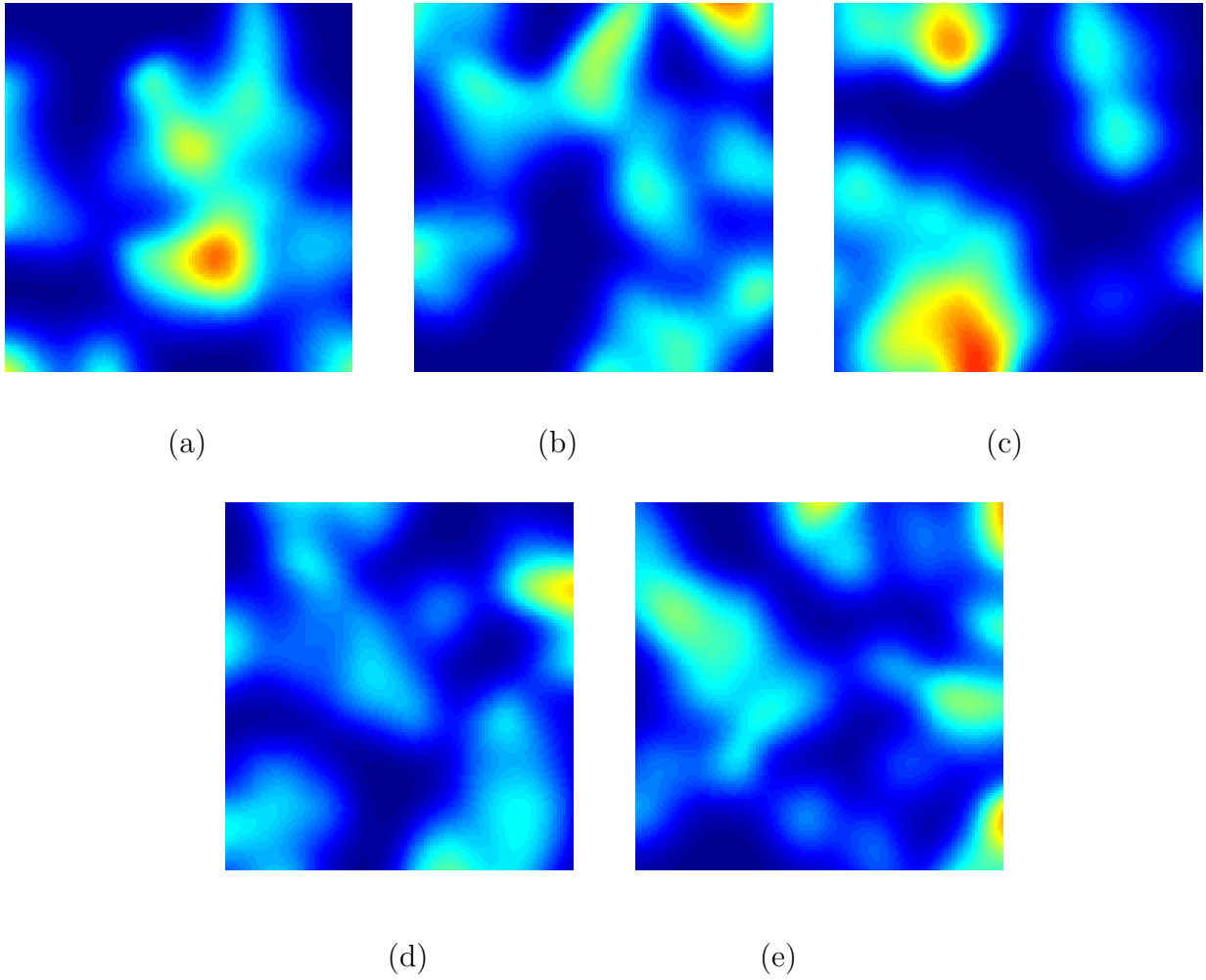


Figure 6.3 – Cartes d’abondance $\bar{\mathbf{V}}$. Image correspondant à la première ligne (a), deuxième ligne (b), troisième ligne (c), quatrième ligne (d) et cinquième ligne (e) de la matrice $\bar{\mathbf{V}}$.

6.3 RECONSTRUCTION DE PHASE

6.3.1 Formulation du problème

6.3.1.1 État de l’art

Bien que la reconstruction de phase ne soit pas un problème inverse nouveau, il reste l’un de ceux les plus ardues à résoudre en traitement d’images. Il consiste à produire une estimation $\hat{\mathbf{y}}$ d’un

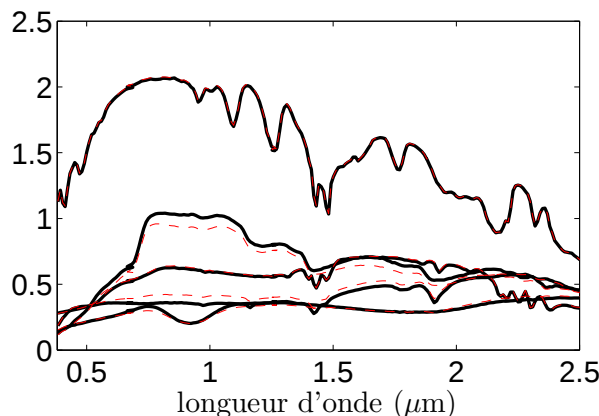


Figure 6.4 – Spectres exacts $\bar{\mathbf{U}} = \Omega \bar{\mathbf{T}}$ (lignes continues noires) et reconstruits $\hat{\mathbf{U}} = \Omega \hat{\mathbf{T}}$ (lignes discontinues rouges).

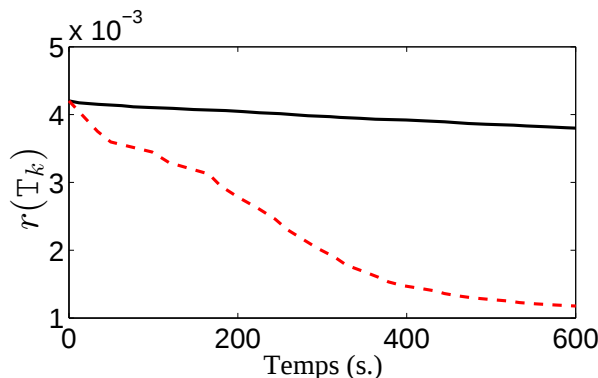


Figure 6.5 – Comparaison de la vitesse de convergence de l'algorithme 6.1 (ligne discontinue rouge) et de l'algorithme PALM (ligne continue noire).

signal original $\bar{\mathbf{y}}$ à partir de l'acquisition de son module dégradé $|\mathbf{H}\bar{\mathbf{y}}|$ (pouvant être bruité) où $\mathbf{H}$ est un opérateur linéaire à valeurs complexes. Un tel problème peut se rencontrer dans beaucoup de domaines d'applications en traitement d'images, en particulier en cristallographie [Harrison, 1993], en imagerie optique [Walther, 1963], en tomographie par contraste de phase [Bauschke et al., 2005], et en imagerie par diffraction cohérente [Shechtman et al., 2013].

Les méthodes les plus populaires permettant d'estimer $\bar{\mathbf{y}}$ sont probablement l'algorithme de Gerchberg-Saxton (GS) [Gerchberg et Saxton, 1972] ainsi que sa version relaxée, l'algorithme de Fienup [Fienup, 1982]. Introduites initialement pour résoudre le cas particulier où $\mathbf{H}$ est une matrice modélisant une transformée de Fourier, ces algorithmes alternent entre une projection sur l'image de $\mathbf{H}$ et une projection sur l'ensemble non convexe des vecteurs ayant un module égal à $|\mathbf{H}\bar{\mathbf{y}}|$. Remarquons que, puisque cet ensemble est non convexe, les algorithmes de GS et Fienup ne bénéficient pas des garanties de convergence de l'algorithme POCS (*Projections Onto Convex Set*) [Bauschke et al., 2002]. Des versions plus générales de ces algorithmes, plus rapides et donnant de meilleurs résultats de reconstruction, sont présentées dans [Bauschke et al., 2003, 2005].

Des relaxations convexes du problème de reconstruction de phase basées sur des formulation de programmation semi-définie (*Semi-Definite-Programming*, SDP) ont été récemment proposées dans [Candès *et al.*, 2013; Waldspurger *et al.*, 2013]. Les algorithmes résultant sont appelés algorithme *PhaseLift* [Candès *et al.*, 2013] et algorithme *PhaseCut* [Waldspurger *et al.*, 2013]. Ces deux méthodes permettent d'obtenir de meilleurs résultats que les algorithmes GS et Fienup dans un certain nombre d'applications [Fogel *et al.*, 2013].

Les méthodes citées ci-dessus ont cependant tendance à ne pas être robustes lorsque les données sont bruitées et/ou à perdre en efficacité lorsque l'on traite des cas sous-déterminés [Shechtman *et al.*, 2014]. De telles difficultés sont souvent dues au caractère mal posé du problème de reconstruction de phase, et une méthode permettant de les contourner (ou les amoindrir) est d'incorporer des informations *a priori* dans le processus de reconstruction. Certains algorithmes basés sur la relaxation SDP [Jaganathan *et al.*, 2012], sur les projections alternées [Mukherjee et Seelamantula, 2012] et sur les méthodes gloutonnes [Shechtman *et al.*, 2014], ont récemment été proposées, pour résoudre le problème de reconstruction de phase sous l'hypothèse que le signal original $\bar{\mathbf{y}}$ a une représentation parcimonieuse dans un certain dictionnaire (qui peut être redondant). Cependant, à notre connaissance, ces approches peuvent rapidement devenir trop coûteuses lorsque la dimension du problème augmente. Cela est particulièrement vrai pour les méthodes SDP et de projections alternées lorsque la pseudo-inverse de $\mathbf{H}$ n'a pas de forme explicite (comme cela est souvent le cas pour des mesures qui ne sont pas faites dans le domaine de Fourier), et pour la méthode gloutonne de [Shechtman *et al.*, 2014] lorsque le degré de parcimonie des données n'est pas suffisant.

Dans cette section, nous allons introduire une nouvelle stratégie de reconstruction de phase reposant sur la minimisation d'un critère pénalisé. Ce critère est composé de (i) une différence de fonctions convexes réalisant une approximation lisse du terme d'attache aux données non convexe de moindres carrés usuel, et (ii) d'un terme de régularisation convexe non nécessairement lisse additivement séparable par bloc. Pour résoudre le problème variationnel résultant, nous utilisons l'algorithme 5.4 donné dans le chapitre 5.

6.3.1.2 Problème d'optimisation

Soient $\bar{\mathbf{y}} \in \mathbb{R}^M$ un signal original inconnu et $\mathbf{H} \in \mathbb{C}^{S \times M}$ une matrice d'observations dont les éléments sont à valeurs complexes. On suppose que les observations $\mathbf{z} \in [0, +\infty]^S$ sont liées au signal d'origine suivant le modèle :

$$\mathbf{z} = |\mathbf{H}\bar{\mathbf{y}}| + \mathbf{w}, \quad (6.14)$$

où $|\cdot|$ désigne l'opérateur de module terme à terme, et $\mathbf{w} \in [0, +\infty]^S$ est une réalisation d'un bruit additif.

Remarque 6.1. Il est important de remarquer que le modèle (6.14) permet aussi de traiter le cas où le signal original $\bar{\mathbf{y}}$ est à valeurs complexes. En effet, dans ce cas, le vecteur d'observations s'écrit

$$\mathbf{z} = \left| \begin{bmatrix} \mathbf{H}_{\mathcal{R}} + i\mathbf{H}_{\mathcal{I}} & -\mathbf{H}_{\mathcal{I}} + i\mathbf{H}_{\mathcal{R}} \end{bmatrix} \begin{bmatrix} \bar{\mathbf{y}}_{\mathcal{R}} \\ \bar{\mathbf{y}}_{\mathcal{I}} \end{bmatrix} \right| + \mathbf{w}, \quad (6.15)$$

où $i^2 = -1$ et $(\cdot)_{\mathcal{R}}$ (resp. $(\cdot)_{\mathcal{I}}$) désigne la partie réelle (resp. imaginaire) de son argument.

L'objectif est alors de produire une estimée $\hat{\mathbf{y}} \in \mathbb{R}^M$ du signal original $\mathbf{y}$ à partir des données observées $\mathbf{z}$.

Soit $\mathbf{W} \in \mathbb{R}^{M \times N}$, $M \leq N$, un opérateur de trame à la synthèse (par exemple, un opérateur d'ondelettes qui peut être redondant) [Mallat, 2009] tel que $\mathbf{y} = \mathbf{W}\mathbf{x}$, avec $\mathbf{x} \in \mathbb{R}^N$ le vecteur contenant les coefficients de trame. En considérant une approche à la synthèse (voir section 2.2.4.2 du chapitre 2) le signal estimé s'écrit $\hat{\mathbf{y}} = \mathbf{W}\hat{\mathbf{x}}$ où $\hat{\mathbf{x}} \in \mathbb{R}^N$ est obtenu en minimisant la somme d'une fonction d'attache aux données $\mathbf{h}$ et d'une fonction de régularisation $\mathbf{g}$, i.e.

$$\hat{\mathbf{x}} \in \underset{\mathbf{x} \in \mathbb{R}^N}{\text{Argmin}} \{f(\mathbf{x}) = \mathbf{h}(\mathbf{x}) + \mathbf{g}(\mathbf{x})\}. \quad (6.16)$$

Dans un contexte de reconstruction de phase, un choix usuel de terme d'attache aux données $\mathbf{h}$ est le critère non lisse et non convexe de moindres carrés [Fienup, 1982; Shechtman et al., 2014] :

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \mathbf{h}(\mathbf{x}) = \sum_{s=1}^S \varphi_s(|[\mathbf{H}\mathbf{W}\mathbf{x}]^{(s)}|), \quad (6.17)$$

où, pour tout $s \in \{1, \dots, S\}$,

$$(\forall u \in [0, +\infty[) \quad \varphi_s(u) = \frac{1}{2}(u - z^{(s)})^2 = \frac{1}{2}(u^2 + (z^{(s)})^2) - uz^{(s)}. \quad (6.18)$$

Nous nous proposons de remplacer cette fonction par une approximation de façon à générer une approximation lisse de la fonction non lisse $\mathbf{h}$. Nous définissons l'approximation de φ_s comme la somme d'une fonction convexe et d'une fonction concave, paramétrée par une constante $\delta > 0$:

$$(\forall u \in [0, +\infty[) \quad \hat{\varphi}_s(u) = \varphi_{s,1}(u) + \varphi_{s,2}(u), \quad (6.19)$$

où, pour tout $s \in \{1, \dots, S\}$,

$$(\forall u \in [0, +\infty[) \quad \varphi_{s,1}(u) = \frac{1}{2}(u^2 + (z^{(s)})^2), \quad (6.20)$$

$$\varphi_{s,2}(u) = -z^{(s)}(u^2 + \delta^2)^{1/2}. \quad (6.21)$$

Pour tout $s \in \{1, \dots, S\}$, les dérivées premières et secondes de $\varphi_{s,1}$ et $\varphi_{s,2}$ en $u \in [0, +\infty[$ sont données par

$$\begin{aligned} \dot{\varphi}_{s,1}(u) &= u, & \ddot{\varphi}_{s,1}(u) &= 1, \\ \dot{\varphi}_{s,2}(u) &= -z^{(s)}u(u^2 + \delta^2)^{-1/2}, & \ddot{\varphi}_{s,2}(u) &= -z^{(s)}\delta^2(u^2 + \delta^2)^{-3/2}. \end{aligned} \quad (6.22)$$

Ainsi, on obtient une approximation de la fonction d'attache aux données (6.17) séparée en deux fonctions $\hat{\mathbf{h}} = \mathbf{h}_1 + \mathbf{h}_2$, définies par

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \mathbf{h}_1(\mathbf{x}) = \sum_{s=1}^S \varphi_{s,1}(|[\mathbf{H}\mathbf{W}\mathbf{x}]^{(s)}|), \quad (6.23)$$

$$\mathbf{h}_2(\mathbf{x}) = \sum_{s=1}^S \varphi_{s,2}(|[\mathbf{H}\mathbf{W}\mathbf{x}]^{(s)}|). \quad (6.24)$$

2. On considèrera la dérivée à droite lorsque $u = 0$.

Nous pouvons remarquer que dans le cas limite où $\delta = 0$, on retrouve la fonction de moindres carrés standard (6.18).

Concernant le terme de régularisation, nous nous focalisons sur le cas où $\mathbf{g}$ est une fonction additivement séparable par bloc. Plus précisément, définissons une partition $(\mathbb{J}_j)_{1 \leq j \leq J}$ de l'ensemble des indices des coefficients de trames $\{1, \dots, N\}$ en $J \geq 2$ sous-ensembles de dimensions (non nulles) $(N_j)_{1 \leq j \leq J}$. On suppose alors que la fonction $\mathbf{g}$ peut s'écrire

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \mathbf{g}(\mathbf{x}) = \sum_{j=1}^J \mathbf{g}_j(\mathbf{x}^{(j)}), \quad (6.25)$$

où, pour tout $j \in \{1, \dots, J\}$, $\mathbf{g}_j: \mathbb{R}^{N_j} \rightarrow]-\infty, +\infty]$ est une fonction semi-algébrique satisfaisant l'hypothèse 5.1(i).

La structure du problème (6.16) nous permet donc d'utiliser l'algorithme 5.4 pour résoudre le problème de reconstruction de phase décrit dans la section 6.3.1.2. Dans la suite nous utiliserons les mêmes notations que dans la section 5.2. Par ailleurs, nous introduisons l'opérateur linéaire $\mathbf{T} = \mathbf{H}\mathbf{W} = (\mathbf{T}^{(s,n)})_{1 \leq s \leq S, 1 \leq n \leq N} \in \mathbb{C}^{S \times N}$.

6.3.2 Mise en œuvre de l'algorithme

Comme nous l'avons vu dans le chapitre 4, la vitesse de convergence de l'algorithme explicite-implicite à métrique variable repose sur le choix des matrices induisant la métrique. Il en va de même pour sa version alternée par bloc donnée par l'algorithme 5.4. L'efficacité des simulations numériques repose sur l'utilisation de majorantes quadratiques donnant, à chaque itération $k \in \mathbb{N}$ de l'algorithme 5.4, de bonnes approximations de la fonction $\mathbf{h}_{j_k}(\cdot, \mathbf{x}_k^{(\bar{j}_k)})^3$ en $\mathbf{x}_k^{(j_k)}$, dont les matrices de courbure $(\mathbf{A}_{j_k}(\mathbf{x}_k))_{k \in \mathbb{N}}$ sont simples à implémenter.

Définissons, pour tout $k \in \mathbb{N}$, les fonctions $\mathbf{h}_{1,j_k}(\cdot, \mathbf{x}_k^{(\bar{j}_k)})$ et $\mathbf{h}_{2,j_k}(\cdot, \mathbf{x}_k^{(\bar{j}_k)})$ respectivement associées à $\mathbf{h}_1$ et $\mathbf{h}_2$. Nous proposons de majorer ces deux fonctions séparément.

D'une part, nous avons remarqué que, pour tout $s \in \{1, \dots, S\}$, $\varphi_{s,2}$ est concave. Ainsi, pour tout $k \in \mathbb{N}$, $\mathbf{h}_{2,j_k}(\cdot, \mathbf{x}_k^{(\bar{j}_k)})$ peut être majorée en $\mathbf{x}_k^{(j_k)}$ par son approximation de premier ordre :

$$(\forall \mathbf{y} \in \mathbb{R}^{N_{j_k}}) \quad \mathbf{q}_{2,j_k}(\mathbf{y}, \mathbf{x}_k) = \mathbf{h}_2(\mathbf{x}_k) + \langle \mathbf{y} - \mathbf{x}_k^{(j_k)}, \nabla_{j_k} \mathbf{h}_2(\mathbf{x}_k) \rangle. \quad (6.26)$$

Il nous reste donc à construire des matrices SDP $(\mathbf{A}_{j_k}(\mathbf{x}_k))_{k \in \mathbb{N}}$ telles que, pour tout $k \in \mathbb{N}$,

$$(\forall \mathbf{y} \in \mathbb{R}^{N_{j_k}}) \quad \mathbf{q}_{1,j_k}(\mathbf{y}, \mathbf{x}_k) = \mathbf{h}_1(\mathbf{x}_k) + \langle \mathbf{y} - \mathbf{x}_k^{(j_k)}, \nabla_{j_k} \mathbf{h}_1(\mathbf{x}_k) \rangle + \frac{1}{2} \|\mathbf{y} - \mathbf{x}_k^{(j_k)}\|_{\mathbf{A}_{j_k}(\mathbf{x}_k)}^2 \quad (6.27)$$

soit une fonction majorant $\mathbf{h}_{1,j_k}(\cdot, \mathbf{x}_k^{(\bar{j}_k)})$. Dans la proposition suivante, nous donnons une matrice $\mathbf{B} \in \mathcal{S}_N^+$ nous permettant de construire des approximations majorantes de $\mathbf{h}_1$ en $\mathbf{x}_k$ pour tout $k \in \mathbb{N}$.

Proposition 6.2. *Soit $\tilde{\mathbf{x}} \in \mathbb{R}^N$. Une fonction quadratique majorant $\mathbf{h}_1$ en $\tilde{\mathbf{x}}$ est donnée par*

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \mathbf{q}_1(\mathbf{x}, \tilde{\mathbf{x}}) = \mathbf{h}_1(\tilde{\mathbf{x}}) + \langle \mathbf{x} - \tilde{\mathbf{x}}, \nabla \mathbf{h}_1(\tilde{\mathbf{x}}) \rangle + \frac{1}{2} \|\mathbf{x} - \tilde{\mathbf{x}}\|_{\mathbf{B}}^2, \quad (6.28)$$

3. Pour tout $k \in \mathbb{N}$, la fonction $\mathbf{h}_{j_k}(\cdot, \mathbf{x}_k^{(\bar{j}_k)})$ est définie par l'équation (5.4)

où $\mathbf{B} = \text{Diag}(\boldsymbol{\Omega}^\top \mathbf{1}_S) + \varepsilon \mathbf{I}_N$, $\mathbf{1}_S$ étant le vecteur unité de $\mathbb{R}^S$, $\varepsilon \geq 0$, et $\boldsymbol{\Omega} = (\Omega^{(s,n)})_{1 \leq s \leq S, 1 \leq n \leq N} \in \mathbb{R}^{S \times N}$ est définie par

$$(\forall s \in \{1, \dots, S\})(\forall n \in \{1, \dots, N\})$$

$$\Omega^{(s,n)} = |\mathbf{T}_{\mathcal{R}}^{(s,n)}| \sum_{n'=1}^N |\mathbf{T}_{\mathcal{R}}^{(s,n')}| + |\mathbf{T}_{\mathcal{I}}^{(s,n)}| \sum_{n'=1}^N |\mathbf{T}_{\mathcal{I}}^{(s,n')}|. \quad (6.29)$$

Démonstration. Soit $\tilde{\mathbf{x}} \in \mathbb{R}^N$. Pour tout $s \in \{1, \dots, S\}$, on a,

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \varphi_{1,s}(|\mathbf{T}^{(s)} \mathbf{x}|) = \varphi_{1,s}(|\mathbf{T}^{(s)} \tilde{\mathbf{x}}|) + \langle \mathbf{x} - \tilde{\mathbf{x}}, [(\mathbf{T}^{(s)})^* \mathbf{T}^{(s)}]_{\mathcal{R}} \tilde{\mathbf{x}} \rangle + \frac{1}{2} |\mathbf{T}^{(s)}(\mathbf{x} - \tilde{\mathbf{x}})|^2,$$

où $\mathbf{T}^{(s)}$ désigne la ligne s de la matrice $\mathbf{T}$, et $(\cdot)^*$ est l'opérateur de conjugué de matrice. En sommant sur $s \in \{1, \dots, S\}$, on obtient alors

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad h_1(\mathbf{x}) = h_1(\tilde{\mathbf{x}}) + \langle \mathbf{x} - \tilde{\mathbf{x}}, \nabla h_1(\tilde{\mathbf{x}}) \rangle + \frac{1}{2} ||| \mathbf{T}(\mathbf{x} - \tilde{\mathbf{x}}) |||^2, \quad (6.30)$$

où $||| \cdot |||$ est la norme hermitienne de $\mathbb{C}^S$.

Soient $(V_{\mathcal{R}}^{(s,n)})_{1 \leq s \leq S, 1 \leq n \leq N} \in [0, +\infty]^{S \times N}$ et $(V_{\mathcal{I}}^{(s,n)})_{1 \leq s \leq S, 1 \leq n \leq N} \in [0, +\infty]^{S \times N}$ définis tels que, pour tout $s \in \{1, \dots, S\}$,

$$\sum_{n \in \mathcal{S}_{\mathcal{R}}^{(s)}} V_{\mathcal{R}}^{(s,n)} = 1, \quad \text{et} \quad \sum_{n \in \mathcal{S}_{\mathcal{I}}^{(s)}} V_{\mathcal{I}}^{(s,n)} = 1$$

où

$$\begin{aligned} \mathcal{S}_{\mathcal{R}}^{(s)} &= \{n \in \{1, \dots, N\} \mid V_{\mathcal{R}}^{(s,n)} \neq 0\} = \{n \in \{1, \dots, N\} \mid T_{\mathcal{R}}^{(s,n)} \neq 0\}, \\ \mathcal{S}_{\mathcal{I}}^{(s)} &= \{n \in \{1, \dots, N\} \mid V_{\mathcal{I}}^{(s,n)} \neq 0\} = \{n \in \{1, \dots, N\} \mid T_{\mathcal{I}}^{(s,n)} \neq 0\}. \end{aligned}$$

D'après l'inégalité de Jensen, pour tout $s \in \{1, \dots, S\}$, on a

$$\begin{aligned} & \left| \sum_{n=1}^N \mathbf{T}^{(s,n)}(\mathbf{x}^{(n)} - \tilde{\mathbf{x}}^{(n)}) \right|^2 \\ &= \left(\sum_{n=1}^N \mathbf{T}_{\mathcal{R}}^{(s,n)}(\mathbf{x}^{(n)} - \tilde{\mathbf{x}}^{(n)}) \right)^2 + \left(\sum_{n=1}^N \mathbf{T}_{\mathcal{I}}^{(s,n)}(\mathbf{x}^{(n)} - \tilde{\mathbf{x}}^{(n)}) \right)^2 \\ &= \left(\sum_{n \in \mathcal{S}_{\mathcal{R}}^{(s)}} V_{\mathcal{R}}^{(s,n)} \left(\frac{\mathbf{T}_{\mathcal{R}}^{(s,n)}}{V_{\mathcal{R}}^{(s,n)}}(\mathbf{x}^{(n)} - \tilde{\mathbf{x}}^{(n)}) \right) \right)^2 + \left(\sum_{n \in \mathcal{S}_{\mathcal{I}}^{(s)}} V_{\mathcal{I}}^{(s,n)} \left(\frac{\mathbf{T}_{\mathcal{I}}^{(s,n)}}{V_{\mathcal{I}}^{(s,n)}}(\mathbf{x}^{(n)} - \tilde{\mathbf{x}}^{(n)}) \right) \right)^2 \\ &\leq \sum_{n \in \mathcal{S}_{\mathcal{R}}^{(s)}} \frac{(\mathbf{T}_{\mathcal{R}}^{(s,n)})^2}{V_{\mathcal{R}}^{(s,n)}} (\mathbf{x}^{(n)} - \tilde{\mathbf{x}}^{(n)})^2 + \sum_{n \in \mathcal{S}_{\mathcal{I}}^{(s)}} \frac{(\mathbf{T}_{\mathcal{I}}^{(s,n)})^2}{V_{\mathcal{I}}^{(s,n)}} (\mathbf{x}^{(n)} - \tilde{\mathbf{x}}^{(n)})^2. \end{aligned} \quad (6.31)$$

En particulier, nous pouvons choisir

$$(\forall (s, n) \in \{1, \dots, S\} \times \{1, \dots, N\}) \quad \mathbf{V}_{\mathcal{R}}^{(s,n)} = \begin{cases} 0 & \text{si } \mathbf{T}_{\mathcal{R}}^{(s,n)} = 0, \\ \frac{|\mathbf{T}_{\mathcal{R}}^{(s,n)}|}{\sum_{n'=1}^N |\mathbf{T}_{\mathcal{R}}^{(s,n')}|} & \text{sinon,} \end{cases}$$

$$\mathbf{V}_{\mathcal{I}}^{(s,n)} = \begin{cases} 0 & \text{si } \mathbf{T}_{\mathcal{I}}^{(s,n)} = 0, \\ \frac{|\mathbf{T}_{\mathcal{I}}^{(s,n)}|}{\sum_{n'=1}^N |\mathbf{T}_{\mathcal{I}}^{(s,n')}|} & \text{sinon.} \end{cases}$$

On peut alors d  duire de (6.31) que, pour tout $s \in \{1, \dots, S\}$,

$$\left| \sum_{n=1}^N \mathbf{T}^{(s,n)} (\mathbf{x}^{(n)} - \tilde{\mathbf{x}}^{(n)}) \right|^2 \leq \sum_{n=1}^N \left(|\mathbf{T}_{\mathcal{R}}^{(s,n)}| \sum_{n'=1}^N |\mathbf{T}_{\mathcal{R}}^{(s,n')}| \right) (\mathbf{x}^{(n)} - \tilde{\mathbf{x}}^{(n)})^2$$

$$+ \sum_{n=1}^N \left(|\mathbf{T}_{\mathcal{I}}^{(s,n)}| \sum_{n'=1}^N |\mathbf{T}_{\mathcal{I}}^{(s,n')}| \right) (\mathbf{x}^{(n)} - \tilde{\mathbf{x}}^{(n)})^2.$$

On obtient ainsi

$$|||\mathbf{T}(\mathbf{x} - \tilde{\mathbf{x}})|||^2 \leq \langle \mathbf{x} - \tilde{\mathbf{x}}, \text{Diag}(\mathbf{\Omega}^\top \mathbf{1}_S) (\mathbf{x} - \tilde{\mathbf{x}}) \rangle, \quad (6.32)$$

o   $\mathbf{\Omega}$ est d  finie par (6.29). Pour finir, en combinant les   quations (6.30) et (6.32), on en d  duit que la fonction $\mathbf{h}_1$ peut   tre born  e par la majorante (6.28). $\square$

En utilisant le lemme 6.2 et la remarque 5.3(ii), nous pouvons maintenant construire, pour tout $k \in \mathbb{N}$, une majorante quadratique de $\mathbf{h}_{1,j_k}(\cdot, \mathbf{x}_k^{j_k})$ en $\mathbf{x}_k$ de la forme de (6.27) avec

$$(\forall k \in \mathbb{N}) \quad \mathbf{A}_{j_k}(\mathbf{x}_k) = \text{Diag}(\mathbf{\Omega}_{j_k}^\top \mathbf{1}_S) + \varepsilon \mathbf{I}_{N_{j_k}}, \quad (6.33)$$

o   $\mathbf{\Omega}_{j_k} \in \mathbb{R}^{S \times N_{j_k}}$ est la matrice obtenue en extrayant de la matrice $\mathbf{\Omega}$ d  finie par (6.29) les colonnes dont les indices sont dans $\mathbb{J}_{j_k}$. Remarquons que l'hypoth  se 5.2(ii) est satisfaite pour les matrices (6.33) avec

$$\begin{cases} \underline{\nu} = \varepsilon + \min_{n \in \mathbb{J}_{j_k}} \sum_{s=1}^S \Omega^{(s,n)}, \\ \overline{\nu} = \varepsilon + \max_{n \in \mathbb{J}_{j_k}} \sum_{s=1}^S \Omega^{(s,n)}. \end{cases} \quad (6.34)$$

Si aucune colonne de $\mathbf{T}$ n'est nulle, alors on peut choisir $\varepsilon = 0$ dans (6.34). Sinon, il faut prendre $\varepsilon > 0$.

6.3.3 R  sultats de simulation

Nous allons maintenant pr  senter nos r  sultats de simulation. Consid  rons une image originale $\bar{\mathbf{y}} \in \mathbb{C}^M$    valeurs complexes (voir figure 6.6) de dimension $M = 128 \times 128$. Nous observons une version d  grad  e $\mathbf{z} \in \mathbb{R}^S$ de cette image, suivant le mod  le (6.15). La matrice de d  gradation $\mathbf{H} \in \mathbb{C}^{S \times M}$ correspond    la composition de deux op  rateurs, le premier   tant une matrice de projection

modélisant $S = 23400$ projections de Radon à partir de 128 lignes d'acquisition parallèles et 180 angles équidistants sur $[0, \pi[$, et le second correspond à un opérateur de flou à valeurs complexes. Pour modéliser l'opérateur de flou, nous avons utilisé un noyau de convolution 1D de longueur 3, agissant séparément sur chaque angle d'acquisition. Ce modèle s'inspire de la version linéarisée d'un problème de tomographie par contraste de phase proposé dans [Davidoiu *et al.*, 2012; Guigay *et al.*, 2007] dans le cas 3D. D'autre part, on suppose que le vecteur $\mathbf{w}$ est une réalisation d'un vecteur aléatoire suivant une loi normale centrée et de variance égale à 0.025 (tronquée de façon à garantir la positivité des observations).

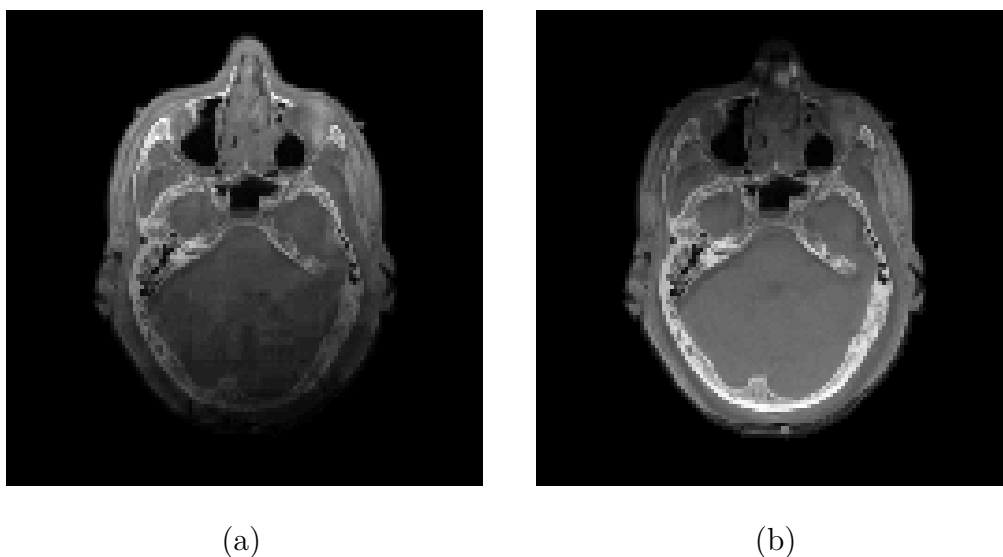


Figure 6.6 – Parties réelle (a) et imaginaire (b) de l'image originale $\bar{\mathbf{y}}$

L'opérateur de trame $\mathbf{W} \in \mathbb{R}^{M \times N}$ est choisi de façon à ce que $N = 8M$ et que chaque vecteur $\mathbf{x} \in \mathbb{R}^N$ soit la concaténation de deux vecteurs $(\mathbf{x}_{\mathcal{R}}, \mathbf{x}_{\mathcal{I}}) \in \mathbb{R}^{4M} \times \mathbb{R}^{4M}$, où $\mathbf{x}_{\mathcal{R}} = (\mathbf{x}_{\mathcal{R}}^{(p)})_{1 \leq p \leq 4M}$ (resp. $\mathbf{x}_{\mathcal{I}} = (\mathbf{x}_{\mathcal{I}}^{(p)})_{1 \leq p \leq 4M}$) correspond à une décomposition de Haar redondante de $\mathbf{y}_{\mathcal{R}}$ (resp. $\mathbf{y}_{\mathcal{I}}$) sur un niveau de résolution. Un exemple de décomposition de la partie réelle de $\bar{\mathbf{y}}$ est donné dans la figure 6.7(a).

La fonction de régularisation $\mathbf{g}$ est définie par

$$(\forall \mathbf{x} \in \mathbb{R}^{8M}) \quad \mathbf{g}(\mathbf{x}) = \sum_{p=1}^{4M} \mathbf{g}_p(\mathbf{x}_{\mathcal{R}}^{(p)}, \mathbf{x}_{\mathcal{I}}^{(p)}), \quad (6.35)$$

avec, pour tout $p \in \{1, \dots, 4M\}$,

$$(\forall \mathbf{u} \in \mathbb{R}^2) \quad \mathbf{g}_p(\mathbf{u}) = \begin{cases} \vartheta_p \|\mathbf{u} - \boldsymbol{\omega}_p\|_2^{\kappa_p} & \text{si } p \notin \mathbb{E}, \\ 0 & \text{si } p \in \mathbb{E} \text{ et } \mathbf{u} = \mathbf{0}, \\ +\infty & \text{sinon,} \end{cases} \quad (6.36)$$

où $\mathbb{E}$ correspond à l'arrière plan de l'objet (donné dans la figure 6.7(b) pour les coefficients correspondants à la partie réelle de l'image). Afin de promouvoir la parcimonie des coefficients de

détails de $\hat{\mathbf{x}}$ nous choisissons $\kappa_p = 1$ et $\vartheta_p = \vartheta^d \in]0, +\infty[$ lorsque $p \in \{M+1, \dots, 4M\}$. Pour les coefficients d'approximation, nous choisissons $\kappa_p = 2$ et $\vartheta_p = \vartheta^a \in]0, +\infty[$ lorsque $p \in \{1, \dots, M\}$. Le paramètre $\omega_p \in \mathbb{R}^2$ correspond à une valeur moyenne des coefficients de trame. Dans nos simulations nous avons pris $\omega_p = (0.4, 0.6)$ pour les coefficients d'approximation, et $\omega_p = \mathbf{0}$ sinon. Ainsi, la fonction de régularisation donnée par (6.35)-(6.36) satisfait l'hypothèse 5.1(i) et est semi-algébrique. De plus, son opérateur proximal a une forme explicite (voir section 2.3.4).

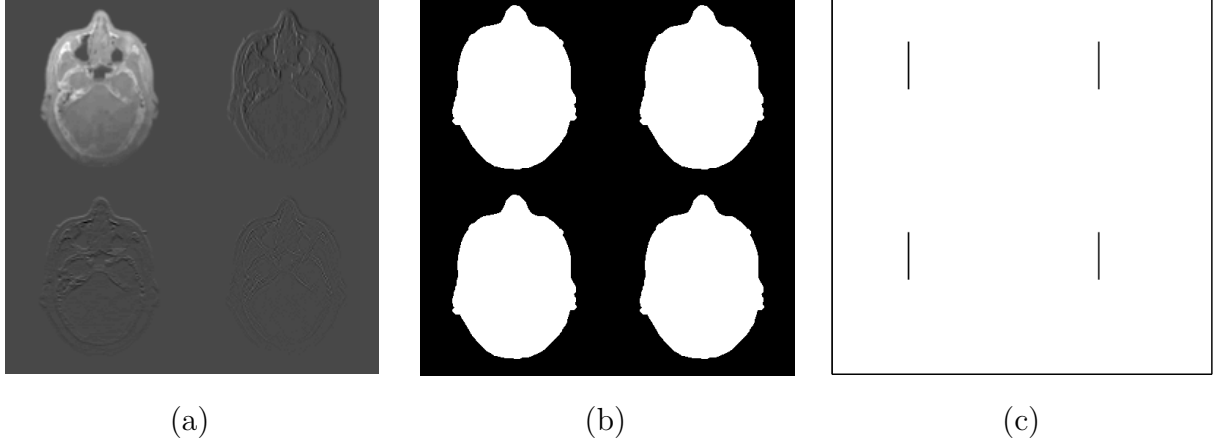


Figure 6.7 – Seuls les coefficients de trame correspondants à la partie réelle de l'image $\bar{\mathbf{y}}_{\mathcal{R}}$ sont représentés dans les trois figures. (a) Exemple de décomposition de trame de $\bar{\mathbf{y}}_{\mathcal{R}}$. (b) Masque $\mathbb{E}$ permettant de forcer les coefficients de l'arrière plan de l'objet à être égaux à 0. (c) Indices d'un bloc $\mathbf{x}^{(j)}$ pour $Q = 32$.

Dans nos simulations, les paramètres ϑ^a , ϑ^d et δ sont ajustés de façon à maximiser le RSB entre l'image originale $\bar{\mathbf{y}}$ et l'image reconstruite $\hat{\mathbf{y}}$. Dans l'algorithme 5.4, pour tout j , $\mathbf{x}^{(j)}$ correspond à un vecteur, de dimension fixe $8Q$, faisant intervenir 8 blocs extraits à la fois des approximations et des détails des parties réelles et imaginaires de l'image. Une illustration pour la partie réelle est donnée dans la figure 6.7(c). À chaque itération $k \in \mathbb{N}$, j_k est choisi aléatoirement de façon à ce que chaque bloc soit mis à jour au moins une fois toutes les J itérations. En pratique, nous observons que le temps de reconstruction varie en fonction du paramètre Q . La figure 6.8(a) donne le temps de reconstruction nécessaire à l'algorithme 5.4 pour atteindre le critère d'arrêt $\|\mathbf{x}_k - \hat{\mathbf{x}}\|/\|\hat{\mathbf{x}}\| \leq 0.025$ en fonction du paramètre Q , où la solution optimale $\hat{\mathbf{x}}$ a été calculée préalablement pour un grand nombre d'itérations. Dans cet exemple, le meilleur choix en terme de vitesse de convergence semble être $Q = 2$.

La figure 6.8(b) illustre la décroissance de $(\|\mathbf{x}_k - \hat{\mathbf{x}}\|/\|\hat{\mathbf{x}}\|)_k$ (en échelle logarithmique) obtenue en utilisant l'algorithme 5.4 et l'algorithme PALM de [Bolte et al., 2013] avec $\gamma_k \equiv 1.9$ et $Q = 2$, où $\hat{\mathbf{x}}$ correspond à la solution optimale obtenue pour chaque algorithme. Nous pouvons observer que, dans cet exemple, le préconditionnement permet d'accélérer la vitesse de convergence des itérées.

Dans la figure 6.9(a)-(b) nous donnons les parties réelle et imaginaire de l'image reconstruite en utilisant la méthode proposée dans cette section. Nous donnons aussi dans la figure 6.9(c)-(d) les parties réelle et imaginaire de l'image reconstruite obtenue en utilisant une version améliorée

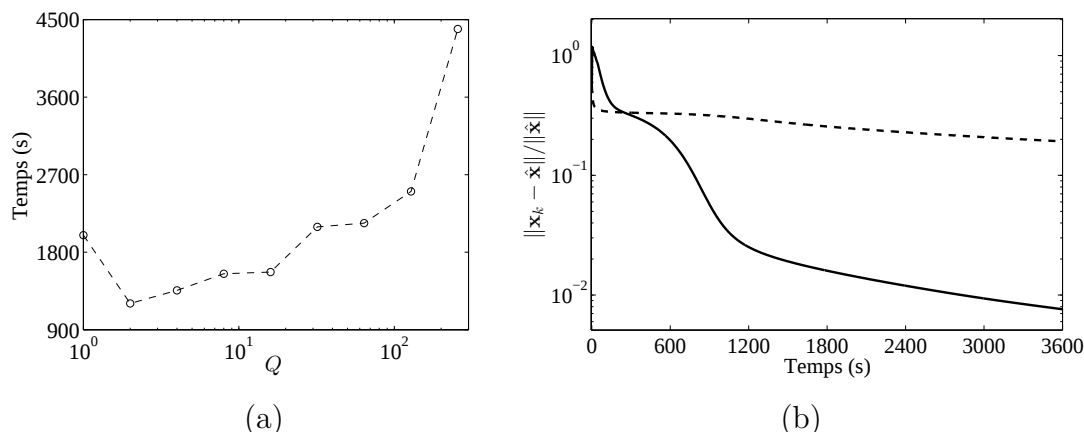


Figure 6.8 – (a) Temps de reconstruction nécessaire pour satisfaire le critère d'arrêt $\|\mathbf{x}_k - \hat{\mathbf{x}}\| / \|\hat{\mathbf{x}}\| \leq 0.025$ avec l'algorithme 5.4 en fonction de Q . (b) Profil de convergence de l'algorithme 5.4 (ligne continue) et de l'algorithme PALM de [Bolte et al., 2013] (ligne discontinue), obtenu en utilisant Matlab 7 sur une machine Intel(R) Core(TM) i7-3520M @ 2.9GHz.

la méthode de projections alternée proposée par [Mukherjee et Seelamantula, 2012], où la non-inversibilité de la matrice $\mathbf{H}$ a été prise en compte. Remarquons de plus qu'il faut compter un temps 10 fois plus long avec la méthode de [Mukherjee et Seelamantula, 2012] que le temps mis avec l'algorithme 5.4, pour obtenir une valeur de RSB stable. Notons qu'en raison de la grande dimension du problème traité, les approches SDP de [Candès et al., 2013; Jaganathan et al., 2012; Waldspurger et al., 2013] ainsi que la méthode gloutonne de [Shechtman et al., 2014] ne peuvent pas être mises en œuvre. D'autre part, l'algorithme standard de GS ne donne pas des résultats de qualité acceptable.

6.4 DÉCONVOLUTION AVEUGLE DE SIGNAUX SISMQUES

6.4.1 Formulation du problème

6.4.1.1 État de l'art

Dans cette section, nous cherchons à estimer un signal sismique inconnu $\bar{\mathbf{x}} \in \mathbb{R}^N$ à partir d'observations $\mathbf{z} \in \mathbb{R}^N$ obtenues par le modèle :

$$\mathbf{z} = \bar{\mathbf{k}} * \bar{\mathbf{x}} + \mathbf{w}, \quad (6.37)$$

où $\bar{\mathbf{k}} \in \mathbb{R}^S$ représente une réponse impulsionnelle (par exemple, la réponse linéaire d'un capteur, ou un noyau de convolution modélisant un flou), et $\mathbf{w} \in \mathbb{R}^N$ est une réalisation d'un vecteur aléatoire modélisant un bruit additif. Un exemple de signal sismique $\bar{\mathbf{x}} \in \mathbb{R}^N$ ainsi que de son observation faite au travers du modèle (6.37) sont donnés dans la figure 6.10.

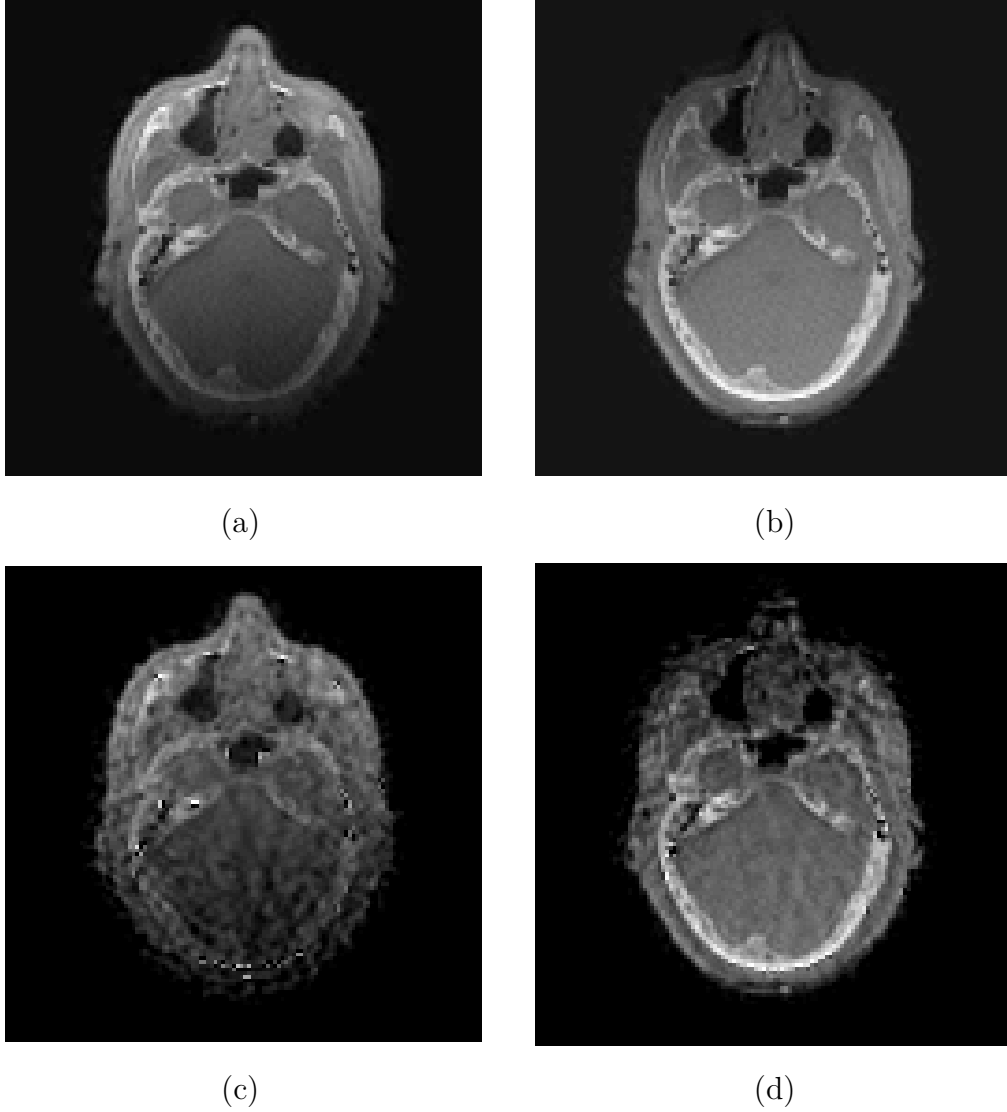


Figure 6.9 – Parties réelle (a) (resp. (c)) et imaginaire (b) (resp. (d)) de l'image reconstruite $\hat{\mathbf{y}}$ en utilisant l'algorithme 5.4, $RSB = 21.27$ dB (resp. l'algorithme de projections alternées régularisé de [Mukherjee et Seelamantula, 2012], $RSB = 14.45$ dB).

L'utilisation d'approches standards, telles que le filtre de Wiener [Wiener, 1949] ou des extensions de type SURELET [Pesquet *et al.*, 2009], permettent de trouver une estimée $\hat{\mathbf{x}} \in \mathbb{R}^N$ du signal original $\mathbf{x}$ en minimisant un critère basé sur le carré de la norme euclidienne $\|\cdot\|^2$. Cependant, d'une part l'utilisation seule d'une fonction d'attache aux données de moindres carrés ne permet pas de débruiter correctement le signal, et d'autre part, en ajoutant une régularisation $\|\cdot\|^2$ on obtient souvent des estimations trop lisses du signal de départ.

Dans cette section, nous nous intéresserons au cas de la déconvolution aveugle, c'est à dire lorsque le noyau $\bar{\mathbf{k}}$ est lui aussi inconnu, et doit être estimé conjointement au signal $\mathbf{x}$. Ce type

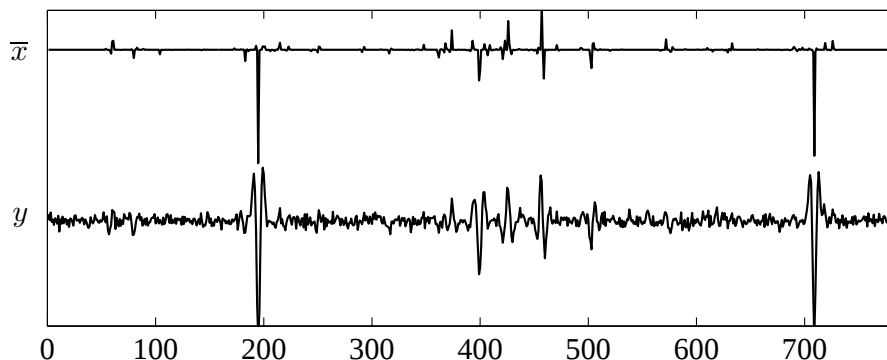


Figure 6.10 – (Haut) Signal sismique inconnu $\bar{\mathbf{x}}$. (Bas) Observations $\mathbf{z}$.

de problème peut se rencontrer dans diverses applications telles que les communications (égalisation ou estimation de canal) [Haykin, 1994], le contrôle non destructif [Nandi *et al.*, 1997], la géophysique [Kaaresen et Taxt, 1998; Pham *et al.*, 2014; Takahata *et al.*, 2012], le traitement des images [Ahmed *et al.*, 2014; Kato *et al.*, 1999; Kundur et Hatzinakos, 1996a,b], l'imagerie médicale et la télédétection Campisi et Egiazarian [2007]. Le problème de déconvolution aveugle étant un problème très sous-déterminé, il requiert l'utilisation d'hypothèses supplémentaires sur le signal et le noyau à estimer. Une approche usuelle est de définir les estimées $(\hat{\mathbf{x}}, \hat{\mathbf{k}}) \in \mathbb{R}^N \times \mathbb{R}^S$ de $(\bar{\mathbf{x}}, \bar{\mathbf{k}})$ comme les minimiseurs de la somme d'une fonction d'attache aux données et des termes de régularisation sur le signal et sur le noyau de convolution (voir section 2.2.1). Le problème de déconvolution aveugle étant sujet à des ambiguïtés d'échelles, et les fonctions de régularisation servant à la fois à prendre en compte des connaissances sur les objets originaux et assurant la stabilité de la solution, il peut être intéressant de les choisir invariantes par échelle [Comon, 1996; Moreau et Pesquet, 1997].

La fonction ℓ_1/ℓ_2 , introduite dans [Gray, 1978] pour traiter des problèmes de déconvolution en géophysique, possède de telles propriétés. Elle a été utilisée plus récemment comme mesure de parcimonie dans [Barak *et al.*, 2014; Hoyer, 2004; Hurley et Rickard, 2004; Zibulevsky et Pearlmutter, 2001] pour résoudre des problèmes de NMF (*Nonnegative Matrix Factorization*) [Mørup *et al.*, 2008], des problèmes de traitement d'image en utilisant une décomposition en ondelettes [Ji *et al.*, 2012], ou pour la reconstruction de signaux parcimonieux [Demanet et Hand, 2014]. Il est de plus souligné dans [Benichoux *et al.*, 2013] que cette fonction de régularisation peut être efficace pour les problèmes de déconvolution aveugle de signaux parcimonieux.

Récemment, [Krishnan *et al.*, 2011] a proposé un algorithme de minimisation alternée permettant d'utiliser la fonction de régularisation ℓ_1/ℓ_2 . Les auteurs proposent de transformer le problème régularisé par la fonction non convexe ℓ_1/ℓ_2 en une famille de problèmes régularisés par la fonction convexe ℓ_1 . Pour cela, à chaque itération, le résultat obtenu est pondéré en le divisant par la norme ℓ_2 de l'itérée précédente. Le problème régularisé par ℓ_1 peut être résolu en utilisant l'algorithme explicite-implicite (aussi appelé ISTA (*Iterative Shrinkage-Thresholding Algorithm*)) lorsque la fonction non lisse est la norme ℓ_1 [Beck et Teboulle, 2009]. Bien que la convergence de cette méthode n'ait pas été établie, elle semble très efficace en pratique. Une autre méthode, basée sur un algorithme de gradient projeté avec changement d'échelle, a été proposée dans [Esser *et al.*, 2013] pour résoudre une approximation lisse de la fonction ℓ_1/ℓ_2 . Cependant, cette approche ne permet de résoudre que des problèmes traitant des signaux parcimonieux à

valeurs positives.

L'objectif de cette section est de proposer un algorithme efficace permettant de prendre en compte une approximation lisse de la pénalisation ℓ_1/ℓ_2 pour des problèmes traitant des données à valeurs réelles signées.

6.4.1.2 Problème d'optimisation

On définit les estimées $(\hat{\mathbf{x}}, \hat{\mathbf{k}}) \in \mathbb{R}^N \times \mathbb{R}^S$ de $(\bar{\mathbf{x}}, \bar{\mathbf{k}})$ comme les minimiseurs de la fonction

$$(\forall (\mathbf{x}, \mathbf{k}) \in \mathbb{R}^N \times \mathbb{R}^S) \quad f(\mathbf{x}, \mathbf{k}) = h_1(\mathbf{x}, \mathbf{k}) + h_2(\mathbf{x}) + g_1(\mathbf{x}) + g_2(\mathbf{k}), \quad (6.38)$$

où $h_1: \mathbb{R}^N \times \mathbb{R}^S \rightarrow]-\infty, +\infty]$ correspond au critère de moindres carrés associé au modèle (6.37) défini par

$$(\forall (\mathbf{x}, \mathbf{k}) \in \mathbb{R}^N \times \mathbb{R}^S) \quad h_1(\mathbf{x}, \mathbf{k}) = \frac{1}{2} \|\mathbf{k} * \mathbf{x} - \mathbf{z}\|^2, \quad (6.39)$$

$g_1: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ et $g_2: \mathbb{R}^S \rightarrow]-\infty, +\infty]$ sont respectivement des termes de régularisation sur $\mathbf{x}$ et $\mathbf{k}$ non nécessairement lisses, et $h_2: \mathbb{R}^N \rightarrow]-\infty, +\infty]$ est une fonction de régularisation différentiable de gradient Lipschitz sur $\mathbf{x}$. Nous supposons de plus que g_1 et g_2 sont des fonctions semi-algébriques satisfaisant l'hypothèse 5.1(i) donnée dans le chapitre 5, et nous notons, pour tout $(\mathbf{x}, \mathbf{k}) \in \mathbb{R}^N \times \mathbb{R}^S$, $h(\mathbf{x}, \mathbf{k}) = h_1(\mathbf{x}, \mathbf{k}) + h_2(\mathbf{x})$ la partie différentiable du critère (6.38).

Concernant la fonction de régularisation h_2 , cette dernière doit à la fois promouvoir la parcimonie de $\bar{\mathbf{x}}$, et être simple à implémenter numériquement. Nous avons vu que la fonction de régularisation non convexe ℓ_1/ℓ_2 [Slavakis *et al.*, 2013], définie comme le quotient des normes

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \ell_1(\mathbf{x}) = \sum_{n=1}^N |x^{(n)}| \quad \text{et} \quad \ell_2(\mathbf{x}) = \left(\sum_{n=1}^N (x^{(n)})^2 \right)^{1/2},$$

permet de traiter efficacement les problèmes de déconvolution aveugle. Cependant, n'étant ni différentiable, ni convexe, elle se révèle compliquée à optimiser.

Nous proposons donc de remplacer cette fonction non lisse par une approximation lisse de celle-ci qui sera plus facilement maniable. Pour cela, rappelons que des approximations lisses de ℓ_1 et ℓ_2 , notées $\ell_{1,\alpha}$ (parfois appelée pénalisation ℓ_1 - ℓ_2 -hybride ou hyperbolique) et $\ell_{2,\eta}$, sont définies par

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \ell_{1,\alpha}(\mathbf{x}) = \sum_{n=1}^N \left(\sqrt{(x^{(n)})^2 + \alpha^2} - \alpha \right) \quad \text{et} \quad \ell_{2,\eta}(\mathbf{x}) = \sqrt{\sum_{n=1}^N (x^{(n)})^2 + \eta^2}, \quad (6.40)$$

où $(\alpha, \eta) \in]0, +\infty[^2$. Remarquons que dans le cas limite où $\alpha = \eta = 0$, on retrouve les normes usuelles ℓ_1 et ℓ_2 . Nous définissons alors notre approximation lisse de la fonction ℓ_1/ℓ_2 de la façon suivante :

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad h_2(\mathbf{x}) = \lambda \log \left(\frac{\ell_{1,\alpha}(\mathbf{x}) + \beta}{\ell_{2,\eta}(\mathbf{x})} \right), \quad (6.41)$$

avec $(\lambda, \beta, \alpha, \eta) \in]0, +\infty[^4$.

D'une part, remarquons que la fonction log permet non seulement de rendre la fonction plus facile à manier, mais aussi de mieux faire ressortir la parcimonie. Un exemple en est donné dans la figure 6.11 où h_2 est comparée à d'autres fonctions permettant de promouvoir la parcimonie.

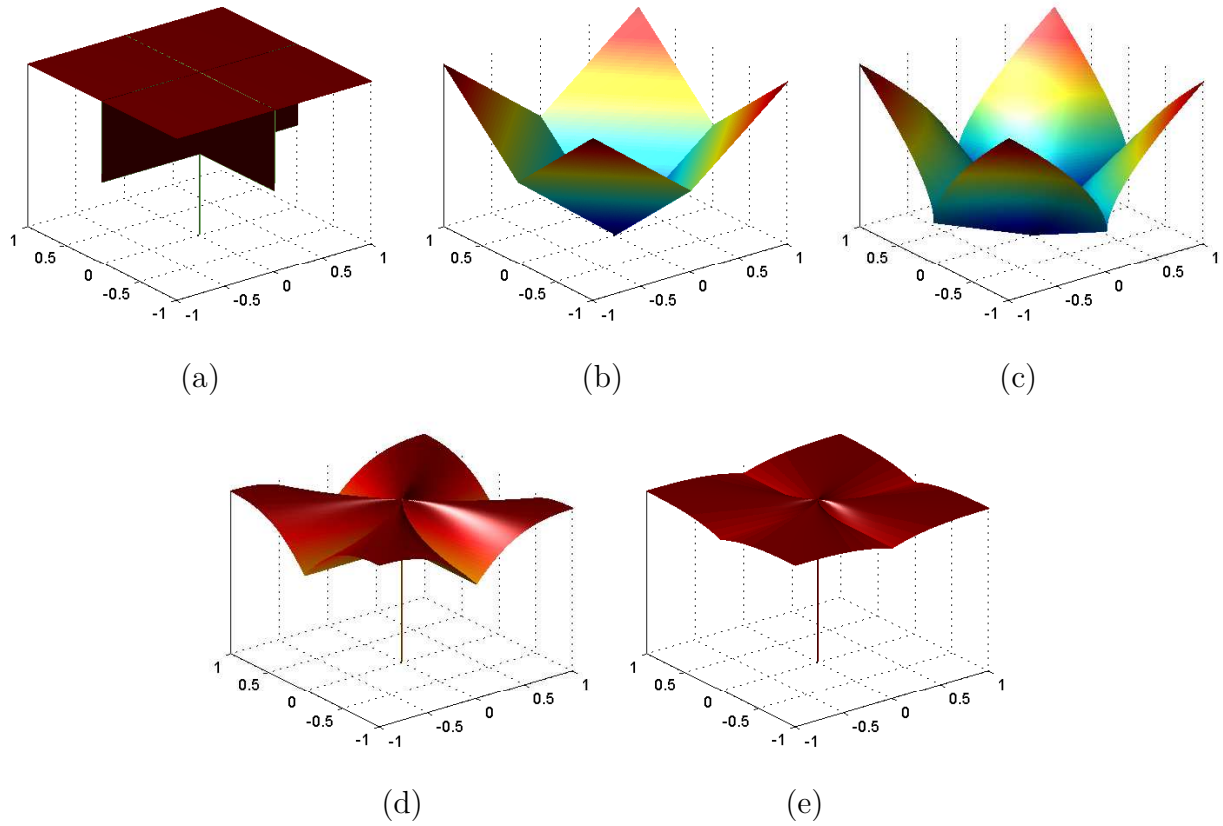


Figure 6.11 – (a) Fonction ℓ_0 . (b) Fonction ℓ_1 . (c) Fonction $\ell_{1/2}$. (d) Fonction $\ell_{1,\alpha}/\ell_{2,\eta}$. (e) Fonction $h_2 = \log \left((\ell_{1,\alpha} + \beta)/\ell_{2,\eta} \right)$.

D'autre part, notons que f correspond à la fonction Lagrangienne associée à la minimisation de $h_1 + g_1 + g_2$ sous la contrainte

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \log \left(\frac{\ell_{1,\alpha}(\mathbf{x}) + \beta}{\ell_{2,\eta}(\mathbf{x})} \right) \leq \log(\vartheta), \quad (6.42)$$

pour un certain $\vartheta > 0$. D'après la monotonie de la fonction $\log$, (6.42) est équivalente à

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \frac{\ell_{1,\alpha}(\mathbf{x}) + \beta}{\ell_{2,\eta}(\mathbf{x})} \leq \vartheta, \quad (6.43)$$

qui, d'après (6.41), peut être interprétée comme une formulation sous contrainte d'une approximation lisse de la fonction ℓ_1/ℓ_2 , dès que β est assez petit.

Pour finir, soulignons que la fonction h_2 est de gradient Lipschitz sur tout sous-ensemble compact de $\mathbb{R}^N$. Ainsi, la fonction h satisfait l'hypothèse 5.1(ii).

6.4.2 Mise en œuvre de l'algorithme

Pour minimiser la fonction (6.38), nous allons utiliser l'algorithme 5.4 en alternant sur les variables $\mathbf{x}$ et $\mathbf{k}$. Pour cela, nous avons besoin de définir, pour tout $(\tilde{\mathbf{x}}, \tilde{\mathbf{k}}) \in \mathbb{R}^N \times \mathbb{R}^S$, des matrices $\mathbf{A}_1(\tilde{\mathbf{x}}, \tilde{\mathbf{k}}) \in \mathcal{S}_N^+$ et $\mathbf{A}_2(\tilde{\mathbf{x}}, \tilde{\mathbf{k}}) \in \mathcal{S}_S^+$ satisfaisant l'hypothèse 5.2. De telles matrices sont données dans la proposition suivante.

Proposition 6.3. *Pour tout $(\tilde{\mathbf{x}}, \tilde{\mathbf{k}}) \in \mathbb{R}^N \times \mathbb{R}^S$, soient*

$$\begin{aligned}\mathbf{A}_1(\tilde{\mathbf{x}}, \tilde{\mathbf{k}}) &= \left(L_1(\tilde{\mathbf{k}}) + \frac{9\lambda}{8\eta^2} \right) \mathbf{I}_N + \frac{\lambda}{\ell_{1,\alpha}(\tilde{\mathbf{x}}) + \beta} \mathbf{A}_{\ell_{1,\alpha}}(\tilde{\mathbf{x}}), \\ \mathbf{A}_2(\tilde{\mathbf{x}}, \tilde{\mathbf{k}}) &= (L_2(\tilde{\mathbf{x}}) + \varepsilon) \mathbf{I}_S,\end{aligned}$$

où

$$\mathbf{A}_{\ell_{1,\alpha}}(\tilde{\mathbf{x}}) = \text{Diag} \left(\left(((\tilde{x}^{(n)})^2 + \alpha^2)^{-1/2} \right)_{1 \leq n \leq N} \right), \quad (6.44)$$

$\varepsilon > 0$ et $L_1(\tilde{\mathbf{k}})$ (resp. $L_2(\tilde{\mathbf{x}})$) est une constante de Lipschitz de $\nabla_1 \mathbf{h}_1(\cdot, \tilde{\mathbf{k}})$ (resp. $\nabla_2 \mathbf{h}_1(\tilde{\mathbf{x}}, \cdot)$).⁴ La fonction définie par

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \mathbf{q}(\mathbf{x} | \tilde{\mathbf{x}}, \tilde{\mathbf{k}}) = \mathbf{h}(\tilde{\mathbf{x}}, \tilde{\mathbf{k}}) + \langle \mathbf{x} - \tilde{\mathbf{x}}, \nabla_1 \mathbf{h}(\tilde{\mathbf{x}}, \tilde{\mathbf{k}}) \rangle + \frac{1}{2} \|\mathbf{x} - \tilde{\mathbf{x}}\|_{\mathbf{A}_1(\tilde{\mathbf{x}}, \tilde{\mathbf{k}})}^2, \quad (6.45)$$

est une fonction quadratique majorant $\mathbf{h}(\cdot, \tilde{\mathbf{k}})$ en $\tilde{\mathbf{x}}$, et

$$(\forall \mathbf{k} \in \mathbb{R}^S) \quad \mathbf{q}(\mathbf{k} | \tilde{\mathbf{x}}, \tilde{\mathbf{k}}) = \mathbf{h}(\tilde{\mathbf{x}}, \tilde{\mathbf{k}}) + \langle \mathbf{k} - \tilde{\mathbf{k}}, \nabla_1 \mathbf{h}(\tilde{\mathbf{x}}, \tilde{\mathbf{k}}) \rangle + \frac{1}{2} \|\mathbf{k} - \tilde{\mathbf{k}}\|_{\mathbf{A}_2(\tilde{\mathbf{x}}, \tilde{\mathbf{k}})}^2, \quad (6.46)$$

est une fonction quadratique majorant $\mathbf{h}(\tilde{\mathbf{x}}, \cdot)$ en $\tilde{\mathbf{k}}$.

Démonstration. Soit $\tilde{\mathbf{x}} \in \mathbb{R}^N$. Décomposons $\mathbf{h}_2$ en une somme de deux fonctions $\mathbf{h}_2^{(1)}$ et $\mathbf{h}_2^{(2)}$ où

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \begin{cases} \mathbf{h}_2^{(1)}(\mathbf{x}) = \lambda \log(\ell_{1,\alpha}(\mathbf{x}) + \beta), \\ \mathbf{h}_2^{(2)}(\mathbf{x}) = -\lambda \log(\ell_{2,\eta}(\mathbf{x})). \end{cases} \quad (6.47)$$

Il nous suffit alors de montrer que

(i) la fonction

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \tilde{\mathbf{q}}(\mathbf{x} | \tilde{\mathbf{x}}) = \mathbf{h}_2^{(1)}(\tilde{\mathbf{x}}) + \langle \mathbf{x} - \tilde{\mathbf{x}}, \nabla \mathbf{h}_2^{(1)}(\tilde{\mathbf{x}}) \rangle + \frac{1}{2} \|\mathbf{x} - \tilde{\mathbf{x}}\|_{\mathbf{A}_{\mathbf{h}_2^{(1)}}(\tilde{\mathbf{x}})}^2, \quad (6.48)$$

est une fonction majorante de $\mathbf{h}_2^{(1)}$ en $\tilde{\mathbf{x}}$ pour

$$\mathbf{A}_{\mathbf{h}_2^{(1)}}(\tilde{\mathbf{x}}) = \frac{\lambda}{\ell_{1,\alpha}(\tilde{\mathbf{x}}) + \beta} \mathbf{A}_{\ell_{1,\alpha}}(\tilde{\mathbf{x}}), \quad (6.49)$$

4. La fonction $\mathbf{h}_1$ étant choisie quadratique, ces constantes de Lipschitz s'obtiennent facilement.

(ii) une constante de Lipschitz du gradient de $h_2^{(2)}$ est égale à $\mu = \frac{9\lambda}{8\eta^2}$.

D'une part, posons

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad l(\mathbf{x}) = \ell_{1,\alpha}(\mathbf{x}) + \beta. \quad (6.50)$$

D'après [Allain et al., 2006], on a alors, pour tout $\mathbf{x} \in \mathbb{R}^N$,

$$l(\mathbf{x}) \leq l(\tilde{\mathbf{x}}) + (\mathbf{x} - \tilde{\mathbf{x}})^\top \nabla l(\tilde{\mathbf{x}}) + \frac{1}{2} \|\mathbf{x} - \tilde{\mathbf{x}}\|_{\mathbf{A}_{\ell_{1,\alpha}}(\tilde{\mathbf{x}})}^2, \quad (6.51)$$

où $\mathbf{A}_{\ell_{1,\alpha}}(l)$ est donnée par (6.44).

D'autre part, pour tout $(u, v) \in]0, +\infty[^2$,

$$\log v \leq \log u + \frac{v}{u} - 1 = \log u + \frac{v - u}{u}. \quad (6.52)$$

En posant $v = l(\mathbf{x}) > 0$ et $u = l(\tilde{\mathbf{x}}) > 0$, et en combinant (6.51) et (6.52), on obtient

$$h_2^{(1)}(\mathbf{x}) \leq h_2^{(1)}(\tilde{\mathbf{x}}) + \langle \mathbf{x} - \tilde{\mathbf{x}}, \frac{\lambda}{\tau(\tilde{\mathbf{x}})} \nabla l(\tilde{\mathbf{x}}) \rangle + \frac{1}{2} \langle \mathbf{x} - \tilde{\mathbf{x}}, \frac{\lambda}{\tau(\tilde{\mathbf{x}})} \mathbf{A}_{\ell_{1,\alpha}}(\tilde{\mathbf{x}})(\mathbf{x} - \tilde{\mathbf{x}}) \rangle.$$

Ainsi, l'assertion (i) est vérifiée en remarquant que

$$\nabla h_2^{(1)}(\tilde{\mathbf{x}}) = \frac{\lambda}{l(\tilde{\mathbf{x}})} \nabla l(\tilde{\mathbf{x}}) \quad \text{et} \quad \mathbf{A}_{h_2^{(1)}}(\tilde{\mathbf{x}}) = \frac{\lambda}{l(\tilde{\mathbf{x}})} \mathbf{A}_{\ell_{1,\alpha}}(\tilde{\mathbf{x}}).$$

D'autre part, le Hessien de $h_2^{(2)}$ est donné par

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \nabla^2 h_2^{(2)}(\mathbf{x}) = \frac{2\lambda}{\ell_{2,\eta}^4(\mathbf{x})} \mathbf{x} \mathbf{x}^\top - \frac{\lambda}{\ell_{2,\eta}^2(\mathbf{x})} \mathbf{I}_N.$$

En remarquant que, pour tout $\mathbf{x} \in \mathbb{R}^N$, $\ell_{2,\eta}^2(\mathbf{x}) = \|\mathbf{x}\|^2 + \eta^2$, et en appliquant l'inégalité triangulaire, on obtient

$$\|\nabla^2 h_2^{(2)}(\mathbf{x})\| \leq \frac{2\lambda \|\mathbf{x}\|^2}{(\|\mathbf{x}\|^2 + \eta^2)^2} + \frac{\lambda}{\|\mathbf{x}\|^2 + \eta^2} = \chi(\|\mathbf{x}\|),$$

où

$$\chi: u \in [0, +\infty[\mapsto \lambda \frac{3u^2 + \eta^2}{(u^2 + \eta^2)^2}.$$

La dérivée de χ est donnée, pour tout $u \in [0, +\infty[$, par

$$\dot{\chi}(u) = \lambda \frac{2u}{(u^2 + \eta^2)^3} (\eta^2 - 3u^2),$$

donc χ est une fonction croissante sur $[0, \eta/\sqrt{3}]$ et décroissante sur $] \eta/\sqrt{3}, +\infty[$. De plus, on a $\sup_{u \in [0, +\infty[} \chi(u) = \chi(\eta/\sqrt{3}) = \frac{9\lambda}{8\eta^2}$. On obtient alors l'assertion (ii). $\square$

Dans le cas particulier de la minimisation de la fonction (6.38), l'algorithme 5.4 se réduit à l'algorithme 6.2. Les propriétés de convergence de cet algorithme sont données dans le théorème 5.1. Remarquons que dans l'algorithme 6.2, pour tout $k \in \mathbb{N}$, les entiers I_k et J_k permettent de mettre à jour plusieurs fois de suite les vecteurs $\mathbf{x}_k$ et $\mathbf{k}_k$. Cela est possible grâce à l'hypothèse de mise à jour quasi-cyclique des blocs donnée dans la définition 5.1(ii). Nous montrerons l'intérêt de cette règle de mis à jour dans la section suivante.

Algorithme 6.2 Algorithme explicite-implicite à métrique variable alterné par bloc appliqué au problème de déconvolution aveugle.

Initialisation : Pour tout $k \in \mathbb{N}$, soient $J_k \in \mathbb{N}^*$, $I_k \in \mathbb{N}^*$ et soient $(\gamma_{\mathbf{x}}^{k,j})_{0 \leq j \leq J_k-1} \in]0, +\infty[^{J_k}$ et $(\gamma_{\mathbf{k}}^{k,i})_{0 \leq i \leq I_k-1} \in]0, +\infty[^{I_k}$. Soient $\mathbf{x}_0 \in \text{dom } \mathbf{g}_1$ et $\mathbf{k}_0 \in \text{dom } \mathbf{g}_2$.

Iterations :

Pour $k = 0, 1, \dots$

$$\begin{array}{l}
 \left[\begin{array}{l}
 \mathbf{x}_{k,0} = \mathbf{x}_k, \mathbf{k}_{k,0} = \mathbf{k}_k, \\
 \text{Pour } j = 0, \dots, J_k - 1 \\
 \left[\begin{array}{l}
 \tilde{\mathbf{x}}_{k,j} = \mathbf{x}_{k,j} - \gamma_{\mathbf{x}}^{k,j} \mathbf{A}_1(\mathbf{x}_{k,j}, \mathbf{k}_k)^{-1} \nabla_1 h(\mathbf{x}_{k,j}, \mathbf{k}_k), \\
 \mathbf{x}_{k,j+1} = \text{prox}_{(\gamma_{\mathbf{x}}^{k,j})^{-1} \mathbf{A}_1(\mathbf{x}_{k,j}, \mathbf{k}_k), \mathbf{g}_1}(\tilde{\mathbf{x}}_{k,j}), \\
 \mathbf{x}_{k+1} = \mathbf{x}_{k,J_k}.
 \end{array} \right. \\
 \text{Pour } i = 0, \dots, I_k - 1 \\
 \left[\begin{array}{l}
 \tilde{\mathbf{k}}_{k,i} = \mathbf{k}_{k,i} - \gamma_{\mathbf{k}}^{k,i} \mathbf{A}_2(\mathbf{x}_{k+1}, \mathbf{k}_{k,i})^{-1} \nabla_2 h(\mathbf{x}_{k+1}, \mathbf{k}_{k,i}), \\
 \mathbf{k}_{k,i+1} = \text{prox}_{(\gamma_{\mathbf{k}}^{k,i})^{-1} \mathbf{A}_2(\mathbf{x}_{k+1}, \mathbf{k}_{k,i}), \mathbf{g}_2}(\tilde{\mathbf{k}}_{k,i}), \\
 \mathbf{k}_{k+1} = \mathbf{k}_{k,I_k}.
 \end{array} \right.
 \end{array} \right.
 \end{array}$$

6.4.3 Résultats de simulation

Le signal sismique parcimonieux $\bar{\mathbf{x}}$, de longueur $N = 784$, donné en haut de la figure 6.10, est composé d'une suite de "pics" appelés coefficients de réflectivité primaire [Walden et Hosken, 1986]. Ces séries de réflectivités indiquent le temps de parcours d'une onde sismique entre deux réflecteurs sismiques, et l'amplitude des événements sismiques réfléchis vers le capteur. La trace sismique observée $\mathbf{z}$, donnée en bas de la figure 6.10, suit le modèle (6.37). Dans ce contexte, le noyau $\bar{\mathbf{k}}$ est lié à la source sismique générée. Nous utilisons ici une onde sismique passe-bande de "Ricker" (appelée aussi chapeau Mexicain [Ricker, 1940]) de taille $S = 41$ (voir figure 6.13(2)) dont les fréquences de spectre sont concentrées entre 10 et 40 Hz. Le bruit additif $\mathbf{w}$ est une réalisation d'un bruit blanc gaussien de variance σ^2 .

Les séries de réflectivité étant parcimonieuses, mais limitées en amplitude, nous choisissons comme fonction de régularisation

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \mathbf{g}_1(\mathbf{x}) = \iota_{[\mathbf{x}_{\min}, \mathbf{x}_{\max}]^N}(\mathbf{x}), \quad (6.53)$$

où $(\mathbf{x}_{\min}, \mathbf{x}_{\max}) \in \mathbb{R}^2$ sont les amplitudes minimale et maximale de $\bar{\mathbf{x}}$. De même, l'onde sismique étant d'énergie finie, on choisit

$$(\forall \mathbf{k} \in \mathbb{R}^S) \quad \mathbf{g}_2(\mathbf{k}) = \iota_{\mathcal{C}}(\mathbf{k}), \quad (6.54)$$

où $\mathcal{C} = \{\mathbf{k} \in [\mathbf{k}_{\min}, \mathbf{k}_{\max}]^S \mid \|\mathbf{k}\| \leq \delta\}$, avec $\delta > 0$, et $\mathbf{k}_{\min}$ (resp. $\mathbf{k}_{\max}$) la valeur minimale (resp. maximale) de $\bar{\mathbf{k}}$.

La figure 6.12 donne les temps de reconstruction obtenus, en secondes, en utilisant l'algorithme 6.2, en fonction du nombre d'itérations internes $J_k \equiv J$, en fixant $I_k \equiv 1$, $\gamma_{\mathbf{x}}^{k,i} \equiv 1$, $\gamma_{\mathbf{k}}^{k,j} \equiv 1.9$ et le niveau de bruit $\sigma = 0.03$. Les temps de reconstruction donnés correspondent

au temps nécessaire pour que le critère d'arrêt $\|\mathbf{x}_k - \mathbf{x}_{k-1}\| \leq \sqrt{N} \times 10^{-6}$ soit vérifié. On peut observer que le meilleur compromis en terme de vitesse de convergence est obtenu pour une valeur intermédiaire d'itérations internes sur $\mathbf{x}_k$ égale à $J = 71$. D'autre part, nous avons pu remarquer que le choix de J influe très peu sur la qualité de reconstruction.

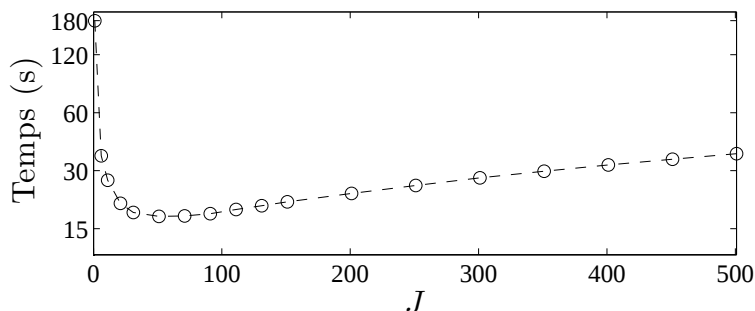


Figure 6.12 — Temps de reconstruction en fonction du nombre d'itérations internes $J_k \equiv J$ effectuées dans l'algorithme 6.2 (moyennes obtenues pour 30 réalisations de bruit).

La méthode proposée dans cette section est comparée à celle de [Krishnan et al., 2011]. Les deux algorithmes sont initialisés de la même manière : $\mathbf{x}_0$ est un signal constant vérifiant $\|\mathbf{x}_0\| \leq \max\{|x_{\min}|, |x_{\max}|\}$, et $\mathbf{k}_0$ est un filtre gaussien centré appartenant à l'ensemble $\mathcal{C}$. Dans les tables 6.1 et 6.2 nous présentons, respectivement, les moyennes et variances obtenues pour trois niveaux de bruit $\sigma \in \{0.01, 0.02, 0.03\}$, moyennés sur 200 réalisations de bruit chacun. Les paramètres de régularisation de [Krishnan et al., 2011] et $(\lambda, \alpha, \beta, \eta) \in]0, +\infty[^4$ apparaissant dans (6.41) sont choisis de manière à minimiser la norme ℓ_1 entre le signal original $\bar{\mathbf{x}}$ et le signal estimé $\hat{\mathbf{x}}$. De plus, pour l'algorithme 6.2, nous prenons, pour tout $k \in \mathbb{N}$, $J_k = 71$ et $I_k = 1$. Bien que les deux méthodes donnent de bons résultats de reconstruction en terme d'erreurs ℓ_1 et ℓ_2 , l'algorithme 6.2 donne de meilleurs résultats, pour tous les niveaux de bruits, autant pour la reconstruction de $\bar{\mathbf{x}}$ que de $\bar{\mathbf{h}}$. D'autre part, d'après la table 6.2, nous pouvons remarquer que les variances des erreurs ℓ_1 et ℓ_2 sont du même ordre de grandeur pour les deux méthodes (autour de 10% de l'erreur moyenne), ce qui permet de confirmer statistiquement la stabilité des résultats présentés dans la table 6.1. On peut remarquer de plus que, bien que les variances soient légèrement plus élevées avec l'algorithme 6.2 qu'avec [Krishnan et al., 2011] pour l'estimation de $\bar{\mathbf{x}}$, le comportement inverse est observé pour l'estimation de $\bar{\mathbf{k}}$. Enfin, l'algorithme 6.2 est significativement plus rapide en moyenne, et sa vitesse de convergence est plus stable que [Krishnan et al., 2011].

Finalement, nous donnons dans la figure 6.13 un résultat de reconstruction pour $\sigma = 0.03$. La figure 6.13(1) illustre l'erreur résiduelle de l'estimation parcimonieuse du signal $\bar{\mathbf{x}} - \hat{\mathbf{x}}$, pour une réalisation de bruit donnée, où $\hat{\mathbf{x}}$ est estimée par [Krishnan et al., 2011] dans (a), et avec l'algorithme 6.2 dans (b). De même, nous donnons dans la figure 6.13(2) l'estimation de $\bar{\mathbf{k}}$ pour la même réalisation de bruit.

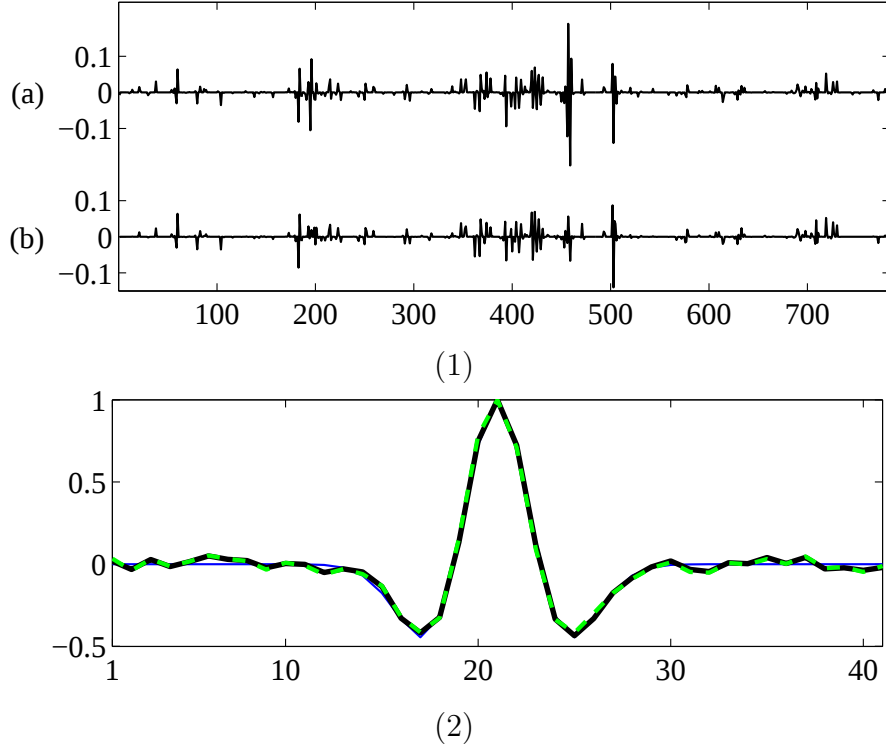


Figure 6.13 — (1) *Erreur résiduelle $\bar{\mathbf{x}} - \hat{\mathbf{x}}$ pour le signal estimé $\hat{\mathbf{x}}$ obtenu en utilisant [Krishnan et al., 2011] (a) et l’algorithme 6.2 (b).* (2) *Noyau original $\bar{\mathbf{k}}$ (ligne fine continue bleue), estimé $\hat{\mathbf{k}}$ par l’algorithme 6.2 (ligne épaisse continue noire) et par [Krishnan et al., 2011] (ligne discontinue épaisse verte).*

6.5 CONCLUSION

Dans ce chapitre nous avons traité trois problèmes inverses en traitement du signal, que nous avons résolus en utilisant l’algorithme explicite-implicite à métrique variable alterné par bloc présenté dans le chapitre 5. Nous avons vu que cet algorithme permet de résoudre des problèmes où deux objets sont à estimer conjointement (NMF et déconvolution aveugle), et permet également de traiter des problèmes de grande dimension tel que celui de reconstruction de phase présenté dans la section 6.3. Nous avons montré, sur les trois exemples de simulation, que notre méthode se compare favorablement à celles existantes. Nous avons montré de plus que, la métrique variable permet d’accélérer la convergence de l’algorithme explicite-implicite alterné par bloc dans les exemples traités. Pour finir, nous avons vu dans le dernier exemple l’intérêt d’utiliser la règle essentiellement cyclique pour la mise à jour des blocs, lorsque ces derniers ne sont pas de même nature.

Niveau de bruit (σ)		0.01	0.02	0.03
Erreur sur le signal observé $\mathbf{z}$	$\ell_2 (\times 10^{-2})$	7.14	7.35	7.68
	$\ell_1 (\times 10^{-2})$	2.85	3.44	4.09
Erreur sur le signal estimé $\hat{\mathbf{x}}$	[Krishnan et al., 2011]	$\ell_2 (\times 10^{-2})$	1.23	1.66
		$\ell_1 (\times 10^{-3})$	3.79	4.69
	Algorithme 6.2	$\ell_2 (\times 10^{-2})$	1.09	1.63
		$\ell_1 (\times 10^{-3})$	3.42	4.30
Erreur sur le noyau estimé $\hat{\mathbf{k}}$	[Krishnan et al., 2011]	$\ell_2 (\times 10^{-2})$	1.88	2.51
		$\ell_1 (\times 10^{-2})$	1.44	1.96
	Algorithme 6.2	$\ell_2 (\times 10^{-2})$	1.62	2.26
		$\ell_1 (\times 10^{-2})$	1.22	1.77
Temps de reconstruction (s)	[Krishnan et al., 2011]	106	61	56
	Algorithme 6.2	56	22	18

Table 6.1 – Moyennes sur 200 réalisations de bruit, pour chaque niveau de bruit, des résultats obtenus pour l'estimation de $\bar{\mathbf{x}}$ et $\bar{\mathbf{k}}$ en utilisant [Krishnan et al., 2011] et l'algorithme 6.2 (Matlab 8 sur une machine Intel(R) Xeon(R) CPU E5-2609 v2@2.5GHz).

Niveau de bruit (σ)		0.01	0.02	0.03
Erreur sur le signal observé $\mathbf{z}$	$\ell_2 (\times 10^{-4})$	3.52	7.02	10.42
	$\ell_1 (\times 10^{-4})$	2.98	5.55	8.06
Erreur sur le signal estimé $\hat{\mathbf{x}}$	[Krishnan et al., 2011]	$\ell_2 (\times 10^{-3})$	1.35	2.35
		$\ell_1 (\times 10^{-4})$	2.17	3.36
	Algorithme 6.2	$\ell_2 (\times 10^{-3})$	2.92	3.44
		$\ell_1 (\times 10^{-4})$	3.88	4.81
Erreur sur le noyau estimé $\hat{\mathbf{k}}$	[Krishnan et al., 2011]	$\ell_2 (\times 10^{-3})$	2.65	6.24
		$\ell_1 (\times 10^{-3})$	1.99	3.96
	Algorithme 6.2	$\ell_2 (\times 10^{-3})$	2.94	2.75
		$\ell_1 (\times 10^{-3})$	2.04	2.25
Temps de reconstruction (s)	[Krishnan et al., 2011]	54	45	46
	Algorithme 6.2	11	4	6

Table 6.2 – Variances sur 200 réalisations de bruit, pour chaque niveau de bruit, des résultats obtenus pour l'estimation de $\bar{\mathbf{x}}$ et $\bar{\mathbf{k}}$ en utilisant [Krishnan et al., 2011] et l'algorithme 6.2.

CHAPITRE 7

Algorithmes primaux-duaux alternés et distribués pour les problèmes d'inclusion monotone et d'optimisation convexe

Sommaire

7.1	Introduction	144
7.2	Opérateurs monotones et optimisation	145
7.2.1	Notations	145
7.2.2	Algorithme explicite-implicite alterné aléatoirement par bloc	147
7.3	Algorithmes primaux-duaux alternés aléatoirement par bloc	151
7.3.1	Problème	152
7.3.2	Première sous-classe d'algorithmes	153
7.3.3	Seconde sous-classe d'algorithmes	159
7.3.4	Application aux problèmes variationnels	163
7.4	Application de l'algorithme primal-dual alterné aléatoirement par bloc	168
7.4.1	Description du problème	168
7.4.2	Méthodes proposées	170
7.4.3	Résultats numériques	172
7.5	Algorithmes distribués	178
7.5.1	Problème et notations	178
7.5.2	Algorithme distribués asynchrones pour les problèmes d'in- clusion monotone	181
7.5.3	Application aux problèmes variationnels	187
7.6	Conclusion	192
Annexe 7.A	Démonstration du lemme 7.1	194
Annexe 7.B	Démonstration de la proposition 7.3	195
Annexe 7.C	Démonstration de la proposition 7.5	197
Annexe 7.D	Démonstration de la proposition 7.7	198
Annexe 7.E	Démonstration de la proposition 7.11	199
Annexe 7.F	Démonstration de la proposition 7.12	202

7.1 INTRODUCTION

Ces dernières années, les méthodes primales-duales ont suscité un intérêt croissant. Elles permettent de minimiser des sommes de fonctions convexes (voir section 2.4.4 et [Komodakis et Pesquet, 2014]) mais aussi de trouver le zéro d'une somme d'opérateurs monotones [Bauschke et Combettes, 2011]. Lorsque plusieurs opérateurs linéaires apparaissent dans la formulation du problème étudié, résoudre conjointement les formes primales et duales permet de construire des stratégies d'optimisation où les opérateurs linéaires n'ont pas besoin d'être inversés, ce qui offre un avantage indéniable en terme de complexité algorithmique lorsque l'on traite des problèmes de grande taille (voir par exemple [Becker et Combettes, 2014; Couprie *et al.*, 2013; Harizanov *et al.*, 2013; Jezierska *et al.*, 2012; Pustelnik *et al.*, 2014; Repetti *et al.*, 2012; Teuber *et al.*, 2013]).

Dans la section 2.4.4, nous avons vu qu'il existe plusieurs classes d'algorithmes primaux-duaux. Nous nous concentrerons sur la première classe, basée sur l'algorithme explicite-implicite, déjà étudiée dans de nombreux travaux [Chambolle et Pock, 2010; Chen *et al.*, 2013; Combettes *et al.*, 2014, 2010; Combettes et Vũ, 2014; Condat, 2013; Esser *et al.*, 2010; Goldstein *et al.*, 2013; He et Yuan, 2012; Loris et Verhoeven, 2011; Pock et Chambolle, 2011; Vũ, 2013]. Nous avons vu que pour minimiser une somme de fonctions convexes, ces algorithmes exploitent les gradients des fonctions différentiables et les opérateurs proximaux des autres fonctions. Dans le cas des opérateurs monotones, cela signifie que l'opérateur est appliqué de façon explicite s'il est cocoercif, sinon sa résolvante est utilisée.

La plupart des algorithmes proposés dans les travaux mentionnés ci-dessus consistent à éclater le problème original en une somme de termes plus simples dont les opérateurs associés peuvent être traités individuellement, de façon parallèle, à chaque itération de l'algorithme. Notre objectif dans ce chapitre est de rendre plus flexible les algorithmes primaux-duaux existants, en s'autorisant à activer, à chaque itération, une sous-partie seulement des opérateurs impliqués. Pour cela, nous baserons sur les travaux récents de [Combettes et Pesquet, 2015], proposant des approches basées sur des techniques de balayage aléatoire applicables à des algorithmes générant des suites (quasi-)Fejér monotones. Les algorithmes résultants de ces approches aléatoires sont tolérants à des erreurs stochastiques satisfaisant une condition de sommabilité.

Dans ce chapitre, nous donnerons des versions alternées par bloc de deux variantes d'algorithmes primaux-duaux basés sur l'approche explicite-implicite. Comme nous l'avons vu dans les chapitres 5 et 6, les méthodes par bloc présentent des avantages algorithmiques en terme de vitesse de convergence et d'allocation de mémoire. Nous verrons de plus dans ce chapitre qu'elles permettent de développer des stratégies distribuées. Plus précisément, nous allons nous intéresser à des problèmes multi-agents où les mises à jours peuvent être effectuées de manière asynchrone sur le voisinage d'un nombre limité d'agents. Nous montrerons que les algorithmes développés permettent de traiter des problèmes variationnels et d'inclusion monotone, et nous ferons le lien avec les méthodes existantes. Notons que, dans le cas variationnel, des méthodes primales-duales distribuées existent, mettant en œuvre des étapes de sous-gradients [Chang *et al.*, 2014; Yuan *et al.*, 2011] (voir aussi [Towfic et Sayed, 2014] pour des applications à des réseaux de données). D'une manière générale, les méthodes de sous-gradients convergent sous la condition que le pas tend vers zéro (voir section 2.4.1). Cependant, l'utilisation de l'opérateur proximal, qui peut être vu comme une étape implicite du sous-gradient, permet d'avoir des hypothèses moins restrictives sur le pas. Par exemple, la convergence des itérées peut être établie pour des pas constants.

La suite du chapitre est organisée de la façon suivante : dans la section 7.2 nous ferons des rappels sur la théorie des opérateurs monotones et donnerons nos notations. Dans la section 7.3, nous proposerons des algorithmes primaux-duaux alternés par bloc appliqués à des problèmes d'inclusion monotone et à des problèmes variationnels. Dans la section 7.4 nous traiterons une application en reconstruction de maillage 3D. Nous développerons ensuite des algorithmes distribués dans la section 7.5, permettant de résoudre des problèmes d'inclusion monotone et des problèmes variationnels. Finalement, nous donnerons quelques conclusions dans la section 7.6.

7.2 OPÉRATEURS MONOTONES ET OPTIMISATION

Nous ne donnerons dans cette section que les définitions et propriétés qui seront utiles pour notre étude. Pour de plus amples informations sur les opérateurs monotones, le lecteur peut se référer à [Bauschke et Combettes, 2011], et pour les notions de probabilités, à [Fortet, 1995].

7.2.1 Notations

Dans ce chapitre, nous notons $(\Omega, \mathcal{F}, \mathbb{P})$ l'espace de probabilité sous-jacent. Soit $\mathcal{H}$ un espace de Hilbert réel séparable. On note $\mathcal{B}$ sa σ -algèbre de Borel. Une variable aléatoire de $\mathcal{H}$ est une application mesurable $x: (\Omega, \mathcal{F}) \rightarrow (\mathcal{H}, \mathcal{B})$. La plus petite σ -algèbre générée par une famille Φ de variables aléatoires est noté $\sigma(\Phi)$. De plus, l'espérance est notée $\mathbb{E}(\cdot)$.

Soit $\mathcal{G}$ un espace de Hilbert réel. Soit $L \in \mathcal{B}(\mathcal{H}, \mathcal{G})^1$, on note son opérateur adjoint L^* , son inverse L^{-1} et sa racine carrée $L^{1/2}$. Si $L \in \mathcal{S}^+(\mathcal{H})^2$, alors L est un isomorphisme et son inverse L^{-1} est aussi dans $\mathcal{S}^+(\mathcal{H})$.

Nous notons $2^{\mathcal{H}}$ l'ensemble des parties de l'ensemble $\mathcal{H}$. Soit $A: \mathcal{H} \rightarrow 2^{\mathcal{H}}$. Si, pour tout $x \in \mathcal{H}$, Ax est un singleton, alors A est un opérateur univalué.

Définition 7.1. Soit $A: \mathcal{H} \rightarrow 2^{\mathcal{H}}$.

(i) L'ensemble des zéros de A est donné par

$$\text{zer } A = \{x \in \mathcal{H} \mid 0 \in Ax\}.$$

(ii) L'inverse de A est défini par

$$A^{-1}: \mathcal{H} \mapsto 2^{\mathcal{H}}: u \mapsto \{x \in \mathcal{H} \mid u \in Ax\}.$$

(iii) On dit que A est un opérateur monotone si

$$(\forall (x, y) \in \mathcal{H}^2)(\forall u \in Ax)(\forall v \in Ay) \quad \langle x - y \mid u - v \rangle \geq 0.$$

1. Rappelons que $\mathcal{B}(\mathcal{H}, \mathcal{G})$ correspond à l'ensemble des opérateurs linéaires bornés de $\mathcal{H}$ vers $\mathcal{G}$.
2. Rappelons que $\mathcal{S}(\mathcal{H})$ correspond à l'ensemble des opérateurs $L \in \mathcal{B}(\mathcal{H}) = \mathcal{B}(\mathcal{H}, \mathcal{H})$ auto-adjoints. De plus, $\mathcal{S}^+(\mathcal{H})$ est l'ensemble des opérateurs $L \in \mathcal{S}(\mathcal{H})$ fortement positifs.

(iv) On dit que A est un opérateur maximal monotone s'il est monotone et s'il n'existe pas d'autre opérateur monotone dont le graphe est inclus dans celui de A .

(v) On dit que A est un opérateur β -fortement monotone, pour $\beta \in]0, +\infty[$, si

$$(\forall (x, y) \in \mathcal{H}^2)(\forall u \in Ax)(\forall v \in Ay) \quad \langle x - y \mid u - v \rangle \geq \beta \|x - y\|^2.$$

(vi) La résolvante de A est définie par

$$J_A: \mathcal{H} \rightarrow 2^{\mathcal{H}}: x \mapsto J_A x = (\text{Id} + A)^{-1}x.$$

(vii) Soit $C: \mathcal{H} \rightarrow 2^{\mathcal{H}}$. La somme parallèle entre A et C est définie par

$$A \square C = (A^{-1} + C^{-1})^{-1}.$$

L'élément neutre de la somme parallèle est le cône normal à $\{0\}$, noté $N_{\{0\}}$.

Définition 7.2. Soit B un opérateur univalué de $\mathcal{H}$ dans $\mathcal{H}$.

(i) L'opérateur B est dit β -cocoercif, pour $\beta \in]0, +\infty[$, si

$$(\forall (x, y) \in \mathcal{H}^2) \quad \langle x - y \mid Bx - By \rangle \geq \beta \|Bx - By\|^2.$$

(ii) L'opérateur B est dit fermement contractant s'il est 1-cocoercif.

(iii) L'opérateur B est dit α -moyenné, pour $\alpha \in]0, 1[$, si

$$(\forall (x, y) \in \mathcal{H}^2) \quad \|Bx - By\|^2 \leq \|x - y\|^2 - \frac{1 - \alpha}{\alpha} \|(\text{Id} - B)x - (\text{Id} - B)y\|^2$$

La propriété suivante permet de faire le lien entre certaines des définitions données ci-dessus.

Propriété 7.1.

(i) Soient $B: \mathcal{H} \rightarrow \mathcal{H}$ et $\beta \in]0, +\infty[$. L'opérateur B est β -cocoercif si et seulement si $B^{-1}: \mathcal{H} \rightarrow 2^{\mathcal{H}}$ est β -fortement monotone.

(ii) Un opérateur $A: \mathcal{H} \rightarrow 2^{\mathcal{H}}$ est maximal monotone si et seulement si sa résolvante est un opérateur fermement contractant de $\mathcal{H}$ vers $\mathcal{H}$.

Dans le chapitre 2 nous avons énoncé le théorème de Moreau pour les fonctions de $\Gamma_0(\mathcal{H})$. Le théorème suivant le généralise au cas des opérateurs maximaux monotones.

Théorème 7.1. (Formule de décomposition de Moreau)[Combettes et Vũ, 2014, Ex. 3.9]

Soient $A: \mathcal{H} \rightarrow 2^{\mathcal{H}}$ un opérateur maximal monotone, $U \in \mathcal{S}^+(\mathcal{H})$ et $\gamma \in]0, +\infty[$. On note alors $J_{\gamma U A}: \mathcal{H} \rightarrow \mathcal{H}$ la résolvante de A relativement à γU , et on a

$$J_{\gamma U A} = U^{1/2} J_{\gamma U^{1/2} A U^{1/2}} U^{-1/2} = \text{Id} - \gamma U J_{\gamma^{-1} U^{-1} A^{-1}} (\gamma^{-1} U^{-1} \cdot). \quad (7.1)$$

Nous allons maintenant nous intéresser aux liens existants entre les opérateurs monotones et les fonctions convexes.

Propriété 7.2. *Soit $f: \mathcal{H} \rightarrow]-\infty, +\infty]$.*

- (i) *Si $f \in \Gamma_0(\mathcal{H})$, le sous-différentiel de Moreau ∂f donné par la Définition 2.8 est un opérateur maximal monotone.*
- (ii) *Si f est une fonction propre et β -fortement convexe, pour $\beta \in]0, +\infty[$, alors ∂f est β -fortement monotone.*
- (iii) **Théorème de Baillon-Haddad.** *Soit $\beta \in]0, +\infty[$. Si f est différentiable et convexe, alors le gradient de f , ∇f , est β^{-1} -Lipschitz si et seulement s'il est β -cocoercif.*
- (iv) *Soit $U \in \mathcal{S}^+(\mathcal{H})$. Si $f \in \Gamma_0(\mathcal{H})$, alors $\text{prox}_{U,f} = J_{U^{-1}\partial f}$.*

Soient $(\mathcal{G}_i)_{1 \leq i \leq m}$ des espaces de Hilbert réels. Nous notons $\mathcal{G} = \mathcal{G}_1 \oplus \dots \oplus \mathcal{G}_m$ leur somme directe de Hilbert, i.e. l'espace produit muni du produit scalaire :

$$(\forall \mathbf{x} = (x^{(i)})_{1 \leq i \leq m} \in \mathcal{G})(\forall \mathbf{y} = (y^{(i)})_{1 \leq i \leq m} \in \mathcal{G}) \quad (\mathbf{x}, \mathbf{y}) \mapsto \sum_{i=1}^m \langle x^{(i)} | y^{(i)} \rangle,$$

où, pour tout $i \in \{1, \dots, m\}$, $x^{(i)} \in \mathcal{G}_i$ et $y^{(i)} \in \mathcal{G}_i$.

De plus, nous définissons $\mathbb{D}_m = \{0, 1\}^m \setminus \{\mathbf{0}\}$ l'ensemble des codes binaires non nuls de longueur m .

7.2.2 Algorithme explicite-implicite alterné aléatoirement par bloc

Dans un premier temps nous allons rappeler les résultats principaux présentés dans [Combettes et Pesquet, 2015] sur lesquels nous nous baserons dans ce chapitre.

On suppose que $\mathcal{K}_1, \dots, \mathcal{K}_m$ sont des espaces de Hilbert réels séparables, où $m \in \mathbb{N}^*$, et $\mathcal{K} = \mathcal{K}_1 \oplus \dots \oplus \mathcal{K}_m$ correspond à leur somme directe de Hilbert.

On définit, pour tout $k \in \mathbb{N}$,

- $\mathbf{S}_k: \mathcal{H} \rightarrow \mathcal{H}$ un opérateur β_k -moyenné, pour $\beta_k \in]0, 1[$,
- et $\mathbf{T}_k: \mathcal{H} \rightarrow \mathcal{H}: \mathbf{z} \mapsto (\mathbf{T}_{i,k}\mathbf{z})_{1 \leq i \leq m}$ un opérateur α_k -moyenné, pour $\alpha_k \in]0, 1[$, avec $(\forall i \in \{1, \dots, m\}) \mathbf{T}_{i,k}: \mathcal{H} \rightarrow \mathcal{H}_i$.

L'objectif est de trouver

$$\mathbf{z} \in \tilde{\mathbf{Z}} = \bigcap_{k \in \mathbb{N}} \text{Fix}(\mathbf{T}_k \circ \mathbf{S}_k). \quad (7.2)$$

Une méthode permettant de résoudre le problème (7.2) est donnée par l'algorithme 7.1.

Algorithme 7.1 [Combettes et Pesquet, 2015, Algo. (4.1)]

Initialisation : Soient $\mathbf{z}_0$, $(\mathbf{a}_k)_{k \in \mathbb{N}}$ et $(\mathbf{b}_k)_{k \in \mathbb{N}}$ des variables aléatoires à valeurs dans $\mathcal{K}$. Soient $(\varepsilon_k)_{k \in \mathbb{N}}$ des variables aléatoires identiquement distribuées à valeurs dans $\mathbb{D}_m$. Soit, pour tout $k \in \mathbb{N}$, $\lambda_k \in]0, +\infty[$.

Itérations :

Pour $k = 0, 1, \dots$

$$\left| \begin{array}{l} \mathbf{r}_k = \mathbf{S}_k \mathbf{z}_k + \mathbf{b}_k, \\ \text{pour } i = 1, \dots, m \\ \quad z_{k+1}^{(i)} = z_k^{(i)} + \lambda_k \varepsilon_k^{(i)} \left(\mathbf{T}_{i,k} \mathbf{r}_k + \mathbf{a}_k^{(i)} - z_k^{(i)} \right). \end{array} \right.$$

Le théorème suivant, issu de [Combettes et Pesquet, 2015, Thm. 4.1], donne les garanties de convergence de l'algorithme 7.1.

Théorème 7.2. [Combettes et Pesquet, 2015, Thm. 4.1] *Soit $(\mathbf{z}_k)_{k \in \mathbb{N}}$ la suite générée par l'algorithme 7.1. Supposons que $\tilde{\mathbf{Z}}$ est non vide, $\sup_{k \in \mathbb{N}} \alpha_k < 1$ et $\sup_{k \in \mathbb{N}} \beta_k < 1$. De plus, supposons que, pour tout $k \in \mathbb{N}$, $\lambda_k \in]0, 1]$ et que $\inf_{k \in \mathbb{N}} \lambda_k > 0$. Posons, pour tout $k \in \mathbb{N}$, $\mathcal{E}_k = \sigma(\varepsilon_k)$ et $\mathcal{Z}_k = \sigma(\mathbf{z}_0, \dots, \mathbf{z}_k)$. Si*

- (i) $\sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{a}_k\|^2 | \mathcal{Z}_k)} < +\infty$ et $\sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{b}_k\|^2 | \mathcal{Z}_k)} < +\infty$ P-presque sûrement,
- (ii) pour tout $k \in \mathbb{N}$, $\mathcal{E}_k$ et $\mathcal{Z}_k$ sont indépendants, et $(\forall i \in \{1, \dots, m\}) \mathbb{P}[\varepsilon_0^{(i)} = 1] > 0$,

on a alors P-presque sûrement

$$(\forall \mathbf{z} \in \tilde{\mathbf{Z}}) \quad \mathbf{T}_k(\mathbf{S}_k \mathbf{z}_k) - \mathbf{S}_k \mathbf{z}_k + \mathbf{S}_k \mathbf{z} \rightarrow \mathbf{z}, \quad (7.3)$$

$$(\forall \mathbf{z} \in \tilde{\mathbf{Z}}) \quad \mathbf{z}_k - \mathbf{S}_k \mathbf{z}_k + \mathbf{S}_k \mathbf{z} \rightarrow \mathbf{z}. \quad (7.4)$$

De plus, si

- (iii) l'ensemble des points de faible adhérence séquentielle de $(\mathbf{z}_k)_{k \in \mathbb{N}}$ est inclus dans $\tilde{\mathbf{Z}}$ P-presque sûrement,

alors, $(\mathbf{z}_k)_{k \in \mathbb{N}}$ converge faiblement P-presque sûrement vers une variable aléatoire à valeurs dans $\tilde{\mathbf{Z}}$.

Les algorithmes présentés dans ce chapitre sont issus des itérations de l'algorithme explicite-implicite pour la résolution de problèmes d'inclusions monotones [Combettes et Wajs, 2005] (voir [Attouch et al., 2010b] pour des exemples de problèmes pouvant être résolus par cette méthode). Une version alternée aléatoirement par bloc a récemment été proposée dans [Combettes et Pesquet, 2015, Sec. 5.2] (une version pour le cas variationnel a aussi été proposée dans [Necoara et Patrascu, 2014; Richtárik et Talác, 2014]). Nous allons montrer comment inclure des opérateurs de préconditionnement dans l'algorithme explicite-implicite alterné aléatoirement par bloc grâce à un changement de métrique.

Soient $\mathbf{Q}: \mathcal{K} \rightarrow 2^{\mathcal{K}}$ un opérateur maximal monotone et $\mathbf{R}: \mathcal{K} \rightarrow \mathcal{K}$ un opérateur cocoercif. On veut trouver

$$\mathbf{z} \in \mathbf{Z} = \text{zer}(\mathbf{Q} + \mathbf{R}).$$

Nous proposons d'utiliser l'algorithme 7.2 ci-dessous pour résoudre ce problème d'inclusion, où nous posons, pour tout $\mathbf{V} \in \mathcal{S}^+(\mathcal{K})$ et pour tout $k \in \mathbb{N}$, $\mathbf{J}_{\gamma_k} \mathbf{V} \mathbf{Q}: \mathbf{z} \mapsto (\mathbf{T}_{i,k} \mathbf{z})_{1 \leq i \leq m}$ avec $(\forall i \in \{1, \dots, m\}) \mathbf{T}_{i,k}: \mathcal{K} \rightarrow \mathcal{K}_i$.

Algorithme 7.2 Algorithme préconditionné explicite-implicite alterné aléatoirement par bloc

Initialisation : Soient $\mathbf{z}_0, (\mathbf{s}_k)_{k \in \mathbb{N}}$ et $(\mathbf{t}_k)_{k \in \mathbb{N}}$ des variables aléatoires à valeurs dans $\mathcal{K}$ et $\mathbf{V} \in \mathcal{S}^+(\mathcal{K})$. Soient $(\varepsilon_k)_{k \in \mathbb{N}}$ des variables aléatoires identiquement distribuées à valeurs dans $\mathbb{D}_m$. Soient, pour tout $k \in \mathbb{N}$, $\lambda_k \in]0, +\infty[$ et $\gamma_k \in]0, +\infty[$.

Itérations :

Pour $k = 0, 1, \dots$

$$\left| \begin{array}{l} \mathbf{r}_k = \mathbf{V} \mathbf{R} \mathbf{z}_k, \\ \text{pour } i = 1, \dots, m \\ \quad \left| \begin{array}{l} z_{k+1}^{(i)} = z_k^{(i)} + \lambda_k \varepsilon_k^{(i)} \left(\mathbf{T}_{i,k}(\mathbf{z}_k - \gamma_k \mathbf{r}_k + \mathbf{s}_k) + \mathbf{t}_k^{(i)} - z_k^{(i)} \right). \end{array} \right. \end{array} \right.$$

Les garanties de convergence de l'algorithme 7.2 sont données dans la proposition suivante, dont la démonstration se déduit du théorème 7.2 et de la démonstration de [Combettes et Pesquet, 2015, Prop. 5.9].

Proposition 7.1. Soit $(\mathbf{z}_k)_{k \in \mathbb{N}}$ la suite générée par l'algorithme 7.2. Supposons que $\mathbf{Z}$ est non vide, et que $\mathbf{V}^{1/2} \mathbf{R} \mathbf{V}^{1/2}$ est ϑ -cocoercif, avec $\vartheta \in]0, +\infty[$. Supposons, de plus, que $\inf_{k \in \mathbb{N}} \gamma_k > 0$, $\sup_{k \in \mathbb{N}} \gamma_k < 2\vartheta$, et que, pour tout $k \in \mathbb{N}$, $\lambda_k \in]0, 1]$ et $\inf_{k \in \mathbb{N}} \lambda_k > 0$. Posons $(\forall k \in \mathbb{N}) \mathcal{E}_k = \sigma(\varepsilon_k)$ et $\mathcal{Z}_k = \sigma(\mathbf{z}_0, \dots, \mathbf{z}_k)$. Si

- (i) $\sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{s}_k\|^2 | \mathcal{Z}_k)} < +\infty$ et $\sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{t}_k\|^2 | \mathcal{Z}_k)} < +\infty$ P-presque sûrement,
- (ii) pour tout $k \in \mathbb{N}$, $\mathcal{E}_k$ et $\mathcal{Z}_k$ sont indépendants et $(\forall i \in \{1, \dots, m\}) \mathbb{P}[\varepsilon_0^{(i)} = 1] > 0$,

alors $(\mathbf{z}_k)_{k \in \mathbb{N}}$ converge faiblement P-presque sûrement vers une variable aléatoire à valeurs dans $\mathbf{Z}$.

Démonstration. On a $\mathbf{Z} = \text{zer}(\mathbf{V} \mathbf{Q} + \mathbf{V} \mathbf{R}) \neq \emptyset$. Puisque $\mathbf{V} \in \mathcal{S}^+(\mathcal{K})$, on peut changer la métrique de la norme de $\mathcal{K}$:

$$(\forall \mathbf{z} \in \mathcal{K}) \quad \|\mathbf{z}\|_{\mathbf{V}^{-1}} = \sqrt{\langle \mathbf{z} | \mathbf{V}^{-1} \mathbf{z} \rangle}.$$

Notons de même $\langle \cdot | \cdot \rangle_{\mathbf{V}^{-1}}$ le produit scalaire associé. Dans cet espace renormé, $\mathbf{V} \mathbf{Q}$ est maximal

monotone. De plus,

$$\begin{aligned}
 (\forall (\mathbf{z}, \mathbf{z}') \in \mathcal{K}^2) \quad \|\mathbf{V}\mathbf{R}\mathbf{z} - \mathbf{V}\mathbf{R}\mathbf{z}'\|_{\mathbf{V}^{-1}}^2 &= \|\mathbf{V}^{1/2}\mathbf{R}\mathbf{z} - \mathbf{V}^{1/2}\mathbf{R}\mathbf{z}'\|^2 \\
 &\leq \vartheta^{-1} \langle \mathbf{V}^{-1/2}\mathbf{z} - \mathbf{V}^{-1/2}\mathbf{z}' \mid \mathbf{V}^{1/2}\mathbf{R}\mathbf{z} - \mathbf{V}^{1/2}\mathbf{R}\mathbf{z}' \rangle \\
 &= \vartheta^{-1} \langle \mathbf{z} - \mathbf{z}' \mid \mathbf{R}\mathbf{z} - \mathbf{R}\mathbf{z}' \rangle \\
 &= \vartheta^{-1} \langle \mathbf{z} - \mathbf{z}' \mid \mathbf{V}\mathbf{R}\mathbf{z} - \mathbf{V}\mathbf{R}\mathbf{z}' \rangle_{\mathbf{V}^{-1}},
 \end{aligned}$$

ce qui montre que $\mathbf{V}\mathbf{R}$ est ϑ -cocoercif dans $(\mathcal{K}, \|\cdot\|_{\mathbf{V}^{-1}})$. On peut alors utiliser un algorithme explicite-implicite pour trouver un élément de $\mathbf{Z}$ en itérant, pour tout $k \in \mathbb{N}$,

$$\mathbf{z}_{k+1} = \mathbf{J}_{\gamma_k \mathbf{V}\mathbf{Q}}(\mathbf{z}_k - \gamma_k \mathbf{V}\mathbf{R}\mathbf{z}_k).$$

Dans $(\mathcal{K}, \|\cdot\|_{\mathbf{V}^{-1}})$, l'opérateur $\mathbf{T}_k = \mathbf{J}_{\gamma_k \mathbf{V}\mathbf{Q}}$ est fermement contractant (et donc 1/2-moyenné), l'opérateur $\mathbf{S}_k = \mathbf{Id} - \gamma_k \mathbf{V}\mathbf{R}$ est $\gamma_k/(2\vartheta)$ -moyenné [Bauschke et Combettes, 2011, Prop. 4.33], et on a $\mathbf{Z} = \bigcap_{k \in \mathbb{N}} \text{Fix } \mathbf{T}_k \circ \mathbf{S}_k$ [Bauschke et Combettes, 2011, Prop. 25.1(iv)]. L'algorithme 7.2 apparaît alors comme un cas particulier de l'algorithme 7.1. Remarquons que la condition (i) induit

$$\begin{aligned}
 \sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{s}_k\|_{\mathbf{V}^{-1}}^2 \mid \mathcal{Z}_k)} &\leq \sqrt{\|\mathbf{V}^{-1}\|} \sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{s}_k\|^2 \mid \mathcal{Z}_k)} < +\infty \\
 \sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{t}_k\|_{\mathbf{V}^{-1}}^2 \mid \mathcal{Z}_k)} &\leq \sqrt{\|\mathbf{V}^{-1}\|} \sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{t}_k\|^2 \mid \mathcal{Z}_k)} < +\infty
 \end{aligned}$$

et que la convergence faible au sens de $\langle \cdot \mid \cdot \rangle$ et $\langle \cdot \mid \cdot \rangle_{\mathbf{V}^{-1}}$ sont équivalentes.

Afin d'appliquer le théorème 7.2, il reste à démontrer la condition (iii) de ce théorème. Pour cela, il suffit d'utiliser un raisonnement analogue à celui de la démonstration de [Combettes et Pesquet, 2015, Thm. 5.9]. Les résultats (7.3) et (7.4) assurent l'existence de $\tilde{\Omega} \in \mathcal{F}$ vérifiant $\mathbb{P}(\tilde{\Omega}) = 1$ et

$$(\forall \omega \in \tilde{\Omega})(\forall \mathbf{z} \in \tilde{\mathbf{Z}}) \quad \begin{cases} \mathbf{T}_k(\mathbf{S}_k \mathbf{z}_k(\omega)) - \mathbf{S}_k \mathbf{z}_k(\omega) + \mathbf{S}_k \mathbf{z} \rightarrow \mathbf{z} \\ \mathbf{S}_k \mathbf{z}_k(\omega) - \mathbf{z}_k(\omega) - \mathbf{S}_k \mathbf{z} \rightarrow -\mathbf{z}. \end{cases}$$

Par conséquent, puisque $\inf_{k \in \mathbb{N}} \gamma_k > 0$ et $\mathbf{V} \in \mathcal{S}^+(\mathcal{K})$,

$$(\forall \omega \in \tilde{\Omega})(\forall \mathbf{z} \in \tilde{\mathbf{Z}}) \quad \begin{cases} \mathbf{J}_{\gamma_k \mathbf{V}\mathbf{Q}}(\mathbf{z}_k(\omega) - \gamma_k \mathbf{V}\mathbf{R}\mathbf{z}_k(\omega)) - \mathbf{z}_k(\omega) = \mathbf{T}_k(\mathbf{S}_k \mathbf{z}_k(\omega)) - \mathbf{z}_k(\omega) \rightarrow \mathbf{0} \\ \mathbf{R}\mathbf{z}_k(\omega) \rightarrow \mathbf{R}\mathbf{z}. \end{cases} \quad (7.5)$$

Posons, de plus,

$$(\forall k \in \mathbb{N}) \quad \begin{cases} \mathbf{y}_k = \mathbf{J}_{\gamma_k \mathbf{V}\mathbf{Q}}(\mathbf{z}_k - \gamma_k \mathbf{V}\mathbf{R}\mathbf{z}_k) \\ \mathbf{u}_k = \gamma_k^{-1} \mathbf{V}^{-1}(\mathbf{z}_k - \mathbf{y}_k) - \mathbf{R}\mathbf{z}_k. \end{cases} \quad (7.6)$$

Nous pouvons alors déduire de (7.5) que

$$(\forall \omega \in \tilde{\Omega})(\forall \mathbf{z} \in \tilde{\mathbf{Z}}) \quad \begin{cases} \mathbf{z}_k(\omega) - \mathbf{y}_k(\omega) \rightarrow \mathbf{0} \\ \mathbf{u}_k(\omega) \rightarrow -\mathbf{R}\mathbf{z}. \end{cases} \quad (7.7)$$

Pour établir la condition (iii) il est alors suffisant de fixer $\mathbf{z} \in \mathbf{Z}$, $\hat{\mathbf{z}} \in \mathcal{K}$, de définir une sous-suite $(n_k)_{k \in \mathbb{N}}$ de $\mathbb{N}$ strictement croissante, de poser $\omega \in \tilde{\Omega}$ tel que $\mathbf{z}_{n_k}(\omega) \rightarrow \hat{\mathbf{z}}$ et de montrer que $\hat{\mathbf{z}} \in \mathbf{Z}$.

Nous déduisons de (7.5) que $\mathbf{R}\mathbf{z}_{n_k}(\omega) \rightarrow \mathbf{R}\mathbf{z}$. Ainsi, d'après [Bauschke et Combettes, 2011, Ex. 20.28] $\mathbf{R}$ étant maximal monotone, on peut déduire de [Bauschke et Combettes, 2011, Prop. 20.33(ii)] que $\mathbf{R}\mathbf{z} = \mathbf{R}\hat{\mathbf{z}}$. Nous pouvons de plus déduire de (7.7) que $\mathbf{y}_{n_k}(\omega) \rightarrow \hat{\mathbf{z}}$ et $\mathbf{u}_{n_k} \rightarrow -\mathbf{R}\hat{\mathbf{z}} = -\mathbf{R}\mathbf{z}$. L'équation (7.6) assure que $(\forall k \in \mathbb{N}) (\mathbf{y}_{n_k}(\omega), \mathbf{u}_{n_k}(\omega)) \in \text{Graph } \mathbf{Q}$. On déduit alors de [Bauschke et Combettes, 2011, Prop. 20.33(ii)] que $-\mathbf{R}\hat{\mathbf{z}} \in \mathbf{Q}\hat{\mathbf{z}}$, i.e. $\hat{\mathbf{z}} \in \mathbf{Z}$. $\square$

Remarque 7.1.

- (i) Soient $\tilde{\mathcal{K}}$ un espace de Hilbert réel séparable, $\mathbf{L} \in \mathcal{B}(\mathcal{K}, \tilde{\mathcal{K}})$, et $\tilde{\mathbf{R}}: \tilde{\mathcal{K}} \rightarrow \tilde{\mathcal{K}}$ un opérateur $\tilde{\vartheta}$ -cocoercif, avec $\tilde{\vartheta} \in]0, +\infty[$. Si $\mathbf{R} = \mathbf{L}^* \tilde{\mathbf{R}} \mathbf{L}$, alors $\mathbf{V}^{1/2} \mathbf{R} \mathbf{V}^{1/2}$ est ϑ -cocoercif pour tout $\mathbf{V} \in \mathcal{S}^+(\mathcal{K})$ vérifiant

$$\vartheta \|\mathbf{L} \mathbf{V} \mathbf{L}^*\| = \tilde{\vartheta}.$$

En effet, si $\tilde{\mathbf{R}}$ est $\tilde{\vartheta}$ -cocoercif alors pour tout $(\mathbf{x}, \mathbf{y}) \in \mathcal{K}^2$,

$$\begin{aligned} \langle \mathbf{x} - \mathbf{y} \mid \tilde{\mathbf{R}}\mathbf{x} - \tilde{\mathbf{R}}\mathbf{y} \rangle &\geq \tilde{\vartheta} \|\tilde{\mathbf{R}}\mathbf{x} - \tilde{\mathbf{R}}\mathbf{y}\|^2 \\ \Leftrightarrow \langle \mathbf{L} \mathbf{V}^{1/2} \mathbf{x} - \mathbf{L} \mathbf{V}^{1/2} \mathbf{y} \mid \tilde{\mathbf{R}} \mathbf{L} \mathbf{V}^{1/2} \mathbf{x} - \tilde{\mathbf{R}} \mathbf{L} \mathbf{V}^{1/2} \mathbf{y} \rangle &\geq \tilde{\vartheta} \|\tilde{\mathbf{R}} \mathbf{L} \mathbf{V}^{1/2} \mathbf{x} - \tilde{\mathbf{R}} \mathbf{L} \mathbf{V}^{1/2} \mathbf{y}\|^2. \end{aligned}$$

On obtient donc

$$\begin{aligned} \langle \mathbf{x} - \mathbf{y} \mid \mathbf{V}^{1/2} \mathbf{R} \mathbf{V}^{1/2} \mathbf{x} - \mathbf{V}^{1/2} \mathbf{R} \mathbf{V}^{1/2} \mathbf{y} \rangle \\ \geq \vartheta \|\mathbf{V}^{1/2} \mathbf{L}^*\|^2 \|\tilde{\mathbf{R}} \mathbf{L} \mathbf{V}^{1/2} \mathbf{x} - \tilde{\mathbf{R}} \mathbf{L} \mathbf{V}^{1/2} \mathbf{y}\|^2 \\ \geq \vartheta \|\mathbf{V}^{1/2} \mathbf{L}^* (\tilde{\mathbf{R}} \mathbf{L} \mathbf{V}^{1/2} \mathbf{x} - \tilde{\mathbf{R}} \mathbf{L} \mathbf{V}^{1/2} \mathbf{y})\|^2 = \vartheta \|\mathbf{V}^{1/2} \mathbf{R} \mathbf{V}^{1/2} \mathbf{x} - \mathbf{V}^{1/2} \mathbf{R} \mathbf{V}^{1/2} \mathbf{y}\|^2. \end{aligned}$$

- (ii) Soit $k \in \mathbb{N}$ une itération de l'algorithme 7.2. Les variables aléatoires $\mathbf{s}_k$ et $\mathbf{t}_k$ peuvent être vues comme des termes d'erreurs apparaissant dans les calculs respectifs de $\mathbf{R}$ et $\mathbf{J}_{\gamma_k} \mathbf{v}_k$. Notons que la possibilité de considérer des termes d'erreurs stochastiques aléatoires offre des degrés de liberté plus importants que l'hypothèse d'erreur déterministe sommable généralement faite dans la littérature (voir section 2.4.2). Remarquons cependant qu'il existe aussi des modèles prenant en compte des erreurs relatives [Monteiro et Svaiter, 2010; Solodov et Svaiter, 2001; Svaiter, 2014].
- (iii) Soit $k \in \mathbb{N}^*$. D'après l'algorithme 7.2, $\mathbf{e}_k$ et $\mathbf{z}_k$ sont indépendants si $\mathbf{e}_k$ est indépendant de $(\mathbf{z}_0, (\mathbf{e}_{k'}, \mathbf{s}_{k'}, \mathbf{t}_{k'})_{0 \leq k' < k})$.

7.3 ALGORITHMES PRIMAUX-DUAUX ALTERNÉS ALÉATOIREMENT PAR BLOC

Soient p et q des entiers strictement positifs, et $(\mathcal{H}_j)_{1 \leq j \leq p}$, $(\mathcal{G}_l)_{1 \leq l \leq q}$ des espaces de Hilbert réels séparables. De plus, on note respectivement $\mathcal{H} = \mathcal{H}_1 \oplus \dots \oplus \mathcal{H}_p$ et $\mathcal{G} = \mathcal{G}_1 \oplus \dots \oplus \mathcal{G}_q$ les

sommes directes de Hilbert de $(\mathcal{H}_j)_{1 \leq j \leq p}$ et $(\mathcal{G}_l)_{1 \leq l \leq q}$. Enfin, notons $\mathcal{K}$ l'espace produit $\mathcal{H} \oplus \mathcal{G}$.

7.3.1 Problème

Le problème suivant, mettant en œuvre des opérateurs monotones, a été beaucoup étudié ces dernières années (voir, par exemple dans [Boţ et Hendrich, 2013; Briceño-Arias, 2012; Combettes, 2013; Combettes et Pesquet, 2011; Pesquet et Pustelnik, 2012; Raguet *et al.*, 2013.]). Il jouera un rôle central tout au long de ce chapitre.

Problème 7.1. *Pour tout $j \in \{1, \dots, p\}$ et pour tout $l \in \{1, \dots, q\}$, soient $A_j: \mathcal{H}_j \rightarrow 2^{\mathcal{H}_j}$ un opérateur maximal monotone, $C_j: \mathcal{H}_j \rightarrow \mathcal{H}_j$ un opérateur cocoercif, $B_l: \mathcal{G}_l \rightarrow 2^{\mathcal{G}_l}$ un opérateur maximal monotone, $D_l: \mathcal{G}_l \rightarrow 2^{\mathcal{G}_l}$ un opérateur maximal et fortement monotone, et $L_{l,j} \in \mathcal{B}(\mathcal{H}_j, \mathcal{G}_l)$. Supposons que*

$$(\forall l \in \{1, \dots, q\}) \quad \mathbb{L}_l = \{j \in \{1, \dots, p\} \mid L_{l,j} \neq 0\} \neq \emptyset, \quad (7.8)$$

$$(\forall j \in \{1, \dots, p\}) \quad \mathbb{L}_j^* = \{l \in \{1, \dots, q\} \mid L_{l,j} \neq 0\} \neq \emptyset, \quad (7.9)$$

et que l'ensemble $\mathbf{F}$ de solutions du problème

trouver $\mathbf{x}^{(1)} \in \mathcal{H}_1, \dots, \mathbf{x}^{(p)} \in \mathcal{H}_p$ tels que

$$(\forall j \in \{1, \dots, p\}) \quad 0 \in A_j \mathbf{x}^{(j)} + C_j \mathbf{x}^{(j)} + \sum_{l=1}^q L_{l,j}^* (B_l \square D_l) \left(\sum_{j'=1}^p L_{l,j'} \mathbf{x}^{(j')} \right) \quad (7.10)$$

est non vide. On note, de plus, $\mathbf{F}^*$ l'ensemble de solutions du problème dual :

trouver $\mathbf{v}^{(1)} \in \mathcal{G}_1, \dots, \mathbf{v}^{(q)} \in \mathcal{G}_q$ tels que

$$(\forall l \in \{1, \dots, q\}) \quad 0 \in - \sum_{j=1}^p L_{l,j} (A_j^{-1} \square C_j^{-1}) \left(- \sum_{l'=1}^q L_{l',j}^* \mathbf{v}^{(l')} \right) + B_l^{-1} \mathbf{v}^{(l)} + D_l^{-1} \mathbf{v}^{(l)}. \quad (7.11)$$

L'objectif est de trouver un couple $(\hat{\mathbf{x}}, \hat{\mathbf{v}})$ de variables aléatoires à valeurs dans $\mathbf{F} \times \mathbf{F}^*$.

Le problème 7.1 peut être réécrit comme la recherche du zéro d'une somme de deux opérateurs maximaux monotones dans l'espace produit $\mathcal{K}$ comme indiqué ci-dessous.

Proposition 7.2. [Vũ, 2013] *Posons*

$$\begin{aligned} \mathbf{A}: \mathcal{H} \rightarrow 2^{\mathcal{H}}: \mathbf{x} &\mapsto \bigtimes_{j=1}^p A_j \mathbf{x}^{(j)} & \mathbf{B}: \mathcal{G} \rightarrow 2^{\mathcal{G}}: \mathbf{v} &\mapsto \bigtimes_{l=1}^q B_l \mathbf{v}^{(l)} \\ \mathbf{C}: \mathcal{H} \rightarrow \mathcal{H}: \mathbf{x} &\mapsto (C_j \mathbf{x}^{(j)})_{1 \leq j \leq p} & \mathbf{D}: \mathcal{G} \rightarrow 2^{\mathcal{G}}: \mathbf{v} &\mapsto \bigtimes_{l=1}^q D_l \mathbf{v}^{(l)} \\ \mathbf{L}: \mathcal{H} \rightarrow \mathcal{G}: \mathbf{x} &\mapsto \left(\sum_{j=1}^p L_{l,j} \mathbf{x}^{(j)} \right)_{1 \leq l \leq q} \end{aligned}$$

et introduisons les opérateurs

$$\begin{aligned} \mathbf{Q}: \quad \mathcal{K} &\rightarrow 2^{\mathcal{K}} \\ (\mathbf{x}, \mathbf{v}) &\mapsto (\mathbf{Ax} + \mathbf{L}^* \mathbf{v}) \times (-\mathbf{Lx} + \mathbf{B}^{-1} \mathbf{v}) \end{aligned} \quad (7.12)$$

et

$$\begin{aligned} \mathbf{R}: \quad \mathcal{K} &\rightarrow \mathcal{K} \\ (\mathbf{x}, \mathbf{v}) &\mapsto (\mathbf{Cx}, \mathbf{D}^{-1} \mathbf{v}). \end{aligned} \quad (7.13)$$

Les assertions suivantes sont alors vérifiées

- (i) $\mathbf{Q}$ est un opérateur maximal monotone et $\mathbf{R}$ est un opérateur cocoercif.
- (ii) $\mathbf{Z} = \text{zer}(\mathbf{Q} + \mathbf{R})$ est non vide.
- (iii) Un couple $(\hat{\mathbf{x}}, \hat{\mathbf{v}})$ de variables aléatoires est une solution du problème 7.1 si et seulement si $(\hat{\mathbf{x}}, \hat{\mathbf{v}})$ est à valeurs dans $\mathbf{Z}$.

D'après la propriété ci-dessus, nous pouvons donc utiliser l'algorithme 7.2 pour résoudre le problème 7.1. De plus, plusieurs algorithmes peuvent en être déduits selon le choix fait pour les opérateurs de préconditionnement.

Dans la suite, $\mathbf{L} \in \mathcal{B}(\mathcal{H}, \mathcal{G})$ sera défini comme dans la proposition 7.2.

7.3.2 Première sous-classe d'algorithmes

Nous donnons dans un premier temps deux résultats préliminaires qui nous seront utiles pour développer les algorithmes proposés dans cette section.

Lemme 7.1. Soient $\mathbf{W} \in \mathcal{S}^+(\mathcal{H})$ et $\mathbf{U} \in \mathcal{S}^+(\mathcal{G})$ tels que $\|\mathbf{U}^{1/2} \mathbf{LW}^{1/2}\| < 1$.

- (i) L'opérateur défini par

$$\begin{aligned} \mathbf{V}': \quad \mathcal{K} &\rightarrow \mathcal{K} \\ (\mathbf{x}, \mathbf{v}) &\mapsto (\mathbf{W}^{-1} \mathbf{x} - \mathbf{L}^* \mathbf{v}, -\mathbf{Lx} + \mathbf{U}^{-1} \mathbf{v}) \end{aligned} \quad (7.14)$$

appartient à $\mathcal{S}^+(\mathcal{K})$. Son inverse donné par

$$\begin{aligned} \mathbf{V}: \quad \mathcal{K} &\rightarrow \mathcal{K} \\ (\mathbf{x}, \mathbf{v}) &\mapsto ((\mathbf{W}^{-1} - \mathbf{L}^* \mathbf{U} \mathbf{L})^{-1} \mathbf{x} + \mathbf{W} \mathbf{L}^* (\mathbf{U}^{-1} - \mathbf{LW} \mathbf{L}^*)^{-1} \mathbf{v}, \\ &\quad (\mathbf{U}^{-1} - \mathbf{LW} \mathbf{L}^*)^{-1} (\mathbf{LW} \mathbf{x} + \mathbf{v})) \end{aligned} \quad (7.15)$$

est aussi un opérateur de $\mathcal{S}^+(\mathcal{K})$.

- (ii) Soient $\mathbf{C}: \mathcal{H} \rightarrow \mathcal{H}$, $\mathbf{D}: \mathcal{G} \rightarrow 2^{\mathcal{G}}$, et $\mathbf{R}: \mathcal{K} \rightarrow \mathcal{K}$ les opérateurs définis dans la proposition 7.2. Si $\mathbf{W}^{1/2}\mathbf{C}\mathbf{W}^{1/2}$ est μ -cocoercif, avec $\mu \in]0, +\infty[$, et $\mathbf{U}^{1/2}\mathbf{D}^{-1}\mathbf{U}^{1/2}$ est ν -cocoercif, avec $\nu \in]0, +\infty[$, alors, pour tout $\alpha \in]0, +\infty[$, $\mathbf{V}^{1/2}\mathbf{R}\mathbf{V}^{1/2}$ est ϑ_α -cocoercif, avec

$$\vartheta_\alpha = (1 - \|\mathbf{U}^{1/2}\mathbf{L}\mathbf{W}^{1/2}\|^2) \min \{ \mu(1 + \alpha\|\mathbf{U}^{1/2}\mathbf{L}\mathbf{W}^{1/2}\|)^{-1}, \nu(1 + \alpha^{-1}\|\mathbf{U}^{1/2}\mathbf{L}\mathbf{W}^{1/2}\|)^{-1} \}. \quad (7.16)$$

La démonstration du lemme 7.1 est donnée dans l'annexe 7.A.

Remarque 7.2.

- (i) Dans (7.16), un choix simple est de prendre $\alpha = 1$, ce qui conduit à la constante de cocoercivité

$$\vartheta_1 = (1 - \|\mathbf{U}^{1/2}\mathbf{L}\mathbf{W}^{1/2}\|) \min\{\mu, \nu\}.$$

Un autre choix possible de cette constante est $\vartheta_{\hat{\alpha}}$ où $\hat{\alpha}$ est le maximiseur de $\alpha \mapsto \vartheta_\alpha$ sur $]0, +\infty[$. On obtient alors

$$\hat{\alpha} = \frac{\mu - \nu + \sqrt{(\mu - \nu)^2 + 4\mu\nu\|\mathbf{U}^{1/2}\mathbf{L}\mathbf{W}^{1/2}\|^2}}{2\nu\|\mathbf{U}^{1/2}\mathbf{L}\mathbf{W}^{1/2}\|}.$$

- (ii) Lorsque $\mathbf{D}^{-1} = \mathbf{0}$, la constante positive ν peut être choisie arbitrairement grande. Une constante de cocoercivité de $\mathbf{V}^{1/2}\mathbf{R}\mathbf{V}^{1/2}$ est alors donnée par

$$\lim_{\substack{\alpha \rightarrow 0 \\ \alpha > 0}} \vartheta_\alpha = (1 - \|\mathbf{U}^{1/2}\mathbf{L}\mathbf{W}^{1/2}\|^2)\mu.$$

Lemme 7.2. Soient $\mathbf{A}: \mathcal{H} \rightarrow 2^{\mathcal{H}}$, $\mathbf{B}: \mathcal{G} \rightarrow 2^{\mathcal{G}}$, $\mathbf{C}: \mathcal{H} \rightarrow \mathcal{H}$, $\mathbf{D}: \mathcal{G} \rightarrow 2^{\mathcal{G}}$, $\mathbf{Q}: \mathcal{K} \rightarrow 2^{\mathcal{K}}$, et $\mathbf{R}: \mathcal{K} \rightarrow \mathcal{K}$. Soient $\mathbf{W} \in \mathcal{S}^+(\mathcal{H})$ et $\mathbf{U} \in \mathcal{S}^+(\mathcal{G})$ tels que $\|\mathbf{U}^{1/2}\mathbf{L}\mathbf{W}^{1/2}\| < 1$. Soit $\mathbf{V} \in \mathcal{S}^+(\mathcal{K})$ défini par (7.15). Pour tout $\mathbf{z} = (\mathbf{x}, \mathbf{v}) \in \mathcal{K}$ et $(\mathbf{c}, \mathbf{e}) \in \mathcal{K}$, on pose

$$\begin{cases} \mathbf{y} = \mathbf{J}_{\mathbf{W}\mathbf{A}}(\mathbf{x} - \mathbf{W}(\mathbf{L}^*\mathbf{v} + \mathbf{C}\mathbf{x} + \mathbf{c})) \\ \mathbf{u} = \mathbf{J}_{\mathbf{U}\mathbf{B}^{-1}}(\mathbf{v} + \mathbf{U}(\mathbf{L}(2\mathbf{y} - \mathbf{x}) - \mathbf{D}^{-1}\mathbf{v} + \mathbf{e})). \end{cases} \quad (7.17)$$

On a alors, $(\mathbf{y}, \mathbf{u}) = \mathbf{J}_{\mathbf{V}\mathbf{Q}}(\mathbf{z} - \mathbf{V}\mathbf{R}\mathbf{z} + \mathbf{s})$ où

$$\mathbf{s} = ((\mathbf{W}^{-1} - \mathbf{L}^*\mathbf{U}\mathbf{L})^{-1}(\mathbf{L}^*\mathbf{U}\mathbf{e} - \mathbf{c}), (\mathbf{U}^{-1} - \mathbf{L}\mathbf{W}\mathbf{L}^*)^{-1}(\mathbf{e} - \mathbf{L}\mathbf{W}\mathbf{c})).$$

Démonstration. Soient $\mathbf{z} = (\mathbf{x}, \mathbf{v}) \in \mathcal{K}$ et $\mathbf{s} = (\mathbf{c}', \mathbf{e}') \in \mathcal{K}$. On a les équivalences suivantes :

$$\begin{aligned}
 & (\mathbf{y}, \mathbf{u}) = \mathbf{J}_{\mathbf{VQ}}(\mathbf{z} - \mathbf{VRz} + \mathbf{s}) \\
 \Leftrightarrow & \mathbf{z} - \mathbf{VRz} + \mathbf{s} \in (\mathbf{Id} + \mathbf{VQ})(\mathbf{y}, \mathbf{u}) \\
 \Leftrightarrow & \mathbf{V}^{-1}(\mathbf{z} + \mathbf{s} - (\mathbf{y}, \mathbf{u})) - \mathbf{Rz} \in \mathbf{Q}(\mathbf{y}, \mathbf{u}) \\
 \Leftrightarrow & \begin{cases} \mathbf{W}^{-1}(\mathbf{x} - \mathbf{y} + \mathbf{c}') - \mathbf{L}^*(\mathbf{v} + \mathbf{e}') - \mathbf{Cx} \in \mathbf{Ay} \\ \mathbf{U}^{-1}(\mathbf{v} - \mathbf{u} + \mathbf{e}') + \mathbf{L}(2\mathbf{y} - \mathbf{x} - \mathbf{c}') - \mathbf{D}^{-1}\mathbf{v} \in \mathbf{B}^{-1}\mathbf{u} \end{cases} \quad (7.18) \\
 \Leftrightarrow & \begin{cases} \mathbf{x} + \mathbf{c}' - \mathbf{W}(\mathbf{L}^*(\mathbf{v} + \mathbf{e}') + \mathbf{Cx}) \in (\mathbf{Id} + \mathbf{WA})\mathbf{y} \\ \mathbf{v} + \mathbf{e}' + \mathbf{U}(\mathbf{L}(2\mathbf{y} - \mathbf{x} - \mathbf{c}') - \mathbf{D}^{-1}\mathbf{v}) \in (\mathbf{Id} + \mathbf{UB}^{-1})\mathbf{u} \end{cases} \\
 \Leftrightarrow & \begin{cases} \mathbf{y} = \mathbf{J}_{\mathbf{WA}}(\mathbf{x} + \mathbf{c}' - \mathbf{W}(\mathbf{L}^*(\mathbf{v} + \mathbf{e}') + \mathbf{Cx})) \\ \mathbf{u} = \mathbf{J}_{\mathbf{UB}^{-1}}(\mathbf{v} + \mathbf{e}' + \mathbf{U}(\mathbf{L}(2\mathbf{y} - \mathbf{x} - \mathbf{c}') - \mathbf{D}^{-1}\mathbf{v})), \end{cases}
 \end{aligned}$$

où, pour obtenir (7.18), on a utilisé l'expression de $\mathbf{Q}$ donnée par (7.12) et l'expression de $\mathbf{V}$ donnée par (7.14).

Afin de conclure, remarquons que, puisque $\|\mathbf{U}^{1/2}\mathbf{LW}^{1/2}\| < 1$, nous avons montré dans la démonstration du lemme 7.1(i) que $\mathbf{W}^{-1} - \mathbf{L}^*\mathbf{UL}$ et $\mathbf{U}^{-1} - \mathbf{LW}\mathbf{L}^*$ sont des isomorphismes (en conséquence de (7.73) et (7.74)). Ainsi, pour tout $(\mathbf{c}, \mathbf{e}) \in \mathcal{K}$,

$$\begin{cases} \mathbf{Wc} = \mathbf{WL}^*\mathbf{e}' - \mathbf{c}' \\ \mathbf{Ue} = \mathbf{e}' - \mathbf{ULc}' \end{cases} \Leftrightarrow \begin{cases} \mathbf{c}' = (\mathbf{W}^{-1} - \mathbf{L}^*\mathbf{UL})^{-1}(\mathbf{L}^*\mathbf{Ue} - \mathbf{c}) \\ \mathbf{e}' = (\mathbf{U}^{-1} - \mathbf{LW}\mathbf{L}^*)^{-1}(\mathbf{e} - \mathbf{LWc}). \end{cases}$$

□

Nous pouvons maintenant donner un premier algorithme primal-dual alterné aléatoirement par bloc pour résoudre le problème 7.1. Ses itérations sont décrites dans l'algorithme 7.3.

Algorithme 7.3 Algorithme primal-dual alterné aléatoirement par bloc (version 1)

Initialisation : Soient $\mathbf{x}_0, (\mathbf{a}_k)_{k \in \mathbb{N}}$ et $(\mathbf{c}_k)_{k \in \mathbb{N}}$ des variables aléatoires à valeurs dans $\mathcal{H}$, et soient $\mathbf{v}_0, (\mathbf{b}_k)_{k \in \mathbb{N}}$, et $(\mathbf{d}_k)_{k \in \mathbb{N}}$ des variables aléatoires à valeurs dans $\mathcal{G}$. Pour tout $j \in \{1, \dots, p\}$, soit $\mathbf{W}_j \in \mathcal{S}^+(\mathcal{H}_j)$, et, pour tout $l \in \{1, \dots, q\}$, soit $\mathbf{U}_l \in \mathcal{S}^+(\mathcal{G}_l)$. Soient $(\varepsilon_k)_{k \in \mathbb{N}}$ des variables aléatoires identiquement distribuées à valeurs dans $\mathbb{D}_{p+q}$. Soit, pour tout $k \in \mathbb{N}$, $\lambda_k \in]0, +\infty[$.

Itérations :

Pour $k = 0, 1, \dots$

$$\left[\begin{array}{l} \text{pour } j = 1, \dots, p \\ \quad \left[\begin{array}{l} y_k^{(j)} = \varepsilon_k^{(j)} \left(\mathbf{J}_{\mathbf{W}_j \mathbf{A}_j} \left(x_k^{(j)} - \mathbf{W}_j \left(\sum_{l \in \mathbb{L}_j^*} \mathbf{L}_{l,j}^* v_k^{(l)} + \mathbf{C}_j x_k^{(j)} + c_k^{(j)} \right) \right) + a_k^{(j)} \right), \\ x_{k+1}^{(j)} = x_k^{(j)} + \lambda_k \varepsilon_k^{(j)} (y_k^{(j)} - x_k^{(j)}), \end{array} \right. \\ \text{pour } l = 1, \dots, q \\ \quad \left[\begin{array}{l} u_k^{(l)} = \varepsilon_k^{(p+l)} \left(\mathbf{J}_{\mathbf{U}_l \mathbf{B}_l^{-1}} \left(v_k^{(l)} + \mathbf{U}_l \left(\sum_{j \in \mathbb{L}_l} \mathbf{L}_{l,j} (2y_k^{(j)} - x_k^{(j)}) - \mathbf{D}_l^{-1} v_k^{(l)} + d_k^{(l)} \right) \right) + b_k^{(l)} \right), \\ v_{k+1}^{(l)} = v_k^{(l)} + \lambda_k \varepsilon_k^{(p+l)} (u_k^{(l)} - v_k^{(l)}). \end{array} \right. \end{array} \right.$$

En utilisant les lemmes 7.1 et 7.2, nous obtenons la proposition suivante établissant la convergence de l'algorithme 7.3.

Proposition 7.3. *Considérons l'algorithme 7.3. Supposons que, pour tout $j \in \{1, \dots, p\}$, $\mathbf{W}_j^{1/2} \mathbf{C}_j \mathbf{W}_j^{1/2}$ est μ_j -cocoercif, avec $\mu_j \in]0, +\infty[$, et, pour tout $l \in \{1, \dots, q\}$, $\mathbf{U}_l^{1/2} \mathbf{D}_l^{-1} \mathbf{U}_l^{1/2}$ est ν_l -cocoercif, avec $\nu_l \in]0, +\infty[$. Supposons que*

$$(\exists \alpha \in]0, +\infty[) \quad 2\vartheta_\alpha > 1 \quad (7.19)$$

où ϑ_α est défini par (7.16) avec $\mu = \min\{\mu_1, \dots, \mu_p\}$, $\nu = \min\{\nu_1, \dots, \nu_q\}$ et

$$\mathbf{W}: \mathcal{H} \rightarrow \mathcal{H}: \mathbf{x} \mapsto (\mathbf{W}_1 \mathbf{x}^{(1)}, \dots, \mathbf{W}_p \mathbf{x}^{(p)}) \quad \text{et} \quad \mathbf{U}: \mathcal{G} \rightarrow \mathcal{G}: \mathbf{v} \mapsto (\mathbf{U}_1 \mathbf{v}^{(1)}, \dots, \mathbf{U}_q \mathbf{v}^{(q)}). \quad (7.20)$$

Supposons que les éléments de la suite $(\lambda_k)_{k \in \mathbb{N}}$ sont dans $]0, 1]$ et que $\inf_{k \in \mathbb{N}} \lambda_k > 0$. Posons $(\forall k \in \mathbb{N}) \mathcal{E}_k = \sigma(\varepsilon_k)$ et $\mathcal{X}_k = \sigma(\mathbf{x}_{k'}, \mathbf{v}_{k'})_{0 \leq k' \leq k}$, et supposons de plus que les assertions suivantes sont vérifiées :

- (i) $\sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{a}_k\|^2 | \mathcal{X}_k)} < +\infty$, $\sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{b}_k\|^2 | \mathcal{X}_k)} < +\infty$, $\sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{c}_k\|^2 | \mathcal{X}_k)} < +\infty$,
et $\sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{d}_k\|^2 | \mathcal{X}_k)} < +\infty$ P-presque sûrement,
- (ii) pour tout $k \in \mathbb{N}$, $\mathcal{E}_k$ et $\mathcal{X}_k$ sont indépendants, et $(\forall l \in \{1, \dots, q\}) \mathbb{P}[\varepsilon_0^{(p+l)} = 1] > 0$,
- (iii) pour tout $j \in \{1, \dots, p\}$ et $k \in \mathbb{N}$, $\bigcup_{l \in \mathbb{L}_j^*} \{\omega \in \Omega \mid \varepsilon_k^{(p+l)}(\omega) = 1\} \subset \{\omega \in \Omega \mid \varepsilon_k^{(j)}(\omega) = 1\}$.

La suite $(\mathbf{x}_k)_{k \in \mathbb{N}}$ converge alors faiblement P-presque sûrement vers une variable aléatoire à valeurs dans $\mathbf{F}$, et $(\mathbf{v}_k)_{k \in \mathbb{N}}$ converge faiblement P-presque sûrement vers une variable aléatoire à valeurs dans $\mathbf{F}^*$.

La démonstration de la proposition 7.3 est donnée dans l'annexe 7.B. Un certain nombre d'observations peuvent être faites sur cette proposition.

Remarque 7.3.

- (i) Les variables aléatoires booléennes $(\varepsilon_k^{(i)})_{1 \leq i \leq p+q}$ indiquent les variables $(x_k^{(j)})_{1 \leq j \leq p}$ et $(v_k^{(l)})_{1 \leq l \leq q}$ qui sont activées à chaque itération k de l'algorithme 7.3. D'un point de vue algorithmique, lorsque certaines de ces variables sont nulles, les variables associées n'ont pas besoin d'être mises à jour. Remarquons que, en accord avec la condition (iii), pour tout $j \in \{1, \dots, p\}$, $y_k^{(j)}$ n'est pas seulement mis à jour lorsque $x_k^{(j)}$ est activée, mais aussi lorsqu'il existe $l \in \{1, \dots, q\}$ tel que $v_k^{(l)}$ est activé et $L_{l,j} \neq 0$.
- (ii) Pour tout $n \in \mathbb{N}$, $j \in \{1, \dots, p\}$, et $l \in \{1, \dots, q\}$, $a_k^{(j)}$, $b_k^{(l)}$, $c_k^{(j)}$, et $d_k^{(l)}$ sont des erreurs stochastiques pouvant apparaître, à l'itération k de l'algorithme 7.3, lors des calculs respectifs de $J_{W_j A_j}$, $J_{U_l B_l^{-1}}$, C_j , et D_l^{-1} .
- (iii) En utilisant l'inégalité triangulaire et l'inégalité de Cauchy-Schwarz, on obtient

$$\begin{aligned}
 (\mathbf{x} \in \mathcal{H}) \quad \|\mathbf{U}^{1/2} \mathbf{L} \mathbf{W}^{1/2} \mathbf{x}\|^2 &= \sum_{l=1}^q \left\| \sum_{j=1}^p U_l^{1/2} L_{l,j} W_j^{1/2} \mathbf{x}^{(j)} \right\|^2 \\
 &\leq \sum_{l=1}^q \left(\sum_{j=1}^p \|U_l^{1/2} L_{l,j} W_j^{1/2}\| \|\mathbf{x}^{(j)}\| \right)^2 \\
 &\leq \sum_{l=1}^q \left(\sum_{j=1}^p \|U_l^{1/2} L_{l,j} W_j^{1/2}\|^2 \right) \left(\sum_{j=1}^p \|\mathbf{x}^{(j)}\|^2 \right),
 \end{aligned}$$

ce qui implique que

$$\|\mathbf{U}^{1/2} \mathbf{L} \mathbf{W}^{1/2}\| \leq \left(\sum_{j=1}^p \sum_{l=1}^q \|U_l^{1/2} L_{l,j} W_j^{1/2}\|^2 \right)^{1/2}. \quad (7.21)$$

Pour tout $j \in \{1, \dots, p\}$, soit $\tilde{\mu}_j \in]0, +\infty[$ une constante de cocoercivité de $\mathbb{C}_j$ et, pour tout $l \in \{1, \dots, q\}$, soit $\tilde{\nu}_l \in]0, +\infty[$ une constante de forte monotonie de $\mathbb{D}_l$. On peut alors choisir

$$\mu = \min\{(\|\mathbb{W}_j\|^{-1}\tilde{\mu}_j)_{1 \leq j \leq p}\}, \quad \nu = \min\{(\|\mathbb{U}_l\|^{-1}\tilde{\nu}_l)_{1 \leq l \leq q}\}. \quad (7.22)$$

Par conséquent, en utilisant la remarque 7.2(i), une condition suffisante pour que (7.19) soit satisfaite avec $\alpha = 1$ est

$$\left(1 - \left(\sum_{j=1}^p \sum_{l=1}^q \|\mathbb{U}_l^{1/2} \mathbb{L}_{l,j} \mathbb{W}_j^{1/2}\|^2\right)^{1/2}\right) \min\{(\|\mathbb{W}_j\|^{-1}\tilde{\mu}_j)_{1 \leq j \leq p}, (\|\mathbb{U}_l\|^{-1}\tilde{\nu}_l)_{1 \leq l \leq q}\} > \frac{1}{2}. \quad (7.23)$$

Lorsque $\mathbf{D}^{-1} = \mathbf{0}$, en accord avec la remarque 7.2(ii), cette condition peut être remplacée par

$$\left(1 - \sum_{j=1}^p \sum_{l=1}^q \|\mathbb{U}_l^{1/2} \mathbb{L}_{l,j} \mathbb{W}_j^{1/2}\|^2\right) \min\{(\|\mathbb{W}_j\|^{-1}\tilde{\mu}_j)_{1 \leq j \leq p}\} > \frac{1}{2}. \quad (7.24)$$

- (iv) L'algorithme 7.3 généralise plusieurs résultats existants dans le cas déterministe, lorsque $p = 1$ et que la mise à jour des variables duales n'est pas aléatoire. Dans la plupart de ces travaux, $\mathbb{W}_1 = \tau \text{Id}$ et, pour tout $l \in \{1, \dots, q\}$, $\mathbb{U}_l = \rho_l \text{Id}$ avec $(\tau, \rho_1, \dots, \rho_q) \in]0, +\infty[^{q+1}$. En particulier, dans [Vũ, 2013], une condition suffisante pour que (7.23) soit satisfaite est utilisée, tandis que dans [Condat, 2013] une condition similaire à (7.24) est utilisée dans un cadre variationnel qui correspondrait au cas où $\mathbf{D}^{-1} = \mathbf{0}$. L'algorithme proposé généralise aussi les résultats donnés dans [Combettes et Vũ, 2014, Sec. 6] lorsqu'une métrique constante est considérée.

Les problèmes 7.10 et 7.11 étant symétriques, les rôles des variables primales et des variables duales des deux problèmes peuvent être échangées. Cela nous conduit à une forme symétrique de l'algorithme 7.3, donnée dans l'algorithme 7.4.

Algorithme 7.4 Algorithme primal-dual alterné aléatoirement par bloc (version 1 symétrique)

Initialisation : Soient $\mathbf{x}_0, (\mathbf{a}_k)_{k \in \mathbb{N}}$ et $(\mathbf{c}_k)_{k \in \mathbb{N}}$ des variables aléatoires à valeurs dans $\mathcal{H}$, et soient $\mathbf{v}_0, (\mathbf{b}_k)_{k \in \mathbb{N}}$, et $(\mathbf{d}_k)_{k \in \mathbb{N}}$ des variables aléatoires à valeurs dans $\mathcal{G}$. Pour tout $j \in \{1, \dots, p\}$, soit $\mathbf{W}_j \in \mathcal{S}^+(\mathcal{H}_j)$, et, pour tout $l \in \{1, \dots, q\}$, soit $\mathbf{U}_l \in \mathcal{S}^+(\mathcal{G}_l)$. Soient $(\varepsilon_k)_{k \in \mathbb{N}}$ des variables aléatoires identiquement distribuées à valeurs dans $\mathbb{D}_{p+q}$. Soit, pour tout $k \in \mathbb{N}$, $\lambda_k \in]0, +\infty[$.

Itérations :

Pour $k = 0, 1, \dots$

$$\left[\begin{array}{l} \text{pour } l = 1, \dots, q \\ \quad \left[\begin{array}{l} u_k^{(l)} = \varepsilon_k^{(p+l)} \left(\mathbf{J}_{\mathbf{U}_l \mathbf{B}_l^{-1}} \left(v_k^{(l)} + \mathbf{U}_l \left(\sum_{j \in \mathbb{L}_l} \mathbf{L}_{l,j} x_k^{(j)} - \mathbf{D}_l^{-1} v_k^{(l)} + d_k^{(l)} \right) \right) + b_k^{(l)} \right), \\ v_{k+1}^{(l)} = v_k^{(l)} + \lambda_k \varepsilon_k^{(p+l)} (u_k^{(l)} - v_k^{(l)}), \end{array} \right. \\ \text{pour } j = 1, \dots, p \\ \quad \left[\begin{array}{l} y_k^{(j)} = \varepsilon_k^{(j)} \left(\mathbf{J}_{\mathbf{W}_j \mathbf{A}_j} \left(x_k^{(j)} - \mathbf{W}_j \left(\sum_{l \in \mathbb{L}_j^*} \mathbf{L}_{l,j}^* (2u_k^{(l)} - v_k^{(l)}) + \mathbf{C}_j x_k^{(j)} + c_k^{(j)} \right) \right) + a_k^{(j)} \right), \\ x_{k+1}^{(j)} = x_k^{(j)} + \lambda_k \varepsilon_k^{(j)} (y_k^{(j)} - x_k^{(j)}). \end{array} \right. \end{array} \right.$$

Proposition 7.4. Considérons l'algorithme 7.4, $\mathbf{W}$ et $\mathbf{U}$ définis par (7.20) et (μ, ν) donnés dans la proposition 7.3. Supposons que (7.19) est vérifiée où ϑ_α est défini par (7.16). Supposons que les éléments de la suite $(\lambda_k)_{k \in \mathbb{N}}$ sont dans $]0, 1]$ et que $\inf_{k \in \mathbb{N}} \lambda_k > 0$. De plus, supposons que la condition (i) de la proposition 7.3 est vérifiée où $(\forall k \in \mathbb{N}) \mathcal{E}_k = \sigma(\varepsilon_k)$ et $\mathcal{X}_k = \sigma(\mathbf{x}_{k'}, \mathbf{v}_{k'})_{0 \leq k' \leq k}$, et que les assertions suivantes sont satisfaites :

- (ii) pour tout $k \in \mathbb{N}$, $\mathcal{E}_k$ et $\mathcal{X}_k$ sont indépendants, et $(\forall j \in \{1, \dots, p\}) \mathbb{P}[\varepsilon_0^{(j)} = 1] > 0$.
- (iii) pour tout $l \in \{1, \dots, q\}$ et $k \in \mathbb{N}$, $\bigcup_{j \in \mathbb{L}_l} \{\omega \in \Omega \mid \varepsilon_k^{(j)}(\omega) = 1\} \subset \{\omega \in \Omega \mid \varepsilon_k^{(p+l)}(\omega) = 1\}$.

La suite $(\mathbf{x}_k)_{k \in \mathbb{N}}$ converge alors faiblement $\mathbf{P}$ -presque sûrement vers une variable aléatoire à valeurs dans $\mathbf{F}$, et $(\mathbf{v}_k)_{k \in \mathbb{N}}$ converge faiblement $\mathbf{P}$ -presque sûrement vers une variable aléatoire à valeurs dans $\mathbf{F}^*$.

7.3.3 Seconde sous-classe d'algorithmes

Dans cette section, nous considérons que l'opérateur $\mathbf{V}$ a une forme diagonale, et nous procédons par une approche similaire à celle suivie dans la section 7.3.2.

Lemme 7.3. Soient $\mathbf{W} \in \mathcal{S}^+(\mathcal{H})$ et $\mathbf{U} \in \mathcal{S}^+(\mathcal{G})$ tels que $\|\mathbf{U}^{1/2}\mathbf{LW}^{1/2}\| < 1$.

(i) L'opérateur défini par

$$\begin{aligned} \mathbf{V}: \quad \mathcal{K} &\rightarrow \mathcal{K} \\ (\mathbf{x}, \mathbf{v}) &\mapsto (\mathbf{Wx}, (\mathbf{U}^{-1} - \mathbf{LWL}^*)^{-1}\mathbf{v}) \end{aligned} \quad (7.25)$$

est dans $\mathcal{S}^+(\mathcal{K})$.

(ii) Soient $\mathbf{C}: \mathcal{H} \rightarrow \mathcal{H}$, $\mathbf{D}: \mathcal{G} \rightarrow 2^{\mathcal{G}}$, et $\mathbf{R}: \mathcal{K} \rightarrow \mathcal{K}$ les opérateurs définis dans la proposition 7.2. Si $\mathbf{W}^{1/2}\mathbf{C}\mathbf{W}^{1/2}$ est μ -cocoercif, avec $\mu \in]0, +\infty[$, et $\mathbf{U}^{1/2}\mathbf{D}^{-1}\mathbf{U}^{1/2}$ est ν -cocoercif, avec $\nu \in]0, +\infty[$, alors $\mathbf{V}^{1/2}\mathbf{R}\mathbf{V}^{1/2}$ est ϑ -cocoercif, où

$$\vartheta = \min \{ \mu, \nu(1 - \|\mathbf{U}^{1/2}\mathbf{LW}^{1/2}\|^2) \}. \quad (7.26)$$

Démonstration.

(i) Tout d'abord, remarquons que, comme nous l'avons montré dans (7.74), si $\|\mathbf{U}^{1/2}\mathbf{LW}^{1/2}\| < 1$, alors $\mathbf{U}^{-1} - \mathbf{LWL}^*$ est un opérateur fortement positif de $\mathcal{B}(\mathcal{G})$ et est donc un isomorphisme. Ainsi, pour tout $(\mathbf{x}, \mathbf{v}) \in \mathcal{K}$,

$$\begin{aligned} \langle (\mathbf{x}, \mathbf{v}) \mid \mathbf{V}(\mathbf{x}, \mathbf{v}) \rangle &= \langle \mathbf{x} \mid \mathbf{Wx} \rangle + \left\langle \mathbf{v} \mid (\mathbf{U}^{-1} - \mathbf{LWL}^*)^{-1}\mathbf{v} \right\rangle \\ &\geq \|\mathbf{W}^{-1}\|^{-1}\|\mathbf{x}\|^2 + \|\mathbf{U}^{-1} - \mathbf{LWL}^*\|^{-1}\|\mathbf{v}\|^2 \\ &\geq \min \{ \|\mathbf{W}^{-1}\|^{-1}, \|\mathbf{U}^{-1} - \mathbf{LWL}^*\|^{-1} \} (\|\mathbf{x}\|^2 + \|\mathbf{v}\|^2). \end{aligned}$$

Donc, $\mathbf{V} \in \mathcal{S}^+(\mathcal{K})$.

(ii) Soient $\mathbf{z} = (\mathbf{x}, \mathbf{v}) \in \mathcal{K}$ et $\mathbf{z}' = (\mathbf{x}', \mathbf{v}') \in \mathcal{K}$. On a

$$\begin{aligned} \|\mathbf{Rz} - \mathbf{Rz}'\|_{\mathbf{V}}^2 &= \langle \mathbf{Cx} - \mathbf{Cx}' \mid \mathbf{W}(\mathbf{Cx} - \mathbf{Cx}') \rangle + \left\langle \mathbf{D}^{-1}\mathbf{v} - \mathbf{D}^{-1}\mathbf{v}' \mid (\mathbf{U}^{-1} - \mathbf{LWL}^*)^{-1}(\mathbf{D}^{-1}\mathbf{v} - \mathbf{D}^{-1}\mathbf{v}') \right\rangle \\ &\leq \|\mathbf{Cx} - \mathbf{Cx}'\|_{\mathbf{W}}^2 + \|(\mathbf{Id} - \mathbf{U}^{1/2}\mathbf{LWL}^*\mathbf{U}^{1/2})^{-1}\| \|\mathbf{D}^{-1}\mathbf{v} - \mathbf{D}^{-1}\mathbf{v}'\|_{\mathbf{U}}^2 \\ &\leq \mu^{-1}\langle \mathbf{x} - \mathbf{x}' \mid \mathbf{Cx} - \mathbf{Cx}' \rangle + \nu^{-1}(1 - \|\mathbf{U}^{1/2}\mathbf{LW}^{1/2}\|^2)^{-1} \langle \mathbf{v} - \mathbf{v}' \mid \mathbf{D}^{-1}\mathbf{v} - \mathbf{D}^{-1}\mathbf{v}' \rangle \\ &\leq \max \{ \mu^{-1}, \nu^{-1}(1 - \|\mathbf{U}^{1/2}\mathbf{LW}^{1/2}\|^2)^{-1} \} \langle \mathbf{z} - \mathbf{z}' \mid \mathbf{Rz} - \mathbf{Rz}' \rangle, \end{aligned}$$

ce qui, en utilisant la remarque faite au début de la démonstration du lemme 7.1(ii), montre que $\mathbf{V}^{1/2}\mathbf{R}\mathbf{V}^{1/2}$ est ϑ -cocoercif. $\square$

Lemme 7.4. Soient $\mathbf{B}: \mathcal{G} \rightarrow 2^{\mathcal{G}}$, $\mathbf{C}: \mathcal{H} \rightarrow \mathcal{H}$, $\mathbf{D}: \mathcal{G} \rightarrow 2^{\mathcal{G}}$, $\mathbf{Q}: \mathcal{K} \rightarrow 2^{\mathcal{K}}$, et $\mathbf{R}: \mathcal{K} \rightarrow \mathcal{K}$. On suppose que l'opérateur $\mathbf{A}$ défini dans la proposition 7.2 est nul. Soient $\mathbf{W} \in \mathcal{S}^+(\mathcal{H})$ et $\mathbf{U} \in \mathcal{S}^+(\mathcal{G})$ tels que $\|\mathbf{U}^{1/2}\mathbf{LW}^{1/2}\| < 1$. Soit $\mathbf{V} \in \mathcal{S}^+(\mathcal{K})$ défini par (7.25). Pour tout $\mathbf{z} = (\mathbf{x}, \mathbf{v}) \in \mathcal{K}$ et $(\mathbf{e}_1, \mathbf{e}_2) \in \mathcal{K}$, posons

$$\begin{cases} \mathbf{u} = \mathbf{J}_{\mathbf{UB}^{-1}}(\mathbf{v} + \mathbf{U}(\mathbf{L}(\mathbf{x} - \mathbf{W}(\mathbf{Cx} + \mathbf{L}^*\mathbf{v})) - \mathbf{D}^{-1}\mathbf{v} + \mathbf{e}_2)) \\ \mathbf{y} = \mathbf{x} - \mathbf{W}(\mathbf{Cx} + \mathbf{L}^*\mathbf{u} + \mathbf{e}_1). \end{cases} \quad (7.27)$$

On a alors $(\mathbf{y}, \mathbf{u}) = \mathbf{J}_{\mathbf{VQ}}(\mathbf{z} - \mathbf{VRz} + \mathbf{s})$ où

$$\mathbf{s} = (-\mathbf{W}\mathbf{e}_1, (\mathbf{U}^{-1} - \mathbf{LWL}^*)^{-1}(\mathbf{e}_2 + \mathbf{LW}\mathbf{e}_1)). \quad (7.28)$$

Démonstration. Soient $\mathbf{z} = (\mathbf{x}, \mathbf{v}) \in \mathcal{K}$ et $\mathbf{s} = (\mathbf{e}'_1, \mathbf{e}'_2) \in \mathcal{K}$. On a les équivalences suivantes :

$$\begin{aligned} & (\mathbf{y}, \mathbf{u}) = \mathbf{J}_{\mathbf{VQ}}(\mathbf{z} - \mathbf{VRz} + \mathbf{s}) \\ \Leftrightarrow & \mathbf{V}^{-1}(\mathbf{z} + \mathbf{s} - (\mathbf{y}, \mathbf{u})) - \mathbf{Rz} \in \mathbf{Q}(\mathbf{y}, \mathbf{u}) \\ \Leftrightarrow & \begin{cases} \mathbf{W}^{-1}(\mathbf{x} - \mathbf{y} + \mathbf{e}'_1) - \mathbf{L}^*\mathbf{u} - \mathbf{Cx} = \mathbf{0} \\ (\mathbf{U}^{-1} - \mathbf{LWL}^*)(\mathbf{v} - \mathbf{u} + \mathbf{e}'_2) + \mathbf{Ly} - \mathbf{D}^{-1}\mathbf{v} \in \mathbf{B}^{-1}\mathbf{u} \end{cases} \\ \Leftrightarrow & \begin{cases} \mathbf{y} = \mathbf{x} - \mathbf{W}(\mathbf{Cx} + \mathbf{L}^*\mathbf{u}) + \mathbf{e}'_1 \\ \mathbf{v} + \mathbf{e}'_2 + \mathbf{U}(\mathbf{L}(\mathbf{x} - \mathbf{W}(\mathbf{Cx} + \mathbf{L}^*\mathbf{v} + \mathbf{L}^*\mathbf{e}'_2) + \mathbf{e}'_1) - \mathbf{D}^{-1}\mathbf{v}) \in (\mathbf{Id} + \mathbf{UB}^{-1})\mathbf{u} \end{cases} \\ \Leftrightarrow & \begin{cases} \mathbf{u} = \mathbf{J}_{\mathbf{UB}^{-1}}(\mathbf{v} + \mathbf{e}'_2 + \mathbf{U}(\mathbf{L}(\mathbf{x} - \mathbf{W}(\mathbf{Cx} + \mathbf{L}^*\mathbf{v} + \mathbf{L}^*\mathbf{e}'_2) + \mathbf{e}'_1) - \mathbf{D}^{-1}\mathbf{v})) \\ \mathbf{y} = \mathbf{x} - \mathbf{W}(\mathbf{Cx} + \mathbf{L}^*\mathbf{u}) + \mathbf{e}'_1, \end{cases} \end{aligned}$$

qui implique (7.27) pour

$$\begin{cases} -\mathbf{W}\mathbf{e}_1 = \mathbf{e}'_1 \\ \mathbf{U}\mathbf{e}_2 = \mathbf{UL}(\mathbf{e}'_1 - \mathbf{WL}^*\mathbf{e}'_2) + \mathbf{e}'_2. \end{cases}$$

Puisque $\mathbf{U}^{-1} - \mathbf{LWL}^*$ est un isomorphisme, ces égalités sont équivalentes à (7.28). $\square$

D'après les lemmes 7.3 et 7.4, on peut déduire un second type d'algorithme pour résoudre le problème 7.1 lorsque $\mathbf{A} = \mathbf{0}$.

Soient $\mathbf{W}$ et $\mathbf{U}$ définis par (7.20). Nous proposons de résoudre le problème 7.1 lorsque $\mathbf{A} = \mathbf{0}$ en utilisant l'algorithme 7.5.

Algorithme 7.5 Algorithme primal-dual alterné aléatoirement par bloc (version 2)

Initialisation : Soient $\mathbf{x}_0$, et $(\mathbf{c}_k)_{k \in \mathbb{N}}$ des variables aléatoires à valeurs dans $\mathcal{H}$, et soient $\mathbf{v}_0$, $(\mathbf{b}_k)_{k \in \mathbb{N}}$, et $(\mathbf{d}_k)_{k \in \mathbb{N}}$ des variables aléatoires à valeurs dans $\mathcal{G}$. Pour tout $j \in \{1, \dots, p\}$, soit $\mathbf{W}_j \in \mathcal{S}^+(\mathcal{H}_j)$, et, pour tout $l \in \{1, \dots, q\}$, soit $\mathbf{U}_l \in \mathcal{S}^+(\mathcal{G}_l)$. Soient $(\varepsilon_k)_{k \in \mathbb{N}}$ des variables aléatoires identiquement distribuées à valeurs dans $\mathbb{D}_{p+q}$. Soit, pour tout $k \in \mathbb{N}$, $\lambda_k \in]0, +\infty[$.

Itérations :

Pour $k = 0, 1, \dots$

pour $j = 1, \dots, p$ $\eta_k^{(j)} = \max \left\{ \varepsilon_k^{(p+l)} \mid l \in \mathbb{L}_j^* \right\}$ $w_k^{(j)} = \eta_k^{(j)} \left(x_k^{(j)} - \mathbf{W}_j (\mathbf{C}_j x_k^{(j)} + c_k^{(j)}) \right)$ pour $l = 1, \dots, q$ $u_k^{(l)} = \varepsilon_k^{(p+l)} \left(\mathbf{J}_{\mathbf{U}_l \mathbf{B}_l^{-1}} \left(v_k^{(l)} + \mathbf{U}_l \left(\sum_{j \in \mathbb{L}_l} \mathbf{L}_{l,j} (w_k^{(j)} - \mathbf{W}_j \sum_{l' \in \mathbb{L}_j^*} \mathbf{L}_{l',j}^* v_k^{(l')}) - \mathbf{D}_l^{-1} v_k^{(l)} + d_k^{(l)} \right) \right) + b_k^{(l)} \right)$ $v_{k+1}^{(l)} = v_k^{(l)} + \lambda_k \varepsilon_k^{(p+l)} (u_k^{(l)} - v_k^{(l)})$ pour $j = 1, \dots, p$ $x_{k+1}^{(j)} = x_k^{(j)} + \lambda_k \varepsilon_k^{(j)} \left(w_k^{(j)} - \mathbf{W}_j \sum_{l \in \mathbb{L}_j^*} \mathbf{L}_{l,j}^* u_k^{(l)} - x_k^{(j)} \right).$
--

Proposition 7.5. Considérons l'algorithme 7.5, (μ, ν) donnés dans la proposition 7.3, et $\mathbf{W}$ défini par (7.20). Supposons que

$$\min \{ \mu, \nu(1 - \|\mathbf{U}^{1/2} \mathbf{L} \mathbf{W}^{1/2}\|^2) \} > \frac{1}{2}, \quad (7.29)$$

et que les éléments de la suite $(\lambda_k)_{k \in \mathbb{N}}$ sont dans $]0, 1]$ et vérifient $\inf_{k \in \mathbb{N}} \lambda_k > 0$. Posons $(\forall k \in \mathbb{N})$ $\mathbf{E}_k = \sigma(\varepsilon_k)$ et $\mathbf{X}_k = \sigma(\mathbf{x}_{k'}, \mathbf{v}_{k'})_{0 \leq k' \leq k}$, et supposons de plus que l'assertion suivante est vérifiée

- (i) $\sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{b}_k\|^2 | \mathbf{X}_k)} < +\infty$, $\sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{c}_k\|^2 | \mathbf{X}_k)} < +\infty$, et $\sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{d}_k\|^2 | \mathbf{X}_k)} < +\infty$ $\mathbf{P}$ -presque sûrement,

et que les condition (ii)-(iii) de la proposition 7.4 sont satisfaites.

Si, dans le problème 7.1, $(\forall j \in \{1, \dots, p\}) \mathbf{A}_j = \mathbf{0}$, alors $(\mathbf{x}_k)_{k \in \mathbb{N}}$ converge faiblement $\mathbf{P}$ -presque sûrement vers une variable aléatoire à valeurs dans $\mathbf{F}$, et $(\mathbf{v}_k)_{k \in \mathbb{N}}$ converge faiblement $\mathbf{P}$ -presque sûrement vers une variable aléatoire à valeurs dans $\mathbf{F}^*$.

La démonstration de la proposition 7.5 est donnée dans l'annexe 7.C.

Remarque 7.4.

- (i) Pour tout $j \in \{1, \dots, p\}$, une constante de cocoercivité de C_j est donnée par $\tilde{\mu}_j \in]0, +\infty[$ et, pour tout $l \in \{1, \dots, q\}$, une constante de forte monotonie de D_l est donnée par $\tilde{\nu}_l \in]0, +\infty[$. En utilisant (7.21)-(7.22), une condition nécessaire pour que (7.29) soit satisfaite est

$$\min \left\{ (\|W_j\|^{-1} \tilde{\mu}_j)_{1 \leq j \leq p}, \left(1 - \sum_{j=1}^p \sum_{l=1}^q \|U_l^{1/2} L_{l,j} W_j^{1/2}\|^2\right) (\|U_l\|^{-1} \tilde{\nu}_l)_{1 \leq l \leq q} \right\} > \frac{1}{2}. \quad (7.30)$$

- (ii) Dans le cas où, pour tout $l \in \{1, \dots, q\}$, $D_l^{-1} = 0$, nous pouvons utiliser une condition moins restrictive que la condition (7.29) donnée par

$$\begin{cases} \|\mathbf{U}^{1/2} \mathbf{L} \mathbf{W}^{1/2}\|^2 < 1 \\ \mu > \frac{1}{2}. \end{cases} \quad (7.31)$$

De plus, une condition nécessaire pour que (7.30) soit satisfaite est alors

$$\begin{cases} \sum_{j=1}^p \sum_{l=1}^q \|U_l^{1/2} L_{l,j} W_j^{1/2}\|^2 < 1 \\ \min \{ (\|W_j\|^{-1} \tilde{\mu}_j)_{1 \leq j \leq p} \} > \frac{1}{2}. \end{cases}$$

Notons que cette condition est moins restrictive que (7.24).

De la même façon que dans la section 7.3.2, remarquons nous pourrions trouver une forme symétrique de l'algorithme 7.5. En effet, les rôles des variables primales et duales des problèmes (7.10) et (7.11) peuvent être échangées en choisissant, dans le problème dual (7.11), $(\forall l \in \{1, \dots, q\}) B_l^{-1} = 0$ à la place de $(\forall j \in \{1, \dots, p\}) A_j = 0$ dans le problème primal (7.10).

7.3.4 Application aux problèmes variationnels

Dans cette section nous allons montrer que les résultats obtenus dans la section 7.3 nous permettent de construire des algorithmes proximaux primaux-duaux alternés par bloc pour résoudre un ensemble de problèmes d'optimisation convexe non nécessairement lisses.

Nous considérons la classe de problèmes d'optimisation suivante :

Problème 7.2. Pour tout $j \in \{1, \dots, p\}$, soit $f_j \in \Gamma_0(\mathcal{H}_j)$, soit $h_j \in \Gamma_0(\mathcal{H}_j)$ une fonction différentiable de gradient Lipschitz, et, pour tout $l \in \{1, \dots, q\}$, soit $g_l \in \Gamma_0(\mathcal{G}_l)$, soit $l_l \in \Gamma_0(\mathcal{G}_l)$ une fonction fortement convexe, et soit $L_{l,j} \in \mathcal{B}(\mathcal{H}_j, \mathcal{G}_l)$. Supposons que (7.8) et (7.9) sont satisfaites, et qu'il existe $(\bar{x}^{(1)}, \dots, \bar{x}^{(p)}) \in \mathcal{H}_1 \oplus \dots \oplus \mathcal{H}_p$ tels que

$$(\forall j \in \{1, \dots, p\}) \quad 0 \in \partial f_j(\bar{x}^{(j)}) + \nabla h_j(\bar{x}^{(j)}) + \sum_{l=1}^q L_{l,j}^* (\partial g_l \square \partial l_l) \left(\sum_{j'=1}^p L_{l,j'} \bar{x}^{(j')} \right). \quad (7.32)$$

Soit $\tilde{\mathbf{F}}$ l'ensemble défini par

$$\tilde{\mathbf{F}} = \underset{\mathbf{x}^{(1)} \in \mathcal{H}_1, \dots, \mathbf{x}^{(p)} \in \mathcal{H}_p}{\text{Argmin}} \sum_{j=1}^p (f_j(\mathbf{x}^{(j)}) + h_j(\mathbf{x}^{(j)})) + \sum_{l=1}^q (\mathbf{g}_l \square \mathbf{l}_l) \left(\sum_{j=1}^p \mathbf{L}_{l,j} \mathbf{x}^{(j)} \right) \quad (7.33)$$

et soit $\tilde{\mathbf{F}}^*$ l'ensemble des solutions du problème dual associé défini par :

$$\tilde{\mathbf{F}}^* = \underset{\mathbf{v}^{(1)} \in \mathcal{G}_1, \dots, \mathbf{v}^{(q)} \in \mathcal{G}_q}{\text{Argmin}} \sum_{j=1}^p (f_j^* \square h_j^*) \left(- \sum_{l=1}^q \mathbf{L}_{l,j}^* \mathbf{v}^{(l)} \right) + \sum_{l=1}^q (\mathbf{g}_l^*(\mathbf{v}^{(l)}) + \mathbf{l}_l^*(\mathbf{v}^{(l)})). \quad (7.34)$$

L'objectif est de trouver un couple de variables aléatoires $(\hat{\mathbf{x}}, \hat{\mathbf{v}})$ tel que $\hat{\mathbf{x}}$ soit à valeurs dans $\tilde{\mathbf{F}}$ et $\hat{\mathbf{v}}$ soit à valeurs dans $\tilde{\mathbf{F}}^*$.

Notons que sous certaines hypothèses peu contraignantes, on peut assurer que la condition d'inclusion dans le problème 7.2 est satisfaite :

Proposition 7.6. [Combettes, 2013, Prop. 5.3] *Considérons le problème 7.2, et supposons que (7.33) admet une solution. L'existence de $(\bar{\mathbf{x}}^{(1)}, \dots, \bar{\mathbf{x}}^{(p)}) \in \mathcal{H}_1 \oplus \dots \oplus \mathcal{H}_p$ vérifiant (7.32) est garantie dans chacun des cas suivants :*

- (i) *Pour tout $j \in \{1, \dots, p\}$, f_j est à valeurs réelles et, pour tout $l \in \{1, \dots, q\}$, $(\mathbf{x}^{(j)})_{1 \leq j \leq p} \mapsto \sum_{j=1}^p \mathbf{L}_{l,j} \mathbf{x}^{(j)}$ est surjective.*
- (ii) *Pour tout $l \in \{1, \dots, q\}$, $\mathbf{g}_l$ ou $\mathbf{l}_l$ est à valeurs réelles.*
- (iii) *$(\mathcal{H}_j)_{1 \leq j \leq p}$ et $(\mathcal{G}_l)_{1 \leq l \leq q}$ sont de dimension finie et $(\forall j \in \{1, \dots, p\}) (\exists \mathbf{x}^{(j)} \in \text{ri dom } f_j)$ tel que $(\forall l \in \{1, \dots, q\}) \sum_{j=1}^p \mathbf{L}_{l,j} \mathbf{x}^{(j)} \in \text{ri dom } \mathbf{g}_l + \text{ri dom } \mathbf{l}_l$.*

Nous proposons de résoudre le problème 7.2 en utilisant l'algorithme 7.6 déduit de l'algorithme 7.3.

Algorithme 7.6 Algorithme primal-dual (version 1) appliqué dans un cadre variationnel

Initialisation : Soient $\mathbf{x}_0$, $(\mathbf{a}_k)_{k \in \mathbb{N}}$ et $(\mathbf{c}_k)_{k \in \mathbb{N}}$ des variables aléatoires à valeurs dans $\mathcal{H}$, et soient $\mathbf{v}_0$, $(\mathbf{b}_k)_{k \in \mathbb{N}}$, et $(\mathbf{d}_k)_{k \in \mathbb{N}}$ des variables aléatoires à valeurs dans $\mathcal{G}$. Pour tout $j \in \{1, \dots, p\}$, soit $\mathbf{W}_j \in \mathcal{S}^+(\mathcal{H}_j)$, et, pour tout $l \in \{1, \dots, q\}$, soit $\mathbf{U}_l \in \mathcal{S}^+(\mathcal{G}_l)$. Soient $(\varepsilon_k)_{k \in \mathbb{N}}$ des variables aléatoires identiquement distribuées à valeurs dans $\mathbb{D}_{p+q}$. Soit, pour tout $k \in \mathbb{N}$, $\lambda_k \in]0, +\infty[$.

Itérations :

Pour $k = 0, 1, \dots$

$$\left[\begin{array}{l} \text{pour } j = 1, \dots, p \\ \quad \left[\begin{array}{l} y_k^{(j)} = \varepsilon_k^{(j)} \left(\text{prox}_{\mathbf{W}_j^{-1}, \mathbf{f}_j} \left(x_k^{(j)} - \mathbf{W}_j \left(\sum_{l \in \mathbb{L}_j^*} \mathbf{L}_{l,j}^* v_k^{(l)} + \nabla \mathbf{h}_j(x_k^{(j)}) + c_k^{(j)} \right) \right) + a_k^{(j)} \right), \\ x_{k+1}^{(j)} = x_k^{(j)} + \lambda_k \varepsilon_k^{(j)} (y_k^{(j)} - x_k^{(j)}), \end{array} \right. \\ \text{pour } l = 1, \dots, q \\ \quad \left[\begin{array}{l} u_k^{(l)} = \varepsilon_k^{(p+l)} \left(\text{prox}_{\mathbf{U}_l^{-1}, \mathbf{g}_l^*} \left(v_k^{(l)} + \mathbf{U}_l \left(\sum_{j \in \mathbb{L}_l} \mathbf{L}_{l,j} (2y_k^{(j)} - x_k^{(j)}) - \nabla l_l^*(v_k^{(l)}) + d_k^{(l)} \right) \right) + b_k^{(l)} \right), \\ v_{k+1}^{(l)} = v_k^{(l)} + \lambda_k \varepsilon_k^{(p+l)} (u_k^{(l)} - v_k^{(l)}). \end{array} \right. \end{array} \right.$$

Le résultat de convergence suivant, concernant l'algorithme 7.6, se déduit de la proposition 7.3 :

Proposition 7.7. *Considérons l'algorithme 7.6 et $\mathbf{W}$ et $\mathbf{U}$ définis par (7.20). Supposons que, pour tout $j \in \{1, \dots, p\}$, le gradient de $\mathbf{h}_j \circ \mathbf{W}_j^{1/2}$ est μ_j^{-1} -Lipschitz, avec $\mu_j \in]0, +\infty[$, et que, pour tout $l \in \{1, \dots, q\}$, le gradient de $l_l^* \circ \mathbf{U}_l^{1/2}$ est ν_l^{-1} -Lipschitz, avec $\nu_l \in]0, +\infty[$. Supposons de plus que (7.19) est satisfaite où ϑ_α est défini par (7.16), $\mu = \min\{\mu_1, \dots, \mu_p\}$, et $\nu = \min\{\nu_1, \dots, \nu_q\}$. Supposons que les éléments de $(\lambda_k)_{k \in \mathbb{N}}$ sont dans $]0, 1]$ et vérifient $\inf_{k \in \mathbb{N}} \lambda_k > 0$. Posons $(\forall k \in \mathbb{N}) \mathbf{E}_k = \sigma(\varepsilon_k)$ et $\mathbf{X}_k = \sigma(\mathbf{x}_{k'}, \mathbf{v}_{k'})_{0 \leq k' \leq k}$ et supposons que les conditions (i)-(iii) de la proposition 7.3 sont vérifiées.*

La suite $(\mathbf{x}_k)_{k \in \mathbb{N}}$ converge alors faiblement P-presque sûrement vers une variable aléatoire à valeurs dans $\tilde{\mathbf{F}}$, et $(\mathbf{v}_k)_{k \in \mathbb{N}}$ converge faiblement P-presque sûrement vers une variable aléatoire à valeurs dans $\tilde{\mathbf{F}}^$.*

La démonstration de la proposition 7.7 est donnée dans l'annexe 7.D.

De même que pour obtenir l'algorithme 7.4, les problèmes (7.33) et (7.34) étant symétriques, les rôles des variables primales et duales peuvent être échangées dans les deux problèmes. On obtient alors l'algorithme 7.7, dont la démonstration de convergence se déduit en remplaçant les conditions (ii)-(iii) de la proposition 7.3 par les conditions (ii)-(iii) de la proposition 7.4.

Algorithme 7.7 Algorithme primal-dual (version 1 symétrique) appliqué dans un cadre variationnel

Initialisation : Soient $\mathbf{x}_0, (\mathbf{a}_k)_{k \in \mathbb{N}}$ et $(\mathbf{c}_k)_{k \in \mathbb{N}}$ des variables aléatoires à valeurs dans $\mathcal{H}$, et soient $\mathbf{v}_0, (\mathbf{b}_k)_{k \in \mathbb{N}}$, et $(\mathbf{d}_k)_{k \in \mathbb{N}}$ des variables aléatoires à valeurs dans $\mathcal{G}$. Pour tout $j \in \{1, \dots, p\}$, soit $\mathbf{W}_j \in \mathcal{S}^+(\mathcal{H}_j)$, et, pour tout $l \in \{1, \dots, q\}$, soit $\mathbf{U}_l \in \mathcal{S}^+(\mathcal{G}_l)$. Soient $(\varepsilon_k)_{k \in \mathbb{N}}$ des variables aléatoires identiquement distribuées à valeurs dans $\mathbb{D}_{p+q}$. Soit, pour tout $k \in \mathbb{N}$, $\lambda_k \in]0, +\infty[$.

Itérations :

Pour $k = 0, 1, \dots$

$$\left[\begin{array}{l} \text{pour } l = 1, \dots, q \\ \quad \left[\begin{array}{l} u_k^{(l)} = \varepsilon_k^{(p+l)} \left(\text{prox}_{\mathbf{U}_l^{-1}, \mathbf{g}_l^*} \left(v_k^{(l)} + \mathbf{U}_l \left(\sum_{j \in \mathbb{L}_l} \mathbf{L}_{l,j} x_k^{(j)} - \nabla \mathbf{l}_l^* v_k^{(l)} + d_k^{(l)} \right) \right) + b_k^{(l)} \right), \\ v_{k+1}^{(l)} = v_k^{(l)} + \lambda_k \varepsilon_k^{(p+l)} (u_k^{(l)} - v_k^{(l)}), \end{array} \right. \\ \text{pour } j = 1, \dots, p \\ \quad \left[\begin{array}{l} y_k^{(j)} = \varepsilon_k^{(j)} \left(\text{prox}_{\mathbf{W}_j^{-1}, \mathbf{f}_j} \left(x_k^{(j)} - \mathbf{W}_j \left(\sum_{l \in \mathbb{L}_j^*} \mathbf{L}_{l,j}^* (2u_k^{(l)} - v_k^{(l)}) + \nabla \mathbf{h}_j x_k^{(j)} + c_k^{(j)} \right) \right) + a_k^{(j)} \right), \\ x_{k+1}^{(j)} = x_k^{(j)} + \lambda_k \varepsilon_k^{(j)} (y_k^{(j)} - x_k^{(j)}). \end{array} \right. \end{array} \right.$$

Proposition 7.8. Considérons l'algorithme 7.7, et $\mathbf{W}$ et $\mathbf{U}$ définis par (7.20) et supposons que, pour tout $j \in \{1, \dots, p\}$, le gradient de $\mathbf{h}_j \circ \mathbf{W}_j^{1/2}$ est μ_j^{-1} -Lipschitz, avec $\mu_j \in]0, +\infty[$, et que, pour tout $l \in \{1, \dots, q\}$, le gradient de $\mathbf{l}_l^* \circ \mathbf{U}_l^{1/2}$ est ν_l^{-1} -Lipschitz, avec $\nu_l \in]0, +\infty[$. Supposons de plus que (7.19) est satisfaite où ϑ_α est défini par (7.16), $\mu = \min\{\mu_1, \dots, \mu_p\}$, et $\nu = \min\{\nu_1, \dots, \nu_q\}$. Supposons que les éléments de $(\lambda_k)_{k \in \mathbb{N}}$ sont dans $]0, 1]$ et vérifient $\inf_{k \in \mathbb{N}} \lambda_k > 0$. Posons $(\forall k \in \mathbb{N}) \mathbf{E}_k = \sigma(\varepsilon_k)$ et $\mathbf{X}_k = \sigma(\mathbf{x}_{k'}, \mathbf{v}_{k'})_{0 \leq k' \leq k}$ et supposons que la condition (i) de la proposition 7.3 et les conditions (ii)-(iii) de la proposition 7.4 sont vérifiées.

La suite $(\mathbf{x}_k)_{k \in \mathbb{N}}$ converge alors faiblement P-presque sûrement vers une variable aléatoire à valeurs dans $\tilde{\mathbf{F}}$, et $(\mathbf{v}_k)_{k \in \mathbb{N}}$ converge faiblement P-presque sûrement vers une variable aléatoire à valeurs dans $\tilde{\mathbf{F}}^*$.

De façon similaire, l'algorithme 7.5 nous permet d'obtenir l'algorithme 7.8, et nous pouvons déduire de la proposition 7.5 le résultat de convergence suivant.

Proposition 7.9. Considérons l'algorithme 7.8, $\mathbf{W}$ et $\mathbf{U}$ définis par (7.20) et (μ, ν) définis par la proposition 7.7. Supposons que la condition (7.29) est satisfaite et que les éléments de $(\lambda_k)_{k \in \mathbb{N}}$ sont dans $]0, 1]$ et vérifient $\inf_{k \in \mathbb{N}} \lambda_k > 0$. Posons $(\forall k \in \mathbb{N}) \mathbf{E}_k = \sigma(\varepsilon_k)$ et $\mathbf{X}_k = \sigma(\mathbf{x}_{k'}, \mathbf{v}_{k'})_{0 \leq k' \leq k}$ et supposons que la condition (i) de la proposition 7.5, les conditions (ii)-(iii) de la proposition 7.4 sont satisfaites.

Si, dans le problème 7.2, $(\forall j \in \{1, \dots, p\}) \mathbf{f}_j = 0$, alors $(\mathbf{x}_k)_{k \in \mathbb{N}}$ converge faiblement P-presque

sûrement vers une variable aléatoire à valeurs dans $\tilde{\mathbf{F}}$, et $(\mathbf{v}_k)_{k \in \mathbb{N}}$ converge faiblement $\mathbf{P}$ -presque sûrement vers une variable aléatoire à valeurs dans $\tilde{\mathbf{F}}^*$.

Algorithme 7.8 Algorithme primal-dual (version 2) appliqué dans un cadre variationnel

Initialisation : Soient $\mathbf{x}_0$, et $(\mathbf{c}_k)_{k \in \mathbb{N}}$ des variables aléatoires à valeurs dans $\mathcal{H}$, et soient $\mathbf{v}_0$, $(\mathbf{b}_k)_{k \in \mathbb{N}}$, et $(\mathbf{d}_k)_{k \in \mathbb{N}}$ des variables aléatoires à valeurs dans $\mathcal{G}$. Pour tout $j \in \{1, \dots, p\}$, soit $\mathbf{W}_j \in \mathcal{S}^+(\mathcal{H}_j)$, et, pour tout $l \in \{1, \dots, q\}$, soit $\mathbf{U}_l \in \mathcal{S}^+(\mathcal{G}_l)$. Soient $(\varepsilon_k)_{k \in \mathbb{N}}$ des variables aléatoires identiquement distribuées à valeurs dans $\mathbb{D}_{p+q}$. Soit, pour tout $k \in \mathbb{N}$, $\lambda_k \in]0, +\infty[$.

Itérations :

Pour $k = 0, 1, \dots$

$$\begin{array}{l}
 \text{pour } j = 1, \dots, p \\
 \left[\begin{array}{l} \eta_k^{(j)} = \max \left\{ \varepsilon_k^{(p+l)} \mid l \in \mathbb{L}_j^* \right\} \\ w_k^{(j)} = \eta_k^{(j)} \left(x_k^{(j)} - \mathbf{W}_j (\nabla \mathbf{h}_j(x_k^{(j)}) + c_k^{(j)}) \right) \end{array} \right. \\
 \text{pour } l = 1, \dots, q \\
 \left[\begin{array}{l} u_k^{(l)} = \varepsilon_k^{(p+l)} \left(\text{prox}_{\mathbf{U}_l^{-1}, \mathbf{g}_l^*} \left(v_k^{(l)} + \mathbf{U}_l \left(\sum_{j \in \mathbb{L}_l} \mathbf{L}_{l,j}(w_k^{(j)} - \mathbf{W}_j \sum_{l' \in \mathbb{L}_j^*} \mathbf{L}_{l',j}^* v_k^{(l')}) \right. \right. \right. \\ \left. \left. \left. - \nabla \mathbf{l}_l^*(v_k^{(l)}) + d_k^{(l)} \right) \right) + b_k^{(l)} \right) \\ v_{k+1}^{(l)} = v_k^{(l)} + \lambda_k \varepsilon_k^{(p+l)} (u_k^{(l)} - v_k^{(l)}) \end{array} \right. \\
 \text{pour } j = 1, \dots, p \\
 \left[\begin{array}{l} x_{k+1}^{(j)} = x_k^{(j)} + \lambda_k \varepsilon_k^{(j)} \left(w_k^{(j)} - \mathbf{W}_j \sum_{l \in \mathbb{L}_j^*} \mathbf{L}_{l,j}^* u_k^{(l)} - x_k^{(j)} \right). \end{array} \right.
 \end{array}$$

Dans la remarque suivante nous donnons les liens existants entre les algorithmes proximaux alternés par bloc proposés dans cette section et les travaux existants dans la littérature.

Remarque 7.5.

- (i) En pratique, il est souvent intéressant de considérer le cas particulier du problème (7.2) où $(\forall l \in \{1, \dots, q\}) \mathbf{l}_l = \iota_{\{0\}}$, i.e. $\mathbf{l}_l^* = 0$. Dans ce cas, l'ensemble $\tilde{\mathbf{F}}$ se réécrit

$$\tilde{\mathbf{F}} = \underset{\mathbf{x}^{(1)} \in \mathcal{H}_1, \dots, \mathbf{x}^{(p)} \in \mathcal{H}_p}{\text{Argmin}} \sum_{j=1}^p (\mathbf{f}_j(\mathbf{x}^{(j)}) + \mathbf{h}_j(\mathbf{x}^{(j)})) + \sum_{l=1}^q \mathbf{g}_l \left(\sum_{j=1}^p \mathbf{L}_{l,j} \mathbf{x}^{(j)} \right).$$

- (ii) L'algorithme 7.6 généralise les approches déterministes proposées dans [Chambolle et Pock, 2010; Condat, 2013; Esser et al., 2010; He et Yuan, 2012; Vũ, 2013] présentant le cas où $p = 1$, en y introduisant un balayage aléatoire des blocs et en considérant des erreurs stochastiques.

De façon similaire, l'algorithme 7.8 généralise les algorithmes présentés dans [Chen et al., 2013; Loris et Verhoeven, 2011], développés dans le cas déterministe ne prenant en compte aucun terme d'erreur, dans le cas particulier où $p = q = 1$, $\mathcal{H}_1$ et $\mathcal{G}_1$ sont des espaces de dimension finie, $\mathbf{l}_1 = \iota_{\{0\}}$, $\mathbf{W}_1 = \tau \text{Id}$ $\mathbf{U}_1 = \rho \text{Id}$, avec $(\tau, \rho) \in]0, +\infty[^2$, sans relaxation ($\lambda_k \equiv 1$) ou bien avec une relaxation constante ($\lambda_k \equiv \lambda_0 < 1$).

Récemment, ces travaux ont été généralisés au cas où les espaces de Hilbert peuvent être de dimension infinie, pour $p = 1$ et $q > 1$, en utilisant des opérateurs de préconditionnement arbitraires, et en prenant en compte des termes d'erreurs déterministes et sommables [Combettes et al., 2014].

L'intérêt pratique d'utiliser des opérateurs de préconditionnement pour des méthodes primales-duales, afin d'en accélérer la convergence, a été démontré dans [Combettes et al., 2014; Pock et Chambolle, 2011; Repetti et al., 2012].

- (iii) Dans [Combettes et Pesquet, 2015, Cor. 5.5], un autre algorithme primal-dual alterné aléatoirement par bloc a été proposé pour résoudre un cas particulier du problème 7.2, obtenu lorsque $(\forall j \in \{1, \dots, p\}) \mathbf{h}_j = 0$ et $(\forall l \in \{1, \dots, q\}) \mathbf{l}_l = \iota_{\{0\}}$. Cet algorithme est obtenu en se basant sur les itérations de l'algorithme de Douglas-Rachford. On peut aussi mentionner l'algorithme ADMM aléatoire développé pour des espaces de dimensions finies dans [Iutzeler et al., 2013]. Remarquons cependant que l'algorithme donné dans [Combettes et Pesquet, 2015, Cor. 5.5] requiert l'inversion de $\text{Id} + \mathbf{L}\mathbf{L}^*$ ou $\text{Id} + \mathbf{L}^*\mathbf{L}$ (voir [Combettes et Pesquet, 2015, Rmq. 5.4]), contrairement aux algorithmes 7.6 et 7.8 qui ne nécessitent aucune inversion d'opérateur linéaire.

Le diagramme donné dans la figure 7.5, présentée à la fin de la section 7.4, résume les liens existants entre les algorithmes proposés dans les sections 7.3 et 7.4.

7.4 APPLICATION DE L'ALGORITHME PRIMAL-DUAL ALTERNÉ ALÉATOIREMENT PAR BLOC

7.4.1 Description du problème

Considérons un maillage surfacique représenté par ses p nœuds de coordonnées spatiales $\bar{\mathbf{x}} = (\bar{\mathbf{x}}^{(j)})_{1 \leq j \leq p} \in \mathbb{R}^{p \times 3}$, où, pour tout $j \in \{1, \dots, p\}$, $\bar{\mathbf{x}}^{(j)} = (\bar{\mathbf{x}}^{(\mathcal{X}_j)}, \bar{\mathbf{x}}^{(\mathcal{Y}_j)}, \bar{\mathbf{x}}^{(\mathcal{Z}_j)}) \in \mathbb{R}^3$ représente les

coordonnées spatiales du j -ème nœud, ainsi que par sa matrice d'adjacence³ $\mathbf{A} \in \mathbb{R}^{p \times p}$.

Nous nous intéressons au problème d'estimation de ce maillage à partir d'une observation inexacte de celui-ci, ayant la même "topologie" (i.e. la même matrice d'adjacence $\mathbf{A}$), mais dont la position des nœuds est connue de façon imprécise selon le modèle

$$\mathbf{z} = (\mathbf{z}^{(\mathcal{X})}, \mathbf{z}^{(\mathcal{Y})}, \mathbf{z}^{(\mathcal{Z})}) = \bar{\mathbf{x}} + \mathbf{b}, \quad (7.35)$$

où $\mathbf{b} = (\mathbf{b}_j)_{1 \leq j \leq p} \in \mathbb{R}^{p \times 3}$ est une réalisation d'une variable aléatoire modélisant les incertitudes des positions spatiales des nœuds du maillage. Une méthode permettant de produire une estimée $\hat{\mathbf{x}} \in \mathbb{R}^{p \times 3}$ de $\bar{\mathbf{x}}$, est de la définir comme une solution du problème (7.2), en prenant $(\forall l \in \{1, \dots, q\})$ $l_l \equiv \iota_{\{0\}}$, avec les fonctions $(f_j)_{1 \leq j \leq p}$, $(h_j)_{1 \leq j \leq p}$ et $(g_l)_{1 \leq l \leq q}$ choisies de façon à incorporer des informations concernant le modèle d'observation ainsi que des connaissances *a priori* sur la position des nœuds du maillage [Couprie *et al.*, 2013; Shuman *et al.*, 2013].

Nous proposons d'utiliser un terme d'attache aux données nous permettant d'obtenir une estimation robuste, correspondant à la somme sur $j \in \{1, \dots, p\}$ des fonctions définies par :

$$(\forall \mathbf{x}^{(j)} = (\mathbf{x}^{(\mathcal{D}_j)})_{\mathcal{D} \in \{\mathcal{X}, \mathcal{Y}, \mathcal{Z}\}} \in \mathbb{R}^3) \quad h_j(\mathbf{x}^{(j)}) = \sum_{\mathcal{D} \in \{\mathcal{X}, \mathcal{Y}, \mathcal{Z}\}} \varphi_j(\mathbf{x}^{(\mathcal{D}_j)} - \mathbf{z}^{(\mathcal{D}_j)}), \quad (7.36)$$

où, pour tout $j \in \{1, \dots, p\}$, φ_j est la fonction de Huber paramétrée par $\delta^{(j)} \in]0, +\infty]$, définie par (2.28), qui est différentiable et de gradient Lipschitz.

De plus, afin d'obtenir un maillage estimé lisse, nous proposons d'utiliser comme fonction de régularisation une variation totale isotrope [Elmoataz *et al.*, 2008; Grady et Polimeni, 2010] définie, pour tout $\mathbf{x} \in \mathbb{R}^{p \times 3}$, par

$$\sum_{l=1}^p g_l(\mathbf{L}_l \mathbf{x}) = \sum_{l=1}^p \beta^{(l)} \sum_{\mathcal{D} \in \{\mathcal{X}, \mathcal{Y}, \mathcal{Z}\}} \|(\mathbf{x}^{(\mathcal{D}_l)} - \mathbf{x}^{(\mathcal{D}_i)})_{i \in \mathcal{V}_l}\|_2, \quad (7.37)$$

avec, en particulier,

$$(\forall l \in \{1, \dots, p\}) \quad \mathbf{L}_l: \mathbb{R}^{p \times 3} \rightarrow \mathbb{R}^{\kappa_l \times 3}: \mathbf{x} \mapsto \sum_{j=1}^p \mathbf{L}_{l,j} \mathbf{x}^{(j)} = (\mathbf{x}^{(l)} - \mathbf{x}^{(i)})_{i \in \mathcal{V}_l}, \quad (7.38)$$

$$g_l: \mathbb{R}^{\kappa_l \times 3} \rightarrow \mathbb{R}: \mathbf{u} \mapsto \beta^{(l)} \sum_{\mathcal{D} \in \{\mathcal{X}, \mathcal{Y}, \mathcal{Z}\}} \|\mathbf{u}^{\mathcal{D}}\|_2, \quad (7.39)$$

où, pour tout $l \in \{1, \dots, p\}$, $\beta^{(l)} \in]0, +\infty[$, et $\mathcal{V}_l$, de cardinal κ_l , désigne l'ensemble des voisins du nœud l (c'est-à-dire l'ensemble des indices $i \in \{1, \dots, p\}$ tels que $\mathbf{A}^{(i,l)} = 1$). Les ensembles définis par (7.8) et (7.9) sont alors

$$\begin{aligned} (\forall l \in \{1, \dots, p\}) \quad \mathbb{L}_l &= \{l\} \cup \mathcal{V}_l, \\ (\forall j \in \{1, \dots, p\}) \quad \mathbb{L}_j^* &= \{l \in \{1, \dots, p\} \mid l = j \text{ ou } j \in \mathcal{V}_l\}. \end{aligned}$$

Finalement, nous ajoutons un terme de contrainte de façon à s'assurer que les nœuds estimés soient compris dans un hypercube paramétré par des positions spatiales minimales et maximales données $(\mathbf{x}_{\min}^{(\mathcal{D})}, \mathbf{x}_{\max}^{(\mathcal{D})})_{\mathcal{D} \in \{\mathcal{X}, \mathcal{Y}, \mathcal{Z}\}} \in (\mathbb{R}^3)^2$:

$$(\forall j \in \{1, \dots, p\}) \quad f_j(\mathbf{x}^{(j)}) = \sum_{\mathcal{D} \in \{\mathcal{X}, \mathcal{Y}, \mathcal{Z}\}} \iota_{[\mathbf{x}_{\min}^{(\mathcal{D})}, \mathbf{x}_{\max}^{(\mathcal{D})}]}(\mathbf{x}^{(\mathcal{D}_j)}). \quad (7.40)$$

3. La matrice d'adjacence $\mathbf{A} = (\mathbf{A}^{(j,j')})_{1 \leq j, j' \leq p} \in \mathbb{R}^{p \times p}$ vérifie $\mathbf{A}^{(j,j')} = 1$ lorsque les nœuds j et j' (pour $j \neq j'$) sont connectés, et $\mathbf{A}^{(j,j')} = 0$ sinon.

7.4.2 Méthodes proposées

Nous proposons de tester indépendamment les deux versions des algorithmes que nous avons proposées dans la section 7.3.4 afin de résoudre le problème décrit ci-dessus. Pour cela, nous allons proposer des versions simplifiées des algorithmes 7.4 et 7.8 définies respectivement par les algorithmes 7.9 et 7.10.

7.4.2.1 Première méthode

Dans un premier temps, nous proposons d'utiliser l'algorithme 7.9.

Algorithme 7.9 Algorithme primal-dual (version 1 symétrique simplifiée) appliqué dans un cadre variationnel

Initialisation : Soient $\mathbf{x}_0$, $(\mathbf{a}_k)_{k \in \mathbb{N}}$ et $(\mathbf{c}_k)_{k \in \mathbb{N}}$ des variables aléatoires à valeurs dans $\mathcal{H}$, et soient $\mathbf{v}_0$, et $(\mathbf{b}_k)_{k \in \mathbb{N}}$ des variables aléatoires à valeurs dans $\mathcal{G}$. Pour tout $j \in \{1, \dots, p\}$, soit $\mathbf{W}_j \in \mathcal{S}^+(\mathcal{H}_j)$, et, pour tout $l \in \{1, \dots, q\}$, soit $\mathbf{U}_l \in \mathcal{S}^+(\mathcal{G}_l)$. Soient $(\varepsilon_k)_{k \in \mathbb{N}}$ des variables aléatoires identiquement distribuées à valeurs dans $\mathbb{D}_p$. Soit, pour tout $k \in \mathbb{N}$, $\lambda_k \in]0, +\infty[$.

Itérations :

Pour $k = 0, 1, \dots$

$$\left[\begin{array}{l} \text{pour } l = 1, \dots, q \\ \quad \eta_k^{(l)} = \max_{1 \leq j \leq p} \{ \varepsilon_k^{(j)} \mid l \in \mathbb{L}_j^* \} \\ \quad u_k^{(l)} = \eta_k^{(l)} \left(\text{prox}_{\mathbf{U}_l^{-1}, \mathbf{g}_l^*} \left(v_k^{(l)} + \mathbf{U}_l \sum_{j \in \mathbb{L}_l} \mathbf{L}_{l,j} x_k^{(j)} \right) + b_k^{(l)} \right), \\ \quad v_{k+1}^{(l)} = v_k^{(l)} + \lambda_k \eta_k^{(l)} (u_k^{(l)} - v_k^{(l)}), \\ \text{pour } j = 1, \dots, p \\ \quad y_k^{(j)} = \varepsilon_k^{(j)} \left(\text{prox}_{\mathbf{W}_j^{-1}, \mathbf{f}_j} \left(x_k^{(j)} - \mathbf{W}_j \left(\sum_{l \in \mathbb{L}_j^*} \mathbf{L}_{l,j}^* (2u_k^{(l)} - v_k^{(l)}) + \nabla \mathbf{h}_j x_k^{(j)} + c_k^{(j)} \right) \right) + a_k^{(j)} \right), \\ \quad x_{k+1}^{(j)} = x_k^{(j)} + \lambda_k \varepsilon_k^{(j)} (y_k^{(j)} - x_k^{(j)}). \end{array} \right.$$

Notons que, pour tout $k \in \mathbb{N}$ et pour tout $l \in \{1, \dots, q\}$, le choix

$$\eta_k^{(l)} = \max_{1 \leq j \leq p} \{ \varepsilon_k^{(j)} \mid l \in \mathbb{L}_j^* \}$$

dans l'algorithme 7.9 permet d'assurer que la condition (iii) de la Proposition 7.4 soit satisfaite. Ainsi, la convergence de l'algorithme 7.9 est assurée par la Proposition 7.8.

Cet algorithme requiert de choisir les matrices $(\mathbf{W}_j)_{1 \leq j \leq p}$ et $(\mathbf{U}_l)_{1 \leq l \leq q}$. Dans notre exemple nous prendrons simplement $q = p$, pour tout $j \in \{1, \dots, p\}$, $\mathbf{W}_j \equiv \tau \mathbf{I}_3$ et, pour tout $l \in \{1, \dots, p\}$, $\mathbf{U}_l \equiv \rho \mathbf{I}_{\kappa_l}$, où (τ, ρ) sont des constantes positives vérifiant la condition de convergence

$$2 \min\{(\mu_j)_{1 \leq j \leq p}\} \left(1 - \|\mathbf{W}^{1/2} \mathbf{L} \mathbf{U}^{1/2}\|^2 \right) = 2 \min\{(\mu_j)_{1 \leq j \leq p}\} \left(1 - \tau \rho \|\mathbf{L}\|^2 \right) > 1, \quad (7.41)$$

obtenue lorsque $l_l \equiv l_{\{0\}}$ et où $\mathbf{L} : \mathbf{x} \in \mathbb{R}^{p \times 3} \mapsto (\mathbf{L}_l \mathbf{x})_{1 \leq l \leq p}$. De plus, si, pour tout $j \in \{1, \dots, p\}$, $\tilde{\mu}_j^{-1}$ est la constante de Lipschitz de h_j , alors la contrainte ci-dessus est satisfaite en choisissant $(\tau, \rho) \in]0, +\infty[^2$ vérifiant

$$\min\{(\tilde{\mu}_j)_{1 \leq j \leq p}\}(\tau^{-1} - \rho\|\mathbf{L}\|^2) > 1/2. \quad (7.42)$$

7.4.2.2 Seconde méthode

Algorithme 7.10 Algorithme primal-dual (version 2 simplifiée) appliqué dans un cadre variationnel

Initialisation : Soient $\mathbf{x}_0$ et $(\mathbf{c}_k)_{k \in \mathbb{N}}$ des variables aléatoires à valeurs dans $\mathcal{H}$, et soient $\mathbf{v}_0$ et $(\mathbf{b}_k)_{k \in \mathbb{N}}$ des variables aléatoires à valeurs dans $\mathcal{G}$. Pour tout $j \in \{1, \dots, p\}$, soit $\mathbf{W}_j \in \mathcal{S}^+(\mathcal{H}_j)$, et, pour tout $l \in \{1, \dots, q\}$, soit $\mathbf{U}_l \in \mathcal{S}^+(\mathcal{G}_l)$. Soient $(\varepsilon_k)_{k \in \mathbb{N}}$ des variables aléatoires identiquement distribuées à valeurs dans $\mathbb{D}_{p+q}$. Soit, pour tout $k \in \mathbb{N}$, $\lambda_k \in]0, +\infty[$.

Itérations :

Pour $k = 0, 1, \dots$

$$\left[\begin{array}{l} \text{pour } j = 1, \dots, p \\ \quad \left[\begin{array}{l} \eta_k^{(j)} = \max \left\{ \varepsilon_k^{(p+l)} \mid l \in \mathbb{L}_j^* \right\} \\ w_k^{(j)} = \eta_k^{(j)} \left(x_k^{(j)} - \mathbf{W}_j (\nabla h_j(x_k^{(j)}) + c_k^{(j)}) \right) \end{array} \right. \\ \text{pour } l = 1, \dots, q \\ \quad \left[\begin{array}{l} \varepsilon_k^{(p+l)} = \max \left\{ \varepsilon_k^{(j)} \mid j \in \mathbb{L}_l \right\} \\ u_k^{(l)} = \varepsilon_k^{(p+l)} \left(\text{prox}_{\mathbf{U}_l^{-1}, \mathbf{g}_l^*} \left(v_k^{(l)} + \mathbf{U}_l \left(\sum_{j \in \mathbb{L}_l} \mathbf{L}_{l,j} (w_k^{(j)} - \mathbf{W}_j \sum_{l' \in \mathbb{L}_j^*} \mathbf{L}_{l',j}^* v_k^{(l')}) \right) + b_k^{(l)} \right) \right) \\ v_{k+1}^{(l)} = v_k^{(l)} + \lambda_k \varepsilon_k^{(p+l)} (u_k^{(l)} - v_k^{(l)}) \end{array} \right. \\ \text{pour } j = 1, \dots, p \\ \quad \left[\begin{array}{l} x_{k+1}^{(j)} = x_k^{(j)} + \lambda_k \varepsilon_k^{(j)} \left(w_k^{(j)} - \mathbf{W}_j \sum_{l \in \mathbb{L}_j^*} \mathbf{L}_{l,j}^* u_k^{(l)} - x_k^{(j)} \right). \end{array} \right. \end{array} \right.$$

Dans l'algorithme 7.10, pour tout $k \in \mathbb{N}$ et pour tout $l \in \{1, \dots, q\}$, la condition $\varepsilon_k^{(p+l)} = \max \{ \varepsilon_k^{(j)} \mid j \in \mathbb{L}_l \}$ permet d'assurer que la condition (iii) de la proposition 7.4 est satisfaite. Ainsi, la convergence de l'algorithme 7.10 est assurée par la proposition 7.9.

Afin d'appliquer ce second algorithme au problème décrit dans la section 7.4.1, nous posons $q = 2p$ et nous définissons, pour tout $l \in \{1, \dots, p\}$, $(\forall \mathbf{x} \in \mathbb{R}^{p \times 3})$ $\mathbf{g}_{p+l}(\sum_{j=1}^p \mathbf{L}_{p+l,j} \mathbf{x}^{(j)}) = \mathbf{f}_l(\mathbf{x}^{(l)})$, où $\mathbf{f}_l$ est définie par (7.40). Donc, pour tout $l \in \{1, \dots, p\}$,

$$\mathbf{g}_{p+l} = \mathbf{f}_l \quad \text{et} \quad (\forall j \in \{1, \dots, p\}) \quad \mathbf{L}_{p+l,j} = \begin{cases} \mathbf{I}_3 & \text{si } j = l \\ \mathbf{0}_3 & \text{sinon.} \end{cases}$$

Dans cet algorithme, nous choisissons, pour tout $j \in \{1, \dots, p\}$, $W_j \equiv \tau \mathbf{I}_3$, $U_j \equiv \rho_1 \mathbf{I}_{\kappa_j}$ et $U_{p+j} \equiv \rho_2 \mathbf{I}_3$, où (τ, ρ_1, ρ_2) sont des constantes positives vérifiant la condition de convergence (7.31) donnée dans la remarque 7.4, obtenue lorsque $l \equiv l_{\{0\}}$:

$$\begin{cases} \frac{\tau}{2} < \min\{(\tilde{\mu}_j)_{1 \leq j \leq p}\} \\ \tau (\rho_1 \|\mathbf{L}\|^2 + \rho_2) < 1, \end{cases} \quad (7.43)$$

où $(\forall l \in \{1, \dots, p\}) \mathbf{L}_l$ est donnée par (7.38) et, $(\forall j \in \{1, \dots, p\}) \tilde{\mu}_j^{-1}$ est la constante de Lipschitz de h_j .

7.4.2.3 Implémentation

Puisque, pour tout $l \in \{1, \dots, p\}$, g_l correspond à une somme de normes ℓ_2 , $\text{prox}_{U_l^{-1}, g_l^*}$ a une forme explicite [Combettes et Pesquet, 2010]. De plus, pour tout $j \in \{1, \dots, p\}$, l'opérateur proximal de f_j relatif à la métrique induite par W_j^{-1} correspond à un simple opérateur de projection sur $[x_{\min}^{(\mathcal{X})}, x_{\max}^{(\mathcal{X})}] \times [x_{\min}^{(\mathcal{Y})}, x_{\max}^{(\mathcal{Y})}] \times [x_{\min}^{(\mathcal{Z})}, x_{\max}^{(\mathcal{Z})}]$.

Remarquons que, pour les deux algorithmes, les boucles internes sur les $l \in \{1, \dots, q\}$ et les $j \in \{1, \dots, q\}$ peuvent être calculées en parallèle. D'autre part, lorsque l'on traite de très grandes quantités de données, nous risquons d'être rapidement limités en terme de mémoire, et le temps de calcul peut devenir secondaire. Cette limitation de mémoire implique que seule une partie des variables globales peuvent être manipulées à chaque itération de l'algorithme.

7.4.3 Résultats numériques

7.4.3.1 Premier exemple

Dans ce premier exemple numérique, nous avons utilisé l'algorithme 7.13 pour résoudre le problème décrit dans la section 7.4.1. Dans nos simulations nous avons utilisé le maillage surfacique *Bunny*, disponible sur <http://graphics.stanford.edu/data/3Dscanrep/>, composé de $p = 8171$ nœuds, illustré en Figure 7.1(a). Le maillage bruité $\mathbf{z}$, représenté dans la Figure 7.1(b), résulte du modèle d'observation décrit dans la section 7.4.1. Plus précisément, nous définissons deux ensembles de nœuds $\mathcal{N}_1$ et $\mathcal{N}_2$, de cardinaux respectifs N_1 et N_2 , vérifiant $N_1 + N_2 = p$. On suppose que les éléments de $(\mathbf{b}_j)_{j \in \mathcal{N}_1}$ modélisent un bruit blanc gaussien de moyenne nulle et d'écart-type $\sigma_1 = 10^{-3}$, tandis que les éléments de $(\mathbf{b}_j)_{j \in \mathcal{N}_2}$ sont indépendants et identiquement distribués selon une loi normale mixte, comprenant deux composantes de mélange pondérées par des poids $(\pi, 1 - \pi)$ ($\pi = 0.98$), de moyennes nulles, et d'écart-types respectifs $\sigma_2 = 5 \times 10^{-3}$ et $\sigma_2' = 1.5 \times 10^{-2}$. Dans notre application, nous avons choisi $\mathcal{N}_2$, de cardinal $N_2 = 1327$, comme étant l'ensemble des nœuds situés au niveau des oreilles et de la queue du lapin. Nous montrons dans la Figure 7.1(c) le maillage restauré obtenu en utilisant l'algorithme 7.9 avec $\lambda_k \equiv 1$. Les paramètres de régularisation $(\beta^{(l)})_{1 \leq l \leq p}$ et $(\delta^{(j)})_{1 \leq j \leq p}$ sont choisis différemment suivant qu'ils correspondent à l'ensemble $\mathcal{N}_1$ ou $\mathcal{N}_2$, de façon à minimiser l'EQM (défini par (2.20)) du maillage restauré. De plus, nous avons comparé le maillage obtenu après reconstruction par la méthode proposée, avec le maillage restauré par une méthode de filtrage de l'état de l'art [Taubin et al., 1996], donné dans la Figure 7.1(d). Notons que, comme il l'a été remarqué dans [Couprie et al.,

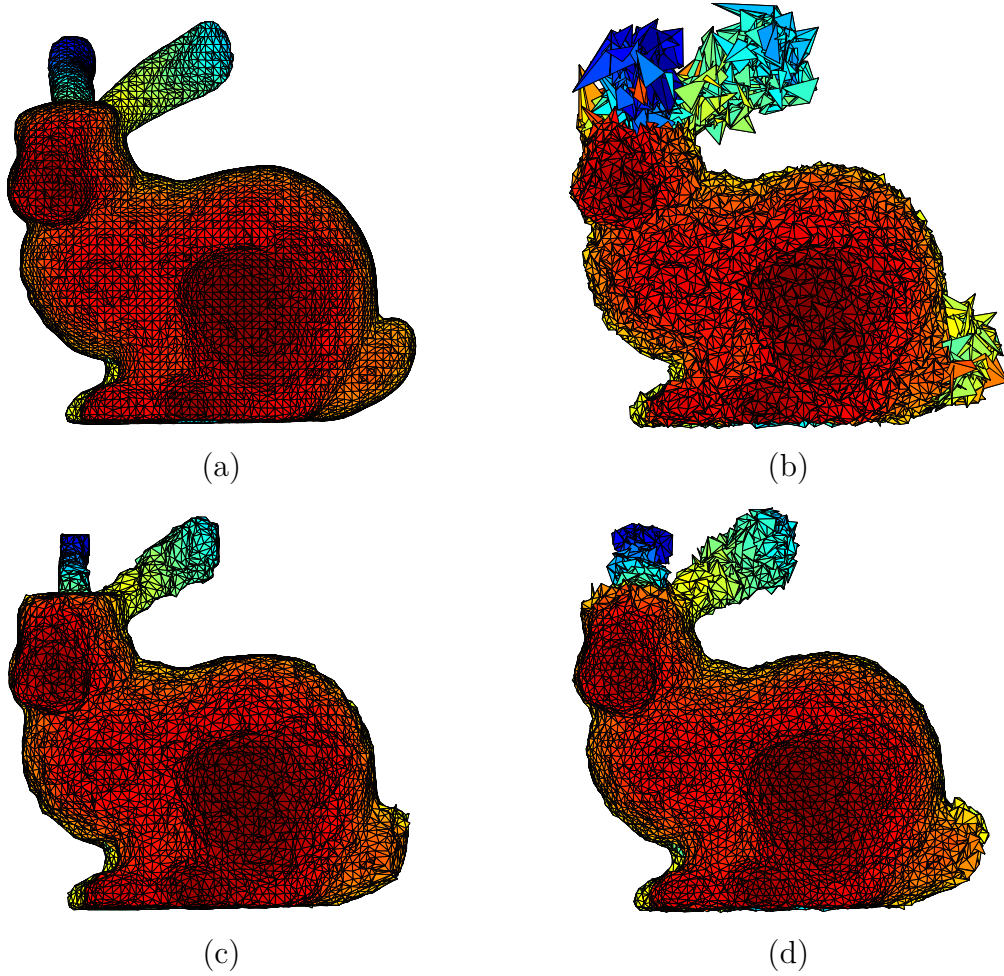


Figure 7.1 – (a) Maillage original $\bar{\mathbf{x}}$, (b) maillage bruité $\mathbf{z}$ avec $EQM = 5.45 \times 10^{-6}$, et maillage reconstruit $\hat{\mathbf{x}}$ en utilisant (c) l'algorithme 7.9 avec $EQM = 8.89 \times 10^{-7}$, et (d) une méthode de lissage du Laplacien avec $EQM = 1.29 \times 10^{-6}$.

2013], l'utilisation d'une méthode régularisée permet d'obtenir une amélioration en terme d'EQM par rapport à la méthode de [Taubin *et al.*, 1996].

L'algorithme 7.9 permet de sélectionner, à chaque itération, de façon flexible, les nœuds qui sont mis à jour. Afin de prendre en compte les variations spatiales du bruit, nous proposons de choisir, pour tout $k \in \mathbb{N}$ et $j \in \{1, \dots, p\}$, $\varepsilon_k^{(j)}$ comme étant une variable aléatoire suivant une loi de Bernoulli de probabilité $\mathbf{p}_j$ telle que $\mathbf{p}_j = 1$ si $j \in \mathcal{N}_2$, et $\mathbf{p}_j = \mathbf{p} \in]0, 1]$ sinon. Remarquons que, pour $\mathbf{p} = 1$, on retrouve l'algorithme déterministe primal-dual de [Condat, 2013; Vũ, 2013]. Pour différentes valeurs de $\mathbf{p}$, on évalue l'indicateur de performance défini par

$$C(\mathbf{p}) = \bar{k} \frac{\mathbf{p}\mathcal{N}_1 + \mathcal{N}_2}{\mathcal{N}_1 + \mathcal{N}_2}, \quad (7.44)$$

où $\bar{k} \in \mathbb{N}^*$ correspond à la première itération de l'algorithme 7.9 pour laquelle

$$\|\mathbf{x}_{\bar{k}} - \mathbf{x}_{\bar{k}-1}\| \leq 10^{-6} \sqrt{3\mathbf{p}}.$$

Pour une probabilité p fixée, $C(p)$ peut être interprété comme le nombre d'itérations nécessaires pour atteindre la convergence, pondéré par le nombre de nœuds mis à jour. Nous montrons dans la Figure 7.2 l'évolution de l'indicateur $C(p)$, moyenné sur dix exécutions de l'algorithme 7.9 (correspondant à dix réalisations différentes des variables aléatoires $(\varepsilon_k)_{k \in \mathbb{N}}$). Nous observons alors que la valeur minimale de $C(p)$ est obtenue pour $p = 0.33$, ce qui montre que la stratégie optimale de l'algorithme 7.9 en terme de profil de convergence est obtenue lorsque seule une fraction des nœuds de l'ensemble $\mathcal{N}_1$ (correspondant aux nœuds dont la position est affectée par le bruit le plus faible) est mise à jour à chaque itération.

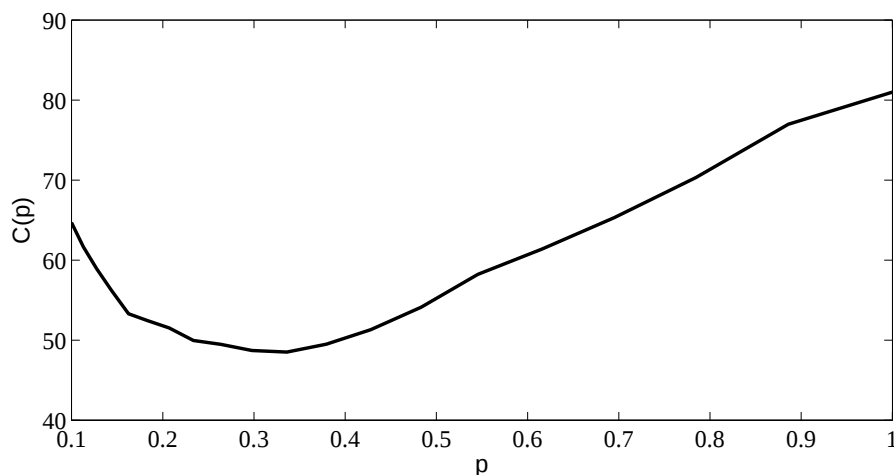


Figure 7.2 – Indicateur $C(p)$ obtenu pour différentes valeurs de la probabilité p (moyenné sur dix exécutions de l'algorithme 7.9).

7.4.3.2 Second exemple

Dans ce second exemple numérique, nous avons utilisé l'algorithme 7.10 pour résoudre le problème décrit dans la section 7.4.1. Nous avons considéré le maillage surfacique **Dragon** de plus grande dimension disponible sur <http://graphics.stanford.edu/data/3Dscanrep/>, composé de $p = 100250$ nœuds, illustré en Figure 7.3(a). Les coordonnées du maillage original sont dégradées par un bruit distribué suivant un mélange de gaussiennes. Le maillage observé est donné dans la figure 7.3(b). Le maillage restauré obtenu en utilisant l'algorithme 7.10, avec $\lambda_k \equiv 1$ est représenté dans la figure 7.3(c). Les paramètres de régularisation $(\beta^{(l)})_{1 \leq l \leq q}$ et $(\delta^{(j)})_{1 \leq j \leq p}$ sont choisis de façon à minimiser l'EQM du maillage restauré. Afin d'avoir un élément de comparaison en terme de qualité de restauration, nous donnons dans la figure 7.3(d) le maillage restauré en utilisant la méthode de [Taubin *et al.*, 1996].

Dans cette application, nous divisons le maillage en une partition de p/r blocs, où chaque bloc est composé de $r \leq p$ nœuds. Nous choisissons, à chaque itération $k \in \mathbb{N}$, les variables booléennes $(\varepsilon_k^{(j)})_{1 \leq j \leq p}$ de façon à ce que seuls les éléments liés aux r nœuds sélectionnés soient mis à jour à la k -ème itération. Ainsi, $(\varepsilon_k^{(j)})_{1 \leq j \leq p}$ satisfait $\sum_{j=1}^p \varepsilon_k^{(j)} = r$.

La figure 7.4 représente à la fois le temps nécessaire pour atteindre le critère d'arrêt $\|x_k - \hat{x}\| \leq 10^{-3} \|\hat{x}\|$, où $\hat{x}$ a été calculé au préalable pour un grand nombre d'itérations, et la mémoire

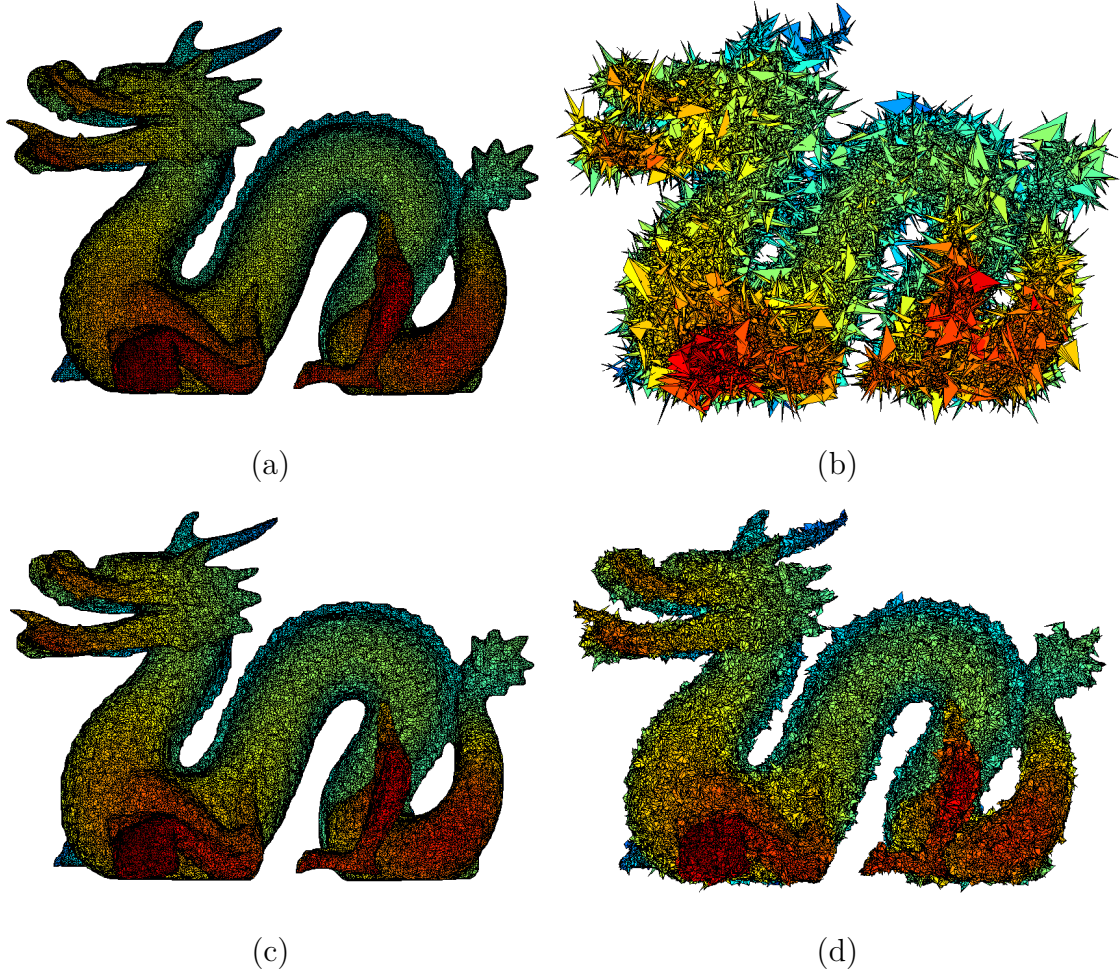


Figure 7.3 – (a) Maillage original $\bar{\mathbf{x}}$, (b) maillage bruité $\mathbf{z}$ avec $EQM = 2.89 \times 10^{-6}$, et maillage reconstruit $\hat{\mathbf{x}}$ en utilisant (c) l'algorithme 7.10 avec $EQM = 8.09 \times 10^{-8}$, et (d) une méthode de lissage du Laplacien avec $EQM = 5.23 \times 10^{-7}$.

nécessaire pour l'exécution de l'algorithme avec Matlab R2013a, en fonction du nombre de blocs p/r . Nous pouvons remarquer que le temps minimum de reconstruction (égal à 5s.) est obtenu lorsque $p = 1$, i.e. avec une stratégie déterministe. Cependant, choisir des blocs plus petits permet de réduire l'occupation mémoire, au dépend d'une vitesse d'exécution plus lente. L'avantage de la méthode présentée ici pour des données de grande dimension, est qu'elle permet à l'utilisateur de choisir la stratégie qui lui permettra d'obtenir un temps de reconstruction minimal, sans dépasser la capacité de mémoire disponible sur sa propre architecture d'ordinateur.

Remarque 7.6. Le diagramme donné dans la figure 7.5 résume les liens existants entre les différents algorithmes développés dans les sections 7.3 et 7.4. Bien que seul l'algorithme 7.9 ait été présenté avec un choix simplifié des variables booléennes $(\epsilon_k)_{k \in \mathbb{N}}$, un choix similaire peut être appliqué aux autres algorithmes. Les cadres verts renvoient aux problèmes traités (problèmes 7.1

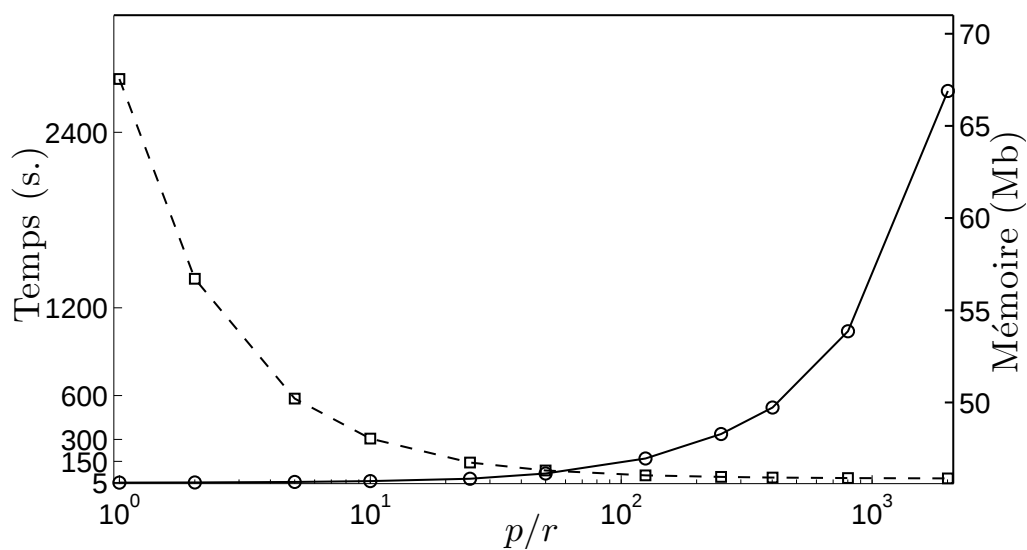


Figure 7.4 — Temps de reconstruction (cercles) et mémoire nécessaire (carrés) pour différents nombres de blocs p/r .

et 7.2) ainsi qu'à leurs cas particuliers (vert clair), et les ellipses bleues donnent les algorithmes permettant de les traiter.

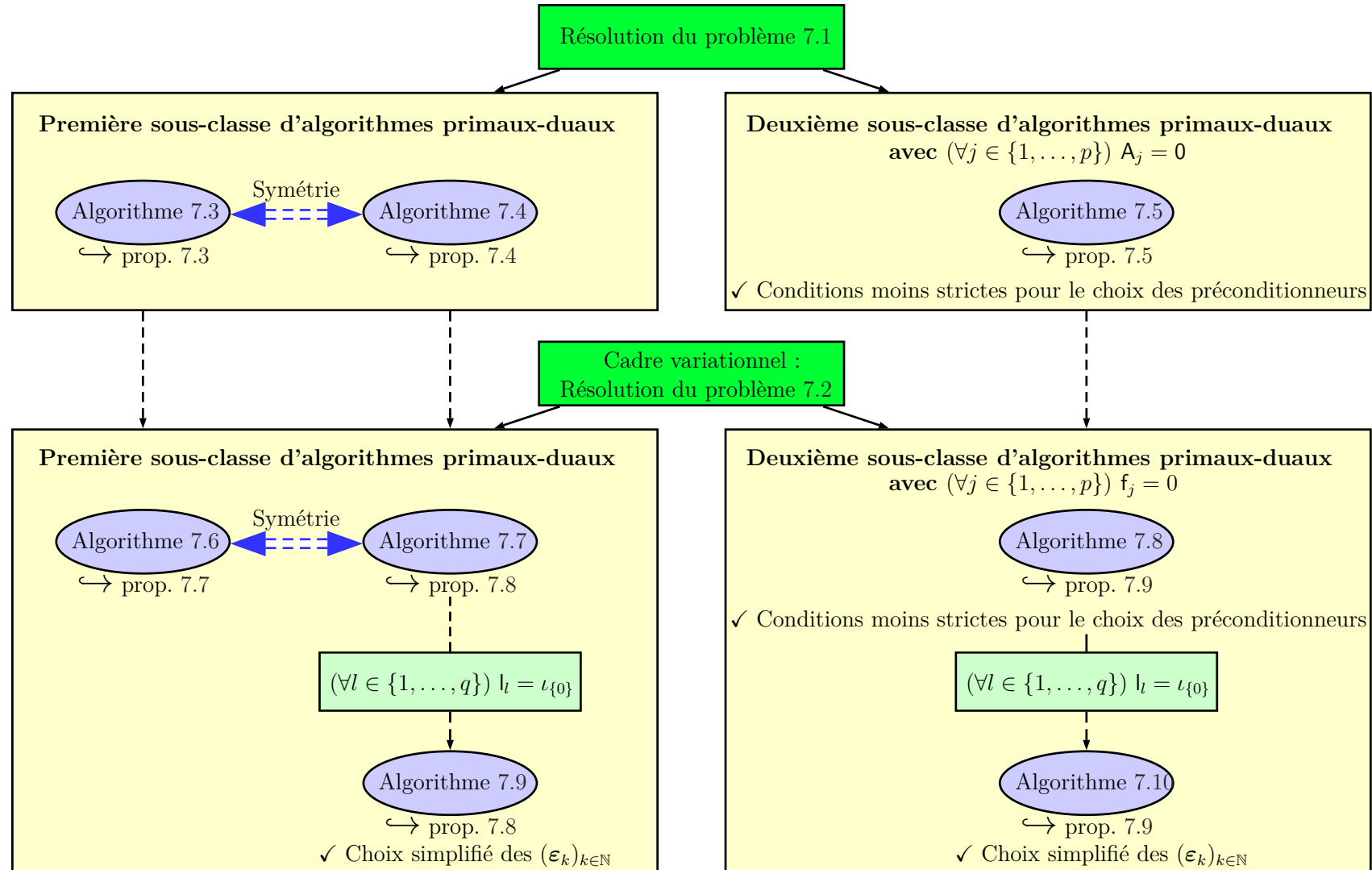


Figure 7.5 – Diagramme résumant les algorithmes primal-duaux alternés par bloc développés dans les sections 7.3 et 7.4 pour résoudre les problèmes 7.1 et 7.2.

7.5 ALGORITHMES DISTRIBUÉS

7.5.1 Problème et notations

Dans cette section, $\mathcal{H}, \mathcal{G}_1, \dots, \mathcal{G}_m$ sont des espaces de Hilbert réels séparables, et on note $\mathcal{G} = \mathcal{G}_1 \oplus \dots \oplus \mathcal{G}_m$. On considère le problème suivant :

Problème 7.3. *Pour tout $i \in \{1, \dots, m\}$, soient $A_i: \mathcal{H} \rightarrow 2^{\mathcal{H}}$ un opérateur maximal monotone, $C_i: \mathcal{H} \rightarrow \mathcal{H}$ un opérateur cocoercif, $B_i: \mathcal{G}_i \rightarrow 2^{\mathcal{G}_i}$ un opérateur maximal monotone, $D_i: \mathcal{G}_i \rightarrow 2^{\mathcal{G}_i}$ un opérateur maximal et fortement monotone, et M_i un opérateur non nul de $\mathcal{B}(\mathcal{H}, \mathcal{G}_i)$. On suppose que l'ensemble $\hat{F}$ de solutions du problème :*

$$\text{trouver } \mathbf{x} \in \mathcal{H} \text{ tel que } 0 \in \sum_{i=1}^m A_i \mathbf{x} + C_i \mathbf{x} + M_i^*(B_i \square D_i)(M_i \mathbf{x}) \quad (7.45)$$

est non vide. Notre objectif est de trouver une variable aléatoire $\hat{x}$ à valeurs dans $\hat{F}$.

Afin d'obtenir des algorithmes parallèles servant à trouver le zéro d'une somme d'opérateurs maximaux monotones [Combettes et Pesquet, 2008; Pesquet et Pustelnik, 2012], ou bien pour traiter des problèmes de consensus [Boyd et al., 2011; Nedić et Ozdaglar, 2010], le problème (7.45) peut être reformulé dans l'espace produit $\mathcal{H}^m$ comme suit

$$\text{trouver } (\mathbf{x}^{(1)}, \dots, \mathbf{x}^{(m)}) \in \Lambda_m \text{ tel que } 0 \in \sum_{i=1}^m A_i \mathbf{x}^{(i)} + C_i \mathbf{x}^{(i)} + M_i^*(B_i \square D_i)(M_i \mathbf{x}^{(i)}) \quad (7.46)$$

où

$$\Lambda_m = \{(\mathbf{x}^{(1)}, \dots, \mathbf{x}^{(m)}) \in \mathcal{H}^m \mid \mathbf{x}^{(1)} = \dots = \mathbf{x}^{(m)}\}.$$

Pour construire des algorithmes distribués, la contrainte linéaire apparaissant dans le problème (7.46) est divisée en un ensemble de contraintes similaires, chacune d'entre elles faisant intervenir un sous ensemble réduit de variables. Chaque indice $i \in \{1, \dots, m\}$ correspond alors à un agent distinct, et il s'agit de modéliser les relations topologiques existantes entre chaque agent. Pour ce faire, on définit des sous-ensembles non vides $(\mathbb{V}_\ell)_{1 \leq \ell \leq r}$ de $\{1, \dots, m\}$, de cardinaux $(\kappa^{(\ell)})_{1 \leq \ell \leq r}$, vérifiant l'hypothèse suivante :

Hypothèse 7.1. *Pour tout $\mathbf{x} = (\mathbf{x}^{(i)})_{1 \leq i \leq m} \in \mathcal{H}^m$,*

$$\mathbf{x} \in \Lambda_m \quad \Leftrightarrow \quad (\forall \ell \in \{1, \dots, r\}) \quad (\mathbf{x}^{(i)})_{i \in \mathbb{V}_\ell} \in \Lambda_{\kappa^{(\ell)}}.$$

Remarque 7.7. *Notons que l'hypothèse 7.1 est satisfaite dans les cas particuliers suivants*

- si $r = 1$ et $\mathbb{V}_1 = \{1, \dots, m\}$,
- si $r = m - 1$ et $(\forall \ell \in \{1, \dots, m - 1\}) \mathbb{V}_\ell = \{\ell, \ell + 1\}$.

De manière plus générale, si les ensembles $(\mathbb{V}_\ell)_{1 \leq \ell \leq r}$ correspondent aux hyper-arêtes d'un hypergraphe de sommets $\{1, \dots, m\}$, alors l'hypothèse 7.1 est équivalente à supposer que l'hypergraphe est connecté.

Pour la suite de cette section, nous avons besoin d'introduire les notations suivantes :

$$\begin{aligned}
 \mathcal{H} &= \mathcal{H}^{\kappa(1)} \oplus \cdots \oplus \mathcal{H}^{\kappa(r)}, & \mathbf{\Lambda} &= \mathbf{\Lambda}_{\kappa(1)} \oplus \cdots \oplus \mathbf{\Lambda}_{\kappa(r)}, \\
 \mathbf{A} &= \bigtimes_{i=1}^m \mathbf{A}_i, & \mathbf{C} &= \bigtimes_{i=1}^m \mathbf{C}_i, \\
 \mathbf{B} &= \bigtimes_{i=1}^m \mathbf{B}_i, & \mathbf{D} &= \bigtimes_{i=1}^m \mathbf{D}_i, \\
 \mathbf{S}: \mathcal{H}^m &\rightarrow \mathcal{H}: \mathbf{x} \mapsto (\mathbf{S}_\ell \mathbf{x})_{1 \leq \ell \leq r}, & \mathbf{M}: \mathcal{H}^m &\rightarrow \mathcal{G}: \mathbf{x} \mapsto (\mathbf{M}_i \mathbf{x}^{(i)})_{1 \leq i \leq m},
 \end{aligned}$$

où, pour tout $\ell \in \{1, \dots, r\}$,

$$\mathbf{S}_\ell: \mathcal{H}^m \rightarrow \mathcal{H}^{\kappa(\ell)}: \mathbf{x} \mapsto (\mathbf{x}^{(i)})_{i \in \mathbb{V}_\ell} = (\mathbf{x}^{(i(\ell, j))})_{1 \leq j \leq \kappa(\ell)} \quad (7.47)$$

et $i(\ell, 1), \dots, i(\ell, \kappa(\ell))$ désignent les éléments de $\mathbb{V}_\ell$ ordonnés dans un ordre croissant. Notons que, pour tout $\ell \in \{1, \dots, r\}$, l'opérateur adjoint de $\mathbf{S}_\ell$ est donné par

$$\mathbf{S}_\ell^*: \mathcal{H}^{\kappa(\ell)} \rightarrow \mathcal{H}^m: \mathbf{z}^{(\ell)} = (\mathbf{z}^{(\ell, j)})_{1 \leq j \leq \kappa(\ell)} \mapsto (\mathbf{x}^{(i)})_{1 \leq i \leq m} \quad (7.48)$$

où

$$(\forall i \in \{1, \dots, m\}) \quad \mathbf{x}^{(i)} = \begin{cases} \mathbf{z}^{(\ell, j)} & \text{si } i = i(\ell, j) \text{ avec } j \in \{1, \dots, \kappa(\ell)\} \\ 0 & \text{sinon.} \end{cases} \quad (7.49)$$

L'opérateur adjoint de $\mathbf{S}$ est alors donné par

$$\mathbf{S}^*: \mathcal{H} \rightarrow \mathcal{H}^m: (\mathbf{z}^{(\ell)})_{1 \leq \ell \leq r} \mapsto \sum_{\ell=1}^r \mathbf{S}_\ell^* \mathbf{z}^{(\ell)} = (\mathbf{x}^{(i)})_{1 \leq i \leq m} \quad (7.50)$$

où, pour tout $i \in \{1, \dots, m\}$,

$$\mathbf{x}^{(i)} = \sum_{(\ell, j) \in \mathbb{V}_i^*} \mathbf{z}^{(\ell, j)} \quad (7.51)$$

avec

$$\mathbb{V}_i^* = \{(\ell, j) \in \{1, \dots, r\} \times \{1, \dots, \kappa(\ell)\} \mid i(\ell, j) = i\}.$$

En conséquence de l'hypothèse 7.1, le cardinal de $\mathbb{V}_i^*$ (correspondant au nombre d'ensembles $(\mathbb{V}_\ell)_{1 \leq \ell \leq r}$ qui contiennent l'indice i) est non nul.

Un exemple d'hypergraphe connecté est donné dans la figure 7.6. Il est composé de $m = 6$ nœuds (ou agents) et $r = 4$ hyper-arêtes modélisant les relations topologiques entre ces agents.

Le lien entre les problèmes 7.3 et 7.1 est mis en évidence dans la proposition suivante :

Proposition 7.10. *Sous l'hypothèse 7.1, le problème (7.46) est équivalent à*

$$\text{trouver } \mathbf{x} \in \mathcal{H}^m \text{ tel que } \mathbf{0} \in \mathbf{A}\mathbf{x} + \mathbf{C}\mathbf{x} + \mathbf{M}^*(\mathbf{B} \square \mathbf{D})(\mathbf{M}\mathbf{x}) + \mathbf{S}^*\mathbf{N}_\Lambda(\mathbf{S}\mathbf{x}). \quad (7.52)$$

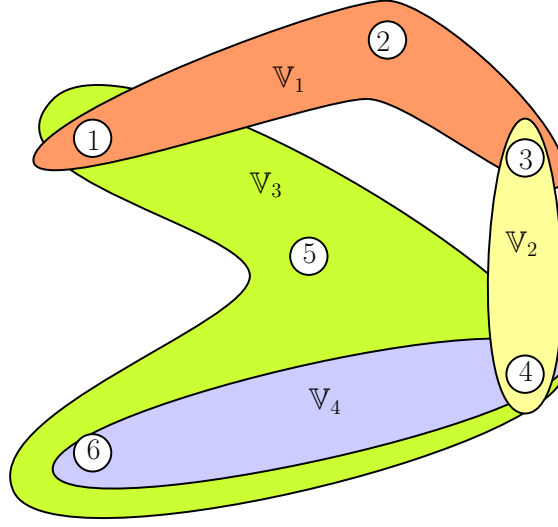


Figure 7.6 – Exemple d'hypergraphe connecté avec $r = 4$ et $m = 6$. Ici, par exemple, le nœud 3 est à la fois dans les hyper-arêtes $\mathbb{V}_1 = \{i(1,1), i(1,2), i(1,3)\}$ et $\mathbb{V}_2 = \{i(2,1), i(2,2)\}$. De plus, les voisins du nœud $3 = i(1,3) = i(2,1)$ sont les nœuds $1 = i(1,1) = i(3,1)$, $2 = i(1,2)$ et $4 = i(2,2) = i(3,2) = i(4,1)$, et on a $\mathbb{V}_3^* = \{(1,3), (2,1)\}$.

Démonstration. Pour tout $\mathbf{x} \in \mathcal{H}^m$, on a les équivalences suivantes :

$$\begin{aligned}
 & \begin{cases} 0 \in \sum_{i=1}^m \mathbf{A}_i \mathbf{x}^{(i)} + \mathbf{C}_i \mathbf{x}^{(i)} + \mathbf{M}_i^* (\mathbf{B}_i \square \mathbf{D}_i) (\mathbf{M}_i \mathbf{x}^{(i)}) \\ \mathbf{x} \in \Lambda_m \end{cases} \\
 \Leftrightarrow & \begin{cases} (\forall i \in \{1, \dots, m\}) \quad 0 \in \mathbf{A}_i \mathbf{x}^{(i)} + \mathbf{C}_i \mathbf{x}^{(i)} + \mathbf{M}_i^* (\mathbf{B}_i \square \mathbf{D}_i) (\mathbf{M}_i \mathbf{x}^{(i)}) + \mathbf{u}^{(i)} \\ \mathbf{x} \in \Lambda_m \\ \sum_{i=1}^m \mathbf{u}^{(i)} = 0 \end{cases} \\
 \Leftrightarrow & \begin{cases} \mathbf{0} \in \mathbf{A} \mathbf{x} + \mathbf{C} \mathbf{x} + \mathbf{M}^* (\mathbf{B} \square \mathbf{D}) (\mathbf{M} \mathbf{x}) + \mathbf{u} \\ \mathbf{x} \in \Lambda_m \\ \mathbf{u} \in \Lambda_m^\perp \end{cases} \\
 \Leftrightarrow & \mathbf{0} \in \mathbf{A} \mathbf{x} + \mathbf{C} \mathbf{x} + \mathbf{M}^* (\mathbf{B} \square \mathbf{D}) (\mathbf{M} \mathbf{x}) + \mathbf{S}^* \mathbf{N}_\Lambda (\mathbf{S} \mathbf{x}),
 \end{aligned}$$

où l'on a utilisé le fait que $\mathbf{N}_{\Lambda_m} = \partial \iota_{\Lambda_m} = \partial (\iota_\Lambda \circ \mathbf{S}) = \mathbf{S}^* \partial \iota_\Lambda \mathbf{S} = \mathbf{S}^* \mathbf{N}_\Lambda \mathbf{S}$ puisque $\Lambda + \text{ran}(\mathbf{S})$ est un sous-espace fermé de $\mathcal{H}$ [Bauschke et Combettes, 2011, Prop. 6.19 & 16.42]. $\square$

Remarque 7.8. Si nous reformulons (7.52) avec les notations du problème 7.1, nous pouvons voir que les problèmes 7.3 et 7.1 sont équivalents en posant $p = m$, $q = m+r$, $\mathcal{H}_1 = \dots = \mathcal{H}_m = \mathcal{H}$,

$(\forall \ell \in \{1, \dots, r\}) \mathcal{G}_{m+\ell} = \mathcal{H}^{\kappa^{(\ell)}}, \mathbf{B}_{m+\ell} = \mathbf{N}_{\Lambda_{\kappa^{(\ell)}}}, \mathbf{D}_{m+\ell} = \mathbf{N}_{\{0\}}, \text{ et}$

$$(\forall (k, i) \in \{1, \dots, m\}^2) \quad \mathbf{L}_{k,i} = \begin{cases} \mathbf{M}_i & \text{si } k = i \\ 0 & \text{sinon,} \end{cases}$$

$$(\forall \ell \in \{1, \dots, r\})(\forall \mathbf{x} \in \mathcal{H}^m) \quad \sum_{i=1}^m \mathbf{L}_{m+\ell,i} \mathbf{x}^{(i)} = \mathbf{S}_\ell \mathbf{x}$$

(ainsi, (7.8) et (7.9) sont satisfaites).

7.5.2 Algorithme distribués asynchrones pour les problèmes d'inclusion monotone

Notre objectif est maintenant de résoudre le problème 7.3. Pour cela, nous allons développer des algorithmes distribués asynchrones au sens où, à chaque itération, seul un nombre limité d'opérateurs $(\mathbf{A}_i)_{1 \leq i \leq m}, (\mathbf{B}_i)_{1 \leq i \leq m}, (\mathbf{C}_i)_{1 \leq i \leq m}, \text{ et } (\mathbf{D}_i)_{1 \leq i \leq m}$ est activé, et ce, de façon aléatoire. Nous définissons donc l'algorithme 7.11.

La convergence faible presque sûre de l'algorithme 7.11 vers une solution du problème 7.3 est assurée par la proposition 7.11. Ce résultat est obtenu en se basant sur la remarque 7.8 et en utilisant l'algorithme 7.4 et la proposition 7.4.

Proposition 7.11. *Considérons l'algorithme 7.11. Supposons que, pour tout $i \in \{1, \dots, m\}$, $\mathbf{W}_i^{1/2} \mathbf{C}_i \mathbf{W}_i^{1/2}$ est μ_i -cocoercif, avec $\mu_i \in]0, +\infty[$, et $\mathbf{U}_i^{1/2} \mathbf{D}_i^{-1} \mathbf{U}_i^{1/2}$ est ν_i -cocoercif, avec $\nu_i \in]0, +\infty[$. Posons*

$$\bar{\theta}_i = \sum_{\ell \in \{\ell' \in \{1, \dots, r\} \mid i \in \mathbb{V}_{\ell'}\}} \theta^{(\ell)}. \quad (7.53)$$

Supposons que

$$(\exists \alpha \in]0, +\infty[) \quad (1 - \chi) \min\{\mu(1 + \alpha\sqrt{\chi})^{-1}, \nu(1 + \alpha^{-1}\sqrt{\chi})^{-1}\} > \frac{1}{2} \quad (7.54)$$

où

$$\chi = \max_{i \in \{1, \dots, m\}} \|\mathbf{U}_i^{1/2} \mathbf{M}_i \mathbf{W}_i^{1/2}\|^2 + \bar{\theta}_i \|\mathbf{W}_i\|, \quad (7.55)$$

avec $\mu = \min\{\mu_1, \dots, \mu_m\}$, et $\nu = \min\{\nu_1, \dots, \nu_m\}$. Supposons que les éléments de la suite $(\lambda_k)_{k \in \mathbb{N}}$ soient dans $]0, 1]$ et vérifient $\inf_{k \in \mathbb{N}} \lambda_k > 0$. Posons $(\forall k \in \mathbb{N}) \mathcal{E}_k = \sigma(\varepsilon_k)$ et $\mathcal{X}_k = \sigma(\mathbf{x}_{k'}, \mathbf{v}_{k'})_{0 \leq k' \leq k}$, et supposons de plus que les assertions suivantes sont vérifiées :

- (i) $\sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{a}_k\|^2 | \mathcal{X}_k)} < +\infty$, $\sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{b}_k\|^2 | \mathcal{X}_k)} < +\infty$, $\sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{c}_k\|^2 | \mathcal{X}_k)} < +\infty$,
et $\sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{d}_k\|^2 | \mathcal{X}_k)} < +\infty$ P-presque sûrement.
- (ii) Pour tout $k \in \mathbb{N}$, $\mathcal{E}_k$ et $\mathcal{X}_k$ sont indépendants, et $(\forall i \in \{1, \dots, m\}) \mathbb{P}[\varepsilon_0^{(i)} = 1] > 0$.
- (iii) Pour tout $k \in \mathbb{N}$,

$$(\forall i \in \{1, \dots, m\}) \quad \{\omega \in \Omega \mid \varepsilon_k^{(i)}(\omega) = 1\} \subset \{\omega \in \Omega \mid \varepsilon_k^{(m+i)}(\omega) = 1\} \quad (7.56)$$

Algorithme 7.11 Algorithme distribué asynchrone (version 1)

Initialisation : Soient $\mathbf{x}_0$, $(\mathbf{a}_k)_{k \in \mathbb{N}}$ et $(\mathbf{c}_k)_{k \in \mathbb{N}}$ des variables aléatoires à valeurs dans $\mathcal{H}^m$, et soient $(v_0^{(i)})_{1 \leq i \leq m}$, $(\mathbf{b}_k)_{k \in \mathbb{N}}$, et $(\mathbf{d}_k)_{k \in \mathbb{N}}$ des variables aléatoires à valeurs dans $\mathcal{G}$.

Pour tout $\ell \in \{1, \dots, r\}$, soit $v_0^{(m+\ell)} = (v_0^{(m+\ell, j)})_{1 \leq j \leq \kappa^{(\ell)}}$ une variable aléatoire à valeurs dans $\mathcal{H}^{\kappa^{(\ell)}}$, $\bar{x}_0^{(\ell)} = \frac{1}{\kappa^{(\ell)}} \sum_{i \in \mathbb{V}_\ell} x_0^{(i)}$, et $\theta^{(\ell)} \in]0, +\infty[$.

Soient, pour tout $i \in \{1, \dots, m\}$, $\mathbf{W}_i \in \mathcal{S}^+(\mathcal{H})$ et $\mathbf{U}_i \in \mathcal{S}^+(\mathcal{G}_i)$.

Soient $(\varepsilon_k)_{k \in \mathbb{N}}$ des variables aléatoires identiquement distribuées à valeurs dans $\mathbb{D}_{2m+r}$.

Soit, pour tout $k \in \mathbb{N}$, $\lambda_k \in]0, +\infty[$.

Itérations :

Pour $k = 0, 1, \dots$

pour $\ell = 1, \dots, r$ $\bar{u}_k^{(\ell)} = \varepsilon_k^{(2m+\ell)} \left(\frac{1}{\kappa^{(\ell)}} \sum_{j=1}^{\kappa^{(\ell)}} v_k^{(m+\ell, j)} + \theta^{(\ell)} \bar{x}_k^{(\ell)} \right),$ pour $j = 1, \dots, \kappa^{(\ell)}$ $w_k^{(\ell, j)} = \varepsilon_k^{(2m+\ell)} \left(2(\theta^{(\ell)} x_k^{(i(\ell, j))}) - \bar{u}_k^{(\ell)} + v_k^{(m+\ell, j)} \right),$ pour $i = 1, \dots, m$ $u_k^{(i)} = \varepsilon_k^{(m+i)} \left(\mathbf{J}_{\mathbf{U}_i \mathbf{B}_i^{-1}} \left(v_k^{(i)} + \mathbf{U}_i \left(\mathbf{M}_i x_k^{(i)} - \mathbf{D}_i^{-1} v_k^{(i)} + d_k^{(i)} \right) \right) + b_k^{(i)} \right),$ $y_k^{(i)} = \varepsilon_k^{(i)} \left(\mathbf{J}_{\mathbf{W}_i \mathbf{A}_i} \left(x_k^{(i)} - \mathbf{W}_i \left(\mathbf{M}_i^* (2u_k^{(i)} - v_k^{(i)}) + \sum_{(\ell, j) \in \mathbb{V}_i^*} w_k^{(\ell, j)} + \mathbf{C}_i x_k^{(i)} + c_k^{(i)} \right) \right) + a_k^{(i)} \right),$ $v_{k+1}^{(i)} = v_k^{(i)} + \lambda_k \varepsilon_k^{(m+i)} (u_k^{(i)} - v_k^{(i)}),$ $x_{k+1}^{(i)} = x_k^{(i)} + \lambda_k \varepsilon_k^{(i)} (y_k^{(i)} - x_k^{(i)}),$ pour $\ell = 1, \dots, r$ $v_{k+1}^{(m+\ell)} = v_k^{(m+\ell)} + \frac{\lambda_k}{2} \varepsilon_k^{(2m+\ell)} (w_k^{(\ell)} - v_k^{(m+\ell)}),$ $\eta_k^{(\ell)} = \max \left\{ \varepsilon_k^{(i)} \mid i \in \mathbb{V}_\ell \right\},$ $\bar{x}_{k+1}^{(\ell)} = \bar{x}_k^{(\ell)} + \eta_k^{(\ell)} \left(\frac{1}{\kappa^{(\ell)}} \sum_{i \in \mathbb{V}_\ell} x_{k+1}^{(i)} - \bar{x}_k^{(\ell)} \right).$	
---	--

et

$$(\forall \ell \in \{1, \dots, r\}) \quad \bigcup_{i \in \mathbb{V}_\ell} \{ \omega \in \Omega \mid \varepsilon_k^{(i)}(\omega) = 1 \} \subset \{ \omega \in \Omega \mid \varepsilon_k^{(2m+\ell)}(\omega) = 1 \}. \quad (7.57)$$

Sous l'hypothèse 7.1, pour tout $i \in \{1, \dots, m\}$, $(x_k^{(i)})_{k \in \mathbb{N}}$ converge faiblement P-presque sûrement vers une variable aléatoire $\hat{x}$ à valeurs dans $\hat{\mathbf{F}}$ et, pour tout $\ell \in \{1, \dots, r\}$, $(\bar{x}_k^{(\ell)})_{k \in \mathbb{N}}$ converge faiblement P-presque sûrement vers $\hat{x}$.

La démonstration de la proposition 7.11 est donnée dans l'annexe 7.E.

Remarque 7.9.

- (i) La k -ème itération de l'algorithme (7.11) est composé de deux types d'opérations : d'une part, la mise à jour de certaines des variables $(x_k^{(i)})_{1 \leq i \leq m}$ et $(v_k^{(i)})_{1 \leq i \leq m}$ en utilisant les opérateurs $(J_{W_i A_i})_{1 \leq i \leq m}$, $(J_{U_i B_i^{-1}})_{1 \leq i \leq m}$, $(C_i)_{1 \leq i \leq m}$, et $(D_i^{-1})_{1 \leq i \leq m}$, d'autre part, des étapes de fusion effectuées sur les ensembles $(V_\ell)_{1 \leq \ell \leq r}$. Dans ce contexte, un choix simple pour que les variables aléatoires booléennes $(\varepsilon_k^{(l)})_{m+1 \leq l \leq 2m+r}$ satisfassent la condition (iii) est de prendre, pour tout $k \in \mathbb{N}$,

$$(\forall i \in \{1, \dots, m\}) \quad \varepsilon_k^{(m+i)} = \varepsilon_k^{(i)}, \quad (7.58)$$

$$(\forall \ell \in \{1, \dots, r\}) \quad \varepsilon_k^{(2m+\ell)} = \eta_k^{(\ell)} = \max \{ \varepsilon_k^{(i)} \mid i \in V_\ell \}. \quad (7.59)$$

- (ii) D'après (7.92), nous pouvons remarquer que, pour tout $k \in \mathbb{N}$ et $\ell \in \{1, \dots, r\}$, $\Pi_{\Lambda_{\kappa^{(\ell)}}} u_k^{(m+\ell)} = 0$, ce qui implique que la relation de récursivité suivante est satisfaite :

$$\sum_{j=1}^{\kappa^{(\ell)}} v_{k+1}^{(m+\ell, j)} = (1 - \lambda_k \varepsilon_k^{(2m+\ell)}) \sum_{j=1}^{\kappa^{(\ell)}} v_k^{(m+\ell, j)}. \quad (7.60)$$

En particulier, si l'initialisation $(v_0^{(m+\ell)})_{1 \leq \ell \leq r}$ est choisie de façon à satisfaire la condition suivante :

$$(\forall \ell \in \{1, \dots, r\}) \quad \sum_{j=1}^{\kappa^{(\ell)}} v_0^{(m+\ell, j)} = 0, \quad (7.61)$$

alors l'algorithme 7.11 peut se simplifier et on obtient l'algorithme 7.12.

- (iii) De la même manière que pour la remarque 7.3(iii), une condition suffisante pour que (7.54) soit satisfaite est obtenue en posant $\alpha = 1$:

$$(1 - \sqrt{\chi}) \min\{\mu, \nu\} > \frac{1}{2}. \quad (7.62)$$

- (iv) Lorsque, pour tout $i \in \{1, \dots, m\}$, $D_i^{-1} = 0$, une condition moins contraignante est de prendre

$$(1 - \chi) \mu > \frac{1}{2}. \quad (7.63)$$

Algorithme 7.12 Algorithme distribué asynchrone (version 1 simplifiée)

Initialisation : Soient $\mathbf{x}_0$, $(\mathbf{a}_k)_{k \in \mathbb{N}}$ et $(\mathbf{c}_k)_{k \in \mathbb{N}}$ des variables aléatoires à valeurs dans $\mathcal{H}^m$, et soient $(v_0^{(i)})_{1 \leq i \leq m}$, $(\mathbf{b}_k)_{k \in \mathbb{N}}$, et $(\mathbf{d}_k)_{k \in \mathbb{N}}$ des variables aléatoires à valeurs dans $\mathcal{G}$.

Pour tout $\ell \in \{1, \dots, r\}$, soit $v_0^{(m+\ell)} = (v_0^{(m+\ell, j)})_{1 \leq j \leq \kappa^{(\ell)}}$ une variable aléatoire à valeurs dans $\mathcal{H}^{\kappa^{(\ell)}}$ vérifiant $\sum_{j=1}^{\kappa^{(\ell)}} v_0^{(m+\ell, j)} = 0$, soit $\bar{x}_0^{(\ell)} = \frac{1}{\kappa^{(\ell)}} \sum_{i \in \mathbb{V}_\ell} x_0^{(i)}$, et soit $\theta^{(\ell)} \in]0, +\infty[$.

Soient, pour tout $i \in \{1, \dots, m\}$, $\mathbf{W}_i \in \mathcal{S}^+(\mathcal{H})$ et $\mathbf{U}_i \in \mathcal{S}^+(\mathcal{G}_i)$.

Soient $(\varepsilon_k)_{k \in \mathbb{N}}$ des variables aléatoires identiquement distribuées à valeurs dans $\mathbb{D}_{2m+r}$.

Soit, pour tout $k \in \mathbb{N}$, $\lambda_k \in]0, +\infty[$.

Itérations :

Pour $k = 0, 1, \dots$

pour $\ell = 1, \dots, r$

pour $j = 1, \dots, \kappa^{(\ell)}$

$$w_k^{(\ell, j)} = \varepsilon_k^{(2m+\ell)} \left(2\theta^{(\ell)} (x_k^{(i(\ell, j))}) - \bar{x}_k^{(\ell)} \right) + v_k^{(m+\ell, j)},$$

pour $i = 1, \dots, m$

$$u_k^{(i)} = \varepsilon_k^{(m+i)} \left(\mathbf{J}_{\mathbf{U}_i \mathbf{B}_i^{-1}} \left(v_k^{(i)} + \mathbf{U}_i \left(\mathbf{M}_i x_k^{(i)} - \mathbf{D}_i^{-1} v_k^{(i)} + \mathbf{d}_k^{(i)} \right) \right) + \mathbf{b}_k^{(i)} \right),$$

$$y_k^{(i)} = \varepsilon_k^{(i)} \left(\mathbf{J}_{\mathbf{W}_i \mathbf{A}_i} \left(x_k^{(i)} - \mathbf{W}_i \left(\mathbf{M}_i^* (2u_k^{(i)} - v_k^{(i)}) + \sum_{(\ell, j) \in \mathbb{V}_i^*} w_k^{(\ell, j)} + \mathbf{C}_i x_k^{(i)} + \mathbf{c}_k^{(i)} \right) \right) + \mathbf{a}_k^{(i)} \right),$$

$$v_{k+1}^{(i)} = v_k^{(i)} + \lambda_k \varepsilon_k^{(m+i)} (u_k^{(i)} - v_k^{(i)}),$$

$$x_{k+1}^{(i)} = x_k^{(i)} + \lambda_k \varepsilon_k^{(i)} (y_k^{(i)} - x_k^{(i)}),$$

pour $\ell = 1, \dots, r$

$$v_{k+1}^{(m+\ell)} = v_k^{(m+\ell)} + \frac{\lambda_k}{2} \varepsilon_k^{(2m+\ell)} (w_k^{(\ell)} - v_k^{(m+\ell)}),$$

$$\eta_k^{(\ell)} = \max \left\{ \varepsilon_k^{(i)} \mid i \in \mathbb{V}_\ell \right\},$$

$$\bar{x}_{k+1}^{(\ell)} = \bar{x}_k^{(\ell)} + \eta_k^{(\ell)} \left(\frac{1}{\kappa^{(\ell)}} \sum_{i \in \mathbb{V}_\ell} x_{k+1}^{(i)} - \bar{x}_k^{(\ell)} \right).$$

De plus, si $(\forall i \in \{1, \dots, m\}) \mathbf{B}_i = 0$, les $(\|\mathbf{M}_i\|)_{1 \leq i \leq m}$ peuvent être choisis aussi petits que l'on veut, et ainsi on peut poser

$$\chi = \max_{i \in \{1, \dots, m\}} \bar{\theta}_i \|\mathbf{W}_i\|. \quad (7.64)$$

Dans ce cas, l'algorithme 7.12 peut être simplifié, en remarquant que, pour tout $k \in \mathbb{N}$, les variables $(u_k^{(i)})_{1 \leq i \leq m}$ et $(v_k^{(i)})_{1 \leq i \leq m}$ ne sont plus utiles. En imposant les conditions (7.59) et (7.61), et en posant

$$(\forall n \in \mathbb{N})(\forall \ell \in \{1, \dots, r\}) \quad \tilde{v}_k^{(\ell)} = v_k^{(m+\ell)}, \quad (7.65)$$

on obtient l'algorithme 7.13

Algorithme 7.13 Algorithme distribué asynchrone (version 1 simplifiée dans le cas où $(\forall i \in \{1, \dots, m\}) D_i^{-1} = 0$ et $B_i = 0$)

Initialisation : Soient $\mathbf{x}_0$, $(\mathbf{a}_k)_{k \in \mathbb{N}}$ et $(\mathbf{c}_k)_{k \in \mathbb{N}}$ des variables aléatoires à valeurs dans $\mathcal{H}^m$. Pour tout $\ell \in \{1, \dots, r\}$, soit $\tilde{v}_0^{(\ell)} = (\tilde{v}_0^{(\ell, j)})_{1 \leq j \leq \kappa^{(\ell)}}$ une variable aléatoire à valeurs dans $\mathcal{H}^{\kappa^{(\ell)}}$ vérifiant $\sum_{j=1}^{\kappa^{(\ell)}} \tilde{v}_0^{(\ell, j)} = 0$, soit $\bar{x}_0^{(\ell)} = \frac{1}{\kappa^{(\ell)}} \sum_{i \in \mathbb{V}_\ell} x_0^{(i)}$, et soit $\theta^{(\ell)} \in]0, +\infty[$.

Soient, pour tout $i \in \{1, \dots, m\}$, $W_i \in \mathcal{S}^+(\mathcal{H})$.

Soient $(\varepsilon_k)_{k \in \mathbb{N}}$ des variables aléatoires identiquement distribuées à valeurs dans $\mathbb{D}_m$. Soit, pour tout $k \in \mathbb{N}$, $\lambda_k \in]0, +\infty[$.

Itérations :

Pour $k = 0, 1, \dots$

pour $\ell = 1, \dots, r$ $\eta_k^{(\ell)} = \max \left\{ \varepsilon_k^{(i)} \mid i \in \mathbb{V}_\ell \right\}$ pour $j = 1, \dots, \kappa^{(\ell)}$ $w_k^{(\ell, j)} = \eta_k^{(\ell)} \left(2\theta^{(\ell)} (x_k^{(i(\ell, j))}) - \bar{x}_k^{(\ell)} \right) + \tilde{v}_k^{(\ell, j)}$	pour $i = 1, \dots, m$ $y_k^{(i)} = \varepsilon_k^{(i)} \left(J_{W_i A_i} \left(x_k^{(i)} - W_i \left(\sum_{(\ell, j) \in \mathbb{V}_i^*} w_k^{(\ell, j)} + C_i x_k^{(i)} + c_k^{(i)} \right) \right) + a_k^{(i)} \right)$ $x_{k+1}^{(i)} = x_k^{(i)} + \lambda_k \varepsilon_k^{(i)} (y_k^{(i)} - x_k^{(i)})$
pour $\ell = 1, \dots, r$ $\tilde{v}_{k+1}^{(\ell)} = \tilde{v}_k^{(\ell)} + \frac{\lambda_k}{2} \eta_k^{(\ell)} (w_k^{(\ell)} - \tilde{v}_k^{(\ell)})$ $\bar{x}_{k+1}^{(\ell)} = \bar{x}_k^{(\ell)} + \eta_k^{(\ell)} \left(\frac{1}{\kappa^{(\ell)}} \sum_{i \in \mathbb{V}_\ell} x_{k+1}^{(i)} - \bar{x}_k^{(\ell)} \right).$	

Nous pouvons déduire de la section 7.3.3 un autre type d'algorithme distribué, lorsque, pour tout $i \in \{1, \dots, m\}$, $A_i = 0$. En se basant sur l'algorithme 7.5, nous obtenons alors l'algorithme 7.14, dont la convergence vers une solution du problème 7.3, lorsque $A_i = 0$ pour tout $i \in \{1, \dots, m\}$, est assurée par la proposition 7.12. La démonstration de ce résultat est basée sur la proposition 7.5.

Proposition 7.12. *Considérons l'algorithme 7.14. Supposons que, pour tout $i \in \{1, \dots, m\}$, $W_i^{1/2} C_i W_i^{1/2}$ est μ_i -cocoercif, avec $\mu_i \in]0, +\infty[$, et $U_i^{1/2} D_i^{-1} U_i^{1/2}$ est ν_i -cocoercif, avec $\nu_i \in]0, +\infty[$. Soient μ , ν , et χ définis par la proposition 7.11, et supposons que*

$$\min \{ \mu, \nu(1 - \chi) \} > \frac{1}{2}. \quad (7.66)$$

Supposons que les éléments de la suite $(\lambda_k)_{k \in \mathbb{N}}$ sont dans $]0, 1]$ et vérifient $\inf_{k \in \mathbb{N}} \lambda_k > 0$. De plus, posons $(\forall k \in \mathbb{N}) \mathcal{E}_k = \sigma(\varepsilon_k)$ and $\mathcal{X}_k = \sigma(\mathbf{x}_{k'}, \mathbf{v}_{k'})_{0 \leq n' \leq n}$ et supposons que

Algorithme 7.14 Algorithme distribué asynchrone (version 2 simplifiée)

Initialisation : Soient $\mathbf{x}_0$, et $(\mathbf{c}_k)_{k \in \mathbb{N}}$ des variables aléatoires à valeurs dans $\mathcal{H}^m$, et soient $(v_0^{(i)})_{1 \leq i \leq m}$, $(\mathbf{b}_k)_{k \in \mathbb{N}}$, et $(\mathbf{d}_k)_{k \in \mathbb{N}}$ des variables aléatoires à valeurs dans $\mathcal{G}$.

Pour tout $\ell \in \{1, \dots, r\}$, soit $v_0^{(m+\ell)} = (v_0^{(m+\ell, j)})_{1 \leq j \leq \kappa^{(\ell)}}$ une variable aléatoire à valeurs dans $\mathcal{H}^{\kappa^{(\ell)}}$ vérifiant $\sum_{j=1}^{\kappa^{(\ell)}} v_0^{(m+\ell, j)} = 0$, et soit $\theta^{(\ell)} \in]0, +\infty[$.

Soient, pour tout $i \in \{1, \dots, m\}$, $\mathbf{W}_i \in \mathcal{S}^+(\mathcal{H})$ et $\mathbf{U}_i \in \mathcal{S}^+(\mathcal{G}_i)$.

Soient $(\varepsilon_k)_{k \in \mathbb{N}}$ des variables aléatoires identiquement distribuées à valeurs dans $\mathbb{D}_{2m+r}$.

Soit, pour tout $k \in \mathbb{N}$, $\lambda_k \in]0, +\infty[$.

Itérations :

Pour $k = 0, 1, \dots$

pour $i = 1, \dots, m$ $\eta_k^{(i)} = \max \left\{ \varepsilon_k^{(m+i)}, (\varepsilon_k^{(2m+\ell)})_{\ell \in \{\ell' \in \{1, \dots, r\} \mid i \in \mathbb{V}_{\ell'}\}} \right\}$ $w_k^{(i)} = \eta_k^{(i)} \left(x_k^{(i)} - \mathbf{W}_i(\mathbf{C}_i x_k^{(i)} + c_k^{(i)}) \right)$ $\tilde{w}_k^{(i)} = \eta_k^{(i)} \left(w_k^{(i)} - \mathbf{W}_i(\mathbf{M}_i^* v_k^{(i)} + \sum_{(\ell, j) \in \mathbb{V}_i^*} v_k^{(m+\ell, j)}) \right)$ $u_k^{(i)} = \varepsilon_k^{(m+i)} \left(\mathbf{J}_{\mathbf{U}_i \mathbf{B}_i^{-1}} \left(v_k^{(i)} + \mathbf{U}_i(\mathbf{M}_i \tilde{w}_k^{(i)} - \mathbf{D}_i^{-1} v_k^{(i)} + d_k^{(i)}) \right) + b_k^{(i)} \right)$ $v_{k+1}^{(i)} = v_k^{(i)} + \lambda_k \varepsilon_k^{(m+i)} (u_k^{(i)} - v_k^{(i)})$ pour $\ell = 1, \dots, r$ $\bar{w}_k^{(\ell)} = \varepsilon_k^{(2m+\ell)} \frac{\theta^{(\ell)}}{\kappa^{(\ell)}} \sum_{i \in \mathbb{V}_\ell} \tilde{w}_k^{(i)}$ pour $j = 1, \dots, \kappa^{(\ell)}$ $u_k^{(m+\ell, j)} = \varepsilon_k^{(2m+\ell)} \left(v_k^{(m+\ell, j)} + \theta^{(\ell)} \tilde{w}_k^{(i(j, \ell))} - \bar{w}_k^{(\ell)} \right)$ $v_{k+1}^{(m+\ell)} = v_k^{(m+\ell)} + \lambda_k \varepsilon_k^{(2m+\ell)} (u_k^{(m+\ell)} - v_k^{(m+\ell)})$ pour $i = 1, \dots, m$ $x_{k+1}^{(i)} = x_k^{(i)} + \lambda_k \varepsilon_k^{(i)} \left(w_k^{(i)} - \mathbf{W}_i \left(\mathbf{M}_i^* u_k^{(i)} + \sum_{(\ell, j) \in \mathbb{V}_i^*} u_k^{(m+\ell, j)} \right) - x_k^{(i)} \right).$

- (i) $\sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{b}_k\|^2 | \mathcal{X}_k)} < +\infty$, $\sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{c}_k\|^2 | \mathcal{X}_k)} < +\infty$, et $\sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{d}_k\|^2 | \mathcal{X}_k)} < +\infty$
P-presque sûrement

et que les conditions (ii)-(iii) de la proposition 7.11 sont satisfaites.

Sous l'hypothèse 7.1, si $(\forall i \in \{1, \dots, m\}) \mathbf{A}_i = 0$ dans le problème 7.3, alors, pour tout $i \in \{1, \dots, m\}$, $(x_k^{(i)})_{k \in \mathbb{N}}$ converge faiblement P-presque sûrement vers une variable aléatoire $\hat{x}$ à valeurs dans $\hat{\mathbf{F}}$.

La démonstration de la proposition 7.12 est donnée dans l'annexe 7.F.

Remarque 7.10. Lorsque $(\forall i \in \{1, \dots, m\}) \mathbf{D}_i^{-1} = 0$, la condition (7.66) peut être réécrite de la

façon suivante

$$(\forall i \in \{1, \dots, m\}) \quad \|U_i^{1/2} M_i W_i^{1/2}\|^2 + \bar{\theta}_i \|W_i\| < 1 \quad \text{et} \quad \mu_i > 1/2. \quad (7.67)$$

7.5.3 Application aux problèmes variationnels

Nous proposons maintenant d'appliquer les résultats présentés dans la section 7.5.2 à la résolution de problèmes variationnels pouvant s'exprimer de la façon suivante :

Problème 7.4. *Pour tout $i \in \{1, \dots, m\}$, soient $f_i \in \Gamma_0(\mathcal{H})$, $h_i \in \Gamma_0(\mathcal{H})$ une fonction différentiable de gradient Lipschitz, $g_i \in \Gamma_0(\mathcal{G}_i)$, $l_i \in \Gamma_0(\mathcal{G}_i)$ une fonction fortement convexe, et M_i un opérateur non nul de $\mathcal{B}(\mathcal{H}, \mathcal{G}_i)$. Supposons qu'il existe $\bar{x} \in \mathcal{H}$ vérifiant*

$$0 \in \sum_{i=1}^m \partial f_i(\bar{x}) + \nabla h_i(\bar{x}) + M_i^*(\partial g_i \square \partial l_i)(M_i \bar{x}). \quad (7.68)$$

Soit $\check{F}$ l'ensemble défini par

$$\check{F} = \operatorname{Argmin}_{x \in \mathcal{H}} \sum_{i=1}^m f_i(x) + h_i(x) + (g_i \square l_i)(M_i x). \quad (7.69)$$

L'objectif est de trouver une variable aléatoire $\hat{x}$ à valeurs dans $\check{F}$.

Une méthode proximale, déduite de l'algorithme 7.12, permettant de résoudre le problème 7.4 est donnée dans l'algorithme 7.15. Ses résultats de convergence, donnés dans la proposition 7.13, découlent de la proposition 7.11.

Proposition 7.13. *Considérons l'algorithme 7.15. Supposons que, pour tout $i \in \{1, \dots, m\}$, le gradient $h_i \circ W_i^{1/2}$ est μ_i^{-1} -Lipschitz, avec $\mu_i \in]0, +\infty[$, et $l_i^* \circ U_i^{1/2}$ est ν_i^{-1} -Lipschitz, avec $\nu_i \in]0, +\infty[$. Soient $\mu = \min\{\mu_1, \dots, \mu_m\}$, $\nu = \min\{\nu_1, \dots, \nu_m\}$, soit χ défini par la proposition 7.11, et supposons que (7.54) est vérifiée. Supposons que les éléments de la suite $(\lambda_k)_{k \in \mathbb{N}}$ sont dans $]0, 1]$ et vérifient $\inf_{k \in \mathbb{N}} \lambda_k > 0$. Posons, $(\forall k \in \mathbb{N})$ $\mathcal{E}_k = \sigma(\varepsilon_k)$ et $\mathcal{X}_k = \sigma(x_{k'}, v_{k'})_{0 \leq k' \leq k}$ et supposons que les conditions (i)-(iii) de la proposition 7.11 sont vérifiées.*

Sous l'hypothèse 7.1, pour tout $i \in \{1, \dots, m\}$, $(x_k^{(i)})_{k \in \mathbb{N}}$ converge faiblement P-presque sûrement vers une variable aléatoire $\hat{x}$ à valeurs dans $\check{F}$ et, pour tout $\ell \in \{1, \dots, r\}$, $(\bar{x}_k^{(\ell)})_{k \in \mathbb{N}}$ converge faiblement P-presque sûrement vers $\hat{x}$.

Démonstration. Dans la proposition 7.11, pour tout $i \in \{1, \dots, m\}$, posons $A_i = \partial f_i$, $B_i = \partial g_i$, $C_i = \nabla h_i$, et $D_i^{-1} = \nabla l_i^*$. D'après (7.68) et [Bauschke et Combettes, 2011, Prop. 16.5], on a

$$0 \in \sum_{i=1}^m A_i \bar{x} + C_i \bar{x} + M_i^*(B_i \square D_i)(M_i \bar{x}) \subset \partial \left(\sum_{i=1}^m f_i + h_i + (g_i \square l_i) \circ M_i \right)(\bar{x}), \quad (7.70)$$

ce qui montre que $\emptyset \neq \hat{F} \subset \check{F}$. Nous pouvons donc conclure en appliquant la proposition 7.11 et en utilisant la remarque 7.9(ii). $\square$

Algorithme 7.15 Algorithme distribué asynchrone (version 1 simplifiée) appliqué dans un cadre variationnel

Initialisation : Soient $\mathbf{x}_0, (\mathbf{a}_k)_{k \in \mathbb{N}}$ et $(\mathbf{c}_k)_{k \in \mathbb{N}}$ des variables aléatoires à valeurs dans $\mathcal{H}^m$, et soient $(v_0^{(i)})_{1 \leq i \leq m}, (\mathbf{b}_k)_{k \in \mathbb{N}}$, et $(\mathbf{d}_k)_{k \in \mathbb{N}}$ des variables aléatoires à valeurs dans $\mathcal{G}$. Pour tout $\ell \in \{1, \dots, r\}$, soit $v_0^{(m+\ell)} = (v_0^{(m+\ell, j)})_{1 \leq j \leq \kappa^{(\ell)}}$ une variable aléatoire à valeurs dans $\mathcal{H}^{\kappa^{(\ell)}}$ vérifiant $\sum_{j=1}^{\kappa^{(\ell)}} v_0^{(m+\ell, j)} = 0$, soit $\bar{x}_0^{(\ell)} = \frac{1}{\kappa^{(\ell)}} \sum_{i \in \mathbb{V}_\ell} x_0^{(i)}$, et soit $\theta^{(\ell)} \in]0, +\infty[$. Soient, pour tout $i \in \{1, \dots, m\}$, $\mathbf{W}_i \in \mathcal{S}^+(\mathcal{H})$ et $\mathbf{U}_i \in \mathcal{S}^+(\mathcal{G}_i)$. Soient $(\varepsilon_k)_{k \in \mathbb{N}}$ des variables aléatoires identiquement distribuées à valeurs dans $\mathbb{D}_{2m+r}$. Soit, pour tout $k \in \mathbb{N}$, $\lambda_k \in]0, +\infty[$.

Itérations :

Pour $k = 0, 1, \dots$

pour $\ell = 1, \dots, r$ pour $j = 1, \dots, \kappa^{(\ell)}$ $w_k^{(\ell, j)} = \varepsilon_k^{(2m+\ell)} \left(2\theta^{(\ell)} (x_k^{(i(\ell, j))} - \bar{x}_k^{(\ell)}) + v_k^{(m+\ell, j)} \right),$ pour $i = 1, \dots, m$ $u_k^{(i)} = \varepsilon_k^{(m+i)} \left(\text{prox}_{\mathbf{U}_i^{-1}, \mathcal{G}_i^*} \left(v_k^{(i)} + \mathbf{U}_i \left(\mathbf{M}_i x_k^{(i)} - \nabla l_i^*(v_k^{(i)}) + d_k^{(i)} \right) \right) + b_k^{(i)} \right),$ $y_k^{(i)} = \varepsilon_k^{(i)} \left(\text{prox}_{\mathbf{W}_i^{-1}, \mathcal{F}_i} \left(x_k^{(i)} - \mathbf{W}_i \left(\mathbf{M}_i^* (2u_k^{(i)} - v_k^{(i)}) + \sum_{(\ell, j) \in \mathbb{V}_i^*} w_k^{(\ell, j)} + \nabla h_i(x_k^{(i)}) + c_k^{(i)} \right) \right) + a_k^{(i)} \right),$ $v_{k+1}^{(i)} = v_k^{(i)} + \lambda_k \varepsilon_k^{(m+i)} (u_k^{(i)} - v_k^{(i)}),$ $x_{k+1}^{(i)} = x_k^{(i)} + \lambda_k \varepsilon_k^{(i)} (y_k^{(i)} - x_k^{(i)}),$ pour $\ell = 1, \dots, r$ $v_{k+1}^{(m+\ell)} = v_k^{(m+\ell)} + \frac{\lambda_k}{2} \varepsilon_k^{(2m+\ell)} (w_k^{(\ell)} - v_k^{(m+\ell)}),$ $\eta_k^{(\ell)} = \max \left\{ \varepsilon_k^{(i)} \mid i \in \mathbb{V}_\ell \right\},$ $\bar{x}_{k+1}^{(\ell)} = \bar{x}_k^{(\ell)} + \eta_k^{(\ell)} \left(\frac{1}{\kappa^{(\ell)}} \sum_{i \in \mathbb{V}_\ell} x_{k+1}^{(i)} - \bar{x}_k^{(\ell)} \right).$

Remarque 7.11.

- (i) De façon similaire, un second algorithme distribué d'optimisation convexe peut être déduit de la proposition 7.12.
- (ii) Dans le cas où, pour tout $\ell \in \{1, \dots, r\}$, $\kappa^{(\ell)} = 2$, les variables $(\bar{x}_k^{(\ell)})_{1 \leq \ell \leq r}$ sont mises à jour à chaque itération $k \in \mathbb{N}$, i.e.

$$(\forall k \in \mathbb{N})(\forall \ell \in \{1, \dots, r\}) \quad \bar{x}_k^{(\ell)} = \frac{1}{2} (x_k^{(i(\ell, 1))} + x_k^{(i(\ell, 2))}),$$

et (7.58)-(7.59) sont satisfaites, une version simplifiée de l'algorithme 7.15 est donnée par l'algorithme 7.16.

Algorithme 7.16 Algorithme distribué (version 1 simplifiée) sur un graphe non orienté

Initialisation : Soient $\mathbf{x}_0$, $(\mathbf{a}_k)_{k \in \mathbb{N}}$ et $(\mathbf{c}_k)_{k \in \mathbb{N}}$ des variables aléatoires à valeurs dans $\mathcal{H}^m$, et soient $(v_0^{(i)})_{1 \leq i \leq m}$, $(\mathbf{b}_k)_{k \in \mathbb{N}}$, et $(\mathbf{d}_k)_{k \in \mathbb{N}}$ des variables aléatoires à valeurs dans $\mathcal{G}$.

Pour tout $\ell \in \{1, \dots, r\}$, soient $\theta^{(\ell)} \in]0, +\infty[$, $v_0^{(m+\ell,1)}$ une variable aléatoire à valeurs dans $\mathcal{H}$ et $v_0^{(m+\ell,2)} = -v_0^{(m+\ell,1)}$.

Soient, pour tout $i \in \{1, \dots, m\}$, $\mathbf{W}_i \in \mathcal{S}^+(\mathcal{H})$, $\mathbf{U}_i \in \mathcal{S}^+(\mathcal{G}_i)$, et $\bar{\theta}_i$ défini par (7.53).

Soient $(\varepsilon_k)_{k \in \mathbb{N}}$ des variables aléatoires identiquement distribuées à valeurs dans $\mathbb{D}_{2m+r}$.

Soit, pour tout $k \in \mathbb{N}$, $\lambda_k \in]0, +\infty[$.

Itérations :

Pour $k = 0, 1, \dots$

$$\begin{array}{l}
 \text{pour } \ell = 1, \dots, r \\
 \quad \left[\begin{array}{l}
 \eta_k^{(\ell)} = \max \left\{ \varepsilon_k^{(i)} \mid i \in \mathbb{V}_\ell \right\}, \\
 v_{k+1}^{(\ell,1)} = v_k^{(m+\ell,1)} + \frac{\lambda_k}{2} \eta_k^{(\ell)} \theta^{(\ell)} (x_k^{(i(\ell,1))} - x_k^{(i(\ell,2))}), \\
 v_{k+1}^{(m+\ell,2)} = -v_{k+1}^{(m+\ell,1)},
 \end{array} \right. \\
 \text{pour } i = 1, \dots, m \\
 \quad \left[\begin{array}{l}
 u_k^{(i)} = \varepsilon_k^{(i)} \left(\text{prox}_{\mathbf{U}_i^{-1}, \mathbf{g}_i^*} \left(v_k^{(i)} + \mathbf{U}_i (\mathbf{M}_i x_k^{(i)} - \nabla l_i^*(v_k^{(i)}) + d_k^{(i)}) \right) + b_k^{(i)} \right), \\
 y_k^{(i)} = \varepsilon_k^{(i)} \left(\text{prox}_{\mathbf{W}_i^{-1}, \mathbf{f}_i} \left((1 - \mathbf{W}_i \bar{\theta}_i) x_k^{(i)} - \mathbf{W}_i (\mathbf{M}_i^* (2u_k^{(i)} - v_k^{(i)}) \right. \right. \right. \\
 \quad \left. \left. \left. + \sum_{(\ell,j) \in \mathbb{V}_i^*} (v_k^{(m+\ell,j)} - \theta^{(\ell)} x_k^{(i(\ell,\bar{j}))}) + \nabla h_i(x_k^{(i)}) + c_k^{(i)}) \right) + a_k^{(i)} \right), \\
 v_{k+1}^{(i)} = v_k^{(i)} + \lambda_k \varepsilon_k^{(i)} (u_k^{(i)} - v_k^{(i)}), \\
 x_{k+1}^{(i)} = x_k^{(i)} + \lambda_k \varepsilon_k^{(i)} (y_k^{(i)} - x_k^{(i)}),
 \end{array} \right.
 \end{array}$$

En effet dans ce cas, pour tout $k \in \mathbb{N}$ et $\ell \in \{1, \dots, r\}$, on a

$$(\forall j \in \{1, 2\}) \quad w_k^{(\ell,j)} = \eta_k^{(\ell)} \left(\theta^{(\ell)} (x_k^{(i(\ell,j))} - x_k^{(i(\ell,\bar{j}))}) + v_k^{(m+\ell,j)} \right),$$

où, pour tout $j \in \{1, 2\}$, $\bar{j} = 3 - j$. Par conséquent, puisque (7.61) est satisfaite,

$$\begin{aligned}
 (\forall k \in \mathbb{N})(\forall \ell \in \{1, \dots, r\}) \quad v_{k+1}^{(m+\ell,1)} &= v_k^{(m+\ell,1)} + \frac{\lambda_k}{2} \eta_k^{(\ell)} \theta^{(\ell)} (x_k^{(i(\ell,1))} - x_k^{(i(\ell,2))}), \\
 v_{k+1}^{(m+\ell,2)} &= -v_{k+1}^{(m+\ell,1)}.
 \end{aligned}$$

En remarquant que puisque $\kappa_\ell \equiv 2$, pour tout $i \in \{1, \dots, m\}$, pour tout $(\ell, j) \in \mathbb{V}_i^*$, $x_k^{(i(\ell,j))} =$

$x_k^{(i)}$, et en utilisant (7.53), on a

$$\begin{aligned} (\forall i \in \{1, \dots, m\}) \quad \varepsilon_k^{(i)} \sum_{(\ell, j) \in \mathbb{V}_i^*} w_k^{(\ell, j)} &= \varepsilon_k^{(i)} \left(\sum_{(\ell, j) \in \mathbb{V}_i^*} (v_k^{(m+\ell, j)} + \theta^{(\ell)}(x_k^{(i(\ell, j))} - x_k^{(i(\ell, \bar{j}))})) \right) \\ &= \varepsilon_k^{(i)} \left(\bar{\theta}_i x_k^{(i)} + \sum_{(\ell, j) \in \mathbb{V}_i^*} (v_k^{(m+\ell, j)} - \theta^{(\ell)} x_k^{(i(\ell, \bar{j}))}) \right). \end{aligned}$$

Ainsi, en réarrangeant les différentes étapes de l'algorithme 7.15, on obtient l'algorithme 7.16. Dans ce cas, les ensembles $(\mathbb{V}_\ell)_{1 \leq \ell \leq r}$ peuvent être vus comme les arêtes d'un graphe connecté non orienté dont les nœuds sont indicés par $i \in \{1, \dots, m\}$.

- (iii) Si $(\forall i \in \{1, \dots, m\}) \mathbf{g}_i = \mathbf{0}$ et $\mathbf{l}_i = \iota_{\{0\}}$, $(\forall \ell \in \{1, \dots, r\}) \kappa^{(\ell)} = 2$, et (7.59) est satisfaite, alors on peut prendre les $(\mathbf{M}_i)_{1 \leq i \leq m}$ aussi petit qu'on le souhaite et l'algorithme 7.16 se réduit à l'algorithme 7.17. Pour cela, nous posons $(\forall i \in \{1, \dots, m\}) v_0^{(i)} = 0$, $(\forall k \in \mathbb{N}) \tilde{\mathbf{v}}_k = (\tilde{v}_k^{(\ell)})_{1 \leq \ell \leq r} = (v_k^{(m+\ell)})_{1 \leq \ell \leq r}$, et $\mathbf{b}_k = \mathbf{d}_k = \mathbf{0}$. Notons que, dans ce cas particulier,

Algorithme 7.17 Algorithme distribué (version 1 simplifiée) sur un graphe non orienté lorsque $(\forall i \in \{1, \dots, m\}) \mathbf{g}_i = \mathbf{0}$ et $\mathbf{l}_i = \iota_{\{0\}}$

Initialisation : Soient $\mathbf{x}_0$, $(\mathbf{a}_k)_{k \in \mathbb{N}}$ et $(\mathbf{c}_k)_{k \in \mathbb{N}}$ des variables aléatoires à valeurs dans $\mathcal{H}^m$. Pour tout $\ell \in \{1, \dots, r\}$, soient $\theta^{(\ell)} \in]0, +\infty[$, $\tilde{v}_0^{(\ell, 1)}$ une variable aléatoire à valeurs dans $\mathcal{H}$ et $\tilde{v}_0^{(\ell, 2)} = -\tilde{v}_0^{(\ell, 1)}$.

Soient, pour tout $i \in \{1, \dots, m\}$, $\mathbf{W}_i \in \mathcal{S}^+(\mathcal{H})$ et $\bar{\theta}_i$ défini par (7.53).

Soient $(\varepsilon_k)_{k \in \mathbb{N}}$ des variables aléatoires identiquement distribuées à valeurs dans $\mathbb{D}_m$. Soit, pour tout $k \in \mathbb{N}$, $\lambda_k \in]0, +\infty[$.

Itérations :

Pour $k = 0, 1, \dots$

$$\left| \begin{array}{l} \text{pour } \ell = 1, \dots, r \\ \quad \left| \begin{array}{l} \eta_k^{(\ell)} = \max \left\{ \varepsilon_k^{(i)} \mid i \in \mathbb{V}_\ell \right\}, \\ \tilde{v}_{k+1}^{(\ell, 1)} = \tilde{v}_k^{(\ell, 1)} + \frac{\lambda_k}{2} \eta_k^{(\ell)} \theta^{(\ell)} (x_k^{(i(\ell, 1))} - x_k^{(i(\ell, 2))}), \\ \tilde{v}_{k+1}^{(\ell, 2)} = -\tilde{v}_{k+1}^{(\ell, 1)}, \end{array} \right. \\ \text{pour } i = 1, \dots, m \\ \quad \left| \begin{array}{l} y_k^{(i)} = \varepsilon_k^{(i)} \left(\text{prox}_{\mathbf{W}_i^{-1}, \mathbf{f}_i} \left((1 - \mathbf{W}_i \bar{\theta}_i) x_k^{(i)} - \mathbf{W}_i \left(\sum_{(\ell, j) \in \mathbb{V}_i^*} (\tilde{v}_k^{(\ell, j)} - \theta^{(\ell)} x_k^{(i(\ell, \bar{j}))}) \right. \right. \right. \\ \quad \left. \left. \left. + \nabla \mathbf{h}_i(x_k^{(i)}) + \mathbf{c}_k^{(i)} \right) \right) + \mathbf{a}_k^{(i)} \right) \\ x_{k+1}^{(i)} = x_k^{(i)} + \lambda_k \varepsilon_k^{(i)} (y_k^{(i)} - x_k^{(i)}). \end{array} \right. \end{array} \right.$$

d'après la remarque 7.9(iv), il suffit que les conditions (7.63) et (7.64) soient satisfaites pour que (7.54) soit vérifiée.

Notons, pour tout $i \in \{1, \dots, m\}$, κ_i^* le cardinal de l'ensemble $\mathbb{V}_i^*$ (ce qui correspond au degré du nœud i). Dans le cas particulier où $\mathcal{H}$ est un espace euclidien, qu'aucun terme d'erreur n'est pris en compte, que

$$\begin{aligned} (\forall k \in \mathbb{N}) \quad \lambda_k &= 1, \\ (\forall \ell \in \{1, \dots, r\}) \quad \theta^{(\ell)} &= \theta \in]0, +\infty[, \\ (\forall i \in \{1, \dots, m\}) \quad \mathbf{W}_i &= \tau_i \text{Id} \end{aligned}$$

avec

$$(\forall i \in \{1, \dots, m\}) \quad \tau_i = \tau / \kappa_i^*,$$

et $\tau \in]0, +\infty[$, on retrouve l'algorithme distribué développé dans [Bianchi et al., 2014]. D'après (7.53), pour tout $i \in \{1, \dots, m\}$, $\bar{\theta}_i = \kappa_i^* \theta$. L'expression (7.64) s'écrit alors

$$\chi = \max_{i \in \{1, \dots, m\}} (\theta \kappa_i^* \tau_i) = \theta \max_{i \in \{1, \dots, m\}} (\kappa_i^* \tau / \kappa_i^*) = \theta \tau.$$

De plus, pour tout $i \in \{1, \dots, m\}$, μ_i^{-1} étant une constante de Lipschitz du gradient de $\mathbf{h}_i \circ \mathbf{W}_i^{1/2}$, la constante de Lipschitz de $\mathbf{h}_i$ est alors égale à $\bar{\mu}_i = \kappa_i^* / (\mu_i \tau)$. Ainsi $\mu_i = \kappa_i^* / (\bar{\mu}_i \tau)$, et la condition (7.63) se réécrit

$$(\forall i \in \{1, \dots, m\}) \quad \tau^{-1} - \theta > \frac{1}{2 \min_{i \in \{1, \dots, m\}} (\kappa_i^* / \bar{\mu}_i)}. \quad (7.71)$$

La condition de convergence donnée dans [Bianchi et al., 2014, Thm. 6] est

$$\tau^{-1} - \theta > \frac{\bar{\mu}}{2\underline{\kappa}^*}, \quad (7.72)$$

où $\bar{\mu} = \max_{i \in \{1, \dots, m\}} \bar{\mu}_i$ et $\underline{\kappa}^* = \min_{i \in \{1, \dots, m\}} \kappa_i^*$. Or, puisque

$$\min_{i \in \{1, \dots, m\}} (\kappa_i^* / \bar{\mu}_i) \geq \left(\min_{i \in \{1, \dots, m\}} \kappa_i^* \right) \left(\min_{i \in \{1, \dots, m\}} 1 / \bar{\mu}_i \right) = \frac{\underline{\kappa}^*}{\bar{\mu}},$$

on a alors

$$\frac{\bar{\mu}}{2\underline{\kappa}^*} \geq \frac{1}{2 \min_{i \in \{1, \dots, m\}} (\kappa_i^* / \bar{\mu}_i)}.$$

On en déduit ainsi que la condition (7.72) donnée dans [Bianchi et al., 2014, Thm. 6] est plus restrictive que la condition (7.71) que nous déduisons de nos résultats de convergence.

Remarque 7.12. Le diagramme représenté dans la figure 7.7 donne les différents algorithmes distribués développés dans cette section. Les cadres verts foncés renvoient aux problèmes traités (problèmes 7.3 et 7.4), les cadres verts clairs indiquent les cas particuliers de ces problèmes, et les ellipses bleues donnent les algorithmes permettant de les traiter.

7.6 CONCLUSION

Dans ce chapitre nous avons proposé plusieurs algorithmes permettant de résoudre des problèmes d'inclusion monotone ou des problèmes de minimisation de fonctions convexes. Ces algorithmes combinent des méthodes primales-duales, utiles pour traiter des sommes d'opérateurs (ou de fonctions) composés avec des opérateurs linéaires, et des stratégies alternées aléatoirement par blocs proposées dans [Combettes et Pesquet, 2015]. Notons que dans aucun des algorithmes proposés il n'est nécessaire d'inverser des opérateurs linéaires, ce qui peut se révéler utile en pratique lorsque ces derniers sont mal conditionnés ou de très grande dimension. De plus, les algorithmes primaux-duaux proposés peuvent être accélérés grâce à l'utilisation d'opérateurs de préconditionnement. Nous avons montré les performances de la stratégie alternée aléatoirement par bloc au travers d'un exemple applicatif de reconstruction d'un maillage 3D. Les divers algorithmes alternés aléatoirement par bloc sont donnés dans le diagramme de la figure 7.5 Enfin, nous avons déduit de ces algorithmes primaux-duaux alternés par bloc des algorithmes distribués permettant de résoudre des problèmes d'inclusion monotone (ou des problèmes de minimisation de fonctions convexe) en utilisant une topologie générale d'hypergraphes. Un diagramme résumant les liens existants entre les différents algorithmes distribués proposés est donné dans le figure 7.7.

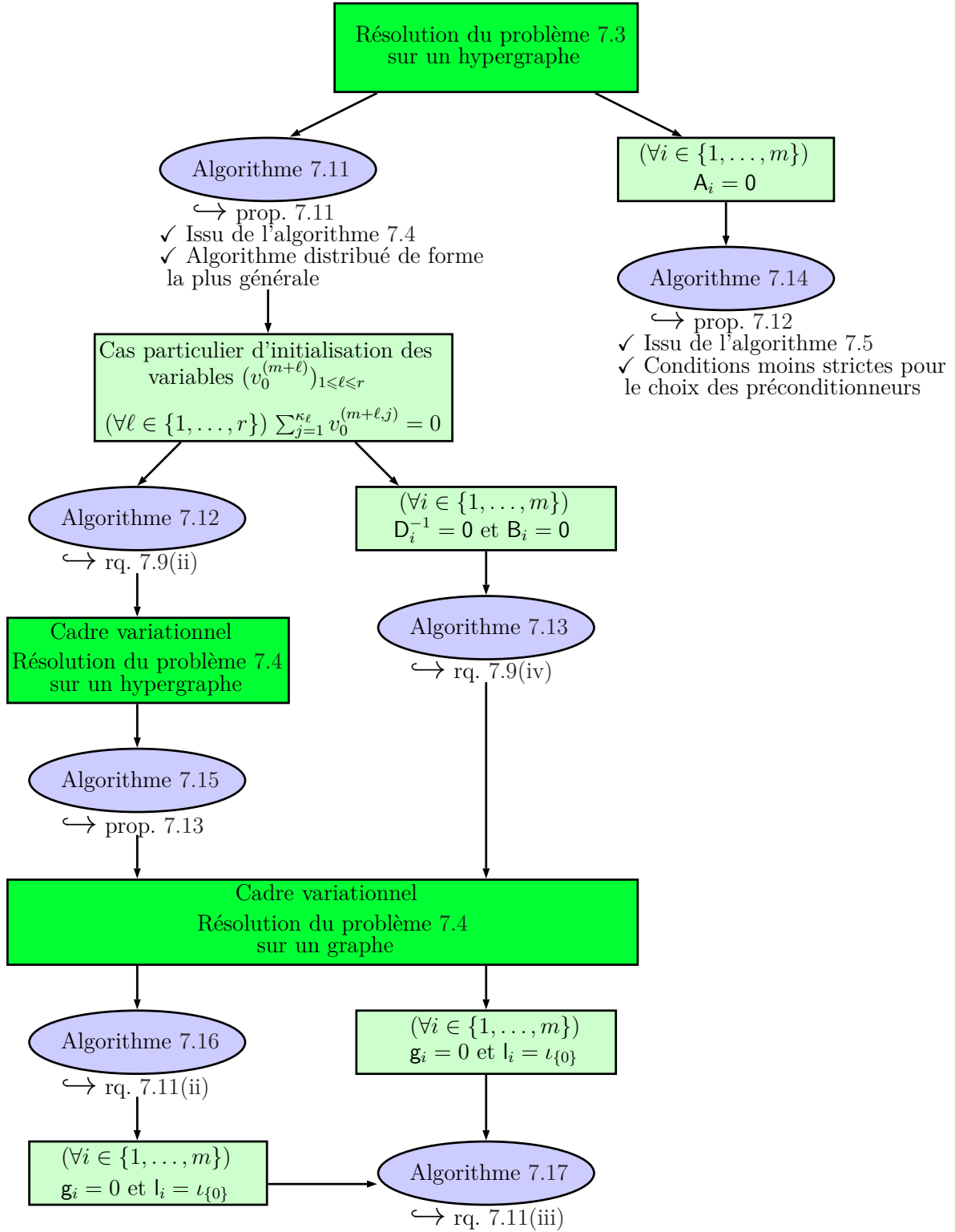


Figure 7.7 – Diagramme résumant les algorithmes développés dans la section 7.5.

7.A DÉMONSTRATION DU LEMME 7.1

(i) Les opérateurs $\mathbf{W}^{-1}$ et $\mathbf{U}^{-1}$ étant linéaires, bornés et auto-adjoints, $\mathbf{V}'$ est aussi linéaire, borné et auto-adjoint. De plus, pour tout $(\mathbf{x}, \mathbf{v}) \in \mathcal{K}$,

$$\begin{aligned} \langle \mathbf{x} \mid (\mathbf{W}^{-1} - \mathbf{L}^* \mathbf{U} \mathbf{L}) \mathbf{x} \rangle &= \langle \mathbf{W}^{-1/2} \mathbf{x} \mid (\mathbf{Id} - \mathbf{W}^{1/2} \mathbf{L}^* \mathbf{U} \mathbf{L} \mathbf{W}^{1/2}) \mathbf{W}^{-1/2} \mathbf{x} \rangle \\ &= \langle \mathbf{x} \mid \mathbf{W}^{-1} \mathbf{x} \rangle - \langle \mathbf{W}^{-1/2} \mathbf{x} \mid \mathbf{W}^{1/2} \mathbf{L}^* \mathbf{U} \mathbf{L} \mathbf{W}^{1/2} \mathbf{W}^{-1/2} \mathbf{x} \rangle \\ &\geq (1 - \|\mathbf{W}^{1/2} \mathbf{L}^* \mathbf{U} \mathbf{L} \mathbf{W}^{1/2}\|) \langle \mathbf{x} \mid \mathbf{W}^{-1} \mathbf{x} \rangle \\ &\geq (1 - \|\mathbf{U}^{1/2} \mathbf{L} \mathbf{W}^{1/2}\|^2) \|\mathbf{W}\|^{-1} \|\mathbf{x}\|^2 \end{aligned} \quad (7.73)$$

et

$$\begin{aligned} \langle \mathbf{v} \mid (\mathbf{U}^{-1} - \mathbf{L} \mathbf{W} \mathbf{L}^*) \mathbf{v} \rangle &\geq (1 - \|\mathbf{U}^{1/2} \mathbf{L} \mathbf{W} \mathbf{L}^* \mathbf{U}^{1/2}\|) \langle \mathbf{v} \mid \mathbf{U}^{-1} \mathbf{v} \rangle \\ &\geq (1 - \|\mathbf{U}^{1/2} \mathbf{L} \mathbf{W}^{1/2}\|^2) \|\mathbf{U}\|^{-1} \|\mathbf{v}\|^2. \end{aligned} \quad (7.74)$$

On peut donc en déduire que

$$\begin{aligned} \langle (\mathbf{x}, \mathbf{v}) \mid \mathbf{V}'(\mathbf{x}, \mathbf{v}) \rangle &= \langle \mathbf{x} - \mathbf{W} \mathbf{L}^* \mathbf{v} \mid \mathbf{W}^{-1}(\mathbf{x} - \mathbf{W} \mathbf{L}^* \mathbf{v}) \rangle + \langle \mathbf{v} \mid (\mathbf{U}^{-1} - \mathbf{L} \mathbf{W} \mathbf{L}^*) \mathbf{v} \rangle \\ &\geq (1 - \|\mathbf{U}^{1/2} \mathbf{L} \mathbf{W}^{1/2}\|^2) \|\mathbf{U}\|^{-1} \|\mathbf{v}\|^2 \end{aligned}$$

et, de façon similaire,

$$\langle (\mathbf{x}, \mathbf{v}) \mid \mathbf{V}'(\mathbf{x}, \mathbf{v}) \rangle \geq (1 - \|\mathbf{U}^{1/2} \mathbf{L} \mathbf{W}^{1/2}\|^2) \|\mathbf{W}\|^{-1} \|\mathbf{x}\|^2.$$

Les deux inégalités ci-dessus nous donnent

$$\begin{aligned} \langle (\mathbf{x}, \mathbf{v}) \mid \mathbf{V}'(\mathbf{x}, \mathbf{v}) \rangle &\geq (1 - \|\mathbf{U}^{1/2} \mathbf{L} \mathbf{W}^{1/2}\|^2) \min\{\|\mathbf{W}\|^{-1}, \|\mathbf{U}\|^{-1}\} \max\{\|\mathbf{x}\|^2, \|\mathbf{v}\|^2\} \\ &\geq \frac{1}{2} (1 - \|\mathbf{U}^{1/2} \mathbf{L} \mathbf{W}^{1/2}\|^2) \min\{\|\mathbf{W}\|^{-1}, \|\mathbf{U}\|^{-1}\} (\|\mathbf{x}\|^2 + \|\mathbf{v}\|^2), \end{aligned}$$

Ce qui montre que $\mathbf{V}'$ est un opérateur fortement positif. C'est donc un isomorphisme et son inverse appartient à $\mathcal{S}^+(\mathcal{K})$.

D'autre part, (7.73) (resp. (7.74)) montre que $\mathbf{W}^{-1} - \mathbf{L}^* \mathbf{U} \mathbf{L}$ (resp. $\mathbf{U}^{-1} - \mathbf{L} \mathbf{W} \mathbf{L}^*$) est un isomorphisme puisqu'il appartient à $\mathcal{S}^+(\mathcal{H})$ (resp. $\mathcal{S}^+(\mathcal{G})$). L'expression de l'inverse de $\mathbf{V}'$ peut être obtenue par un calcul direct.

(ii) Soit $\alpha \in]0, +\infty[$. Montrer que $\mathbf{V}^{1/2} \mathbf{R} \mathbf{V}^{1/2}$ est ϑ_α -cocoercif est équivalent à montrer que

$$\begin{aligned} (\forall (\mathbf{z}, \mathbf{z}') \in \mathcal{K}^2) \quad &\langle \mathbf{z} - \mathbf{z}' \mid \mathbf{V}^{1/2} \mathbf{R} \mathbf{V}^{1/2} \mathbf{z} - \mathbf{V}^{1/2} \mathbf{R} \mathbf{V}^{1/2} \mathbf{z}' \rangle \\ &\geq \vartheta_\alpha \|\mathbf{V}^{1/2} \mathbf{R} \mathbf{V}^{1/2} \mathbf{z} - \mathbf{V}^{1/2} \mathbf{R} \mathbf{V}^{1/2} \mathbf{z}'\|^2 \\ \Leftrightarrow (\forall (\mathbf{z}, \mathbf{z}') \in \mathcal{K}^2) \quad &\langle \mathbf{z} - \mathbf{z}' \mid \mathbf{R} \mathbf{z} - \mathbf{R} \mathbf{z}' \rangle \geq \vartheta_\alpha \|\mathbf{R} \mathbf{z} - \mathbf{R} \mathbf{z}'\|_{\mathbf{V}}^2. \end{aligned}$$

Soit $\mathbf{z} = (\mathbf{x}, \mathbf{v}) \in \mathcal{K}$ et $\mathbf{z}' = (\mathbf{x}', \mathbf{v}') \in \mathcal{K}$. On a, d'une part

$$\begin{aligned} \|\mathbf{Rz} - \mathbf{Rz}'\|_{\mathbf{V}}^2 &= \langle \mathbf{Cx} - \mathbf{Cx}' \mid (\mathbf{W}^{-1} - \mathbf{L}^* \mathbf{UL})^{-1} (\mathbf{Cx} - \mathbf{Cx}') \rangle \\ &\quad + \langle \mathbf{D}^{-1} \mathbf{v} - \mathbf{D}^{-1} \mathbf{v}' \mid (\mathbf{U}^{-1} - \mathbf{LWL}^*)^{-1} (\mathbf{D}^{-1} \mathbf{v} - \mathbf{D}^{-1} \mathbf{v}') \rangle \\ &\quad + 2 \langle \mathbf{Cx} - \mathbf{Cx}' \mid \mathbf{WL}^* (\mathbf{U}^{-1} - \mathbf{LWL}^*)^{-1} (\mathbf{D}^{-1} \mathbf{v} - \mathbf{D}^{-1} \mathbf{v}') \rangle. \end{aligned} \quad (7.75)$$

D'autre part,

$$\begin{aligned} &\langle \mathbf{Cx} - \mathbf{Cx}' \mid (\mathbf{W}^{-1} - \mathbf{L}^* \mathbf{UL})^{-1} (\mathbf{Cx} - \mathbf{Cx}') \rangle \\ &= \langle \mathbf{W}^{1/2} (\mathbf{Cx} - \mathbf{Cx}') \mid (\mathbf{Id} - \mathbf{W}^{1/2} \mathbf{L}^* \mathbf{ULW}^{1/2})^{-1} \mathbf{W}^{1/2} (\mathbf{Cx} - \mathbf{Cx}') \rangle \\ &\leq \|(\mathbf{Id} - \mathbf{W}^{1/2} \mathbf{L}^* \mathbf{ULW}^{1/2})^{-1}\| \|\mathbf{Cx} - \mathbf{Cx}'\|_{\mathbf{W}}^2 \\ &= (1 - \|\mathbf{U}^{1/2} \mathbf{LW}^{1/2}\|^2)^{-1} \|\mathbf{Cx} - \mathbf{Cx}'\|_{\mathbf{W}}^2, \end{aligned} \quad (7.76)$$

$$\begin{aligned} &\langle \mathbf{D}^{-1} \mathbf{v} - \mathbf{D}^{-1} \mathbf{v}' \mid (\mathbf{U}^{-1} - \mathbf{LWL}^*)^{-1} (\mathbf{D}^{-1} \mathbf{v} - \mathbf{D}^{-1} \mathbf{v}') \rangle \\ &\leq (1 - \|\mathbf{U}^{1/2} \mathbf{LW}^{1/2}\|^2)^{-1} \|\mathbf{D}^{-1} \mathbf{v} - \mathbf{D}^{-1} \mathbf{v}'\|_{\mathbf{U}}^2, \end{aligned} \quad (7.77)$$

$$\begin{aligned} &\langle \mathbf{Cx} - \mathbf{Cx}' \mid \mathbf{WL}^* (\mathbf{U}^{-1} - \mathbf{LWL}^*)^{-1} (\mathbf{D}^{-1} \mathbf{v} - \mathbf{D}^{-1} \mathbf{v}') \rangle \\ &\leq \|\mathbf{W}^{1/2} (\mathbf{Cx} - \mathbf{Cx}')\| \|\mathbf{W}^{1/2} \mathbf{L}^* \mathbf{U}^{1/2} (\mathbf{Id} - \mathbf{U}^{1/2} \mathbf{LWL}^* \mathbf{U}^{1/2})^{-1}\| \\ &\quad \times \|\mathbf{U}^{1/2} (\mathbf{D}^{-1} \mathbf{v} - \mathbf{D}^{-1} \mathbf{v}')\| \\ &\leq \|\mathbf{U}^{1/2} \mathbf{LW}^{1/2}\| (1 - \|\mathbf{U}^{1/2} \mathbf{LW}^{1/2}\|^2)^{-1} \|\mathbf{Cx} - \mathbf{Cx}'\|_{\mathbf{W}} \|\mathbf{D}^{-1} \mathbf{v} - \mathbf{D}^{-1} \mathbf{v}'\|_{\mathbf{U}} \\ &\leq \frac{1}{2} \|\mathbf{U}^{1/2} \mathbf{LW}^{1/2}\| (1 - \|\mathbf{U}^{1/2} \mathbf{LW}^{1/2}\|^2)^{-1} \\ &\quad \times (\alpha \|\mathbf{Cx} - \mathbf{Cx}'\|_{\mathbf{W}}^2 + \alpha^{-1} \|\mathbf{D}^{-1} \mathbf{v} - \mathbf{D}^{-1} \mathbf{v}'\|_{\mathbf{U}}^2). \end{aligned} \quad (7.78)$$

En combinant les équations (7.75) à (7.78) et en utilisant la cocoercivité de $\mathbf{W}^{1/2} \mathbf{C} \mathbf{W}^{1/2}$ et de $\mathbf{U}^{1/2} \mathbf{D}^{-1} \mathbf{U}^{1/2}$, on obtient la suite d'inégalités :

$$\begin{aligned} \|\mathbf{Rz} - \mathbf{Rz}'\|_{\mathbf{V}}^2 &\leq (1 - \|\mathbf{U}^{1/2} \mathbf{LW}^{1/2}\|^2)^{-1} ((1 + \alpha \|\mathbf{U}^{1/2} \mathbf{LW}^{1/2}\|) \|\mathbf{Cx} - \mathbf{Cx}'\|_{\mathbf{W}}^2 \\ &\quad + (1 + \alpha^{-1} \|\mathbf{U}^{1/2} \mathbf{LW}^{1/2}\|) \|\mathbf{D}^{-1} \mathbf{v} - \mathbf{D}^{-1} \mathbf{v}'\|_{\mathbf{U}}^2) \\ &\leq (1 - \|\mathbf{U}^{1/2} \mathbf{LW}^{1/2}\|^2)^{-1} (\mu^{-1} (1 + \alpha \|\mathbf{U}^{1/2} \mathbf{LW}^{1/2}\|) \langle \mathbf{x} - \mathbf{x}' \mid \mathbf{Cx} - \mathbf{Cx}' \rangle \\ &\quad + \nu^{-1} (1 + \alpha^{-1} \|\mathbf{U}^{1/2} \mathbf{LW}^{1/2}\|) \langle \mathbf{v} - \mathbf{v}' \mid \mathbf{D}^{-1} \mathbf{v} - \mathbf{D}^{-1} \mathbf{v}' \rangle) \\ &\leq \vartheta_{\alpha}^{-1} \langle \mathbf{z} - \mathbf{z}' \mid \mathbf{Rz} - \mathbf{Rz}' \rangle. \end{aligned}$$

7.B DÉMONSTRATION DE LA PROPOSITION 7.3

D'après la condition (iii), pour tout $j \in \{1, \dots, p\}$, $\max \{\varepsilon_k^{(j)}, (\varepsilon_k^{(p+l)})_{l \in \mathbb{L}_j^*}\} = \varepsilon_k^{(j)}$. De plus, pour tout $l \in \{1, \dots, q\}$, $j \in \mathbb{L}_l \Leftrightarrow l \in \mathbb{L}_j^*$. Les itérations de l'algorithme 7.3 sont donc équivalentes

à

$$\begin{aligned}
 & \text{pour } k = 0, 1, \dots \\
 & \left[\begin{array}{l}
 \text{pour } j = 1, \dots, p \\
 \left[\begin{array}{l}
 \eta_k^{(j)} = \max \{ \varepsilon_k^{(j)}, (\varepsilon_k^{(p+l)})_{l \in \mathbb{L}_j^*} \} \\
 y_k^{(j)} = \eta_k^{(j)} \left(\mathbf{J}_{\mathbf{W}_j \mathbf{A}_j} (x_k^{(j)} - \mathbf{W}_j (\sum_{l=1}^q \mathbf{L}_{l,j}^* v_k^{(l)} + \mathbf{C}_j x_k^{(j)} + c_k^{(j)})) + a_k^{(j)} \right) \\
 x_{k+1}^{(j)} = x_k^{(j)} + \lambda_k \varepsilon_k^{(j)} (y_k^{(j)} - x_k^{(j)})
 \end{array} \right. \\
 \text{pour } l = 1, \dots, q \\
 \left[\begin{array}{l}
 u_k^{(l)} = \varepsilon_k^{(p+l)} \left(\mathbf{J}_{\mathbf{U}_l \mathbf{B}_l^{-1}} (v_k^{(l)} + \mathbf{U}_l (\sum_{j \in \mathbb{L}_l} \mathbf{L}_{l,j} (2y_k^{(j)} - x_k^{(j)}) - \mathbf{D}_l^{-1} v_k^{(l)} + d_k^{(l)})) + b_k^{(l)} \right) \\
 v_{k+1}^{(l)} = v_k^{(l)} + \lambda_k \varepsilon_k^{(p+l)} (u_k^{(l)} - v_k^{(l)}).
 \end{array} \right.
 \end{array} \right. \quad (7.79)
 \end{aligned}$$

D'autre part, d'après la proposition 7.2(i)-(ii), $\mathbf{Q}$ est maximal monotone, $\mathbf{R}$ est cocoercif, et $\mathbf{Z} = \text{zer}(\mathbf{Q} + \mathbf{R}) \neq \emptyset$. On peut remarquer que (7.16) et (7.19) impliquent que $\|\mathbf{U}^{1/2} \mathbf{L} \mathbf{W}^{1/2}\| < 1$. Ainsi, d'après le lemme 7.2, l'algorithme (7.79) peut être réécrit sous la forme de l'algorithme 7.2, où $m = p + q$, $\mathbf{V}$ est défini par (7.15) et, pour tout $k \in \mathbb{N}$,

$$\mathbf{z}_k = (\mathbf{x}_k, \mathbf{v}_k), \quad (7.80)$$

$$\gamma_k = 1, \quad (7.81)$$

$$\mathbf{J}_{\mathbf{V}\mathbf{Q}}: \mathbf{z} \mapsto (\mathbf{T}_{i,k} \mathbf{z})_{1 \leq i \leq m}, \quad (7.82)$$

$$(\forall j \in \{1, \dots, p\}) \quad \mathbf{T}_{j,k}: \mathcal{K} \rightarrow \mathcal{H}_j, \quad (7.83)$$

$$(\forall l \in \{1, \dots, q\}) \quad \mathbf{T}_{p+l,k}: \mathcal{K} \rightarrow \mathcal{G}_l, \quad (7.84)$$

$$\mathbf{t}_k = (\mathbf{a}_k, \mathbf{b}_k), \quad (7.85)$$

$$\mathbf{s}_k = ((\mathbf{W}^{-1} - \mathbf{L}^* \mathbf{U} \mathbf{L})^{-1} (\mathbf{L}^* \mathbf{U} \mathbf{e}_k - \mathbf{c}_k), (\mathbf{U}^{-1} - \mathbf{L} \mathbf{W} \mathbf{L}^*)^{-1} (\mathbf{e}_k - \mathbf{L} \mathbf{W} \mathbf{c}_k)), \quad (7.86)$$

$$\mathbf{e}_k = 2\mathbf{L} \mathbf{a}_k + \mathbf{d}_k. \quad (7.87)$$

Puisque $\mathbf{W} \in \mathcal{S}^+(\mathcal{H})$ et $\mathbf{U} \in \mathcal{S}^+(\mathcal{G})$ vérifient $\|\mathbf{U}^{1/2} \mathbf{L} \mathbf{W}^{1/2}\| < 1$, le lemme 7.1(i) assure que $\mathbf{V} \in \mathcal{S}^+(\mathcal{K})$. De plus, pour tout $(\mathbf{x}, \tilde{\mathbf{x}}) \in \mathcal{H}^2$,

$$\begin{aligned}
 \langle \mathbf{x} - \tilde{\mathbf{x}} \mid \mathbf{W}^{1/2} \mathbf{C} \mathbf{W}^{1/2} \mathbf{x} - \mathbf{W}^{1/2} \mathbf{C} \mathbf{W}^{1/2} \tilde{\mathbf{x}} \rangle &= \sum_{j=1}^p \langle \mathbf{x}^{(j)} - \tilde{\mathbf{x}}^{(j)} \mid \mathbf{W}_j^{1/2} \mathbf{C}_j \mathbf{W}_j^{1/2} \mathbf{x}^{(j)} - \mathbf{W}_j^{1/2} \mathbf{C}_j \mathbf{W}_j^{1/2} \tilde{\mathbf{x}}^{(j)} \rangle \\
 &\geq \sum_{j=1}^p \mu_j \|\mathbf{W}_j^{1/2} \mathbf{C}_j \mathbf{W}_j^{1/2} \mathbf{x}^{(j)} - \mathbf{W}_j^{1/2} \mathbf{C}_j \mathbf{W}_j^{1/2} \tilde{\mathbf{x}}^{(j)}\|^2 \\
 &\geq \mu \|\mathbf{W}^{1/2} \mathbf{C} \mathbf{W}^{1/2} \mathbf{x} - \mathbf{W}^{1/2} \mathbf{C} \mathbf{W}^{1/2} \tilde{\mathbf{x}}\|^2.
 \end{aligned}$$

Donc $\mathbf{W}^{1/2} \mathbf{C} \mathbf{W}^{1/2}$ est μ -cocoercif, et en procédant de façon similaire, on déduit que $\mathbf{U}^{1/2} \mathbf{D}^{-1} \mathbf{U}^{1/2}$ est ν -cocoercif. En conséquence du lemme 7.1(ii), $\mathbf{V}^{1/2} \mathbf{R} \mathbf{V}^{1/2}$ est donc ϑ_α -cocoercif, et nos hypothèses garantissent que $1 = \sup_{k \in \mathbb{N}} \gamma_k < 2\vartheta_\alpha$. D'autre part, nous pouvons déduire de la condi-

tion (i) et des équations (7.85)-(7.87) que

$$\begin{aligned}
 & \sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{t}_k\|^2 | \mathcal{X}_k)} \leq \sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{a}_k\|^2 | \mathcal{X}_k)} + \sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{b}_k\|^2 | \mathcal{X}_k)} < +\infty, \\
 & \sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{s}_k\|^2 | \mathcal{X}_k)} \\
 & \leq \|(\mathbf{W}^{-1} - \mathbf{L}^* \mathbf{U} \mathbf{L})^{-1}\| \left(\sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{c}_k\|^2 | \mathcal{X}_k)} + 2\|\mathbf{L}^* \mathbf{U} \mathbf{L}\| \sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{a}_k\|^2 | \mathcal{X}_k)} \right) \\
 & + \|\mathbf{L}^* \mathbf{U}\| \sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{d}_k\|^2 | \mathcal{X}_k)} + \|(\mathbf{U}^{-1} - \mathbf{L} \mathbf{W} \mathbf{L}^*)^{-1}\| \left(2\|\mathbf{L}\| \sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{a}_k\|^2 | \mathcal{X}_k)} \right) \\
 & + \sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{d}_k\|^2 | \mathcal{X}_k)} + \|\mathbf{L} \mathbf{W}\| \sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{c}_k\|^2 | \mathcal{X}_k)} < +\infty.
 \end{aligned}$$

De plus, puisque nous avons supposé que, pour tout $j \in \{1, \dots, p\}$, $\mathbb{L}_j^* \neq \emptyset$, (ii) et (iii) garantissent que la condition (ii) de la proposition 7.1 est aussi vérifiée. Par conséquent, toutes les hypothèses de la proposition 7.1 sont satisfaites, ce qui nous permet d'établir la convergence presque sûre de $(\mathbf{x}_k, \mathbf{v}_k)_{k \in \mathbb{N}}$ vers une variable aléatoire à valeurs dans $\mathbf{Z}$. Finalement, la proposition 7.2(iii) nous assure que la limite est à valeurs dans $\mathbf{F} \times \mathbf{F}^*$.

7.C DÉMONSTRATION DE LA PROPOSITION 7.5

Tout d'abord, remarquons que, d'après la condition (iii) de la proposition 7.4, puisque $(\forall j \in \{1, \dots, p\}) \mathbb{L}_j^* \neq \emptyset$, les itérations de l'algorithme 7.5 sont équivalentes à

$$\begin{aligned}
 & \text{pour } k = 0, 1, \dots \\
 & \quad \left| \begin{array}{l} \text{pour } l = 1, \dots, q \\ \quad \left[\zeta_k^{(l)} = \max \{ \varepsilon_k^{(p+l)}, (\varepsilon_k^{(j)})_{j \in \mathbb{L}_l} \} \right. \\ \text{pour } j = 1, \dots, p \\ \quad \left[\begin{array}{l} \eta_k^{(j)} = \max \{ \varepsilon_k^{(j)}, (\zeta_k^{(l)})_{l \in \mathbb{L}_j^*} \} \\ w_k^{(j)} = \eta_k^{(j)} (x_k^{(j)} - \mathbf{W}_j (\mathbf{C}_j x_k^{(j)} + c_k^{(j)})) \end{array} \right. \\ \text{pour } l = 1, \dots, q \\ \quad \left[\begin{array}{l} u_k^{(l)} = \zeta_k^{(l)} \left(\mathbf{J}_{\mathbf{U}_l \mathbf{B}_l^{-1}} (v_k^{(l)} + \mathbf{U}_l (\sum_{j \in \mathbb{L}_l} \mathbf{L}_{l,j} (w_k^{(j)} - \mathbf{W}_j \sum_{l' \in \mathbb{L}_j^*} \mathbf{L}_{l',j}^* v_k^{(l')}) \right. \right. \\ \quad \left. \left. - \mathbf{D}_l^{-1} v_k^{(l)} + d_k^{(l)})) + b_k^{(l)} \right) \right. \\ v_{k+1}^{(l)} = v_k^{(l)} + \lambda_k \varepsilon_k^{(p+l)} (u_k^{(l)} - v_k^{(l)}) \\ \text{pour } j = 1, \dots, p \\ \quad \left[x_{k+1}^{(j)} = x_k^{(j)} + \lambda_k \varepsilon_k^{(j)} \left(w_k^{(j)} - \mathbf{W}_j \sum_{l \in \mathbb{L}_j^*} \mathbf{L}_{l,j}^* u_k^{(l)} - x_k^{(j)} \right) \right. \end{array} \right. \end{array} \right. \quad (7.88)
 \end{aligned}$$

De plus, la condition (7.29) implique que $\|\mathbf{U}^{1/2} \mathbf{L} \mathbf{W}^{1/2}\| < 1$. Ainsi, le lemme 7.4 nous assure l'équivalence entre les algorithmes (7.88) et 7.2 lorsque $\mathbf{V}$ est donné par (7.25), $\mathbf{Q}$ est donné par

(7.12) (avec $\mathbf{A} = \mathbf{0}$), et $\mathbf{R}$ est donné par (7.13), dès lors que, pour tout $k \in \mathbb{N}$, (7.80)-(7.84) sont satisfaites et que

$$\begin{aligned} \mathbf{t}_k &= (\mathbf{0}, \mathbf{b}_k), \\ \mathbf{s}_k &= (-\mathbf{W}\mathbf{e}_{1,k}, (\mathbf{U}^{-1} - \mathbf{L}\mathbf{W}\mathbf{L}^*)^{-1}(\mathbf{L}\mathbf{W}\mathbf{L}^*\mathbf{b}_k + \mathbf{d}_k)), \\ \mathbf{e}_{1,k} &= \mathbf{L}^*\mathbf{b}_k + \mathbf{c}_k. \end{aligned}$$

Dans la démonstration de la proposition 7.3, nous avons vu que $\mathbf{W}^{1/2}\mathbf{C}\mathbf{W}^{1/2}$ est μ -cocoercif et $\mathbf{U}^{1/2}\mathbf{D}^{-1}\mathbf{U}^{1/2}$ est ν -cocoercif. D'après le lemme 7.3(ii), $\mathbf{V}^{1/2}\mathbf{R}\mathbf{V}^{1/2}$ est donc ϑ -cocoercif, où ϑ est donné par (7.26), et (7.29) assure que $1 = \sup_{k \in \mathbb{N}} \gamma_k < 2\vartheta$. De plus,

$$\begin{aligned} \sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{t}_k\|^2 | \mathcal{X}_k)} &= \sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{b}_k\|^2 | \mathcal{X}_k)} < +\infty, \\ \sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{s}_k\|^2 | \mathcal{X}_k)} &\leq \|\mathbf{W}\mathbf{L}^*\| \sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{b}_k\|^2 | \mathcal{X}_k)} + \|\mathbf{W}\| \sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{c}_k\|^2 | \mathcal{X}_k)} \\ &\quad + \|(\mathbf{U}^{-1} - \mathbf{L}\mathbf{W}\mathbf{L}^*)^{-1}\| \left(\|\mathbf{L}\mathbf{W}\mathbf{L}^*\| \sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{b}_k\|^2 | \mathcal{X}_k)} + \sum_{k \in \mathbb{N}} \sqrt{\mathbb{E}(\|\mathbf{d}_k\|^2 | \mathcal{X}_k)} \right) < +\infty. \end{aligned}$$

Puisque nous avons supposé que, pour tout $l \in \{1, \dots, q\}$, $\mathbb{L}_l \neq \emptyset$, les conditions (ii)-(iii) de la proposition 7.4 assurent que la condition (ii) de la proposition 7.1 est satisfaite. La convergence est obtenue en appliquant cette proposition.

7.D DÉMONSTRATION DE LA PROPOSITION 7.7

Posons

$$\begin{aligned} (\forall j \in \{1, \dots, p\}) \quad \mathbf{A}_j &= \partial \mathbf{f}_j, & \mathbf{C}_j &= \nabla \mathbf{h}_j, \\ (\forall l \in \{1, \dots, q\}) \quad \mathbf{B}_l &= \partial \mathbf{g}_l, & \mathbf{D}_l^{-1} &= \nabla \mathbf{l}_l^*. \end{aligned}$$

On peut alors remarquer que, pour tout $j \in \{1, \dots, p\}$ et $l \in \{1, \dots, q\}$,

$$\mathbf{J}_{\mathbf{W}_j \mathbf{A}_j} = \text{prox}_{\mathbf{W}_j^{-1}, \mathbf{f}_j} \quad \text{et} \quad \mathbf{J}_{\mathbf{U}_l \mathbf{B}_l^{-1}} = \text{prox}_{\mathbf{U}_l^{-1}, \mathbf{g}_l^*},$$

et que les hypothèses de Lipschitz-différentiabilité faites sur $\mathbf{h}_j$ et $\mathbf{l}_l^*$ sont équivalentes aux hypothèses que $\mathbf{W}_j^{1/2}\mathbf{C}_j\mathbf{W}_j^{1/2}$ est μ_j -cocoercif et $\mathbf{U}_l^{1/2}\mathbf{D}_l^{-1}\mathbf{U}_l^{1/2}$ est ν_l -cocoercif [Bauschke et Combettes, 2011, Cor. 16.42 & 18.16]. La proposition 7.3 nous assure donc que $(\mathbf{x}_k)_{k \in \mathbb{N}}$ converge faiblement P-presque sûrement vers une variable aléatoire à valeurs dans $\mathbf{F}$ et que $(\mathbf{v}_k)_{k \in \mathbb{N}}$ converge faiblement P-presque sûrement vers une variable aléatoire à valeurs dans $\mathbf{F}^*$, où $\mathbf{F}$ et $\mathbf{F}^*$ sont définis dans le problème 7.1.

Il nous reste à montrer que la première limite est une variable aléatoire à valeurs dans $\tilde{\mathbf{F}}$, et que la seconde est une variable aléatoire à valeurs dans $\tilde{\mathbf{F}}^*$. On définit les fonctions séparables

$\mathbf{f} \in \Gamma_0(\mathcal{H})$, $\mathbf{h} \in \Gamma_0(\mathcal{H})$, $\mathbf{g} \in \Gamma_0(\mathcal{G})$, et $\mathbf{l} \in \Gamma_0(\mathcal{G})$ par

$$\begin{aligned} \mathbf{f}: \mathbf{x} &\mapsto \sum_{j=1}^p \mathbf{f}_j(\mathbf{x}^{(j)}), & \mathbf{h}: \mathbf{x} &\mapsto \sum_{j=1}^p \mathbf{h}_j(\mathbf{x}^{(j)}), \\ \mathbf{g}: \mathbf{v} &\mapsto \sum_{l=1}^q \mathbf{g}_l(\mathbf{v}^{(l)}), & \mathbf{l}: \mathbf{v} &\mapsto \sum_{l=1}^q \mathbf{l}_l(\mathbf{v}^{(l)}). \end{aligned}$$

D'après [Bauschke et Combettes, 2011, Prop. 16.8], (7.32) peut être réécrite comme suit

$$\mathbf{0} \in \partial \mathbf{f}(\bar{\mathbf{x}}) + \nabla \mathbf{h}(\bar{\mathbf{x}}) + \mathbf{L}^*(\partial \mathbf{g} \square \partial \mathbf{l})(\mathbf{L}\bar{\mathbf{x}}). \quad (7.89)$$

Par [Bauschke et Combettes, 2011, Prop. 16.38 & 17.26], puisque $\text{dom } \mathbf{h} = \mathcal{H}$, on a $\partial \mathbf{f} + \nabla \mathbf{h} = \partial(\mathbf{f} + \mathbf{h})$, et puisque $\text{dom } \mathbf{l}^* = \mathcal{G}$, d'après [Bauschke et Combettes, 2011, Prop. 24.27] on a $\partial \mathbf{g} \square \partial \mathbf{l} = \partial(\mathbf{g} \square \mathbf{l})$. L'équation (7.89) implique que $\mathbf{L}(\text{dom}(\mathbf{f} + \mathbf{h})) \cap \text{dom}(\mathbf{g} \square \mathbf{l}) \neq \emptyset$ (d'après [Bauschke et Combettes, 2011, Prop. 16.3(i)]) et on peut déduire de [Bauschke et Combettes, 2011, Prop. 16.5] que

$$(\forall \mathbf{x} \in \mathcal{H}) \quad \partial \mathbf{f}(\mathbf{x}) + \nabla \mathbf{h}(\mathbf{x}) + \mathbf{L}^*(\partial \mathbf{g} \square \partial \mathbf{l})(\mathbf{L}\mathbf{x}) \subset \partial(\mathbf{f} + \mathbf{h} + (\mathbf{g} \square \mathbf{l}) \circ \mathbf{L})(\mathbf{x}).$$

En utilisant (7.10) et la règle de Fermat [Bauschke et Combettes, 2011, Thm. 16.2], nous pouvons conclure que

$$\mathbf{F} = \text{zer}(\partial \mathbf{f} + \nabla \mathbf{h} + \mathbf{L}^*(\partial \mathbf{g} \square \partial \mathbf{l})\mathbf{L}) \subset \text{zer}(\partial(\mathbf{f} + \mathbf{h} + (\mathbf{g} \square \mathbf{l}) \circ \mathbf{L})) = \tilde{\mathbf{F}}.$$

Par des arguments similaires, le fait que $\mathbf{F}^* = \text{zer}(-\mathbf{L}(\partial \mathbf{f}^* \square \partial \mathbf{h}^*)(-\mathbf{L}^*) + \partial \mathbf{g}^* + \nabla \mathbf{l}^*) \neq \emptyset$ nous permet de montrer que $\mathbf{F}^* \subset \tilde{\mathbf{F}}^*$.

7.E DÉMONSTRATION DE LA PROPOSITION 7.11

En utilisant la proposition 7.10, la remarque 7.8, les équations (7.47), (7.50)-(7.51), en posant

$$(\forall \ell \in \{1, \dots, r\}) \quad \mathbf{U}_{m+\ell} = \theta^{(\ell)} \text{Id} \quad (7.90)$$

$$(\forall k \in \mathbb{N}) \quad b_k^{(m+\ell)} = d_k^{(m+\ell)} = 0, \quad (7.91)$$

et en remarquant que $J_{U_{m+\ell} N_{\Lambda_{\kappa(\ell)}}^{-1}} = \text{Id} - \theta^{(\ell)} \Pi_{\Lambda_{\kappa(\ell)}} (\cdot / \theta^{(\ell)}) = \text{Id} - \Pi_{\Lambda_{\kappa(\ell)}}$ (d'après la formule de Moreau (7.1)), l'algorithme 7.4 pour résoudre le problème (7.46) se réécrit

$$\begin{aligned}
 & \text{pour } k = 0, 1, \dots \\
 & \quad \left[\begin{array}{l} \text{pour } i = 1, \dots, m \\ \quad \left[\begin{array}{l} u_k^{(i)} = \varepsilon_k^{(m+i)} \left(J_{U_i B_i^{-1}} (v_k^{(i)} + U_i (M_i x_k^{(i)} - D_i^{-1} v_k^{(i)} + d_k^{(i)})) + b_k^{(i)} \right) \\ v_{k+1}^{(i)} = v_k^{(i)} + \lambda_k \varepsilon_k^{(m+i)} (u_k^{(i)} - v_k^{(i)}) \end{array} \right. \\ \text{pour } \ell = 1, \dots, r \\ \quad \left[\begin{array}{l} u_k^{(m+\ell)} = \varepsilon_k^{(2m+\ell)} (v_k^{(m+\ell)} + \theta^{(\ell)} (x_k^{(i)})_{i \in \mathbb{V}_\ell} - \Pi_{\Lambda_{\kappa(\ell)}} (v_k^{(m+\ell)} + \theta^{(\ell)} (x_k^{(i)})_{i \in \mathbb{V}_\ell})) \\ v_{k+1}^{(m+\ell)} = v_k^{(m+\ell)} + \lambda_k \varepsilon_k^{(2m+\ell)} (u_k^{(m+\ell)} - v_k^{(m+\ell)}) \end{array} \right. \\ \text{pour } i = 1, \dots, m \\ \quad \left[\begin{array}{l} y_k^{(i)} = \varepsilon_k^{(i)} \left(J_{W_i A_i} (x_k^{(i)} - W_i (M_i^* (2u_k^{(i)} - v_k^{(i)}) + \sum_{(\ell,j) \in \mathbb{V}_i^*} (2u_k^{(m+\ell,j)} - v_k^{(m+\ell,j)}) \right. \\ \quad \left. + C_i x_k^{(i)} + c_k^{(i)})) + a_k^{(i)} \right) \\ x_{k+1}^{(i)} = x_k^{(i)} + \lambda_k \varepsilon_k^{(i)} (y_k^{(i)} - x_k^{(i)}). \end{array} \right. \end{array} \right. \quad (7.92)
 \end{aligned}$$

Or, pour tout $\ell \in \{1, \dots, r\}$, la projection dans l'espace vectoriel $\Lambda_{\kappa(\ell)}$ a une forme explicite donnée par

$$(\forall k \in \mathbb{N}) \quad \Pi_{\Lambda_{\kappa(\ell)}} (v_k^{(m+\ell)} + \theta^{(\ell)} (x_k^{(i)})_{i \in \mathbb{V}_\ell}) = \frac{1}{\kappa(\ell)} \left(\sum_{j=1}^{\kappa(\ell)} v_k^{(m+\ell,j)} + \theta^{(\ell)} \sum_{i \in \mathbb{V}_\ell} x_k^{(i)} \right).$$

Ainsi, l'algorithme (7.92) se réécrit

$$\begin{aligned}
 & \text{pour } k = 0, 1, \dots \\
 & \quad \left[\begin{array}{l} \text{pour } i = 1, \dots, m \\ \quad \left[\begin{array}{l} u_k^{(i)} = \varepsilon_k^{(m+i)} \left(J_{U_i B_i^{-1}} (v_k^{(i)} + U_i (M_i x_k^{(i)} - D_i^{-1} v_k^{(i)} + d_k^{(i)})) + b_k^{(i)} \right) \\ v_{k+1}^{(i)} = v_k^{(i)} + \lambda_k \varepsilon_k^{(m+i)} (u_k^{(i)} - v_k^{(i)}) \end{array} \right. \\ \text{pour } \ell = 1, \dots, r \\ \quad \left[\begin{array}{l} \bar{u}_k^{(\ell)} = \varepsilon_k^{(2m+\ell)} \frac{1}{\kappa(\ell)} \left(\sum_{j=1}^{\kappa(\ell)} v_k^{(m+\ell,j)} + \theta^{(\ell)} \sum_{i \in \mathbb{V}_\ell} x_k^{(i)} \right) \\ \text{pour } j = 1, \dots, \kappa(\ell) \\ \quad \left[\begin{array}{l} u_k^{(m+\ell,j)} = \varepsilon_k^{(2m+\ell)} (v_k^{(m+\ell,j)} + \theta^{(\ell)} x_k^{(i(\ell,j))} - \bar{u}_k^{(\ell)}) \\ v_{k+1}^{(m+\ell,j)} = v_k^{(m+\ell,j)} + \lambda_k \varepsilon_k^{(2m+\ell)} (u_k^{(m+\ell,j)} - v_k^{(m+\ell,j)}) \end{array} \right. \\ \text{pour } i = 1, \dots, m \\ \quad \left[\begin{array}{l} y_k^{(i)} = \varepsilon_k^{(i)} \left(J_{W_i A_i} (x_k^{(i)} - W_i (M_i^* (2u_k^{(i)} - v_k^{(i)}) + \sum_{(\ell,j) \in \mathbb{V}_i^*} (2u_k^{(m+\ell,j)} - v_k^{(m+\ell,j)}) \right. \\ \quad \left. + C_i x_k^{(i)} + c_k^{(i)})) + a_k^{(i)} \right) \\ x_{k+1}^{(i)} = x_k^{(i)} + \lambda_k \varepsilon_k^{(i)} (y_k^{(i)} - x_k^{(i)}). \end{array} \right. \end{array} \right. \quad (7.93)
 \end{aligned}$$

Posons, pour tout $k \in \mathbb{N}$ et pour tout $\ell \in \{1, \dots, r\}$,

$$\bar{x}_k^{(\ell)} = \frac{1}{\kappa^{(\ell)}} \sum_{i \in \mathbb{V}_\ell} x_k^{(i)}, \quad (7.94)$$

$$w_k^{(\ell)} = \varepsilon_k^{(2m+\ell)} (2u_k^{(m+\ell)} - v_k^{(m+\ell)}), \quad (7.95)$$

$$\eta_k^{(\ell)} = \max \{ \varepsilon_k^{(i)} \mid i \in \mathbb{V}_\ell \}. \quad (7.96)$$

En utilisant (7.57) et l'étape de mise à jour

$$\bar{x}_{k+1}^{(\ell)} = \bar{x}_k^{(\ell)} + \eta_k^{(\ell)} \left(\frac{1}{\kappa^{(\ell)}} \sum_{i \in \mathbb{V}_\ell} x_{k+1}^{(i)} - \bar{x}_k^{(\ell)} \right),$$

les itérations de l'algorithme 7.11 sont obtenues en réordonnant les étapes dans (7.93).

Afin d'appliquer la proposition 7.4, il nous faut maintenant démontrer que la condition (7.19), où ϑ_α est défini par (7.16), est satisfaite. Soit

$$\mathbf{W}: \mathcal{H}^m \rightarrow \mathcal{H}^m: \mathbf{x} \mapsto (\mathbf{W}_i \mathbf{x}^{(i)})_{1 \leq i \leq m} \quad \text{et} \quad \mathbf{U}: \mathcal{G} \oplus \mathcal{H} \rightarrow \mathcal{G} \oplus \mathcal{H}: \mathbf{v} \mapsto (\mathbf{U}_l \mathbf{v}^{(l)})_{1 \leq l \leq m+r}.$$

D'après la remarque 7.8 et l'équation (7.90), on a

$$(\forall \mathbf{x} \in \mathcal{H}^m) \quad \mathbf{U}^{1/2} \mathbf{L} \mathbf{W}^{1/2} \mathbf{x} = (\mathbf{U}_1^{1/2} \mathbf{M} \mathbf{W}^{1/2} \mathbf{x}, \mathbf{U}_2^{1/2} \mathbf{S} \mathbf{W}^{1/2} \mathbf{x})$$

où $\mathbf{U}_1: \mathcal{G} \rightarrow \mathcal{G}: (\mathbf{v}^{(i)})_{1 \leq i \leq m} \mapsto (\mathbf{U}_i \mathbf{v}^{(i)})_{1 \leq i \leq m}$ et $\mathbf{U}_2: \mathcal{H} \rightarrow \mathcal{H}: (\mathbf{v}^{(m+\ell)})_{1 \leq \ell \leq r} \mapsto (\theta^{(\ell)} \mathbf{v}^{(m+\ell)})_{1 \leq \ell \leq r}$. On en déduit que

$$\begin{aligned} \|\mathbf{U}^{1/2} \mathbf{L} \mathbf{W}^{1/2} \mathbf{x}\|^2 &= \|\mathbf{U}_1^{1/2} \mathbf{M} \mathbf{W}^{1/2} \mathbf{x}\|^2 + \|\mathbf{U}_2^{1/2} \mathbf{S} \mathbf{W}^{1/2} \mathbf{x}\|^2 \\ &= \sum_{i=1}^m \|\mathbf{U}_i^{1/2} \mathbf{M}_i \mathbf{W}_i^{1/2} \mathbf{x}^{(i)}\|^2 + \left\langle \mathbf{W}^{1/2} \mathbf{x} \mid \mathbf{S}^* \mathbf{U}_2 \mathbf{S} \mathbf{W}^{1/2} \mathbf{x} \right\rangle. \end{aligned}$$

En utilisant (7.47) et (7.50)-(7.51), on peut de plus remarquer que

$$\mathbf{S}^* \mathbf{U}_2 \mathbf{S}: \mathcal{H}^m \rightarrow \mathcal{H}^m: (\mathbf{x}^{(i)})_{1 \leq i \leq m} \mapsto (\bar{\theta}_i \mathbf{x}^{(i)})_{1 \leq i \leq m}$$

ce qui implique

$$\begin{aligned} \|\mathbf{U}^{1/2} \mathbf{L} \mathbf{W}^{1/2} \mathbf{x}\|^2 &= \sum_{i=1}^m \|\mathbf{U}_i^{1/2} \mathbf{M}_i \mathbf{W}_i^{1/2} \mathbf{x}^{(i)}\|^2 + \sum_{i=1}^m \bar{\theta}_i \|\mathbf{W}_i^{1/2} \mathbf{x}^{(i)}\|^2 \\ &\leq \sum_{i=1}^m (\|\mathbf{U}_i^{1/2} \mathbf{M}_i \mathbf{W}_i^{1/2}\|^2 + \bar{\theta}_i \|\mathbf{W}_i\|) \|\mathbf{x}^{(i)}\|^2 \leq \chi \|\mathbf{x}\|^2, \end{aligned}$$

et on en déduit donc que $\|\mathbf{U}^{1/2} \mathbf{L} \mathbf{W}^{1/2}\|^2 \leq \chi$. Ainsi, (7.54) implique (7.19).

De plus, la condition (iii) de la proposition 7.4 se réécrit sous la forme de la condition (iii) de l'actuelle proposition. On peut donc déduire des propositions 7.4 et 7.10 que, pour tout $i \in \{1, \dots, m\}$, $(x_k^{(i)})_{k \in \mathbb{N}}$ converge faiblement P-presque sûrement vers une variable aléatoire $\hat{x}$ à valeurs dans $\hat{F}$. En conséquence de (7.94), pour tout $\ell \in \{1, \dots, r\}$, $(\bar{x}_k^{(\ell)})_{k \in \mathbb{N}}$ converge aussi faiblement P-presque sûrement vers $\hat{x}$.

7.F DÉMONSTRATION DE LA PROPOSITION 7.12

En choisissant $(\mathbf{U}_{m+\ell})_{1 \leq \ell \leq r}$ comme dans (7.90) et en considérant des termes d'erreur nuls comme dans (7.91), l'algorithme 7.5 pour résoudre le problème (7.46) peut être réécrit de la façon suivante :

$$\begin{array}{l}
 \text{pour } k = 0, 1, \dots \\
 \left[\begin{array}{l}
 \text{pour } i = 1, \dots, m \\
 \left[\begin{array}{l}
 \eta_k^{(i)} = \max \{ \varepsilon_k^{(m+i)}, (\varepsilon_k^{(2m+\ell)})_{\ell \in \{\ell' \in \{1, \dots, r\} | i \in \mathbb{V}_{\ell'}\}} \} \\
 w_k^{(i)} = \eta_k^{(i)} (x_k^{(i)} - \mathbf{W}_i(\mathbf{C}_i x_k^{(i)} + c_k^{(i)})) \\
 \tilde{w}_k^{(i)} = \eta_k^{(i)} (w_k^{(i)} - \mathbf{W}_i(\mathbf{M}_i^* v_k^{(i)} + \sum_{(\ell, j) \in \mathbb{V}_i^*} v_k^{(m+\ell, j)})) \\
 u_k^{(i)} = \varepsilon_k^{(m+i)} (\mathbf{J}_{\mathbf{U}_i \mathbf{B}_i^{-1}}(v_k^{(i)} + \mathbf{U}_i(\mathbf{M}_i \tilde{w}_k^{(i)} - \mathbf{D}_i^{-1} v_k^{(i)} + d_k^{(i)})) + b_k^{(i)}) \\
 v_{k+1}^{(i)} = v_k^{(i)} + \lambda_k \varepsilon_k^{(m+i)} (u_k^{(i)} - v_k^{(i)})
 \end{array} \right. \\
 \text{pour } \ell = 1, \dots, r \\
 \left[\begin{array}{l}
 u_k^{(m+\ell)} = \varepsilon_k^{(2m+\ell)} (v_k^{(m+\ell)} + \theta^{(\ell)} (\tilde{w}_k^{(i)})_{i \in \mathbb{V}_\ell} - \Pi_{\Lambda_{\kappa(\ell)}}(v_k^{(m+\ell)} + \theta^{(\ell)} (\tilde{w}_k^{(i)})_{i \in \mathbb{V}_\ell})) \\
 v_{k+1}^{(m+\ell)} = v_k^{(m+\ell)} + \lambda_k \varepsilon_k^{(2m+\ell)} (u_k^{(m+\ell)} - v_k^{(m+\ell)})
 \end{array} \right. \\
 \text{pour } i = 1, \dots, m \\
 \left[\begin{array}{l}
 x_{k+1}^{(i)} = x_k^{(i)} + \lambda_k \varepsilon_k^{(i)} (w_k^{(i)} - \mathbf{W}_i(\mathbf{M}_i^* u_k^{(i)} + \sum_{(\ell, j) \in \mathbb{V}_i^*} u_k^{(m+\ell, j)}) - x_k^{(i)}).
 \end{array} \right.
 \end{array}
 \right.
 \end{array}$$

Le reste de la démonstration est identique à celle de la proposition 7.11.

CHAPITRE 8

Opérateur proximal de fonctions quotient

Sommaire

8.1	Introduction	204
8.2	Fonctions quotient	205
8.3	Opérateur proximal de fonctions quotient	206
8.3.1	Opérateur proximal de la somme des fonctions quotient . . .	206
8.3.2	Opérateur proximal du maximum des fonctions quotient . . .	210
8.4	Application à l'estimation de la sélectivité	214
8.4.1	Estimation de la sélectivité en SGBD	214
8.4.2	Formulation du problème	214
8.4.3	Solution du problème en SGBD	216
8.5	Conclusion	219

8.1 INTRODUCTION

Le but du travail présenté dans ce chapitre est l'optimisation des requêtes dans les *Systèmes de Gestion de Bases de Données* (SGBD ou DBMS en anglais pour *Database Management Systems*). Dans ce contexte, le plan d'exécution des requêtes optimales dépend de la précision de l'estimation des sélectivités (la sélectivité correspond à la proportion de n -uplets satisfaisants les prédicats d'une requête). Les modèles pour l'estimation des n -uplets, tels que ceux présentés dans [Markl *et al.*, 2007], requièrent la résolution d'un problème de faisabilité. Plus précisément, soient $\mathbf{A} \in \mathbb{R}^{N \times N}$, $\mathbf{x} \in \mathbb{R}^N$ et $\mathbf{b} \in \mathbb{R}^N$. Nous considérons un système linéaire sous-déterminé $\mathbf{Ax} = \mathbf{b}$ n'admettant pas de solution positive ou nulle. Nous allons alors rechercher à remplacer $\mathbf{b}$ par un élément $\hat{\mathbf{b}}$ qui soit "correct" au sens où le système $\mathbf{Ax} = \hat{\mathbf{b}}$ admette une solution positive ou nulle. Bien que les problèmes de faisabilité aient été étudiés dans divers travaux (voir [Chinneck, 2008; Combettes, 1996], et les références associées), notre approche sera différente de celles existantes par le critère que l'on va proposer de minimiser. Les approches existantes s'intéressent à la minimisation des différences des éléments $\hat{\mathbf{b}}^{(n)}$ et $\mathbf{b}^{(n)}$, pour tout $n \in \{1, \dots, N\}$, tandis que nous minimiserons leurs quotients $\max\{\hat{\mathbf{b}}^{(n)}/\mathbf{b}^{(n)}, \mathbf{b}^{(n)}/\hat{\mathbf{b}}^{(n)}\}$. L'intérêt d'utiliser les quotients par rapport aux différences a été mis en évidence dans [Moerkotte *et al.*, 2009]. Ainsi, dans ce chapitre, nous proposons d'étudier la somme de ces quotients, désignée par $\mathbf{q}_1$, ainsi que leur maximum, noté $\mathbf{q}_\infty$.

Nous proposons d'appliquer un algorithme primal-dual pour résoudre les problèmes faisant intervenir les fonctions quotient $\mathbf{q}_\nu$, $\nu = 1, \infty$. Pour cela, nous allons dans un premier temps analyser les opérateurs proximaux des fonctions quotient $\mathbf{q}_\nu$, $\nu = 1, \infty$. Nous démontrerons que l'opérateur proximal de la somme des quotients, $\mathbf{q}_1$, est une fonction de seuillage. Nous verrons que cette fonction est un peu plus complexe que la fonction de seuillage doux, correspondant à l'opérateur proximal de la norme ℓ_1 , car son calcul fait intervenir une équation du troisième degré. Concernant l'opérateur proximal du maximum des quotients $\mathbf{q}_\infty$, nous montrerons qu'il peut être calculé de façon itérative grâce à un algorithme primal-dual en utilisant des projections épigraphiques terme à terme [Chierchia *et al.*, 2013a]. Nous utiliserons les résultats obtenus pour résoudre le problème de faisabilité décrit ci-dessus, et donnerons un exemple d'application numérique.

Le reste du chapitre sera organisé comme suit : Dans la section 8.2, nous introduirons les distances quotient entre des éléments positifs puis nous les utiliserons pour définir les fonctions quotient de vecteurs d'éléments positifs. Ensuite, dans la section 8.3, nous nous intéresserons à la détermination de l'opérateur proximal des fonctions quotient. Nous proposerons dans la section 8.4 une méthode utilisant ces opérateurs proximaux pour résoudre des problèmes de faisabilité apparaissant, notamment, dans les estimations de n -uplets nécessaires pour l'optimisation de requêtes en SGBD. Après avoir décrit le problème d'estimation des n -uplets, nous proposerons un algorithme de minimisation primal-dual pour le résoudre et illustrerons ses performances par un exemple numérique. Nous concluons dans la section 8.5.

8.2 FONCTIONS QUOTIENT

Nous définissons la fonction quotient q par

$$\begin{aligned} q: [0, +\infty[^2 &\rightarrow [0, +\infty[\\ (x, y) &\mapsto q(x, y) := \frac{\max(x, y)}{\min(x, y)}. \end{aligned} \quad (8.1)$$

Remarquons que la fonction $(x, y) \in [0, +\infty[^2 \mapsto q(x, y) - 1$ vérifie deux des conditions définissant une distance :

- elle est symétrique ($(\forall (x, y) \in [0, +\infty[^2) \quad q(x, y) = q(y, x)$),
- elle vérifie $q(x, y) - 1 = 0$ si et seulement si $x = y$ (puisque $q(x, y) - 1 = \frac{|x-y|}{\min(x, y)}$).

Cependant, elle ne satisfait pas l'inégalité triangulaire. Une fonction similaire à la fonction quotient, appelée *distance relative généralisée*, définie par

$$(x, y) \in \mathbb{R}^* \times \mathbb{R}^* \mapsto \frac{|x - y|}{\max(x, y)}, \quad (8.2)$$

a été utilisée dans [Coutinho et Figueiredo, 2005; Griesel, 1974; Metcalf, 1976; Ziv, 1982]. Le lien entre ces deux fonctions (8.1) et (8.2) a été étudié dans [Setzer *et al.*, 2010]. D'autre part, la fonction quotient (8.1) a été utilisée comme une mesure de contraste en traitement d'image dans [Palma-Amestoy *et al.*, 2009] en raison de sa propriété d'homogénéité, i.e.

$$(\forall (x, y) \in [0, +\infty[^2) (\forall \lambda > 0) \quad q(\lambda x, \lambda y) = q(x, y) \quad (8.3)$$

Soit $b \in]0, +\infty[$ fixé, on se propose d'étendre la fonction $q(\cdot, b)$ à l'ensemble de l'axe des réels de la façon suivante :

$$\begin{aligned} q(\cdot, b): \mathbb{R} &\rightarrow [0, +\infty[\\ x &\mapsto q(x, b) := \begin{cases} \frac{x}{b} & \text{si } b \leq x, \\ \frac{b}{x} & \text{si } 0 < x < b, \\ +\infty & \text{sinon.} \end{cases} \end{aligned} \quad (8.4)$$

Notons que cette fonction est convexe et semi-continue inférieurement, et que, de plus, d'après (8.3), on a

$$(\forall x \in \mathbb{R}) \quad q(x, b) = q(x/b, 1). \quad (8.5)$$

Dans la suite du chapitre nous noterons

$$(\forall x \in \mathbb{R}) \quad \theta(x) := q(x, 1). \quad (8.6)$$

Soit $\mathbf{b} = (b^{(n)})_{1 \leq n \leq N} \in]0, +\infty[^N$. Nous nous intéressons d'une part à la somme de fonctions quotient $\mathbf{q}_1(\cdot, \mathbf{b}): \mathbb{R}^N \rightarrow [0, +\infty]$ et d'autre part à leur maximum $\mathbf{q}_\infty(\cdot, \mathbf{b}): \mathbb{R}^N \rightarrow [0, +\infty]$. Ces deux fonctions sont définies par

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \mathbf{q}_1(\mathbf{x}, \mathbf{b}) := \sum_{n=1}^N q(x^{(n)}, b^{(n)}), \quad (8.7)$$

$$\mathbf{q}_\infty(\mathbf{x}, \mathbf{b}) := \max_{1 \leq n \leq N} q(x^{(n)}, b^{(n)}). \quad (8.8)$$

De plus, pour $\nu \in \{1, \infty\}$, nous posons $(\forall \mathbf{x} \in \mathbb{R}^N) \theta_\nu(\mathbf{x}) := \mathbf{q}_\nu(\mathbf{x}, \mathbf{1}_N)$, où $\mathbf{1}_N$ est le vecteur unitaire de $\mathbb{R}^N$. Notons $\mathbf{B} = \text{Diag}(\mathbf{b})$ la matrice diagonale de $\mathbb{R}^{N \times N}$ dont les éléments diagonaux correspondent aux éléments du vecteur $\mathbf{b} = (\mathbf{b}^{(n)})_{1 \leq n \leq N}$. D'après (8.5), on a alors

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \mathbf{q}_\nu(\mathbf{x}, \mathbf{b}) = \theta_\nu(\mathbf{B}^{-1}\mathbf{x}), \quad (8.9)$$

8.3 OPÉRATEUR PROXIMAL DE FONCTIONS QUOTIENT

Dans cette section, nous nous intéressons aux opérateurs proximaux des fonctions convexes $\mathbf{q}_1(\cdot, \mathbf{b})$ et $\mathbf{q}_\infty(\cdot, \mathbf{b})$, pour $\mathbf{b} \in]0, +\infty[^N$ fixé. D'après l'équation (8.9) on a, pour $\nu \in \{1, \infty\}$,

$$\begin{aligned} (\forall \gamma \in]0, +\infty[)(\forall \mathbf{x} \in \mathbb{R}^N) \quad \text{prox}_{\gamma \mathbf{q}_\nu(\cdot, \mathbf{b})}(\mathbf{x}) &= \underset{\mathbf{y} \in \mathbb{R}^N}{\text{argmin}} \quad \mathbf{q}_\nu(\mathbf{y}, \mathbf{b}) + \frac{1}{2\gamma} \|\mathbf{x} - \mathbf{y}\|^2 \\ &= \underset{\mathbf{y} \in \mathbb{R}^N}{\text{argmin}} \quad \theta_\nu(\mathbf{B}^{-1}\mathbf{y}) + \frac{1}{2\gamma} \|\mathbf{x} - \mathbf{y}\|^2 \\ &= \mathbf{B} \underset{\mathbf{y}' \in \mathbb{R}^N}{\text{argmin}} \quad \theta_\nu(\mathbf{y}') + \frac{1}{2\gamma} \left\langle \mathbf{B}^{-1}\mathbf{x} - \mathbf{y}' \mid \mathbf{B}^2(\mathbf{B}^{-1}\mathbf{x} - \mathbf{y}') \right\rangle \\ &= \mathbf{B} \underset{\mathbf{y}' \in \mathbb{R}^N}{\text{argmin}} \quad \theta_\nu(\mathbf{y}') + \frac{1}{2\gamma} \|\mathbf{B}^{-1}\mathbf{x} - \mathbf{y}'\|_{\mathbf{B}^2}^2 \\ &= \mathbf{B} \text{prox}_{\gamma^{-1}\mathbf{B}^2, \theta_\nu}(\mathbf{B}^{-1}\mathbf{x}), \end{aligned} \quad (8.10)$$

où, pour tout $\alpha \in \mathbb{R}$, $\mathbf{B}^\alpha = \text{Diag}(((\mathbf{b}^{(n)})^\alpha)_{1 \leq n \leq N})$. Par conséquent, il nous suffit d'étudier l'opérateur proximal :

$$(\forall \gamma \in]0, +\infty[)(\forall \mathbf{x} \in \mathbb{R}^N) \quad \text{prox}_{\mathbf{U}, \theta_\nu}(\mathbf{x}) = \underset{\mathbf{y} \in \mathbb{R}^N}{\text{argmin}} \quad \theta_\nu(\mathbf{y}) + \frac{1}{2} \|\mathbf{x} - \mathbf{y}\|_{\mathbf{U}}^2, \quad (8.11)$$

où $\nu \in \{1, \infty\}$ et $\mathbf{U} = \text{Diag}(\mathbf{u})$ pour $\mathbf{u} \in]0, +\infty[^N$.

8.3.1 Opérateur proximal de la somme des fonctions quotient

Pour $\nu = 1$, la fonction θ_1 étant additivement séparable, l'argument minimal de (8.11) peut être calculé terme à terme, i.e.,

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \text{prox}_{\mathbf{U}, \theta_1}(\mathbf{x}) = \left(\text{prox}_{\theta/u^{(n)}}(\mathbf{x}^{(n)}) \right)_{1 \leq n \leq N}. \quad (8.12)$$

Ainsi, il nous suffit d'étudier

$$(\forall x \in \mathbb{R}) \quad \text{prox}_{\gamma \theta}(x) = \underset{y \in \mathbb{R}}{\text{argmin}} \quad \theta(y) + \frac{1}{2\gamma} (x - y)^2, \quad (8.13)$$

pour $\gamma > 0$. L'expression de cet opérateur proximal est donnée dans la proposition suivante.

Proposition 8.1. Soient $\gamma > 0$ et $x \in \mathbb{R}$. L'opérateur proximal de la fonction quotient définie par (8.4)-(8.6) est donné par

$$\text{prox}_{\gamma\theta}(x) = \begin{cases} \zeta^* & \text{si } x < 1 - \gamma, \\ 1 & \text{si } x \in [1 - \gamma, 1 + \gamma], \\ x - \gamma & \text{si } x > 1 + \gamma, \end{cases} \quad (8.14)$$

où ζ^* est l'unique solution dans $]0, 1]$ du polynôme de degré trois

$$r(\zeta) = \zeta^3 - x\zeta^2 - \gamma = 0. \quad (8.15)$$

Démonstration. Soit $\gamma \in]0, +\infty[$. Nous allons appliquer la proposition 2.10 affirmant que, pour toute fonction $\psi \in \Gamma_0(\mathbb{R})$ différentiable en 0 telle que $\dot{\psi}(0) = 0$, on a $\text{prox}_{\psi+\gamma|\cdot|} = \text{prox}_{\psi} \circ \text{soft}_{\gamma}^{-1}$. Pour cela, nous allons décomposer la fonction θ de la façon suivante :

$$(\forall y \in \mathbb{R}) \quad \theta(y) = 1 + \phi(y - 1), \quad (8.16)$$

où $(\forall z \in \mathbb{R}) \quad \phi(z) := \psi(z) + |z|$ avec

$$\psi(z) := \begin{cases} +\infty & \text{si } z \leq -1, \\ z + \frac{1}{1+z} - 1 & \text{si } z \in]-1, 0], \\ 0 & \text{si } z > 0. \end{cases}$$

Notons que la fonction ψ est dans $\Gamma_0(\mathbb{R})$, elle est différentiable en zéro et vérifie $\dot{\psi}(0) = 0$. Ainsi, $(\forall z \in \mathbb{R}) \quad \psi(z) \geq 0$. Intéressons nous à l'opérateur proximal de $\gamma\psi$:

$$(\forall x \in \mathbb{R}) \quad \text{prox}_{\gamma\psi}(x) = \underset{z \in \mathbb{R}}{\text{argmin}} \left\{ \varphi(z) := \psi(z) + \frac{1}{2\gamma}(z - x)^2 \right\}. \quad (8.17)$$

- Étudions le cas où $x \in]0, +\infty[$. D'une part, la fonction $z \in \mathbb{R} \mapsto \psi(z)$ est convexe et sa valeur minimale, égale à 0, est atteinte dès que $z \in]0, +\infty[$. D'autre part, la fonction $z \in \mathbb{R} \mapsto \frac{1}{2}(z - x)^2$ atteint son minimum en $z = x \in]0, +\infty[$, et sa valeur minimale est égale à zéro. Ainsi, la fonction $z \in \mathbb{R} \mapsto \varphi(z)$ est strictement convexe et $(\forall z \in \mathbb{R}) \quad \varphi(z) \geq 0 = \varphi(x)$.
- Étudions le cas où $x \in]-\infty, 0]$. L'expression de φ est donnée par

$$(\forall z \in \mathbb{R}) \quad \varphi(z) = \begin{cases} +\infty & \text{si } z \leq -1, \\ z + \frac{1}{1+z} - 1 + \frac{1}{2\gamma}(z - x)^2 & \text{si } -1 < z \leq 0, \\ \frac{1}{2\gamma}(z - x)^2 & \text{si } 0 < z. \end{cases} \quad (8.18)$$

1. soft_{γ} désigne la fonction de seuillage doux définie par (2.56).

Cette fonction est strictement convexe, et sa dérivée est donnée par

$$(\forall z \in \mathbb{R}) \quad \dot{\varphi}(z) = \begin{cases} +\infty & \text{si } z \leq -1, \\ 1 - \frac{1}{(1+z)^2} + \frac{z-x}{\gamma} & \text{si } -1 < z \leq 0, \\ \frac{z-x}{\gamma} & \text{si } 0 < z. \end{cases} \quad (8.19)$$

Ainsi φ est strictement décroissante sur $] -1, z^*[$, strictement croissante sur $]z^*, 0]$, et strictement croissante sur $]0, +\infty[$, où $z^* \in] -1, 0]$ est la solution de

$$0 = (1 + z^*)^3 - (x + 1 - \gamma)(1 + z^*)^2 - \gamma. \quad (8.20)$$

On peut donc en déduire que

$$(\forall x \in \mathbb{R}) \quad z^* = \text{prox}_{\gamma\psi}(x) = \underset{z \in]-1, 0]}{\operatorname{argmin}} \varphi(z) \quad (8.21)$$

est l'unique solution de (8.20) sur $] -1, 0]$.

On peut donc en conclure que

$$(\forall x \in \mathbb{R}) \quad \text{prox}_{\gamma\psi}(x) = \begin{cases} x & \text{si } x \geq 0, \\ z^* & \text{sinon,} \end{cases} \quad (8.22)$$

où $z^* \in] -1, 0]$ est l'unique solution de (8.20).

D'après la proposition 2.10, on peut alors en déduire que

$$(\forall x \in \mathbb{R}) \quad \text{prox}_{\gamma\phi}(x) = \begin{cases} x - \gamma & \text{si } x > \gamma, \\ 0 & \text{si } x \in [-\gamma, \gamma], \\ t^* & \text{si } x < -\gamma, \end{cases} \quad (8.23)$$

où t^* est solution de

$$(1 + t^*)^3 - (x + 1)(1 + t^*)^2 - \gamma = 0. \quad (8.24)$$

Par conséquent, en utilisant l'équation (8.16) on obtient

$$\begin{aligned} \text{prox}_{\gamma\theta}(x) &= \underset{y \in \mathbb{R}}{\operatorname{argmin}} 1 + \phi(y - 1) + \frac{1}{2\gamma}(y - x)^2 \\ &= 1 + \underset{y \in \mathbb{R}}{\operatorname{argmin}} \phi(y) + \frac{1}{2\gamma}(y - (x - 1))^2 \\ &= 1 + \text{prox}_{\gamma\phi}(x - 1) \\ &= \begin{cases} x - \gamma & \text{si } x > 1 + \gamma, \\ 1 & \text{si } x \in [1 - \gamma, 1 + \gamma], \\ 1 + y^* & \text{si } x < 1 - \gamma, \end{cases} \end{aligned} \quad (8.25)$$

où y^* est l'unique solution dans $]0, 1]$ de

$$(1 + y^*)^3 - (1 + y^*)^2 x - \gamma = 0.$$

En posant $\zeta := 1 + y$ on retrouve alors (8.14)-(8.15). $\square$

Remarque 8.1.

- (i) La racine $\zeta^* \in]0, 1[$ du polynôme de troisième degré r , donné dans (8.15), peut être calculée en utilisant la méthode de Newton [Hanke-Bourgeois, 2002], ou bien, de façon alternative, en utilisant la formule de Cardan [Tignol, 2001].
- (ii) Une illustration graphique de la fonction $\text{prox}_{\gamma\theta}$, pour plusieurs valeurs de γ , est donnée en figure 8.1.

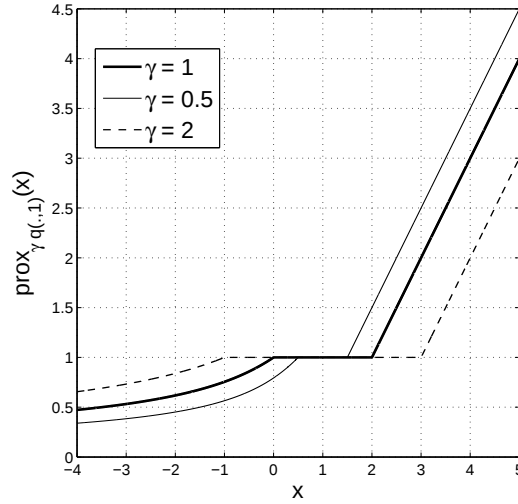


Figure 8.1 — Fonction $\text{prox}_{\gamma q(\cdot, 1)}$ pour différents γ .

Notons que, par un raisonnement similaire à celui utilisé pour obtenir (8.10), on a

$$(\forall \gamma \in]0, +\infty[)(\forall x \in \mathbb{R}) \quad \text{prox}_{\gamma q(\cdot, b)}(x) = b \text{prox}_{\gamma b^{-2}\theta}(x/b), \quad (8.26)$$

où $b \in]0, +\infty[$ est fixé. Ainsi, le corollaire suivant, qui donne l'opérateur proximal de $q(\cdot, b)$, pour $b \in]0, +\infty[$, est une conséquence directe de la proposition 8.1 :

Corollaire 8.1. Soient $\gamma > 0$, $b \in]0, +\infty[$ et $x \in \mathbb{R}$. L'opérateur proximal de $q(\cdot, b)$ en x est donné par

$$\text{prox}_{\gamma q(\cdot, b)}(x) = \begin{cases} x - \frac{\gamma}{b} & \text{si } x > b + \frac{\gamma}{b}, \\ b & \text{si } x \in [b - \frac{\gamma}{b}, b + \frac{\gamma}{b}], \\ \zeta^* & \text{si } x < b - \frac{\gamma}{b}, \end{cases}$$

où ζ^* est l'unique solution sur $]0, b]$ du polynôme de troisième degré

$$\zeta^3 - x\zeta^2 - \gamma b = 0.$$

8.3.2 Opérateur proximal du maximum des fonctions quotient

Nous allons maintenant nous intéresser à

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \text{prox}_{\mathbf{U}, \theta_\infty}(\mathbf{x}) = \underset{\mathbf{y} \in \mathbb{R}^N}{\operatorname{argmin}} \quad \theta_\infty(\mathbf{y}) + \frac{1}{2} \|\mathbf{x} - \mathbf{y}\|_{\mathbf{U}}^2. \quad (8.27)$$

Nous proposons de réécrire ce problème de minimisation sous la forme d'un problème de minimisation sous contrainte épigraphique. Ce type d'approche a été introduite dans [Chierchia *et al.*, 2013a] et a été utilisée, par exemple, dans [Chierchia *et al.*, 2013b; Harizanov *et al.*, 2013].

En utilisant la définition 2.3(vi) de l'épigraphe, le problème de minimisation (8.27) peut se réécrire comme suit :

$$(\forall \mathbf{x} \in \mathbb{R}^N) \quad \text{prox}_{\mathbf{U}, \theta_\infty}(\mathbf{x}) = \underset{\mathbf{y} \in \mathbb{R}^N, \xi \in \mathbb{R}}{\operatorname{argmin}} \quad \xi + \frac{1}{2} \|\mathbf{x} - \mathbf{y}\|_{\mathbf{U}}^2$$

$$\text{tel que } (\forall n \in \{1, \dots, N\}) \quad (\mathbf{y}^{(n)}, \xi) \in \operatorname{epi} \theta. \quad (8.28)$$

Le problème sous contraintes épigraphiques (8.28) peut être réécrit en faisant intervenir la fonction indicatrice donnée dans la définition 2.3(viii) :

$$\text{prox}_{\mathbf{U}, \theta_\infty}(\mathbf{x}) = \underset{(\mathbf{y}, \xi) \in \mathbb{R}^{N+1}, (\mathbf{s}, \boldsymbol{\eta}) \in \mathbb{R}^{2N}}{\operatorname{argmin}} \quad \xi + \frac{1}{2} \|\mathbf{x} - \mathbf{y}\|_{\mathbf{U}}^2 + \sum_{n=1}^N \iota_{\operatorname{epi} \theta}(\mathbf{s}^{(n)}, \eta^{(n)})$$

$$\text{tel que } \mathbf{s} = \mathbf{y} \text{ et } \xi \mathbf{1}_N = \boldsymbol{\eta}. \quad (8.29)$$

Le problème (8.29) peut être résolu, par exemple, en utilisant ADMM, décrit dans le chapitre 2 par l'algorithme 2.9. Les itérations de cette méthode appliquées au problème (8.29) sont données dans l'algorithme 8.1.

Algorithme 8.1 ADMM pour le calcul de $\text{prox}_{\gamma \theta_\infty}(\mathbf{x})$, où $\mathbf{x} \in \mathbb{R}^N$

Initialisation : Soient $\mu > 0$, $\begin{pmatrix} \mathbf{s}_0 \\ \boldsymbol{\eta}_0 \end{pmatrix} \in \mathbb{R}^{2N}$ et $\begin{pmatrix} \mathbf{p}_{\mathbf{y},0} \\ \mathbf{p}_{\xi,0} \end{pmatrix} \in \mathbb{R}^{2N}$.

Itérations :

Pour $k = 0, 1, \dots$

$$\left[\begin{array}{l} \begin{pmatrix} \mathbf{y}_{k+1} \\ \xi_{k+1} \end{pmatrix} = \underset{(\mathbf{y}, \xi) \in \mathbb{R}^{N+1}}{\operatorname{argmin}} \quad \xi + \frac{1}{2} \|\mathbf{x} - \mathbf{y}\|_{\mathbf{U}}^2 + \frac{\mu}{2} \left(\|\mathbf{y} - \mathbf{s}_k + \mathbf{p}_{\mathbf{y},k}\|^2 + \|\xi \mathbf{1}_N - \boldsymbol{\eta}_k + \mathbf{p}_{\xi,k}\|^2 \right), \\ \text{pour } n = 1, \dots, N \\ \left[\begin{array}{l} \begin{pmatrix} \mathbf{s}_{k+1}^{(n)} \\ \eta_{k+1}^{(n)} \end{pmatrix} = \Pi_{\operatorname{epi} \theta} \left(\begin{pmatrix} \mathbf{y}_{k+1}^{(n)} + \mathbf{p}_{\mathbf{y},k+1}^{(n)} \\ \xi_{k+1}^{(n)} + \mathbf{p}_{\xi,k+1}^{(n)} \end{pmatrix} \right), \\ \begin{pmatrix} \mathbf{p}_{\mathbf{y},k+1} \\ \mathbf{p}_{\xi,k+1} \end{pmatrix} = \begin{pmatrix} \mathbf{p}_{\mathbf{y},k} \\ \mathbf{p}_{\xi,k} \end{pmatrix} + \begin{pmatrix} \mathbf{y}_{k+1} \\ \xi_{k+1} \mathbf{1}_N \end{pmatrix} - \begin{pmatrix} \mathbf{s}_{k+1} \\ \boldsymbol{\eta}_{k+1} \end{pmatrix}. \end{array} \right. \end{array} \right.$$

Notons que dans la première étape de l'algorithme 8.1, on peut effectuer la minimisation séparément par rapport à la variable $\mathbf{y}$ et par rapport à la variable ξ :

$$(\forall k \in \mathbb{N}) \quad \begin{cases} \mathbf{y}_{k+1} = \operatorname{argmin}_{\mathbf{y} \in \mathbb{R}^N} \frac{1}{2} \|\mathbf{x} - \mathbf{y}\|_{\mathbf{U}}^2 + \frac{\mu}{2} \|\mathbf{y} - \mathbf{s}_k + \mathbf{p}_{\mathbf{y},k}\|^2 \\ \xi_{k+1} = \operatorname{argmin}_{\xi \in \mathbb{R}} \xi + \frac{\mu}{2} \|\xi \mathbf{1}_N - \boldsymbol{\eta}_k + \mathbf{p}_{\xi,k}\|^2. \end{cases} \quad (8.30)$$

On veut donc trouver, pour tout $k \in \mathbb{N}$, $(\mathbf{y}_{k+1}, \xi_{k+1})$ tels que

$$\begin{cases} \mathbf{0}_N = \mathbf{U}(\mathbf{y}_{k+1} - \mathbf{x}) + \mu(\mathbf{y}_{k+1} - \mathbf{s}_k + \mathbf{p}_{\mathbf{y},k}) \\ 0 = 1 + N\mu\xi_{k+1} - \mu \sum_{n=1}^N (\eta_k^{(n)} - \mathbf{p}_{\xi,k}^{(n)}), \end{cases} \quad (8.31)$$

où $\mathbf{0}_N$ est le vecteur nul de $\mathbb{R}^N$. On obtient alors

$$\begin{cases} \mathbf{y}_{k+1} = (\mathbf{I}_N + \mu\mathbf{U}^{-1})^{-1} (\mathbf{x} + \gamma\mu(\mathbf{s}_k - \mathbf{p}_{\mathbf{y},k})) \\ \xi_{k+1} = \frac{1}{N} \left(\sum_{n=1}^N (\eta_k^{(n)} - \mathbf{p}_{\xi,k}^{(n)}) - \frac{1}{\mu} \right). \end{cases} \quad (8.32)$$

Rappelons que $\mathbf{U} = \operatorname{Diag}(\mathbf{u})$ avec $\mathbf{u} = (u^{(n)})_{1 \leq n \leq N} \in]0, +\infty[^N$. Ainsi, pour tout $n \in \{1, \dots, N\}$, on a

$$\mathbf{y}_{k+1}^{(n)} = \frac{1}{1 + \mu(u^{(n)})^{-1}} \left(\mathbf{x}^{(n)} + \mu(u^{(n)})^{-1} (\mathbf{s}_k^{(n)} - \mathbf{p}_{\mathbf{y},k}^{(n)}) \right). \quad (8.33)$$

D'autre part, pour tout $k \in \mathbb{N}$, et pour chacune des composantes $n \in \{1, \dots, N\}$, nous devons projeter $(\mathbf{y}_{k+1}^{(n)} + \mathbf{p}_{\mathbf{y},k}^{(n)}, \xi_{k+1} + \mathbf{p}_{\xi,k}^{(n)})$ dans l'ensemble $\operatorname{epi} \theta$. Cette projection est donnée par la proposition suivante :

Proposition 8.2. Soit $\mathbf{b} > 0$. La projection $\Pi_{\operatorname{epi} \mathbf{q}(\cdot, \mathbf{b})}(\mathbf{u}, \zeta)$ de $(\mathbf{u}, \zeta) \in \mathbb{R}^2$ dans l'épigraphe de $\mathbf{q}(\cdot, \mathbf{b})$ est donné par

$$\Pi_{\operatorname{epi} \mathbf{q}(\cdot, \mathbf{b})}(\mathbf{u}, \zeta) := \begin{cases} (\mathbf{u}, \zeta) & \text{si } \mathbf{u} > 0 \text{ et } \max\{\frac{\mathbf{u}}{\mathbf{b}}, \frac{\mathbf{b}}{\mathbf{u}}\} \leq \zeta, \\ \frac{\mathbf{b}\mathbf{u} + \zeta}{1 + \mathbf{b}^2}(\mathbf{b}, 1) & \text{si } 1 + \mathbf{b}^2 - \mathbf{b}\mathbf{u} < \zeta < \frac{\mathbf{u}}{\mathbf{b}}, \\ (\mathbf{b}, 1) & \text{si } \zeta \leq \min\{1 + \mathbf{b}^2 - \mathbf{b}\mathbf{u}, 1 - \mathbf{b}^2 + \mathbf{b}\mathbf{u}\}, \\ (\mathbf{t}^*, \frac{\mathbf{b}}{\mathbf{t}^*}) & \text{si } 1 - \mathbf{b}^2 + \mathbf{b}\mathbf{u} < \zeta \text{ et } \zeta < \frac{\mathbf{b}}{\mathbf{u}} \text{ pour } \mathbf{u} > 0, \end{cases} \quad (8.34)$$

où $\mathbf{t}^*$ est l'unique racine dans $]0, \mathbf{b}[$ du polynôme du quatrième ordre

$$\mathbf{r}(\mathbf{t}) = \mathbf{t}^4 - \mathbf{u}\mathbf{t}^3 + \zeta\mathbf{b}\mathbf{t} - \mathbf{b}^2. \quad (8.35)$$

Démonstration. Soit $\mathbf{b} > 0$. Afin de définir la projection sur l'épigraphe de $\mathbf{q}(\cdot, \mathbf{b})$, nous devons trouver les fonctions perpendiculaires² aux fonctions

$$\begin{cases} \mathbf{u} \in]0, \mathbf{b}] & \mapsto \mathbf{q}_1(\mathbf{u}, \mathbf{b}) := \mathbf{b}/\mathbf{u} \\ \mathbf{u} \in [\mathbf{b}, +\infty[& \mapsto \mathbf{q}_2(\mathbf{u}, \mathbf{b}) := \mathbf{u}/\mathbf{b} \end{cases} \quad (8.36)$$

au point $(\mathbf{b}, 1)$.

2. Une fonction π est dite perpendiculaire à une fonction $\mathbf{q}: E \rightarrow \mathbb{R}$, pour $E \subset \mathbb{R}$, en un point $(\mathbf{b}, \zeta) \in E \times \mathbb{R}$ si les vecteurs directeurs de leurs tangentes au point $(\mathbf{b}, \zeta)$ sont orthogonaux.

- (i) Soit $u \in]0, b]$. La fonction perpendiculaire à $q_1(u, b)$ en $(b, 1)$, que nous noterons $\pi_1(u) := \alpha_1 u + \beta_1$, où $(\alpha_1, \beta_1) \in \mathbb{R}^2$, satisfait

$$\begin{cases} \langle (u, \zeta) - (t, \varrho), (1, -b/t^2) \rangle = 0, \\ (t, \varrho) = (b, 1), \\ \zeta = \pi_1(u), \end{cases}$$

ce qui implique

$$\langle (u, \zeta) - (b, 1), (1, -1/b) \rangle = 0,$$

soit, de façon équivalente,

$$bu - b^2 + 1 = \zeta = \pi_1(u).$$

Ainsi, la perpendiculaire π_1 est donnée par

$$\pi_1(u) = bu + 1 - b^2.$$

- (ii) Soit $u \in [b, +\infty[$. La fonction perpendiculaire à $q_2(u, b)$ en $(b, 1)$, notée $\pi_2(u)$ satisfait

$$\begin{cases} \pi_2(u) = -bu + \beta_2, \\ \pi_2(b) = 1, \end{cases}$$

où $\beta_2 \in \mathbb{R}$. Donc

$$\begin{cases} \pi_2(u) = -bu + \beta_2, \\ -b^2 + \beta_2 = 1, \end{cases}$$

et de façon équivalente

$$\pi_2(u) = -bu + 1 + b^2.$$

D'après (i) et (ii), les ensembles sont alors définis par

$$\begin{aligned} A_1 &= \{(u, \zeta) \mid u > 0 \text{ et } \zeta \leq \max\{u/b, b/u\}\}, \\ A_2 &= \{(u, \zeta) \mid 1 + b^2 - bu < \zeta < u/b\}, \\ A_3 &= \{(u, \zeta) \mid \zeta \leq \min\{1 - b^2 + bu, 1 + b^2 - bu\}\}, \\ A_3 &= \{(u, \zeta) \mid \zeta > 1 - b^2 + bu \text{ et } (\zeta < b/u \text{ si } u > 0)\}. \end{aligned}$$

Ces ensembles sont représentés sur la figure 8.2 dans le cas où $b = 1$.

- **Cas où** $(u, \zeta) \in A_1$. On remarque que $A_1 = \text{epi } q(\cdot, b)$. Ainsi,

$$(\forall (u, \zeta) \in A_1) \quad \Pi_{\text{epi } q(\cdot, b)}(u, \zeta) = (u, \zeta). \quad (8.37)$$

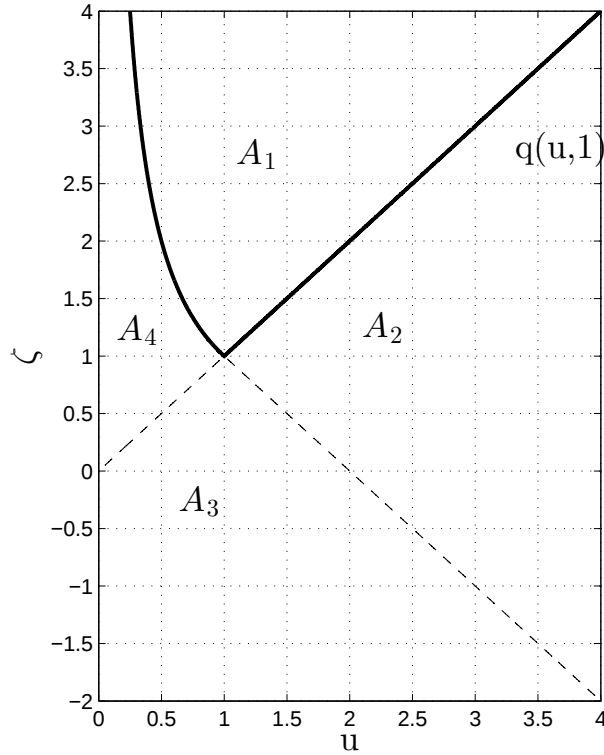


Figure 8.2 – Division de $\mathbb{R}^2$ en quatre sous-espaces pour calculer la projection épigraphique dans $\text{epi } q(\cdot, b)$ pour $b = 1$.

- **Cas où** $(u, \zeta) \in A_2$. La projection orthogonale $(t, \varrho) = \Pi_{\text{epi } q(\cdot, b)}(u, \zeta)$ satisfait

$$\begin{cases} \varrho = t/b, \\ \left\langle \begin{pmatrix} u \\ \zeta \end{pmatrix} - \begin{pmatrix} t \\ \varrho \end{pmatrix}, \begin{pmatrix} 1 \\ 1/b \end{pmatrix} \right\rangle = 0. \end{cases}$$

Donc, $t = \frac{b}{1+b^2}(bu + \zeta)$.

- **Cas où** $(u, \zeta) \in A_3$. On remarque que A_3 correspond au cône normal de $\text{epi } q(\cdot, b)$ en $(b, 1)$. Ainsi,

$$(\forall (u, \zeta) \in A_3) \quad \Pi_{\text{epi } q(\cdot, b)}(u, \zeta) = (b, 1). \quad (8.38)$$

- **Cas où** $(u, \zeta) \in A_4$. Dans ce cas, la projection orthogonale se fait sur $t \in]0, b[\mapsto \tau(t) := (t, b/t)$. La direction de la tangente de τ est $(1, -b/t^2)$, donc $(t, \varrho) = \Pi_{\text{epi } q(\cdot, b)}(u, \zeta)$ satisfait

$$\begin{cases} \varrho = t/b, \\ \left\langle \begin{pmatrix} u \\ \zeta \end{pmatrix} - \begin{pmatrix} t \\ \varrho \end{pmatrix}, \begin{pmatrix} 1 \\ 1/b \end{pmatrix} \right\rangle = 0. \end{cases}$$

Par conséquent, $\Pi_{\text{epi } \mathbf{q}(\cdot, \mathbf{b})}(\mathbf{u}, \zeta) = (\mathbf{t}, \mathbf{t}/\mathbf{b})$, où $\mathbf{t} \in]0, \mathbf{b}[$ est la solution unique de $r(\mathbf{t}) = 0$, r étant le polynôme de quatrième ordre défini par (8.35). L'unicité est garantie puisque $\mathbf{q}(\cdot, \mathbf{b})$ étant convexe, son épigraphe est un ensemble convexe fermé non vide (proposition 2.1). Par conséquent, la projection sur cet ensemble est unique (proposition 2.9).

□

8.4 APPLICATION À L'ESTIMATION DE LA SÉLECTIVITÉ

Le but de cette section est de résoudre le problème de minimisation suivant

$$\text{trouver } \hat{\mathbf{x}} = \underset{\mathbf{x} \in \mathcal{C}}{\text{argmin}} \quad \mathbf{q}_\nu(\mathbf{A}\mathbf{x}, \mathbf{b}), \quad (8.39)$$

où $\nu \in \{1, \infty\}$, $\mathbf{A} \in \mathbb{R}^{M \times N}$, $\mathbf{b} \in \mathbb{R}^M$, et

$$\mathcal{C} = \{\mathbf{x} = (\mathbf{x}^{(n)})_{1 \leq n \leq N} \in [0, +\infty[^N \mid \sum_{n=1}^N \mathbf{x}^{(n)} \leq 1\}. \quad (8.40)$$

Ce type de problème apparaît, notamment, dans l'estimation des sélectivités pour l'optimisation des requêtes basées sur les coûts en SGBD (c.f. [Markl *et al.*, 2007; Moerkotte *et al.*, 2009]). Une brève description de cette application est donnée dans le paragraphe suivant.

8.4.1 Estimation de la sélectivité en SGBD

La sélectivité indique la proportion de n -uplets dans une base de donnée satisfaisant les prédicats dans une requête. La précision de l'estimation des sélectivités est cruciale pour la conception de plans d'exécution des requêtes optimales. Cependant, en pratique, cette estimation est faite de façon inexacte, introduisant des erreurs dans le modèle. Une question se pose alors : Comment ces erreurs influencent l'optimisation des plans de requêtes ? Dans [Moerkotte *et al.*, 2009], l'erreur de propagation de mauvaises estimations de sélectivités a été examinée au moyen d'une requête optimisée appropriée. Il est apparu que l'erreur de propagation est multiplicative (voir aussi [Ioannidis et Christodoulakis, 1991]). Il a de plus été démontré que les pires bornes d'erreurs pouvant être atteintes par la fonction de coût d'un plan optimal basé sur des sélectivités erronées, sont liées à la fois à la fonction de coût d'un plan optimal basé sur des sélectivités exactes, et au quotient des sélectivités erronées et exactes. Par conséquent, il semble que la fonction quotient d'erreur de sélectivités donnera de meilleures estimations que d'autres mesures d'erreurs (en particulier, on aura de meilleurs résultats que pour les erreurs additives).

8.4.2 Formulation du problème

Soit N le nombre de prédicats possibles d'une requête, $\mathcal{P}_N = 2^{\mathcal{I}_N}$ l'ensemble des parties de $\mathcal{I}_N = \{1, \dots, N\}$. On suppose que l'on connaît une famille $\{\mathbf{p}_n \mid n \in \mathcal{I}_N\}$ de prédicats simples.

De la même manière que dans [Markl et al., 2007], nous allons modéliser les sélectivités de prédicats par une distribution de probabilités. Nous allons donc les représenter sous *Forme Normale Disjonctive* (FND), dont les principaux symboles sont résumés dans la Table 8.1.

Symbole logique	définition
$\wedge$	et
$\vee$	ou
$\neg$	non

Table 8.1 – Grammaire pour la FND

Exemple 8.1. Soit $N = 3$, les prédicats $\mathbf{p}_1$ et $\mathbf{p}_1 \wedge \mathbf{p}_2$ ont pour FND :

$$\begin{aligned}\mathbf{p}_1 &= (\mathbf{p}_1 \wedge \neg \mathbf{p}_2 \wedge \neg \mathbf{p}_3) \vee (\mathbf{p}_1 \wedge \mathbf{p}_2 \wedge \neg \mathbf{p}_3) \vee (\mathbf{p}_1 \wedge \neg \mathbf{p}_2 \wedge \mathbf{p}_3) \vee (\mathbf{p}_1 \wedge \mathbf{p}_2 \wedge \mathbf{p}_3), \\ \mathbf{p}_1 \wedge \mathbf{p}_2 &= (\mathbf{p}_1 \wedge \mathbf{p}_2 \wedge \neg \mathbf{p}_3) \vee (\mathbf{p}_1 \wedge \mathbf{p}_2 \wedge \mathbf{p}_3).\end{aligned}$$

Pour la FND du prédicat $\mathbf{p}_1 \wedge \mathbf{p}_2$, nous rappelons que :

- $\mathbf{p}_1 \wedge \mathbf{p}_2$ est appelée conjonction de prédicats (ici des prédicats $\mathbf{p}_1$ et $\mathbf{p}_2$).
- $\mathbf{v} = (\mathbf{p}_1 \wedge \mathbf{p}_2 \wedge \neg \mathbf{p}_3)$ est une proposition de la FND du prédicat $\mathbf{p}_1 \wedge \mathbf{p}_2$.

Nous utilisons les notations suivantes :

- Nous représentons chaque ensemble de $\mathcal{P}_N$ par une chaîne binaire $\beta \in \{0, 1\}^N$, où, pour tout $n \in \mathcal{I}_N$, $\beta^{(n)} = 1$ si et seulement si n est dans l'ensemble considéré. Alors, $\beta = \mathbf{0}_N$ correspond à l'ensemble vide, tandis que $\beta = \mathbf{1}_N$ correspond à l'ensemble $\mathcal{I}_N$ complet. On définit de plus $\mathcal{P}_N^* = \{0, 1\}^N \setminus \{\mathbf{0}_N\}$.
- Nous indexons les sélectivités $\mathbf{s}_\beta$ de prédicats par les étiquettes binaires $\beta \in \{0, 1\}^N$, où, pour tout $n \in \mathcal{I}_N$, $\beta^{(n)} = 1$ si et seulement si le prédicat $\mathbf{p}_n$ est dans la conjonction.
- Les sélectivités $\mathbf{x}_\beta$ des propositions $\mathbf{v}$ de la FND sont indexées par des chaînes binaires $\beta \in \{0, 1\}^N$, où, pour tout $n \in \mathcal{I}_N$, $\beta^{(n)} = 1$ si et seulement si $\mathbf{p}_n \in \mathbf{v}$ et $\beta^{(n)} = 0$ si $\neg \mathbf{p}_n \in \mathbf{v}$.

Exemple 8.2. Les FND de l'exemple 8.1 sont représentées de la façon suivante :

$$\begin{aligned}\mathbf{s}_{100} &= \mathbf{x}_{100} + \mathbf{x}_{110} + \mathbf{x}_{101} + \mathbf{x}_{111}, \\ \mathbf{s}_{110} &= \mathbf{x}_{110} + \mathbf{x}_{111}.\end{aligned}$$

Les prédicats étant des probabilités, nous avons alors

$$\sum_{\beta \in \{0, 1\}^N} \mathbf{x}_\beta = 1.$$

Les données $\mathbf{x}_\beta$ peuvent être interprétées comme les probabilités d'apparition des propositions correspondantes. Notons que $\mathbf{x}_{\mathbf{0}_N}$ n'est un terme de la somme de $\mathbf{s}_\beta$ que lorsque $\beta = \mathbf{0}_N$.

Remarquons que seule une petite partie des sélectivités $\mathbf{s}_\beta$, où $\beta \in \mathcal{J} \subset \mathcal{P}_N^*$, tel que $M = \#\mathcal{J} \ll 2^N$, des conjonctions de prédicats peut être mémorisée comme statistiques multivariées en SGBD.

Définissons alors $\mathbf{b} = (\mathbf{s}_\beta)_{\beta \in \mathcal{J}}$, $\mathbf{x} = (\mathbf{x}_\beta)_{\beta \in \mathcal{P}_N^*}$ et $\mathbf{A} = (\mathbf{a}_{\beta, \beta'})_{\beta \in \mathcal{J}, \beta' \in \mathcal{P}_N^*} \in \{0, 1\}^{M \times N}$, où

$$(\forall (\beta \in \mathcal{J}, \beta' \in \mathcal{P}_N^*)) \quad \mathbf{a}_{\beta, \beta'} = \begin{cases} 1 & \text{si } (\forall i \in \mathcal{I}_n) \quad \beta_i = 1 \Rightarrow \beta'_i = 1, \\ 0 & \text{sinon.} \end{cases} \quad (8.41)$$

Si, pour tout $\beta \in \mathcal{J}$, $\mathbf{s}_\beta$ était connu de façon précise, les $\mathbf{x}_\beta$ satisferaient alors $\mathbf{Ax} = \mathbf{b}$. Dans [Markl et al., 2007], les auteurs proposent d'estimer $(\mathbf{x}_\beta)_{\beta \in \mathcal{P}_N^*}$ (et par conséquent, toutes les sélectivités) en maximisant l'entropie

$$\max_{\mathbf{x} \in \mathbb{R}^N} \sum_{\beta \in \{0,1\}^N} -\mathbf{x}_\beta \log \mathbf{x}_\beta \quad \text{tel que} \quad \begin{cases} \mathbf{Ax} = \mathbf{b}, \\ \mathbf{x} \geq \mathbf{0}, \\ \sum_{\beta \in \{0,1\}^N} \mathbf{x}_\beta = 1. \end{cases} \quad (8.42)$$

Si ce problème d'optimisation convexe est réalisable, i.e. admet une solution, il est alors possible de le résoudre en utilisant diverses méthodes. Par exemple, on peut utiliser la méthode de Newton en l'appliquant au problème dual (c.f. [Fletcher, 2000, p. 222-223]), ou bien une méthode itérative avec changement d'échelle comme il l'a été proposé dans [Markl et al., 2007].

Cependant, en pratique, les inexactitudes présentes dans les sélectivités $\mathbf{s}_\beta$ mémorisées rendent le problème (8.42) non réalisable, au sens où $\mathbf{Ax} = \mathbf{b}$ n'admet pas de solution $\mathbf{x} \in \mathbb{R}^N$ vérifiant $\mathbf{x} \geq \mathbf{0}$. Une solution pour pallier ce problème pourrait être d'ajouter un nouveau terme de régularisation au terme d'entropie à maximiser, pénalisant l'erreur entre $\mathbf{Ax}$ et $\mathbf{b}$. Néanmoins, cela introduirait un paramètre de régularisation qu'il faudrait optimiser par la suite. Par conséquent, nous proposons de traiter ce problème en deux étapes. Dans un premier temps, pour les raisons que nous avons données dans la section 8.4.1, nous allons minimiser le quotient de $\mathbf{Ax}$ par $\mathbf{b}$, i.e., nous allons résoudre le problème (8.39) de façon à rendre le problème réalisable. Nous pourrions alors résoudre le problème (8.42) dans un deuxième temps en remplaçant la variable $\mathbf{b}$ apparaissant dans ce problème par $\mathbf{Ax}$, où $\hat{\mathbf{x}}$ est la solution du problème (8.39), et en utilisant les méthodes précédemment citées. Dans ce chapitre, notre étude se focalisera sur la première étape.

8.4.3 Solution du problème en SGBD

Nous allons maintenant nous intéresser à la résolution du problème (8.39). Une approche, déjà présentée dans [Setzer et al., 2010], est d'utiliser une méthode SOCP (*Second Order Cone Programming* [Lobo et al., 1998]). Cependant, dans cette section nous proposons de résoudre ce problème à l'aide d'un algorithme primal-dual. Ce type d'algorithme présente l'avantage de pouvoir effectuer certaines étapes, telles que le calcul de l'opérateur de seuillage $\text{prox}_{\mathbf{q}(\cdot, \mathbf{b})}$ ou des projections épigraphiques, en parallèle à chaque itération.

8.4.3.1 Algorithmes proposés

Pour $\nu \in \{1, +\infty\}$, nous reformulons le problème (8.39) comme suit :

$$\text{trouver } \hat{\mathbf{x}} = \underset{\mathbf{x} \in \mathbb{R}^N, \mathbf{y} \in \mathbb{R}^M}{\text{argmin}} \quad \mathbf{q}_\nu(\mathbf{y}, \mathbf{b}) + \iota_{\mathcal{C}}(\mathbf{x}) \quad \text{tel que } \mathbf{Ax} = \mathbf{y}. \quad (8.43)$$

- Lorsque $\nu = 1$, le problème d'optimisation (8.43) peut être résolu par divers algorithmes praux-duaux (voir section 2.4.4).

Dans notre approche, nous proposons d'utiliser l'algorithme 2.14 donné dans la section 2.4.4, qui se réécrit sous la forme de l'algorithme 8.2 pour résoudre le problème

$$\text{trouver } \hat{\mathbf{x}} = \underset{\mathbf{x} \in \mathbb{R}^N, \mathbf{y} \in \mathbb{R}^M}{\operatorname{argmin}} \quad \mathbf{q}_1(\mathbf{y}, \mathbf{b}) + \iota_{\mathcal{C}}(\mathbf{x}) \text{ tel que } \mathbf{Ax} = \mathbf{y}. \quad (8.44)$$

Algorithme 8.2 Algorithme primal-dual pour résoudre le problème (8.44)

Initialisation : Soient $\mathbf{x}_0 \in \mathbb{R}^N$, $\mathbf{p}_0 = \mathbf{z}_0 \in \mathbb{R}^M$, $(\mu, \sigma) \in]0, +\infty[^2$, et $\lambda \in]0, 1]$.

Itérations :

Pour $k = 0, 1, \dots$

$$\left\{ \begin{array}{l} \mathbf{x}_{k+1} = \Pi_{\mathcal{C}}(\mathbf{x}_k - \mu\sigma \mathbf{A}^\top \mathbf{z}_k) \\ \mathbf{y}_{k+1} = \operatorname{prox}_{\mathbf{q}_1(\cdot, \mathbf{b})/\sigma}(\mathbf{p}_k + \mathbf{Ax}_{k+1}) \\ \mathbf{p}_{k+1} = \mathbf{p}_k + \mathbf{Ax}_{k+1} - \mathbf{y}_{k+1} \\ \mathbf{z}_{k+1} = \mathbf{p}_{k+1} + \lambda(\mathbf{p}_{k+1} - \mathbf{p}_k) \end{array} \right.$$

Les propriétés de convergence de cet algorithme ont été énoncées dans le corollaire 2.1. L'étape de mise à jour des $(\mathbf{y}_k)_{k \in \mathbb{N}}$ nécessite le calcul de l'opérateur proximal de la fonction $\mathbf{q}_1(\cdot, \mathbf{b})$. Ce dernier peut-être effectué en calculant terme à terme, pour tout $m \in \{1, \dots, M\}$, le seuillage de $\mathbf{q}(\cdot, \mathbf{b}^{(m)})$ décrit dans le corollaire 8.1.

- Lorsque $\nu = \infty$, pour $\mathbf{b} = (\mathbf{b}^{(m)})_{1 \leq m \leq M}$, le problème (8.43) se réécrit comme suit

$$\begin{aligned} \text{trouver } (\hat{\mathbf{x}}, \hat{\xi}, \hat{\mathbf{y}}, \hat{\boldsymbol{\eta}}) \in \underset{(\mathbf{x}, \xi) \in \mathbb{R}^{N+1}, (\mathbf{y}, \boldsymbol{\eta}) \in \mathbb{R}^{2M}}{\operatorname{Argmin}} \quad & \xi + \sum_{m=1}^M \iota_{\operatorname{epi} \mathbf{q}(\cdot, \mathbf{b}^{(m)})}(\mathbf{y}^{(m)}, \boldsymbol{\eta}^{(m)}) + \iota_{\mathcal{C}}(\mathbf{x}) \\ \text{tel que } & \begin{cases} \mathbf{Ax} = \mathbf{y}, \\ \xi \mathbf{1}_M = \boldsymbol{\eta}. \end{cases} \end{aligned} \quad (8.45)$$

Pour résoudre ce problème, on peut de nouveau utiliser l'algorithme 2.14. On obtient alors les itérations données dans l'algorithme 8.3, où

$$(\forall (\mathbf{x}, \xi) \in \mathbb{R}^{N+1}) \quad f(\mathbf{x}, \xi) = \xi + \iota_{\mathcal{C}}(\mathbf{x}).$$

Algorithme 8.3 Algorithme primal-dual pour résoudre le problème (8.45)

Initialisation : Soient $\mathbf{x}_0 \in \mathbb{R}^N$, $\xi_0 \in \mathbb{R}$, $\mathbf{p}_{\mathbf{x},0} = \mathbf{z}_{\mathbf{x},0} \in \mathbb{R}^M$, $\mathbf{p}_{\xi,0} = \mathbf{z}_{\xi,0} \in \mathbb{R}^M$, $(\mu, \sigma) \in]0, +\infty[^2$, et $\lambda \in]0, 1]$.

Itérations :

Pour $k = 0, 1, \dots$

$$\left[\begin{array}{l} \begin{pmatrix} \mathbf{x}_{k+1} \\ \xi_{k+1} \end{pmatrix} = \text{prox}_{\mu f} \left(\begin{pmatrix} \mathbf{x}_k \\ \xi_k \end{pmatrix} - \mu \sigma \begin{pmatrix} \mathbf{A}^\top \mathbf{z}_{\mathbf{x},k} \\ \mathbf{1}_M^\top \mathbf{z}_{\xi,k} \end{pmatrix} \right) \\ \text{pour } m = 1, \dots, M \\ \left[\begin{pmatrix} \mathbf{y}_{k+1}^{(m)} \\ \eta_{k+1}^{(m)} \end{pmatrix} = \Pi_{\text{epi } \mathbf{q}(\cdot, \mathbf{b}^{(m)})/\sigma} \left(\begin{pmatrix} \mathbf{p}_{\mathbf{x},k}^{(m)} + [\mathbf{A}\mathbf{x}_{k+1}]^{(m)} \\ \mathbf{p}_{\xi,k}^{(m)} + \xi_{k+1} \end{pmatrix} \right) \\ \begin{pmatrix} \mathbf{p}_{\mathbf{x},k+1} \\ \mathbf{p}_{\xi,k+1} \end{pmatrix} = \begin{pmatrix} \mathbf{p}_{\mathbf{x},k} + \mathbf{A}\mathbf{x}_{k+1} - \mathbf{y}_{k+1} \\ \mathbf{p}_{\xi,k} + \mathbf{1}_M \xi_{k+1} - \eta_{k+1} \end{pmatrix} \\ \begin{pmatrix} \mathbf{z}_{\mathbf{x},k+1} \\ \mathbf{z}_{\xi,k+1} \end{pmatrix} = \begin{pmatrix} \mathbf{p}_{\mathbf{x},k} + \lambda (\mathbf{p}_{\mathbf{x},k} - \mathbf{p}_{\mathbf{x},k+1}) \\ \mathbf{p}_{\xi,k} + \lambda (\mathbf{p}_{\xi,k} - \mathbf{p}_{\xi,k+1}) \end{pmatrix} \end{array} \right.$$

Soit $k \in \mathbb{N}$ une itération de l'algorithme 8.3. Dans la première étape, $\mathbf{x}_{k+1}$ est donné par la projection de $\mathbf{x}_k - \mu \sigma \mathbf{A}^\top \mathbf{z}_{\mathbf{x},k}$ sur le simplexe $\mathcal{C}$, et ξ_{k+1} est donné par

$$\xi_{k+1} = \xi_k - \mu \sigma \mathbf{1}_M^\top \mathbf{z}_{\xi,k} - \mu. \quad (8.46)$$

La deuxième étape requiert la projection épigraphique calculée terme à terme du n -uplet $([\mathbf{p}_{\mathbf{x},k} + \mathbf{A}\mathbf{x}_{k+1}]^{(m)}, [\mathbf{p}_{\xi,k} + \mathbf{1}_M \xi_{k+1}]^{(m)})$ sur l'épigraphe de $\mathbf{q}(\cdot, \mathbf{b}^{(m)})$, pour tout $m \in \{1, \dots, M\}$. Cette dernière peut être calculée en exploitant le corollaire 8.2.

8.4.3.2 Exemple numérique

Nous allons maintenant présenter un exemple numérique afin d'illustrer notre méthode. Considérons le système d'équations suivant :

$$\begin{pmatrix} 1 & 0 & 1 & 0 & 1 & 0 & 1 \\ 0 & 1 & 1 & 0 & 0 & 1 & 1 \\ 0 & 0 & 0 & 1 & 1 & 1 & 1 \\ 0 & 0 & 1 & 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 & 1 & 0 & 1 \\ 0 & 0 & 0 & 0 & 0 & 1 & 1 \end{pmatrix} \begin{pmatrix} x_{001} \\ x_{010} \\ x_{011} \\ x_{100} \\ x_{101} \\ x_{110} \\ x_{111} \end{pmatrix} = \begin{pmatrix} 0.2114 \\ 0.6331 \\ 0.6312 \\ 0.5182 \\ 0.9337 \\ 0.0035 \end{pmatrix} = \begin{pmatrix} s_{001} \\ s_{010} \\ s_{100} \\ s_{011} \\ s_{101} \\ s_{110} \end{pmatrix} \quad (8.47)$$

n'admettant pas de solution $\mathbf{x}$ positive. On résout le problème (8.39) pour $\nu = 1$ en utilisant l'algorithme 8.2, et pour $\nu = \infty$ avec l'algorithme 8.3. Nous utilisons comme critère d'arrêt $\|\mathbf{A}\mathbf{x}_k - \mathbf{y}_k\|_\infty < \epsilon$, où la matrice $\mathbf{A} \in \mathbb{R}^{M \times N}$ est donnée par (8.47) et $\epsilon = 10^{-5}$ (resp. $\epsilon = 10^{-4}$) pour l'algorithme 8.2 (resp. pour l'algorithme 8.3).

Remarquons que dans l'algorithme 8.2, le calcul de l'opérateur proximal de la fonction $\mathbf{q}_1(\cdot, \mathbf{b})$ (resp. dans l'algorithme 8.3, le calcul de la projection épigraphique) nécessite de trouver le zéro d'un polynôme du troisième degré (resp. du quatrième degré). Bien qu'il soit possible de trouver analytiquement ces zéros, dans nos simulations nous utilisons la méthode de Newton qui semble être plus rapide en pratique (comme nous l'avons vu dans la section 2.4.4, l'algorithme primal-dual utilisé est robuste aux erreurs numériques pouvant être introduites par cette méthode de calcul).

Nous comparerons nos solutions à celles obtenues en résolvant le problème suivant, faisant intervenir une norme d'une fonction additive :

$$\text{trouver } \hat{\mathbf{x}} \in \underset{\mathbf{x} \in \mathcal{C}}{\text{Argmin}} \quad \|\mathbf{Ax} - \mathbf{b}\|_\alpha, \quad (8.48)$$

où $\alpha \in \{1, 2\}$ en utilisant, respectivement, des méthodes de programmation linéaire et de programmation quadratique de l'outil MOSEK [MOS].

Pour chaque solution $\hat{\mathbf{x}}$ obtenue avec chacun des quatre algorithmes, nous calculons ensuite $\hat{\mathbf{b}} = \mathbf{A}\hat{\mathbf{x}}$. Notons que le vecteur $\hat{\mathbf{x}}$ n'est ici d'aucun intérêt puisque, comme nous l'avons expliqué à la fin de la section 8.4.2, il sera estimé au cours d'une seconde étape (par exemple, dans une application en base de données, par minimisation entropique), non traitée dans ce chapitre. L'erreur $\mathbf{q}_\infty(\hat{\mathbf{b}}, \mathbf{b})$ pour les quatre vecteurs résultants $\hat{\mathbf{b}}$ est donnée dans la table 8.2.

méthode	(8.39) pour $\nu = \infty$	(8.39) pour $\nu = 1$	(8.48) pour $\alpha = 1$	(8.48) pour $\alpha = 2$
$\mathbf{q}_\infty(\hat{\mathbf{b}}, \mathbf{b})$	2.61	3.65	75.65	74.85

Table 8.2 – Erreur $\mathbf{q}_\infty$ obtenue pour l'exemple (8.47).

Ces résultats reflètent de manière qualitative ce qui a pu être observé sur divers autres exemples. Pour des raisons évidentes, la méthode (8.39) pour $\nu = \infty$ donne la valeur minimale de l'erreur $\mathbf{q}_\infty(\hat{\mathbf{b}}, \mathbf{b})$, puisque la méthode a été conçue dans le but de minimiser cette erreur. Le modèle (8.39) avec $\nu = 1$ donne une erreur $\mathbf{q}_\infty(\hat{\mathbf{b}}, \mathbf{b})$ plus élevée mais cependant très proche de l'erreur minimale. Puisque cette méthode, n'ayant pas recourt aux projections épigraphiques, est plus rapide que la première, elle peut apparaître comme un bon compromis. Pour finir, les méthodes (8.48) pour $\alpha \in \{1, 2\}$ conduisent toutes les deux à des erreurs $\mathbf{q}_\infty(\hat{\mathbf{b}}, \mathbf{b})$ beaucoup plus importantes.

Dans cet exemple, nous ne donnons que la valeur de l'erreur $\mathbf{q}_\infty$, puisque, comme nous l'avons expliqué précédemment, c'est celle qui nous intéresse dans cette application, de par son influence sur la conception des demandes de plans d'exécution.

8.5 CONCLUSION

Dans ce chapitre nous avons d'une part déterminé l'opérateur proximal de la somme de fonctions quotients de vecteurs positifs, et d'autre part donné une méthode pour calculer l'opérateur proximal du maximum de ces fonctions quotients. Ces opérateurs proximaux peuvent être utilisés pour trouver les solutions de problèmes divers. Ici, nous avons considéré un problème de faisabilité intervenant dans l'estimation des sélectivités pour l'optimisation de requêtes. Dans ce type

d'application, les fonctions quotient se révèlent d'un intérêt évident en tant que mesure d'erreur lorsqu'elles sont comparées avec des mesures d'erreur additives. Afin de résoudre le problème sous-jacent, nous avons proposé d'utiliser un algorithme primal-dual, permettant de traiter les opérateurs proximaux des fonctions quotients.

Dans le cadre de l'optimisation des requêtes dans les SGBD, il serait intéressant de mettre en œuvre une implémentation parallèle de certaines étapes de l'algorithme primal-dual ainsi que de l'algorithme permettant de résoudre le problème (8.42). De plus, une comparaison détaillée de plusieurs méthodes, en particulier en ce qui concerne les temps d'exécution, doit être faite. Ces deux points pourront être traités dans de futurs travaux. Nous souhaiterions par ailleurs appliquer les fonctions quotients dans des problèmes pouvant apparaître en traitement d'images, comme par exemple pour la correction d'illumination dans le modèle *Retinex*, dans lesquels une image est supposée être le résultat d'un produit terme à terme de l'illumination et de la réflectivité d'une scène (c.f. [Gonzalez et Woods, 2002]).

CHAPITRE 9

Conclusions et perspectives

9.1 BILAN

Dans cette thèse, nous avons mis en œuvre des méthodes d’optimisation permettant de résoudre des problèmes faisant intervenir de grands volumes de données. En particulier, nous nous sommes intéressés à des problèmes inverses intervenant en traitement du signal et des images. Nous avons vu que ces problèmes peuvent être résolus en passant par la minimisation d’un critère s’écrivant comme la somme de fonctions non nécessairement convexes, non nécessairement différentiables, et pouvant être composées avec des opérateurs linéaires non nécessairement inversibles.

Dans un premier temps, nous nous sommes focalisés sur la minimisation d’une somme de deux fonctions, dont l’une était différentiable et l’autre était convexe.

Une méthode qui permet de résoudre ce problème, pour des données de grande taille, est l’algorithme explicite-implicite, qui alterne, à chaque itération une étape de gradient sur la fonction différentiable et une étape proximale sur la deuxième fonction. Cependant, en pratique, cet algorithme peut converger très lentement. Cela est dû, notamment, au pas apparaissant dans l’algorithme qui dépend de l’inverse de la constante de Lipschitz du gradient du terme différentiable. Ainsi, lorsque cette constante est très grande, le pas devient très petit et freine la convergence.

Une méthode permettant d’accélérer la convergence de l’algorithme explicite-implicite consiste à remplacer, à chaque itération, le pas par une matrice de préconditionnement. L’algorithme résultant correspond alors à l’algorithme explicite-implicite à métrique variable, ainsi nommé car à chaque itération la métrique sous-jacente est adaptée à l’itérée en cours. Ce type d’algorithme a déjà été proposé dans la littérature pour minimiser une somme de fonctions convexes. Cependant nous avons généralisé cette approche au cadre non convexe, et nous avons, de plus, proposé une méthode, basée sur le principe de majoration-minimisation, permettant de choisir efficacement la métrique à chaque itération.

Nous avons démontré la convergence des itérées de l’algorithme proposé vers un point critique du critère grâce à l’utilisation de l’inégalité de Kurdyka-Łojasiewicz. De plus, nous avons montré que la convergence restait assurée lorsqu’une étape de relaxation convexe était incorporée dans la méthode.

Lorsque l'on est amenés à gérer un très grand nombre de données, il peut être intéressant, du point de vue de l'implémentation de l'algorithme, de ne traiter à chaque itération que des sous-parties de ces données, sous la forme de blocs. Les algorithmes adoptant une telle stratégie permettent de ne réaliser, à chaque itération, qu'un nombre réduit d'opérations correspondant à la dimension du bloc traité.

Dans une seconde partie de cette thèse, nous avons proposé de combiner cette stratégie alternée par bloc à l'algorithme explicite-implicite à métrique variable. Nous avons démontré la convergence des itérées de l'algorithme résultant, vers un point critique du critère à minimiser. Dans ce travail, nous nous sommes intéressés au cas où ce critère s'écrit comme la somme de deux fonctions dont l'une est différentiable non nécessairement convexe et l'autre n'est ni nécessairement différentiable, ni nécessairement convexe.

La stratégie de minimisation alternée proposée permet de mettre à jour les blocs selon une règle acyclique, autorisant de choisir les blocs suivant un ordre arbitraire et de mettre à jour certains d'entre eux plus souvent que d'autres. Notons que cette règle de mise à jour est beaucoup plus flexible que la règle cyclique qui est généralement adoptée lors de la mise en place de stratégies alternées par bloc.

Dans une troisième partie de cette thèse, nous nous sommes intéressés au cas où le critère à minimiser s'écrit comme la somme de plusieurs fonctions dont certaines ne sont pas nécessairement différentiables et peuvent être composées avec des opérateurs linéaires.

Ce type de critère ne peut pas être minimisé par un algorithme de type explicite-implicite sans faire appel à des sous-itérations. Cependant, lorsque toutes les fonctions intervenant dans le critère sont convexes, il existe des méthodes efficaces de type primal-dual permettant de résoudre ce problème.

Nous avons proposé de combiner ces algorithmes primaux-duaux à une stratégie alternée par bloc. Les méthodes que nous avons développées permettent de ne mettre à jour, à chaque itération, qu'un certain nombre de blocs, choisis de façon aléatoire suivant une loi arbitraire. L'un des avantages de cette règle de mise à jour aléatoire, par rapport à la règle acyclique, réside en la possibilité de mettre à jour plusieurs blocs en même temps en parallèle, à chaque itération.

Notons de plus que les algorithmes primaux-duaux alternés par bloc que nous avons proposés s'apparentent à des méthodes stochastiques permettant de prendre en compte des erreurs aléatoires.

Nous avons démontré la convergence presque sûre de ces algorithmes vers un couple de variables aléatoires solution du problème primal et du problème dual.

Enfin, les algorithmes primaux-duaux alternés aléatoirement par bloc nous ont permis de développer des algorithmes stochastiques distribués, dans lesquels les étapes de calculs sont réparties sur différents agents communiquant entre eux.

Les contributions méthodologiques de cette thèse ont été étayées par divers tests numériques sur des applications du traitement du signal et des images.

Dans un premier temps nous nous sommes intéressés à la restauration d'images dégradées par un bruit gaussien dépendant du signal, apparaissant dans certains systèmes d'imagerie. Nous avons étudié ce modèle de bruit à travers une interprétation MAP. Nous avons montré que le problème de minimisation qui en résulte n'est pas convexe, et nous avons appliqué l'algorithme explicite-implicite à métrique variable afin de le résoudre. Au travers de cet exemple nous avons pu montrer numériquement que l'introduction d'une métrique variable permet d'accélérer de manière significative la convergence de l'algorithme explicite-implicite.

Nous avons ensuite appliqué la version alternée par bloc de l'algorithme explicite-implicite pour traiter trois exemples différents en traitement du signal.

Dans le premier exemple, nous nous sommes intéressés à un problème de démixage spectral, s'apparentant à un problème de NMF (*nonnegative matrix factorization* ou factorisation de matrices positives), où le but était d'estimer simultanément deux ensembles de données : les spectres apparaissant dans l'image hyperspectrale d'origine ainsi que leur abondance. Nous avons écrit ce problème sous la forme d'une minimisation d'un critère non différentiable et non convexe. Nous avons montré au travers d'un exemple numérique que notre méthode permettait de résoudre efficacement ce problème, et que la métrique variable permettait d'accélérer considérablement la vitesse de convergence de l'algorithme.

Nous avons ensuite présenté une application en reconstruction de phase pour un problème de tomographie. Dans cet exemple, l'image d'origine considérée ainsi que la matrice d'observation étaient à valeurs complexes, cependant seul le module des observations était accessible. Ainsi, le problème associé au modèle considéré était non convexe. De plus, les données étant de grande dimension, nous les avons divisées en blocs de dimensions plus petites. Nous nous sommes intéressés à l'influence du choix de la taille de ces blocs sur la vitesse de convergence. Nous avons montré, grâce à un exemple numérique, que l'utilisation de blocs combinée avec celle d'une stratégie de préconditionnement permettait d'accélérer significativement la vitesse de convergence.

Finalement, nous avons proposé un exemple de déconvolution aveugle pour la restauration de signaux sismiques parcimonieux. Dans cette application nous avons reconstruit de façon simultanée le signal sismique et le noyau de convolution. Afin de modéliser la parcimonie du signal estimé, nous avons proposé d'utiliser une fonction de régularisation non convexe s'écrivant comme le logarithme du quotient d'une norme ℓ_1 et d'une norme ℓ_2 . Dans cette application, nous avons montré, en particulier, l'intérêt pratique d'utiliser une règle acyclique de mise à jour des blocs. En effet, nous avons pu observer que la possibilité de mettre à jour plusieurs fois de suite le signal pour une seule mise à jour du noyau rendait la convergence de l'algorithme plus rapide. Ceci peut s'expliquer, d'une part, par le fait que la taille du signal à estimer était beaucoup plus grande que celle du noyau, et d'autre part par le fait que le signal était plus compliqué à estimer que le noyau du fait de sa parcimonie.

Les algorithmes prismaux-duaux alternés par blocs nous ont permis de développer des méthodes algorithmiques permettant de reconstruire des maillages 3D lorsque la position des nœuds de ces derniers est connue de façon approximative. Ce type de problème peut s'écrire sous la forme d'un problème inverse où l'objectif est d'estimer la position des nœuds du maillage original à partir d'une version dégradée par un bruit additif. Dans cet exemple nous avons montré l'intérêt d'utiliser des méthodes par blocs en terme d'occupation de la mémoire. En effet, ce type de méthode permet de limiter le nombre de calculs qui est fait à chaque itération, en fonction de l'architecture de calcul dont on dispose.

Finalement, nous nous sommes intéressés au problème d'optimisation de requêtes dans les systèmes de gestion de bases de données. Dans ce type d'application, des norme ℓ_1 et ℓ_∞ de quotients de fonctions sont utilisées afin de définir des problèmes réalisables. Nous avons proposé des méthodes basées sur les algorithmes prismaux-duaux pour résoudre ce type de problèmes. Pour cela, nous avons montré que l'opérateur proximal de la norme ℓ_1 d'un quotient de fonctions correspond à un opérateur de seuillage. Concernant la norme ℓ_∞ , son opérateur proximal n'a pas de forme explicite. Nous avons donc reformulé le problème de façon à faire intervenir des projections épigraphiques, et nous avons proposé un algorithme permettant de le résoudre.

9.2 PERSPECTIVES

Certains travaux futurs, se plaçant dans la continuité de ce qui a été présenté dans cette thèse, peuvent être envisagés.

Choisir une métrique dans les algorithmes prismaux-duaux

Dans les chapitres 3 et 5, nous avons proposé des algorithmes de type explicite-implicite accélérés grâce à l'introduction d'une métrique variant au cours des itérations. La méthode que nous avons donné pour choisir la métrique se base sur une interprétation MM de ce dernier de l'algorithme explicite-implicite. Bien que les algorithmes prismaux-duaux proposés dans le chapitre 7 ne fassent intervenir qu'une métrique constante au cours des itérations, il serait possible d'y faire apparaître une métrique variable. Cependant un problème majeur concerne le choix de cette métrique dans ce type d'algorithmes [Repetti *et al.*, 2012].

Une perspective intéressante est alors de proposer une méthode permettant de choisir une métrique (constante ou variant au cours des itérations) appropriée pour le préconditionnement des algorithmes prismaux-duaux. Nous pouvons supposer que trouver un lien entre les algorithmes prismaux-duaux et les stratégies MM permettrait de développer une méthode efficace pour le choix des préconditionneurs dans les méthodes primales-duales.

Proposer une version stochastique de l'algorithme explicite-implicite pour l'optimisation non convexe

Dans cette thèse, nous avons proposé des algorithmes alternés par bloc où les blocs mis à jour à chaque itération sont choisis selon deux règles différentes. Dans le cas de l'algorithme explicite-implicite alterné par bloc, les blocs sont prédéfinis de façon à former une partition des données globales. Au cours des itérations, les blocs sont ensuite choisis un à un pour être mis à jour selon une règle acyclique. Dans le cas des algorithmes stochastiques primaux-duaux, les blocs sont aussi prédéfinis, mais plusieurs d'entre eux peuvent être choisis aléatoirement à chaque itération. Ainsi, la mise à jour peut se faire de façon plus générale qu'avec la règle acyclique puisque les blocs peuvent se superposer et être traités en parallèles pendant une même itération. Il serait donc intéressant d'unifier ces résultats et d'étudier la convergence presque sûre d'algorithmes explicite-implicite stochastiques à métrique variable dans le cadre non convexe, ce qui permettrait de mettre à jour en parallèle les blocs sélectionnés à chaque itération.

Étendre les travaux réalisés sur les algorithmes primaux-duaux

Dans le chapitre 7, deux types d'algorithmes ont été proposés : des algorithmes alternés par blocs, et des algorithmes distribués. Les deux permettent de traiter des problèmes de minimisation convexe faisant intervenir un opérateur linéaire grâce à une formulation primale-duale du problème. Certaines perspectives concernant ces algorithmes peuvent être dressées :

- Dans le chapitre 8, nous avons traité un problème d'optimisation de requêtes dans les systèmes de gestion de bases de données en utilisant un algorithme primal-dual usuel [Chambolle et Pock, 2010]. L'exemple qui a été traité est de très petite dimension, et une extension naturelle serait de considérer un plus grand nombre de requêtes, pour des bases de données plus conséquentes, et d'appliquer un algorithme primal-dual alterné par bloc à ce problème.
- Nous avons montré dans la remarque 7.11 que les algorithmes distribués que nous avons proposés sont plus généraux que les méthodes récemment proposées dans la littérature dans le domaine de l'apprentissage. Ainsi, une seconde perspective concernant ce travail serait d'appliquer nos algorithmes distribués à des problèmes d'apprentissage régularisés tels que ceux abordés dans des cas particuliers dans [Bianchi et al., 2014; Jaggi et al., 2014].
- Les problèmes de minimisation pouvant être résolus par les algorithmes distribués que nous avons proposés sont de la forme

$$\text{trouver } \hat{x} \in \underset{x \in \mathcal{H}}{\text{Argmin}} \sum_{i=1}^m \phi_i(x),$$

où $\mathcal{H}$ est un espace de Hilbert réel, et $(\forall i \in \{1, \dots, m\})$ les fonctions ϕ_i sont convexes, et peuvent s'écrire comme une somme de fonctions (non) différentiables qui peuvent être composées avec des opérateurs linéaires. Ainsi, nous pouvons voir que lorsque les ϕ_i correspondent à des indicatrices d'ensembles convexes fermés non vides $C_i \subset \mathcal{H}$, le problème de

minimisation peut être reformulé de la façon suivante :

$$\text{trouver } \hat{\mathbf{x}} \in \bigcap_{i=1}^m C_i,$$

ou, de façon plus générale,

$$\text{trouver } \hat{\mathbf{x}} \text{ tel que } (\forall i \in \{1, \dots, m\}) \quad L_i \hat{\mathbf{x}} \in C_i,$$

où $(\forall i \in \{1, \dots, m\})$ L_i est un opérateur linéaire. Ce type de problème peut être résolu grâce à des algorithmes de projections alternées. Une perspective intéressante serait donc de faire le lien entre les algorithmes distribués proposés dans le chapitre 7 et les méthodes de projections alternées existantes (e.g. [Bauschke et Borwein, 1996; Combettes, 1994; Pesquet et Combettes, 1996]).

Généraliser la régularisation ℓ_1/ℓ_2

Plusieurs extensions peuvent être envisagées concernant la pénalisation de type ℓ_1/ℓ_2 proposée dans le chapitre 6.

- La première serait d’appliquer notre algorithme de minimisation alternée à métrique variable à des problèmes de déconvolution aveugle dans le cadre de la restauration d’images. La pénalisation ℓ_1/ℓ_2 permettant de promouvoir la parcimonie du signal auquel on l’applique, il suffirait dans ce cas de l’appliquer aux gradients verticaux et horizontaux de l’image à estimer (à l’instar de la variation totale). Cette méthode a été proposée dans [Krishnan et al., 2011]. Les auteurs de cet article ont montré son efficacité sur plusieurs exemples de reconstruction d’image. Notre méthode offrant de solides résultats de convergence et se comparant favorablement à la méthode proposée dans [Krishnan et al., 2011] pour des signaux 1D, il serait intéressant de comparer les deux méthodes pour des signaux 2D. On pourrait de plus envisager d’améliorer la méthode de [Krishnan et al., 2011] en appliquant la pénalisation ℓ_1/ℓ_2 , par exemple, à des gradients de l’image calculés sur un voisinage non local.
- Le problème de déconvolution aveugle est un problème non convexe, et introduire dans celui-ci une pénalisation non convexe rend le problème très difficile à résoudre. Nous avons pu remarquer, en particulier, que certains “pics” du signal estimé sont décalés, indifféremment vers la droite ou vers la gauche, par rapport aux “pics” du signal original. Les signaux reconstruits correspondent alors à des minimiseurs locaux du problème. Ainsi, un second point important qui pourrait être développé concerne l’initialisation de l’algorithme, qui permettrait d’augmenter les chances de trouver des minimiseurs globaux.
- Enfin, il pourrait être intéressant d’étudier une version plus générale de la régularisation ℓ_1/ℓ_2 afin de développer une classe de pénalisations qui pourraient s’adapter au mieux à d’autres applications. L’algorithme de minimisation alternée à métrique variable utilisé dans le chapitre 6 permet de traiter des fonctions diversifiées qui peuvent être ni convexes,

ni différentiables. Le but serait alors de développer une méthode efficace permettant d'utiliser un large spectre de pénalisations servant à promouvoir la parcimonie, en donnant des formules explicites afin de calculer facilement les métriques variables.

Traiter des problèmes de dominance stochastique

Nous avons montré que les algorithmes stochastiques présentés dans le chapitre 7 sont robustes aux erreurs pouvant apparaître lors des calculs des gradients et des opérateurs proximaux. De plus, nous avons montré que ces termes d'erreurs sont des variables aléatoires. Dans les applications qui ont été présentées dans le chapitre 7, l'avantage de ces erreurs aléatoires n'a pas été exploité.

Une perspective intéressante serait d'étudier un problème pratique où ce type d'erreurs est utile. C'est en particulier le cas pour les problèmes de dominance stochastique [Dentcheva et Ruszczyński, 2014; Shapiro *et al.*, 2009] apparaissant, par exemple, en apprentissage ou pour la gestion de portefeuille, où le critère à minimiser fait intervenir des espérances probabilistes.

Bibliographie

The MOSEK Optimization Toolbox. <http://www.mosek.com>.

Absil, P.-A., Mahony, R. et Andrews, B. : Convergence of the iterates of descent methods for analytic cost functions. *SIAM J. Optim.*, 6:531–547, 2005.

Afonso, M., Bioucas-Dias, J. M. et Figueiredo, M. A. T. : An augmented lagrangian approach to the constrained optimization formulation of imaging inverse problems. *IEEE Trans. Image Process.*, 20(3):681–695, Mar. 2011.

Ahmed, A., Recht, B. et Romberg, J. : Blind deconvolution using convex programming. *IEEE Trans. Inform. Theory*, 60(3):1711–1732, Mar. 2014.

Allain, M., Idier, J. et Goussard, Y. : On global and local convergence of half-quadratic algorithms. *IEEE Trans. Image Process.*, 15(5):1130–1142, May 2006.

Alotaibi, A., Combettes, P. L. et Shahzad, N. : Solving coupled composite monotone inclusions by successive Fejér approximations of their Kuhn-Tucker set. *SIAM J. Optim.*, 24(4):2076–2095, Dec. 2014.

Antoniadis, A., Leporini, D. et Pesquet, J.-C. : Wavelet thresholding for some classes of non-gaussian noise. *Statist. Neerlandica*, 56(4):434–453, 2002.

Asner, G. P. et Heidebrecht, K. B. : Spectral unmixing of vegetation, soil and dry carbon cover in arid regions : comparing multispectral and hyperspectral observations. *Int. J. Remote Sens.*, 23(19):3939–3958, Oct. 2002.

Attouch, H. et Bolte, J. : On the convergence of the proximal algorithm for nonsmooth functions involving analytic features. *Math. Program.*, 116:5–16, Jun. 2009.

Attouch, H., Bolte, J., Redont, P. et Soubeyran, A. : Proximal alternating minimization and projection methods for nonconvex problems. An approach based on the Kurdyka-Łojasiewicz inequality. *Math. Oper. Res.*, 35(2):438–457, 2010a.

Attouch, H., Bolte, J. et Svaiter, B. F. : Convergence of descent methods for semi-algebraic and tame problems : proximal algorithms, forward-backward splitting, and regularized Gauss-Seidel methods. *Math. Program.*, 137:91–129, Feb. 2011.

Attouch, H., Briceño Arias, L. M. et Combettes, P. L. : A parallel splitting method for coupled monotone inclusions. *SIAM J. Control Optim.*, 48:3246–3270, 2010b.

- Auslender, A. : Asymptotic properties of the Fenchel dual functional and applications to decomposition problems. *J. Optim. Theory Appl.*, 73(3):427–449, Jun. 1992.
- Barak, B., Kelner, J. et Steurer, D. : Rounding sum-of-squares relaxations. May 31-Jun. 3 2014.
- Bauschke, H. H. et Borwein, J. : On projection algorithms for solving convex feasibility problems. *SIAM Review*, 38(3):367–426, 1996.
- Bauschke, H. H. et Combettes, P. L. : *Convex Analysis and Monotone Operator Theory in Hilbert Spaces*. Springer, New York, 2011.
- Bauschke, H. H., Combettes, P. L. et Luke, D. R. : Phase retrieval, error reduction algorithm, and Fienup variants : a view from convex optimization. *J. Opt. Soc. Amer. A*, 19(7):1334–1345, Jul. 2002.
- Bauschke, H. H., Combettes, P. L. et Luke, D. R. : Hybrid projection-reflection method for phase retrieval. *J. Opt. Soc. Amer. A*, 20(6):1025–1034, June 2003.
- Bauschke, H. H., Combettes, P. L. et Luke, D. R. : A new generation of iterative transform algorithms for phase contrast tomography. In *Proc. IEEE Int. Conf. Acoust., Speech Signal Process.*, vol. 4, p. 89–92, Philadelphia, PA, 19–23 Mar. 2005.
- Bauschke, H. H., Combettes, P. L. et Noll, D. : Joint minimization with alternating Bregman proximity operators. *Pac. J. Optim.*, 2(3):401–424, Sep. 2006.
- Beck, A. et Teboulle, M. : Fast gradient-based algorithms for constrained total variation image denoising and deblurring problems. *IEEE Trans. Image Process.*, 18(11):2419 –2434, Nov. 2009.
- Becker, S. et Fadili, J. : A quasi-Newton proximal splitting method. Rap. tech., Jun. 2012. <http://arxiv.org/abs/1206.1156>.
- Becker, S. R. et Combettes, P. L. : An algorithm for splitting parallel sums of linearly composed monotone operators, with applications to signal recovery. *J. Nonlinear Convex Anal.*, 15(1): 137–159, Jan. 2014.
- Bect, J., Blanc-Féraud, L., Aubert, G. et Chambolle, A. : A l_1 -unified variational framework for image restoration. In T. Pajdla, J. M., éd. : *8th European Conference on Computer Vision (ECCV 2004)*, vol. 3024 de *Lecture Notes in Computer Science*, p. 1–13, Prague, Czech Republic, mai 2004. Springer.
- Benichoux, A., Vincent, E. et Gribonval, R. : A fundamental pitfall in blind deconvolution with sparse and shift-invariant priors. In *Proc. IEEE Int. Conf. Acoust., Speech Signal Process.*, Vancouver, BC, Canada, 26–31 May 2013.
- Bertsekas, D. P. : Projected Newton methods for optimization problems with simple constraints. *SIAM J. Control Optim.*, 20:221–246, 1982.
- Bertsekas, D. P. : *Nonlinear Programming*. Athena Scientific, Belmont, MA, 2nd édn, 1999.

- Bianchi, P., Hachem, W. et Iutzeler, F. : A stochastic coordinate descent primal-dual algorithm and applications to large-scale composite optimization. Rap. tech., 2014. <http://arxiv.org/abs/1407.0898>.
- Bioucas-Dias, J. M. et Figueiredo, M. A. : A new TwIST : Two-step iterative shrinkage/thresholding algorithms for image restoration. *IEEE Trans. Image Process.*, 16:2992–3004, Dec. 2007.
- Bioucas-Dias, J. M., Plaza, A., Dobigeon, N., Parente, M., Qian, D., Gader, P. et Chanussot, J. : Hyperspectral unmixing overview : Geometrical, statistical, and sparse regression-based approaches. *IEEE J. Sel. Topics Appl. Earth Observ.*, 5(2):354–379, Apr. 2012.
- Birgin, E. G., Martínez, J. M. et Raydan, M. : Nonmonotone spectral projected gradient methods on convex sets. *SIAM J. Optim.*, 10(4):1196–1211, 2000.
- Blumensath, T. et Davies, M. : Iterative thresholding for sparse approximations. *J. Fourier Anal. Appl.*, 14(5):629–654, 2008.
- Bochnak, J., Coste, M. et Roy, M.-F. : *Real Algebraic Geometry*. Springer, 1998.
- Boţ, R. I. et Hendrich, C. : A Douglas-Rachford type primal-dual method for solving inclusions with mixtures of composite and parallel-sum type monotone operators. *SIAM J. Optim.*, 23(4):2541–2565, Dec. 2013.
- Boţ, R. I. et Hendrich, C. : Convergence analysis for a primal-dual monotone + skew splitting algorithm with applications to total variation minimization. *J. Math. Imaging Vision*, 49(3):551–568, Jul. 2014.
- Bolte, J., Daniilidis, A. et Lewis, A. : The Łojasiewicz inequality for nonsmooth subanalytic functions with applications to subgradient dynamical systems. *SIAM J. Optim.*, 17:1205–1223, 2006a.
- Bolte, J., Daniilidis, A. et Lewis, A. : A nonsmooth morse-sard theorem for subanalytic functions. *J. Math. Anal. Appl.*, 321:729–740, 2006b.
- Bolte, J., Daniilidis, A., Lewis, A. et Shiota, M. : Clarke subgradients of stratifiable functions. *SIAM J. Optim.*, 18(2):556–572, 2007.
- Bolte, J., Daniilidis, A., Ley, O. et Mazet, L. : Characterizations of Łojasiewicz inequalities : subgradient flows, talweg, convexity. *Trans. Amer. Math. Soc.*, 362(6):3319–3363, 2010.
- Bolte, J. et Pauwels, E. : Majorization-minimization procedures and convergence of sqp methods for semi-algebraic and tame programs. Rap. tech., 2014. <http://arxiv.org/pdf/1409.8147.pdf>.
- Bolte, J., Sabach, S. et Teboulle, M. : Proximal alternating linearized minimization for nonconvex and nonsmooth problems. *To appear in Math. Program.*, 2013.
- Bonettini, S., Zanella, R. et Zanni, L. : A scaled gradient projection method for constrained image deblurring. *Inverse Probl.*, 25(1):015002+, 2009.

- Bonnans, J. F., Gilbert, J. C., Lemaréchal, C. et Sagastizábal, C. A. : A family of variable metric proximal methods. *Math. Program.*, 68:15–47, 1995.
- Bouman, C. et Sauer, K. : A unified approach to statistical tomography using coordinate descent optimization. *IEEE Trans. Image Process.*, 5(3):480–492, 1996.
- Boyd, S., Parikh, N., Chu, E., Peleato, B. et Eckstein, J. : Distributed optimization and statistical learning via the alternating direction method of multipliers. *Found. Trends Machine Learn.*, 8(1):1–122, 2011.
- Boyd, S., Xiao, L. et Mutapcic : Subgradient methods. *Lecture notes of EE392o, Stanford University*, 2003.
- Bregman, L. M. : The method of successive projection for finding a common point of convex sets. *Soviet Math. Dokl.*, 6:688–692, 1965.
- Bresson, X. : A short note for nonlocal tv minimization. Rap. tech., 2009. <http://citeseerx.ist.psu.edu/viewdoc/download;jsessionid=582C006CFE9B89E3DB4BD6FEEB588A5B?doi=10.1.1.210.471&rep=rep1&type=pdf>.
- Briceño-Arias, L. M. : A Douglas-Rachford splitting method for solving equilibrium problems. *Nonlinear Anal.*, 75(16):6053–6059, Nov. 2012.
- Briceños Arias, L. M. et Combettes, P. L. : A monotone + skew splitting model for composite monotone inclusions in duality. *SIAM J. Optim.*, 21(4):1230–1250, Oct. 2011.
- Burger, M., Sawatzky, A. et Steidl, G. : First order algorithms in variational image processing. Rap. tech., 2014. http://www.mathematik.uni-kl.de/fileadmin/image/steidl/publications/algs_book_revision_01.pdf.
- Burke, J. V. et Qian, M. : A variable metric proximal point algorithm for monotone operators. *SIAM J. Control Optim.*, 37:353–375, 1999.
- Campisi, P. et Egiazarian, K., éd. *Blind Image Deconvolution : Theory and Applications*. CRC Press, 2007.
- Candès, E., Strohmer, T. et Voroninski, V. : Phaselift : Exact and stable signal recovery from magnitude measurements via convex programming. *Comm. Pure Appl. Math.*, 66(8):1241–1274, 2013.
- Censor, Y. et Lent, A. : Optimization of $\log x$ entropy over linear equality constraints. *SIAM J. Control Optim.*, 25(4):921–933, juillet 1987.
- Chaâri, L., Chouzenoux, E., Pustelnik, N., Chaux, C. et Moussaoui, S. : OPTIMED : optimisation itérative pour la résolution de problèmes inverses de grande taille. *Traitement du Signal*, 28(3-4):329–374, 2011.
- Chaâri, L., Pustelnik, N., Chaux, C. et Pesquet, J.-C. : Solving inverse problems with overcomplete transforms and convex optimization techniques. In *SPIE, Wavelets XIII*, vol. 7446, San Diego, CA, USA, 2-6 Aug. 2009.

- Chambolle, A. : An algorithm for total variation minimization and applications. *J. Math. Imag. Vision*, 20(1-2):89–97, 2004.
- Chambolle, A. et Pock, T. : A first-order primal-dual algorithm for convex problems with applications to imaging. *J. Math. Imag. Vision*, 40(1):120–145, 2010.
- Chan, T. F. et Mulet, P. : On the convergence of the lagged diffusivity fixed point method in total variation image restoration. *SIAM J. Num. Anal.*, 2(36):354–367, 1999.
- Chang, T., Nedić, A. et Scaglione, A. : Distributed constrained optimization by consensus-based primal-dual perturbation method. *IEEE Trans. Automat. Control*, 59(6):1524–1538, Jun. 2014.
- Chartrand, R. : Nonconvex splitting for regularized low-rank + sparse decomposition. *IEEE Trans. Signal Process.*, 60:5810–5819, 2012.
- Chaux, C., Combettes, P. L., Pesquet, J.-C. et Wajs, V. R. : A variational formulation for frame-based inverse problems. *Inverse Probl.*, 23(4):1495–1518, 2007.
- Chaux, C., Duval, L., Benazza-Benyahia, A. et Pesquet, J.-C. : A nonlinear stein based estimator for multichannel image denoising. *IEEE Trans. Image Process.*, 56(8):3855–3870, Aug. 2008.
- Chen, G. H.-G. et Rockafellar, R. T. : Convergence rates in forward-backward splitting. *SIAM J. Optim.*, 7:421–444, 1997.
- Chen, P., Huang, J. et Zhang, X. : A primal-dual fixed point algorithm for convex separable minimization with applications to image restoration. *Inverse Problems*, 29(2):025011, 2013.
- Chierchia, G., Pustelnik, N., J.-C., P. et Pesquet-Popescu, B. : Epigraphical projection and proximal tools for solving constrained convex optimization problems : Part i. Rap. tech., 2013a. <http://arxiv.org/abs/1210.5844>.
- Chierchia, G., Pustelnik, N., Pesquet, J.-C. et Pesquet-Popescu, B. : An epigraphical convex optimization approach for multicomponent image restoration using non-local structure tensor. *In Proc. IEEE Int. Conf. Acoust., Speech Signal Process.*, p. 1359–1363, 26–31 May 2013b.
- Chinneck, J. W. : *Feasibility and Infeasibility in Optimization*. Springer, New York, 2008.
- Chouzenoux, E., Idier, J. et Moussaoui, S. : A Majorize-Minimize strategy for subspace optimization applied to image restoration. *IEEE Trans. Image Process.*, 20(18):1517–1528, Jun. 2011.
- Chouzenoux, E., Jezierska, A., Pesquet, J.-C. et Talbot, H. : A Majorize-Minimize subspace approach for $\ell_2 - \ell_0$ image regularization. *SIAM J. Imaging Sci.*, 28(1):563–591, 2013.
- Chouzenoux, E., Legendre, M., Moussaoui, S. et Idier, J. : Fast constrained least squares spectral unmixing using primal-dual interior-point optimization. *IEEE J. Sel. Topics Appl. Remote Sens.*, 7(1):59–69, 2014a.

- Chouzenoux, E., Pesquet, J.-C. et Repetti, A. : Variable metric forward-backward algorithm for minimizing the sum of a differentiable function and a convex function. *J. Optim. Theory Appl.*, 162(1), Jul. 2014b.
- Chu, M., Diele, F., Plemmons, R. et Ragni, S. : Optimality, computation, and interpretation of nonnegative matrix factorizations. *SIAM J. Matrix Anal. Appl.*, p. 4–8030, 2004.
- Clark, R. N., Swayze, G. A., Gallagher, A., King, T. V. et Calvin, W. M. : The U.S. geological survey digital spectral library : version 1 : 0.2 to 3.0 μm . *U.S. Geological Survey, Denver, CO, Open File Rep. 93-592*, 1993.
- Combettes, P. L. : Inconsistent signal feasibility problems : least-squares solutions in a product space. *IEEE Trans. Signal Process.*, 42(11):2955–2966, Nov. 1994.
- Combettes, P. L. : The convex feasibility problem in image recovery. In Hawkes, P., éd. : *Advances in Imaging and Electron Physics*, vol 96, p. 155–270. Academic Press, New York, 1996.
- Combettes, P. L. : Systems of structured monotone inclusions : duality, algorithms, and applications. *SIAM J. Optim.*, 23(4):2420–2447, Dec. 2013.
- Combettes, P. L., Condat, L., Pesquet, J.-C. et Vũ, B. C. : A forward-backward view of some primal-dual optimization methods in image recovery. In *Proc. IEEE Int. Conf. Image Process.*, p. 4141–4145, Paris, France, 27-30 Oct. 2014.
- Combettes, P. L., Dũng, D. et Vũ, B. C. : Dualization of signal recovery problems. *Set-Valued Var. Anal.*, 18:373–404, Dec. 2010.
- Combettes, P. L., Dũng, D. et Vũ, B. C. : Proximity for sums of composite functions. *J. Math. Anal. Appl.*, 380(2):680–688, Aug. 2011.
- Combettes, P. L. et Pesquet, J.-C. : A Douglas-Rachford splitting approach to nonsmooth convex variational signal recovery. *IEEE J. Selected Topics Signal Process.*, 1(4):564–574, Dec. 2007a.
- Combettes, P. L. et Pesquet, J.-C. : Proximal thresholding algorithm for minimization over orthonormal bases. *SIAM J. Optim.*, 18(4):1351–1376, Nov. 2007b.
- Combettes, P. L. et Pesquet, J.-C. : A proximal decomposition method for solving convex variational inverse problems. *Inverse Probl.*, 24(6), Dec. 2008.
- Combettes, P. L. et Pesquet, J.-C. : Proximal splitting methods in signal processing. In Bauschke, H. H., Burachik, R., Combettes, P. L., Elser, V., Luke, D. R. et Wolkowicz, H., édés : *Fixed-Point Algorithms for Inverse Problems in Science and Engineering*, p. 185–212. Springer-Verlag, New York, 2010.
- Combettes, P. L. et Pesquet, J.-C. : Primal-dual splitting algorithm for solving inclusions with mixtures of composite, Lipschitzian, and parallel-sum type monotone operators. *Set-Valued Var. Anal.*, p. 1–24, 2011.
- Combettes, P. L. et Pesquet, J.-C. : Stochastic quasi-Fejér block-coordinate fixed point iterations with random sweeping. *To appear in SIAM J. Optim.*, 2015.

- Combettes, P. L. et Vũ, B. C. : Variable metric quasi-Fejér monotonicity. *Nonlinear Anal.*, 78:17–31, 2013.
- Combettes, P. L. et Vũ, B. C. : Variable metric forward-backward splitting with applications to monotone inclusions in duality. *Optimization*, 63(9):1289–1318, Sept. 2014.
- Combettes, P. L. et Wajs, V. R. : Signal recovery by proximal forward-backward splitting. *Multiscale Model. Simul.*, 4(4):1168–1200, Nov. 2005.
- Comon, P. : Contrasts for multichannel blind deconvolution. *Signal Process. Lett.*, 3(7):209–211, Jul. 1996.
- Condat, L. : A primal-dual splitting method for convex optimization involving Lipschitzian, proximable and linear composite terms. *J. Optim. Theory Appl.*, 158(2):460–479, Aug. 2013.
- Coste, M. : An introduction to o-minimal geometry. *RAAG Notes, 81 pages, Institut de Recherche Mathématiques de Rennes*, Nov. 1999.
- Couprrie, C., Grady, L., Najman, L., Pesquet, J.-C. et Talbot, H. : Dual constrained TV-based regularization on graphs. *SIAM J. Imaging Sci.*, 6:1246–1273, 2013.
- Coutinho, D. et Figueiredo, M. : Information theoretic text classification using the Ziv-Merhav method. In *Proc. 2nd Iberian Conf. on Pattern Recognition and Image Analysis (IbPRIA)*, p. 355–362, Estoril, Portugal, 2005.
- Dahlquist, G. et Bjorck, A. : *Numerical Methods*. Dover Publication, Mineola, N.Y., 2003.
- Davidoiu, V., Sixou, B., Langer, M. et Peyrin, F. : Nonlinear phase retrieval using projection operator and iterative wavelet thresholding. *IEEE Signal Process. Lett.*, 19(9):579–582, Sept 2012.
- De Pierro, A. R. : A modified expectation maximization algorithm for penalized likelihood estimation in emission tomography. *IEEE Trans. Med. Imag.*, 14(1):132–137, Mar. 1995.
- Deledalle, C.-A., Vaiteer, S., Fadili, J. M. et Peyré, G. : Stein COnsistent Risk Estimator (SCORE) for hard thresholding. In *Proc. SPARS'13*, 8–11 Jul. 2013.
- Deledalle, C.-A., Vaiteer, S., Fadili, J. M. et Peyré, G. : Stein Unbiased GrAdient estimator of the Risk (SUGAR) for multiple parameter selection. Aug. 2014.
- Demagnet, L. et Hand, P. : Scaling law for recovering the sparsest element in a subspace. *Information and Inference*, 2014.
- Demoment, G. : Image reconstruction and restoration : Overview of common estimation structure and problems. *IEEE Trans. Acoust., Speech, Signal Process.*, 37(12):2024–2036, Dec. 1989.
- Dentcheva, D. et Ruszczyński, A. : Risk-averse portfolio optimization via stochastic dominance constraints. In Lee, C.-F. et Lee, J. C., eds : *Handbook of Financial Econometrics and Statistics*, p. 2281–2302. Springer, 2014.

- Dobigeon, N. et Brun, N. : Spectral mixture analysis of eels spectrum-images. *Ultramicroscopy*, 120:25–34, Sept. 2012.
- Donoho, D. L. : Compressed sensing. *IEEE Trans. Inf. Theory*, 52(4):1289–1306, 2006.
- Donoho, D. L., Johnstone, I. M., Kerkycharian, G. et Picard, D. : Wavelet shrinkage : Asymptopia? *J. R. Stat. Soc. Ser. B Stat. Methodol.*, 57(2):301–369, 1995.
- Eckstein, J. et Bertsekas, D. P. : On the Douglas-Rachford splitting method and the proximal point algorithm for maximal monotone operators. *Math. Program.*, 55(1-3):293–318, 1992.
- Elad, M., Milanfar, P. et Ron, R. : Analysis versus synthesis in signal priors. *Inverse Problems*, 23(3):947–968, Jun. 2007.
- Elmoataz, A., Lezoray, O. et Bogleux, S. : Nonlocal discrete regularization on weighted graphs : A framework for image and manifold processing. *IEEE Trans. Image Process.*, 17:1047–1060, 2008.
- Erdogan, H. et Fessler, J. A. : Monotonic algorithms for transmission tomography. *IEEE Trans. Med. Imag.*, 18(9):801–814, Sept. 1999.
- Esser, E., Lou, Y. et Xin, J. : A method for finding structured sparse solutions to non-negative least squares problems with applications. *SIAM J. Imaging Sci.*, 6(4):2010–2046, 2013.
- Esser, E., Zhang, X. et Chan, T. : A general framework for a class of first order primal-dual algorithms for convex optimization in imaging science. *SIAM J. Imaging Sci.*, 3(4):1015–1046, 2010.
- Fessler, J. A. : Grouped coordinate ascent algorithms for penalized-likelihood transmission image reconstruction. *IEEE Trans. Med. Imag.*, 16:166–175, 1997.
- Fienup, J. R. : Phase retrieval algorithms : A comparison. *Appl. Opt.*, 21:2758–2769, Aug. 1982.
- Figueiredo, M. A. T. et Bioucas-Dias, J. M. : Restoration of Poissonian images using alternating direction optimization. *IEEE Trans. Image Process.*, 19(12):3133–3145, Dec. 2010.
- Figueiredo, M. A. T. et Nowak, R. D. : Deconvolution of Poissonian images using variable splitting and augmented Lagrangian optimization. In *Proc. IEEE Workshop Stat. Sign. Proc.*, p. x+4, Cardiff, United Kingdom, 31 Aug. - 3 Sept. 2009.
- Figueiredo, M. A. T., Nowak, R. D. et Wright, S. : Gradient projection for sparse reconstruction : Application to compressed sensing and other inverse problems. *IEEE J. Sel. Topics Signal Process.*, 1(4):586–597, Dec. 2007.
- Fletcher, R. : *Practical Methods of Optimization*. Wiley, 2nd édn, 2000.
- Fogel, F., Waldspurger, I. et D’Aspremont, A. : Phase retrieval for imaging problems. Rap. tech., Apr. 2013. <http://arxiv.org/abs/1304.7735/>.

- Foi, A., Trimeche, M., Katkovnik, V. et Egiazarian, K. : Practical Poissonian-Gaussian noise modeling and fitting for single-image raw-data. *IEEE Trans. Image Process.*, 17(10):1737–1754, Oct. 2008.
- Fortet, R. M. : *Fonctions et Distributions Aléatoires dans les Espaces de Hilbert*. Hermès, Paris, 1995.
- Fortin, M. et Glowinski, R., éd. *Augmented Lagrangian Methods : Applications to the Numerical Solution of Boundary-Value Problems*. Elsevier Science Ltd, Amsterdam : North-Holland, 1983.
- Frankel, P., Garrigos, G. et Peyrouquet, J. : Splitting methods with variable metric for Kurdyka-Łojasiewicz functions and general convergence rates. *To appear in J. Optim. Theory Appl.*, 2014.
- Fuchs, J.-J. : Convergence of a sparse representations algorithm applicable to real or complex data. *IEEE J. Sel. Topics Signal Process.*, 1(14):598–605, Dec. 2007.
- Gabay, D. et Mercier, B. : A dual algorithm for the solution of nonlinear variational problems via finite elements approximations. *Comput. Math. Appl.*, 2:17–40, 1976.
- Geman, D. et Reynold, G. : Constrained restoration and the recovery of discontinuities. *IEEE Trans. Pattern Anal. Mach. Intell.*, 14(3):367–383, Mar. 1992.
- Geman, D. et Yang, C. : Nonlinear image recovery with half-quadratic regularization. *IEEE Trans. Image Process.*, 4(7):932–946, July 1995.
- Gerchberg, R. W. et Saxton, W. O. : A practical algorithm for the determination of phase from image and diffraction plane pictures. *Optik*, 35:237–246, 1972.
- Gilboa, G. et Osher, S. : Nonlocal operators with applications to image processing. *Multiscale Model. Simul.*, 7(3):1005–1028, 2008.
- Giovannelli, J.-F. et Coulais, A. : Positive deconvolution for superimposed extended source and point sources. *Astron. Astrophys.*, 439:401–412, 2005.
- Goldstein, T., Esser, E. et Baraniuk, R. : Adaptive primal-dual hybrid gradient methods for saddle-point problems. 2013.
- Goldstein, T. et Osher, S. : The split Bregman method for ℓ_1 -regularized problems. *SIAM J. Imaging Sci.*, 2:323–343, 2009.
- Golub, G. H. et Van Loan, C. F. : *Matrix Computations*. Johns Hopkins University Press, Baltimore, Maryland, 3rd édn, 1996.
- Gonzalez, R. C. et Woods, R. E. : *Digital Image Processing*. Addison–Wesley, Reading, second édn, 2002.
- Gowen, A., O'Donnell, C., Cullen, P., Downey, G. et Frias, J. : Hyperspectral imaging : an emerging process analytical tool for food quality and safety control. *Trends in Food Science and Technology*, 18(12):590–598, 2007.

- Grady, L. J. et Polimeni, J. R. : *Discrete Calculus : Applied Analysis on Graphs for Computational Science*. Springer-Verlag, New York, 2010.
- Gray, W. C. : Variable norm deconvolution. Rap. tech., 1978. http://sepwww.stanford.edu/oldreports/sep14/14_19.pdf.
- Griesel, M. A. : A linear Remes-type algorithm for relative error approximation. *SIAM Journal on Numerical Analysis*, 11(1):170–173, 1974.
- Guigay, J. P., Langer, M., Boistel, R. et Cloetens, P. : Mixed transfer function and transport of intensity approach for phase retrieval in the Fresnel region. *Opt. Lett.*, 32(12):1617–1619, Jun 2007.
- Hadjisavvas, N. et Schaible, S. : Generalized monotone multivalued maps. In Floudas, C. A. et Pardalos, P. M., éd. : *Encyclopedia of Optimization*, p. 1193–1197. Springer, New York, 2nd éd., 2009.
- Hanke-Bourgeois, M. : *Grundlagen der Numerischen Mathematik und des Wissenschaftlichen Rechnens*. Teubner, Stuttgart, 2002.
- Harizanov, S., Pesquet, J.-C. et Steidl, G. : Epigraphical projection for solving least squares anscombe transformed constrained optimization problems. In et al., A. K., éd. : *Scale-Space and Variational Methods in Computer Vision. Lecture Notes in Computer Science, SSVM 2013, LNCS 7893*, p. 125–136, Berlin, 2013. Springer.
- Harrison, R. : Phase problem in crystallography. *J. Opt. Soc. Amer. A*, 10(5):1046–1055, 1993.
- Haykin, S., éd. *Blind Deconvolution*. Prentice Hall, 1994.
- He, B. et Yuan, X. : Convergence analysis of primal-dual algorithms for a saddle-point problem : from contraction perspective. *SIAM J. Imaging Sci.*, 5(1):119–149, 2012.
- Healey, G. E. et Kondepudy, R. : Radiometric CCD camera calibration and noise estimation. *IEEE Trans. Pattern Anal. Mach. Intell.*, 16(3):267–276, Mar. 1994.
- Hiriart-Urruty, J.-B. et Lemaréchal, C. : *Convex Analysis and Minimization Algorithms*. Springer-Verlag, New York, 1993.
- Hoyer, P. : Non-negative matrix factorization with sparseness constraints. *J. Mach. Learn. Res.*, 5:1457–1469, 2004.
- Huber, P. J. : *Robust Statistics*. John Wiley, New York, NY, USA, 1981.
- Hunter, D. R. et Lange, K. : A tutorial on MM algorithms. *Amer. Stat.*, 58(1):30–37, Feb. 2004.
- Hurley, N. et Rickard, S. : Comparing measures of sparsity. *IEEE Trans. Inf. Theory*, 55(10):4723–4741, Oct. 2004.
- Idier, J. : Convex half-quadratic criteria and interacting auxiliary variables for image restoration. *IEEE Trans. Image Process.*, 10(7):1001–1009, July 2001.

- Ioannidis, Y. E. et Christodoulakis, S. : On the propagation of errors in the size of join results. *SIGMOD*, 1991.
- Iutzeler, F., Bianchi, P., Ciblat, P. et Hachem, W. : Asynchronous distributed optimization using a randomized alternating direction method of multipliers. *In Proc. 52nd Conf. Decision Control*, p. 3671–3676, Florence, Italy, 10-13 Dec. 2013.
- Jacobson, M. W. et Fessler, J. A. : An expanded theoretical treatment of iteration-dependent majorize-minimize algorithms. *IEEE Trans. Image Process.*, 16(10):2411–2422, Oct. 2007.
- Jaganathan, K., Oymak, S. et Hassibi, B. : Recovery of sparse 1-D signals from the magnitudes of their Fourier transform. *In Proc. IEEE Int. Symp. Inf. Theory (ISIT 2012)*, p. 1473–1477, Cambridge, MA, 1-6 Jul. 2012.
- Jaggi, M., Smith, V., Takáč, M., Terhorst, J., Krishnan, S., Hofmann, T. et Jordan, M. I. : Communication-efficient distributed dual coordinate ascent. *In Proc. Ann. Conf. Neur. Inform. Proc. Syst.*, p. 3068–3076, Montreal, Canada, 8-11 Dec. 2014.
- Janesick, J. R. : *Photon Transfer*, vol. PM170. SPIE Press Book, Bellingham, WA, 2007.
- Jezierska, A., Chouzenoux, E., Pesquet, J.-C. et Talbot, H. : A primal-dual proximal splitting approach for restoring data corrupted with poisson-gaussian noise. *In Proc. Int. Conf. Acoust., Speech Signal Process.*, p. 1085–1088, Kyoto, Japan, 25–30 Mar. 2012.
- Jezierska, A., Chouzenoux, E., Pesquet, J.-C. et Talbot, H. : A convex approach for image restoration with exact poisson-gaussian likelihood. Rap. tech., 2013. <https://hal.archives-ouvertes.fr/hal-00922151>.
- Ji, H., Li, J., Shen, Z. et Wang, K. : Image deconvolution using a characterization of sharp images in wavelet domain. *Appl. Comp. Harm. Analysis*, 32(2):295–304, 2012.
- Jia, S. et Qian, Y. : Constrained nonnegative matrix factorization for hyperspectral unmixing. *IEEE Trans. Geosci. Remote Sens.*, 47(1):161–173, 2009.
- Kaaresen, K. F. et Tøft, T. : Multichannel blind deconvolution of seismic signals. *Geophysics*, 63(6):2093–2107, Nov. 1998.
- Kato, M., Yamada, I. et Sakaniwa, K. : A set-theoretic blind image deconvolution based on hybrid steepest descent method. *IEICE Trans. Fund. Electron. Comm. Comput. Sci.*, E82-A(8):1443–1449, Aug. 1999.
- Keshava, N. et Mustard, J. F. : Spectral unmixing. *IEEE Signal Process. Mag.*, 19(1):44–57, Jan. 2002.
- Komodakis, N. et Pesquet, J.-C. : Playing with duality : An overview of recent primal-dual approaches for solving large-scale optimization problems. *To appear in IEEE Signal Process. Mag.*, 2014.
- Kowalski, M. : Proximal algorithm meets a conjugate descent. Rap. tech., Jun. 2010. <http://hal.archives-ouvertes.fr/docs/00/50/57/33/PDF/proxconj.pdf>.

- Krishnan, D., Tay, T. et Fergus, R. : Blind deconvolution using a normalized sparsity measure. *In Proc. IEEE Conf. Comput. Vis. Pattern Recognit.*, p. 233–240, Colorado Springs, CO, USA, Jun. 21–25, 2011.
- Kundur, D. et Hatzinakos, D. : Blind image deconvolution. *IEEE Signal Process. Mag.*, 13(3):43–64, May 1996a.
- Kundur, D. et Hatzinakos, D. : Blind image deconvolution revisited. *IEEE Signal Process. Mag.*, 13(6):61–63, Nov. 1996b.
- Kurdyka, K. : On gradients of functions definable in o-minimal structures. *Ann. Inst. Fourier*, 48(3):769–783, 1998.
- Lange, K. et Fessler, J. A. : Globally convergent algorithms for maximum a posteriori transmission tomography. *IEEE Trans. Image Process.*, 4(10):1430–1438, Oct. 1995.
- Lee, D. D. et Seung, H. S. : Algorithms for non-negative matrix factorization. *In Advances in Neural and Information Processing Systems*, vol. 13, p. 556–562, 2001.
- Lee, J. D., Sun, Y. et Saunders, M. A. : Proximal Newton-type methods for convex optimization. *In Bartlett, P., Pereira, F., Burges, C., Bottou, L. et Weinberger, K., eds : Advances in Neural Information Processing Systems (NIPS)*, vol. 25, p. 827 – 835, 2012.
- Li, J., Shen, Z., Jin, R. et Zhang, X. : A reweighted ℓ_2 method for image restoration with Poisson and mixed Poisson-Gaussian noise. Rap. tech., 2012. UCLA Preprint, <ftp://ftp.math.ucla.edu/pub/camreport/cam12-84.pdf>.
- Lobo, M. S., Vandenberghe, L., Boyd, S. et Lebrete, H. : Applications of second order cone programming. *Linear Algebra and its Applications*, 284:193–228, 1998.
- Łojasiewicz, S. : *Une propriété topologique des sous-ensembles analytiques réels*, p. 87–89. Editions du centre National de la Recherche Scientifique, 1963.
- Loris, I. et Verhoeven, C. : On a generalization of the iterative soft-thresholding algorithm for the case of non-separable penalty. *Inverse Problems*, 27(12):125007, 2011.
- Lotito, P. A., Parente, L. A. et Solodov, M. V. : A class of variable metric decomposition methods for monotone variational inclusions. *J. Convex Anal.*, 16:857–880, 2009.
- Luenberger, D. G. : *Linear and Nonlinear Programming*. Addison-Wesley, Reading, Massachusetts, 1973.
- Luo, Z. Q. et Tseng, P. : On the convergence of the coordinate descent method for convex differentiable minimization. *J. Optim. Theory Appl.*, 72(1):7–35, 1992a.
- Luo, Z. Q. et Tseng, P. : On the linear convergence of descent methods for convex essentially smooth minimization. *SIAM J. Control Optim.*, 30(2):408–425, 1992b.

- Ma, W.-K., Bioucas-Dias, J. M., Tsung-Han, C., Gillis, N., Gader, P., Plaza, A., Ambikapathi, A. et Chong-Yung, C. : A signal processing perspective on hyperspectral unmixing : Insights from remote sensing. *IEEE Signal Process. Mag.*, 31(1):67–81, Jan. 2014.
- Mallat, S. : *A Wavelet Tour of Signal Processing*. Academic Press, Burlington, MA, 2nd édn, 2009.
- Mangasarian, O. L. : *Nonlinear Programming*. McGraw-Hill, New York, NY, 1969.
- Marjanovic, G. et Solo, V. : On l_q optimization and matrix completion. *IEEE Trans. Image Process.*, 60(11):5714–5724, Nov. 2012.
- Markl, V., Haas, P., Kutsch, M., Megiddo, N., Srivastava, U. et Tran, T. : Consistent selectivity estimation via maximum entropy. *VLDB Journal*, 16(1):55–76, Jan. 2007.
- Metcalf, F. T. : Error measures and their associated means. *Journal of Approximation Theory*, 17:57–65, 1976.
- Michelot, C. : A finite algorithm for finding the projection of a point onto the canonical simplex of $\mathbb{R}^n$. *J. Optim. Theory Appl.*, 50(1):195–200, Jul. 1986.
- Moerkotte, G., Neumann, T. et Steidl, G. : Preventing bad plans by bounding the impact of cardinality estimation errors. *Proc. of the VLDB*, 2(1):982–993, 2009.
- Monteiro, R. D. C. et Svaiter, B. F. : Convergence rate of inexact proximal point methods with relative error criteria for convex optimization. Rap. tech., 2010. http://www.optimization-online.org/DB_HTML/2010/08/2714.html.
- Mordukhovich, B. S. : *Variational Analysis and Generalized Differentiation. Vol. I : Basic theory*, vol. 330 de *Series of Comprehensive Studies in Mathematics*. Springer-Verlag, Berlin-Heidelberg, 2006.
- Moreau, E. et Pesquet, J.-C. : Generalized contrasts for multichannel blind deconvolution of linear systems. *IEEE Signal Process. Lett.*, 4(6):182–183, Jun. 1997.
- Moreau, J. J. : Proximité et dualité dans un espace hilbertien. *Bull. Soc. Math. France*, 93:273–299, 1965.
- Mørup, M., Madsen, K. H. et Hansen, L. K. : Approximate L_0 constrained non-negative matrix and tensor factorization. In *Proc. Int. Symp. Circuits Syst.*, p. 1328–1331, May 2008.
- Moussaoui, S., Chouzenoux, E. et Idier, J. : Primal-dual interior point optimization for penalized least squares estimation of abundance maps in hyperspectral imaging. In *Proc. IEEE Workshop on Hyperspectral Image and Signal Processing : Evolution in Remote Sensing (WHISPERS'12)*, p. 1–4, June 2012.
- Mukherjee, S. et Seelamantula, C. : An iterative algorithm for phase retrieval with sparsity constraints : application to frequency domain optical coherence tomography. In *Proc. IEEE Int. Conf. Acoust., Speech and Signal Process. (ICASSP 2012)*, p. 553–556, Kyoto, Japan, 25–30 Mar. 2012.

- Nandi, A. K., Mampel, D. et Roscher, B. : Blind deconvolution of ultrasonic signals in nondestructive testing applications. *IEEE Trans. Signal Process.*, 45(5):1382–1390, 1997.
- Necoara, I. et Patrascu, A. : A random coordinate descent algorithm for optimization problems with composite objective function and linear coupled constraints. *Comput. Optim. Appl.*, 57(2):307–337, Mar. 2014.
- Nedić, A. et Ozdaglar, A. : Cooperative distributed multi-agent optimization. In Palomar, D. P. et Eldar, Y. C., édés : *Convex Optimization in Signal Processing and Communications*, chap. 10, p. 340–386. Cambridge University Press, Cambridge, UK, 2010.
- Nesterov, Y. : Gradient methods for minimizing composite objective function. 2007.
- Ochs, P., Chen, Y., Brox, T. et Pock, T. : iPiano : inertial proximal algorithm for non-convex optimization. *SIAM J. Imaging Sci.*, 7(2):1388–1419, 2014.
- Ochs, P., Dosovitskiy, A., Brox, T. et Pock, T. : On iteratively reweighted algorithms for non-smooth non-convex optimization in computer vision. *SIAM J. Imaging Sci.*, 8(1):331–372, 2015.
- Ortega, J. M. et Rheinboldt, W. C. : *Iterative Solution of Nonlinear Equations in Several Variables*. Academic Press, New York, NY, 1970.
- Paatero, P. et Tapper, U. : Positive matrix factorization : a nonnegative factor model with optimal utilization of error estimates of data values. *Environmetrics*, 5:111–126, 1994.
- Palis, J. et De Melo, W. : *Geometric Theory of Dynamical Systems. An Introduction*. Springer-Verlag, NewYork–Berlin, 1982.
- Palma-Amestoy, R., Provenzi, E., Bertalmio, M. et Caselles, V. : A perceptually inspired variational framework for color enhancement. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 31(3):458–474, 2009.
- Parente, L. A., Lotito, P. A. et Solodov, M. V. : A class of inexact variable metric proximal point algorithms. *SIAM J. Optim.*, 19:240–260, 2008.
- Parikh, N. et Boyd, S. : Proximity algorithms. *Foundations and Trends in Optimization*, 1(3):123–231, 2013.
- Pesquet, J.-C., Benazza-Benyahia, A. et Chaux, C. : A SURE approach for digital signal/image deconvolution problems. *IEEE Trans. Signal Process.*, 57(12):4616–4632, Dec 2009.
- Pesquet, J.-C. et Combettes, P. : Wavelet synthesis by alternating projections. *IEEE Trans. Signal Process.*, 44(3):728–732, Mar. 1996.
- Pesquet, J.-C. et Pustelnik, N. : A parallel inertial proximal optimization method. *Pac. J. Optim.*, 8(2):273–305, Apr. 2012.
- Peyré, G. : A review of adaptive image representations. *IEEE J. Selected Topics Signal Process.*, 5(5):896–911, 2011.

- Pham, M. Q., Duval, L., Chaux, C. et Pesquet, J.-C. : A primal-dual proximal algorithm for sparse template-based adaptive filtering : Application to seismic multiple removal. *IEEE Trans. Signal Process.*, 62(16):4256–4269, Aug. 2014.
- Pock, T. et Chambolle, A. : Diagonal preconditioning for first order primal-dual algorithms in convex optimization. *In Proc. IEEE Int. Conf. Comput. Vis.*, p. 1762–1769, Barcelona, Spain, 6-13 Nov. 2011.
- Pock, T., Chambolle, A., Cremers, D. et Bischof, H. : A convex relaxation approach for computing minimal partitions. *IEEE Conf. Computer Vision and Pattern Recognition*, p. 810–817, 2009.
- Powell, M. J. D. : On search directions for minimization algorithms. *Math. Program.*, 4:193–201, 1973.
- Pustelnik, N., Borgnat, P. et Flandrin, P. : Empirical Mode Decomposition revisited by multi-component nonsmooth convex optimization. *Signal Process.*, 102:313–331, Sept. 2014.
- Pustelnik, N., Pesquet, J.-C. et Chaux, C. : Relaxing tight frame condition in parallel proximal methods for signal restoration. *IEEE Trans. Signal Process.*, 60(2):968–973, Feb. 2012.
- Raguet, H., Fadili, J., et Peyré, G. : A generalized forward-backward splitting. *SIAM J. Imaging Sci.*, 6(3):1199–1226, 2013.
- Razaviyayn, M., Hong, M. et Luo, Z. : A unified convergence analysis of block successive minimization methods for nonsmooth optimization. *SIAM J. Optim.*, 23(2):1126–1153, 2013.
- Repetti, A., Chouzenoux, E. et Pesquet, J.-C. : A penalized weighted least squares approach for restoring data corrupted with signal-dependent noise. *In Proc. Eur. Sig. and Image Proc. Conference*, p. 1553–1557, Bucharest, Romania, 27-31 Aug. 2012.
- Richtárik, P. et Talác, M. : Iteration complexity of randomized block-coordinate descent methods for minimizing a composite function. *Math. Program.*, 144(1-2):1–38, Apr. 2014.
- Ricker, N. : The form and nature of seismic waves and the structure of seismograms. *Geophysics*, 5(4):348–366, 1940.
- Rockafellar, R. T. : *Convex Analysis*. Princeton University Press, 1970.
- Rockafellar, R. T. : *Conjugate Duality and optimization*. SIAM, 1974.
- Rockafellar, R. T. : Monotone operators and the proximal point algorithm. *SIAM J. Control Optim.*, 14:877–898, 1976.
- Rockafellar, R. T. et Wets, R. J.-B. : *Variational Analysis*. Springer-Verlag, 1st éd., 1997.
- Rudin, L. I., Osher, S. et Fatemi, E. : Nonlinear total variation based noise removal algorithms. *Phys. D*, 60:259–268, 1992.
- Saquist, S., Zheng, J., Bouman, C. A. et Sauer, K. D. : Parallel computation of sequential pixel updates in statistical tomographic reconstruction. *Proc. IEEE Int. Conf. Image Processing*, 2:93–96, 1995.

- Setzer, S., Steidl, G., Teuber, T. et Moerkotte, G. : Approximation related to quotient functionals. *Journal of Approximation Theory*, 162(3):545–558, 2010.
- Shapiro, A., Dentcheva, D. et Ruszczyński, A. : *Lectures on Stochastic Programming : Modeling and Theory*. Society for Industrial and Applied Mathematics and the Mathematical Programming Society, Philadelphia, USA, 2009.
- Shechtman, Y., Beck, A. et Eldar, Y. : GESPAR : Efficient phase retrieval of sparse signals. *IEEE Trans. Signal Process.*, 62(4):928–938, Feb 2014.
- Shechtman, Y., Eldar, Y., Szameit, A. et Segev, M. : Efficient coherent diffractive imaging for sparsely varying objects. *Optics Express*, 21(5):6327–6338, 2013.
- Shiota, M. : *Geometry of Subanalytic and emialgebraic Sets*, vol. 140. Progress in Mathematics, 1997.
- Shor, N. Z. : *Minimization Methods for Non-differentiable Functions*. Springer-Verlag, 1985.
- Shuman, D. I., Narang, S. K., Forssard, P., Ortega, A. et Vandergheynst, P. : The emerging field of signal processing on graphs : Extending high-dimensional data analysis to networks and other irregular domains. *IEEE Signal Processing Magazine*, 30(3):83–98, 2013.
- Slaney, M. et Kak, A. : Principles of computerized tomographic imaging. *SIAM, Philadelphia*, 1988.
- Slavakis, K., Kopsinis, Y., Theodoridis, S. et McLaughlin, S. : Generalized thresholding and online sparsity-aware learning in a union of subspaces. *IEEE Trans. Signal Process.*, 61(15):3760–3773, Aug. 2013.
- Solodov, M. V. et Svaiter, B. F. : A unified framework for some inexact proximal point algorithms. *Numer. Funct. Anal. Optim.*, 22:1013–1035, 2001.
- Sotthivirat, S. et Fessler, J. A. : Image recovery using partitioned-separable paraboloidal surrogate coordinate ascent algorithms. *IEEE Trans. Signal Process.*, 11(3):306–317, 2002.
- Svaiter, B. F. : A class of Fejér convergent algorithms, approximate resolvents and the hybrid proximal-extragradient method. *J. Optim. Theory Appl.*, 162:133–153, July 2014.
- Takahata, A. K., Nadalin, E. Z., Ferrari, R., Duarte, L. T., Suyama, R., Lopes, R. R., Romano, J. M. T. et Tygel, M. : Unsupervised processing of geophysical signals : A review of some key aspects of blind deconvolution and blind source separation. *IEEE Signal Process. Mag.*, 29(4):27–35, Jul. 2012.
- Tappenden, R., Richtárik, P. et Gondzio, J. : Inexact coordinate descent : Complexity and preconditioning. Rap. tech., 2013. <http://arxiv.org/abs/1304.5530>.
- Taubin, G., Zhang, T. et Golub, G. : Optimal surface smoothing as filter design. In Buxton, B. et Cipolla, R., eds : *Computer Vision ECCV '96*, vol. 1064 de *Lecture Notes in Computer Science*, p. 283–292. Springer Berlin Heidelberg, 1996.

- Teuber, T., Steidl, G. et Chan, R.-H. : Minimization and parameter estimation for seminorm regularization models with I -divergence constraints. *Inverse Problems*, 29:035007, Mar. 2013.
- Themelis, K. E., Schmidt, F., Sykioti, O., Rontogiannis, A. A., Koutroumbas, K. D. et Daglis, I. A. : On the unmixing of MEx/OMEGA hyperspectral data. *Planetary and Space Science*, 68(1):34–41, 2012.
- Tian, H., Fowler, B. et Gamal, A. E. : Analysis of temporal noise in CMOS photodiode active pixel sensor. *IEEE Solid State Circuits Mag.*, 36(1):92–101, Jan. 2001.
- Tignol, J. P. : *Galois' Theory of Algebraic Equations*. World Scientific, Singapore, 2001.
- Tomazeli Duarte, L., Moussaoui, S. et Jutten, C. : Source separation in chemical analysis : Recent achievements and perspectives. *IEEE Signal Processing Magazine*, 31(3):135–146, May 2014.
- Towfic, Z. J. et Sayed, A. H. : Stability and performance limits of adaptive primal-dual networks. Rap. tech., 2014. <http://arxiv.org/pdf/1408.3693.pdf>.
- Tran-Dinh, Q., Kyrillidis, A. et Cevher, V. : Composite self-concordant minimization. *To appear in J. Mach. Learn. Res.*, 2014.
- Tseng, P. : A modified forward-backward splitting method for maximal monotone mappings. *SIAM J. Control Optim.*, 38(2):431–446, 1998.
- Tseng, P. : Convergence of a block coordinate descent method for nondifferentiable minimization. *J. Optim. Theory Appl.*, 109(3):475–494, Jun. 2001.
- Vũ, B. C. : A splitting algorithm for dual monotone inclusions involving cocoercive operators. *Advances in Computational Mathematics*, 38(3):667–681, 2013.
- Van den Dries, L. : Tame topology and o-minimal structures. vol. 248 de *London Mathematical Society Lecture Note Series*. Cambridge University Press, 1998.
- Van den Dries, L. et Miller, C. : Geometries categories and o-minimal structures. *Duke Math. J.*, 84:497–540, 1996.
- Walden, A. T. et Hosken, J. W. J. : The nature of the non-Gaussianity of primary reflection coefficients and its significance for deconvolution. *Geophys. Prospect.*, 34(7):1038–1066, 1986.
- Waldspurger, I., d'Aspremont, A. et Mallat, S. : Phase recovery, maxcut and complex semidefinite programming. *To appear in Math. Program.*, 2013.
- Walther, A. : The question of phase retrieval in optics. *J. Modern Opt.*, 10(1):41–49, 1963.
- Werlberger, M., Pock, T. et Bischof, H. : Motion estimation with non-local total variation regularization. In *IEEE Comput. Soc. Conf. Comput. Vision and Pattern Recogn. (CVPR)*, p. 2464–2471, 2010.
- Wiener, N. : *Extrapolation, interpolation, and smoothing of stationary time series*. MIT Press, 1949.

- Xu, Y. et Yin, W. : A block coordinate descent method for regularized multiconvex optimization with applications to nonnegative tensor factorization and completion. *SIAM J. Imaging Sci.*, 6(3):1758–1789, 2013.
- Yuan, D., Xu, S. et Zhao, H. : Distributed primal-dual subgradient method for multiagent optimization via consensus algorithms. *IEEE Trans. Systems, Man, and Cybernetics- Part B*, 41(6):1715 ?–1724, Dec. 2011.
- Zangwill, W. I. : *Nonlinear Programming*. Prentice-Hall, Englewood Cliffs, New Jersey, 1969.
- Zheng, J., Saquib, S. S., Sauer, K. et Bouman, C. A. : Parallelizable bayesian tomography algorithms with rapid, guarenteed convergence. *SIAM J. Optim.*, 14(4):1043–1056, 2004.
- Zibulevsky, M. et Pearlmutter, B. A. : Blind source separation by sparse decomposition in a signal dictionary. *Neural Comput.*, 13(4):863–882, Apr. 2001.
- Ziv, A. : Relative distance - an error measure in round-off error analysis. *Mathematics of Computation*, 39(160):563–569, 1982.
- Zubal, I. G., Harrell, C. R., Smith, E. O., Rattner, Z., Gindi, G. et Hoffer, P. B. : Computerized three-dimensional segmented human anatomy. *Med. Phys.*, 21:299–302, 1994.