



THÈSE
PRÉSENTÉE À
L'UNIVERSITÉ DE PAU ET DES PAYS DE L'ADOUR
ÉCOLE DOCTORALE DES SCIENCES EXACTES
ET DE LEURS APPLICATIONS
PAR
WILLIAM MARÉCHAL
POUR OBTENIR LE GRADE DE
DOCTEUR
SPÉCIALITÉ : ÉNERGÉTIQUE

UTILISATION DE MÉTHODES INVERSES POUR
LA CARACTÉRISATION
DE MATÉRIAUX À CHANGEMENT DE PHASE

Soutenue le **24 septembre 2014**

Devant la Commission d'Examen formée de :		
D. Delaunay	Directeur de recherche - CNRS Laboratoire de Thermocinétique de Nantes	Rapporteur
J. P. Dumas	Professeur émérite - Université de Pau et des Pays de l'Adour	
E. Franquet	Maître de Conférences - Université de Pau et des Pays de l'Adour	
S. Gibout	Maître de Conférences - Université de Pau et des Pays de l'Adour	
D. Rousse	Professeur - École de technologie supérieure de Montréal, Université du Québec	Rapporteur
L. Zalewski	Professeur - Université Lille Nord de France	Président

— 2014 —

Remerciements

J'ai effectué ce travail dans le Laboratoire de Thermique, Energétique et Procédé de l'Université de Pau et des Pays de l'Adour avec un financement ANR dans le cadre du projet stock-E 2010, MICMCP.

Je voudrais commencer par remercier mes Directeurs de Thèse Jean Pierre DUMAS et Stéphane GIBOUT pour m'avoir permis de réaliser ce travail dans de bonnes conditions, pour le soutien et les compétences qu'ils m'ont apportés, sans oublier leur bonne humeur. Ils ont d'abord été mes enseignants à l'Ecole d'Ingénieur, m'ont vu évoluer et mon permis d'augmenter mes compétences. Ils m'ont permis de résoudre les difficultés techniques, à me focaliser sur le projet, à ne pas trop m'égarer et à faire des choix judicieusement. Leur disponibilité et leur réactivité constante sur toute la durée m'a été d'un grand secours.

Ma thèse a commencé 8 mois après le début du projet ANR associé, jusqu'à mon arrivée, Erwin FRANQUET a posé les bases de mon travail avec mes directeurs. Je tiens donc à le remercier pour le travail qu'il m'a transmis ainsi que pour ses conseils qu'il m'a apportés pendant ces années aussi bien à l'Ecole qu'au laboratoire.

Je remercie vivement Daniel ROUSSE de l'Ecole de Technologie Supérieure de Montréal (Québec) ainsi que Didier DELAUNAY, Directeur de Recherche CNRS de l'Université de Nantes, pour avoir accepté d'être Rapporteurs de ma thèse.

Je remercie également les partenaires du projet ANR des laboratoires CETHIL à Lyon et LGCgE à Béthune, particulièrement Laurent ZALEWSKI qui accepta de présider ma soutenance et pour l'aide qu'il m'a apportée.

Je souhaite également remercier Jean-Pierre BEDECARRATS et Didier HAILLOT pour l'aide et les conseils qu'ils m'ont apportés.

Je remercie chaleureusement les équipes du LaTEP et de L'ENSGTI dirigée par Pierre CEZAC et Jacques MERCADIER pour m'avoir accueilli dans leur structure où ils m'ont permis de découvrir les mondes de la recherche et de l'enseignement.

Dans ces structures j'ai commencé comme élève ingénieur à l'école pour gravir jusqu'au grade de Docteur grâce à l'enseignement, le travail administratif et technique de l'ensemble du personnel. Il ne faut pas oublier Emilie DEDIEU, Patricia CAPDET, Stéphanie COURTOIS et Jean Michel SORBET pour leur travail administratif et technique.

Merci à mes collègues de bureau, Adrien LOMONACO et Eric PERNOT pour la détente et les rires que l'on a échangés. Ainsi qu'à tous les autres doctorants et stagiaires qui sont passés pendant ces trois années : Camille, Jean-Baptiste, Romain, Lorenzo, Théo, Sacha, Kévin.....

Enfin, je remercie ma famille pour le soutien qu'elle m'a apporté, sans lequel je ne serais peut-être pas arrivé au bout de ce Doctorat.

Sommaire

Sommaire	8
Liste des figures	13
Liste des tableaux	16
Introduction	21
1 Position du problème	25
1.1 Bilan énergétique	25
1.1.1 c_p équivalent	27
1.1.2 Méthode à suivi d'interface	27
1.1.3 Méthodes enthalpiques	28
1.1.3.1 Ecriture en température	28
1.1.3.2 Ecriture en enthalpie	29
1.2 Differential Scanning Calorimetry (DSC)	29
1.2.1 Présentation de l'appareil	29
1.2.2 Mode dynamique	31
1.2.3 Mode Step	33
1.3 Critiques des résultats	33
1.3.1 Critiques de la DSC	33
1.3.2 La méthode du « c_p équivalent »	34
1.3.3 Exemples de l'influence de l'erreur d'identification sur un cas, le mur [46]	37
1.3.4 Nécessité de prendre en compte le gradient interne	41
1.4 Conclusion	42
2 Modèle direct de la DSC	43
2.1 Description du modèle 2D	43
2.1.1 Modèle générique	43
2.1.2 Enthalpie du corps pur	45
2.1.3 Enthalpie de solutions binaires	46
2.1.4 Résolution numérique des modèles thermodynamiques	53
2.2 Validation expérimentale du modèle direct	59
2.3 Analyses paramétriques	62
2.3.1 Cas du corps pur	62
2.3.1.1 Conductivités thermiques	63
2.3.1.2 Moyenne surfacique des coefficients d'échanges équivalents	65
2.3.1.3 Capacités thermiques massiques	66

2.3.1.4	Température de fusion	68
2.3.1.5	Chaleur latente du corps pur	69
2.3.2	Cas de solutions binaires	70
2.3.2.1	Conductivités thermiques	71
2.3.2.2	Moyenne surfacique des coefficients d'échange équivalent	73
2.3.2.3	Capacités thermiques massiques	74
2.3.2.4	Température eutectique	78
2.3.2.5	Variation d'enthalpie massique à la température eutectique : $\tilde{L}_E$	80
2.3.2.6	Température de liquidus	82
2.3.2.7	Température de fusion du solvant	84
2.3.2.8	Chaleur latente de fusion du solvant	86
2.4	Influence de la forme - surface libre de l'échantillon [67]	88
2.4.1	Situation courante	88
2.4.2	Méthode	88
2.4.3	Résultats	89
2.4.4	Synthèse des résultats	91
2.5	Réduction du modèle	91
2.5.1	Incertitude sur la forme de l'échantillon	91
2.5.2	Réduction de modèle pour un corps pur [40]	93
2.5.3	Réduction de modèle pour un binaire [68]	101
2.6	Conclusion	106
3	Principe des méthodes d'inversion pour la calorimétrie	107
3.1	Qu'est-ce-que l'inversion ?	107
3.2	Objectifs des méthodes inverses	108
3.3	Méthodes d'inversions	109
3.3.1	Les méthodes déterministes ou de descentes	109
3.3.1.1	Méthode du morcellement	111
3.3.1.1.a	La recherche monodimensionnelle	112
3.3.1.2	Développements limités d'ordre 1 : utilisation de l'hyperplan tangent	113
3.3.1.3	Développements limités d'ordre 2 : utilisation du paraboloïde tangent	113
3.3.1.3.a	Méthode de Newton	113
3.3.1.3.b	Algorithme de Gauss - Newton	114
3.3.1.4	Conditionnement - Régularisation de Tikhonov	115
3.3.1.5	La méthode de Levenberg-Marquardt [80] [81]	118
3.3.1.6	Méthode mise en oeuvre	119
3.3.2	Les méthodes stochastiques ou algorithmes évolutionnaires	119
3.3.2.1	L'algorithme génétique	120
3.3.3	A mi-chemin : le Simplexe	123
3.3.3.1	Algorithme du simplexe de Nedler-Mead	124
3.3.3.2	Interprétation géométrique des différentes opérations mathématiques	128
3.3.4	Critère d'arrêt de l'algorithme	129
3.4	Conclusion : adaptation des méthodes inverses à la DSC	129

4	Mise en œuvre des méthodes inverses pour la calorimétrie	131
4.1	Etude des coefficients de sensibilité	131
4.1.1	Coefficient de sensibilité réduit relatif à T_E	133
4.1.2	Coefficient de sensibilité réduit relatif à T_M	134
4.1.3	Coefficient de sensibilité réduit relatif à $T_{M,w}$	134
4.1.4	Coefficient de sensibilité réduit relatif à $\widetilde{L}_E$	136
4.1.5	Coefficient de sensibilité réduit relatif à $L_{M,w}$	137
4.1.6	Coefficient de sensibilité réduit relatif à c_L	138
4.1.7	Coefficient de sensibilité réduit relatif à c_S	139
4.1.8	Coefficient de sensibilité réduit relatif à ρ	140
4.1.9	Coefficient de sensibilité réduit relatif à λ_L	141
4.1.10	Coefficient de sensibilité réduit relatif à λ_S	142
4.1.11	Coefficient de sensibilité réduit relatif à α	143
4.1.12	Superposition des coefficients de sensibilité réduites	144
4.2	Comparaison de trois types d'algorithmes d'inversion	145
4.2.1	Algorithme de Levenberg-Marquardt mis en œuvre : (voir 3.3.1.6)	148
4.2.2	Algorithme Génétique : (voir 3.3.2.1)	148
4.2.3	Algorithme du Simplexe : (voir 3.3.3.1)	148
4.3	Applications de l'inversion à des cas théoriques	149
4.3.1	Identification sur des thermogrammes simulés	149
4.3.2	Influence de l'erreur sur l'inversion : incertitudes sur les résultats	152
4.3.2.1	Erreur déterministe	153
4.3.2.2	Erreurs stochastiques	157
4.3.3	Choix des paramètres d'entrée et intervalle de confiance	159
4.3.4	Erreur sur l'enthalpie identifiée : analyse type « Monte Carlo »	160
4.4	Application à des cas expérimentaux	160
4.4.1	Résultats pour un corps pur	161
4.4.2	Résultats pour des solutions binaires	164
4.4.2.1	solution de $H_2O - NH_4Cl$	164
4.4.2.2	Solution de $H_2O - KCl$	170
4.4.3	Concordance des résultats d'une même solution entre eux : détermination des chaleurs eutectiques L_E	172
4.4.4	Conclusion : incertitude sur l'enthalpie fonction de la température	174
5	Application à un mortier contenant des MCP	175
5.1	Dispositif expérimental du LGCgE [95] [96]	175
5.2	Résultats expérimentaux du LGCE [101]	178
5.3	Modèle direct	180
5.3.1	Géométrie, maillage et conditions aux limites	180
5.3.2	Equation d'état	181
5.4	Caractérisation par méthode inverse	182
5.4.1	Etude de sensibilité	183
5.4.2	Etude de l'erreur	189
5.4.2.1	Erreur déterministe	189
5.4.2.2	Erreur stochastique	190
5.4.2.3	Intervalle de confiance	191
5.5	Résultat d'inversion	191

5.6	Comparaison de la caractérisation par méthode inverse et des caractérisations directes	194
5.7	Conclusion	196
	Conclusion générale	197
A	Recherche de l'optimum géométrique pour la réduction de modèle	207
A.1	Cas du rapport minimal des surfaces 1D et 2D	207
A.2	Conclusion	209

Table des figures

1	Stockage du froid [8]	22
1.1	Volume de contrôle sur lequel s'effectue le bilan	25
1.2	$c_{p,eq}$ pour un échantillon d'eau de 12,52 mg et $5 K \cdot min^{-1}$	27
1.3	Illustration de la résolution par méthode enthalpique, écrite en enthalpie	29
1.4	Schéma global d'un calorimètre	30
1.5	Schéma de la tête du calorimètre	31
1.6	Exemples de thermogrammes expérimentaux dans le cas de l'eau, $m=12,52 mg$	32
1.7	Approximation du c_p équivalent pour un échantillon d'eau de 8,5 mg à $2 K \cdot min^{-1}$	34
1.8	Thermogramme théorique & obtenu par la méthode du c_p équivalent [46]	35
1.9	Influence de la vitesse et de la masse sur le $c_{p,eq}$ [36] [48]	35
1.10	Influence de la vitesse et de la masse sur DTF [46]	36
1.11	Enthalpies et thermogramme pour le cas d'une solution binaire de H_2O-NH_4Cl à 9,02%	36
1.12	Exemple d'inclusion de MCP dans un mur.	37
1.13	Coupe transversale du mur d'épaisseur 200 mm. [46] [49]	38
1.14	Réponse du mur à un réchauffement exponentiel de l'extérieur. [46]	39
1.15	Réponse du mur à une évolution sinusoïdale de la température extérieure. [46]	40
1.16	Thermogramme théorique suivant l'hypothèse température homogène (b) et réel (a)	41
1.17	Evolution du champ de température interne à l'échantillon selon plusieurs rayons et une hauteur correspondant au centre thermique : cas de 11,37 mg d'eau pour $\beta = 5 K \cdot min^{-1}$	41
2.1	Conduction autour de la cellule	44
2.2	Modèle de l'échantillon.	45
2.3	Enthalpie d'un corps pur, l'eau	46
2.4	Diagramme de phase d'une solution binaire avec un solidus vertical et un liquidus rectiligne.	47
2.5	Enthalpie d'une solution binaire ($H_2O - NH_4Cl$), $x_0 = 9,02\%$.	53
2.6	Maillage par symétrie de révolution	54
2.7	Exemples de différents maillages	55
2.8	Illustration du calcul des flux internes (cas général)	57
2.9	Illustration du calcul des flux sur la face latérale (maillage curviligne)	57
2.10	Schéma de résolution du modèle direct	58
2.11	Validation expérimentale du modèle 2D dans le cas du corps pur.	59
2.12	Validation expérimentale du modèle 2D relatif aux solutions de $H_2O - KCl$.	60

2.13 Validation expérimentale du modèle 2D relatif aux solutions de $H_2O - NH_4Cl$.	61
2.14 Influence de la conductivité thermique solide sur le thermogramme	63
2.15 Influence de la conductivité thermique liquide sur le thermogramme	64
2.16 Influence de la moyenne surfacique des coefficients d'échanges équivalents sur le thermogramme	65
2.17 Influence des capacités thermiques spécifiques solides sur le comportement des corps purs.	66
2.18 Influence des capacités thermiques spécifiques liquides sur le comportement des corps purs.	67
2.19 Influence de la température de fusion pour un corps pur.	68
2.20 Influence de la chaleur latente sur le comportement du corps pur.	69
2.21 Influence de la conductivité thermique solide sur le comportement d'une solution de $H_2O - NH_4Cl$	71
2.22 Influence de la conductivité thermique liquide sur le comportement d'une solution de $H_2O - NH_4Cl$	72
2.23 Influence de la moyenne des coefficients d'échanges équivalents sur le comportement d'une solution de $H_2O - NH_4Cl$	73
2.24 Influence de c_S , solution $H_2O - NH_4Cl$	74
2.25 Influence des capacités thermiques spécifiques solides sur le comportement d'une solution de $H_2O - NH_4Cl$	75
2.26 Influence de c_L , solution $H_2O - NH_4Cl$	76
2.27 Influence des capacités thermiques spécifiques liquides sur le comportement d'une solution de $H_2O - NH_4Cl$	77
2.28 Influence de T_F , solution $H_2O - NH_4Cl$	78
2.29 Influence de la température eutectique sur le comportement d'une solution de $H_2O - NH_4Cl$	79
2.30 Influence de $\tilde{L}_E$, solution $H_2O - NH_4Cl$	80
2.31 Influence de la chaleur latente eutectique sur le comportement d'une solution de $H_2O - NH_4Cl$	81
2.32 Influence de T_M , solution $H_2O - NH_4Cl$	82
2.33 Influence de la température de liquidus sur le comportement d'une solution de $H_2O - NH_4Cl$	83
2.34 Influence de $T_{M,w}$, solution $H_2O - NH_4Cl$	84
2.35 Influence de la température de liquidus sur le comportement d'une solution de $H_2O - NH_4Cl$	85
2.36 Influence de $L_{M,w}$, solution $H_2O - NH_4Cl$	86
2.37 Influence de la chaleur latente de fusion sur le comportement d'une solution de $H_2O - NH_4Cl$	87
2.38 Schéma de la surface libre (coubure exagéré)	88
2.39 Modification du maillage pour une surface libre convexe	89
2.40 Simulations de différentes courbures de surface	90
2.41 Déformation du thermogramme selon la courbure de la surface libre, $5 K \cdot min^{-1}$	91
2.42 Géométrie possible de deux corps, l'eau et le mercure. On tâche généralement de remplir une cellule de façon plus correcte.	92
2.43 Géométrie sphérique	93
2.44 Maillage sphérique	94
2.45 Comparaison 1D/2D pour l'eau	95

2.46	Comparaison 1D/2D selon β en adaptant α , λ_S et λ_L	96
2.47	Comparaison des modèles 2D et 1D pour différentes configurations d'un échantillon d'hexadécane, cas de coefficients d'échange non-uniforme à 10 % [40] . .	97
2.48	Comparaison des modèles 2D et 1D pour différentes configurations d'un échantillon d'eau, cas de coefficients d'échange non-uniforme à 10 % [40]	98
2.49	Comparaison 1D et 2D pour $x_0 = 0,57\%$ $\beta = 5 \text{ K} \cdot \text{min}^{-1}$	102
2.50	Comparaison des modèles 1D et 2D pour $x_0 = 2,5\%$ $\beta = 5 \text{ K} \cdot \text{min}^{-1}$	103
2.51	Comparaison des modèles 1D et 2D pour $x_0 = 5\%$ $\beta = 5 \text{ K} \cdot \text{min}^{-1}$	104
2.52	Comparaison des modèles 1D et 2D pour $x_0 = 10\%$ $\beta = 5 \text{ K} \cdot \text{min}^{-1}$	105
3.1	Schéma d'un modèle direct	108
3.2	Schéma méthodologique de l'inversion	109
3.3	Gradient et direction de descente [77]	111
3.4	Méthode de morcellement dans un cas à deux paramètres [75]	112
3.5	Méthode de la section dorée ($F_{obj}(p_1) > F_{obj}(p_2)$) [77]	113
3.6	Méthode de Newton [75]	114
3.7	Levenberg-Marquardt : pondération entre plus grande pente et Gauss-Newton	118
3.8	Exemple d'un croisement BLX	121
3.9	Croisement « BLX- α volumique » [77]	122
3.10	Cycle générationnel	123
3.11	Historique de l'évolution du simplexe de Nedler-Mead pour un problème à deux paramètres [86]	124
3.12	Illustration de l'opération d'expansion [86]	125
3.13	Illustration de l'opération de réflexion [86]	125
3.14	Illustration de l'opération de contraction externe [86]	126
3.15	Illustration de l'opération de réduction faute de contraction externe [86]	126
3.16	Illustration de l'opération de contraction interne [86]	127
3.17	Illustration de l'opération de réduction faute de contraction interne [86]	127
3.18	Algorigramme du simplexe de Nedler-Mead	128
3.19	Schéma de l'inversion appliqué à la D.S.C.	130
4.1	Coefficient de sensibilité réduit relatif à T_E	133
4.2	Coefficient de sensibilité réduit relatif à T_M	134
4.3	Coefficient de sensibilité réduit relatif à $T_{M,w}$	135
4.4	Coefficient de sensibilité réduit relatif à $\tilde{L}_E$	136
4.5	Coefficient de sensibilité réduit relatif à $L_{M,w}$	137
4.6	Coefficient de sensibilité réduit relatif à c_L	138
4.7	Coefficient de sensibilité réduit relatif à c_S	139
4.8	Coefficient de sensibilité réduit relatif à ρ	140
4.9	Coefficient de sensibilité réduit relatif à λ_L	141
4.10	Coefficient de sensibilité réduit relatif à λ_S	142
4.11	Coefficient de sensibilité réduit relatif à α	143
4.12	Superposition des coefficients de sensibilité réduites	144
4.13	Comparaison de la courbe de sensibilité à ρ à la somme de celles à c_S , c_L , $L_{M,w}$ et $\tilde{L}_E$	145
4.14	Evolution de la fonction objectif au cours de l'identification globale.	146

4.15 Evolution de la population de l'algorithme génétique pour la résolution de la fonction de Rosenbrock. [77]	147
4.16 Identifications sur des thermogrammes simulés de corps purs	150
4.17 Identifications sur des thermogrammes simulés de solution binaire H_2O-NH_4Cl	151
4.18 Bruit expérimental du calorimètre	158
4.19 Comparaison thermogramme expérimental/numérique et enthalpie reconstitués de l'eau	161
4.20 Comparaison thermogramme expérimental/numérique et enthalpie reconstituée de l'acide benzoïque	162
4.21 Détermination par mesures directes	163
4.22 $x_0 = 0,480\%$ et $\beta = 5 K \cdot min^{-1}$	164
4.23 $x_0 = 5,00\%$ et $\beta = 5 K \cdot min^{-1}$	165
4.24 $x_0 = 9,02\%$ et $\beta = 5 K \cdot min^{-1}$	165
4.25 $x_0 = 10,00\%$ et $\beta = 5 K \cdot min^{-1}$	166
4.26 Comparaison des thermogrammes expérimentaux et numériques après identification (solution de $H_2O - NH_4Cl$) pour $\beta = 2 K \cdot min^{-1}$.	167
4.27 Thermogrammes et enthalpies identifiés de solutions de $H_2O - NH_4Cl$ pour deux vitesses	169
4.28 T_{onset} de la solution binaire	170
4.29 Comparaison des thermogrammes expérimentaux et numériques après identification (solution de $H_2O - KCl$)	171
4.30 Détermination des chaleurs latentes eutectiques L_E	173
4.31 Comparaison des paramètres identifiés avec leur valeurs exactes sur un diagramme de phase	174
5.1 Maillage 3D	176
5.2 Modélisation 3D de la conduction dans l'échantillon refroidi de manière symétrique - coupe interne	176
5.3 Représentation globale du dispositif [101]	177
5.4 Dimensions et instrumentation du dispositif [101]	178
5.5 Rampes de 3h - $7,73 K \cdot h^{-1}$	179
5.6 Rampes de 4h - $5,8 K \cdot h^{-1}$	179
5.7 Rampes de 6h - $3,87 K \cdot h^{-1}$	180
5.8 Modélisation thermique du mortier	181
5.9 Coefficient de sensibilité réduit relatif à T_M	184
5.10 Coefficient de sensibilité réduit relatif à T_{MCP}	184
5.11 Coefficient de sensibilité réduit relatif à c_S	185
5.12 Coefficient de sensibilité réduit relatif à c_L	185
5.13 Coefficient de sensibilité réduit relatif à $\tilde{L}_{MCP}$	186
5.14 Coefficient de sensibilité réduit relatif à ρ	186
5.15 Coefficient de sensibilité réduit relatif à α	187
5.16 Coefficient de sensibilité réduit relatif à λ	187
5.17 Superposition des coefficients de sensibilité réduits	188
5.18 Comparaison de la courbe de sensibilité au paramètre ρ à la somme de celles aux paramètres c_S , c_L et $\tilde{L}_{M,w}$	188
5.19 Bruit expérimental sur une courbe de flux du mortier, figure 5.7	190
5.20 Rampes de 3h - $7,73 K \cdot h^{-1}$	192

5.21 Rampes de 4h - $5,8 K \cdot h^{-1}$	192
5.22 Rampes de 6h - $3,87 K \cdot h^{-1}$	193
5.23 Comparaison des enthalpies identifiées aux vitesses de 10,3-7,8-5,2 $K \cdot h^{-1}$. . .	193
5.24 Caractérisation de L_{MCP} par calorimétrie DSC à $5 K \cdot min^{-1}$ [101]	195
5.25 Caractérisation de T_M par calorimétrie DSC [101]	196

Liste des tableaux

1.1 Paramètres des simulations	37
2.1 Paramètres des corps purs [63] [64]	62
2.2 Propriétés thermodynamiques relatives aux solutions de $H_2O - NH_4Cl$	70
2.3 Influence de la réduction sur le temps de résolution du modèle de l'eau pour $\beta = 5 K \cdot min^{-1}$, $m = 10,3$ mg et un pas de temps constant de 10^{-4} s	92
2.4 Conductivité ajustée entre les modèles 2D et 1D pour un même échantillon mais pour plusieurs vitesses, $r^{2D} = 1$ mm $z^{2D} = 1.56$ mm	95
2.5 Géométrie de différents échantillons pour les simulations	96
2.6 Conductivité ajustée entre les modèles 2D et 1D (cas du coefficient d'échange uniforme). [40]	99
2.7 Conductivité ajustée entre les modèles 2D et 1D (cas du coefficient d'échange non-uniforme, 1000 et 100 $W \cdot m^{-2} \cdot K^{-1}$).	99
2.8 Conductivités identifiées entre les modèles 2D et 1D (cas du coefficient d'échange non-uniforme, 1000 et 600 $W \cdot m^{-2} \cdot K^{-1}$).	100
2.9 Conductivité ajustée entre les modèles binaires 2D et 1D (cas du coefficient d'échange non-uniforme).	101
4.1 Propriétés de corps purs [63] [64]	149
4.2 Propriétés de solutions binaires. [52]	149
4.3 Propriétés énergétiques des corps purs identifiées	150
4.4 Propriétés énergétiques des solutions binaires identifiées	150
4.5 Influence théorique et expérimentale d'une erreur déterministe de 1% ou 0,1 K sur le résultat de l'inversion mise à part ρ où l'erreur est de 3% - cas d'un corps pur	155
4.6 Influence théorique et expérimentale d'une erreur déterministe de 1% ou 0,1 K sur le résultat de l'inversion mise à part ρ où l'erreur est de 3% - cas d'un binaire	156
4.7 Influence théorique d'une erreur stochastique de 5 μW sur le résultat de l'inversion	159
4.8 Intervalle de confiance sur les paramètres identifiés	160
4.9 Paramètres identifiés sur des thermogrammes expérimentales de corps pur et valeurs déterminées par mesures directes. Valeurs de la littérature [63] [91]	162
4.10 Paramètres identifiés sur des thermogrammes expérimentaux du binaire $H_2O - NH_4Cl$	168
4.11 Paramètres identifiés sur des thermogrammes expérimentaux de binaire $H_2O - NH_4Cl$	172

4.12	Températures corrigées pour le binaire $H_2O - KCl$	172
5.1	Ratio massique de mélange du mortier [96] [101]	178
5.2	Influence théorique d'une erreur déterministe de 3.3% sur ρ et de 1% ou 0.1 K sur les autres paramètres sur le résultat de l'inversion - cas du mortier	190
5.3	Influence théorique d'une erreur stochastique de $0.395 W \cdot m^{-2}$ sur le résultat de l'inversion - cas du mortier	191
5.4	Intervalle de confiance des paramètres identifiés - cas du mortier	191
5.5	Paramètres identifiés sur les courbes de flux expérimentales du mortier à 12,4 % de MCP	194
5.6	Propriétés de l'échantillon à 12.4% de MCP qui ont été déterminé avec le dispositif [101]	195
A.1	Rapport des conductivités pour une géométrie minimale.	208

Nomenclature

Abréviation

MCP	matériaux à changement de phase
DSC	Differential Scanning Calorimetry
GIEC	groupe d'experts intergouvernemental sur l'évolution du climat

Préfixe

d	élément différentiel
δ	petite variation
Δ	différence entre deux états

Opérateur

$\vec{\nabla}()$	gradient
$\vec{\nabla} \cdot$	divergence
∇^2	Laplacien
$\ \cdot\ $	norme
t	transposée

Lettres latines

c	capacité thermique massique	$[J \cdot K^{-1} \cdot kg^{-1}]$
$\vec{d}$	vecteur distance entre deux points	$[m]$
d_1	distance entre le centre 1 et l'interface	$[m]$
d_2	distance entre le centre 2 et l'interface	$[m]$
DTF	largeur du pic de fusion	$[K]$
$E_{n,j}, E_{m,j}$	solution enfant de l'algorithme génétique	
$E_{sensible}$	énergie stockée sous forme de chaleur sensible	$[J]$
$E_{sensible+latente}$	énergie stockée sous forme de chaleur sensible et latente	$[J]$
e	épaisseur	$[m]$
F_{obj}	fonction objectif	
H	enthalpie	$[J]$
h	enthalpie massique	$[J \cdot kg^{-1}]$
L	longueur	$[m]$
L_E	chaleur latente de fusion à l'eutectique	$[J \cdot kg^{-1}]$
L_M	chaleur latente de fusion massique	$[J \cdot kg^{-1}]$
L_{MCP}	chaleur latente du MCP	$[J \cdot kg^{-1}]$

m	masse	[kg]
$\vec{n}$	vecteur normal à une surface	[–]
NI	nombre de maille selon la hauteur	[–]
NJ	nombre de maille selon le rayon	[–]
n_m	nombre de points de la sortie du modèle direct (le thermogramme)	[–]
n_p	nombre de paramètres du modèle direct	[–]
$P_{i,j}$	solution parent de l'algorithme génétique	
$\bar{P}$	barycentre du simplexe	
P_r	point obtenu par réflexion	
P_e	point obtenu par expansion	
P_{ce}	point obtenu par contraction externe	
P_{ci}	point obtenu par contraction interne	
P_{n_p+1}	plus mauvais sommet du simplexe selon F_{obj}	
P_{n_p}	second plus mauvais sommet du simplexe selon F_{obj}	
p	paramètre du modèle	
P	pression	[Pa]
Q	quantité de chaleur	[J]
R_{th}	résistance thermique	[K · m ² · W ^{−1}]
R	rayon de l'échantillon	[m]
r	rayon courant	[m]
S	surface	[m ²]
T	température	[°C ou K]
T_F	température de fusion d'un corps pur	[°C ou K]
T_M	température de liquidus	[°C ou K]
T_P	température de plateau	[°C ou K]
$\bar{T}$	température moyenne au cours d'une journée	[°C]
t	temps	[s]
U	énergie interne	[J]
V	volume	[m ³]
Var	variance de la population	
x_0	fraction massique de soluté	[–]
x_E	fraction massique eutectique	[–]
$x_{a,L}$	fraction massique du soluté dans la phase liquide	[–]
x	position	[m]
Z	hauteur de l'échantillon	[m]
z	hauteur courante	[m]

Matrice et vecteur

$\mathbb{A}$	un modèle mathématique	$n_m \times n_p$
$\mathbf{d}$	vecteur direction de descente	$n_p \times 1$
$\mathbf{D}$	matrice diagonale	$n_p \times n_p$
$\mathbf{e}_p$	vecteur de base	$1 \times n_p$
$\mathbb{I}$	matrice identité	$n_p \times n_p$
$\mathbb{M}$	modèle	
$\mathbf{p}$	vecteur colonne qui contient les paramètres	$n_p \times 1$
$p\mathbb{S}$	matrice de sensibilité réduite	$n_m \times n_p$
$\mathbb{S}$	matrice de sensibilité	$n_m \times n_p$
$\mathbb{S}_\pi$	matrice de sensibilité au paramètre π	$n_m \times 1$
$\mathbb{W}$	matrice de pondération	$n_m \times n_m$
$\mathbf{X}$	matrice des soliciations	
$\mathbf{y}$	vecteur colonne sortie du modèle	$n_m \times 1$
$\mathbf{Y}$	la sortie du modèle $\mathbb{A}$	$n_m \times 1$
$\bar{\bar{\lambda}}$	tenseur de conduction thermique	$3 \times 3, [W \cdot m^{-1} \cdot K^{-1}]$

Lettres Grecques

α	coefficient d'échange équivalent	$[W \cdot m^{-2} K^{-1}]$
α_{BLX}	coefficient de croisement BLX de l'algorithme génétique	$[-]$
β	vitesse de réchauffement	$[K \cdot s^{-1}]$
γ_{NM}	coefficient de contraction du simplexe de Nedler-Mead	$[-]$
$\delta_{courbure}$	paramètre de courbure de la surface libre	$[-]$
δ	terme de mélange de l'algorithme de Levenberg-Marquardt	$[-]$
$\varepsilon, \varepsilon_1, \varepsilon_2$	réels positifs de très faible valeur	
λ	conductivité thermique	$[W \cdot m^{-1} \cdot K^{-1}]$
π	paramètre d'entrée connue du modèle	
ρ	masse volumique	$[kg \cdot m^{-3}]$
ρ_{NM}	coefficient de réflexion du simplexe de Nedler-Mead	$[-]$
χ_L	taux liquide	$[-]$
$\chi_{w,S}$	taux en solvant dans la phase solide	$[-]$
χ_{MCP}	taux massique de MCP dans le mortier	$[-]$
χ_{NM}	coefficient de expansion du simplexe de Nedler-Mead	$[-]$
σ	écart-type	
σ_{NM}	coefficient de réduction du simplexe de Nedler-Mead	$[-]$
τ	durée	$[s]$
ϕ	flux d'énergie	$[W]$
φ	flux surfacique	$[W \cdot m^{-2}]$
χ	taux massique du solvant	$[-]$
Ω	domaine d'intégration	
ω	distance à parcourir dans la direction de descente	$[-]$

Indice

a	soluté
b	bas
$BEST$	la meilleure solution courante
deb	début
D	droit
e	externe
ext	extérieur
E	eutectique
exp	expérience
f	front de fusion
G	gauche
h	haut
i, j	indices numériques
k	numéro de surface de la maille (i,j)
L	liquide
min	correspondant à un minimum
M	fusion ou liquidus
MCP	matériaux à changement de phase
p, eq	c_p équivalent
ref	référence
S	solide
T	grandeur totale sur tout l'échantillon
w	solvant, corps pur, souvent de l'eau
0	valeur initiale
∞	régime final établi

Exposant

—	propriété molaire partielle
★	corps pur
dil	dilution
dis	dissolution
~	pondéré par le taux massique
$1D$	relatif au modèle 1D
$2D$	relatif au modèle 2D
k	itération
T	grandeur totale sur tout l'échantillon

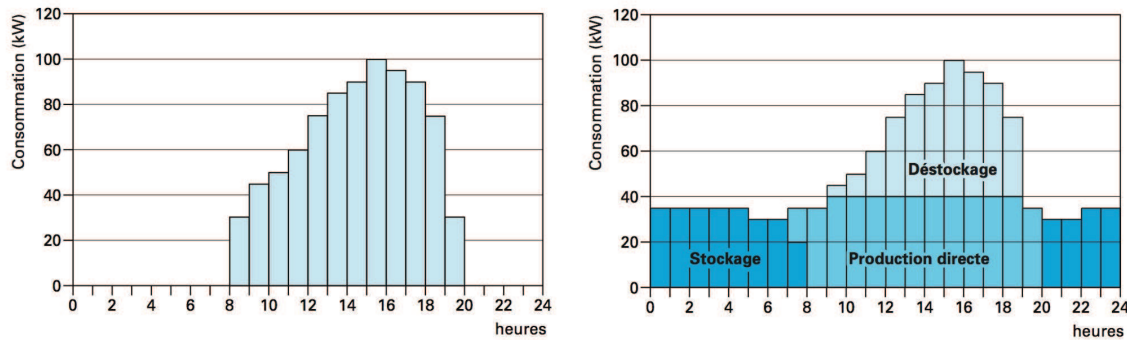
Introduction

La consommation totale d'énergie a été multipliée par 7 depuis 1950, le prix du baril de brut s'est envolé, atteignant plus de 110 \$ en avril 2008, 80 \$ en 2009 et 140 \$ en 2012. Les émissions de carbone issues de combustibles fossiles ont été multipliées par plus de 5 depuis 1950, la concentration de CO_2 dans l'atmosphère a dépassé 380 ppm en 2005. En 2050, les scénarios les plus optimistes prévoient une stabilisation aux alentours de 450 ppm. Le groupe d'experts intergouvernemental sur l'évolution du climat (GIEC) estime pour 2100 le réchauffement climatique à plusieurs degrés sachant qu'en France la température moyenne a augmenté de 1 K sur les quinze dernières années [1] [2] [3].

Pour limiter l'élévation moyenne de la température de 2 à 3 K, il faudrait diviser par 2 avant 2050 les émissions mondiales. Comme les émissions des pays en développement vont continuer de croître, les émissions des pays industrialisés devront être divisées par 4 avant cette échéance. Il s'agit du fameux « facteur 4 » en ligne de mire de notre politique énergétique nationale dans la loi Grenelle 2 [1] : il est devenu nécessaire de revoir la gestion de notre consommation énergétique.

Le secteur du bâtiment consomme près de 44,5% [4] de l'énergie finale française et génère environ 24% [1] des émissions de CO_2 du pays. Les actions visant à réduire la consommation énergétique dans ce secteur seraient les plus efficaces, pour atteindre les objectifs de la France aussi bien en terme environnemental qu'en terme énergétique. Avec ce constat, les lois relatives au Grenelle de l'Environnement 1, adoptées en deuxième lecture par le Parlement le 23 juillet 2009, ont fixé des objectifs ambitieux, notamment qu'à l'horizon 2020 une grande partie du parc résidentiel soit renouvelée pour permettre une baisse de la consommation énergétique de 40%, et en parallèle de diminuer la consommation des bâtiments tertiaires de 38%. Dans un rapport sur les Technologies Clés de 2005 [5], les pouvoirs publics encouragent à porter des efforts sur le domaine du stockage de l'énergie.

L'électricité, la forme d'énergie la plus pratique, ne peut être stockée. Il existe de nombreuses techniques pour transformer l'électricité en une autre forme qui, elle, peut être stockée [6]. Pour le domaine du bâtiment, on diffère son utilisation pour profiter des effets du climat. Cette problématique est depuis longtemps étudiée pour le stockage du froid [7] [8] [9]. Du fait des variations journalières du climat, les besoins et la disponibilité de l'énergie thermique n'est pas la même au cours de la journée. Il en résulte des pics de consommation au moment où l'énergie est chère. En France, où il existe un parc nucléaire important dont la production est peu modulable, un tarif de nuit a été instauré afin d'inciter à décaler la consommation la nuit, d'étaler la consommation du pays. Le stockage permet cet étalement car il consiste à emmagasiner l'énergie pendant qu'elle est peu coûteuse et abondante, pour l'utiliser quand elle est moins disponible et plus chère (figure 1).



(a) Histogramme de consommation d'un système de climatisation

(b) Histogramme avec stockage et déstockage

FIGURE 1 – Stockage du froid [8]

Pour le bâtiment, en France, culturellement il s'agit plus d'un problème de chauffage en hiver que de climatisation en été, mais la problématique reste la même. L'énergie thermique du soleil est disponible le jour et pas la nuit, ce que l'on compense par des systèmes de chauffage.

Dans le secteur de l'habitat, ce qui prime c'est le confort thermique. Ainsi, on ne peut envisager une amélioration des performances énergétiques si on dégrade le confort thermique. Il faudrait au contraire l'améliorer. Il ne faut pas que jouer sur l'aspect quantitatif de la thermique du bâtiment, mais sur l'aspect qualitatif. C'est-à-dire qu'installer un lourd équipement pour réchauffer ou refroidir le bâtiment n'est pas utile pour l'occupant si l'énergie thermique est mal répartie. Il vaut mieux chercher une gestion intelligente des transferts, des pertes thermiques.

Le secteur du bâtiment a des effets de premier plan sur la société, par le biais du confort thermique. Il en découle donc une importance du facteur inertie des bâtiments qui lissera les variations de température, diminuant les consommations de chauffage et de climatisation, et donc nécessitant l'installation de moindres appareils, tout en améliorant le confort pendant les saisons extrêmes.

Le confort thermique d'un bâtiment a deux principaux indicateurs : le niveau de température et sa régularité. La stabilité du champ de température dans un bâtiment dépend majoritairement de l'inertie thermique de ce même bâtiment. Il faut ainsi veiller à l'augmenter et non la réduire. Ceci est possible en mettant en place des systèmes de stockage thermique, afin de lisser les variations journalières et/ou saisonnières. Une solution intéressante est l'emploi de matériaux à changement de phase (MCP) [10]. Ces matériaux présentent le double avantage d'avoir une forte densité de stockage de par leur chaleur latente et de pouvoir fixer le niveau de température par le changement d'état, si la chaleur latente est suffisante.

Pour illustrer ce point, les équations (1), (2) et (3) montrent que pour stocker une même quantité d'énergie par une même masse d'élément, un matériau qui ne fait que du stockage

sensible devrait s'élever d'avantage en température, qu'un MCP.

$$E_{sensible} = \int_{T_1}^{T_2} mcdT \quad (1)$$

$$E_{sensible+latent} = \int_{T_3}^{T_4} mc_S dT + mL_M + \int_{T_5}^{T_6} mc_L dT \quad (2)$$

$$E_{sensible} = E_{sensible+latent} \quad (=) \quad T_2 - T_1 \geq T_6 - T_3 \quad (3)$$

Avec c , c_S , c_L les capacités thermiques massiques du matériau et du MCP solide et liquide en $J \cdot kg^{-1} \cdot K^{-1}$, L_M la chaleur latente du MCP en $J \cdot kg^{-1}$ et m la masse d'élément en kg .

Jusqu'à présent la méthode utilisée pour augmenter l'inertie thermique d'un bâtiment était soit d'augmenter l'épaisseur des murs, soit d'en changer la nature afin de jouer sur sa capacité de stockage sensible. Le problème avec cette approche est une perte significative de place, d'espace à vivre pour l'habitant ou l'utilisateur. Sur ce point, procéder en donnant de la capacité de stockage latent aux murs, en y adjoignant des MCP, augmentera l'inertie thermique tout ne diminuant pas l'espace disponible. Ce qui serait socialement plus acceptable [11].

Augmenter l'inertie thermique du bâtiment présente un autre avantage. En effet, le matériau va permettre d'étaler les effets du climat, les reporter, permettant ainsi aux réchauffements et refroidissements journaliers de se compenser en partie. C'est d'ailleurs tout l'intérêt du stockage thermique dans le bâtiment.

Ainsi donc, avec un matériau adapté, on peut améliorer le confort thermique tout en diminuant la pression énergétique par un effet de lissage de la consommation.

Il est nécessaire de pouvoir modéliser le comportement énergétique du bâtiment qui intégrera ces MCP, afin de pouvoir choisir judicieusement le MCP. Une telle modélisation, pour être une représentation fidèle des phénomènes physiques, passe par la connaissance fine du processus de changement de phase. Ce qui induit de connaître les caractéristiques thermophysiques des matériaux [12].

La commercialisation de cette technologie a déjà commencé bien qu'elle soit freinée par le manque de procédure d'évaluation, de caractérisation et de certification de ce type de matériau. Il est nécessaire de mettre au point une méthode de caractérisation fiable, qui débouchera sur une méthode d'évaluation correcte de ces produits. Ce qui *in-fine* permettra une meilleure industrialisation et diffusion sur le marché, des matériaux à changement de phase.

L'industrie s'est déjà emparée de cette technologie, comme BASF qui propose une plaque de plâtre dans laquelle on a dispersé des micro-capsules de paraffine et DUPONT DE NEMOURS qui propose une plaque composite de MCP. Même si ces deux industriels présentent des secteurs industriels différents, leurs objectifs convergent : promouvoir ces MCP pour l'amélioration de l'inertie thermique des bâtiments. L'intérêt est économique, industriel mais aussi environnemental. De nombreux pays se sont fixés un objectif de réduction des gaz à

effet de serre. Ce projet d'augmentation de l'inertie thermique des bâtiments, pourrait participer à la réduction de l'émission de ces gaz par le fait qu'il permettrait de réduire, voire de supprimer le recours à la climatisation en période estivale et de diminuer la consommation énergétique des systèmes de chauffage l'hiver ou de climatisation l'été (figure 1).

Pour résumer, la solution retenue réside dans le couplage entre une forte intégration des énergies renouvelables au bâti et la réduction significative des consommations d'énergie. La trajectoire vers la haute efficacité énergétique passe, pour une large part, par le développement d'enveloppes très performantes pour le bâtiment (récupération des énergies « gratuites » et leur stockage et restitution en temps voulu). Les matériaux à changement de phase participeront donc à la réduction des besoins en énergie du bâtiment. Afin d'optimiser leur choix, nous montrerons qu'il est nécessaire de déterminer correctement leurs caractéristiques énergétiques. Dans ce travail nous proposons de façon originale l'utilisation de méthodes inverses pour compléter les mesures directes, afin d'aboutir à une meilleure caractérisation des MCP. Notre objectif est d'identifier la courbe d'évolution de l'enthalpie en fonction de la température, indépendamment des conditions opératoires.

De plus, la technologie MCP pouvant s'appliquer à de nombreux types de stockage d'énergie thermique [13] [14] tels que :

- le système de traitement d'air de ventilation ;
- le stockage de l'eau chaude sanitaire par ballon à MCP ;
- le refroidissement des panneaux solaires ;
- le stockage pour pompes à chaleur ;
- le stockage du froid ;
- refroidissement des composants ;
- etc ;

nous pouvons espérer que cette nouvelle méthode par caractérisation inverse puisse être universelle.

Ce travail a été réalisé dans le cadre du projet ANR stock-E 2010 MICMCP pour « méthode inverse de caractérisation des matériaux à changement de phase ». Il comprend cinq chapitres. Dans le chapitre un on mettra en avant le défaut actuel de caractérisation. Les chapitres deux et trois présenteront les outils employés, soit un modèle direct et les méthodes inverses. Les chapitres quatre et cinq présenteront deux cas d'applications.

Chapitre 1

Position du problème

Notre sujet s'inscrit dans le cadre général du stockage thermique par utilisant des matériaux changement de phase (MCP). Depuis des décennies on caractérise au mieux ces matériaux, mais on n'est pas satisfait de la fiabilité des résultats [15] [16].

Ce projet est donc de développer ou d'améliorer une des méthodes choisies de caractérisation des propriétés énergétiques des matériaux à changement de phase, la calorimétrie, par l'utilisation de méthodes inverses que nous aborderons dans le chapitre 3. Pour ce faire il nous faut d'abord établir une stratégie pour résoudre le bilan énergétique [17]. De cette stratégie va découler des erreurs de modélisation plus ou moins importantes.

1.1 Bilan énergétique

Pour ce problème on considère un système immobile dans le référentiel du laboratoire que l'on suppose galiléen. De ce système de pression P , on s'intéresse à une sous-partie de volume de contrôle V quelconque limité par une surface fermée Ω , figure 1.1. Nous traiterons le cas d'un milieu passif sans échange de travail dû à l'électricité ou aux mouvements, nous négligeons donc la convection interne. Dans ce cadre, le travail élémentaire δW se réduit à celui des forces de pression, $-PdV$. Par ailleurs, l'enthalpie est définie par rapport à l'énergie interne U par la relation $H = U + PV$.

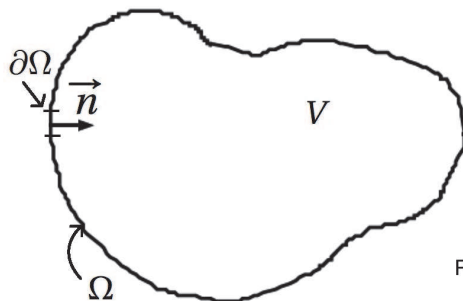


FIGURE 1.1 – Volume de contrôle sur lequel s'effectue le bilan

En notant δQ la quantité de chaleur échangée, le premier principe s'écrit donc pour ce sys-

tème :

$$dU = \delta Q - PdV \quad (1.1)$$

soit pour l'enthalpie :

$$dH = \delta Q + VdP \quad (1.2)$$

L'environnement usuel d'un bâtiment ayant une atmosphère que l'on considère isobare, le terme VdP est nul, ainsi dH se réduit à δQ . L'enveloppe d'un bâtiment étant soumise aux fluctuations temporelles du climat, il convient de prendre en compte la dimension temporelle pour étudier cette évolution dynamique. C'est pourquoi nous allons poursuivre cette réflexion en passant à un bilan de puissance.

En considérant ici que le mode de transfert thermique prépondérant est la conduction, la quantité de chaleur totale échangée en 1 s à travers la surface Ω peut s'écrire en prenant en compte la loi de Fourier :

$$\phi = \oint_{\Omega} \vec{\varphi} \cdot \vec{n} dS = \oint_{\Omega} -\vec{\lambda} \cdot \vec{\nabla}(T) \cdot \vec{n} dS \quad (1.3)$$

Avec $\vec{n}$ le vecteur normal rentrant de la surface et $\vec{\lambda}$ le tenseur de conduction. Nous allons restreindre notre étude à des matériaux isotropes. Dans ce cas, le tenseur se réduit à un simple scalaire λ en $W \cdot m^{-1} \cdot K^{-1}$.

Ainsi le bilan de puissance s'écrit :

$$\iiint_V \rho \frac{\partial h(T)}{\partial t} dV = \oint_{\Omega} \vec{\varphi} \cdot \vec{n} dS \quad (1.4)$$

Grâce au théorème d'Ostrogradski, en tenant compte de la convention qui a été choisie pour les flux, on a :

$$\oint_{\Omega} \vec{\varphi} \cdot \vec{n} dS = - \iiint_V \vec{\nabla} \cdot (\vec{\varphi}) dV \quad (1.5)$$

le bilan devient donc :

$$\iiint_V \left(\rho \frac{\partial h(T)}{\partial t} - \vec{\nabla} \cdot (\lambda \vec{\nabla}(T)) \right) dV = 0 \quad (1.6)$$

Cette intégrale est nulle quel que soit le volume V , donc l'intégrant est nul. Le bilan de puissance s'écrit par conséquent à un niveau local :

$$\rho \frac{\partial h(T)}{\partial t} = \vec{\nabla} \cdot (\lambda \vec{\nabla}(T)) \quad (1.7)$$

Si on considère de plus que la masse volumique ρ , en $kg \cdot m^{-3}$, est constante.

Pour résoudre cette équation en T , on utilise une équation d'état qui caractérise le matériau. C'est dans cette équation d'état que sera pris en compte le changement de phase. Il existe plusieurs méthodes pour prendre en compte le changement de phase.

1.1.1 c_p équivalent

La méthode du « c_p équivalent » étend la notion de capacité thermique au changement de phase en considérant un c_p équivalent, $c_{p,eq}$, incluant le changement de phase et qui donc dépend des capacités thermiques massiques du matériaux mais aussi de son taux liquide. En fait, une expression analytique de $c_{p,eq}$ n'est pas toujours possible, comme dans le cas du corps pur. Ce qui a conduit à considérer que la dérivée par rapport à la température $\frac{\partial h(T)}{\partial T}$ est proportionnelle au thermogramme, comme sur la figure 1.2 déterminée directement par la calorimétrie. On reviendra ultérieurement sur cette méthode pour la critiquer dans la partie 1.3.2.

$$\rho \frac{\partial h(T)}{\partial t} = \rho c_{p,eq} \frac{\partial T}{\partial t} \quad (1.8)$$

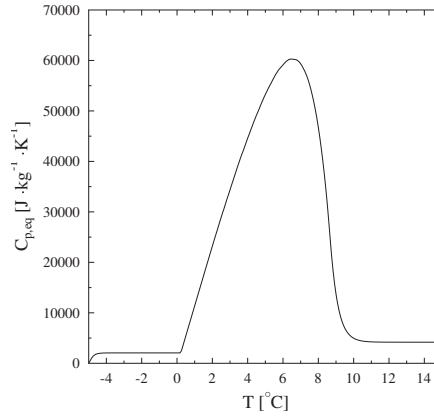


FIGURE 1.2 – $c_{p,eq}$ pour un échantillon d'eau de 12,52 mg et $5 K \cdot min^{-1}$

C'est la méthode la plus utilisée de nos jours, notamment dans des logiciels de simulation comme COMSOL [18] ou TRNSYS [19] [20] [21].

1.1.2 Méthode à suivi d'interface

Le but de cette formulation est de considérer le changement de phase comme une condition à la frontière entre deux domaines. Ce problème de Stefan nécessite donc de pouvoir prévoir la position exacte du front de changement de phase (ici la fusion) à chaque instant. Regardons le cas du corps pur avec son front de fusion caractérisé par sa température constante T_M et étudions ses autres difficultés.

Selon ces méthodes, pour un corps pur, l'équation de l'énergie donne :

$$\rho c_{p,i} \frac{\partial T_i}{\partial t} = \vec{\nabla} \cdot (\lambda_i \vec{\nabla} (T_i)) \quad (1.9)$$

$i = S, L$

et la condition à la frontière repérée par la position du front de fusion x_f est :

$$\left(\lambda_S \vec{\nabla}(T_S) - \lambda_L \vec{\nabla}(T_L) \right) \Big|_{x_f(t)} = \rho L_M \frac{dx_f}{dt} \quad (1.10)$$

Une première approche utilise un maillage fixe [22]. Mais comme l'interface n'est pas forcément sur des nœuds du maillage, cela oblige à interpoler le champ de température dans la zone de transition. Ceci est compliqué à implémenter pour un cas bi ou tri dimensionnel. Pour améliorer cette méthode une solution consiste à adapter le pas de temps afin que l'interface se situe toujours sur un nœud [22]. L'implémentation s'en trouve plus aisée. Mais dans ce cas, pour fidèlement modéliser le changement de phase le nombre de nœuds se doit d'être important. Cette exigence peut causer des problèmes de convergence numérique si le pas de temps est trop grand. Autant dire que le pas de temps risque d'être très petit, ce qui ralentira considérablement le calcul.

L. Landau [22] [23] en a proposé une deuxième : elle consiste à immobiliser l'interface et à faire évoluer le maillage. Le système d'équations qu'il propose est très complexe et nécessite un cadre mathématique très lourd et très contraignant. Et nous ne sommes que dans le cas du corps pur. Les cas qui nous intéressent étant plus complexes, ils vont forcément alourdir le schéma, des exemples d'utilisation de ses méthodes figurent dans les articles suivants [24] [25] [26]. Ce modèle sera ensuite utilisé dans une méthode inverse, ce qui entraînera une multiplication des calculs. Un cadre trop contraignant ne pourrait qu'amplifier les problèmes posés par ces méthodes. Cette formulation du problème de changement de phase ne pourra pas convenir à notre projet qui utilise les méthodes inverses. D'autant plus que dans le cas d'une fusion progressive, l'interface n'est pas clairement définie.

1.1.3 Méthodes enthalpiques

Cette méthode connaît deux écritures différentes.

1.1.3.1 Ecriture en température

La première consiste à écrire le bilan de puissance (1.7) directement en température en explicitant la relation enthalpie-température [27] [28] [29] [30] [31]. Etudions cette écriture dans le cas du corps pur et réécrivons donc ce bilan pour le résoudre en température, équation (1.11).

$$\rho c \frac{\partial T}{\partial t} = \vec{\nabla} \cdot (\lambda \vec{\nabla}(T)) + \frac{d\chi_L}{dt} \times L_M \quad (1.11)$$

Avec $\chi_L = \frac{m_L}{m_S + m_L}$ le taux liquide et L_M la chaleur latente massique du corps pur en $J \cdot kg^{-1}$.

Expliciter ainsi le bilan est une manière d'implémenter la méthode enthalpique. Mais la résoudre en température, sous la forme développée qu'elle a actuellement, oblige à utiliser une résolution itérative pour chaque pas de temps afin de « choisir » une solution physiquement acceptable pour χ_L dans tout le domaine, du fait de l'absence d'équation permettant de la déterminer [27] [28]. En pratique on vérifie que le couple choisi (T, χ_L) redonne une enthalpie conforme à la thermodynamique [31].

1.1.3.2 Ecriture en enthalpie

La seconde manière d'écrire la méthode enthalpique suppose de résoudre directement en enthalpie, de conserver l'équation (1.7) que l'on discretise, équation (1.12) [symbolique], pour déterminer l'enthalpie au pas de temps ultérieur $h_i^{t+\Delta t}$, de résoudre ensuite en se servant l'équation d'état pour trouver la nouvelle température de chaque cellule $T_i^{t+\Delta t}$, figure 1.3(b). Enfin on détermine χ_L en fonction de l'enthalpie de chaque cellule pour pouvoir combiner les propriétés de transfert dans la zone diphasique. C'est ce que nous détaillerons dans le prochain chapitre 2.1.

$$\rho V \frac{h_i^{t+\Delta t} - h_i^t}{\Delta t} = \lambda S \frac{T_{i-1}^t - T_i^t}{\Delta x} + \lambda S \frac{T_{i+1}^t - T_i^t}{\Delta x} \quad (1.12)$$

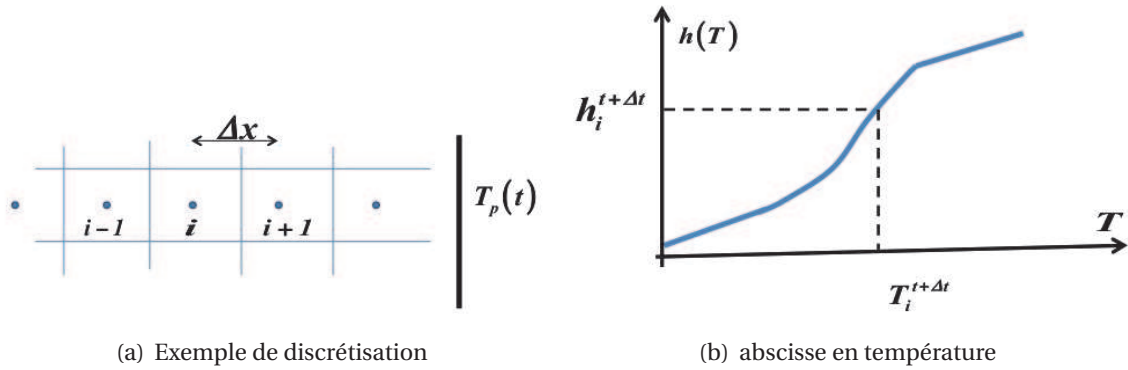


FIGURE 1.3 – Illustration de la résolution par méthode enthalpique, écrite en enthalpie

1.2 Differential Scanning Calorimetry (DSC)

Dans cette section, nous allons présenter un dispositif expérimental qui est très utilisé pour caractériser un changement de phase. Cette étude a pour but de permettre une meilleure interprétation des résultats, et *in fine* de caractériser correctement les propriétés énergétiques des matériaux étudiés. Nous allons commencer par présenter le fonctionnement global de l'appareil.

1.2.1 Présentation de l'appareil

L'appareil est un dispositif de calorimétrie différentielle à balayage par compensation de puissance ou Differential Scanning Calorimetry (DSC). Un calorimètre est un dispositif désigné pour décrire l'évolution thermique d'un échantillon, c'est-à-dire l'énergie échangée lors d'une variation linéaire de la température de plateau, $T_p(t)$.

Pour mesurer ce flux thermique, le calorimètre dispose de deux têtes en platine iridié : la première contient la cellule emplie de l'échantillon à étudier, et la deuxième contient une cellule vide servant de référence. Ces têtes sont en permanence en contact avec une source

froide qui est réalisée soit par un thermostat, soit par un cylindre métallique dont une extrémité plonge dans un bain d'azote liquide, figure 1.4.

L'échange thermique entre les différents éléments est réalisé par un courant d'hélium sec ou d'azote de pression réglable généralement de 1.5 bar [32]. Un courant d'azote ou d'air sec est également amené au niveau des plateaux afin d'éviter tous phénomènes de givrage.

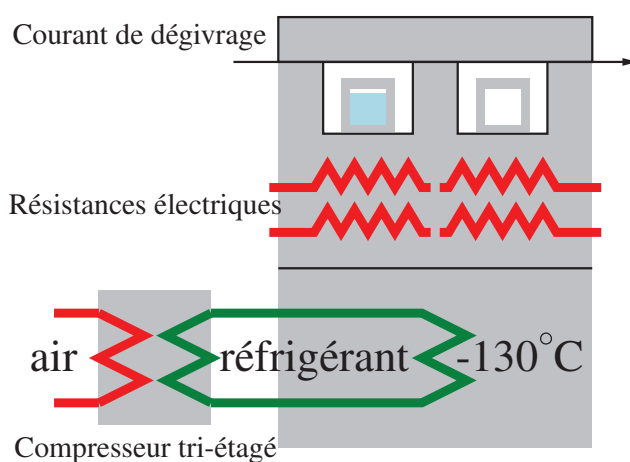


FIGURE 1.4 – Schéma global d'un calorimètre

Un système de capteurs et de régulation permet de faire subir la même évolution thermique aux deux têtes. Comme l'une des têtes contient une cellule pleine alors que l'autre contient la référence vide, les résistances électriques qui génèrent cette même évolution en température auront deux consommations d'énergie électrique différentes, du fait du phénomène thermique dans l'échantillon. C'est cette différence de puissance électrique qui est le flux thermique mesuré. De là provient la notion de compensation de puissance. On appelle thermogramme la courbe donnant l'évolution de ce flux thermique en fonction du temps. Cette courbe est la représentation du phénomène thermique que l'on étudie.

Cet appareil de DSC peut balayer un intervalle de température allant de -170°C à 500°C selon la source froide employée, avec des vitesses constantes aussi bien au réchauffement qu'au refroidissement comprises entre $0,1\text{ K}\cdot\text{min}^{-1}$ à $200\text{ K}\cdot\text{min}^{-1}$.

La figure 1.5 représente schématiquement la tête d'un calorimètre.

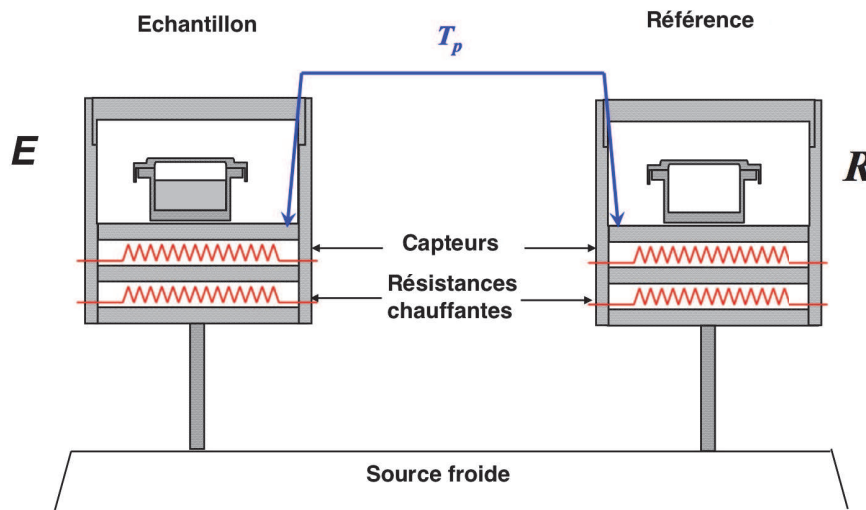


FIGURE 1.5 – Schéma de la tête du calorimètre

Nous utilisons un appareil Pyris Diamond DSC de Perkin-Elmer qui peut contenir des échantillons d'environ 10 mg.

Maintenant nous aborderons les deux modes de fonctionnement de cet appareil.

1.2.2 Mode dynamique

Le mode dynamique fait subir à l'échantillon une rampe de température de la forme :

$$T_p(t) = \beta \times t + T_0 \quad (1.13)$$

$T_p(t)$: Température de plateau à appliquer (K)

β : vitesse de réchauffement, β ($K \cdot \text{min}^{-1}$)

t : temps (s)

T_0 : valeurs de la température au début de la rampe, $t = 0$ s (K)

La principale variable est la vitesse de réchauffement, β . Elle contrôle directement la dynamique des phénomènes étudiés.

Des thermogrammes expérimentaux de la fusion de l'eau réalisés à plusieurs vitesses sont présentés à la figure 1.6. Ces graphiques appellent plusieurs commentaires :

- T_p étant une fonction affine du temps il est correct de l'utiliser comme abscisse, mais attention ce n'est pas la température au sein de l'échantillon.
- En dehors du pic de fusion, on observe que les thermogrammes se stabilisent à des valeurs constantes.
- On voit sur ces figures un intérêt de la représentation en T_p par l'alignement en début de fusion des courbes correspondant à des vitesses différentes. Nous reviendrons sur ces deux derniers points ci-après.

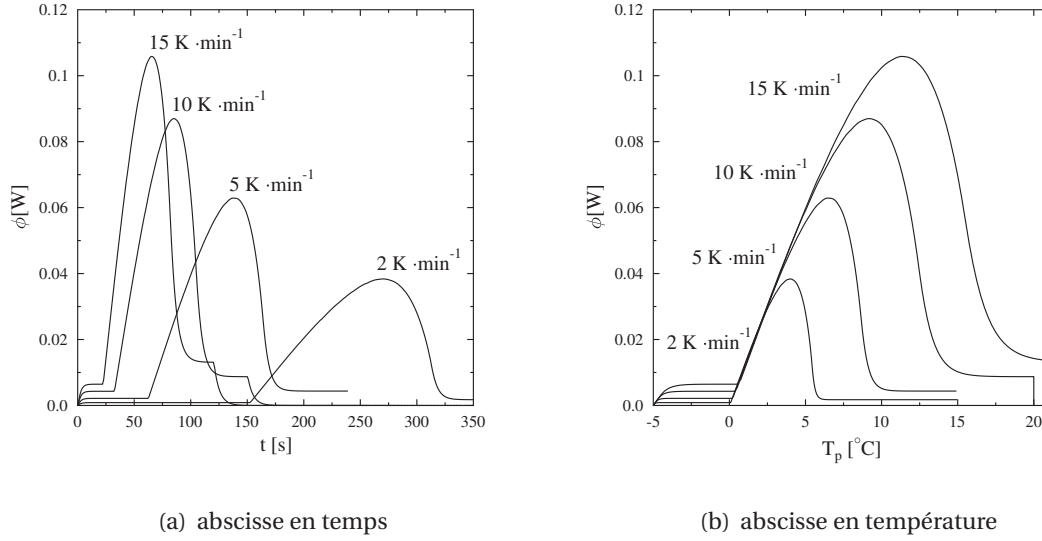


FIGURE 1.6 – Exemples de thermogrammes expérimentaux dans le cas de l’eau, $m=12,52 \text{ mg}$

Dans des domaines de températures où il n’y a pas de changement de phase ou autres transformations et une fois que le régime permanent est établi, l’hypothèse classique d’uniformité du champ de température est vérifiée pour de petits échantillons et de relativement faibles vitesses comme dans cette étude. Sous cette hypothèse le flux de chaleur donnant le thermogramme est [33] [34] [35] :

$$\phi(t) = \beta mc \quad (1.14)$$

avec m , c masse et capacité massique de l’échantillon.

Graduer en température présente l’avantage que, si on superpose plusieurs thermogrammes d’un même échantillon mais réalisés à différentes vitesses, on observe une superposition du début des pics comme dans la figure 1.6(b). Nous verrons que considérant que les creusets en aluminium sont uniformément à la température de plateau T_p [34], le flux thermique peut s’exprimer de la manière suivante :

$$\phi(t) = \sum_{i=1}^3 \alpha_i S_i (T_p(t) - T(t)) \quad (1.15)$$

$T(t)$ est la température de surface de l’échantillon, et $\alpha_i S_i$ est l’inverse de la résistance thermique entre l’environnement et la face i (supérieure, inférieure et latérale) du creuset rempli avec l’échantillon. Pendant les premiers instants d’une transition type corps pur ($t = 0$ s), la température thermodynamique de l’échantillon n’évolue pas sur un laps de temps très court et est égale à la température de fusion T_M . Sa dérivée temporelle peut être considérée nulle ($\left. \frac{dT}{dt}(t) \right|_{t=0s} = 0$). Par conséquent on a :

$$\left. \frac{d\phi}{dT_p}(t) \right|_{t=0s} = \sum_{i=1}^3 \alpha_i S_i \left(1 - \frac{dT}{dT_p} \cdot \left. \frac{dT}{dt}(t) \right|_{t=0s} \right) = \sum_{i=1}^3 \alpha_i S_i \quad (1.16)$$

$\sum_{i=1}^3 \alpha_i S_i$ est l'inverse de la résistance thermique équivalente de l'échantillon, caractéristique de l'échantillon et de l'appareil utilisés. C'est bien ce que l'on observe, la pente à l'origine du thermogramme est indépendante de la vitesse (figure 1.6(b)). L'équation (1.16) indique que cette pente représente la résistance thermique autour de l'échantillon. Ceci est un point important, on fera souvent référence à cette équation dans le reste de la thèse.

1.2.3 Mode Step

Le problème d'un fonctionnement dynamique est que l'évolution dynamique risque de masquer des phénomènes transitoires ou d'y introduire des retards du fait des gradients thermiques au sein du matériau qu'il est nécessaire d'apprécier comme nous le ferons. Une autre voie a été proposée, celle de la « step calorimetry », ou le mode de fonctionnement step du calorimètre. Qui au lieu de faire subir à l'échantillon une rampe de température sur un grand intervalle, lui fait subir une succession de mini-rampes entrecoupées par de grands plateaux où on attend que l'équilibre thermodynamique se rétablisse. L'intérêt du mode step est la lisibilité des températures de début et de fin de phénomène. Un autre intérêt tient au fait que la résolution en température est égale à l'incrément de températures entre les mini-plateaux. L'incertitude sur les températures est donc connue de manière précise. Mais il ne faut pas trop réduire cet incrément car le signal obtenu serait du même ordre de grandeur que le bruit et la précision de la mesure serait perdue [36]. L'inconvénient majeur spécifique à cette méthode est sa durée. Et là, on peut atteindre des durées expérimentales très contraignantes, de l'ordre de plusieurs jours. Mais il y a un autre problème, celui de la multiplication des effets transitoires et leurs effets potentiels sur la ligne de base. Se pose également la question de l'interprétation du thermogramme, qui a été beaucoup moins abordée que pour le mode dynamique. Récemment des comparaisons des deux modes de fonctionnement ont été faites par différents auteurs [37] [38].

1.3 Critiques des résultats

Nous allons maintenant nous intéresser à l'utilisation de cet appareil pour la caractérisation, noter les défauts de la méthode, pour, dans les chapitres suivants, proposer une autre interprétation des résultats.

1.3.1 Critiques de la DSC

Plusieurs auteurs [39] [36] [15] ont étudié les appareils de DSC et en ont montré les limites.

L'utilisation directe des résultats de la DSC pour la caractérisation se traduit par la détermination de l'aire sous le pic qui est égale à la chaleur latente [40], et à la détermination de la température *onset*¹, T_{onset} , que l'on relie à la température de fusion. On veut se servir des thermogrammes que l'on obtient par la DSC pour déterminer la relation enthalpie-température. C'est pourquoi la méthode du « c_p équivalent » a été développée.

1. plusieurs définitions en ont été données impliquant ou non des corrections, sans pour autant parvenir à un consensus [41] [42] [37] [43].

1.3.2 La méthode du « c_p équivalent »

Cette méthode consiste actuellement à considérer que le thermogramme est la dérivée de l'enthalpie par rapport à la température et que cette température thermodynamique est assimilée à la température de plateau T_p .

En effet, s'inspirant de l'équation (1.14), on a défini le c_p équivalent, $c_{p,eq}$, pour prendre en compte le changement de phase et généraliser la notion de capacité thermique.

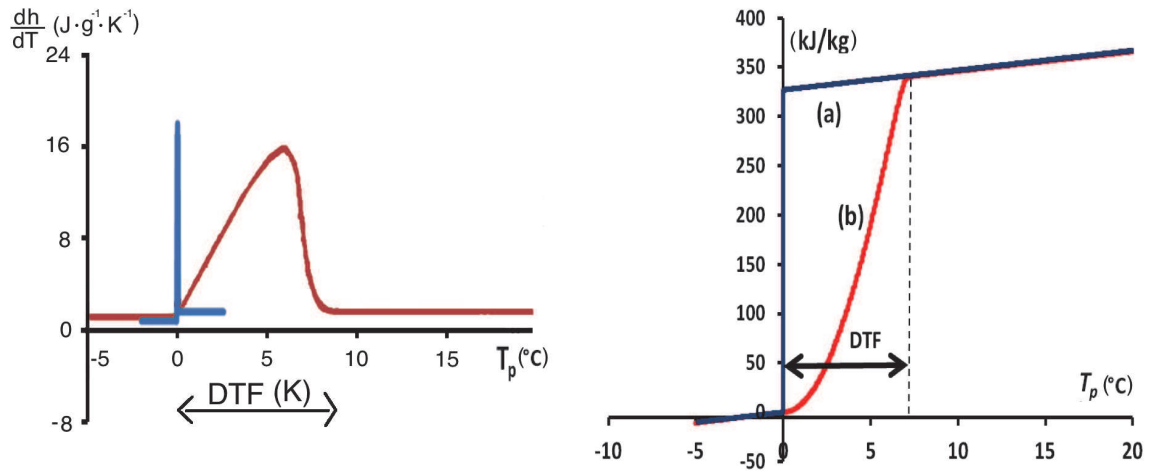
$$c_{p,eq} = \frac{\phi}{\beta m} \quad (1.17)$$

Cette écriture est vraie en dehors du changement de phase, mais c'est une approximation au moment de la fusion. Malgré cela, les auteurs l'utilisent aussi à la fusion [39] [36] [44] [?] [45].

Sous cette approximation, l'équation (1.7) devient avec (1.8) :

$$\rho c_{p,eq}(T) \frac{\partial T}{\partial t} = \vec{\nabla} \cdot (\lambda \vec{\nabla}(T)) \quad (1.18)$$

Plusieurs auteurs ont souligné les défauts de cette méthode [46] [15] [47]. Prenons le cas le plus flagrant que nous connaissons, celui de l'eau pure (figure 1.7). Si on considérait que la dérivée de l'enthalpie par rapport à la température est le thermogramme, on constaterait que le corps pur change de phase sur un intervalle de température, que nous noterions DTF, plutôt qu'à une température fixe, T_M . Cet écart correspondrait à la largeur du pic de fusion (figure 1.7). De nombreuses approximations ont été proposées pour modéliser ce $c_{p,eq}$ et pallier le problème sans pour autant être physiquement satisfaisantes [47].



(a) Dérivée théorique de l'enthalpie (bleu) et approximation du c_p équivalent (rouge) (b) Enthalpie théorique (bleu) et obtenue par la méthode du c_p équivalent (rouge) [46]

FIGURE 1.7 – Approximation du c_p équivalent pour un échantillon d'eau de 8,5 mg à $2 K \cdot min^{-1}$

Si la méthode du c_p équivalent fournissait une approximation satisfaisante de l'enthalpie, un code qui simule le fonctionnement du calorimètre pour la fusion devrait donner le même thermogramme pour les deux enthalpies précédentes (le code en question sera étudié en

détail dans le chapitre suivant 2). Mais ce n'est pas le cas, le thermogramme simulé avec la deuxième enthalpie (en rouge/ (b) sur la figure 1.7(b)) donne un thermogramme (en rouge, (b) sur la figure 1.8) très différent de celui obtenu avec l'enthalpie théorique (en bleu, (a) sur la figure 1.8). Cette méthode n'est pas satisfaisante pour la caractérisation.

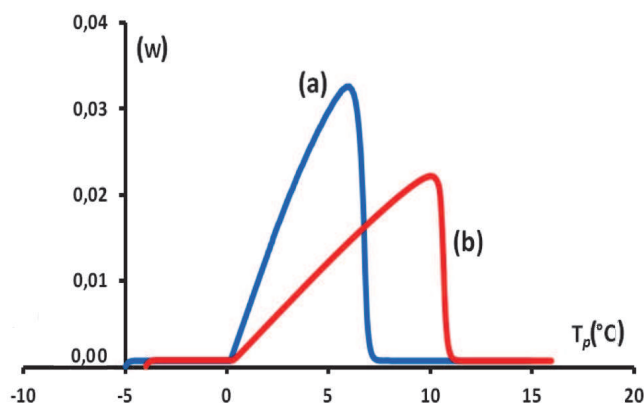


FIGURE 1.8 – Thermogramme théorique & obtenu par la méthode du c_p équivalent [46]

De plus, une des critiques majeures de la caractérisation directe par DSC subsiste, les résultats obtenus par cette méthode varient également selon la masse de l'échantillon ou la vitesse de réchauffement employées (figures 1.9 et 1.10). Sur cette figure, on voit qu'une augmentation de la vitesse augmente sensiblement l'écart DTF et donc l'erreur sur l'enthalpie. On pourrait penser qu'en diminuant suffisamment la vitesse on pourrait négliger cette erreur car le système se rapprochant de l'équilibre thermodynamique, la température de plateau est proche de la température moyenne de l'échantillon. Mais il s'avère que plus β est faible, plus le signal est faible, et donc plus le rapport entre le signal et le bruit est faible. Si la vitesse est trop faible, on n'obtient pas de signal exploitable. Pour augmenter la valeur du flux afin de pouvoir exploiter le signal, on peut augmenter la masse, mais on constate sur cette même figure que l'on doit diminuer la vitesse d'autant plus que la masse est importante.

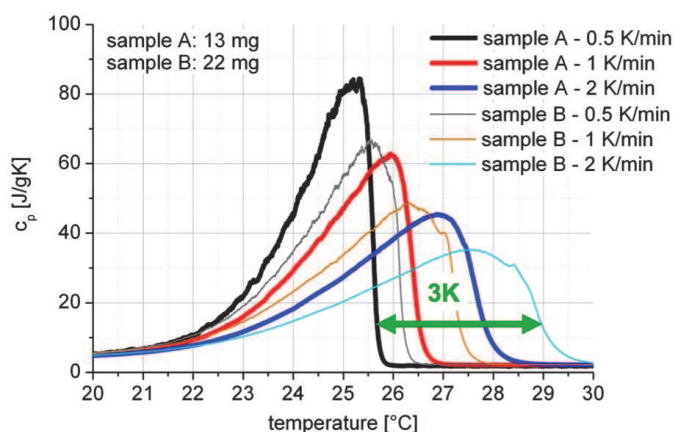


FIGURE 1.9 – Influence de la vitesse et de la masse sur le $c_{p,eq}$ [36] [48]

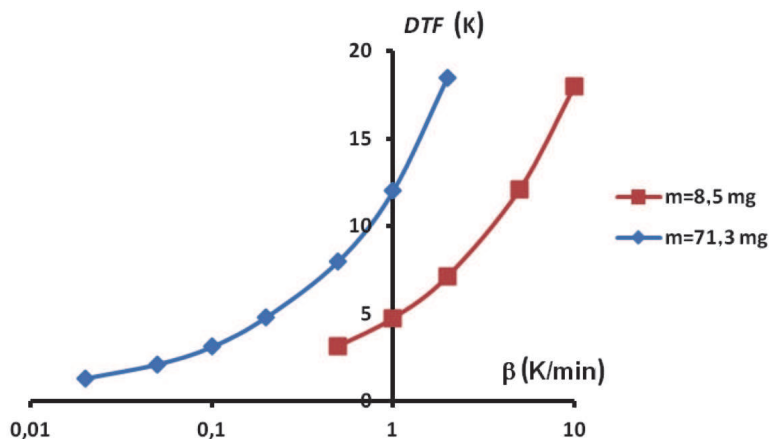
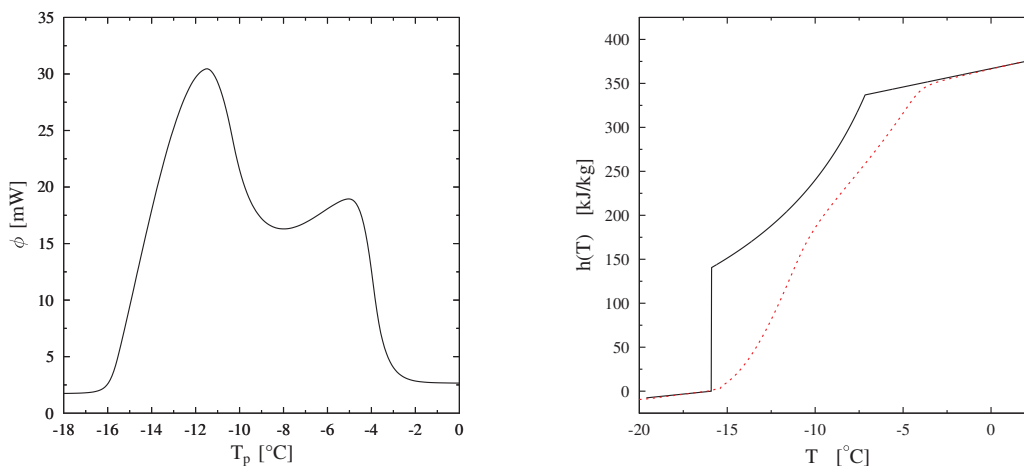


FIGURE 1.10 – Influence de la vitesse et de la masse sur DTF [46]

Par ailleurs, dans le cas d'un corps pur, il est relativement facile de se rendre compte de l'erreur que l'on fait, et de la quantifier puisqu'on assimile une transition brusque à une transition étalée sur un intervalle DTF. Mais dans le cas d'un corps quelconque, comme une solution binaire (figure 1.11), on ne peut pas aussi facilement identifier l'erreur qui est commise par cette approximation car le changement de phase se produit sur un intervalle de température.

(a) Thermogramme à $5 \text{ K} \cdot \text{min}^{-1}$ et $m = 9.6 \text{ mg}$ (b) Enthalpie théorique et obtenue par la méthode du c_p équivalent [46]FIGURE 1.11 – Enthalpies et thermogramme pour le cas d'une solution binaire de $\text{H}_2\text{O} - \text{NH}_4\text{Cl}$ à 9,02%

1.3.3 Exemples de l'influence de l'erreur d'identification sur un cas, le mur [46]

Nous avons vu qu'il y avait mauvaise interprétation des résultats de la DSC pour la caractérisation en thermique. Ces résultats erronés servent ensuite de données d'entrées pour des logiciels de simulations. Le problème qui se pose donc est celui de la fiabilité de ces simulations en considérant une erreur sur les enthalpies. Afin de l'évaluer nous allons étudier le cas d'un mur avec des inclusions de MCP, soumis aux variations quotidiennes de température. Pour simplifier le problème, nous allons considérer qu'il n'y a pas de surfusion.

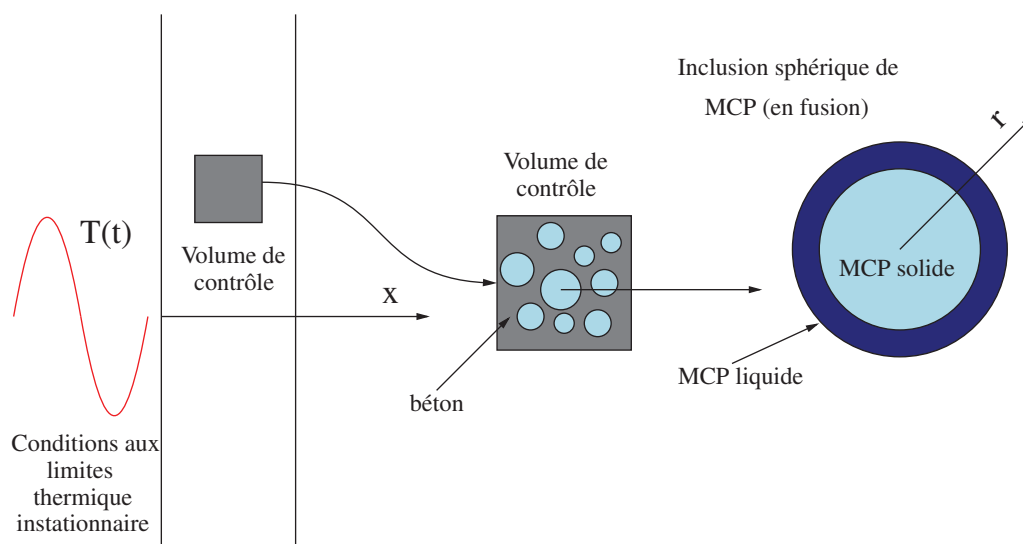


FIGURE 1.12 – Exemple d'inclusion de MCP dans un mur.

Nous allons étudier le cas d'un mur en béton contenant 30 % d'octadécane, $T_M = 28^\circ\text{C}$, soumis sur sa face gauche à différentes variations et sur sa face droite ($x = L$) à une température constante, $T_{air} = 20^\circ\text{C}$. Les paramètres de ce calcul sont :

$T_{air} = 20^\circ\text{C}$	$L = 200\text{ mm}$	$\rho = 1385\text{ kg} \cdot \text{m}^{-3}$
$T_\infty = 50^\circ\text{C}$	$\tau = 5\text{ h}$	$L_F = 72\text{ kJ} \cdot \text{kg}^{-1}$
$T_0 = 20^\circ\text{C}$	$\lambda = 0,65\text{ W} \cdot \text{m}^{-1} \cdot \text{K}^{-1}$	$c_S = c_L$
$T_M = 28^\circ\text{C}$	$\alpha = 50\text{ W} \cdot \text{m}^{-2} \cdot \text{K}^{-1}$	$1690\text{ J} \cdot \text{kg}^{-1} \cdot \text{K}^{-1}$

Tableau 1.1 – Paramètres des simulations

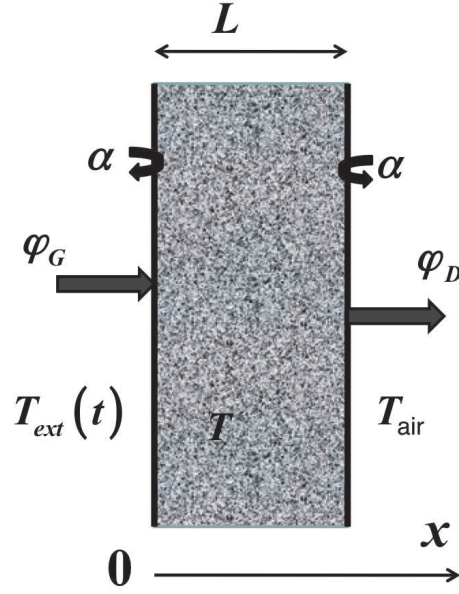


FIGURE 1.13 – Coupe transversale du mur d'épaisseur 200 mm. [46] [49]

Plusieurs simulations ont été faites sur ce mur en considérant soit qu'il n'y avait pas de changement de phase ($L_M = 0 \text{ Jkg}^{-1}$), soit que le changement de phase est bien représenté (DTF=0), soit qu'il y avait une erreur de modélisation sur la relation enthalpie-température déduite d'un thermogramme avec DTF=2,5 ou 10 K (voir figure 1.7).

Pour chacune des simulations, nous allons résoudre l'équation de la chaleur unidimensionnelle :

$$\rho \frac{\partial h}{\partial t}(T(x, t)) = \lambda \frac{\partial^2 T}{\partial x^2}(x, t) \quad (1.19)$$

Où, pour ces calculs, nous avons considéré d'une part la relation enthalpie-température du corps pur, dont la forme générale est donnée dans la sous-section 2.1.2, et d'autre part l'enthalpie obtenue par la méthode du $c_{p,eq}$ dans différentes conditions résultant en différentes valeurs de DTF (voir figures 1.7(b), 1.9 et 1.10).

On considère les flux :

$$\varphi_G = \alpha \times (T(0, t) - T_{ext}(t)) \quad (1.20)$$

$$\varphi_D = \alpha \times (T(L, t) - T_0) \quad (1.21)$$

Dans un premier temps nous avons étudié la réponse du mur à une variation exponentielle de la température extérieure du type ($T_{air} = T_0$) :

$$T_{ext}(t) = T_0 + (T_\infty - T_0)(1 - e^{-\frac{t}{\tau}}) \quad (1.22)$$

On voit à la figure 1.14 que la présence du matériau à changement de phase diminue de moitié, dans notre exemple, la quantité de chaleur qui traverse le mur par rapport à ce que

se serait sans le MCP. Mais on constate aussi que les erreurs de caractérisation conduisent à surestimer les flux sortants du mur de 40% ainsi que de surestimer les températures de plus de 4 K. Une erreur sur les flux devient une erreur sur la consommation énergétique.

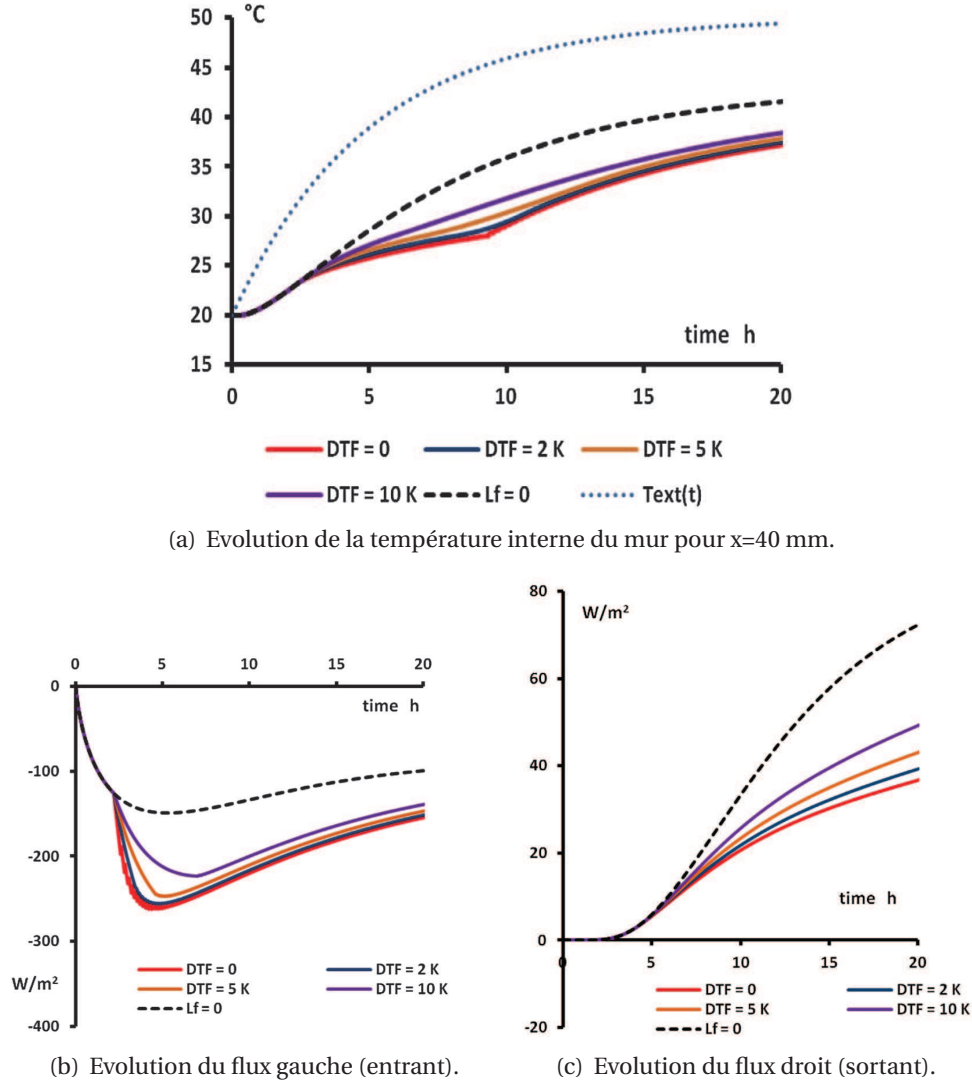


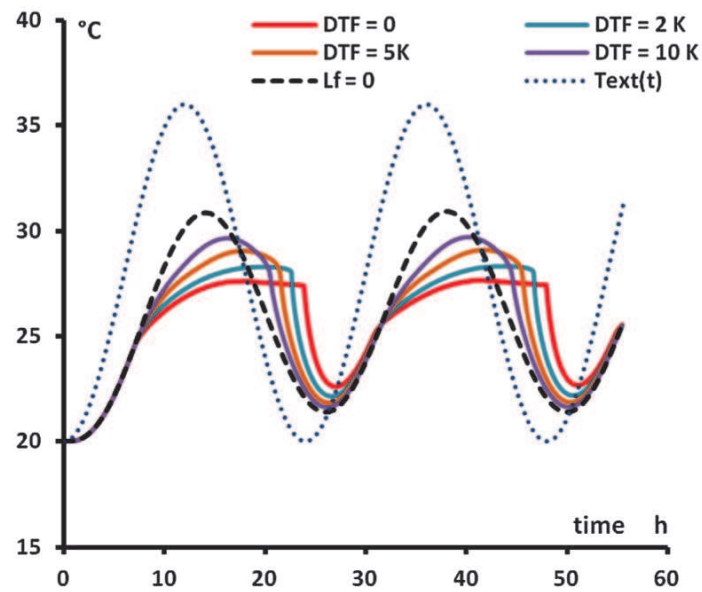
FIGURE 1.14 – Réponse du mur à un réchauffement exponentiel de l'extérieur. [46]

Prenons maintenant le cas d'une variation jour-nuit de la température extérieure en période estivale de la forme :

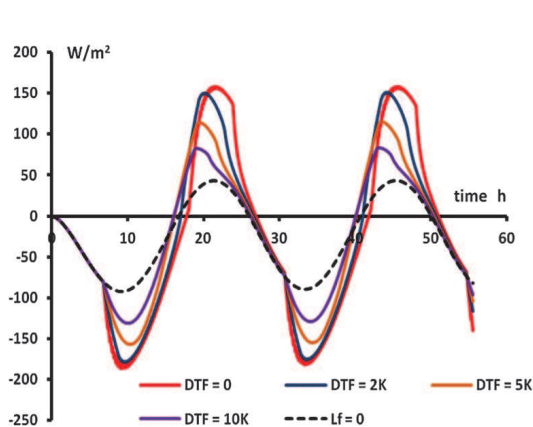
$$T_{ext}(t) = T_0 + (\bar{T} - T_0) \cos\left(\frac{2\pi t}{24}\right) \quad (1.23)$$

$\bar{T}$ est la température moyenne de la journée que nous choisirons égale à la température de fusion du MCP (ici l'octadécane). Les autres propriétés étant les mêmes que dans le cas précédents 1.1

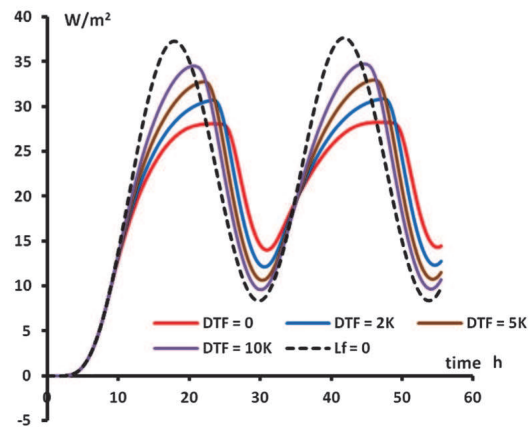
Avec ce cas, on voit sur les figures 1.15(a), 1.15(b) et 1.15(c) que le matériau à changement de phase permet de diminuer l'échange énergétique à travers le mur, et à empêcher un réchauffement excessif au-delà du mur. Comme pour le réchauffement exponentiel, une erreur de modélisation DTF sur l'enthalpie conduit à surestimer les températures et à surestimer les flux sortants du mur. Cela peut amener à des erreurs sur les échanges d'énergie de 10% à 50%, ainsi que des déphasages de quelques heures sur les prévisions [49] [46]. Une erreur (DTF) de 10 K sur la température de fusion résulte en une erreur de 11% sur la puissance d'un dispositif de stockage [36] et donc à un mauvais dimensionnement. Il apparaît donc nécessaire de développer une autre méthode pour exploiter correctement les résultats de la DSC. Nous le ferons dans les chapitres suivants et présenterons les résultats de cette méthode dans le chapitre 4.



(a) Evolution de la température interne du mur pour $x=40$ mm.



(b) Evolution du flux gauche (entrant).



(c) Evolution du flux droit (sortant).

FIGURE 1.15 – Réponse du mur à une évolution sinusoïdale de la température extérieure. [46]

1.3.4 Nécessité de prendre en compte le gradient interne

Il faut connaître l'enthalpie avec précision et la déduire du thermogramme. Les travaux du laboratoire [33] [50] [34] ont montré que l'hypothèse de la température homogène dans l'échantillon n'était pas valable et qu'il était nécessaire de tenir compte des transferts thermiques à l'intérieur de l'échantillon, voir figures 1.16 et 1.17. Sur la première figure, on constate que l'hypothèse de la température homogène ne permet pas de reproduire le thermogramme, et sur la deuxième figure, on voit que cette hypothèse est fautive, qu'il peut y avoir plusieurs degrés d'écart entre les températures internes et la température de plateau T_p . Bien sûr, dans des cas de très grande conductivité comme l'indium cette hypothèse est moins fautive. Mais si la conductivité n'est pas très grande le gradient interne peut atteindre plusieurs K selon la masse et la vitesse β utilisées. On peut adapter ces deux paramètres pour réduire le gradient mais on ne peut pas trop réduire la vitesse car, comme nous l'avons déjà dit dans la sous-section 1.3.2, le bruit deviendrait prépondérant, et de plus diminuer la vitesse reviendrait à augmenter la durée de l'expérience. On pourrait diminuer d'avantage la masse de l'échantillon, mais alors la difficulté de préparer un échantillon et la rigueur des conditions opératoires s'en trouve augmentée. Il faut aussi prendre en compte l'aspect financier, le coût de la balance qui mesure précisément une masse inférieure à quelques milligrammes avec quatre chiffres significatifs nous empêche de considérer leurs utilisations massives.

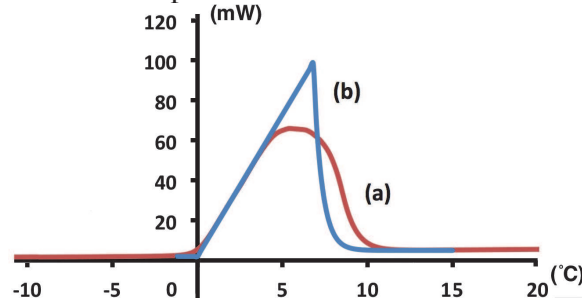


FIGURE 1.16 – Thermogramme théorique suivant l'hypothèse température homogène (b) et réel (a)

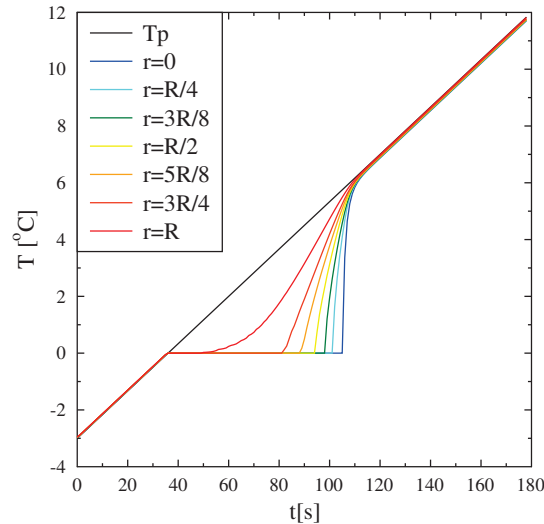


FIGURE 1.17 – Evolution du champ de température interne à l'échantillon selon plusieurs rayons et une hauteur correspondant au centre thermique : cas de 11,37 mg d'eau pour $\beta = 5 \text{ K} \cdot \text{min}^{-1}$

1.4 Conclusion

Les thermogrammes sont actuellement mal interprétés et cette erreur d'interprétation, due au gradient thermique, conduit, en outre, à de mauvais dimensionnements, à une mauvaise utilisation de la technologie de stockage par MCP. Il apparaît nécessaire de développer une meilleure méthode de caractérisation. L'idée est donc de comparer des thermogrammes expérimentaux à des thermogrammes numériques qui, comme nous le verrons dans le chapitre suivant, modélisent le gradient thermique. Pour avoir ensuite recours à une méthode inverse, que nous détaillerons, afin d'ajuster les paramètres du modèle numérique pour obtenir la superposition des deux thermogrammes. Ces paramètres, qui comportent les propriétés thermodynamiques, une fois ajustés permettent de reconstruire l'enthalpie réelle du MCP considéré.

Chapitre 2

Modèle direct de la DSC

2.1 Description du modèle 2D

Ce chapitre, décrit la modélisation du fonctionnement d'un calorimètre, pour pouvoir étudier l'influence des propriétés des matériaux étudiés sur le thermogramme généré par cet appareil. Une fois le modèle direct obtenu et que son fonctionnement aura été étudié, les méthodes employées pour caractériser les matériaux seront abordés dans le chapitre 3.

2.1.1 Modèle générique

Nous utilisons la DSC pour étudier, caractériser des matériaux lors du changement de phase. Il nous faut donc développer un modèle numérique qui reproduit le thermogramme expérimental de DSC dans le cas du changement de phase solide-liquide. L'appareil de DSC utilisé est un Pyris Diamond DSC de Perkin-Elmer qui utilise des échantillons d'environ 10 ml. La cellule qui contient l'échantillon à analyser est de géométrie cylindrique. De nombreux travaux [51] [52] obtenant des résultats probants avec une telle géométrie, nous choisissons donc une géométrie cylindrique à deux dimensions. Une fois la géométrie définie, se pose la problématique de la modélisation du phénomène physique.

Comme dans le problème de Stefan, la modélisation du phénomène à changement de phase présente un déplacement d'interface suivant le temps et l'espace, ce qui pose de nombreux problèmes. De nombreux articles sont parus sur le sujet. Dutil et al. [?] ont proposé un bon inventaire des différentes méthodes existantes pour traiter le problème. Plus de détails commentés sur ces méthodes sont disponibles dans les articles de Sharma et al. [13], Agyenim et al. [53] ou Date [54].

Nous avons choisi de formuler le problème de la zone de transition solide-liquide selon la méthode enthalpique de la sous-section 1.1.3 [55] [56] [57] [13] [17]. En effet, l'approche enthalpique présente de nombreux avantages numériques, l'un d'eux étant que le traitement de l'interface ne requiert pas d'écriture particulière comme l'équation (1.10). L'interface peut être traitée comme le reste du domaine modélisé, il n'y a pas de propriétés particulières de l'interface qui ont besoin d'être connues *a priori*.

On peut exprimer la conduction dans l'échantillon sous la forme (1.7) dont la résolution passe par la connaissance de l'équation d'état correspondante, que nous verrons après cette sous-section. Rappelons cette équation (1.7) :

$$\rho \frac{\partial h}{\partial t} = \vec{\nabla} \cdot (\lambda \vec{\nabla}(T)) \quad (2.1)$$

A cette équation, on adjoint des conditions aux limites et initiales que nous allons maintenant poser.

Comme on peut le voir sur la figure 2.1 l'échantillon à analyser est contenu dans un creuset en aluminium. Le creuset repose sur un plateau en platine, qui est à la température T_p de la sollicitation (1.13) :

$$T_p(t) = \beta \times t + T_0 \quad (2.2)$$

Le creuset étant très conducteur on considère qu'il est uniformément à la température T_p [34] et que tout le flux qui passe par le creuset, provient du plateau du fait que la conductivité de l'air est négligeable devant celle du plateau. On considère donc que l'échantillon se trouve entouré par un thermostat de température T_p .

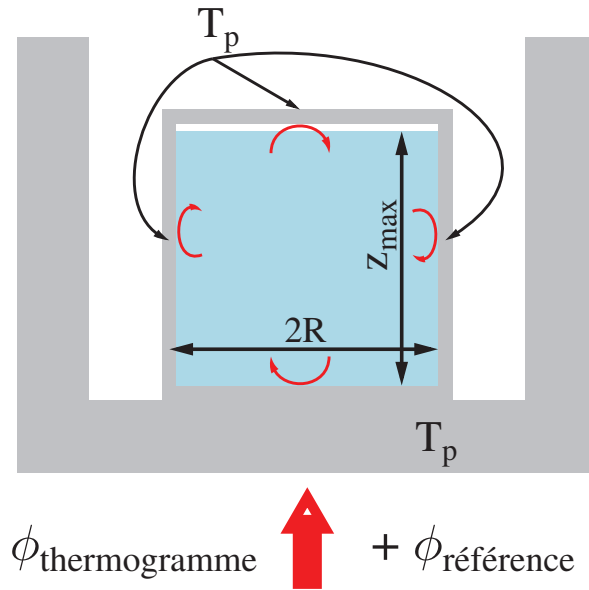


FIGURE 2.1 – Conduction autour de la cellule

Nous prendrons pour représenter l'échange thermique de paroi, un coefficient d'échange équivalent α_i pour chaque surface (voir figure 2.2) qui englobe les différents échanges. Les conditions aux limites de l'échantillon cylindrique s'écrivent donc localement avec l'équation (2.3) :

$$-\lambda \vec{\nabla} T \cdot \vec{n} = \alpha_i \times (T - T_p(t)) \text{ sur } \partial\Omega_i \quad (2.3)$$

Ces conditions aux limites permettent de réexprimer le thermogramme en fonction de ces dernières :

$$\phi(t) = \sum_{i=1}^3 \int_{\partial\Omega_i} -\alpha_i \times (T(r, z, t) - T_p(t)) dS \quad (2.4)$$

Les conditions aux limites sont représentées à la figure 2.2 .

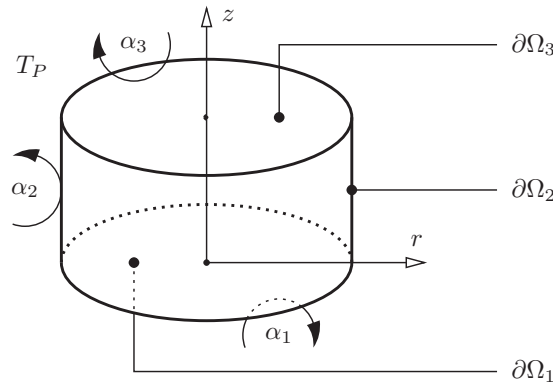


FIGURE 2.2 – Modèle de l'échantillon.

Pour conditions initiales, nous avons choisi de considérer un système homogène en température. Nous partons d'un état thermiquement stable de l'appareil.

Maintenant pour résoudre le bilan (2.1), il est nécessaire d'avoir une relation entre l'enthalpie massique h , qui est une fonction monotone croissante de la température thermodynamique, et la température thermodynamique T afin de fermer le système. Celle-ci est donnée pour les deux cas que nous avons étudiés, le corps pur et la solution binaire présentant un liquidus rectiligne, par leurs équations d'états. Nous allons présenter ces relations.

2.1.2 Enthalpie du corps pur

Si on considère dans notre étude que les capacités massiques sont indépendantes de la température, les lois de la thermodynamique nous montrent que l'enthalpie d'un corps pur s'exprime (voir figure 2.3 pour un graphique) :

$$h = \begin{cases} c_S(T - T_{ref}) + h_{ref} & T < T_M \\ c_S(T_M - T_{ref}) + \chi_L \cdot L_M + h_{ref} & T = T_M \\ c_S(T_M - T_{ref}) + L_M + c_L(T - T_M) + h_{ref} & T > T_M \end{cases} \quad (2.5)$$

avec c_S , c_L les capacités spécifiques des corps aux états solides et liquides en $J \cdot kg^{-1} \cdot K^{-1}$, L_M la chaleur latente en $J \cdot K^{-1}$ et χ_L le taux liquide. Pendant le changement de phase, la

température reste constante ($T = T_M$) et seul le taux liquide évolue, passant de 0 à 1 selon la loi :

$$\chi_L = \frac{h - h_{ref}}{L_M} \text{ si } T = T_M \quad (2.6)$$

la référence, $h_{ref} = 0 \text{ J} \cdot \text{kg}^{-1}$ étant prise à l'état solide et à la température $T_{ref} = 273,15 \text{ K}$ de façon arbitraire.

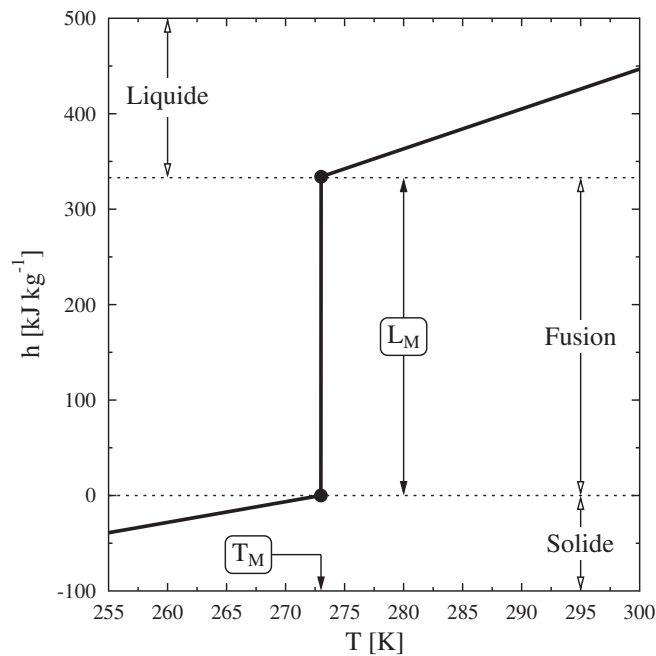


FIGURE 2.3 – Enthalpie d'un corps pur, l'eau

2.1.3 Enthalpie de solutions binaires

Pour cette étude, nous nous restreindrons au cas où le diagramme de phase présente un solidus vertical et un liquidus rectiligne, comme on peut le voir à la figure 2.4. Il s'agit, par exemple, du cas d'une solution saline aqueuse faiblement concentrée. Les solutions $\text{H}_2\text{O} - \text{KCl}$ et $\text{H}_2\text{O} - \text{NH}_4\text{Cl}$ présentent un tel diagramme de phase ; elles seront présentées en exemple. Nous allons reprendre ce calcul classique [58].

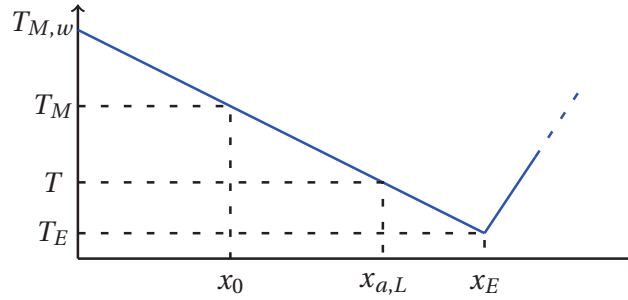


FIGURE 2.4 – Diagramme de phase d'une solution binaire avec un solidus vertical et un liquidus rectiligne.

Le point eutectique (x_E, T_E) correspond à un équilibre entre trois phases (une phase liquide et les deux phases solides pures). L'équilibre est monovariant, c'est-à-dire qu'un seul paramètre intensif est libre d'évoluer sans rompre l'équilibre. Comme on travaille sous l'hypothèse isobare, la température et la composition du liquide sont fixées par l'équilibre. $T_{M,w}$ est la température de fusion du solvant pur. Le solidus vertical s'étend de 0 K à $T_{M,w}$ sur l'axe des ordonnées. Le point (x_0, T_M) situé sur la courbe du liquidus, courbe qui définit l'instant où disparaît le dernier cristal lors d'une fusion, est le point équilibre de la solution étudiée.

Ainsi, la montée en température d'une solution de fraction massique en soluté x_0 inférieure à x_E , se fera d'abord par une fusion du soluté ainsi que d'une partie du solvant quand la température atteindra la température de l'eutectique T_E , suivie de la fusion graduelle du reste du solvant. Elle sera terminée lorsque la température du liquidus T_M sera atteinte. Si la solution est à la concentration x_E , alors tout le solvant aura fondu au plateau eutectique ; la fusion est alors à température constante comme pour un corps pur.

Définissons le taux massique solide du solvant :

$$\chi_{w,S} = \frac{m_{w,S}}{m_w + m_a} \quad (2.7)$$

et rappelons que la fraction massique du soluté s'écrit :

$$x_0 = \frac{m_a}{m_w + m_a} \quad (2.8)$$

avec m_a et m_w les masses totales de départ de soluté, a, et de solvant, w.

Ainsi nous avons quatre cas à considérer :

– $T < T_E$

$$dH = m_{w,S}^* c_{w,S} dT + m_{a,S}^* c_{a,S} dT = m_w c_{w,S}^* dT + m_a c_{a,S}^* dT \quad (2.9)$$

$$\frac{dh}{dT} = \frac{d}{dT} \left(\frac{H}{m_w + m_a} \right) = (1 - x_0) c_{w,S}^* + x_0 c_{a,S}^* \quad (2.10)$$

– $T = T_E$ and $1 - \frac{x_0}{x_E} \leq \chi_{w,S} \leq 1 - x_0$ [59]

$$dH = \overline{h_{w,L}}(T_E, x_E) dm_{w,L} + \overline{h_{a,L}}(T_E, x_E) dm_{a,L} + h_{w,S}^*(T_E) dm_{w,S} + h_{a,S}^*(T_E) dm_{a,S} \quad (2.11)$$

où $\overline{h_{w,L}}(T_E, x_E)$ et $\overline{h_{a,L}}(T_E, x_E)$ sont les enthalpies massiques partielles du solvant et du soluté à la température eutectique T_E et à la fraction massique x_E .

Étant donné que nous avons :

$$m_T = m_a + m_w = m_L + m_S \quad (2.12)$$

$$m_w = m_{w,S} + m_{w,L} \quad (2.13)$$

$$m_a = m_{a,S} + m_{a,L} \quad (2.14)$$

nous avons aussi :

$$dm_{w,S} = -dm_{w,L} \quad (2.15)$$

$$dm_{a,S} = -dm_{a,L} \quad (2.16)$$

Ce qui implique que :

$$dH = \left(\overline{h_{w,L}}(T_E, x_E) - h_{w,S}^*(T_E) \right) dm_{w,L} + \left(\overline{h_{a,L}}(T_E, x_E) - h_{a,S}^*(T_E) \right) dm_{a,L} \quad (2.17)$$

Les termes de l'équation (2.17) s'expriment :

- pour le premier :

$$\overline{h_{w,L}}(T_E, x_E) - h_{w,S}^*(x_E) = \overline{h_{w,L}}(T_E, x_E) - h_{w,L}^*(T_E) + h_{w,L}^*(T_E) - h_{w,S}^*(T_E) \quad (2.18)$$

où :

$$\overline{h_{w,L}}(T_E, x_E) - h_{w,L}^*(T_E) = L_w^{dil}(T_E, x_E) \quad (2.19)$$

est l'enthalpie spécifique de dilution du solvant pour une solution eutectique de fraction massique x_E , et :

$$h_{w,L}^*(T_E) - h_{w,S}^*(T_E) = L_w(T_E) \quad (2.20)$$

est la chaleur latente spécifique de l'eau pure à la température T_E .

Généralement, le terme (2.19) est négligeable. Le terme (2.20) peut être calculé par l'expression :

$$L_w(T_E) = L_w(T_{M,w}) + \int_{T_{M,w}}^{T_E} (c_{w,L}^* - c_{w,S}^*) dT \quad (2.21)$$

Ici encore, si on considère que les capacités spécifiques sont indépendantes de la température et si on suppose que la capacité spécifique du liquide peut être extrapolée à partir de sa valeur au-dessus du point de fusion, nous avons :

$$L_w(T_E) = L_w(T_{M,w}) + (c_{w,L}^* - c_{w,S}^*)(T_E - T_{M,w}) \quad (2.22)$$

- pour le second terme de l'équation (2.17), on peut écrire que :

$$\overline{h_{a,L}}(T_E, x_E) - h_{a,S}^*(T_E) = L_a^{dil}(T_E, x_E) + \Delta h_a^{dis}(T_E) \quad (2.23)$$

où $L_a^{dil}(T_E, x_E)$ est l'enthalpie spécifique du soluté dans la solution eutectique, qui peut toujours être négligée, et $\Delta h_a^{dis}(T_E)$ est l'enthalpie spécifique de dissolution du soluté à la température T_E .

En notant x_E la fraction eutectique massique,

$$x_E = \frac{m_{a,L}^E}{m_{a,L}^E + m_{w,L}^E} \quad (2.24)$$

et $x_{a,L}$ la fraction massique de soluté dans la phase liquide,

$$x_{a,L}(T) = \frac{m_{a,L}}{m_L} \quad (2.25)$$

alors, en tenant compte qu'au palier eutectique la phase liquide est à la fraction eutectique ($x_{a,L} = x_E$), la relation précédente donne :

$$m_{a,L} = m_{w,L} \frac{x_E}{1 - x_E} \quad (2.26)$$

l'équation (2.17) devient :

$$dH = \left(L_w(T_E) + \frac{x_E}{1 - x_E} \Delta h_a^{dis} \right) dm_{w,L} = - \left(L_w(T_E) + \frac{x_E}{1 - x_E} \Delta h_a^{dis} \right) (m_a + m_w) d\chi_{w,S} \quad (2.27)$$

En utilisant les définitions (2.7), (2.13) et (2.15), nous avons donc :

$$\frac{dh}{d\chi_{w,S}} = - \left(L_w(T_E) + \frac{x_E}{1 - x_E} \Delta h_a^{dis} \right) = \frac{-L_E(T_E)}{1 - x_E} \quad (2.28)$$

avec $L_E(T_E)$ l'énergie par unité de masse mise en jeu à l'eutectique, ou chaleur eutectique.

La variation d'enthalpie à la température eutectique est donc :

$$\widetilde{L}_E = \int_{1-x_0}^{1-\frac{x_0}{x_E}} \frac{dh}{d\chi_{w,S}} d\chi_{w,S} = \frac{x_0}{x_E} L_E(T_E) = \frac{T_{M,w} - T_M}{T_{M,w} - T_E} L_E(T_E) \quad (2.29)$$

car, comme on a un liquidus rectiligne :

$$\frac{x_0}{x_E} = \frac{T_{M,w} - T_M}{T_{M,w} - T_E} \quad (2.30)$$

- $T_E < T \leq T_M$ et $0 \leq \chi_{w,S} \leq 1 - \frac{x_0}{x_E}$ [59]

Dans ce cas, nous avons eu une fusion progressive et tout le soluté est dans la phase liquide.

$$m_{a,S} = 0 \quad (2.31)$$

ce qui implique que :

$$m_S = m_{w,S} \quad (2.32)$$

$$m_a = m_{a,L} \quad (2.33)$$

Les relations (2.32) et (2.33) permettent de réécrire x_0 :

$$x_0 = \frac{m_a}{m} = \frac{m_{a,L}}{m} = x_{a,L} \times x_L \Leftrightarrow \chi_{w,S}(T) = 1 - \frac{x_0}{x_{a,L}(T)} \quad (2.34)$$

Dans cet intervalle, la variation de l'enthalpie devient :

$$\begin{aligned} dH = \overline{h_{w,L}}(T, X_{a,L}) dm_{w,L} &+ h_{w,S}^* dm_{w,S} + m_{w,L} \overline{c_{w,L}}(T, X_{a,L}) dT + m_{a,L} \overline{c_{a,L}}(T, X_{a,L}) dT \\ &+ m_{w,S} c_{w,S}^*(T) dT \end{aligned} \quad (2.35)$$

comme précédemment en négligeant l'enthalpie de dilution il vient :

$$\overline{h_{w,L}}(T, x_{a,L}) dm_{w,L} + h_{w,S}^* dm_{w,S} \approx -L_w(T) d\chi_S \quad (2.36)$$

Puis en s'aidant des relations (2.32), (2.33), (2.25) et (2.7) on obtient :

$$\frac{m_{a,L}}{m_T} = \frac{m_{a,L}}{m_L} \times \frac{m_T - m_{w,S}}{m_T} = x_{a,L}(1 - \chi_{w,S}) \quad (2.37)$$

et

$$\begin{aligned} \frac{m_{w,L}}{m_T} &= 1 - \frac{m_{a,L} + m_{w,S}}{m_T} = 1 - \frac{x_{a,L} m_L + \frac{m_T - m_L}{m_L} m_L}{m_T} \\ &= 1 - x_{a,L} \frac{m_L}{m_T} - 1 + \frac{m_L}{m_T} = (1 - x_{a,L}) \frac{m_L}{m_T} \\ &= (1 - x_{a,L}) \left(1 - \frac{m_S}{m_T}\right) = (1 - x_{a,L}) \left(1 - \frac{m_{w,S}}{m_T}\right) \\ &= (1 - x_{a,L})(1 - \chi_{w,S}) \end{aligned} \quad (2.38)$$

Si on définit la capacité massique du liquide par :

$$c_L = (1 - x_{a,L}) \overline{c_{L,w}}(T, x) + x_{a,L} \overline{c_{L,a}}(T, x_{L,a}) \quad (2.39)$$

en combinant les expressions (2.35) (2.36) (2.37) (2.38) et (2.39), nous pouvons obtenir pour la variation d'enthalpie :

$$dh = L_w(T) d\chi_{w,S}(T) + \left(c_{w,S}^* \chi_{w,S}(T) + c_L (1 - \chi_{w,S}(T)) \right) dT \quad (2.40)$$

Ici, la chaleur latente s'exprime :

$$L_w(T) = L_w(T_{M,w}) + \Delta C_{M,w} (T - T_{M,w}) \quad (2.41)$$

avec :

$$\Delta C_{M,w} = c_{w,L}^* - c_{w,S}^* \quad (2.42)$$

Comme nous avons dans notre cas une solution binaire avec un liquidus rectiligne on a :

$$\begin{aligned}\frac{dT_M}{dx_0} &\Leftrightarrow \frac{T_{M,w} - T_M}{0 - x_0} = \frac{T_{M,w} - T}{0 - x_{a,L}} \\ &\Leftrightarrow x_{a,L} = x_0 \frac{T_{M,w} - T}{T_{M,w} - T_M}\end{aligned}\quad (2.43)$$

En différentiant l'équation (2.34) il vient alors :

$$d\chi_{w,S}(T) = \frac{x_0}{x_{a,L}(T)^2} dx_{a,L} = -\frac{T_{M,w} - T_M}{(T_{M,w} - T)^2} dT \quad (2.44)$$

d'où étant donné que pour $T = T_M$ $\chi_{w,S} = 0$:

$$\chi_{w,S}(T) = \frac{T_M - T}{T_{M,w} - T} \quad (2.45)$$

Enfin, cela mène à :

$$\frac{dh}{dT} = -\frac{T_{M,w} - T_M}{(T_{M,w} - T)^2} L_w(T) + c_{w,S}^* \left(1 - \frac{T_{M,w} - T_M}{T_{M,w} - T}\right) + c_L \frac{T_{M,w} - T_M}{T_{M,w} - T} \quad (2.46)$$

– $T > T_M$

$$\frac{dh}{dT} = (1 - x_0) \overline{c_{w,L}} + x_0 \overline{c_{a,L}} = c_L(x_0) \quad (2.47)$$

Avant de passer aux expressions finales de l'enthalpie pour le modèle binaire, nous allons poser quelques hypothèses et changements de variables pour simplifier l'expression de cette enthalpie.

Nous considérerons que la fraction massique de soluté est suffisamment faible et/ou que la capacité calorifique solide du soluté est suffisamment faible, pour que l'on puisse considérer une unique capacité solide égale à la capacité thermique du solvant pur à l'état solide.

$$(1 - x_0) c_{w,S}^* + x_0 c_{a,S}^* \approx c_{w,S}^* = c_S \quad (2.48)$$

De même nous allons considérer que la capacité liquide du mélange ne varie pas avec la concentration et qu'elle est sensiblement égale à celle du solvant.

$$c_L(x_0) = c_L \approx c_{w,L}^* \quad (2.49)$$

Afin de simplifier le système nous n'allons pas utiliser dans le modèle le paramètre L_E , mais la variation d'enthalpie au palier eutectique $\widetilde{L}_E$. Ce qui nous permet de nous affranchir des paramètres x_E et x_0 . Il en résulte que, dans le modèle, la température T_M peut varier indépendamment de la variation d'enthalpie $\widetilde{L}_E$. Ce faisant, x_0 n'est plus liée dans le modèle qu'à la température de liquidus T_M et des autres paramètres indépendants $T_{M,w}$ et T_E . Pour faire varier la fraction massique de soluté x_0 , il faut donc modifier les paramètres T_M et $\widetilde{L}_E$, qui sont indépendants dans ce modèle. On remplace le triplet de variables corrélées x_0 , x_E et L_E par deux variables indépendantes, T_M et $\widetilde{L}_E$.

$$\widetilde{L}_E = \frac{x_0}{x_E} L_E \quad (2.50)$$

$$T_M = T_{M,w} - \frac{x_0}{x_E} \times (T_{M,w} - T_E) \quad (2.51)$$

En intégrant les équations (2.10), (2.28), (2.46) et (2.47) sur leur domaine respectif on obtient les expressions suivantes :

$$- T < T_E$$

$$h(T) = c_S \times (T - T_E) + h_{ref} \quad (2.52)$$

$$- T_E < T \leq T_M$$

$$\begin{aligned} h(T) = \widetilde{L}_E + c_S \times \left[T - T_E + (T_{M,w} - T_M) \ln \left(\frac{T_{M,w} - T}{T_{M,w} - T_E} \right) \right] - c_L (T_{M,w} - T_M) \ln \left(\frac{T_{M,w} - T}{T_{M,w} - T_E} \right) \\ + (T_{M,w} - T_M) \left[\Delta C_{M,w} \ln \left(\frac{T_{M,w} - T}{T_{M,w} - T_E} \right) + L_w(T_{M,w}) \left(\frac{1}{T_{M,w} - T} - \frac{1}{T_{M,w} - T_E} \right) \right] + h_{ref} \end{aligned} \quad (2.53)$$

$$\text{avec } \Delta C_{M,w} = c_{w,L}^* - c_{L,S}^*.$$

$$- T > T_M$$

$$\begin{aligned} h(T) = \widetilde{L}_E + c_S \left[T_M - T_E + (T_{M,w} - T_M) \ln \left(\frac{T_{M,w} - T_M}{T_{M,w} - T_E} \right) \right] \\ - c_L \left[(T_{M,w} - T_M) \ln \left(\frac{T_{M,w} - T_M}{T_{M,w} - T_E} \right) - T + T_M \right] + h_{ref} \\ + (T_{M,w} - T_M) \left[\Delta C_{M,w} \ln \left(\frac{T_{M,w} - T_M}{T_{M,w} - T_E} \right) + L_{M,w}(T_{M,w}) \left(\frac{1}{T_{M,w} - T_M} - \frac{1}{T_{M,w} - T_E} \right) \right] \end{aligned} \quad (2.54)$$

Pour le taux liquide de la solution, χ_L , d'après les relations (2.6), (2.28) et (2.45) on a :

$$\text{Si } T = T_E :$$

$$\chi_L = \frac{m_L}{m_T} = \frac{x_0}{x_E} \frac{h(T_E, \chi) - h(T_E, \chi = 0)}{\frac{x_0}{x_E} L_E} = \frac{h(T_E, \chi) - h(T_E, \chi = 0)}{L_E} \quad (2.55)$$

$$\text{Si } T_E < T \leq T_M$$

$$\chi_L = 1 - \frac{m_S}{m_a + m_w} = 1 - \frac{m_{w,S}}{m_a + m_w} = 1 - \chi_{w,S} = \frac{T_{M,w} - T_M}{T_{M,w} - T} \quad (2.56)$$

Ici la référence en enthalpie est prise arbitrairement à la température T_E et à l'état solide.

$$h_{ref} = h(T_E, \chi_L = 0) = 0 \text{ J} \cdot \text{kg}^{-1} \quad (2.57)$$

Un exemple de profil enthalpique correspondant ici à la figure 2.5.

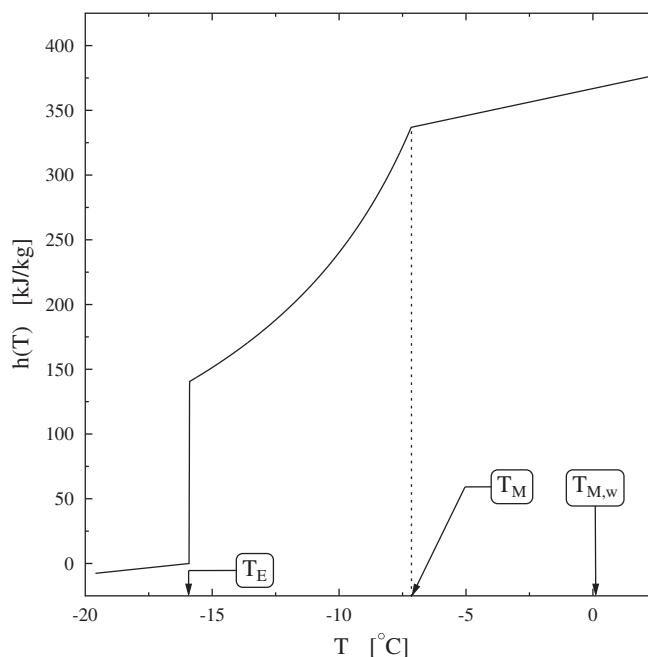


FIGURE 2.5 – Enthalpie d’une solution binaire ($H_2O - NH_4Cl$), $x_0 = 9,02\%$.

2.1.4 Résolution numérique des modèles thermodynamiques

Maintenant que nous avons un modèle mathématique et une géométrie pour représenter notre système, nous allons mettre en œuvre la résolution numérique de notre système d’équations. Pour cela nous avons choisi une approche volume fini car elle est structurellement adaptée aux schémas conservatifs [60].

De manière à simplifier la résolution numérique, nous avons considéré que la masse volumique ρ était constante, que l’on approxime par une valeur de la littérature ou que l’on détermine en pesant un volume donné de la substance à analyser. Cela nous permet de considérer un maillage fixe, car la masse étant constante, le volume est constant. Localement cela se traduit par : $\frac{\partial \rho h}{\partial t} = \rho \frac{\partial h}{\partial t}$. Par cette hypothèse, on fait porter une composante de l’erreur sur la masse ainsi que la non-prise en compte de la variation de masse volumique entre les états solide et liquide sur ρ . En effet, une erreur sur la masse d’un élément va d’un côté modifier son inertie si ρ porte l’erreur. Et si l’erreur vient du volume, alors les surfaces d’échanges seront incorrectes et le transfert thermique en sera modifié à moins que les propriétés de transfert, λ_S , λ_L ou α , ne soient modifiées car elles interviennent par un produit avec les surfaces dans les équations de transferts (2.3) et (2.4). Ce qui nous intéresse ce sont les propriétés énergétiques.

L’enthalpie est une fonction d’état, elle est donc la même si on réchauffe ou si on refroidit. Mais la cristallisation présente, en général, le phénomène de surfusion. Ce phénomène est complexe et relève d’une physique différente de celle des équilibres thermodynamiques. Nous avons décidé de ne pas aborder ce problème dans notre thèse même si c’est une thématique du laboratoire. Nous ne modéliserons donc que le cas de la fusion.

Pour ne pas dépendre des conditions opératoires, l'initialisation se fera donc dans un état homogène à la température $T_{début}$. Ceci est réalisable pour un échantillon en équilibre avec un thermostat de température constante et nous que nous choisirons un état initial en phase solide ($\chi_L^{iter=0} = 0$) à une température inférieure à T_E .

Pour résoudre numériquement le problème, nous avons fait une discrétisation spatiale du domaine, que nous réduisons à une demi-tranche du cylindre du fait de la symétrie de révolution. Cette demi-tranche est subdivisée en un ensemble de quadrilatères que l'on appelle maille. Dans le cas d'un maillage non-déformé, les sommets des quadrilatères ou nœuds du maillage, sont distribués de sorte que les quadrilatères soient des rectangles. Chaque nœud est repéré par sa distance à l'axe de révolution $r_{i,j}$ et sa hauteur $z_{i,j}$. Ce sont ces relations (2.58) et (2.59) que l'on modifie si on veut déformer en partie le maillage. Celles que l'on donne correspondent à un maillage isovolume. Le centre de chaque maille étant le barycentre des 4 nœuds délimitant la maille.

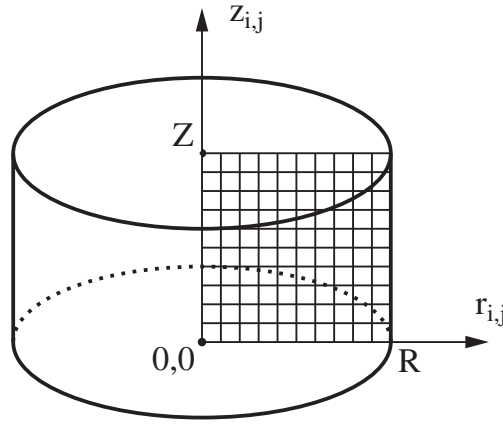


FIGURE 2.6 – Maillage par symétrie de révolution

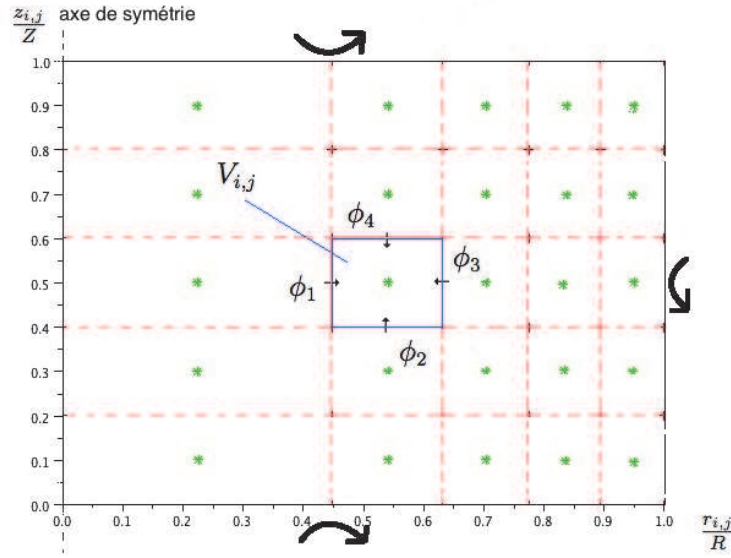
$$z_{i,j} = Z \frac{j-1}{NJ} \quad (2.58)$$

$$r_{i,j} = R \sqrt{\frac{i-1}{NI}} \quad (2.59)$$

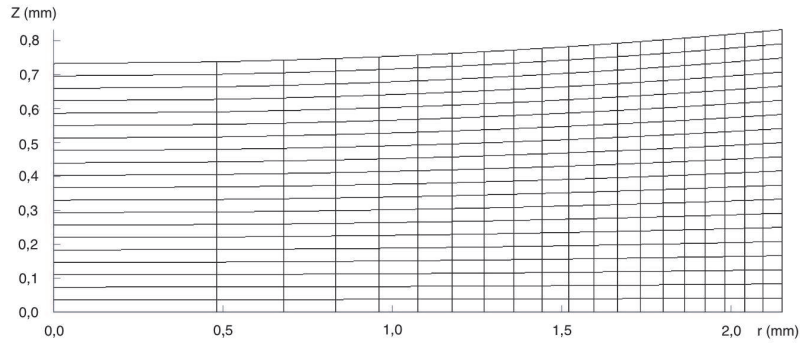
Avec NI et NJ le nombre de nœuds suivant l'axe radial et axial.

Nous voulons pouvoir éventuellement déformer le maillage. Nous définissons donc les propriétés géométriques à partir de relations les plus générales possibles. Les surfaces de chaque face, $S_{i,j,k}$, sont calculées à l'aide du premier théorème de Guldin. Les cosinus venant des produits scalaires entre les vecteurs distances reliant les centres de mailles et les normales

aux surfaces, $\cos(\vec{d}, \vec{n})$, sont calculés avec le théorème d'Al-Kashi. Et les volumes de chaque maille, $V_{i,j}$, seront la somme des volumes générés par révolution des deux triangles contenus dans la maille autour de l'axe centrale. Ces volumes des triangles sont calculés à l'aide du deuxième théorème de Guldin, leurs surfaces par la formule de Héron.



(a) Schéma récapitulatif du maillage



(b) Exemple de maillage déformé

FIGURE 2.7 – Exemples de différents maillages

La discrétisation temporelle du problème est faite selon la méthode d'Euler explicite avec un pas de temps que l'on prend constant Δt entre chaque itération $i \text{ ter}$:

$$\frac{\partial h}{\partial t} \approx \frac{h^{iter+1} - h^{iter}}{\Delta t} \quad (2.60)$$

La résolution par l'approche volume fini implique que l'on intègre l'équation (2.1) sur le volume de chaque maille avant de la résoudre. L'intérêt de cette approche est qu'elle nous

permet d'utiliser l'équation (1.4), qui devient pour chaque maille (i,j) :

$$\rho V_{i,j} \frac{h^{iter+1} - h^{iter}}{\Delta t} = \sum_{k=1}^4 \phi_k \quad (2.61)$$

Ainsi dans chacun de ces volumes, on va résoudre l'équation (2.61) à chaque pas de temps ou itération. Pour ce faire nous découplons le calcul des flux de la résolution du bilan (2.61). Nous présenterons ensuite des calculs dans un cas de maillage plus général. Chaque itération, commence avec l'enthalpie h^{iter} et la température T^{iter} de la maille. De là, on détermine le taux de liquide avec les équations (2.55), (2.56) ou (2.6). Avec ce taux on détermine les conductivités des mailles avec (2.64). Ensuite, on peut calculer les flux entre deux mailles avec la loi de Fourier discrétisée.

$\forall (i, j) \in [2; NI-1] \times [2; NJ-1]$ on a :

$$\begin{aligned} \phi_1 = \phi_{(i-1,j) \Rightarrow (i,j)} &= -\lambda_{S_{i,j,1}} S_{i,j,1} \frac{T_{i-1,j}^{iter} - T_{i,j}^{iter}}{\|\vec{d}\|} \cos(\vec{d}, \vec{n}) \\ \phi_2 = \phi_{(i,j-1) \Rightarrow (i,j)} &= -\lambda_{S_{i,j,2}} S_{i,j,2} \frac{T_{i,j-1}^{iter} - T_{i,j}^{iter}}{\|\vec{d}\|} \cos(\vec{d}, \vec{n}) \\ \phi_3 = \phi_{(i+1,j) \Rightarrow (i,j)} &= -\lambda_{S_{i,j,3}} S_{i,j,3} \frac{T_{i+1,j}^{iter} - T_{i,j}^{iter}}{\|\vec{d}\|} \cos(\vec{d}, \vec{n}) \\ \phi_4 = \phi_{(i,j+1) \Rightarrow (i,j)} &= -\lambda_{S_{i,j,4}} S_{i,j,4} \frac{T_{i,j+1}^{iter} - T_{i,j}^{iter}}{\|\vec{d}\|} \cos(\vec{d}, \vec{n}) \end{aligned} \quad (2.62)$$

Avec $\vec{d}$ le vecteur partant du centre (i, j) allant au centre de l'autre maille, et $\lambda_{S_{i,j,k}}$ la conductivité que l'on prend à la frontière entre les mailles pour satisfaire à l'égalité entre les flux entrant et sortant d'une maille par une même surface. Pour déterminer la conductivité à la frontière $\lambda_{S_{i,j,k}}$, nous allons procéder par une analogie électrique de deux résistances en série entre les centres de deux mailles i_1, j_1 et i_2, j_2 séparées par une distance $d_1 + d_2$:

$$\frac{\|\vec{d}\|}{\lambda_{S_{i,j,k}}} = \frac{d_1}{\lambda_{i_1,j_1}} + \frac{d_2}{\lambda_{i_2,j_2}} \quad (2.63)$$

$\lambda_{i,j}$ est la conductivité thermique de la maille i, j, dans les zones diphasiques on considère qu'elle est linéairement dépendante des conductivités en phase liquide ou solide, proportionnellement à la taux liquide dans la maille (2.64) [61].

$$\lambda_{i,j} = (1 - \chi_{L,i,j}) \lambda_S + \chi_{L,i,j} \lambda_L \quad (2.64)$$

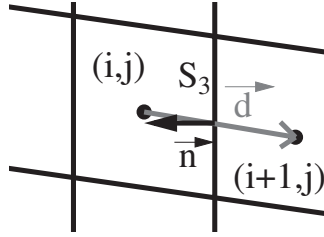


FIGURE 2.8 – Illustration du calcul des flux internes (cas général)

Pour les mailles aux frontières du domaine où il faut aussi tenir compte des conditions aux limites (2.4) entre la maille (i, j) et le creuset, le flux s'exprime alors sous la forme (voir figure 2.9) :

$$\begin{aligned}\phi_{creuset \Rightarrow (i,1)} &= S_2 \frac{T_{i,1}^{iter} - T_p^{iter}}{\frac{\|\vec{d}\| \times \cos(\vec{d}, \vec{n})}{\lambda_{i,1}} + \frac{1}{\alpha_1}} \quad \forall i \in [1; NI] \\ \phi_{creuset \Rightarrow (NI,j)} &= S_3 \frac{T_{NI,j}^{iter} - T_p^{iter}}{\frac{\|\vec{d}\| \times \cos(\vec{d}, \vec{n})}{\lambda_{NI,j}} + \frac{1}{\alpha_2}} \quad \forall j \in [1; NJ] \\ \phi_{creuset \Rightarrow (i,NJ)} &= S_4 \frac{T_{i,NJ}^{iter} - T_p^{iter}}{\frac{\|\vec{d}\| \times \cos(\vec{d}, \vec{n})}{\lambda_{i,NJ}} + \frac{1}{\alpha_3}} \quad \forall j \in [1; NJ]\end{aligned}\tag{2.65}$$

Compte tenu de la symétrie du problème, le flux radial au centre est nul, soit $\phi_1 = 0$ si $i=0$.

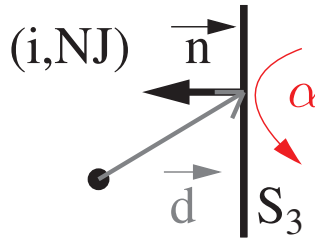


FIGURE 2.9 – Illustration du calcul des flux sur la face latérale (maillage curviligne)

Suivant le principe rappelé au chapitre 1, une fois les flux calculés on résout le bilan pour obtenir h^{iter+1} de la maille qui, en résolvant l'équation d'état correspondante, détermine la nouvelle température T^{iter+1} , ainsi que la taux liquide χ_L^{iter+1} . Ces propriétés serviront à calculer les flux lors de la prochaine itération.

$$h_{i,j}^{iter+1} = h_{i,j}^{iter} + \frac{dt}{\rho \times V_{i,j}} \sum_{k=1}^4 \phi_k \tag{2.66}$$

Le système d'équations est donc :

$$\left\{ \begin{array}{l} h_{i,j}^{iter+1} = h_{i,j}^{iter} + \frac{dt}{\rho \times V_{i,j}} \sum_{k=1}^4 \phi_k \\ \text{avec } \phi_{(k) \Rightarrow (i,j)} = -\lambda_{S_{i,j,k}} S_{i,j,k} \frac{T_k^{iter} - T_{i,j}^{iter}}{\|\vec{d}\|} \cos(\vec{d}, \vec{n}) \\ \text{ou } \phi_{creuset \Rightarrow (i,j)} = S_{i,j,k} \frac{T_{i,j}^{iter} - T_p^{iter}}{\frac{\|\vec{d}\| \times \cos(\vec{d}, \vec{n})}{\lambda_{i,j}} + \frac{1}{\alpha}} \end{array} \right.$$

Un schéma récapitulatif du modèle de fusion est donné figure 2.10.

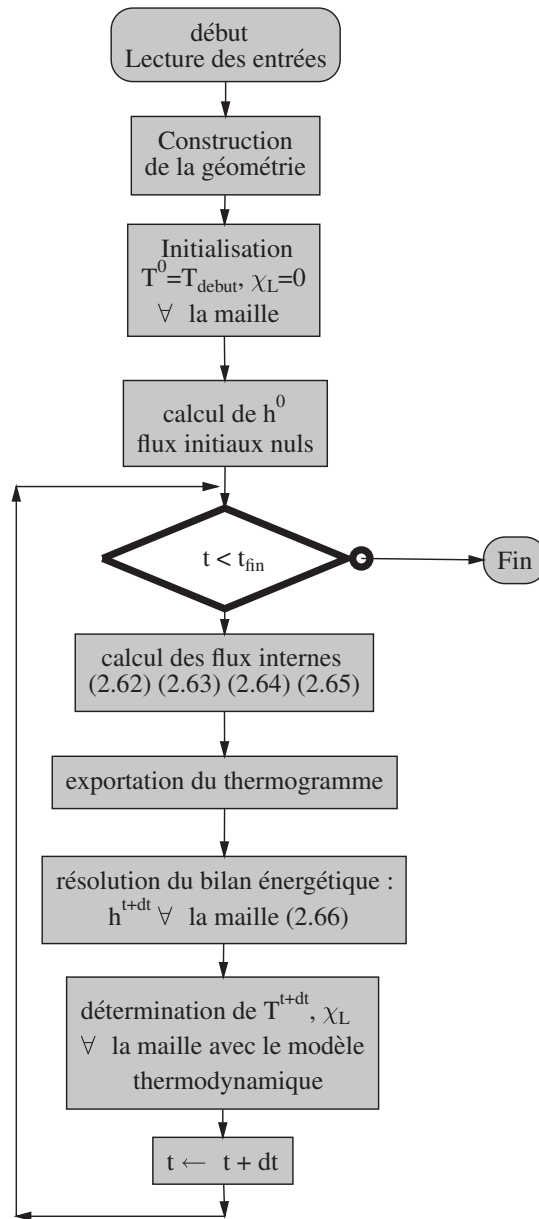


FIGURE 2.10 – Schéma de résolution du modèle direct

2.2 Validation expérimentale du modèle direct

Dans cette section nous allons présenter des cas expérimentaux auxquels nous avons confronté le modèle 2D-cylindrique.

D’abord on traitera le cas du corps pur. Nous avons testé notre modèle pour plusieurs vitesses avec un échantillon d’eau de 11,4 mg et un échantillon d’acide benzoïque de 3,7 mg. On obtient une excellente concordance entre les résultats expérimentaux et les thermogrammes simulés. Les paramètres qui ont servis pour ces simulations directes, proviennent de la littérature ou de nos identifications quand elles n’étaient pas disponibles dans la littérature. Ces valeurs figures dans les tableaux 2.1 et 2.2. On remarque en outre que la vitesse de réchauffement n’affecte presque pas la concordance du modèle avec l’expérience. Il faut noter qu’on ne modélise pas la convection au sein de l’échantillon, ni le mouvement de la partie solide dans le fluide, et comme la convection augmente avec la vitesse de réchauffement, cela pourrait expliquer les petits écarts pour les vitesses les plus élevées, au niveau des sommets et en fin de thermogramme.

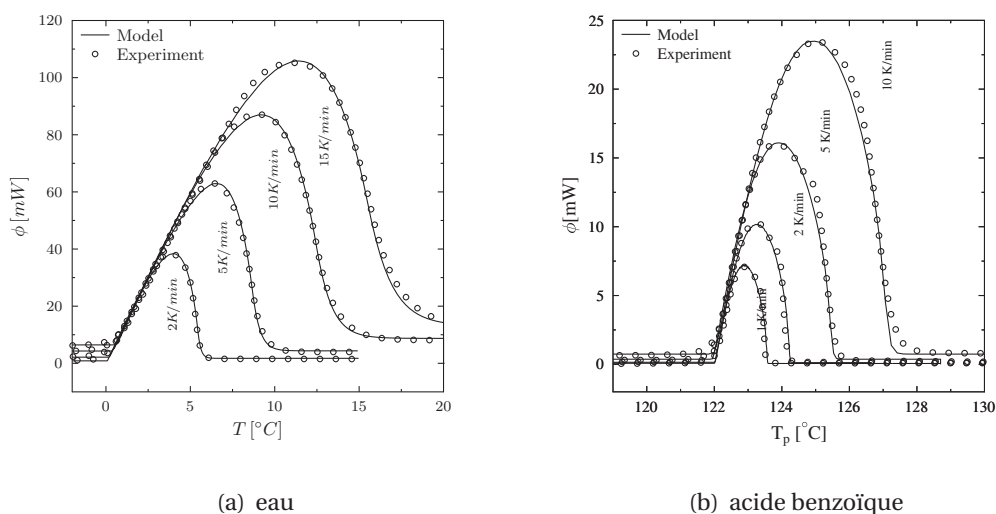
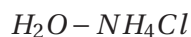


FIGURE 2.11 – Validation expérimentale du modèle 2D dans le cas du corps pur.

Pour le cas du binaire, nous avons comparé le modèle à plusieurs expériences de calorimétrie sur des solutions de $H_2O - NH_4Cl$ et de $H_2O - KCl$ qui obéissent à la relation enthalpie-température de la sous-section 2.1.3, et qui ont été réalisés dans ce laboratoire. Ici aussi on observe une bonne concordance entre résultats numériques et expérimentaux, les écarts au niveau du début du pic eutectique peuvent s’expliquer par la présence d’impuretés. Voici les échantillons et les thermogrammes correspondants :



- une solution à 0,48% de masse 9,2 mg figure 2.13(a) et 2.13(b)
- une solution à 5,0% de masse 11,9 mg figure 2.13(e)

- une solution à 9,02% de masse 9,6 mg figure 2.13(d)
- une solution à 10,0% de masse 9,1 mg figure 2.13(f)

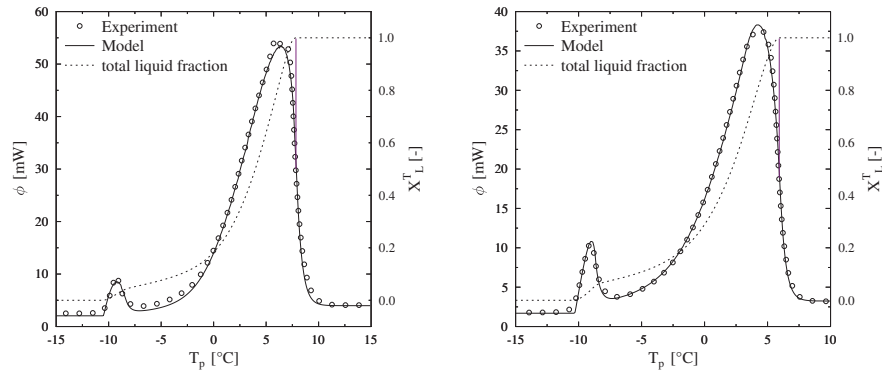
$H_2O - KCl$

- une solution à 0,615% de masse 13,9 mg figure 2.12(a)
- une solution à 1,13% de masse 9,2 mg figure 2.12(b)
- une solution à 2,007% de masse 11,7 mg figure 2.12(c)

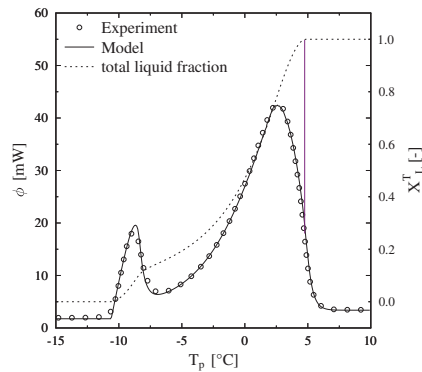
Sur les figures relatives à des solutions binaires, 2.12 et 2.13, nous avons également donné l'évolution du taux liquide total des échantillons qui est calculée par :

$$\chi_L^T = \frac{1}{V} \iiint_V \chi_L dV \quad (2.67)$$

On y remarque que la fin de fusion (dernière particule de glace disparaissant au centre de l'échantillon), qui correspond à $\chi_L^T = 1$, semble être systématiquement atteinte au moment où la partie descendante de la courbe présente un point d'inflexion [62].

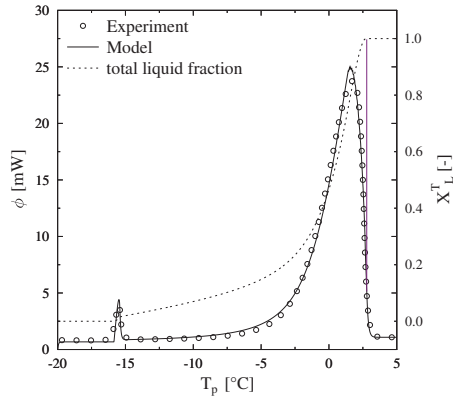
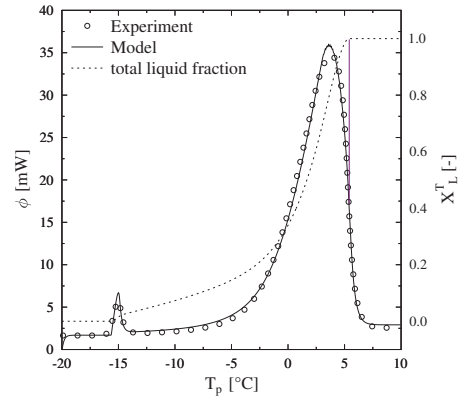
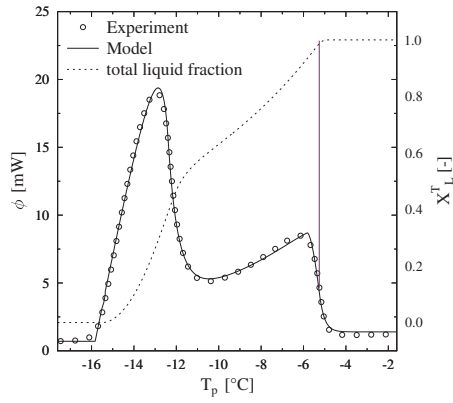
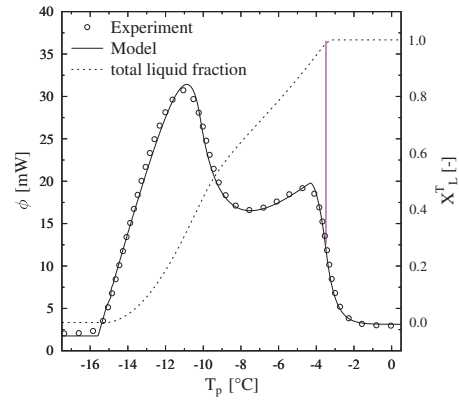
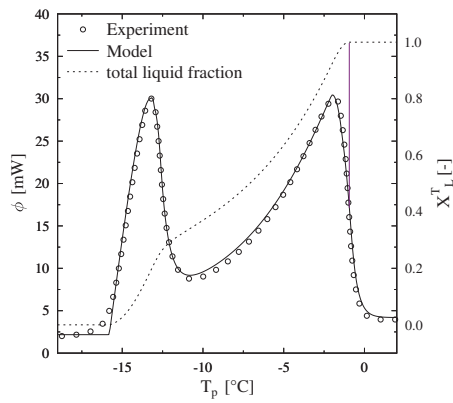
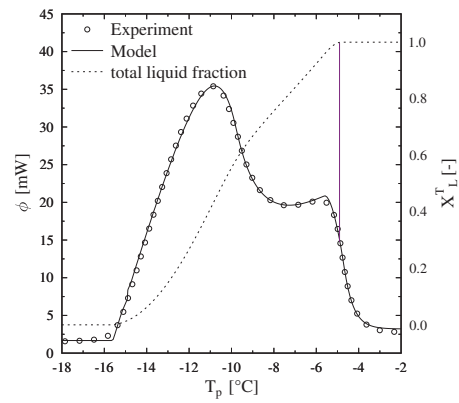


(a) $H_2O - KCl$ à 0,615%, $\beta = 5 \text{ K} \cdot \text{min}^{-1}$ (b) $H_2O - KCl$ à 1,13%, $\beta = 5 \text{ K} \cdot \text{min}^{-1}$



(c) $H_2O - KCl$ à 2,007%, $\beta = 5 \text{ K} \cdot \text{min}^{-1}$

FIGURE 2.12 – Validation expérimentale du modèle 2D relatif aux solutions de $H_2O - KCl$.

(a) $H_2O - NH_4Cl$ à 0,48%, $\beta = 2 \text{ K} \cdot \text{min}^{-1}$ (b) $H_2O - NH_4Cl$ à 0,48%, $\beta = 5 \text{ K} \cdot \text{min}^{-1}$ (c) $H_2O - NH_4Cl$ à 9,02%, $\beta = 2 \text{ K} \cdot \text{min}^{-1}$ (d) $H_2O - NH_4Cl$ à 9,02%, $\beta = 5 \text{ K} \cdot \text{min}^{-1}$ (e) $H_2O - NH_4Cl$ à 5,0%, $\beta = 5 \text{ K} \cdot \text{min}^{-1}$ (f) $H_2O - NH_4Cl$ à 10,0%, $\beta = 5 \text{ K} \cdot \text{min}^{-1}$ FIGURE 2.13 – Validation expérimentale du modèle 2D relatif aux solutions de $H_2O - NH_4Cl$.

Le modèle reflète donc fidèlement le transfert de chaleur au sein d'un échantillon placé dans un calorimètre de DSC, cela quelle que soit la valeur de la vitesse prise dans la gamme usuelle du domaine.

2.3 Analyses paramétriques

Dans cette section, nous allons étudier l'influence de chaque paramètre sur le thermogramme, comment la variation d'un paramètre le déforme, pour les deux cas, le corps pur et la solution binaire, toutes choses égales par ailleurs. Comme le but de ce modèle est de recréer un thermogramme à partir de propriétés physiques, ces informations vont nous permettre d'évaluer le rôle de chaque paramètre. Le problème de erreurs sera abordé dans le chapitre concernant les résultats de l'inversion 4.3.2.

Nous allons considérer un échantillon cylindrique, de rayon $R = 2,15 \text{ mm}$ de hauteur $Z = 0,78 \text{ mm}$ comme celui des cellules de DSC du Perkin-Elmer, et qui est réchauffé à différentes vitesses $\beta = 1, 2, 5$ et $10 \text{ K} \cdot \text{min}^{-1}$.

2.3.1 Cas du corps pur

De manière à obtenir une bonne idée de l'influence de chaque paramètre, nous avons étudié quatre matériaux différents, qui sont tous bien répertoriés et connus dans la littérature, chacun ayant des propriétés très différentes. Ainsi, nous considérons l'eau, avec sa grande chaleur latente et une différence notable entre ses capacités massiques liquide et solide, ensuite vient l'hexadécane qui est une paraffine avec une bonne chaleur latente mais de faibles conductivités thermiques, suivi du mercure, un métal présentant de fortes conductivités thermiques et finalement l'acide palmitique, un acide gras. Les paramètres correspondants que nous avons pris sont donnés dans le tableau 2.1.

Tableau 2.1 – Paramètres des corps purs [63] [64]

Matériaux	λ_S ($\text{W} \cdot \text{m}^{-1} \cdot \text{K}^{-1}$)	λ_L ($\text{W} \cdot \text{m}^{-1} \cdot \text{K}^{-1}$)	c_S ($\text{J} \cdot \text{K}^{-1} \cdot \text{kg}^{-1}$)	c_L ($\text{J} \cdot \text{K}^{-1} \cdot \text{kg}^{-1}$)	T_M (K)	L_M ($\text{J} \cdot \text{kg}^{-1}$)
eau	2,1	0,6	2060	4180	273,15	$3,33 \cdot 10^5$
Hexadecane	0,4	0,21	1753	2000	291,15	$2,30 \cdot 10^5$
Mercure	29	8	140	140	234,35	$1,14 \cdot 10^4$
Acide palmitique	0,5	0,162	1900	2800	337,15	$1,85 \cdot 10^5$

Pour chacune des études qui vont suivre, nous avons modifié la valeur du paramètre étudié proportionnellement à la valeur de référence (voir tableau 2.1). Nous avons choisi des écarts de $\pm 10\text{-}20\%$ ou $1\text{-}2 \text{ K}$, afin de rendre visible l'influence des paramètres sur les thermogrammes.

2.3.1.1 Conductivités thermiques

Notre premier objectif est de déterminer l'influence des transferts thermiques dans l'échantillon. Les résultats figurent dans les figures 2.14 et 2.15 où sont précisées les variations relatives des paramètres. Il est clair que, dans le cas du corps pur, la conductivité solide n'a pas d'influence sur le thermogramme (figure 2.14), car la phase solide est homogène en température, voir figure 2.40. Au contraire, le thermogramme est modifié par la conductivité thermique liquide, figure 2.15. On peut le comprendre car, quand la fusion a lieu, une couche de liquide existe aux frontières de l'échantillon. Pour fondre le reste de l'échantillon, pour faire progresser le front de fusion, il faut tenir compte de la conductivité de cette couche liquide. Si elle est peu conductrice, alors la dynamique de changement de phase entraîne que la valeur du flux thermique échangé s'en trouvera diminuée.

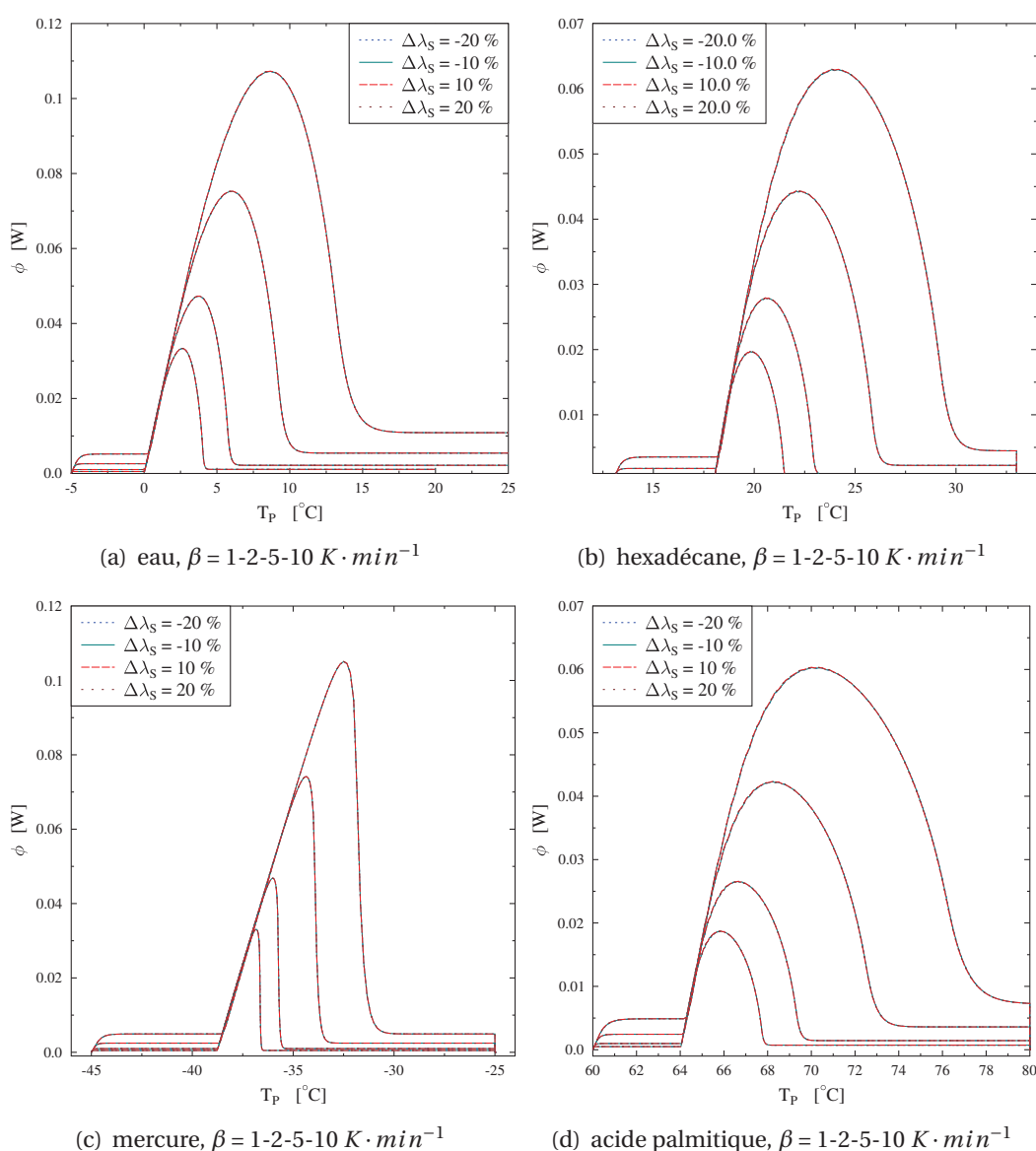


FIGURE 2.14 – Influence de la conductivité thermique solide sur le thermogramme

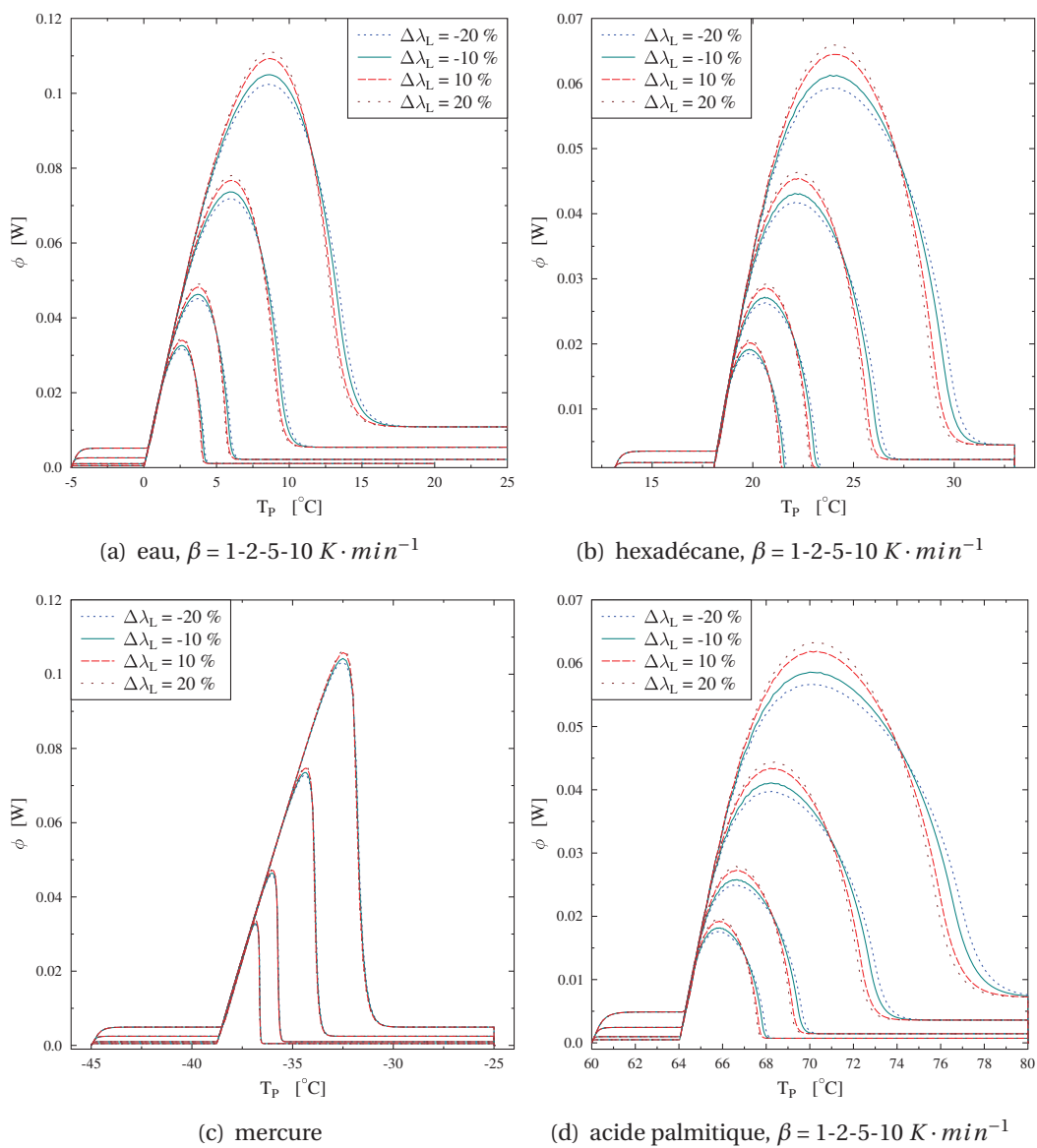


FIGURE 2.15 – Influence de la conductivité thermique liquide sur le thermogramme

2.3.1.2 Moyenne surfacique des coefficients d'échanges équivalents

Pour un corps pur la moyenne surfacique des coefficients équivalents modifie la pente à l'origine selon l'équation (1.16).

$$\alpha = \sum_{i=1}^3 \frac{\alpha_i S_i}{\sum_{j=1}^3 S_j} \quad (2.68)$$

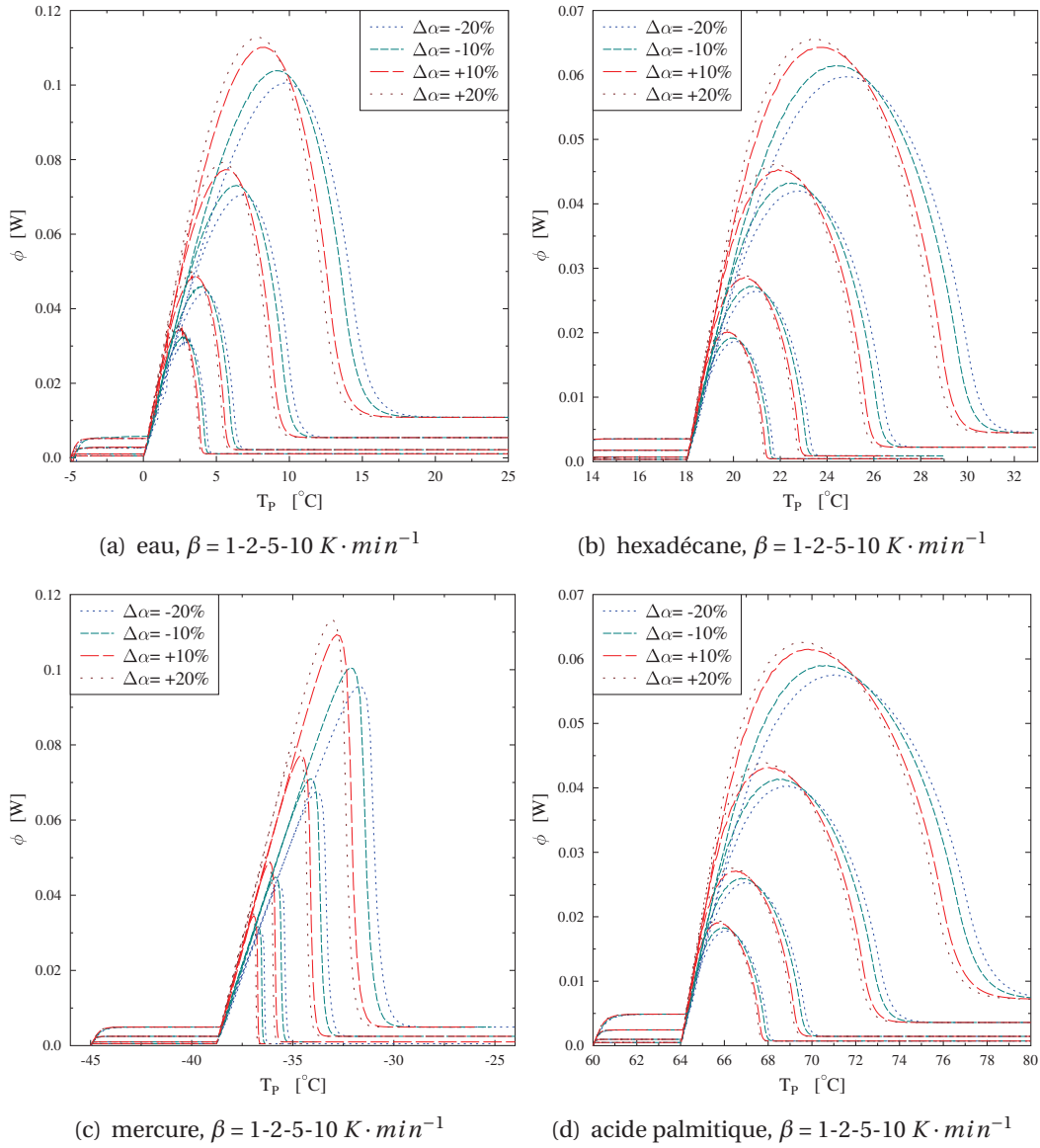


FIGURE 2.16 – Influence de la moyenne surfacique des coefficients d'échanges équivalents sur le thermogramme

2.3.1.3 Capacités thermiques massiques

En appliquant la même procédure, nous allons évaluer l'influence sur le thermogramme de la chaleur sensible. Comme on peut le voir sur la figure 2.17, la capacité solide n'influe qu'avant le pic, car lors de la fusion la phase solide est homogène en température. Il n'y a plus d'évolution thermique dans la phase solide, par lequel c_s pourrait influencer sur le thermogramme. La capacité liquide, figure 2.18, influe principalement après le pic de fusion mais aussi sur la partie descendante de ce pic, car la phase liquide n'y est pas homogène en température. La phase liquide étant soumise à un gradient thermique, la capacité liquide va d'autant influencer sur l'échange d'énergie que le gradient est élevé.

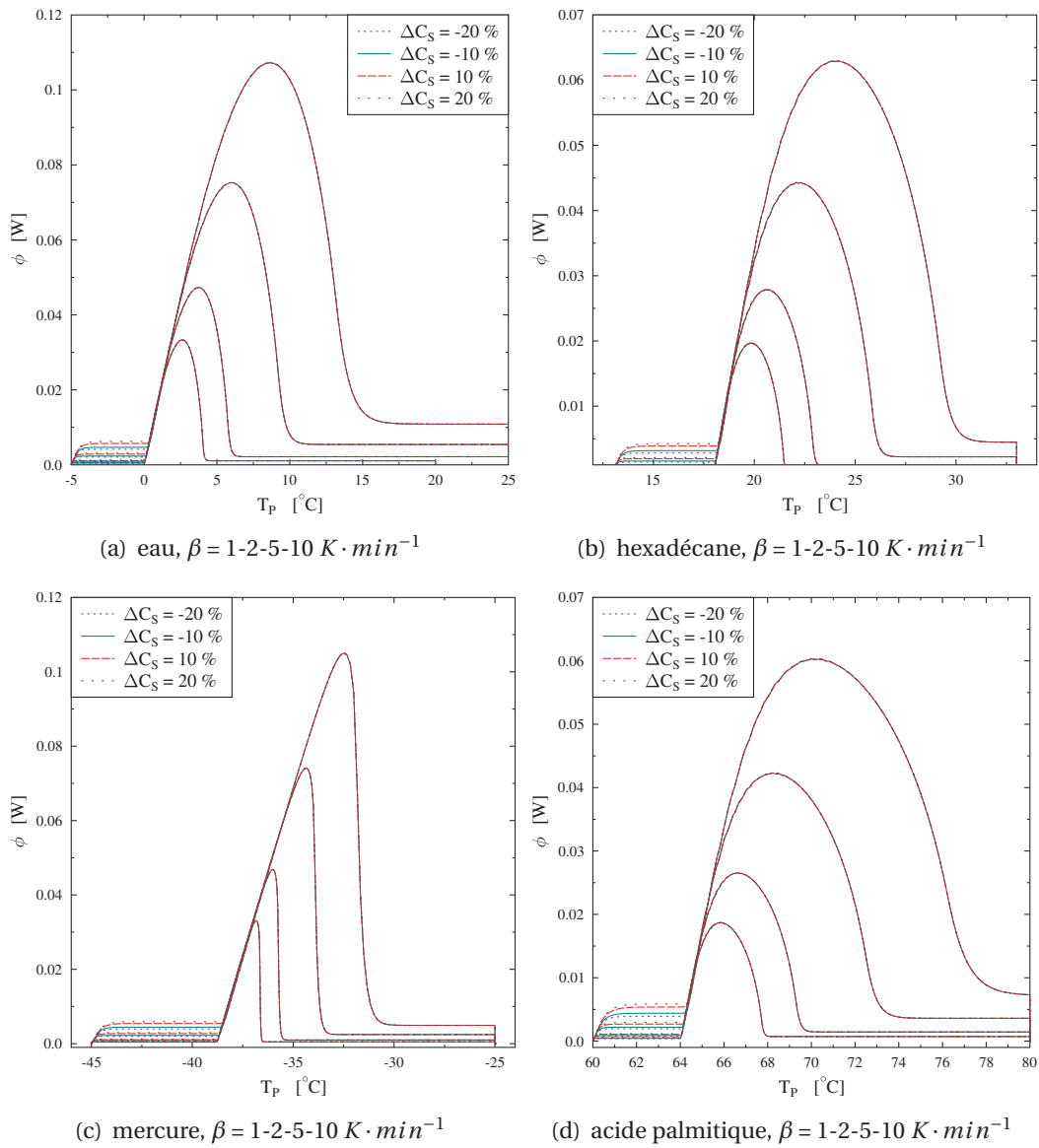


FIGURE 2.17 – Influence des capacités thermiques spécifiques solides sur le comportement des corps purs.

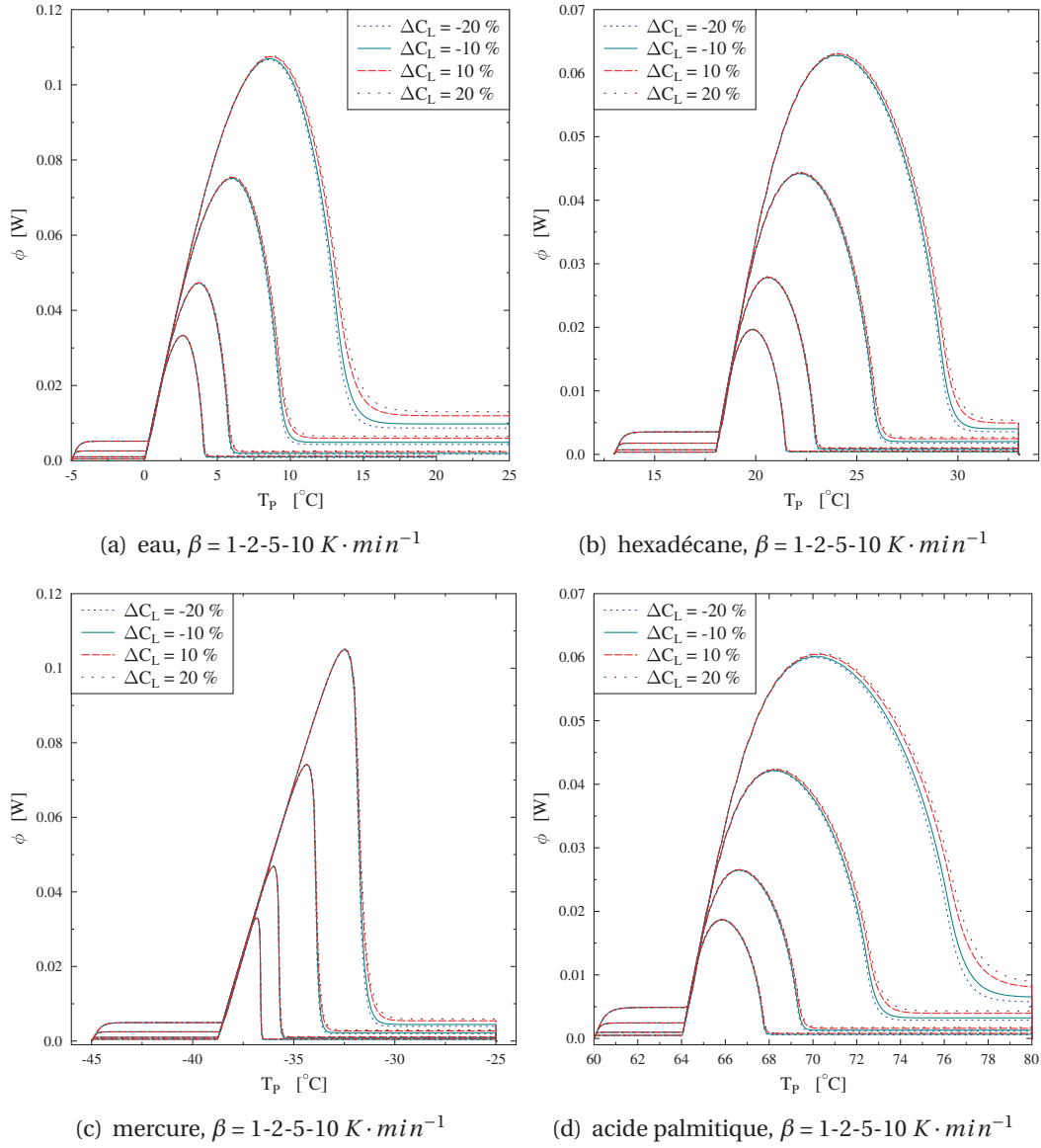


FIGURE 2.18 – Influence des capacités thermiques spécifiques liquides sur le comportement des corps purs.

2.3.1.4 Température de fusion

Ce paramètre est vraiment caractéristique de la fusion, et par conséquent il devrait avoir une grande influence sur le thermogramme. Comme on le constate sur la figure 2.19 ce paramètre engendre une grande modification du thermogramme. Mais dans le cas du corps pur cette influence n'est qu'une simple translation du thermogramme sur l'axe des abscisses. Si la température de fusion est décalée par un mauvais étalonnage (par exemple), on reproduit tout de même la bonne forme du thermogramme.

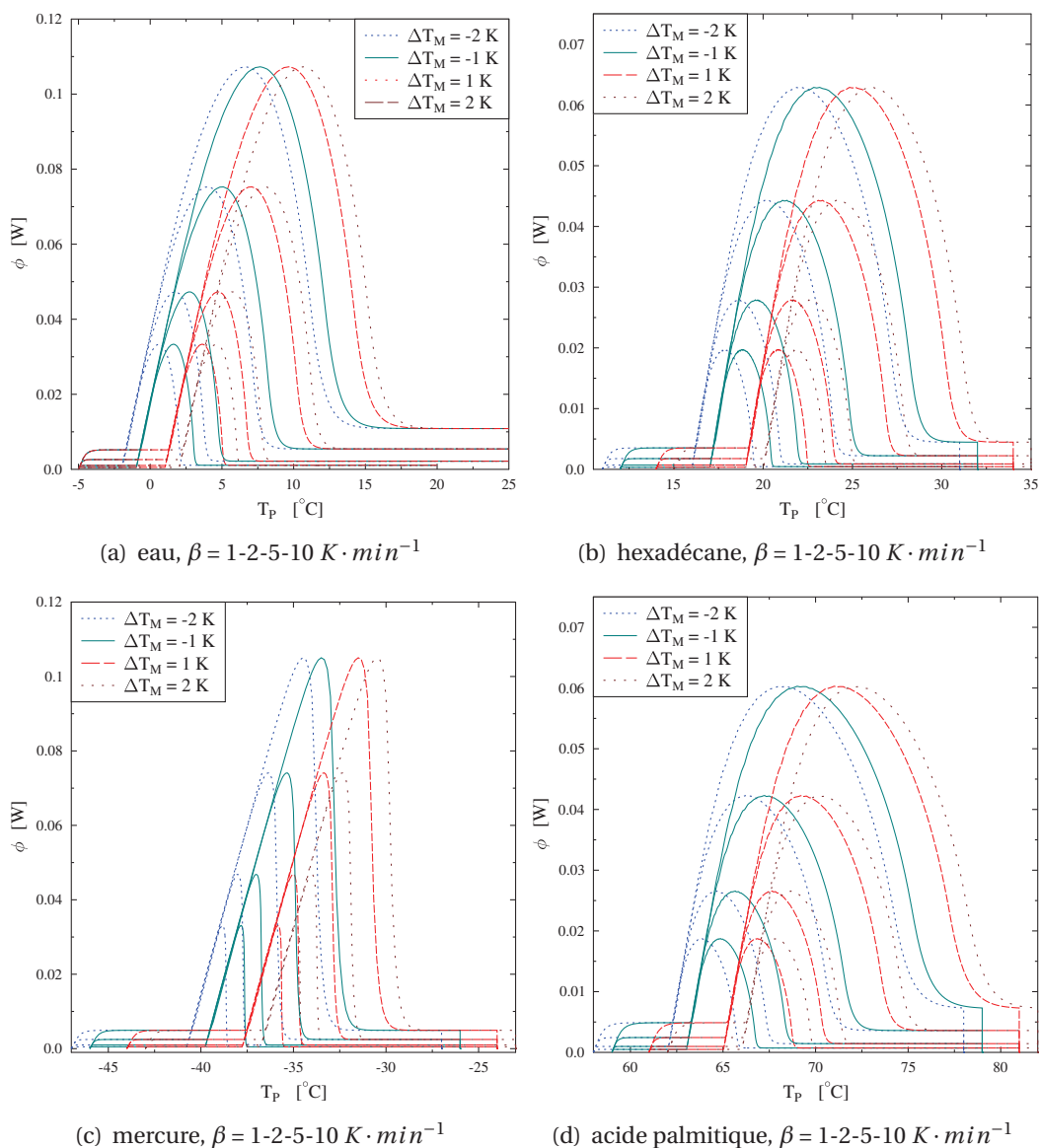


FIGURE 2.19 – Influence de la température de fusion pour un corps pur.

2.3.1.5 Chaleur latente du corps pur

Ici encore, il est connu que ce paramètre est d'une importance primordiale pendant le processus de changement de phase. Comme on peut le voir à la figure 2.20, le thermogramme est grandement influencé par toute modification de la chaleur latente. Ce résultat n'est pas surprenant car le pic représente la quantité de chaleur totale associée au phénomène thermique [34] et la chaleur latente est la quantité de chaleur mise en jeu par le changement de phase. L'égalité entre la surface de ce pic et la valeur de la chaleur latente est admise [65] [66] si ce n'est le problème de délimitation de la limite inférieure de ce pic, dans le cas où les capacités solide et liquide sont différentes. Elle ne modifie pas la pente à l'origine car celle-ci est fixée par la résistance thermique à la surface de l'échantillon (1.16).

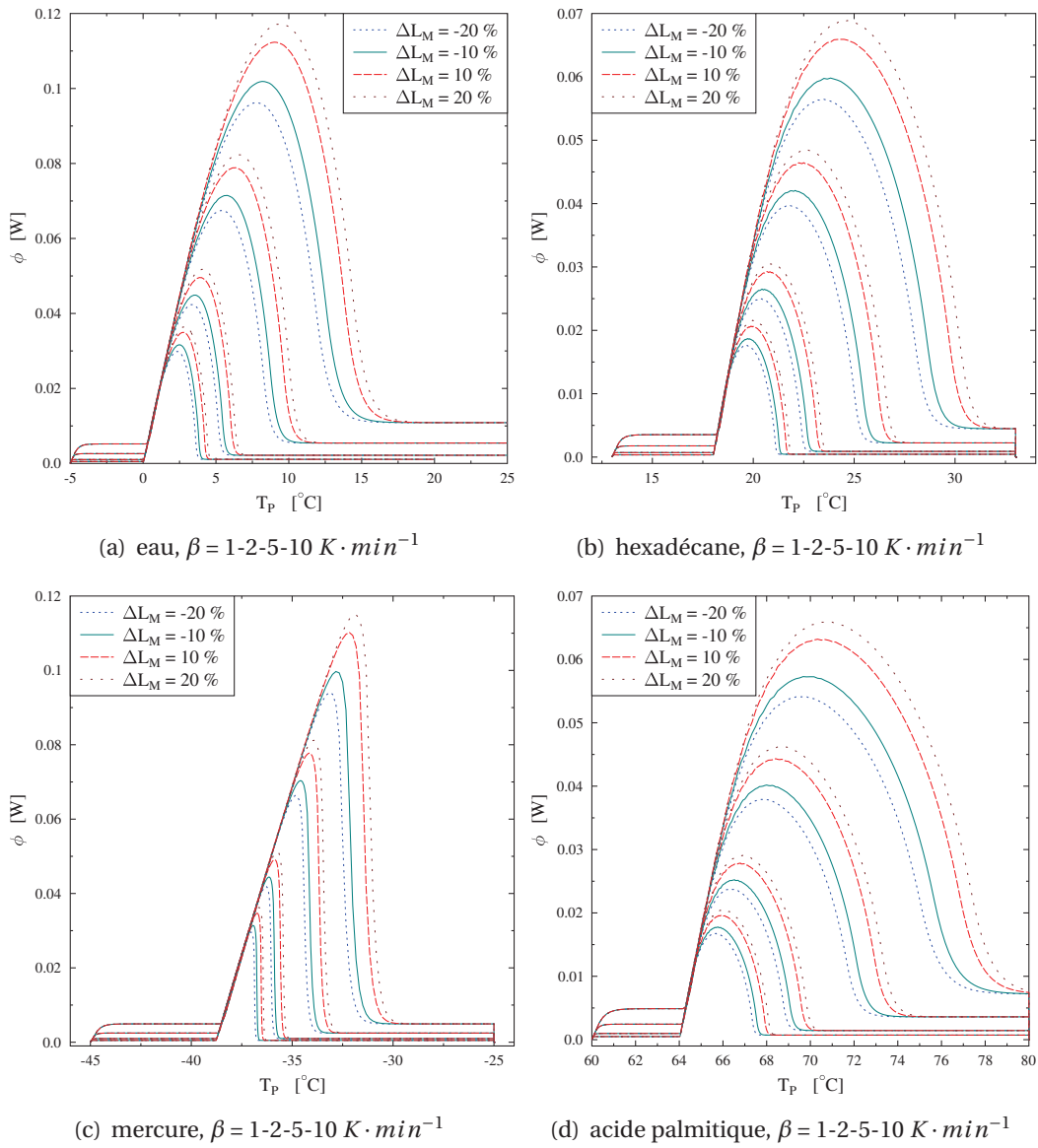


FIGURE 2.20 – Influence de la chaleur latente sur le comportement du corps pur.

2.3.2 Cas de solutions binaires

Nous allons développer maintenant la même analyse, mais pour une solution binaire. Ici, on considère une solution de H_2O-NH_4Cl pour trois fractions massiques en soluté différentes : $x_0 = 2,5\%, 5\%$ and 10% . Ces solutions étant majoritairement aqueuses, nous considérerons les propriétés de l'eau, tableau 2.1, auxquelles on adjoint des propriétés thermodynamiques propres au binaire, tableau 2.2, nous avons pris pour $\widetilde{L}_E$ et T_M les valeurs que nous avons déduites des résultats d'identifications du chapitre 4.

Tableau 2.2 – Propriétés thermodynamiques relatives aux solutions de $H_2O - NH_4Cl$

Matériaux	$L_{M,w}$ ($kJ kg^{-1}$)	$c_{w,S}^*$ ($J \cdot kg^{-1} \cdot K^{-1}$)	$c_{w,L}^*$ ($J \cdot kg^{-1} \cdot K^{-1}$)	T_E ($^{\circ}C$)	$\widetilde{L}_E^1$ ($kJ \cdot kg^{-1}$)	T_M^2 ($^{\circ}C$)
-	333	2060	4180	-15,9	8,2	-0,46
$H_2O/NH_4Cl : 0,57\%$	333	2060	4180	-15,9	35,9	-2,02
$H_2O/NH_4Cl : 2,5\%$	333	2060	4180	-15,9	74,4	-4,03
$H_2O/NH_4Cl : 10\%$	333	2060	4180	-15,9	144	-8,05

1. Ces valeurs ont été déterminées à partir de nos identifications sur des thermogrammes expérimentaux 4.2

2. voir note 1

2.3.2.1 Conductivités thermiques

Les figures 2.21 et 2.22, nous présentent les différents résultats. Il semble que le thermogramme ne soit pas modifié pour des faibles vitesses de réchauffement. En effet, ce n'est que pour une vitesse de réchauffement supérieure à $2 \text{ K} \cdot \text{min}^{-1}$ que chacune des conductivités, solide comme liquide, petite modification des thermogrammes apparaissent sur le sommet d'un des pics. Contrairement au cas du corps pur où seule la conductivité liquide influe, la conductivité solide influe elle aussi sur le sommet du premier pic (eutectique) car la phase solide n'est pas homogène en température et elle est prédominante au cours de la transformation. La conductivité liquide influe sur le deuxième pic car la phase liquide est majoritaire au cours de cette fusion progressive avant le un retour au régime permanent.

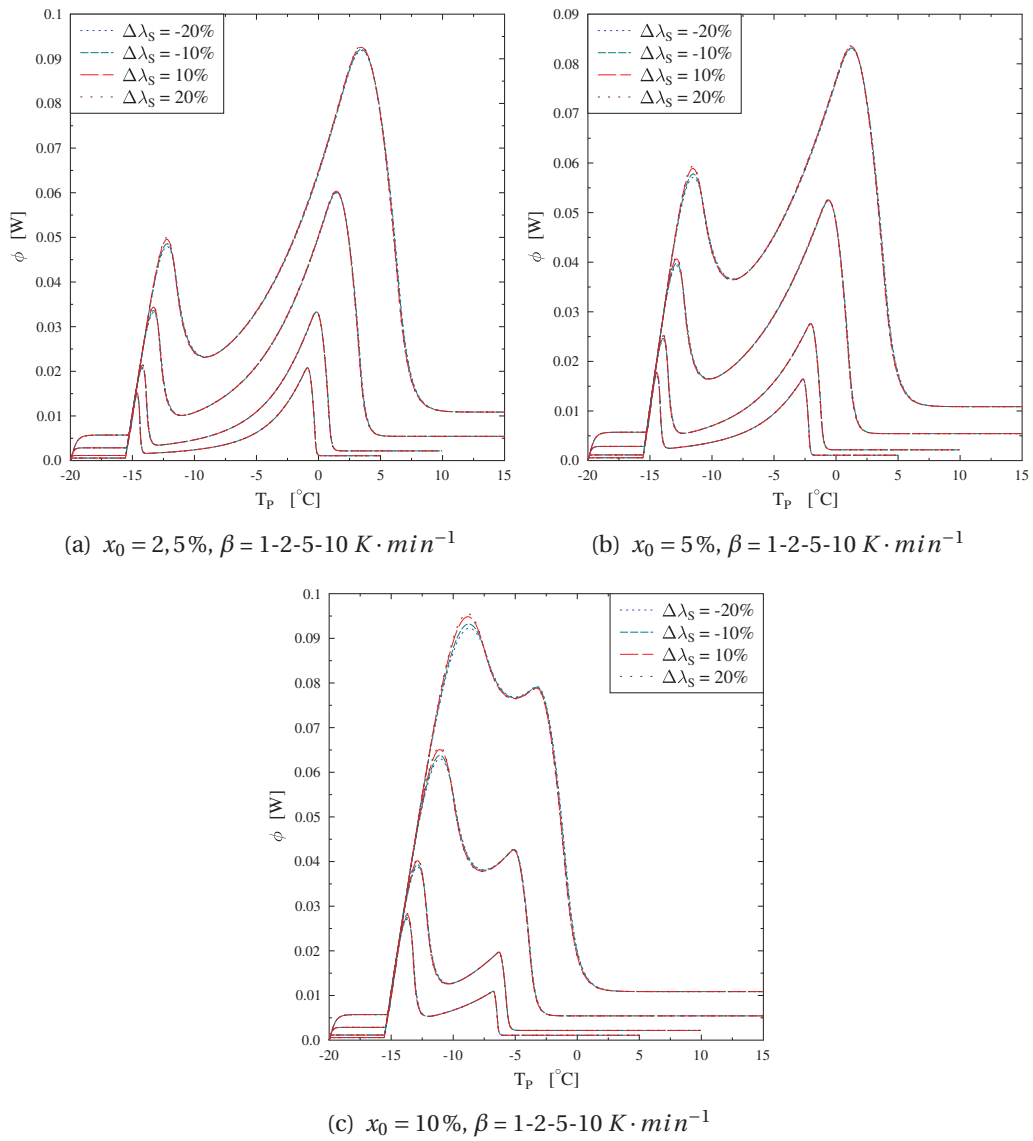


FIGURE 2.21 – Influence de la conductivité thermique solide sur le comportement d'une solution de $\text{H}_2\text{O} - \text{NH}_4\text{Cl}$.

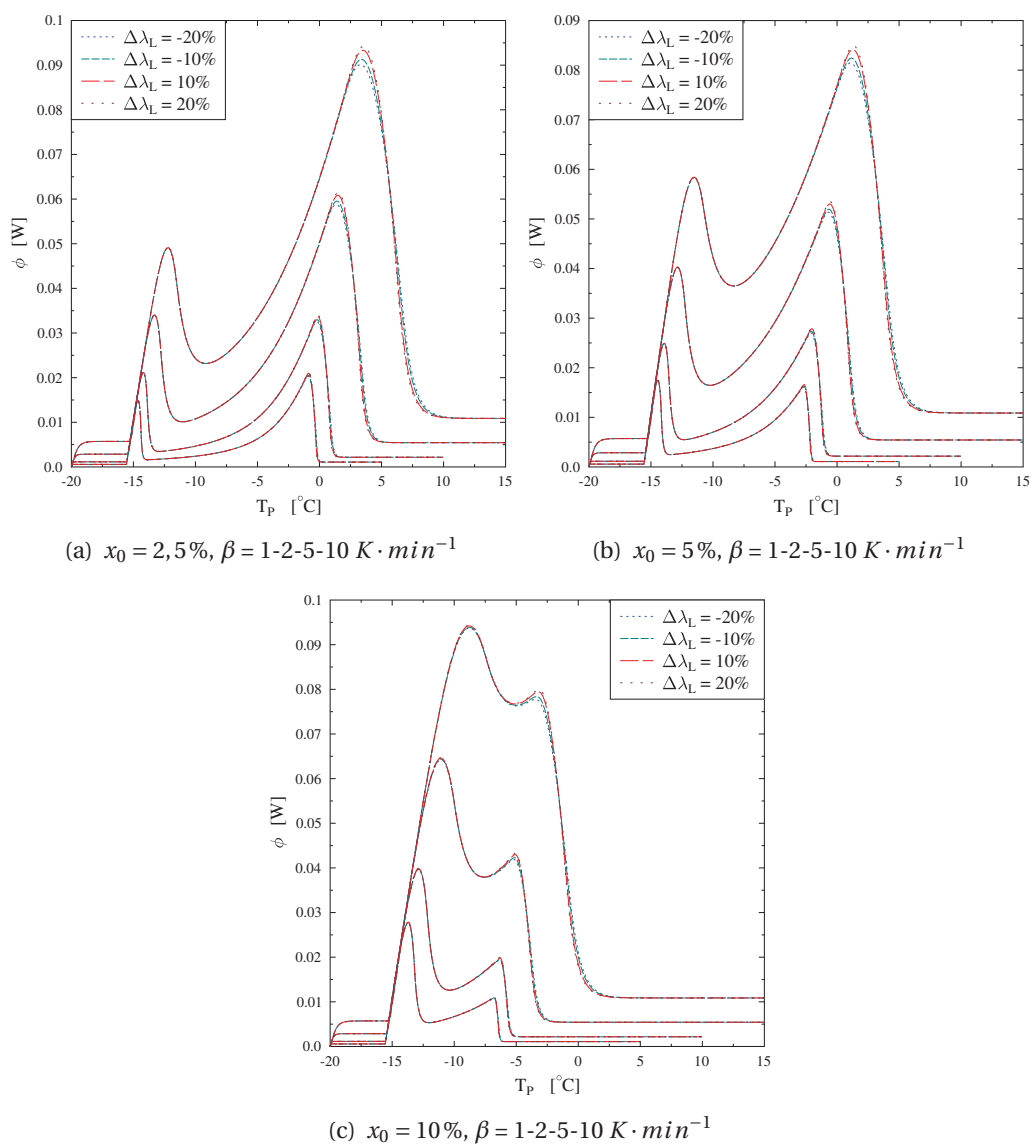


FIGURE 2.22 – Influence de la conductivité thermique liquide sur le comportement d'une solution de $H_2O - NH_4Cl$.

2.3.2.2 Moyenne surfacique des coefficients d'échange équivalent

Comme pour le corps pur, ce paramètre modifie la pente à l'origine du pic eutectique car cette transition se fait à température constante selon l'équation (1.16). Le pic de la fusion progressive du solvant est aussi modifié de manière analogue.

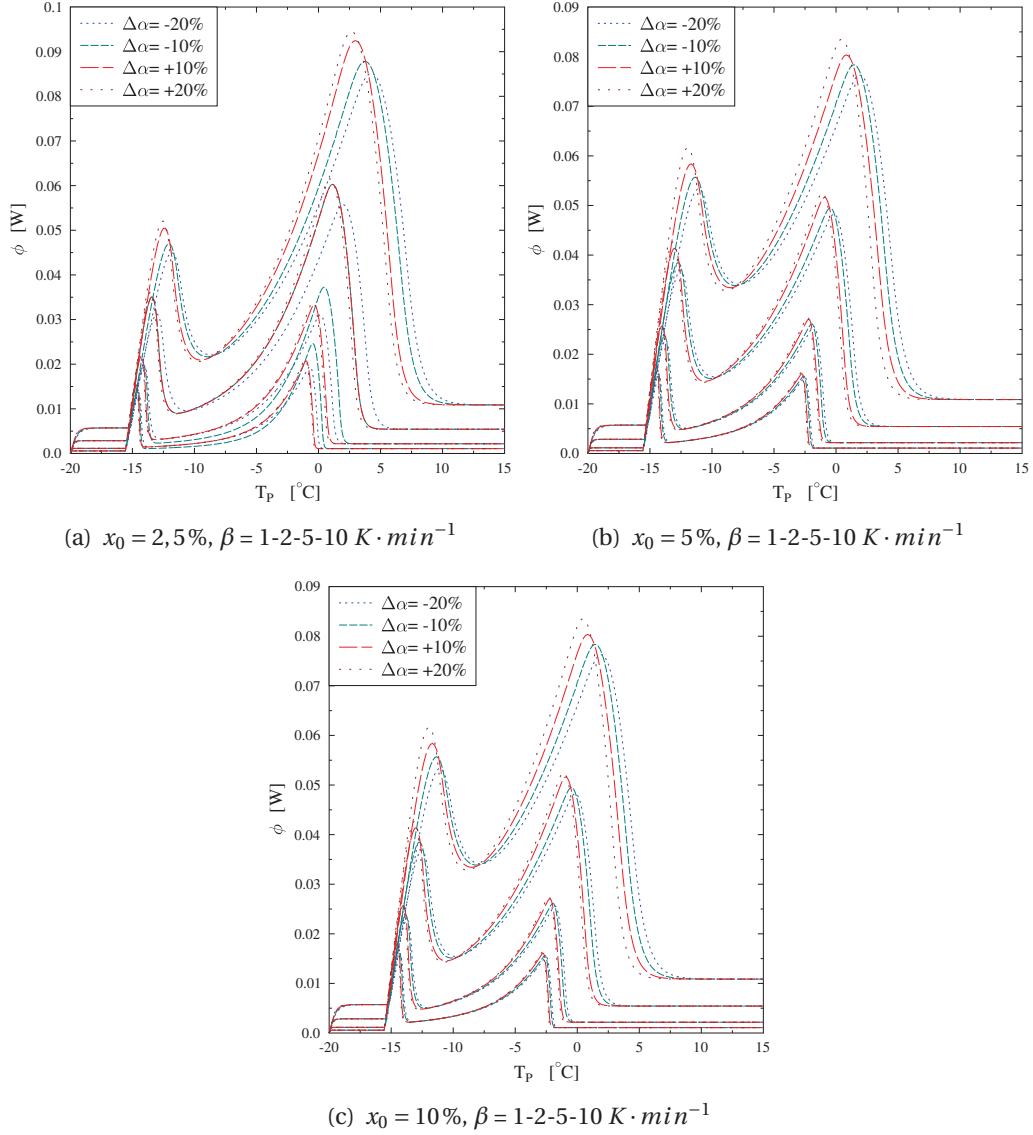
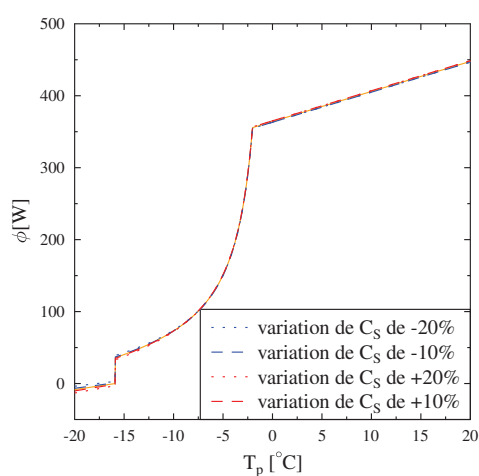


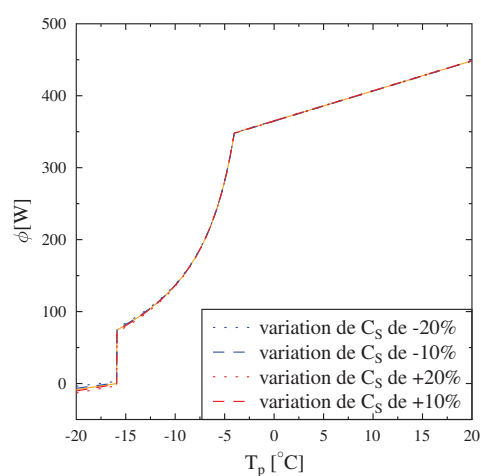
FIGURE 2.23 – Influence de la moyenne des coefficients d'échanges équivalents sur le comportement d'une solution de $H_2O - NH_4Cl$.

2.3.2.3 Capacités thermiques massiques

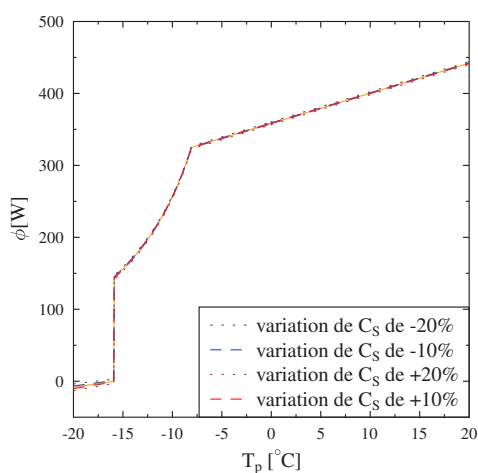
Comme dans le cas du corps pur, les capacités thermiques massiques ont peu d'influence. Sur l'enthalpie, figures 2.24 et 2.26, elles n'influent que en dehors du changement de phase. Globalement, il en résulte que seules les parties du thermogramme avant et après les pics sont modifiées, le changement de phase masquant les autres phénomènes thermiques, comme on peut le voir dans les figures 2.25 et 2.27. Pour les plus faibles concentrations on voit une influence entre les pics des capacités solide et liquide car les termes liés aux capacités dans les équations (2.53) et (2.54) deviennent importants dans ces conditions.



(a) $H_2O - NH_4Cl$ - 2,5 %



(b) $H_2O - NH_4Cl$ - 5 %



(c) $H_2O - NH_4Cl$ - 10 %

FIGURE 2.24 – Influence de c_s , solution $H_2O - NH_4Cl$

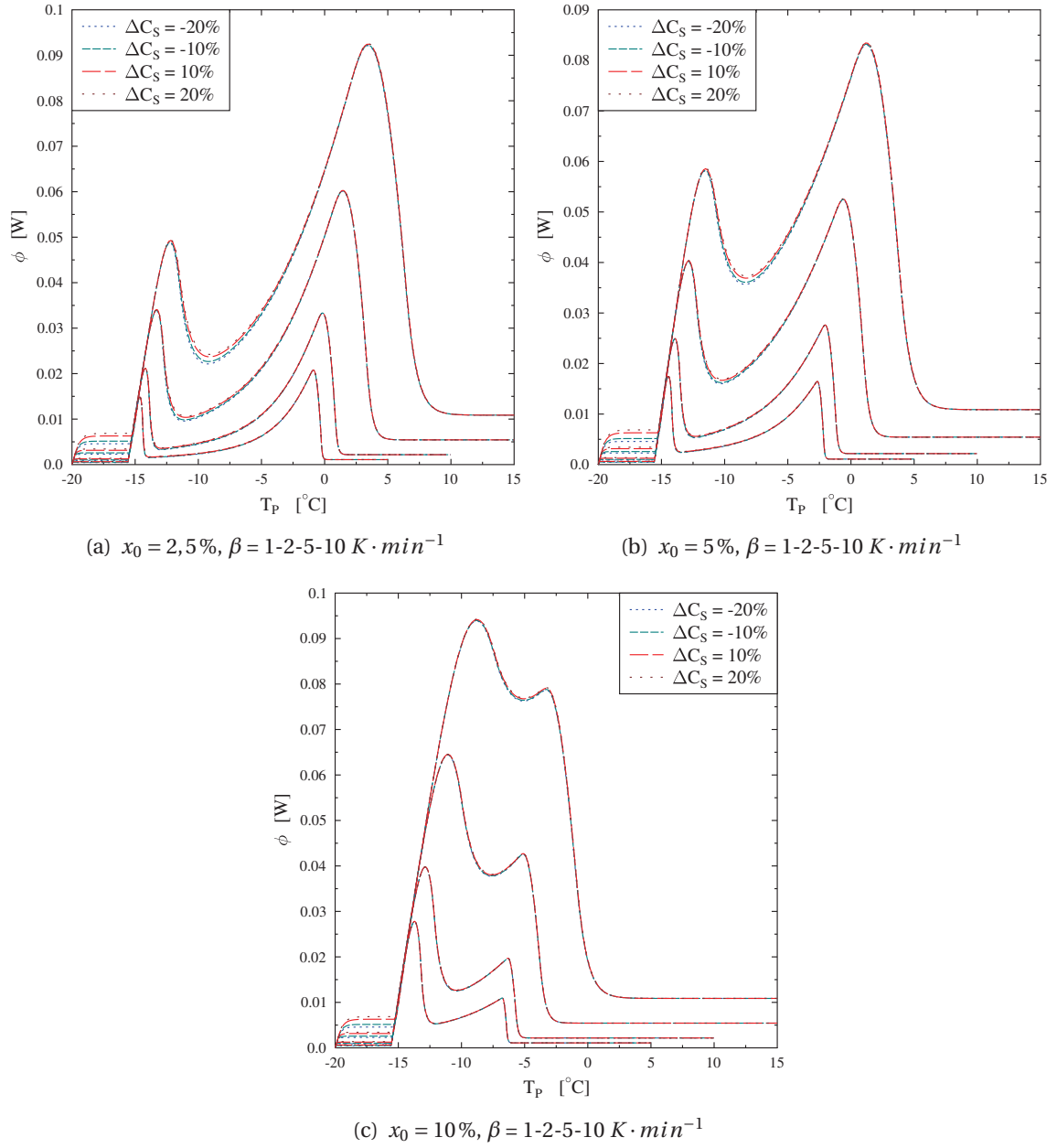
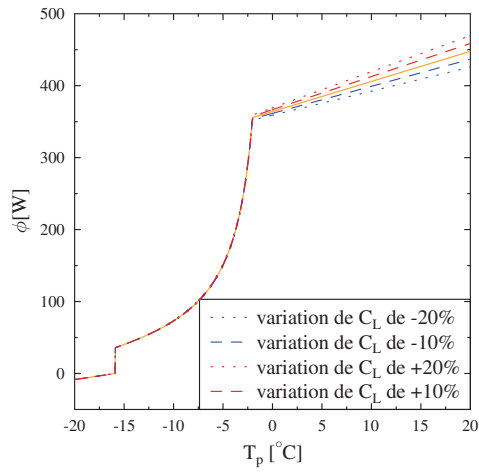
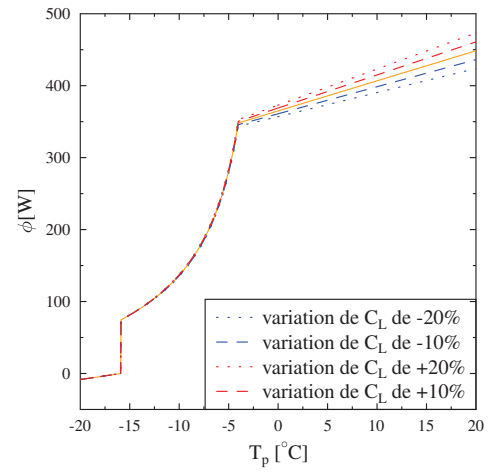
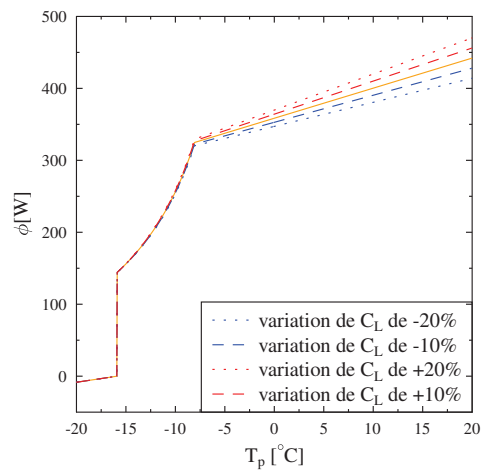


FIGURE 2.25 – Influence des capacités thermiques spécifiques solides sur le comportement d'une solution de $\text{H}_2\text{O} - \text{NH}_4\text{Cl}$.

(a) $H_2O - NH_4Cl$ - 2,5 %(b) $H_2O - NH_4Cl$ - 5 %(c) $H_2O - NH_4Cl$ - 10 %FIGURE 2.26 – Influence de c_L , solution $H_2O - NH_4Cl$

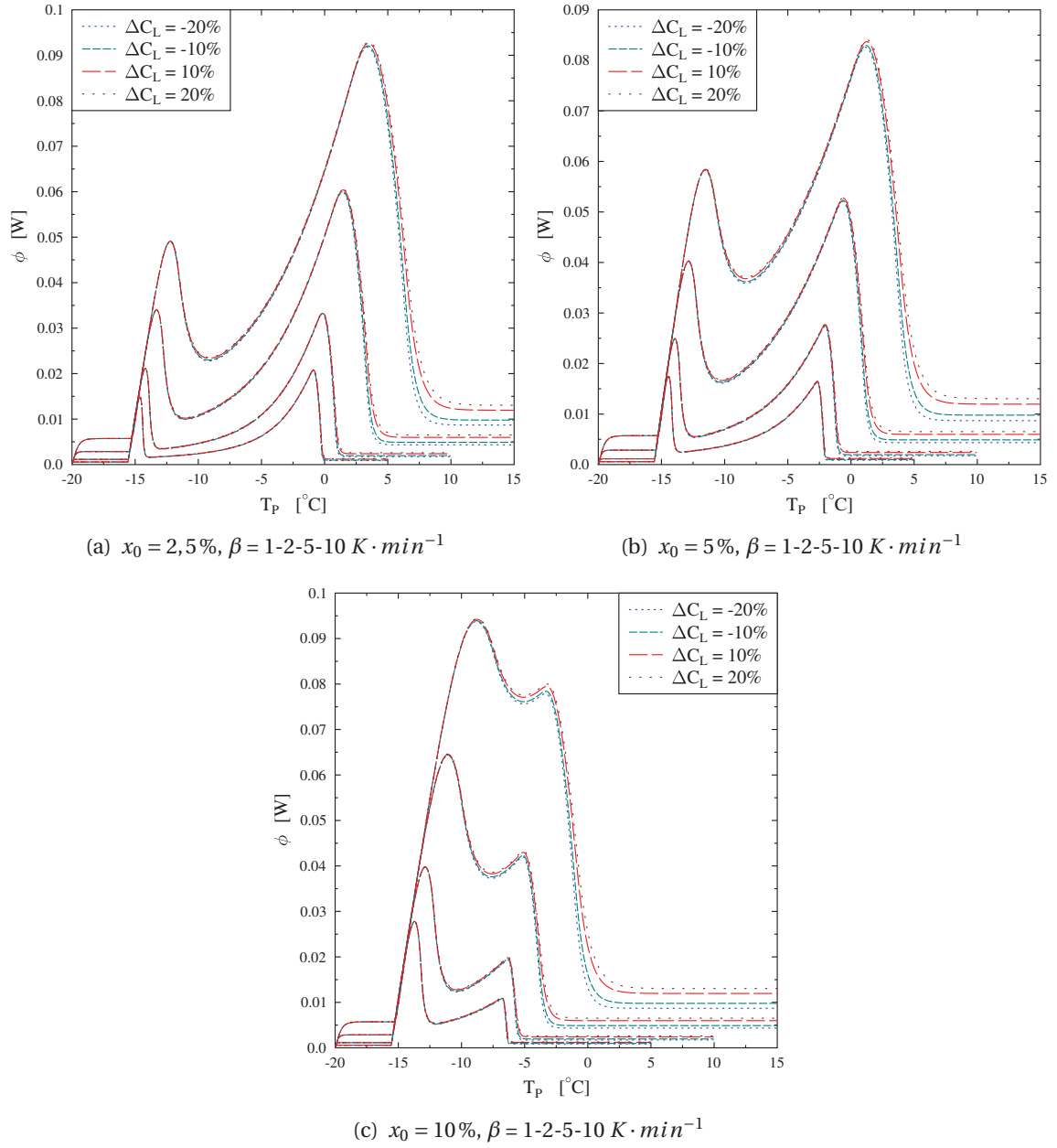


FIGURE 2.27 – Influence des capacités thermiques spécifiques liquides sur le comportement d'une solution de H_2O-NH_4Cl .

2.3.2.4 Température eutectique

A la température eutectique apparaît la première goutte de liquide. C'est cette température qui contrôle le début de la fusion, c'est-à-dire le début du premier pic du thermogramme ou la position du saut d'enthalpie sur les figures de l'ensemble 2.28. Du fait de nos choix pour nos paramètres du modèle (voir équations (2.29) et (2.30)) la modification de T_E toutes choses égales par ailleurs, n'affecte pas la quantité d'énergie mise en jeu au cours de la transformation eutectique. De ce fait, comme pour la température de fusion du corps pur une modification de T_E aura l'effet d'une translation du premier pic, le pic eutectique. Les résultats pour chaque concentration sont présentés dans l'ensemble de figures 2.29.

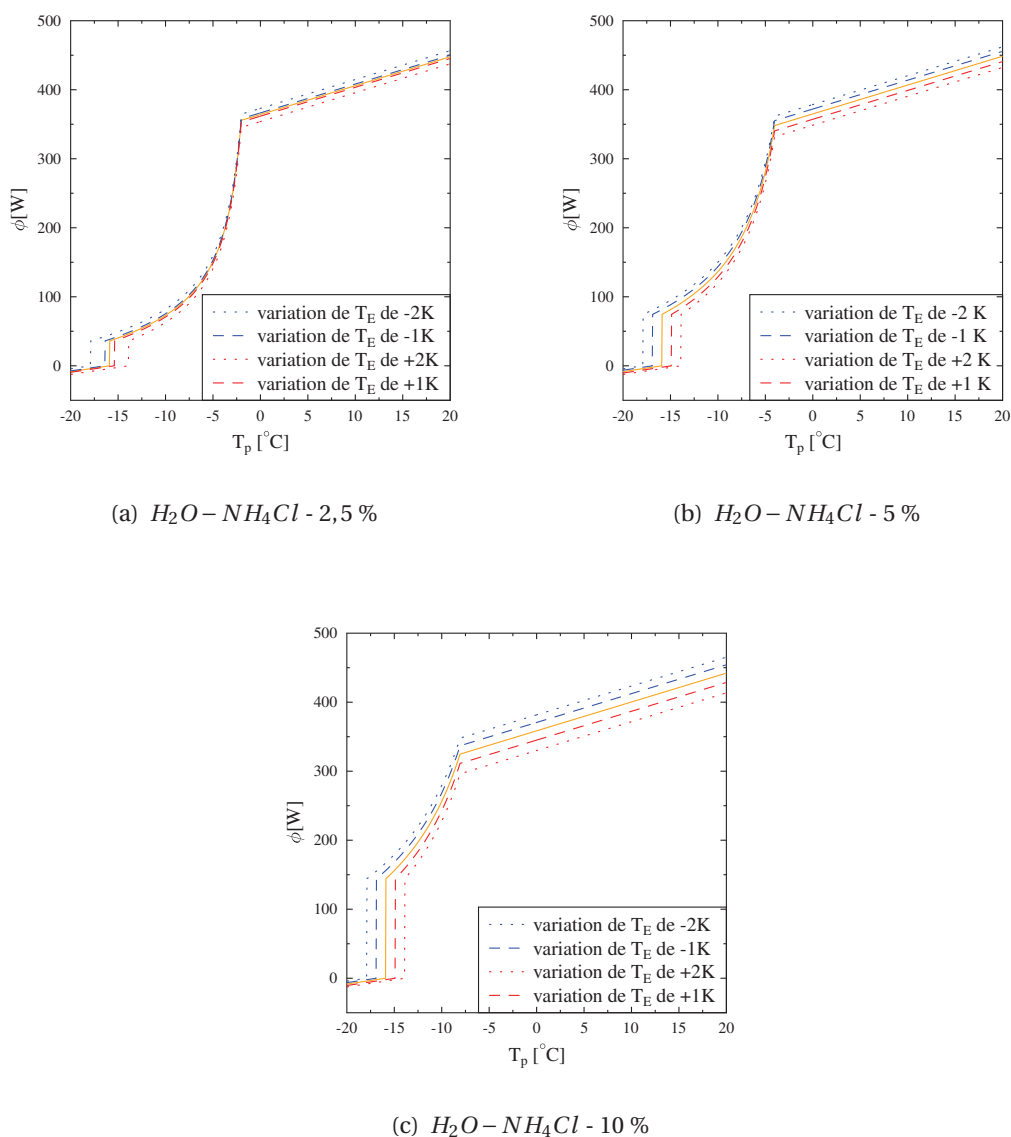


FIGURE 2.28 – Influence de T_E , solution $H_2O - NH_4Cl$

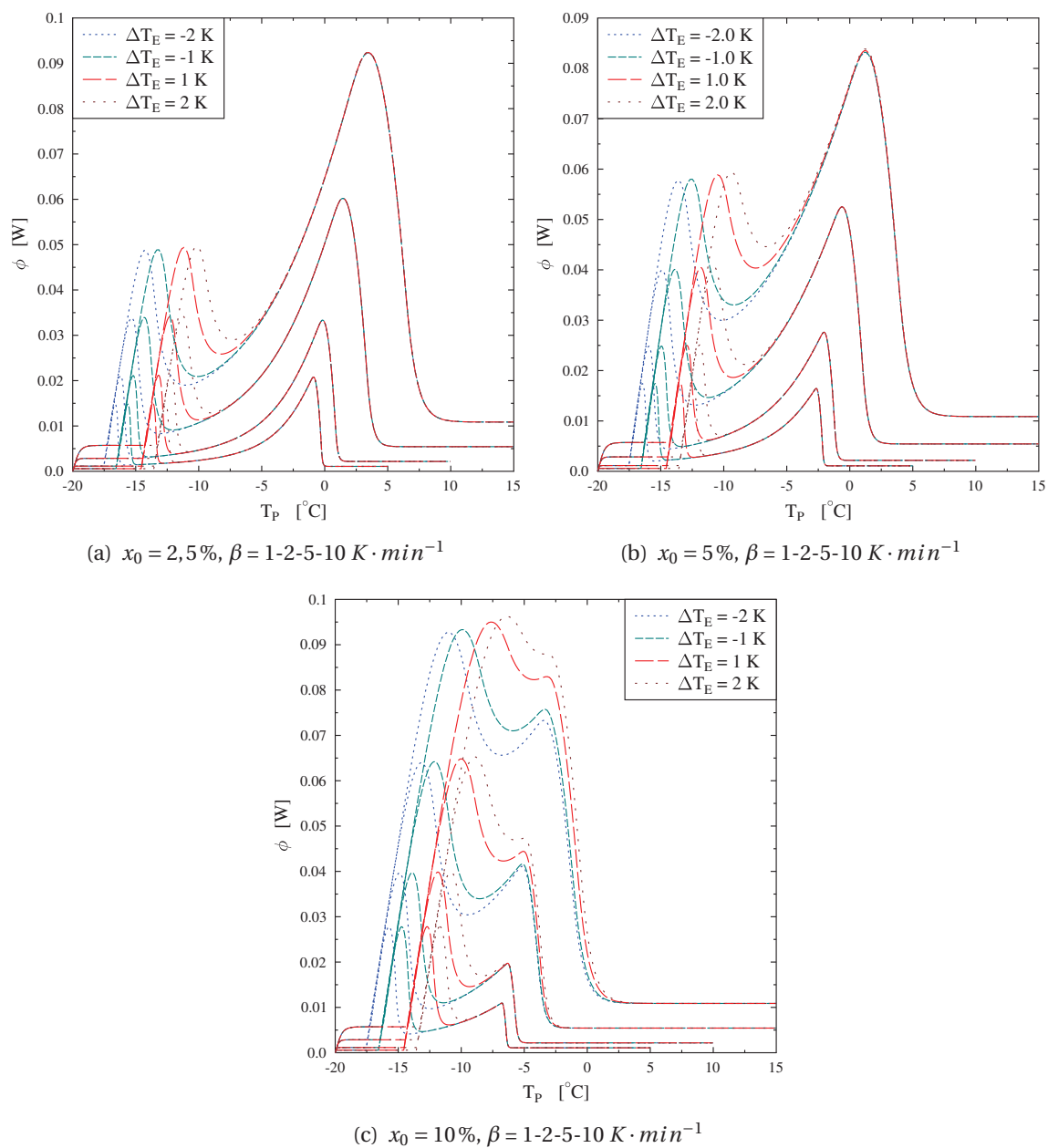


FIGURE 2.29 – Influence de la température eutectique sur le comportement d'une solution de H_2O-NH_4Cl .

2.3.2.5 Variation d'enthalpie massique à la température eutectique : $\tilde{L}_E$

C'est la paramètre que nous avons choisi pour caractériser le saut enthalpique à la température eutectique (voir équation (2.28)) toutes choses égales par ailleurs. L'influence de ce paramètre est montrée sur l'enthalpie puis sur les thermogrammes aux figures 2.30 et 2.31. Pour de faibles fractions massiques en soluté, on remarque qu'il modifie principalement le premier pic du thermogramme comme dans le cas du corps pur. Les pics étant bien séparés, les énergies mises en jeu pour chaque pic ont des ordres de grandeurs différents. Ceci n'est plus vrai pour des fractions massiques en soluté plus importantes où c'est la forme globale du thermogramme qui est modifiée. Les pics y sont mélangés car T_M diminuant avec x_0 et $\tilde{L}_E$ augmentant avec x_0 (voir tableau 2.2), le pic eutectique va s'élargir « vers la droite » et celui de la fusion progressive sera « décalé vers la gauche », entraînant le mélange des pics. Ce qui explique la propagation de la modification de $\tilde{L}_E$ sur le thermogramme. On retrouve ces influences sur les enthalpies, figure 2.30.

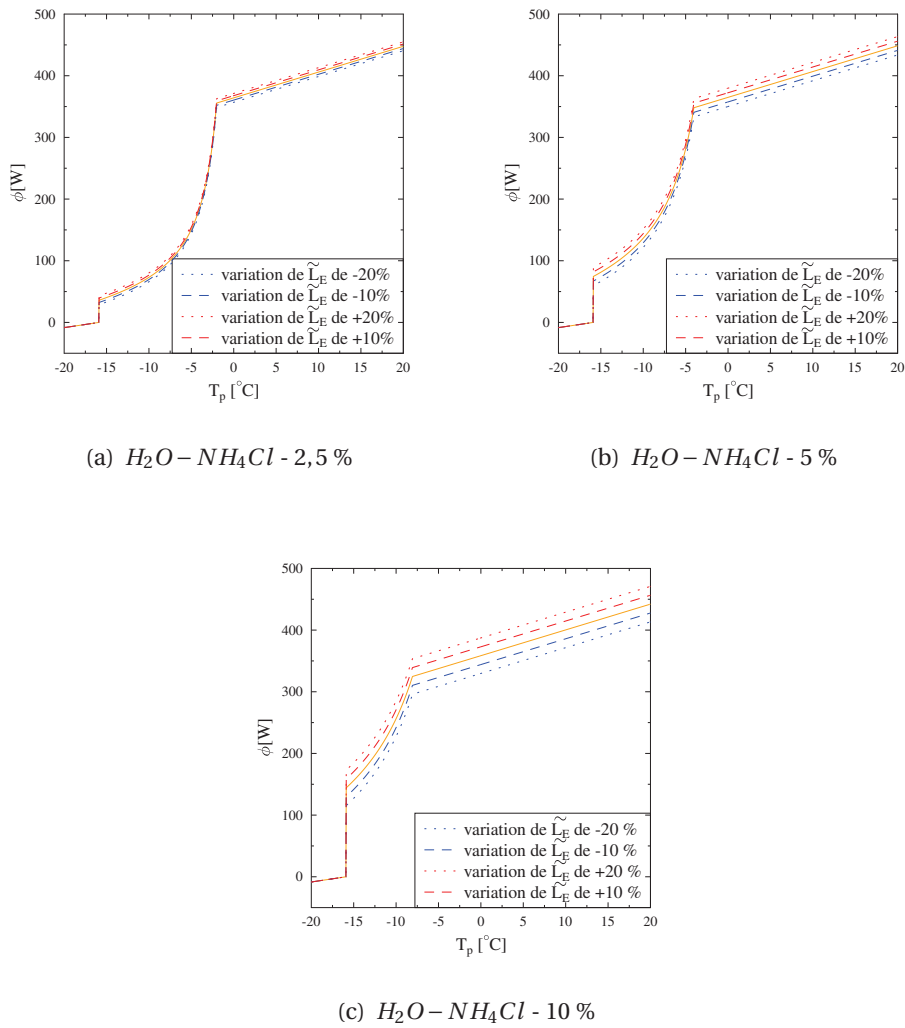


FIGURE 2.30 – Influence de $\tilde{L}_E$, solution $H_2O - NH_4Cl$

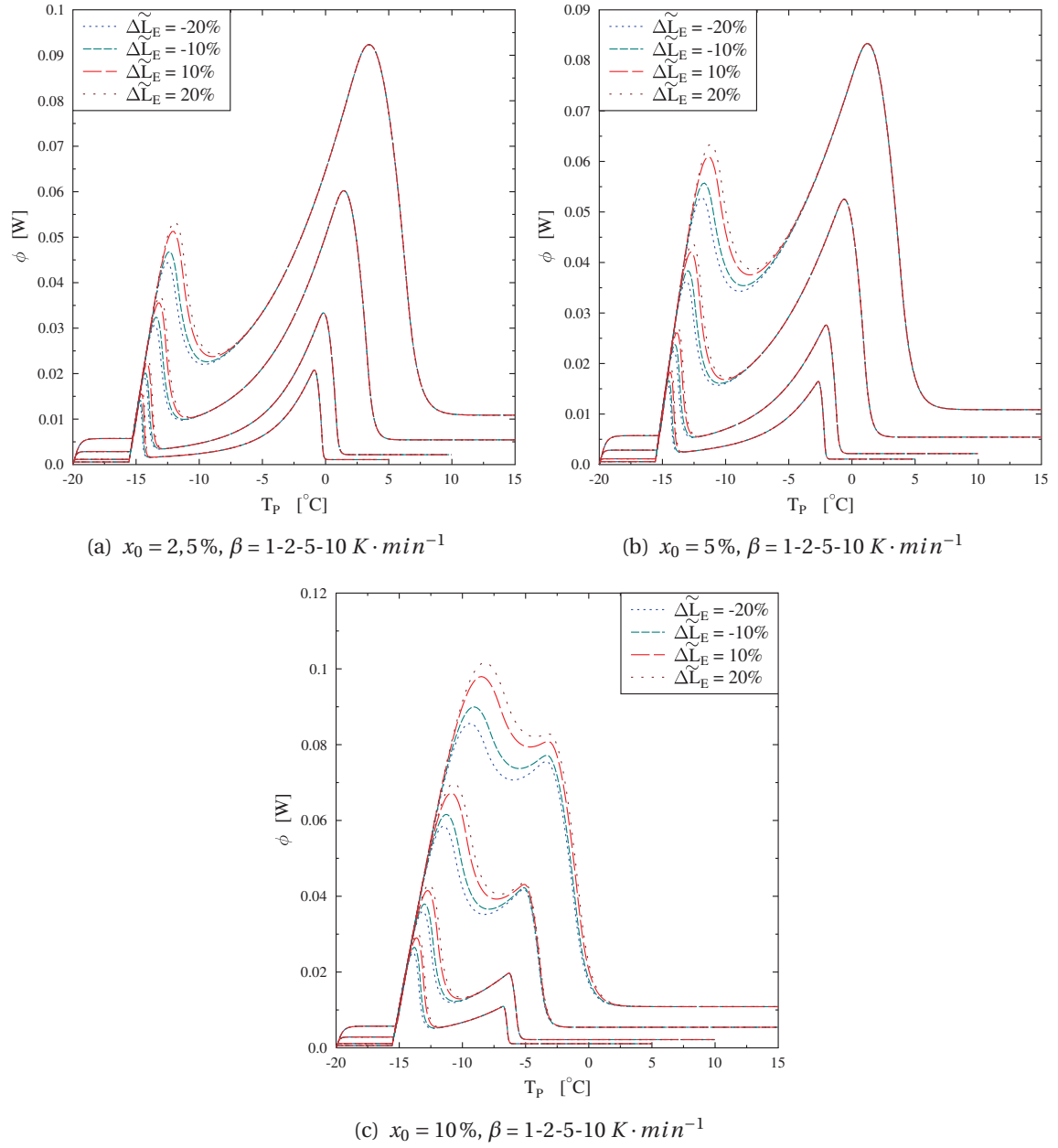
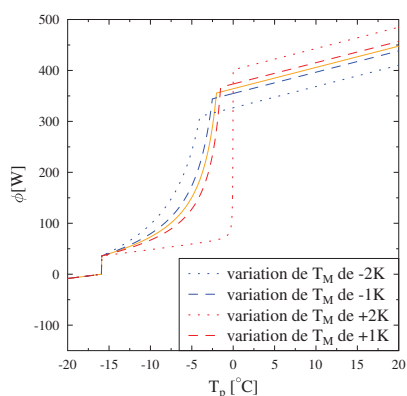


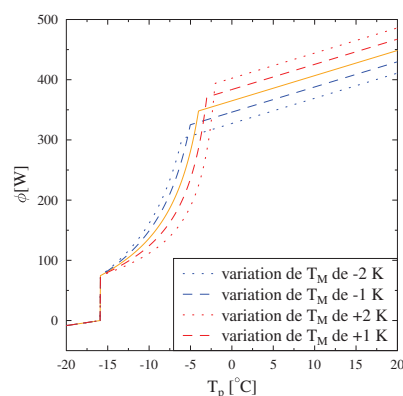
FIGURE 2.31 – Influence de la chaleur latente eutectique sur le comportement d'une solution de H_2O-NH_4Cl .

2.3.2.6 Température de liquidus

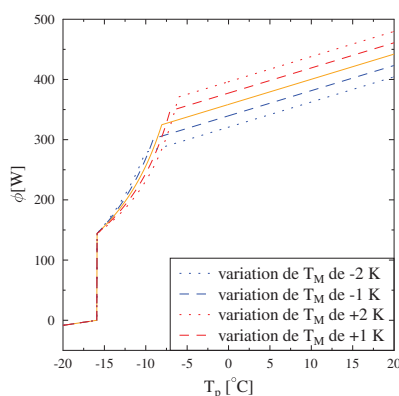
Comme dans le cas de la température eutectique, nous nous attendons à une grande influence de ce paramètre car marquant la fin de fusion au niveau thermodynamique, figure 2.32, il augmente l'intervalle de température couvert par la fusion. Sur le thermogramme il influe sur l'écartement des pics en étalant le deuxième. Dans la figure 2.33, on constate qu'une variation de 1 ou 2 K modifie de façon notable le thermogramme, notamment en décalant le deuxième pic. Une modification de ce paramètre modifie la concentration massique en soluté et devrait modifier la valeur de $\tilde{L}_E$. Mais nous avons choisi de découpler ces variables (2.1.3). Ce faisant comme dans cette étude paramétrique nous ne faisons varier qu'un paramètre à la fois, toutes choses égales par ailleurs. On a pas de variation de $\tilde{L}_E$. Les graphiques sont regroupés par fraction massique de soluté de référence pour pouvoir les comparer. On notera également que un problème d'étalonnage en température va décaler le thermogramme sans en modifier la forme, mais en décalant les paramètres T_E , $T_{M,w}$ et T_M d'une même valeur.



(a) $H_2O - NH_4Cl$ - 2,5 %



(b) $H_2O - NH_4Cl$ - 5 %



(c) $H_2O - NH_4Cl$ - 10 %

FIGURE 2.32 – Influence de T_M , solution $H_2O - NH_4Cl$

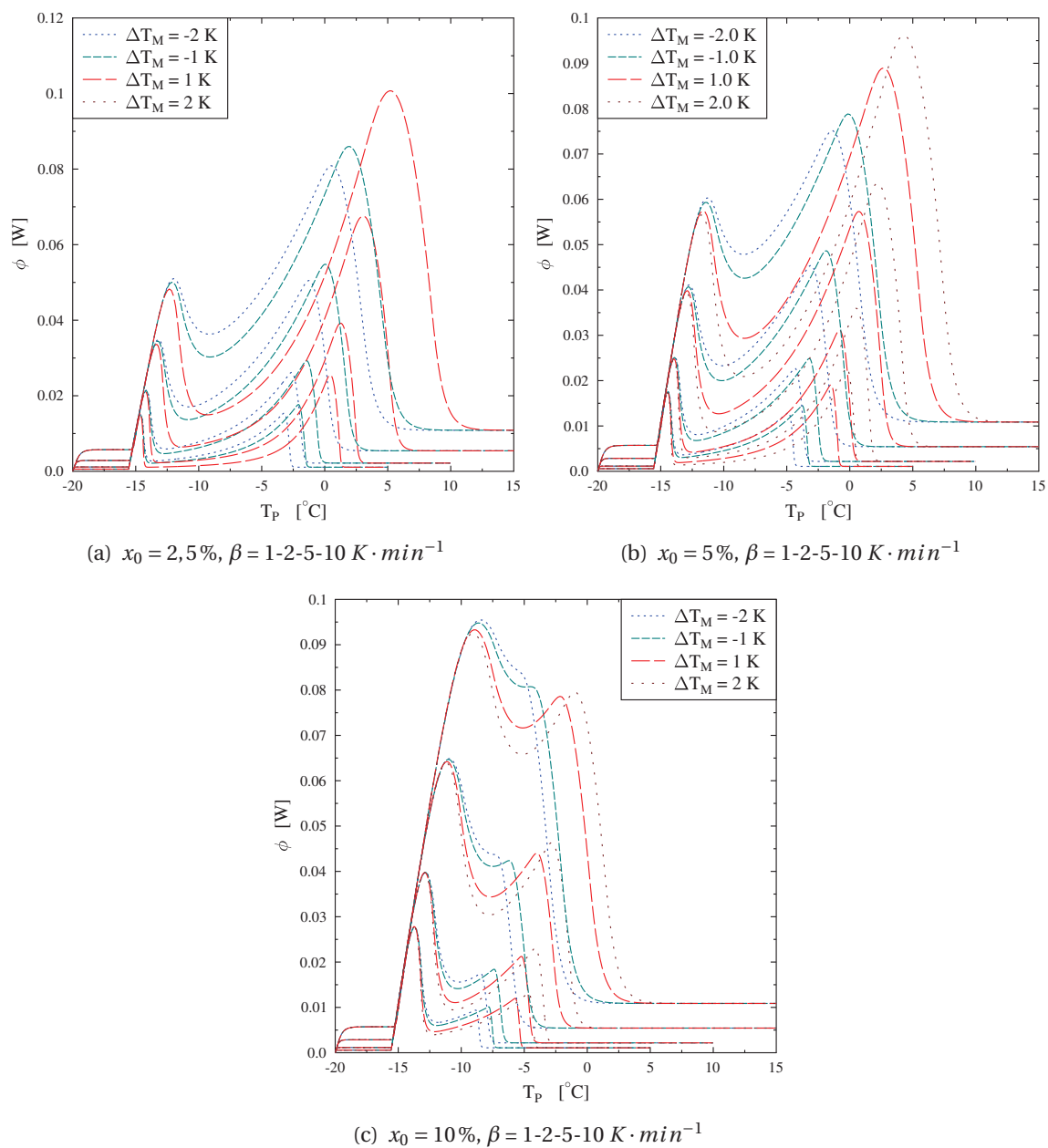
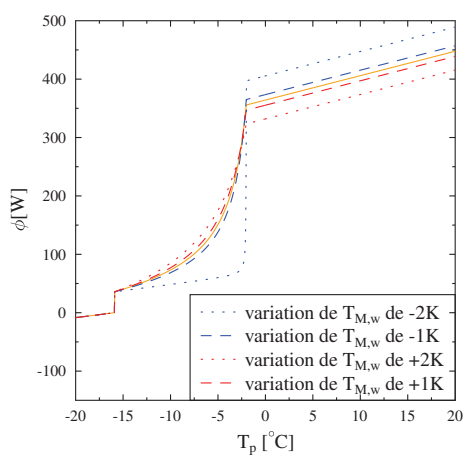


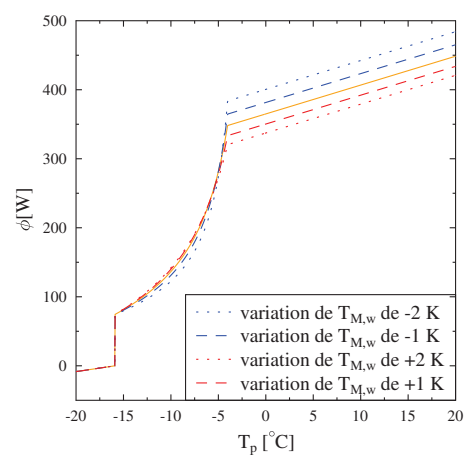
FIGURE 2.33 – Influence de la température de liquidus sur le comportement d'une solution de $\text{H}_2\text{O} - \text{NH}_4\text{Cl}$.

2.3.2.7 Température de fusion du solvant

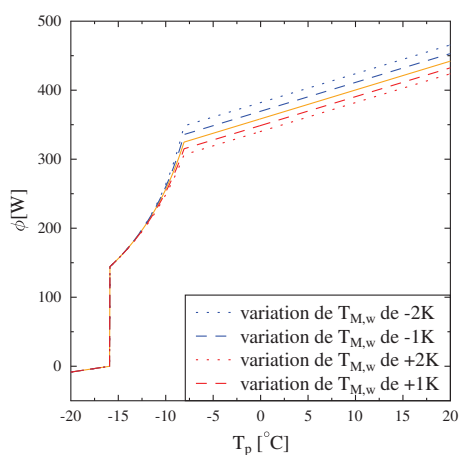
Une augmentation de $T_{M,w}$ diminue la surface du second pic du thermogramme. En effet d'après la figure 2.34, l'enthalpie est décroissante selon ce paramètre. En l'augmentant on diminue l'énergie à apporter pour terminer le changement de phase, figures 2.35 et 2.34.



(a) $H_2O - NH_4Cl$ - 2,5 %



(b) $H_2O - NH_4Cl$ - 5 %



(c) $H_2O - NH_4Cl$ - 10 %

FIGURE 2.34 – Influence de $T_{M,w}$, solution $H_2O - NH_4Cl$

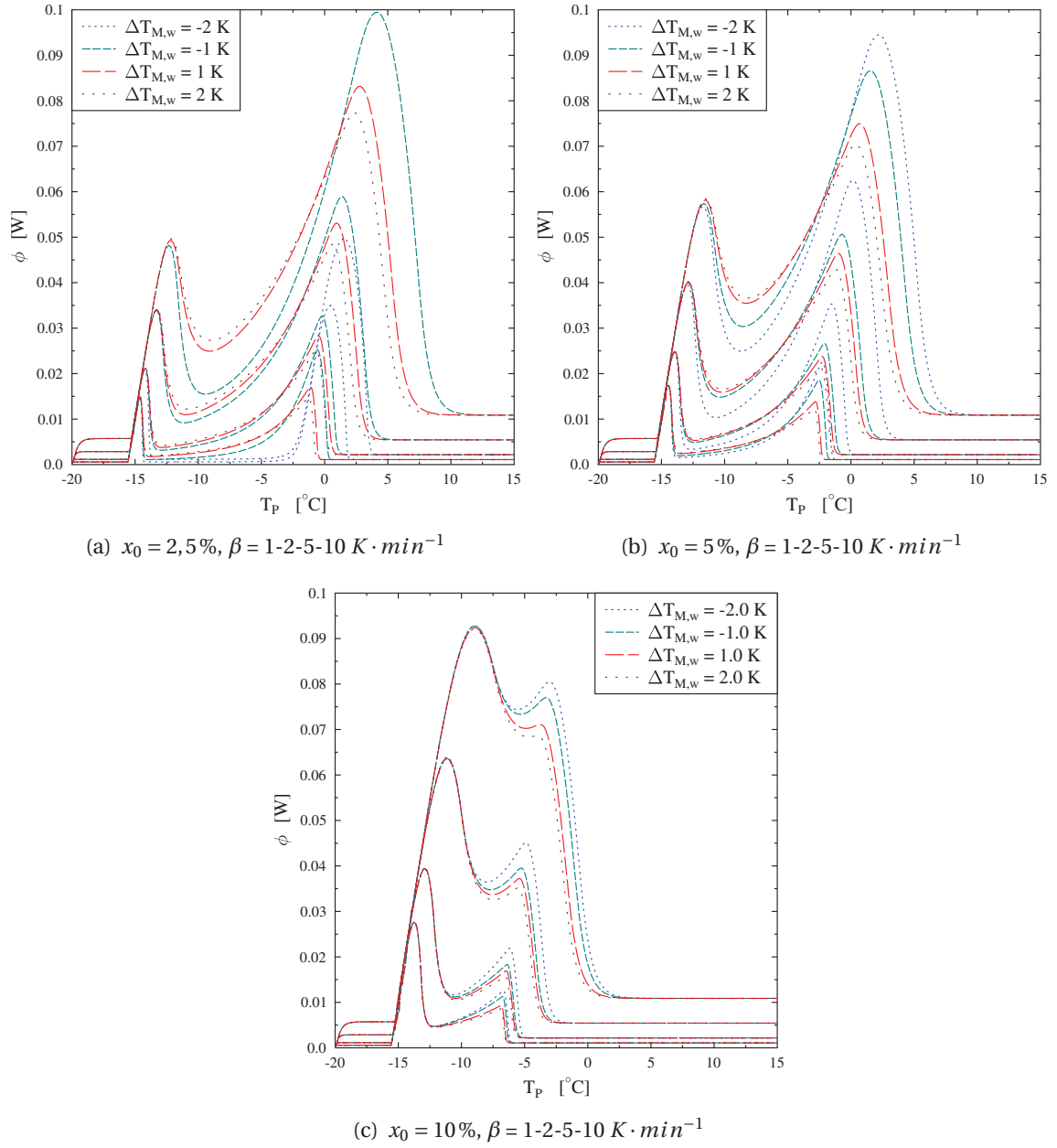
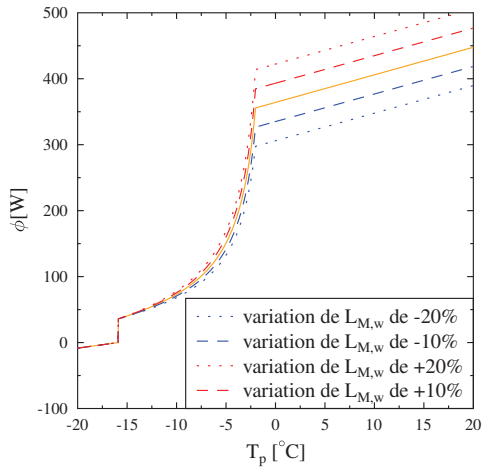


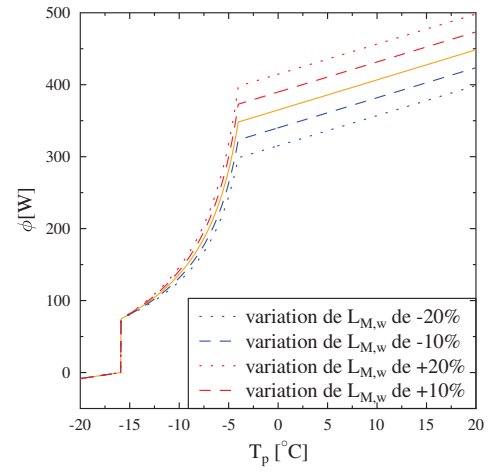
FIGURE 2.35 – Influence de la température de liquidus sur le comportement d'une solution de $H_2O - NH_4Cl$.

2.3.2.8 Chaleur latente de fusion du solvant

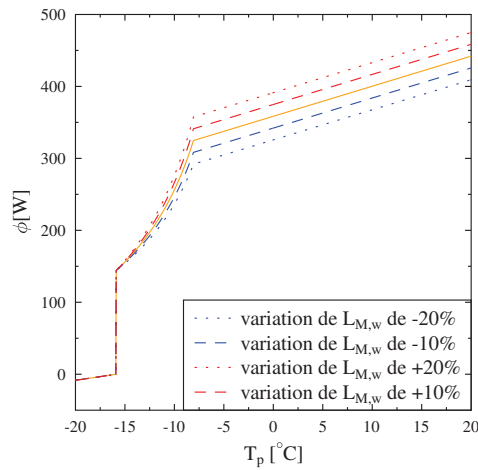
Les résultats correspondants se trouvent dans les figures 2.32 et 2.37. La chaleur latente du solvant, intervenant dans l'équation (2.41), influence le deuxième pic d'une manière semblable à celle dont la chaleur eutectique influence le premier pic.



(a) $H_2O - NH_4Cl$ - 2,5 %



(b) $H_2O - NH_4Cl$ - 5 %



(c) $H_2O - NH_4Cl$ - 10 %

FIGURE 2.36 – Influence de $L_{M,w}$, solution $H_2O - NH_4Cl$

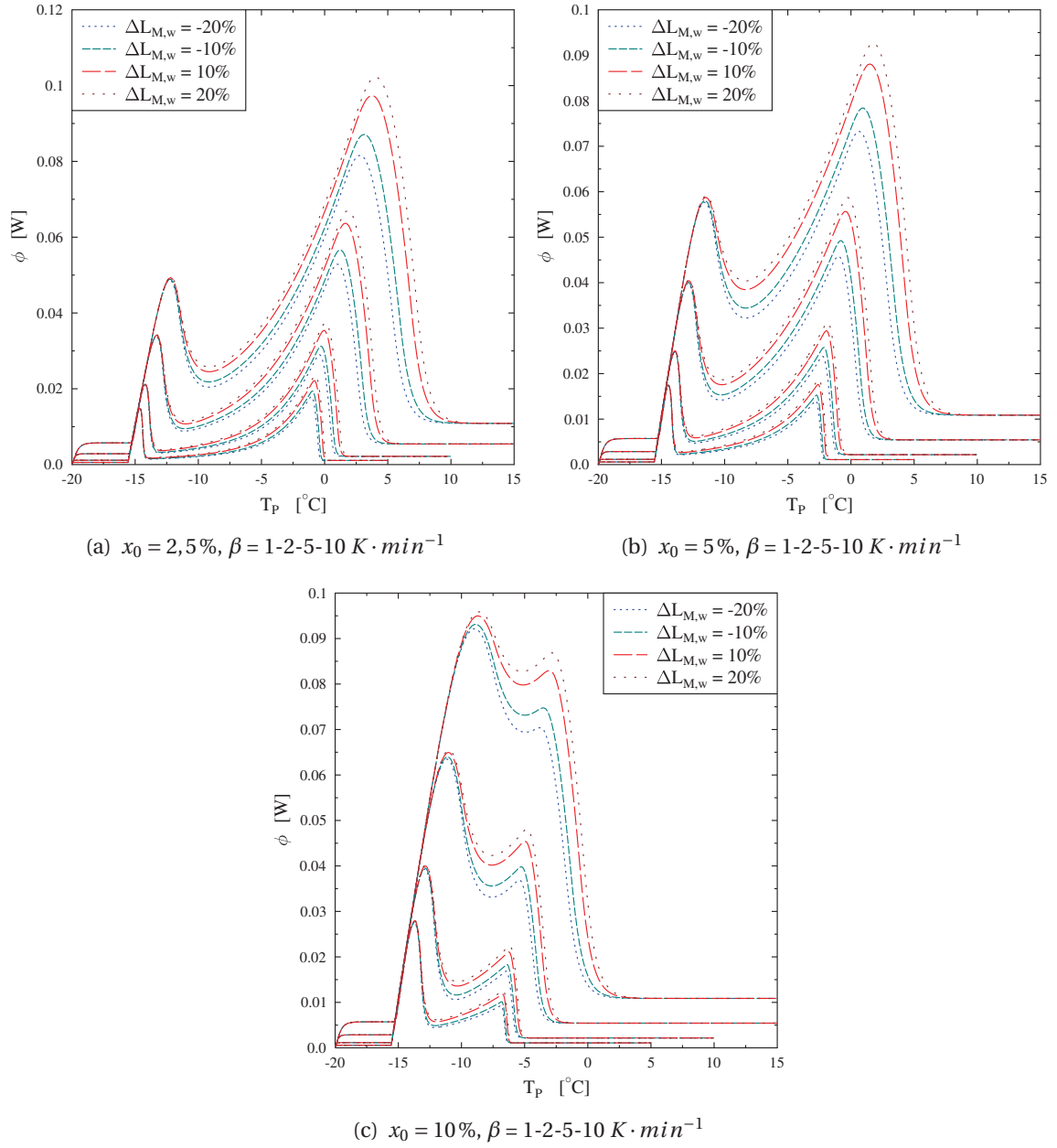


FIGURE 2.37 – Influence de la chaleur latente de fusion sur le comportement d'une solution de H_2O-NH_4Cl .

2.4 Influence de la forme - surface libre de l'échantillon [67]

Nous cherchons à évaluer l'influence qu'a la forme de la surface libre sur le thermogramme, en considérant la même valeur pour le coefficient d'échange équivalent entre chaque simulation.

2.4.1 Situation courante

Jusqu'à présent nous avons considéré que la surface libre de l'échantillon était plane. Mais cela n'est jamais strictement le cas, il y a toujours une surface libre (face supérieure) de l'échantillon. Etant donné que pour un appareil de DSC les échantillons sont de l'ordre de 10 mg, il est fort probable que les effets de capillarité déforment la surface libre (voir figure 2.38), donc modifient la résistance thermique à l'interface, et modifient le thermogramme. Quand il s'agit de remonter uniquement à des grandeurs peu sensibles aux paramètres expérimentaux, comme la température de fusion et l'enthalpie de fusion par exemple, ou lorsque les échantillons analysés ont des dimensions assez grandes pour négliger la courbure de la surface libre, il n'est pas nécessaire d'examiner plus en avant cette question. Mais dernièrement, plusieurs modèles furent proposés pour simuler les transferts thermiques à l'intérieur de l'échantillon lors du changement de phase [56] afin de connaître des grandeurs non accessibles expérimentalement, comme par exemple le champ de température ou la composition des solutions. Pour obtenir ces informations, il est nécessaire de traiter l'échantillon dans son ensemble, *i.e.* à la condition d'en faire une modélisation bi-dimensionnelle (généralement avec une géométrie cylindrique). La plupart des modèles qui ont été développés ne prennent pas en compte la courbure de la surface libre et font intervenir des coefficients d'échange équivalents, α , qui sont ajustés au préalable par le biais d'expériences de calorimétrie différentielle à balayage. Il est intéressant d'étudier l'influence d'une courbure de cette surface libre et si elle va influencer sur certains paramètres du modèle.

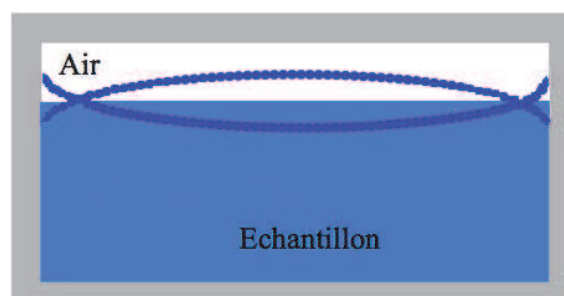


FIGURE 2.38 – Schéma de la surface libre (coubure exagérée)

2.4.2 Méthode

L'étude de la surface libre a été réalisée avec le modèle 2D - cylindrique pour un corps pur. La courbure de la surface libre fut modifiée en redistribuant les nœuds du maillage selon la hauteur du cylindre de manière à conserver le même volume total, donc la même masse

d'échantillon, en remplaçant l'équation (2.58) par :

$$z_{i,j} = Z \frac{j-1}{NJ} \left[1 - \frac{NI-1}{2NI} \delta_{courbure} + \left(\frac{r(i,j)}{R} \right)^2 \delta_{courbure} \right] \quad (2.69)$$

Les nœuds supérieurs sont distribués autour d'une hauteur de référence qui correspond au cas à courbure infinie (surface libre plane), de sorte à conserver la même masse entre les différents cas de courbure (figures 2.38 et 2.39). On ne considère que l'effet géométrique de la variation de la courbure, on considère que cela ne modifie pas le coefficient d'échange.

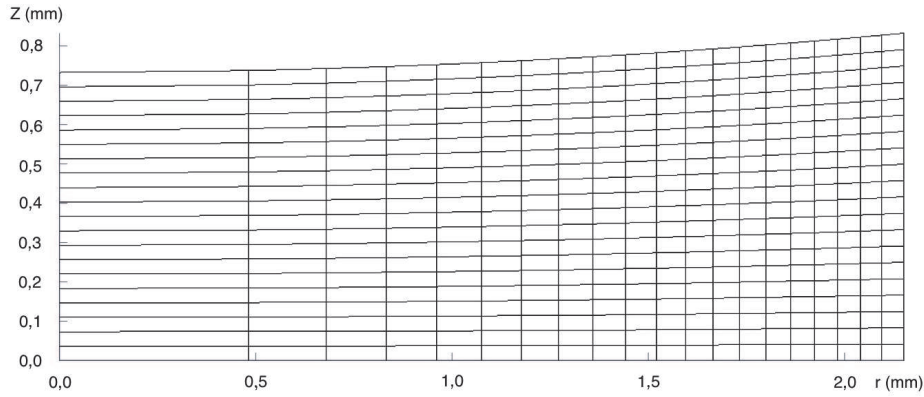


FIGURE 2.39 – Modification du maillage pour une surface libre convexe

2.4.3 Résultats

Nous ne présentons ici qu'une série de résultats correspondant au cas de l'eau. Les coefficients d'échange équivalents utilisés sont : $\alpha_{supérieur} = 100 \text{ W} \cdot \text{m}^{-2} \cdot \text{K}^{-1}$ et $\alpha_{inférieur} = \alpha_{latérale} = 1000 \text{ W} \cdot \text{m}^{-2} \cdot \text{K}^{-1}$.

Initialement, l'échantillon est constitué entièrement de glace (taux liquide $\chi_L = 0$ dans tout le domaine) et la température est fixée à $T_{deb} = -1,0 \text{ } ^\circ\text{C}$.

Le premier essai est effectué pour un chauffage allant de $-1,0 \text{ } ^\circ\text{C}$ à $11 \text{ } ^\circ\text{C}$, à la vitesse de $1 \text{ K} \cdot \text{min}^{-1}$. La figure 2.40 présente les profils de température à l'intérieur de l'échantillon, au bout de 200 s de chauffe, dans les cas où la surface libre est plane 2.40(a), convexe 2.40(c) et concave 2.40(e), le volume de l'échantillon étant identique dans chacun des cas. On peut y voir qu'une fine couche de liquide est présente au-dessus du glaçon du fait d'un apport de chaleur sur la face supérieure et que nous n'avons pas pris en compte la variation de la masse volumique qui aurait fait flotter ce dernier. Néanmoins du fait de la prédominance des échanges inférieur et latéral, la fusion se fait du bas vers le haut et de l'extérieur vers l'intérieur.

Sur la figure 2.40, on peut confirmer que la partie de l'échantillon encore solide est quasi uniformément à la température T_M .

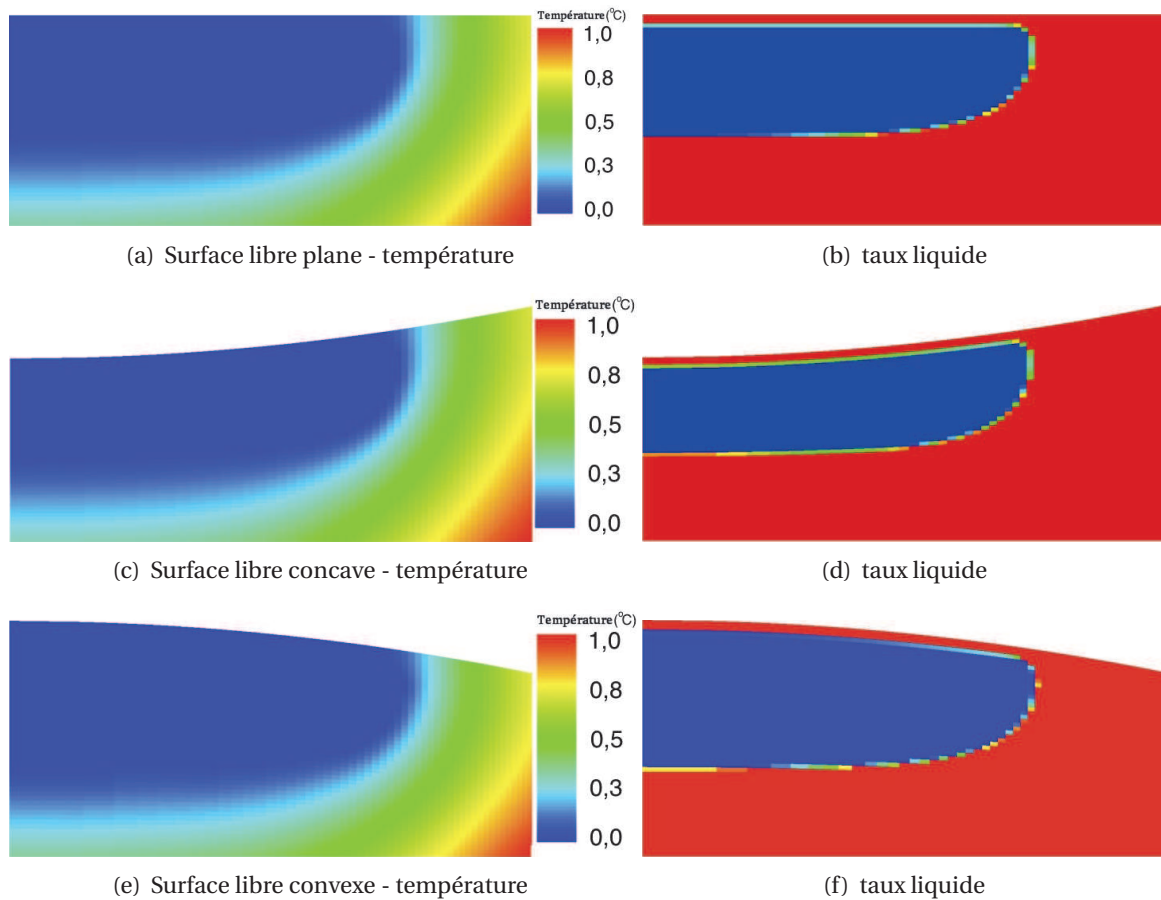


FIGURE 2.40 – Simulations de différentes courbures de surface

La figure 2.41 montre les thermogrammes obtenus pour différentes formes de la surface libre. On remarque une faible influence, avec une fusion d'autant plus rapide que la concavité est prononcée, du fait de l'augmentation du flux latéral résultant de l'augmentation de cette même surface. On observe également que le maximum du pic (flux maximum) est d'autant plus haut en flux que la fusion est rapide du fait de l'augmentation de la surface totale et donc du produit αS qui détermine la pente en début de pic (voir équation (1.16)). Les aires de chaque thermogramme, qui représentent la chaleur latente, restent identiques.

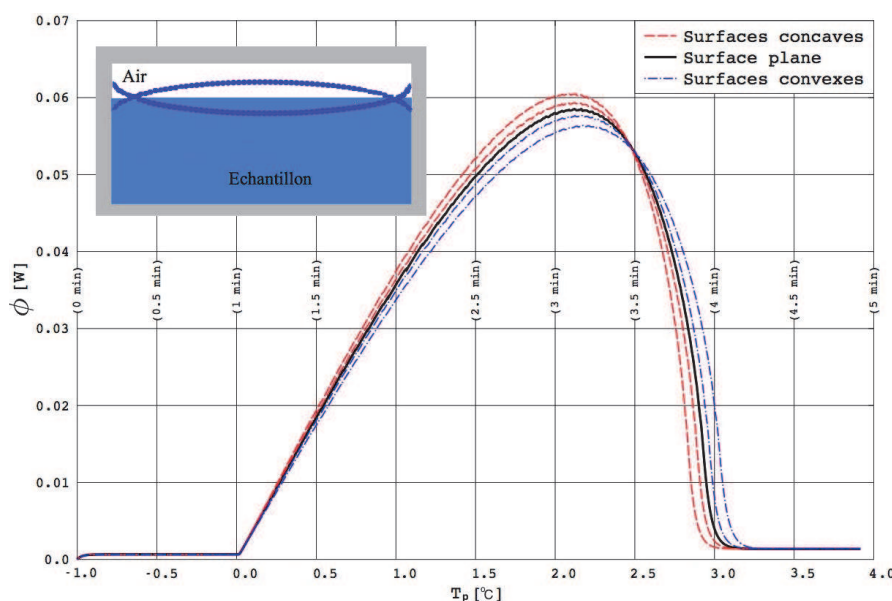


FIGURE 2.41 – Déformation du thermogramme selon la courbure de la surface libre, $5 K \cdot min^{-1}$.

2.4.4 Synthèse des résultats

Les résultats présentés montrent que, lors d'une expérience de calorimétrie, la géométrie de l'échantillon a une influence sur la forme du thermogramme si on considère toujours une même valeur pour le coefficient d'échange équivalent α . Mais on voit sur la figure 2.41 que cette influence est principalement une modification de la pente à l'origine du pic du thermogramme, car nous avons modifié la géométrie, donc les surfaces et donc les flux internes, en conservant les mêmes valeurs pour les paramètres de transfert. La modification de la pente à l'origine induite par la modification de la géométrie, peut donc être compensée par une modification du paramètre α d'après l'équation (1.16) et l'étude sur le paramètre α , paragraphe 2.3.1.2.

Pourtant, le modèle 2D-cylindrique qui propose une surface libre plane (à courbure infinie) fournit de très bons résultats [28] [29] [30], alors que en réalité, du fait des effets de capillarité, cette surface libre présente une légère courbure limitée au ménisque. Comme le modèle à surface libre plane reflète correctement l'expérience, on en conclut que c'est la valeur du coefficient convectif équivalent α utilisé dans le modèle qui s'adapte à la géométrie pour que le résultat coïncide avec l'expérience. Ce coefficient est ici une boîte noire qui permet d'ajuster le modèle du transfert thermique avec l'expérience, car il est difficile à mesurer, et bien souvent on utilise une estimation.

2.5 Réduction du modèle

2.5.1 Incertitude sur la forme de l'échantillon

Comme on peut le voir sur les photographies 2.42, dans le cas de notre appareil Pyris Diamond D.S.C. de Perkin-Elmer comme nous avons des petits échantillons, environ 10 mg, les

effets de capillarité l'emportent et la géométrie réelle n'est pas du tout cylindrique. Et pourtant le modèle coïncide à la réalité.

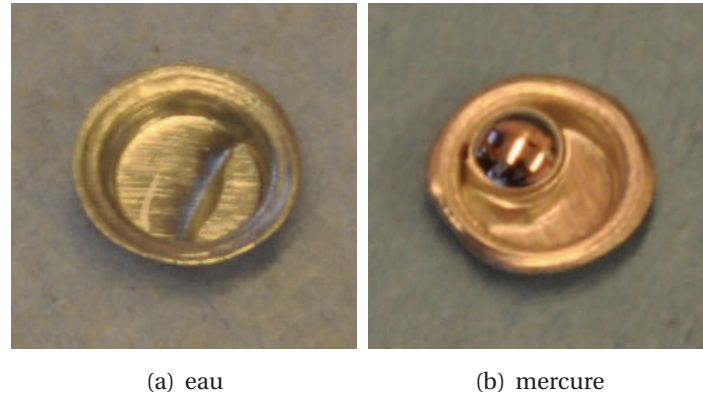


FIGURE 2.42 – Géométrie possible de deux corps, l'eau et le mercure. On tâche généralement de remplir une cellule de façon plus correcte.

Ainsi avec une géométrie écartée de la réalité, car on ne connaît pas la géométrie exacte, on obtient des résultats qui concordent avec la réalité. En sera-t-il de même pour un modèle avec une géométrie plus simple? On a vu dans la section précédente 2.4 qu'une modification de la géométrie entraînait en particulier une modification de la pente à l'origine, ce qui pouvait être compensé par une modification des coefficients d'échange α_i . En comparant les figures 2.14, 2.15, 2.16 et 2.41, on peut se demander si la déformation du thermogramme induite par la réduction de modèle ne pourrait pas être compensée par une adaptation ou ajustement des paramètres de transfert α , λ_S et λ_L . Cela serait intéressant, si comme c'est le cas avec la surface libre, un autre modèle plus simple, comme une sphère à une dimension, pouvait en modifiant les propriétés de transfert, reproduire le comportement énergétique du matériau indépendamment des conditions opératoires. Alors étant donné qu'un modèle 1D, un modèle réduit, est beaucoup plus facile à résoudre, car il nécessite beaucoup moins de noeuds, nous gagnerons en temps de calcul. Car il faut l'avouer, le modèle 2D est très lent, dans notre cas il nécessite le carré du nombre de noeuds nécessaire en 1D. La différence en nombre de calculs nécessaire et donc en temps de simulation, explose rapidement. On le voit d'ailleurs dans le tableau 2.3. Un temps de calcul important est très pénalisant pour une application de cette méthode. Il serait donc intéressant de développer un modèle à une dimension pour le gain en temps de calcul.

Tableau 2.3 – Influence de la réduction sur le temps de résolution du modèle de l'eau pour $\beta = 5 \text{ K} \cdot \text{min}^{-1}$, $m = 10,3 \text{ mg}$ et un pas de temps constant de 10^{-4} s

nombre de maille en 2D	temps 2D (s)	nombre de maille en 1D	temps 1D (s)
10 × 10	76	10	6,1
20 × 20	290	20	9,8
30 × 30	639	30	13,2

2.5.2 Réduction de modèle pour un corps pur [40]

Nous reprenons donc le même modèle que précédemment, mais au lieu d'avoir une géométrie cylindrique, nous avons une géométrie sphérique. Et au lieu d'avoir trois coefficients d'échange équivalents, nous en avons un seul comme le montre la figure 2.43, ce qui est plus facile à identifier. Pour ce qui en est du maillage, nous prendrons un maillage 1D à symétrie sphérique, qui est résolu de la même manière que le modèle 2D en modifiant certaines équations en accord avec la géométrie. Ces modifications sont présentées en suivant.

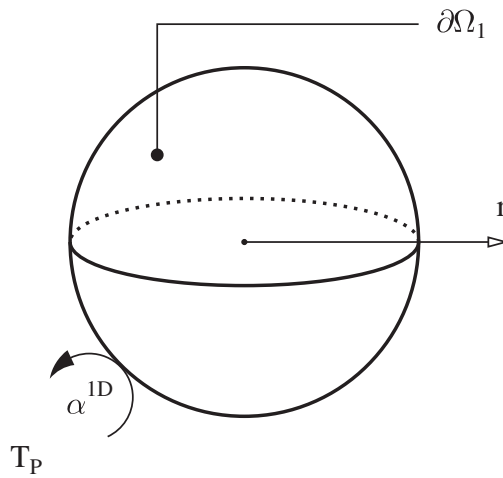


FIGURE 2.43 – Géométrie sphérique

On conserve la relation du maillage radial du cylindre pour le maillage radial de la sphère avec pour seul indice i . Ce faisant les mailles ne sont plus isovolumes.

$$r_i = R \sqrt{\frac{i-1}{NI}} \quad \forall i \in [1; NI+1] \quad (2.70)$$

$$S_{supérieur,i} = 4\pi r_i^2 \quad \forall i \in [1; NI] \quad (2.71)$$

$$S_{inférieur,i} = 4\pi r_{i-1}^2 \quad \forall i \in [2; NI] \quad (2.72)$$

$$V_i = \frac{4}{3}\pi(r_{i+1}^3 - r_i^3) \quad \forall i \in [1; NI] \quad (2.73)$$

Le bilan énergétique sera toujours résolu selon l'équation (2.66). A ceci près qu'il n'y a plus que deux flux rentrant à prendre en compte pour le calcul de chaque maille, le flux supérieur et le flux inférieur qui s'expriment comme suit :

$$\phi_{supérieur,i} = \lambda S_{supérieur,i} \frac{T_{i+1}^{iter} - T_i^{iter}}{r_{i+1} - r_i} \quad \forall i \in [1; NI - 1] \quad (2.74)$$

$$\phi_{inférieur,i} = -\lambda S_{inférieur,i} \frac{T_i^{iter} - T_{i-1}^{iter}}{r_i - r_{i-1}} \quad \forall i \in [2; NI] \quad (2.75)$$

$$\phi_{supérieur,NI} = \lambda S_{supérieur,NI} \frac{T_p^{iter} - T_{NI}^{iter}}{\frac{\|\vec{d}\|}{2} + \frac{1}{\alpha}} \quad (2.76)$$

$$\phi_{inférieur,1} = 0 \quad (2.77)$$

$\vec{d}$ étant ici le vecteur distance (voir figure 2.8) entre le centre de la maille NI et celui d'une maille fictive construite par symétrie plane du centre NI au travers de la surface supérieure NI .

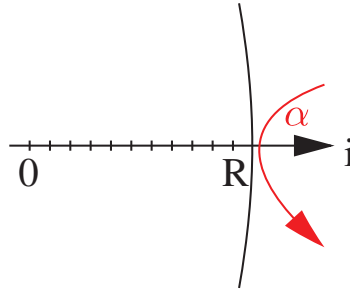


FIGURE 2.44 – Maillage sphérique

Dans le cas 1D sphérique, le coefficient d'échange équivalent, α^{1D} peut être déduit de considérations géométriques, car comme nous l'avons déjà vu, équation (1.16), la pente de début d'une fusion type corps pur est égale au produit αS . C'est-à-dire que on peut l'exprimer par :

$$\alpha^{1D} = \frac{S^{2D}}{S^{1D}} \alpha^{2D} \quad (2.78)$$

$$\text{Avec } \alpha^{2D} = \frac{\alpha_1 S_1^{2D} + \alpha_2 S_2^{2D} + \alpha_3 S_3^{2D}}{S_1^{2D} + S_2^{2D} + S_3^{2D}}.$$

Nous cherchons à déterminer la courbe enthalpie d'un échantillon indépendamment de la vitesse de réchauffement. La réduction de modèle doit donc reproduire les thermogrammes expérimentaux en utilisant les vraies valeurs des paramètres énergétiques des équations d'états. Seuls les paramètres de transfert, α , λ_S et λ_L , dont ne dépend pas l'enthalpie, peuvent être éventuellement modifiés.

Passons maintenant à des simulations numériques pour vérifier l'équivalence entre les deux modèles dans un cas simplifié, celui de l'uniformité des coefficients d'échange équivalent α_i . Considérons d'abord un échantillon d'eau (7,8 mg) réchauffé à une vitesse de $5 \text{ K} \cdot \text{min}^{-1}$, on prend $\rho = 1000 \text{ kg} \cdot \text{m}^{-3}$, $\alpha = 1042 \text{ W} \cdot \text{m}^{-2} \cdot \text{K}^{-1}$, les autres paramètres sont ceux du tableau 2.1. Les thermogrammes obtenus pour un échantillon avec une surface libre plane et un unique

coefficient d'échange, sous les deux modèles, sont comparés dans la figure 2.45. Comme on peut le voir, les deux courbes correspondent. Après identification avec le modèle 1D sur une courbe générée avec le modèle 2D, les seuls paramètres qui ont été modifiés sont les paramètres de transferts α , λ_S et λ_L . Le coefficient d'échange a bien été modifié selon la relation (2.78).

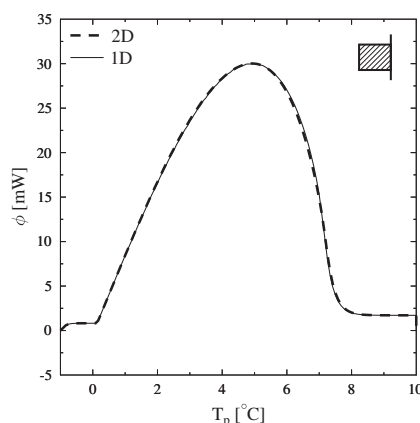


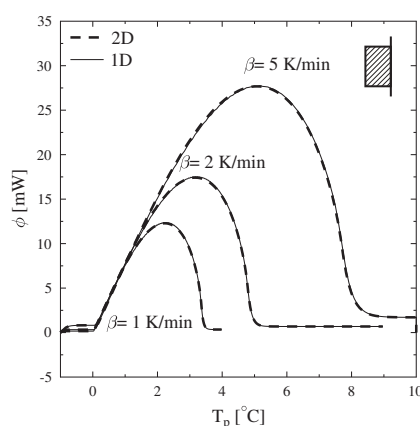
FIGURE 2.45 – Comparaison 1D/2D pour l'eau

Dans l'étude paramétrique du corps pur, nous avons remarqué que la conductivité solide n'avait pas d'influence sur le pic de fusion du corps pur 2.3.1.1, on reproduit le même thermogramme avec des valeurs différentes de λ_S . On ne peut donc rien dire de l'influence de la réduction de modèle sur la conductivité thermique solide.

Nous allons maintenant vérifier si la modification des paramètres α et λ_L est affectée par la vitesse de réchauffement. Dans la figure 2.46 nous avons vérifié cette caractéristique en considérant trois vitesses de réchauffement, en conservant le même ajustement aux paramètres dans chaque cas. A chaque fois nous obtenons une excellente concordance des thermogrammes, les pentes à l'origine sont superposées pour les différentes vitesses preuve qu'il s'agit bien des mêmes coefficients d'échange, et on voit sur le tableau 2.4 que les conductivités liquides ajustées pour le modèle réduit sont les mêmes. La vitesse de réchauffement n'a donc pas d'influence sur les paramètres du modèle.

Tableau 2.4 – Conductivité ajustée entre les modèles 2D et 1D pour un même échantillon mais pour plusieurs vitesses, $r^{2D} = 1 \text{ mm}$ $z^{2D} = 1.56 \text{ mm}$

$\beta (K \cdot \text{min}^{-1})$	$(\frac{\lambda^{1D}}{\lambda^{2D}})_L$
2	0,755
5	0,755
10	0,755

FIGURE 2.46 – Comparaison 1D/2D selon β en adaptant α , λ_S et λ_L

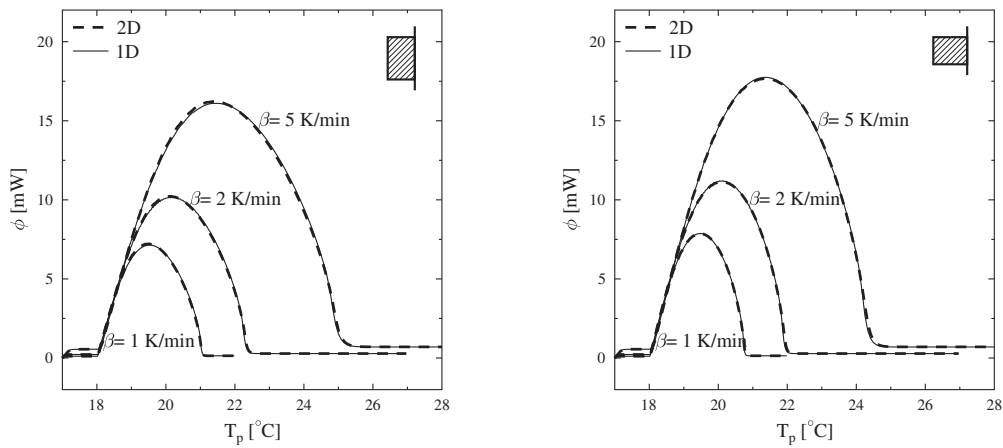
Les conductivités thermiques ajustées étant indépendantes de la vitesse de réchauffement, on suppose une forte dépendance de la modification du rapport des conductivités $(\frac{\lambda^{1D}}{\lambda^{2D}})_L$ par rapport au rapport des surfaces 1D et 2D, $\frac{S^{2D}}{S^{1D}}$. On considère donc quatre formes cylindriques isovolumes pour continuer l'étude de rayon $r = 1$ mm, 1,25 mm, 1,5 mm et $r = 2,15$ mm qui est le rayon interne de notre calorimètre ainsi qu'une dernière forme qui correspond à nos expériences de calorimétrie, elles sont décrites dans le tableau 2.5.

Tableau 2.5 – Géométrie de différents échantillons pour les simulations

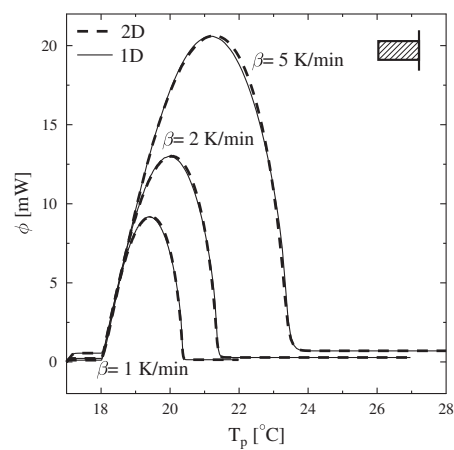
Formes	r^{2D} mm	z^{2D} mm	r^{1D} mm	$\frac{S^{2D}}{S^{1D}}$ —
	1,00	1,5435	1,05	1,153
	1,25	0,98784	1,05	1,265
	1,50	0,686	1,05	1,481
	2,15	0,334	1,05	2,42
	2,15	0,71	1,35	1,696

Enfin, nous avons suivi la même approche pour d'autres matériaux. Nous avons considéré un échantillon de 3,8 mg d'hexadécane dans les mêmes conditions que précédemment. Ici encore nous avons vérifié que seule la géométrie influait. A chaque fois nous avons obtenu une parfaite superposition des thermogrammes 1D et 2D en adaptant les conductivités thermiques, α^{1D} étant fixé par la pente à l'origine (1.16). Ces résultats sont visibles dans les figures 2.47 et 2.48. Les résultats de l'ajustement des paramètres de transferts sont donnés dans le tableau 2.6. On y voit que plus le rapport de géométrie $\frac{S^{2D}}{S^{1D}}$ augmente, donc que le

cylindre est d'avantage aplati ou étiré, au plus la conductivité liquide est modifiée. Il semble que nous ayons un optimum pour lequel le rapport des conductivités serait égal à 1, dans ce cas des coefficients d'échanges uniformes cet optimum est pour une géométrie intermédiaire entre le cylindre haut et le cylindre plat.



(a) échantillon d'hexadécane avec $r^{2D} = 1$ mm. (b) échantillon d'hexadécane avec $r^{2D} = 1,25$ mm.



(c) échantillon d'hexadécane avec $r^{2D} = 1,5$ mm.

FIGURE 2.47 – Comparaison des modèles 2D et 1D pour différentes configurations d'un échantillon d'hexadécane, cas de coefficients d'échange non-uniforme à 10 % [40]

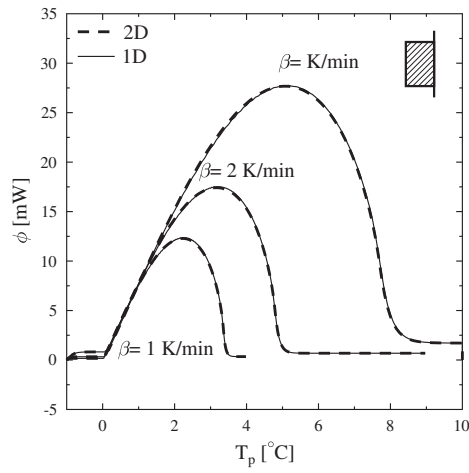
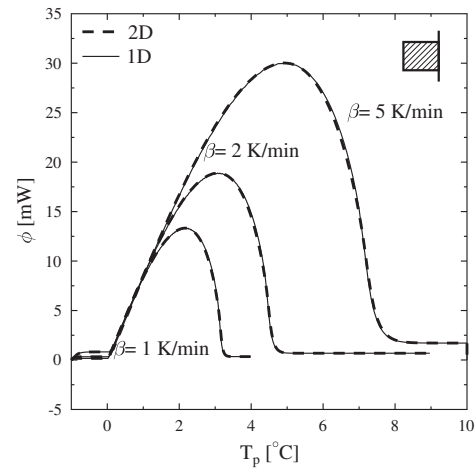
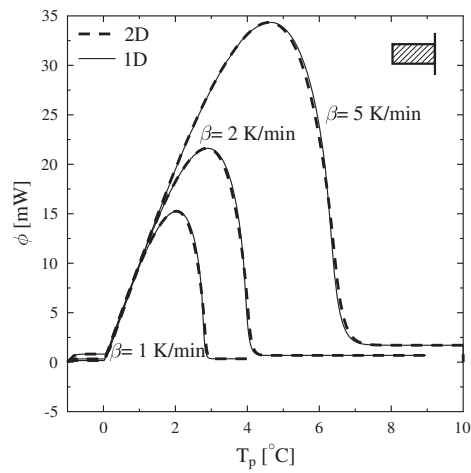
(a) échantillon d'eau avec $r^{2D} = 1$ mm.(b) échantillon d'eau avec $r^{2D} = 1,25$ mm.(c) échantillon d'eau avec $r^{2D} = 1,5$ mm.

FIGURE 2.48 – Comparaison des modèles 2D et 1D pour différentes configurations d'un échantillon d'eau, cas de coefficients d'échange non-uniforme à 10 % [40]

Tableau 2.6 – Conductivité ajustée entre les modèles 2D et 1D (cas du coefficient d'échange uniforme). [40]

Matériaux	masse (mg)	$r^{2D}(mm)$	$\frac{S^{2D}}{S^{1D}}$	$(\frac{\lambda^{1D}}{\lambda^{2D}})_L$
eau	4,85	1,00	1,153	0,98
	4,85	1,25	1,265	1,22
	4,85	1,50	1,481	1,88
	4,85	2,15	2,42	7,83
	10,31	2,15	1,686	3,2
hexadécane	4,85	1,00	1,153	0,98
	4,85	1,25	1,265	1,29
	4,85	1,50	1,481	1,95
	4,85	2,15	2,42	9,3
	10,31	2,15	1,686	3,27

Par la suite, nous avons testé le cas où les coefficients d'échanges équivalents α du cylindre ne sont pas tous égaux. Cette non-uniformité provient de ce que le coefficient d'échange de la face supérieure a une valeur plus faible que celles des autres faces, du fait de l'imperfection du remplissage des creusets. Expérimentalement cela se traduit par la présence d'une couche d'air. Nous avons deux possibilités, soit la couche d'air est importante, soit elle est faible. Nous regarderons deux cas. Un premier cas où on supposera une grosse couche d'air et donc que le coefficient α sur la face supérieur est égale à 10% des α sur les autres faces (voir tableau 2.7). Et un deuxième cas où on supposera une couche d'air moins importante, le rapport des coefficients d'échange valant cette fois 60% (voir tableau 2.8). Selon la même approche, les coefficients de conductivités du modèle 1D sont adaptés afin que les deux thermogrammes se superposent.

Tableau 2.7 – Conductivité ajustée entre les modèles 2D et 1D (cas du coefficient d'échange non-uniforme, 1000 et 100 $W \cdot m^{-2} \cdot K^{-1}$).

Matériaux	masse (mg)	$r^{2D}(mm)$	$\frac{S^{2D}}{S^{1D}}$	$(\frac{\lambda^{1D}}{\lambda^{2D}})_L$
eau	4,85	1,00	1,153	0,78
	4,85	1,25	1,265	0,79
	4,85	1,50	1,481	1,01
	4,85	2,15	2,42	2,58
	10,31	2,15	1,686	1,21
hexadécane	4,85	1,00	1,153	0,84
	4,85	1,25	1,265	0,89
	4,85	1,50	1,481	1,13
	4,85	2,15	2,42	3,24
	10,31	2,15	1,686	1,57

Dans ce tableau 2.7 où l'on considère une importante non-uniformité des coefficients α_i , on remarque une différence par rapport au cas uniforme, tableau 2.6. En effet le rapport $(\frac{\lambda^{1D}}{\lambda^{2D}})_L$ se rapproche de 1 pour une géométrie aplatie. Ceci est dû à ce que le coefficient d'échange supérieur a une valeur plus faible que ceux des autres faces. Or le modèle 1D considère que

l'échange est uniforme sur toute la surface.

Tableau 2.8 – Conductivités identifiées entre les modèles 2D et 1D (cas du coefficient d'échange non-uniforme, 1000 et 600 $W \cdot m^{-2} \cdot K^{-1}$).

Matériaux	masse (mg)	$r^{2D}(mm)$	$\frac{S^{2D}}{S^{1D}}$	$(\frac{\lambda^{1D}}{\lambda^{2D}})_L$
eau	4,85	1,00	1,153	0,96
	4,85	1,25	1,265	1,17
	4,85	1,50	1,481	1,52
	4,85	2,15	2,42	5,45
	10,31	2,15	1,686	2,49
hexadécane	4,85	1,00	1,153	0,98
	4,85	1,25	1,265	1,24
	4,85	1,50	1,481	1,96
	10,31	2,15	1,686	2,91

Dans le tableau 2.8 où la non-uniformité des coefficients α est moins marquée, on retrouve un des résultats du cas uniforme, à savoir qu'il semble y avoir un optimum pour une géométrie intermédiaire au cylindre haut et le cylindre plat au lieu d'un optimum pour un cylindre plat. La validité de l'approximation du modèle 1D comme quoi l'échange est uniforme sur toute la surface de l'échantillon pourrait être acceptable si le rapport des coefficients d'échange est suffisamment proche de 1. Une étude de cet optimum est présentée dans l'annexe A.

Avec ces tableaux 2.6, 2.7 et 2.8, on voit bien que l'on a une forte influence de la géométrie et des valeurs des coefficients d'échange sur le rapport des conductivités. Cette variation est très importante et seules les méthodes inverses, que nous aborderons dans le chapitre 3, peuvent déterminer cette conductivité apparente.

2.5.3 Réduction de modèle pour un binaire [68]

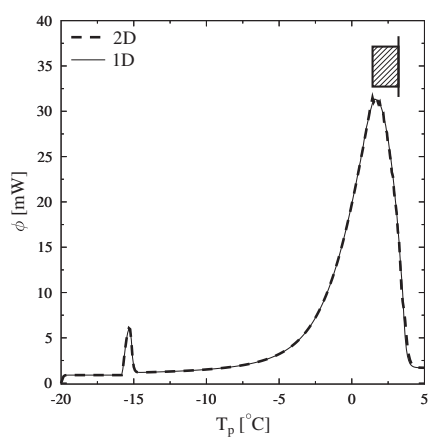
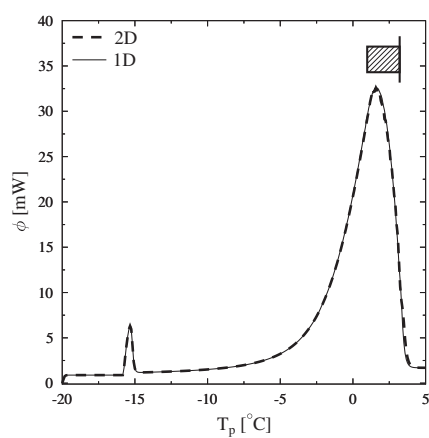
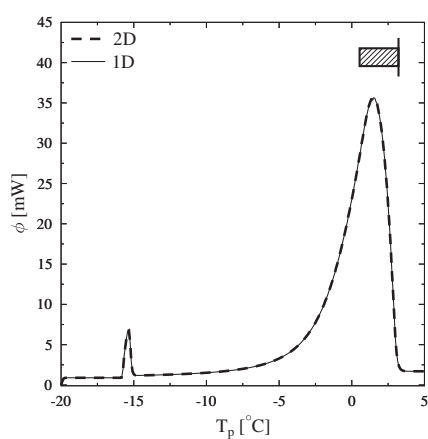
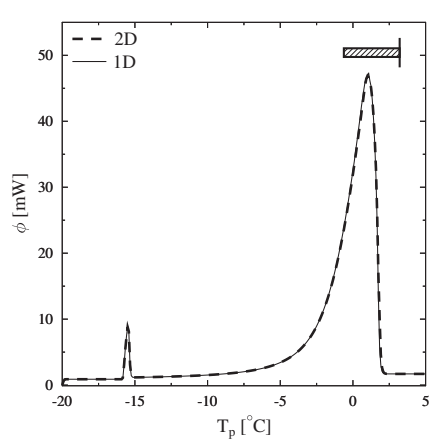
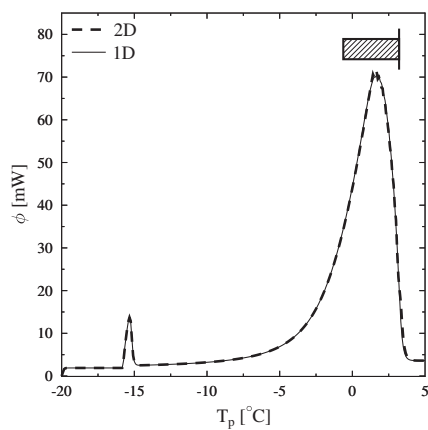
Une étude similaire a été conduite sur le cas d'une solution binaire, de même, le rapport des conductivités ne semble dépendre principalement que du rapport des surfaces, nous n'avons considéré que le cas où le rapport des coefficients α valait 60%. Les conductivités identifiées sont regroupées dans le tableau 2.9.

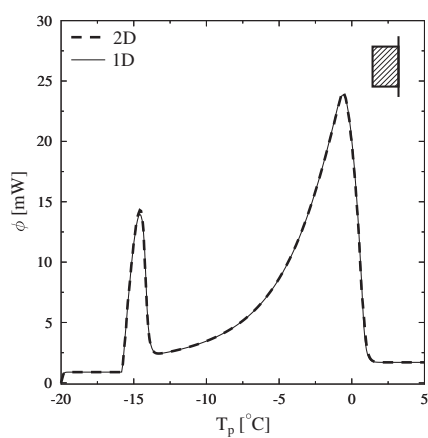
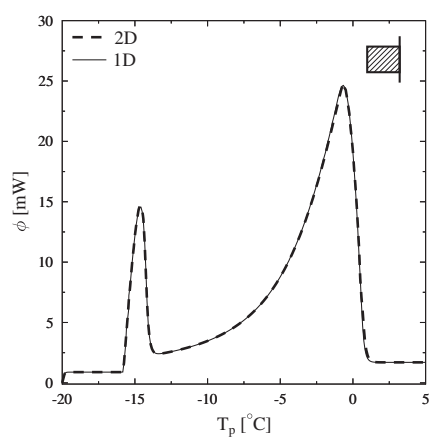
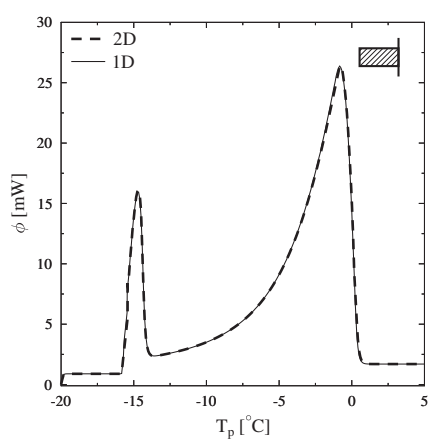
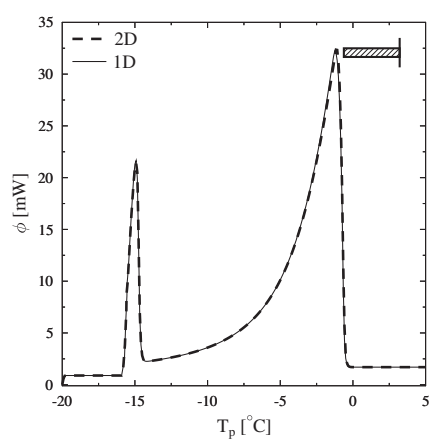
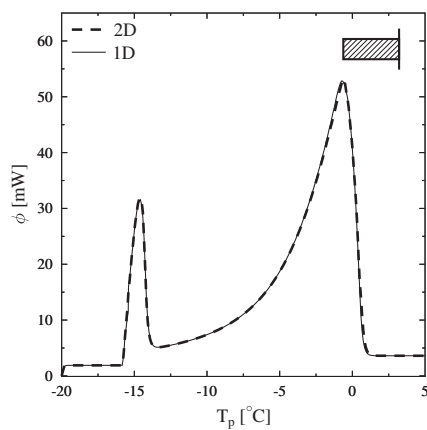
Tableau 2.9 – Conductivité ajustée entre les modèles binaires 2D et 1D (cas du coefficient d'échange non-uniforme).

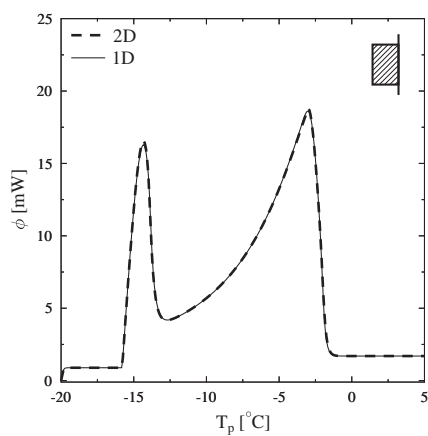
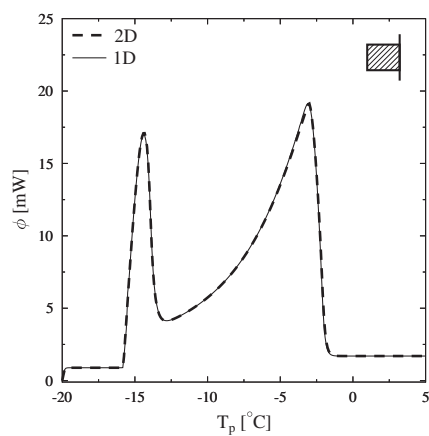
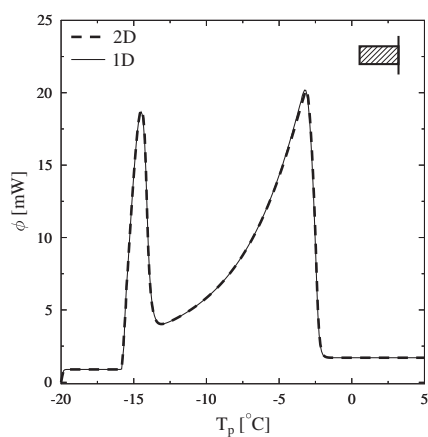
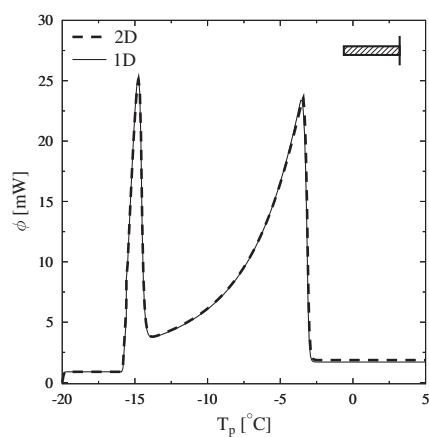
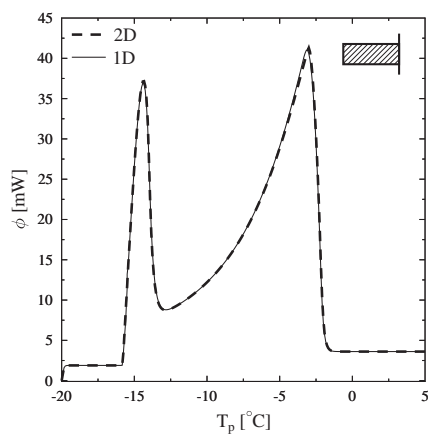
Matériaux	masse (mg)	$r^{2D}(mm)$	$\frac{S^{2D}}{S^{1D}}$	$(\frac{\lambda^{1D}}{\lambda^{2D}})_S$	$(\frac{\lambda^{1D}}{\lambda^{2D}})_L$
$H_2O - NH_4Cl : 0,57\%$	4,85	1,00	1,153	1,1	0,9
	4,85	1,25	1,265	1,1	0,98
	4,85	1,50	1,481	1,53	1,33
	4,85	2,15	2,42	3,75	3,4
	10,31	2,15	1,686	2,15	1,85
$H_2O - NH_4Cl : 2,5\%$	4,85	1,00	1,153	1,1	0,9
	4,85	1,25	1,265	1,1	0,98
	4,85	1,50	1,481	1,47	1,3
	4,85	2,15	2,42	3,77	3,42
	10,31	2,15	1,686	2,09	1,72
$H_2O - NH_4Cl : 5\%$	4,85	1,00	1,153	1,1	0,9
	4,85	1,25	1,265	1,1	0,98
	4,85	1,50	1,481	1,52	1,33
	4,85	2,15	2,42	3,82	3,37
	10,31	2,15	1,686	2,12	1,83
$H_2O - NH_4Cl : 10\%$	4,85	1,00	1,153	1,1	0,9
	4,85	1,25	1,265	1,1	0,98
	4,85	1,50	1,481	1,48	1,34
	4,85	2,15	2,42	3,99	4,43
	10,31	2,15	1,686	2,14	1,92

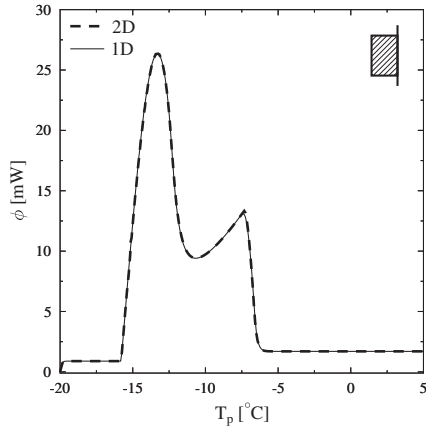
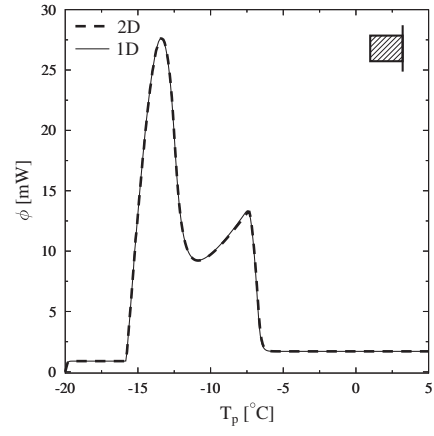
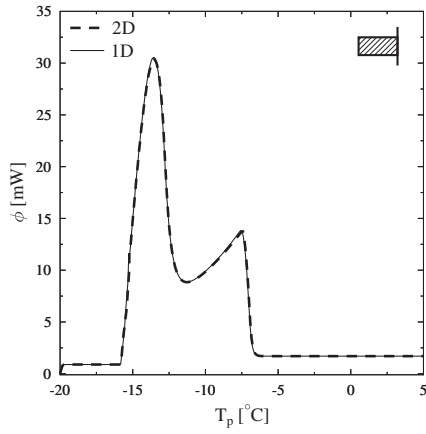
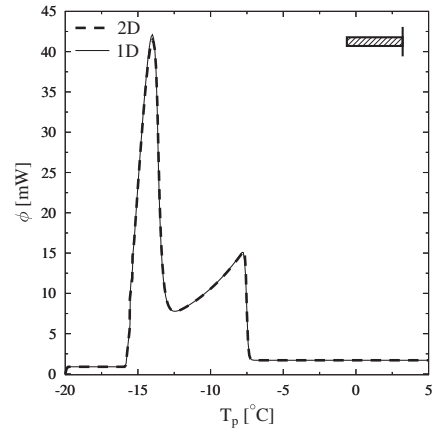
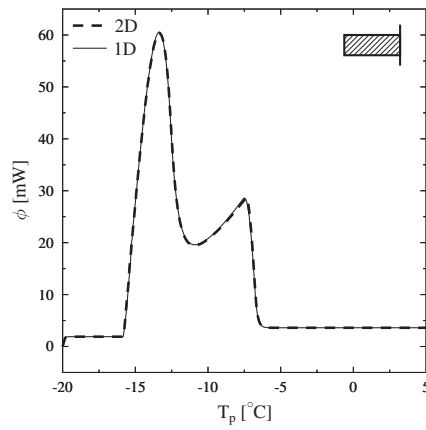
Ici encore il semble y avoir un optimum dans la relation liant le rapport des surfaces 1D et 2D et le rapport des conductivités 1D et 2D si le rapport des coefficients α n'est pas trop faible.

Le modèle 1D permet de reproduire fidèlement les thermogrammes de la DSC comme le modèle 2D mais beaucoup plus rapidement. Cette réduction de modèle a un prix, celui de devoir adapter les propriétés de transferts pour compenser le changement de géométrie. Cette adaptation est indépendante de la vitesse de réchauffement. Ce modèle 1D est mieux adapté aux méthodes inverses que l'on va aborder dans le chapitre suivant car il est plus rapide, comporte moins de paramètres à identifier et il utilise les mêmes valeurs des propriétés énergétiques.

(a) $r^{2D} = 1,00$ mm.(b) $r^{2D} = 1,25$ mm.(c) $r^{2D} = 1,50$ mm.(d) $r^{2D} = 2,15$ mm.(e) $r^{2D} = 2,15$ mm et $m = 10,31$ mg.FIGURE 2.49 – Comparaison 1D et 2D pour $x_0 = 0,57\%$
 $\beta = 5 \text{ K} \cdot \text{min}^{-1}$

(a) $r^{2D} = 1,00$ mm.(b) $r^{2D} = 1,25$ mm.(c) $r^{2D} = 1,50$ mm.(d) $r^{2D} = 2,15$ mm.(e) $r^{2D} = 2,15$ mm et $m = 10,31$ mg.FIGURE 2.50 – Comparaison des modèles 1D et 2D pour $x_0 = 2,5\%$
 $\beta = 5 \text{ K} \cdot \text{min}^{-1}$

(a) $r^{2D} = 1,00$ mm.(b) $r^{2D} = 1,25$ mm.(c) $r^{2D} = 1,50$ mm.(d) $r^{2D} = 2,15$ mm.(e) $r^{2D} = 2,15$ mm et $m = 10,31$ mg.FIGURE 2.51 – Comparaison des modèles 1D et 2D pour $x_0 = 5\%$
 $\beta = 5 \text{ K} \cdot \text{min}^{-1}$

(a) $r^{2D} = 1,00$ mm.(b) $r^{2D} = 1,25$ mm.(c) $r^{2D} = 1,50$ mm.(d) $r^{2D} = 2,15$ mm.(e) $r^{2D} = 2.15$ mm et $m = 10,31$ mg.FIGURE 2.52 – Comparaison des modèles 1D et 2D pour $x_0 = 10\%$
 $\beta = 5 \text{ K} \cdot \text{min}^{-1}$

2.6 Conclusion

Nous disposons d'un modèle numérique 2D qui simule correctement l'évolution thermique d'un corps en changement de phase solide-liquide au sein d'un calorimètre. Il s'agit d'un modèle de transfert thermique couplé à l'équation d'état correspondant à l'échantillon donnant la forme de l'enthalpie en fonction de paramètres du modèle. Ce modèle a été ensuite validé pour des corps purs et un type de solution binaire. Une étude paramétrique de ce modèle a été conduite pour évaluer ou vérifier le rôle de chaque paramètre. Nous nous sommes ensuite posé la question de l'influence de la forme de l'échantillon sur le thermogramme. Et après avoir comparé cette étude à l'étude paramétrique précédente, nous nous sommes orienté vers le développement et la validation d'un modèle réduit afin de gagner en temps de calcul, qui est un des principaux facteurs limitant des méthodes inverses. Mais de ce fait, on identifie pas la vraie valeur de la conductivité, mais de sa transformée par la réduction de modèle et on ne sait pas remonter à la vraie valeur. Il existe une géométrie cylindrique optimale pour laquelle la représentation par le modèle 1D n'affecterait pas ou presque pas les conductivités. Mais la modification des conductivités dépend aussi du degré d'homogénéité des coefficients convectifs équivalent α_i . Plus de détails sont disponibles en annexe [A](#). Maintenant que le modèle est prêt, intéressons-nous à ce qu'est l'inversion et quelle méthode nous allons employer.

Chapitre 3

Principe des méthodes d'inversion pour la calorimétrie

3.1 Qu'est-ce-que l'inversion ?

D'un point de vue « physique » ou « expérimental », on qualifie volontiers de problème inverse toute situation où l'on souhaite évaluer certaines grandeurs physiques ou paramètres $\mathbf{p} = {}^t(p_1, \dots, p_{n_p})$ inaccessibles par l'expérience, à partir de la mesure d'une autre grandeur $\mathbf{y}_{exp}$, qui, elle, est directement accessible par l'expérience. On le fait grâce à un modèle mathématique du problème direct qui calcule explicitement cette même grandeur $\mathbf{y}$ à partir de $\mathbf{p}$ (ce que l'on note symboliquement $\mathbf{y} = \mathbb{M}(\mathbf{p})$, $\mathbb{M}$ est ici le modèle mathématique). En termes mathématiques, cette définition conduit à appeler « problème inverse » la détermination de $\mathbf{p}$ par une équation (algébrique, matricielle, différentielle, aux dérivées partielles, intégrale...) ou par la minimisation de toute fonctionnelle que l'on appelle alors fonction objectif. Une fonction objectif est par exemple le coût d'une installation industrielle dont on essaye d'optimiser le fonctionnement, la distance parcourue si on cherche le plus court chemin sur le trajet d'un commercial, comme ce fut le cas dans le problème posé par le mathématicien viennois, Karl Menger [69], ou la somme des carrés des écarts (méthode des moindres carrés).

Tous ces « problèmes inverses » peuvent être répartis en deux catégories. Ceux qui, pour toute mesure $\mathbf{y}$, admettent une solution unique $\mathbf{p}$ continue par rapport à $\mathbf{y}$ sont dits « bien posés » [70]. Ils sont généralement résolus (de manière analytique ou numérique) par des méthodes classiques. Les autres sont dits « mal posés ». La difficulté de cette deuxième catégorie est qu'elle ne vérifie pas toutes les conditions d'Hadamard [71] soit l'existence de la solution, son unicité et/ou la continuité de la solution par rapport aux mesures. Du point de vue physique, cela signifie qu'une mesure $\mathbf{y}_{exp}$ peut correspondre à un grand nombre de valeurs de $\mathbf{p}$, ces dernières pouvant être fortement éloignées les unes des autres.

Le caractère « mal posé » d'un problème peut avoir plusieurs origines. D'abord l'expérience, les types et points de mesure peuvent ne pas être adaptés au phénomène que étudié. Un exemple pourrait être une mauvaise position des capteurs de températures par rapport au phénomène que étudié. Ensuite la structure mathématique du problème peut ne pas être adaptée à la résolution [70] ou bien un changement de topologie s'impose pour résoudre le problème. Par ailleurs plusieurs types d'erreurs viennent gêner le processus d'inversion [72].

Ils sont abordés dans la sous-section 4.3.2.

La sensibilité des problèmes inverses aux erreurs induit un changement important vis-à-vis du concept de solution, car la recherche des solutions au sens mathématique, associées par le modèle aux mesures $\mathbf{y}$, n'est plus un objectif suffisant. En effet, pour un problème « mal posé » tout $\mathbf{p}$ qui reproduit aux incertitudes près, via le modèle physique, la mesure $\mathbf{y}$ est une réponse a priori possible au problème inverse. Un problème inverse, pour un modèle physique et une mesure donnée, peut n'avoir aucune solution au sens strict mais beaucoup de quasi-solutions « à ε_1 près ».

Toute « théorie de l'inversion » doit donc tenir compte du caractère éventuellement incomplet, imprécis et/ou redondant des données. La stratégie idéale consisterait à inventorier l'ensemble complet des solutions « à ε_1 près » de $\mathbf{y} = \mathbb{M}(\mathbf{p})$, parmi lesquelles on opérerait ensuite un choix suivant des critères additionnels (vraisemblance physique, informations supplémentaires *a priori*) afin de retenir la ou les solutions jugées vraisemblables. Une telle « approche exhaustive » pose des problèmes pratiques insurmontables si le nombre de variables scalaires à identifier n'est pas suffisamment restreint [73] [74].

3.2 Objectifs des méthodes inverses

Le système ou le modèle mathématique dépend habituellement de paramètres, symboliquement notés $\mathbf{p}$: géométrie (la région de l'espace occupée), caractéristiques des matériaux... Le problème direct consiste à calculer la réponse $\mathbf{y}$ à partir de la donnée des sollicitations X qui regroupent les conditions initiales et aux limites du problème, et les paramètres $\mathbf{p}$. Les équations de la physique du modèle donnent, en général, la réponse $\mathbf{y}$ comme fonction implicite de X et $\mathbf{p}$, voir équation (3.1) et figure 3.1.

$$\mathbb{M}(X, \mathbf{p}) = \mathbf{y} \quad (3.1)$$

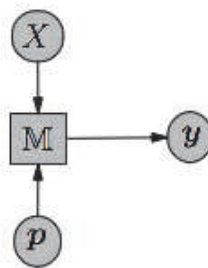


FIGURE 3.1 – Schéma d'un modèle direct

Dans le cas d'un problème inverse, il s'agit généralement de situations où on est dans l'ignorance, au moins partielle, des valeurs de $\mathbf{p}$. En compensation, il faut disposer d'informations, éventuellement partielles, sur la sortie du modèle afin de reconstruire au mieux les informations manquantes. C'est ici qu'interviennent les mesures expérimentales. Ces informations,

qui sont obtenues en comparant la sortie du système réel, y_{exp} , avec la sortie du modèle, y , par le biais d'une norme que l'on choisit $\| \cdot \|$, constituent la fonction objectif :

$$F_{obj}(p) = \|y - y_{exp}\|^2 \leq \varepsilon_2 \quad (3.2)$$

Le terme inverse rappelle qu'on utilise l'information concernant le modèle physique « à l'envers » : connaissant (partiellement) les sorties, on cherche à remonter à certaines caractéristiques, habituellement internes et échappant à la mesure directe.

Un schéma récapitulatif de la formulation du problème inverse est donné dans la figure 3.2.

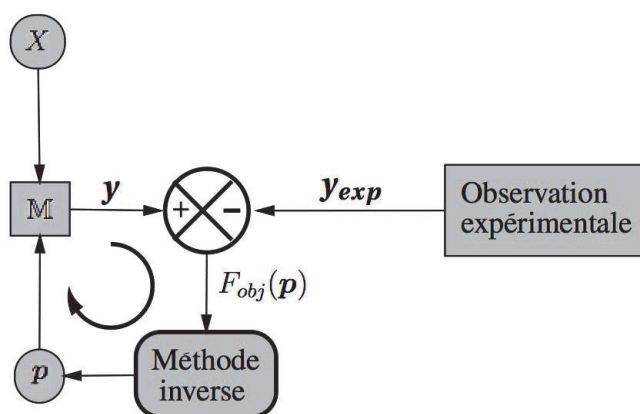


FIGURE 3.2 – Schéma méthodologique de l'inversion

3.3 Méthodes d'inversions

Les méthodes inverses tentent, par modification des paramètres initiaux, de minimiser un critère basé sur la différence entre les sorties mesurées et calculées par le modèle (voir équation (3.2)). Les paramètres optimisés minimisent ainsi cet écart entre l'expérience et la sortie du modèle direct. Les premières méthodes inverses à avoir été développées ont un cadre mathématique très rigoureux, déterministe où, la convergence, l'aboutissement des calculs ont été démontrés. Cette assurance de pouvoir aboutir au résultat a justifié leur développement.

3.3.1 Les méthodes déterministes ou de descentes

Une méthode d'inversion n'est rien d'autre qu'une méthode d'optimisation puisqu'on cherche à minimiser une fonction objectif.

On peut diviser ces méthodes entre les méthodes linéaires et les méthodes non-linéaires par rapport aux paramètres qui nous intéressent, pour résoudre le système d'équations qui relie ces paramètres entre eux. Les méthodes linéaires sont destinées à résoudre des problèmes d'inversion qui peuvent se mettre sous la forme de systèmes linéaires [75]. Ces méthodes, dans le cas d'un problème sans contrainte, reviennent à des méthodes d'inversion de matrice si celles-ci sont inversibles. Sinon, si la matrice associée au problème linéaire n'est

pas inversible, on peut résoudre un problème approché en déterminant la matrice pseudo-inverse ; pour obtenir une pseudo-solution, une méthode utilisée est la décomposition en valeurs singulières ou S.V.D [76].

Notre problème traitant du changement de phase, le système d'équation comportera au moins une non-linéarité. Nous devons donc recourir aux méthodes non-linéaires. Avec les méthodes non-linéaires, l'optimisation des paramètres conduit à minimiser une fonction objectif (F_{obj}). Elle est définie dans le cas d'estimation de paramètres bien souvent comme étant la somme des moindres carrés ou le carré de la norme de l'écart entre les observations et les valeurs calculées numériquement. Par rapport à la définition (3.2) on utilise la norme euclidienne que l'on a pondérée. On considère donc la fonction « objectif » suivante :

$$F_{obj}(\mathbf{p}) = {}^t(\mathbf{y} - \mathbf{y}_{exp}) \cdot \mathbb{W} \cdot (\mathbf{y} - \mathbf{y}_{exp}) \quad (3.3)$$

avec $\mathbf{y} = {}^t(y_1, y_2, \dots, y_{n_m})$ le vecteur sortie ou réponse du modèle. Dans ce travail il s'agit du flux thermique (n_m étant le nombre total d'observations, la taille du vecteur), $\mathbf{y}$ représente la variable d'état calculée en utilisant un modèle numérique. Soit $\mathbf{y}_{exp}$ les valeurs de la variable expérimentale. $\mathbb{W}$ est une matrice qui pondère les éléments de $\mathbf{y} - \mathbf{y}_{exp}$ selon l'importance ou la pertinence qu'on leur accorde, généralement on prend la matrice $\mathbb{W}$ égale à la matrice identité car on n'a pas de raison de discriminer certaines mesures. C'est ce que nous ferons pour le reste de cette étude. Ainsi :

$$F_{obj}(\mathbf{p}) = \|\mathbf{y} - \mathbf{y}_{exp}\|^2 \quad (3.4)$$

On note aussi $\vec{\nabla}(F_{obj})$, $\nabla^2(F_{obj})$ le gradient et la matrice Hessienne ou le Hessian, de la fonction F_{obj} et $\|\cdot\|$ la norme euclidienne.

$$\vec{\nabla}(F_{obj})(\mathbf{p}) = \sum_{i=1}^{n_p} \left(\frac{\partial F_{obj}}{\partial p_i} \right)(\mathbf{p}) \mathbf{e}_{p_i} \quad (3.5)$$

$\mathbf{e}_{p_i}$: vecteurs de base de l'espace des paramètres.

$$\nabla^2(F_{obj})(\mathbf{p}) = \left[\left(\frac{\partial^2 F_{obj}}{\partial p_i \partial p_j} \right)(\mathbf{p}) \right]_{i,j} \quad \forall (i, j) \in [1, \dots, n_p] \times [1, \dots, n_p] \quad (3.6)$$

Comme on cherche à minimiser une fonction objectif, le chemin de convergence emprunte ce que l'on appelle une direction de descente. Une direction de descente $\mathbf{d}$ est une direction de l'espace $\mathbb{R}^n$ vérifie :

$${}^t\vec{\nabla}(F_{obj})(\mathbf{p}) \cdot \mathbf{d} < 0 \quad (3.7)$$

On définit ainsi un demi-espace de descente, l'intersection entre ce demi-espace et le domaine de recherche, soit le domaine de valeurs possibles pour les paramètres, constituant l'ensemble des directions de descentes. On en déduit que $\mathbf{d} = -\mathbb{A} \vec{\nabla}(F_{obj})(\mathbf{p})$ est une direction de descente, si la matrice carrée $\mathbb{A}$ qui modifie la direction par rapport à celle du gradient, est définie positive. En effet l'équation (3.7) devient :

$${}^t\vec{\nabla}(F_{obj})(\mathbf{p}) \cdot \mathbf{d} = -{}^t\vec{\nabla}(F_{obj})(\mathbf{p}) \cdot \mathbb{A} \vec{\nabla}(F_{obj})(\mathbf{p}) < 0 \quad (3.8)$$

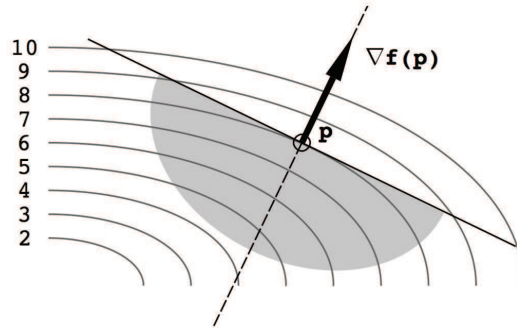


FIGURE 3.3 – Gradient et direction de descente [77]

Par définition de la dérivée, si $\mathbf{d}$ est une direction de descente alors pour tout $\varepsilon > 0$ suffisamment petit, on a :

$$F_{obj}(\mathbf{p} + \varepsilon \times \mathbf{d}) < F_{obj}(\mathbf{p}) \quad (3.9)$$

Au moins localement, la fonction F_{obj} diminue en effectuant un déplacement dans la direction $\mathbf{d}$. Les méthodes à directions de descentes suivent ce principe et construisent une suite d'itérés ($\mathbf{p}^k$) approchant une quasi-solution $\mathbf{p}^*$ du problème de la façon suivante :

$$\mathbf{p}^{k+1} = \mathbf{p}^k + \omega^k \times \mathbf{d}^k \quad (3.10)$$

Le paramètre ω^k est le pas à effectuer le long de la direction de descente $\mathbf{d}^k$ au point courant $\mathbf{p}^k$.

Un algorithme à direction de descente est donc déterminé par les paramètres $\mathbf{d}$ et ω : la façon dont la direction de descente $\mathbf{d}$ est calculée donne son nom à l'algorithme.

Un modèle d'algorithme basé sur des directions de descente est fourni ci-dessous :

- 0) Initialisation $k = 0$: Choix d'un point initial $\mathbf{p}^0 \in \mathbb{R}^n$ et du critère de convergence.
- 1) faire test de convergence pour l'arrêt de l'algorithme
- 2) déterminer une direction de descente $\mathbf{d}^k$ et d'un pas $\omega^k > 0$, tel que la fonction F_{obj} décroisse
- 3) déterminer un nouveau point $\mathbf{p}^{k+1} = \mathbf{p}^k + \omega^k \times \mathbf{d}^k$
- 4) poser $k = k + 1$ et retourner à l'étape 1

La détermination de $\mathbf{d}^k$ et ω^k a donné lieu à de nombreuses méthodes. Nous en aborderons plusieurs par ordre croissant de complexité pour aboutir à la méthode de descente que nous avons utilisée.

3.3.1.1 Méthode du morcellement

En première approche on a voulu déterminer la solution en optimisant chaque paramètre successivement, c'est-à-dire en faisant varier un paramètre tout en bloquant les autres, de sorte à diminuer la fonction objectif dans un cas de minimisation. Puis en recommençant avec un autre paramètre jusqu'à obtenir un point fixe, de sorte que tous les paramètres soient optimisés, c'est la méthode par morcellement [77].

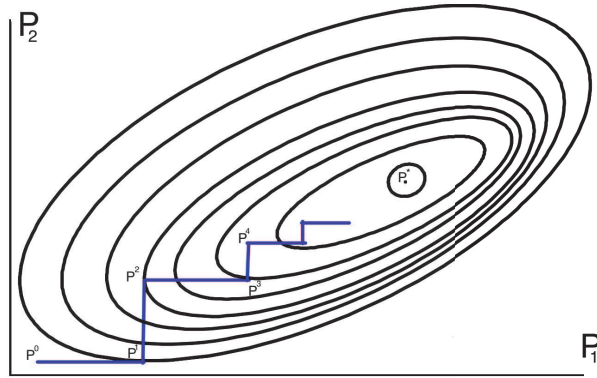


FIGURE 3.4 – Méthode de morcellement dans un cas à deux paramètres [75]

Pour optimiser chaque paramètre on effectue ce qu'on appelle une recherche monodimensionnelle. C'est avec ce type de méthode que l'on cherche ω^k .

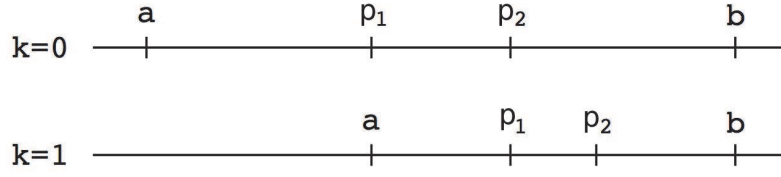
Mais la méthode par morcellement n'est pas forcément la plus rapide car elle n'optimise pas le chemin vers la solution. Ainsi se sont développées les méthodes à gradient, d'ordre 1, qui, elles, optimisent la direction de descente en chaque point.

3.3.1.1.a La recherche monodimensionnelle

La recherche monodimensionnelle consiste à déterminer le pas ω^k à effectuer le long d'une direction de descente $\mathbf{d}^k$, qui minimise la fonction objectif. Soit :

$$\omega^k = \underset{\omega \geq 0}{\text{Min}} \left[F_{obj}(\mathbf{p} + \omega \times \mathbf{d}^k) \right] = \underset{\omega \geq 0}{\text{Min}} [F_{obj}(\omega)] \quad (3.11)$$

On peut repérer le minimum par différentes méthodes comme celle des intervalles constants. Ces méthodes découpent l'intervalle de recherche en sous-intervalles pour en éliminer un selon les valeurs de la fonction objectif aux bornes de chacun des sous-intervalles, réduisant ainsi l'intervalle de recherche à chaque itération. Ces méthodes nécessitent de se donner un intervalle initial dans lequel effectuer la recherche. On a défini cet intervalle avec le point courant et le point qui serait obtenu si $\omega^k = 1$. Nous avons choisi la méthode de la section dorée [77]. Mais une autre méthode plus simple consiste à initialiser l'identification avec $\omega^k = 1$, et à la fin de chaque itération de la méthode inverse. Si le point obtenu améliore le critère alors on augmente ω^{k+1} par rapport à ω^k . Si par contre le calcul donne une moins bonne solution, on recommence l'itération de la méthode inverse avec une valeur inférieure de ω^k , et ainsi de suite jusqu'à ce qu'on obtienne une meilleure solution.

FIGURE 3.5 – Méthode de la section dorée ($F_{obj}(p_1) > F_{obj}(p_2)$) [77]

avec : $p_1 = a \frac{\sqrt{5}-1}{2} + b \frac{3-\sqrt{5}}{2}$ et $p_2 = a \frac{3-\sqrt{5}}{2} + b \frac{\sqrt{5}-1}{2}$.

3.3.1.2 Développements limités d'ordre 1 : utilisation de l'hyperplan tangent

Nous allons présenter cette famille de méthodes avec la méthode de la plus grande pente ou méthode de Cauchy [75] [77]. Elle utilise pour direction de descente l'opposé du gradient de la fonction F_{obj} au point courant $\mathbf{p}$, soit :

$$\mathbf{d}^k = -\vec{\nabla}(F_{obj})(\mathbf{p}^k) \quad (3.12)$$

Cette direction, obtenue par un développement limité du premier ordre est évidemment une direction de descente comme le montre le produit scalaire suivant :

$$\vec{\nabla}(F_{obj})(\mathbf{p}^k) \cdot \mathbf{d}^k = -\|\vec{\nabla}(F_{obj})(\mathbf{p}^k)\|^2 < 0 \quad (3.13)$$

La méthode de la plus grande pente est plus aboutie que la méthode par morcellement mais n'est toujours pas optimum. D'abord car elle nécessite comme la méthode par morcellement, une recherche monodimensionnelle parce qu'elle ne détermine que la direction dans laquelle on doit se déplacer pour améliorer le critère, mais ne fournit pas d'information sur la distance que l'on peut parcourir. Ensuite, comme elle repose sur le gradient de la fonction objectif, on peut s'attendre à ce que proche de l'optimum (minimum si minimalisation) de la fonction objectif le gradient devienne presque nul et donc que la méthode devienne très lente à partir de ce moment là. Cette méthode est plus adaptée aux premières itérations loin de la solution, qu'aux dernières proches de la solution. C'est pourquoi plusieurs variantes ont été développées comme la méthode du gradient conjugué, certaines affinant la détermination en poussant au second ordre comme la méthode de Newton.

3.3.1.3 Développements limités d'ordre 2 : utilisation du paraboloïde tangent

3.3.1.3.a Méthode de Newton

La méthode de Newton utilise un développement limité d'ordre deux qui approxime l'espace de recherche par un paraboloïde qui, lui, dispose d'un minimum. Il n'est donc plus nécessaire de recourir à une recherche monodimensionnelle puisqu'on détermine en même temps $\mathbf{d}^k$ et ω^k connaissant l'extremum (ou minimum) du paraboloïde. Dans la suite on posera donc $\omega^k = 1$. Ainsi, d'un point de vue local, la fonction objectif est approximée par une forme quadratique dépendant du Hessian [77] comme le montre ce développement limité au second ordre :

$$F_{obj}(\mathbf{p}^k + \mathbf{d}^k) \approx F_{obj}(\mathbf{p}^k) + {}^t \mathbf{d}^k \cdot \vec{\nabla}(F_{obj})(\mathbf{p}^k) + \frac{1}{2} {}^t \mathbf{d}^k \cdot \nabla^2(F_{obj})(\mathbf{p}^k) \cdot \mathbf{d}^k \quad (3.14)$$

Pour trouver la direction de descente, il faut trouver $\mathbf{d}^k$ qui minimise le terme de gauche. En dérivant l'expression précédente par rapport à $\mathbf{d}^k$ il vient :

$$\vec{\nabla}(F_{obj})(\mathbf{p}^k) + \nabla^2(F_{obj})(\mathbf{p}^k) \cdot \mathbf{d}^k = 0 \quad (3.15)$$

On a donc :

$$\mathbf{d}^k = -\nabla^2(F_{obj})^{-1} \cdot \vec{\nabla}(F_{obj})(\mathbf{p}^k) \quad (3.16)$$

D'après l'équation (3.8) la matrice Hessienne $\nabla^2(F_{obj})$ doit être définie positive.

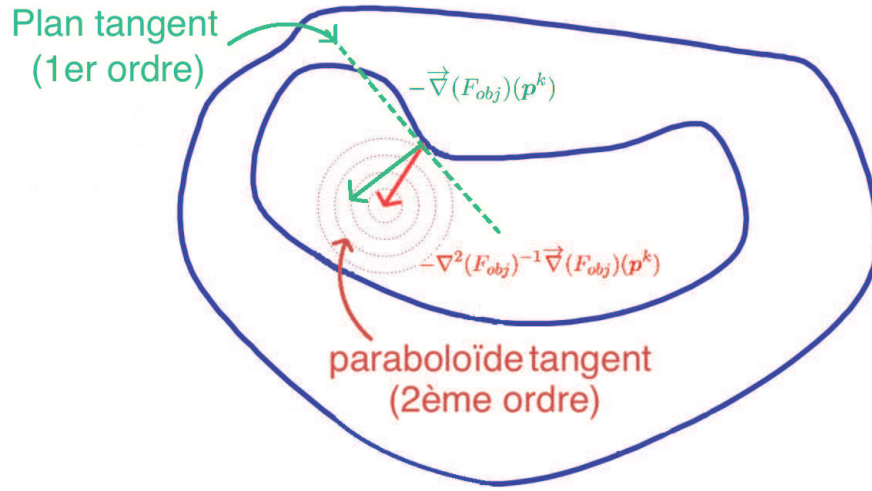


FIGURE 3.6 – Méthode de Newton [75]

Toutes les méthodes précédentes reposent sur la connaissance du gradient et même du Hessian pour certaines. Leurs calculs étant difficiles et leurs valeurs exactes n'étant pas forcément accessibles, on a cherché à les approcher numériquement. Pour notre cas, faute d'avoir une expression analytique de notre fonction objectif nous ne pouvons qu'au mieux les approcher par différences finies.

De nombreuses méthodes ont été développées [71] [78] pour aboutir à la méthode généralisée de Gauss. Nous ne détaillerons pas cette méthode ici mais nous nous attarderons d'avantage sur une version dérivée, avant d'aboutir à une dernière méthode qui tire le meilleur parti des autres méthodes présentées. Parmi les méthodes de descente c'est celle que nous avons utilisée Elle intègre des stratégies de régularisation que nous évoquerons dans la sous-section 3.3.1.4.

3.3.1.3.b Algorithme de Gauss - Newton

Cet algorithme est basé sur celui de Newton en apportant une solution permettant de réduire le coût de calcul de la matrice Hessienne. Le principe est de la remplacer par une approximation moins coûteuses.

En introduisant la matrice Jacobienne ou, comme elle est communément appelée pour ce type d'application, matrice de sensibilité $\mathbb{S}$, qui est la matrice $n_m \times n_p$ des dérivées de la réponse ou sortie du modèle par rapport à ses paramètres, on va chercher à approximer le Hessien.

$$\mathbb{S}(\mathbf{p}) = [s_{i,j}] \text{ avec } s_{i,j} = \frac{\partial y_i(\mathbf{p})}{\partial p_j} \quad \forall (i,j) \in [1, \dots, n_m] \times [1, \dots, n_p] \quad (3.17)$$

Comme $\mathbf{y}_{exp}$ est constante par rapport au modèle, d'après la relation précédente on peut écrire le gradient de la fonction objectif (3.3) sous la forme suivante :

$$\begin{aligned} \vec{\nabla}(F_{obj}(\mathbf{p}^k)) &= \sum_{j=1}^{n_p} \sum_{i=1}^{n_m} 2(y_i - y_{exp}) \frac{\partial y_i}{\partial p_j} \mathbf{e}_j \\ &= 2^t \mathbb{S}(\mathbf{p}^k) \cdot (\mathbf{y} - \mathbf{y}_{exp}) \end{aligned} \quad (3.18)$$

Ainsi en dérivant l'équation (3.18) par rapport à $\mathbf{p}$ on obtient le Hessien :

$$\begin{aligned} \nabla^2(F_{obj})(\mathbf{p}^k) &= 2^t \mathbb{S}(\mathbf{p}^k) \cdot \sum_{j=1}^{n_p} \sum_{i=1}^{n_m} \frac{\partial y_i}{\partial p_j} \mathbf{e}_j + 2^t \vec{\nabla}(\mathbb{S}(\mathbf{p}^k)) \cdot (\mathbf{y} - \mathbf{y}_{exp}) \\ &= 2^t \mathbb{S}(\mathbf{p}^k) \cdot \mathbb{S}(\mathbf{p}^k) + 2^t \vec{\nabla} \left(\sum_{j=1}^{n_p} \sum_{i=1}^{n_m} \frac{\partial y_i - y_{exp,i}}{\partial p_j} \mathbf{e}_j \right) \cdot (\mathbf{y} - \mathbf{y}_{exp}) \\ &= 2^t \mathbb{S}(\mathbf{p}^k) \cdot \mathbb{S}(\mathbf{p}^k) + 2(\nabla^2(\mathbf{y} - \mathbf{y}_{exp})) \cdot (\mathbf{y} - \mathbf{y}_{exp}) \end{aligned} \quad (3.19)$$

car $\mathbf{y}_{exp}$ est constante par rapport aux paramètres.

Dans l'algorithme de Gauss-Newton, contrairement à celui de Newton, on néglige le terme des dérivées secondes. On obtient alors une approximation du Hessien dont l'évaluation ne nécessite plus le calcul de dérivées secondes

$$\nabla^2(F_{obj})(\mathbf{p}^k) \approx 2^t \mathbb{S}(\mathbf{p}^k) \cdot \mathbb{S}(\mathbf{p}^k) \quad (3.20)$$

Comme la méthode de Newton 3.3.1.3.a le pas ω^k est inclus dans le vecteur $\mathbf{d}^k$. Donc en combinant les équations (??), (3.18), (3.19) et (3.20) on obtient la direction de descente de l'algorithme de Gauss-Newton. Il s'agit d'une approximation de celle de la méthode de Newton :

$$\mathbf{d}^k = -[{}^t \mathbb{S}(\mathbf{p}^k) \cdot \mathbb{S}(\mathbf{p}^k)]^{-1} \cdot [{}^t \mathbb{S}(\mathbf{p}^k) \cdot (\mathbf{y} - \mathbf{y}_{exp})] \quad (3.21)$$

Les deux derniers algorithmes nécessitent d'inverser une matrice. Cela nous mène à la question du conditionnement.

3.3.1.4 Conditionnement - Régularisation de Tikhonov

On vient d'aborder une nouvelle notion qui est malheureusement centrale dans les problèmes d'inversion car elle met souvent en échec ces méthodes. Il s'agit du conditionnement. Ce nombre mesure la dépendance de la solution d'un problème numérique par rapport aux données du problème, ceci afin de contrôler la validité d'une solution calculée avec

ces données. Comme nous appliquons ces méthodes sur des mesures expérimentales, un problème avec mauvais nombre de conditionnement sera trop sensible aux erreurs de mesures pour être résolue tel quel.

Les méthodes d'ordre 2 nécessitent d'inverser une matrice afin de déterminer la direction de descente. Le conditionnement c'est l'influence des petites variations autour des valeurs de la sortie du système sur l'entrée identifiée suite à l'inversion matricielle. C'est la qualité de l'inversion du problème. C'est ici en outre qu'intervient le problème posé par les incertitudes évoquées en début de chapitre ainsi que les approximations liées à la méthode d'inversion (différences finies pour la plupart).

Pour illustrer ce propos, nous nous placerons dans un cas d'un système linéaire et inversible de matrice $\mathbb{A}$:

$$\mathbb{A} \cdot \mathbf{p} = \mathbf{Y} \quad (3.22)$$

Supposons maintenant que la sortie $\mathbf{Y}$ subit une légère variation $\delta \mathbf{Y}$. Quelle va être la variation $\delta \mathbf{p}$ que l'on va observer sur les paramètres à identifier $\mathbf{p}$?

$$\mathbf{Y} + \delta \mathbf{Y} = \mathbb{A} \cdot (\mathbf{p} + \delta \mathbf{p}) \quad (3.23)$$

En comparant les équations (3.22) et (3.23), il vient :

$$\begin{aligned} \delta \mathbf{Y} &= \mathbb{A} \cdot \delta \mathbf{p} \\ \Leftrightarrow \delta \mathbf{p} &= \mathbb{A}^{-1} \cdot \delta \mathbf{Y} \\ \Rightarrow \|\delta \mathbf{p}\| &= \|\mathbb{A}^{-1} \cdot \delta \mathbf{Y}\| \end{aligned} \quad (3.24)$$

En conséquence de l'inégalité triangulaire on a de plus :

$$\|\delta \mathbf{p}\| \leq \|\mathbb{A}^{-1}\| \times \|\delta \mathbf{Y}\| \quad (3.25)$$

En appliquant le même raisonnement à l'équation (3.22) on a :

$$\frac{1}{\|\mathbf{p}\|} \leq \|\mathbb{A}\| \times \frac{1}{\|\mathbf{Y}\|} \quad (3.26)$$

En multipliant les termes de gauche des équations (3.22) et (3.23) nous avons :

$$\frac{\|\delta \mathbf{p}\|}{\|\mathbf{p}\|} \leq \|\mathbb{A}\| \times \|\mathbb{A}^{-1}\| \times \frac{\|\delta \mathbf{Y}\|}{\|\mathbf{Y}\|} \quad (3.27)$$

Il en ressort que l'écart sur l'estimation des paramètres est borné par un multiplicatif de l'écart sur la sortie du système. Le scalaire qui relie les deux termes est appelé le nombre de conditionnement. Mathématiquement, le conditionnement d'une matrice est défini pour toute matrice inversible $\mathbb{A}$, on appelle conditionnement de $\mathbb{A}$ relativement à la norme $\|\cdot\|$ le nombre :

$$\text{cond}(\mathbb{A}) = \|\mathbb{A}\| \times \|\mathbb{A}^{-1}\| \quad (3.28)$$

La valeur du conditionnement d'une matrice dépendant en général de la norme choisie, on a coutume de signaler celle-ci en ajoutant un indice dans la notation, par exemple :

$$\text{cond}_{\infty}(\mathbb{A}) = \|\mathbb{A}\|_{\infty} \times \|\mathbb{A}^{-1}\|_{\infty} \quad (3.29)$$

On note que l'on a toujours $\text{cond}(\mathbb{A}) \geq 1$ puisque $1 = \|I_n\| = \|\mathbb{A} \cdot \mathbb{A}^{-1}\| \leq \|\mathbb{A}\| \times \|\mathbb{A}^{-1}\|$. L'égalité est obtenue pour une matrice orthonormale [72].

Dans quel sens le conditionnement influe-t-il sur le calcul? L'équation (3.27) nous montre que l'erreur sur le résultat d'une opération matricielle est bornée par un terme proportionnel au nombre de conditionnement. Dans le cas de l'inversion, si le conditionnement est « mauvais » ($\text{cond}(\mathbb{A}) \gg 1$), le résultat de l'inversion matricielle peut être entaché d'une erreur importante conduisant à une mauvaise direction de recherche...

Mais comment influencer sur le conditionnement d'une matrice? Étudions l'influence des termes diagonaux d'une matrice sur son conditionnement.

Soit la matrice :

$$\mathbb{A} = \begin{bmatrix} \nu & 1 & 1 \\ 1 & \nu & 1 \\ 1 & 1 & \nu \end{bmatrix} \quad (3.30)$$

Associé au nombre de conditionnement :

$$\text{cond}(\mathbb{A}) = \frac{(\nu + 2) \times (\nu + 3)}{|\nu^2 + \nu - 2|} \quad (3.31)$$

3 cas limites sont possibles pour ces termes diagonales : $\nu = (0, 1, \infty)$.

$$\begin{aligned} \lim_{\nu \rightarrow 0} \text{cond}(\mathbb{A}) &= 3 \\ \lim_{\nu \rightarrow 1} \text{cond}(\mathbb{A}) &= \infty \\ \lim_{\nu \rightarrow \infty} \text{cond}(\mathbb{A}) &= 1 \end{aligned} \quad (3.32)$$

On voit que si les termes diagonaux ont des valeurs très proches de celles des autres termes de la matrice, celle-ci sera mal conditionnée. Au contraire si les termes diagonaux sont d'un ordre de grandeur bien supérieur, la matrice est bien conditionnée. Ainsi une méthode pour améliorer le conditionnement d'une matrice consiste à rendre sa diagonale principale dominante. En pratique cela se fait en additionnant la matrice avec une matrice diagonale. Mais ceci présente l'inconvénient de modifier la matrice $\mathbb{A}$ et donc le problème, ajoutant ainsi un biais dans la résolution.

Dans notre cas, la matrice que l'on doit inverser est la matrice Hessien, équations (3.6) et (3.20). Le conditionnement de cette matrice représente celui du problème inverse. En tenant compte de l'équation (3.20) il vient :

$$\text{cond}(\mathbb{A}) = \text{cond}(\nabla^2(F_{obj})) \approx \text{cond}(2^t \mathbb{S} \cdot \mathbb{S}) \quad (3.33)$$

Plusieurs méthodes ont été développées pour régulariser une matrice, c'est-à-dire pour améliorer son conditionnement, et ainsi pouvoir obtenir de bons résultats suite à son inversion. Les plus connues d'entre toutes, sont les méthodes de régularisation de Tikhonov [79] comme l'addition d'un terme sur la diagonale principale de la matrice.

Nous arrivons maintenant à la méthode que nous avons employée dans ce travail. Il s'agit d'une version modifiée de la méthode de Levenberg-Marquardt qui, elle même, combine les méthodes de la plus grande pente et de Gauss-Newton.

3.3.1.5 La méthode de Levenberg-Marquardt [80] [81]

Cette méthode est une combinaison des méthodes précédentes dans l'équation (3.34).

$$\mathbf{p}^{k+1} = \mathbf{p}^k + \omega^k \left[{}^t\mathbb{S}(\mathbf{p}^k) \cdot \mathbb{S}(\mathbf{p}^k) + \delta^k \times \mathbb{I} \right]^{-1} \cdot {}^t\mathbb{S} \cdot (\mathbf{y}_{exp} - \mathbf{y}) \quad (3.34)$$

Le coefficient δ^k appelé « facteur de mélange », ou facteur de régularisation, détermine le mélange entre la méthode des gradients et celle de Gauss-Newton, il est ajusté en cours d'identification :

- Si le nouveau point $\mathbf{p}^{k+1}$ ne diminue pas la fonction objectif ($F_{obj}(\mathbf{p}^{k+1}) > F_{obj}(\mathbf{p}^k)$), alors la valeur de δ^{k+1} est augmentée par rapport à δ^k , donnant plus de poids à la méthode de la plus grande pente. Ensuite le point $\mathbf{p}^{k+1}$ est redéterminé.
- Si au contraire on améliore la solution, δ^{k+1} est diminué par rapport à δ^k afin de donner plus de poids à la méthode de Gauss-Newton.

Lorsque la solution courante est loin de la solution, la méthode se comporte comme la méthode de la plus grande pente. Proche de la solution, il se comporte comme la méthode de Gauss-Newton. La direction de descente de Levenberg-Marquardt est la moyenne pondérée des directions de descente de ces deux méthodes (voir figure 3.7). C'est sur cette pondération qu'intervient la modification que l'on a appliquée, 3.3.1.6. Le pas à parcourir dans la direction de descente ω^k est soit pris comme étant une constante égale à 1, soit elle est déterminée à chaque itération par une recherche monodimensionnelle.

LM : Levenberg-Marquardt
PGP : plus grande pente
GN : Gauss-Newton

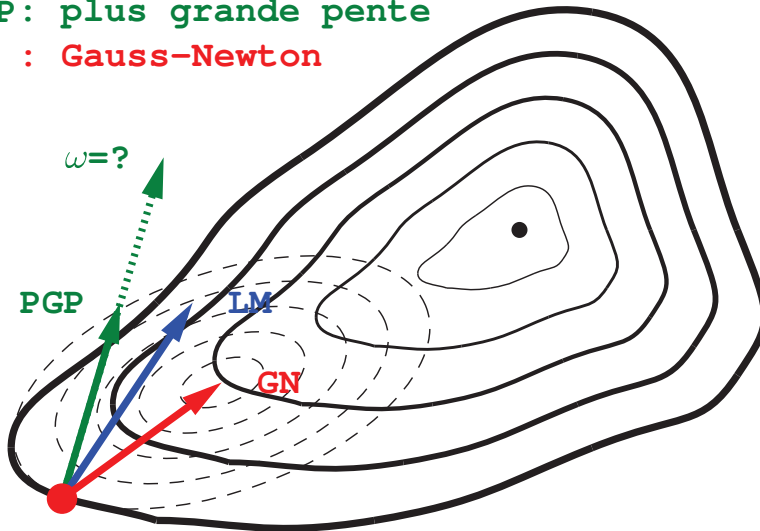


FIGURE 3.7 – Levenbenberg-Marquardt : pondération entre plus grande pente et Gauss-Newton

3.3.1.6 Méthode mise en oeuvre

Cette méthode utilise deux modifications du terme de mélange $\delta^k \times \mathbb{I}$ [61].

En première modification, l'ajout d'un multiple de la matrice identité permet d'améliorer le conditionnement de la matrice à inverser. Les paramètres n'étant pas tous du même ordre de grandeur, l'ajout de la matrice identité va affecter certains paramètres plus que d'autres, ce qui pourrait retarder la convergence. Remplaçons la matrice identité par la diagonale principale de la matrice $[{}^t\mathbb{S}(\mathbf{p}^k) \cdot \mathbb{S}(\mathbf{p}^k)]$, ainsi le biais imposé sera du même ordre de grandeur pour chaque paramètre, ou plutôt sur chaque direction portée par un paramètre.

En deuxième modification, nous avons fait évoluer le facteur de mélange proportionnellement à l'amélioration relative de la fonction objectif depuis l'initialisation.

$$\delta^k = \delta \frac{F_{obj}(\mathbf{p}^k)}{F_{obj}(\mathbf{p}^0)} \quad (3.35)$$

δ est choisi de sorte à améliorer suffisamment le conditionnement de la matrice ${}^t\mathbb{S}(\mathbf{p}^k) \cdot \mathbb{S}(\mathbf{p}^k)$ pour rendre le calcul de la matrice inverse possible.

$$\mathbf{p}^{k+1} = \mathbf{p}^k + \omega^k \left[{}^t\mathbb{S}(\mathbf{p}^k) \cdot \mathbb{S}(\mathbf{p}^k) + \delta \times \frac{F_{obj}(\mathbf{p}^k)}{F_{obj}(\mathbf{p}^0)} \times \mathbb{D} \right]^{-1} \cdot {}^t\mathbb{S} \cdot (\mathbf{y}_{exp} - \mathbf{y}) \quad (3.36)$$

La matrice $\mathbf{D}$ correspond aux termes diagonaux de ${}^t\mathbb{S}(\mathbf{p}^k) \cdot \mathbb{S}(\mathbf{p}^k)$.

3.3.2 Les méthodes stochastiques ou algorithmes évolutionnaires

Les méthodes exposées jusqu'ici sont liées à la mise en oeuvre des algorithmes de minimisation utilisant le gradient de la fonction objectif : plus grande pente, Newton, Levenberg-Marquardt... Ces dernières, bien maîtrisées et d'utilisations classiques, sont « aveugles » par nature : partant d'une valeur initiale des inconnues à identifier, elles fournissent un résultat : minimum global ou local (l'utilisateur ignorant généralement laquelle de ces deux possibilités est atteinte, sauf dans des cas particuliers où la fonction objectif est convexe).

Une autre voie, dont l'exploration a été rendue possible par le progrès en termes de puissance de calcul, consiste à utiliser des méthodes d'optimisation stochastiques. Ces méthodes, plutôt que d'essayer de déterminer le chemin à emprunter pour atteindre la solution, évaluent plusieurs possibilités ou candidats pour la solution et ensuite par des méthodes inspirées par la Nature ou la société, fait évoluer les candidats dans la zone de recherche. Les méthodes stochastiques manipulent une population de solutions potentielles et non-pas une solution unique. Ces méthodes ne n'ont recours qu'à des opérations mathématiques simples, notamment pas de calcul des dérivés ni d'inversion matricielle qui pourraient « bloquer » le calcul. Mais ces méthodes évolutionnaires évaluent un plus grand nombre de fois la fonction objectif, qui peut être coûteuse en temps de calcul.

Dans cette étude nous n'avons utilisé qu'un seul algorithme évolutionnaire, très utilisé dans le domaine de la thermique : l'algorithme génétique. Un inventaire de ses utilisations dans ce domaine est disponible dans l'article [82].

3.3.2.1 L'algorithme génétique

La caractéristique principale de cet algorithme stochastique est qu'il manipule une population de solutions potentielles, chacune représentant un point de l'espace de recherche. Elle est manipulée de la même manière que des populations vivantes évoluent génétiquement. Son avantage principal est sa robustesse. Contrairement aux méthodes de descente, l'algorithme génétique traitant ou explorant simultanément tout le domaine, fait la différence entre l'extremum global et les extrémums locaux. Il finit toujours par converger sur un résultat probant. Cet algorithme évolutionnaire repose sur deux catégories de mécanisme calquées sur la Nature, le croisement des gènes lors de la reproduction et la sélection naturelle.

Pour trouver la solution, on manipule une population d'individus, chaque individu étant une solution possible. Un individu est représenté par un chromosome (génotype) porteur d'information dans ces gènes. Ce chromosome correspond au vecteur paramètre $\mathbf{p}$ (phénotype), et chaque gène correspond à une composante de $\mathbf{p}$, à un des paramètres à identifier. Il existe plusieurs méthodes pour coder ces gènes. Nous utilisons le codage réel qui affecte à chaque gène la valeur réelle du paramètre.

L'évolution de cette population est le passage d'une génération de chromosomes à la suivante. Ce passage se fait suivant plusieurs étapes à travers lesquels l'individu de l'itération précédente avec la meilleure adaptation, le BEST, n'est pas modifié au cours de ces étapes afin d'entraîner la convergence, on appelle cela l'élitisme. Ces étapes sont détaillées ci-après.

Initialisation

Dans cette première étape on se donne une population de départ de N_{ind} individus formant la première génération. C'est ici que l'on peut introduire des estimations *a priori* des paramètres que l'on cherche à identifier ainsi que toutes informations du même type. La population est ensuite distribuée aléatoirement autour de ces estimations afin « d'aider » la convergence.

Evaluation

Ici on note les individus selon la fonction objectif, on calcule leur adaptation (ou *fitness*). Cette étape correspond, dans une population vivante, à l'évaluation des chances d'un individu de pouvoir être choisi pour la reproduction. En général on attribue une note d'autant plus élevée que son adaptation est bonne, il s'agit d'une méthode qui maximise le critère. Ici on veut minimiser l'écart quadratique, donc dans notre cas, on rajoute un signe moins à l'équation (3.4).

$$F_{obj}(\mathbf{p}) = -\|\mathbf{y} - \mathbf{y}_{exp}\|^2 \quad (3.37)$$

Sélection

C'est l'étape de compétition, les individus sont comparés selon leur adaptation afin d'être sélectionnés pour l'étape de croisement. Qui va générer les nouvelles solutions possibles au problème d'inversion. La compétition peut avoir différentes règles :

- La sélection élitiste ne permet qu'aux meilleurs $\frac{N_{ind}}{2}$ individus de se reproduire, mais dans ce cas on réduit trop rapidement la diversité génétique de la population ce qui risque de faire converger trop rapidement, sur un minimum local.
- Une autre option est d'effectuer une sélection probabiliste des individus reproducteurs selon une probabilité proportionnelle à leur adaptation. Laissant ainsi une possibilité même faible aux mauvais candidats de survivre. Un exemple est la méthode RWS (Roulette Wheel Selection) qui a le même principe que la roulette d'un casino où le tirage est avec remise. On découpe la roulette en N_{ind} portions de largeur proportionnelle à l'adaptation de l'individu lui correspondant. Ensuite on fait $\frac{N_{ind}}{2}$ tirages pour obtenir les parents. Avec cette méthode il est possible que le meilleur individu, le meilleur jeu de paramètres, ne soit pas sélectionné, et que le plus mauvais individu soit sélectionné plusieurs fois. Ce genre de situations est peut être néfaste à l'identification. La méthode SUS (Stochastic Universal Sampling) reprend la même méthode mais avec des tirages sans remise ce qui évite les situations extrêmes évoquées précédemment.
- Nous avons choisi encore une autre méthode, celle de la sélection par tournoi. On choisit au hasard deux individus et on conserve celui avec la meilleure adaptation. Les moins mauvais ont ainsi une bonne chance de survivre quelques itérations en se retrouvant confrontés aux plus mauvais. Cette approche est plus fidèle à l'esprit de l'algorithme, véritable analogie de l'évolution génétique d'une population.

Une fois que le processus de sélection est terminé, les individus non-sélectionnés sont « détruits ». Ils seront remplacés par les « enfants » issus de l'étape de croisement. La taille de la population diminue donc de moitié à cette étape, avant de retrouver toute sa taille à la fin de l'étape suivante.

Croisement

Comme pour la reproduction sexuée, chaque gène ou paramètre à identifier des nouveaux individus ou descendants sera une combinaison des gènes correspondants de ses parents. Plus précisément dans notre cas nous avons choisi de prendre deux individus aléatoirement et de combiner leurs gènes linéairement mais en donnant un caractère aléatoire à la prédominance de chaque gène.

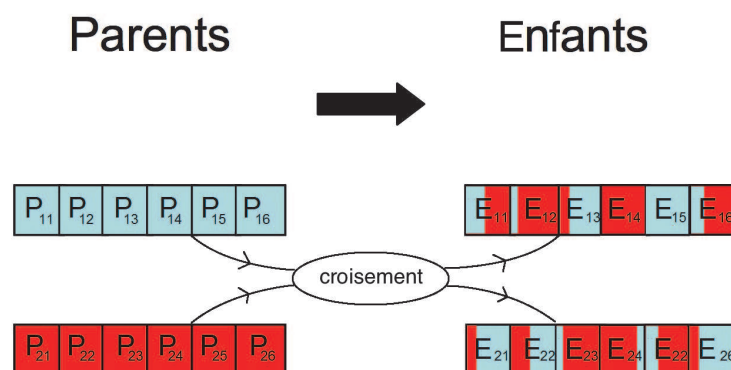


FIGURE 3.8 – Exemple d'un croisement BLX

Nous avons utilisé un croisement de type BLX- α volumique [77] [83] [84]. Ainsi chaque couple de gène parent $P_{i,j}$ et $P_{k,j}$ va donner deux gènes enfants $E_{n,j}$ et $E_{m,j}$ tels que :

$$E_{n,j} = P_{i,j} - \frac{1}{2}(P_{k,j} - P_{i,j}) + (1 + \frac{\alpha_{BLX}}{2})(P_{k,j} - P_{i,j})\text{rand}([0 : 1]) \quad (3.38)$$

$$E_{m,j} = P_{i,j} - \frac{1}{2}(P_{k,j} - P_{i,j}) + (1 + \frac{\alpha_{BLX}}{2})(P_{k,j} - P_{i,j})\text{rand}([0 : 1]) \quad (3.39)$$

$\text{rand}([0 : 1])$ est un nombre aléatoire compris entre 0 et 1.

La variance de la population, Var , varie à cette étape de croisement selon la valeur de α_{BLX} . En effet la variance de la population enfant $\text{Var}_{\text{enfant}}$ peut être exprimée en fonction de la variance de la population parent $\text{Var}_{\text{parent}}$ [77] :

$$\text{Var}_{\text{enfant}} = ((1 + 2 \times \alpha)^2 + 3) \frac{\text{Var}_{\text{parent}}}{6} \quad (3.40)$$

Si α_{BLX} est inférieur à $\frac{\sqrt{3}-1}{2} \approx 0.366$, la variance de la population diminue d'une génération sur l'autre et donc la convergence n'est plus assurée uniquement par l'étape de sélection mais aussi par la réduction du domaine de recherche à chaque étape de croisement. On prendra la valeur $\alpha_{BLX} = \frac{1}{2}$ pour s'assurer que la convergence, qui est caractérisée par la réduction de la variance de la population, n'est pas liée à l'étape de croisement puisque celle-ci l'augmente comme on peut le voir sur la figure 3.9.

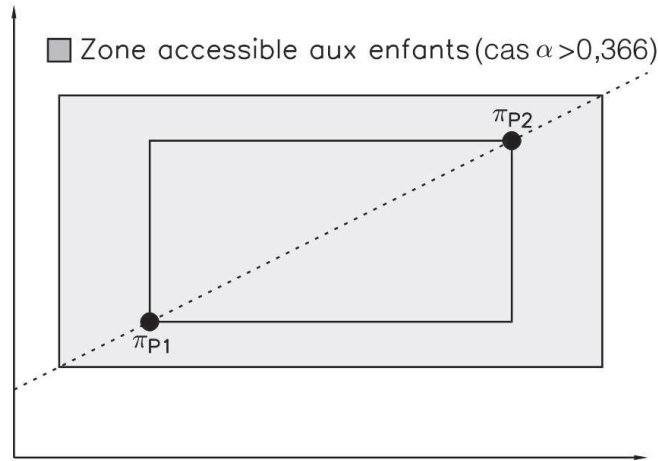


FIGURE 3.9 – Croisement « BLX- α volumique » [77]

Mutation

De façon aléatoire la valeur d'un gène, au sein d'un individu ou chromosome, peut être modifiée. La mutation nous sert à éviter une convergence prématurée de l'algorithme, mais surtout à pouvoir se sortir de minimums locaux. On définit ici un taux ou une probabilité de mutation d'un chromosome qui est généralement compris entre 0,001 et 0,01 ; nous prendrons cette dernière valeur. Il est nécessaire de choisir pour ce taux une valeur relativement

faible de manière à ne pas tomber dans une recherche aléatoire ce qui réduirait à néant l'intérêt de l'algorithme.

– En résumé :

- c'est une méthode stochastique
- elle est robuste car si certains candidats provoquent la divergence du modèle direct, l'algorithme continue de fonctionner avec le reste de la population
- elle est facile à implémenter
- la méthode peut être lente du fait du nombre élevé de candidats à évaluer, son temps d'itération dépendant fortement du temps de calcul de F_{obj}

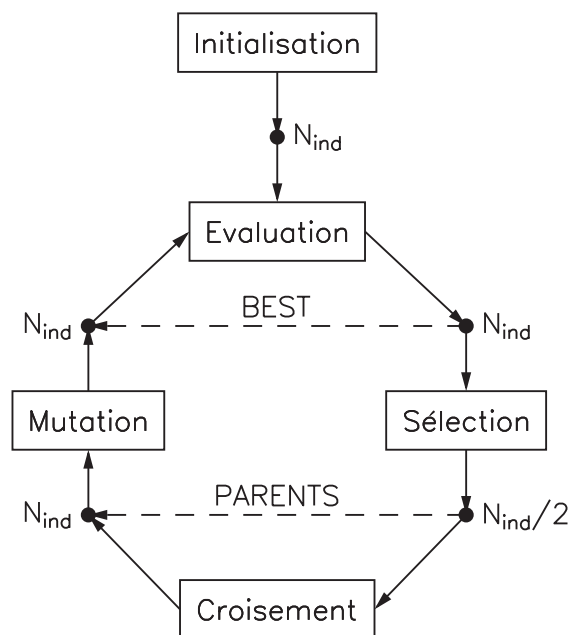


FIGURE 3.10 – Cycle générationnel

3.3.3 A mi-chemin : le Simplexe

Nous avons aussi testé une troisième méthode pour déterminer l'algorithme le plus performant en termes de vitesse de convergence et de simplicité d'utilisation. Nous dirons que cet algorithme est à mi-chemin car bien qu'étant une méthode de descente, il ne détermine pas sa direction de descente grâce à un gradient, mais en évaluant quelques candidats, il définit un hyperplan qui donne une direction de descente. C'est-à-dire qu'il assimile le gradient local au gradient de l'hyperplan passant par les sommets du simplexe. Un simplexe est une figure géométrique à $n_p + 1$ sommets qui évolue dans un espace à n_p dimensions correspondant aux n_p inconnues. Par exemple, pour un problème à deux inconnues ou paramètres, le simplexe sera un triangle qui évoluera en reflétant son plus mauvais sommet selon la fonction objectif, par rapport à l'arête qui lui est opposée [85]. Nedler-Mead [86] ont proposé une

version plus élaborée du simplexe qui, par le biais de plusieurs opérations, déplace le plus mauvais point pour l'améliorer (voir figure 3.11).

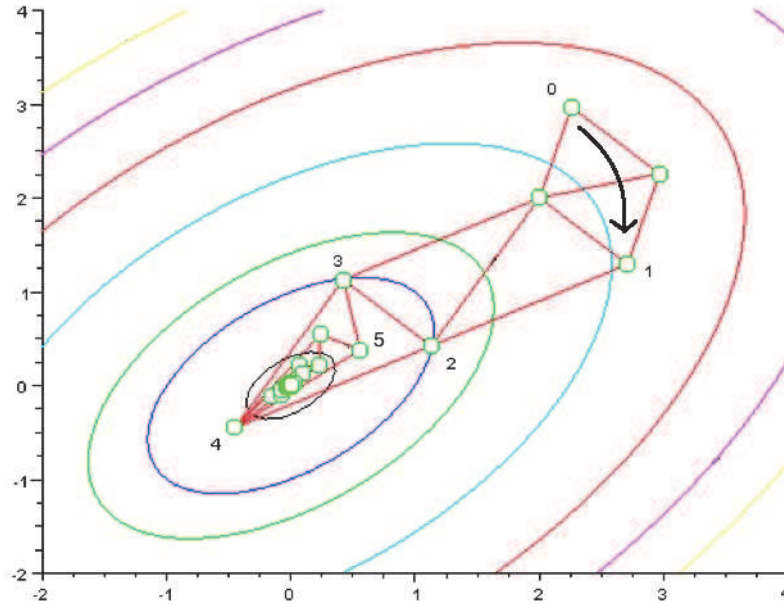


FIGURE 3.11 – Historique de l'évolution du simplexe de Nedler-Mead pour un problème à deux paramètres [86]

3.3.3.1 Algorithme du simplexe de Nedler-Mead

Cet algorithme utilise différentes transformations géométriques afin de déformer le simplexe, afin de pouvoir toujours progresser dans sa recherche de la solution optimum et ne pas rester bloqué en cours de processus. Ces opérations sont déterminées par quatre coefficients. Dans l'ordre celui de réflexion, puis celui d'expansion, de contraction et enfin de réduction. Nous avons pris pour ces coefficients les valeurs standards [86].

$$\rho_{NM} = 1 \quad \chi_{NM} = 2 \quad \gamma_{NM} = \frac{1}{2} \quad \sigma_{NM} = \frac{1}{2} \quad (3.41)$$

D'abord il faut initialiser le simplexe soit $n_p + 1$ points pour n_p paramètres. Pour le premier point nous avons pris des valeurs que nous estimions « proches » de la solution. Pour les n_p suivant nous avons augmenté la valeur de un des n_p paramètres de 1% et si ce paramètre était nul, nous l'avons augmenté de 0,1. Ainsi nous avons une distribution selon chacune des directions du domaine de recherche. Ensuite il faut tous les évaluer avec la fonction objectif F_{obj} de l'équation (3.4), le point qui obtient le meilleur critère est nommé p_{BEST} . Enfin, il faut classer les individus. De ce classement va dépendre les opérations géométriques.

Ce qui suit correspond aux différentes opérations et tests à effectuer pour chaque itération :

Calculer le barycentre des sommets du simplexe mis à part le plus mauvais point $\mathbf{p}_{n_p+1}$:

$$\bar{\mathbf{p}} = \frac{1}{n_p} \sum_{i=1}^{n_p} \mathbf{p}_i \quad (3.42)$$

Calculer le point obtenu par « réflexion » :

$$\mathbf{p}_r = (1 + \rho_{NM})\bar{\mathbf{p}} - \rho_{NM} \cdot \mathbf{p}_{n_p+1} \quad (3.43)$$

Si $F_{obj}(\mathbf{p}_r) < F_{obj}(\mathbf{p}_{BEST})$ alors calculer le point obtenu par « expansion » :

$$\mathbf{p}_e = (1 + \rho_{NM} \cdot \chi_{NM})\bar{\mathbf{p}} - \rho_{NM} \cdot \chi_{NM} \cdot \mathbf{p}_{n_p+1} \quad (3.44)$$

Si de plus $F_{obj}(\mathbf{p}_e) < F_{obj}(\mathbf{p}_r)$ alors il faut effectuer l'opération d'expansion du simplexe

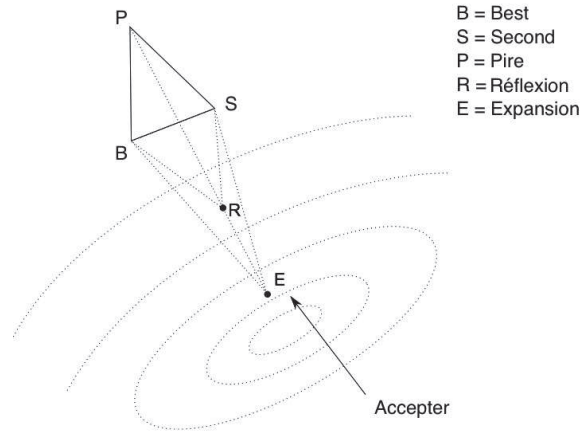


FIGURE 3.12 – Illustration de l'opération d'expansion [86]

sinon effectuer l'opération de « réflexion ».

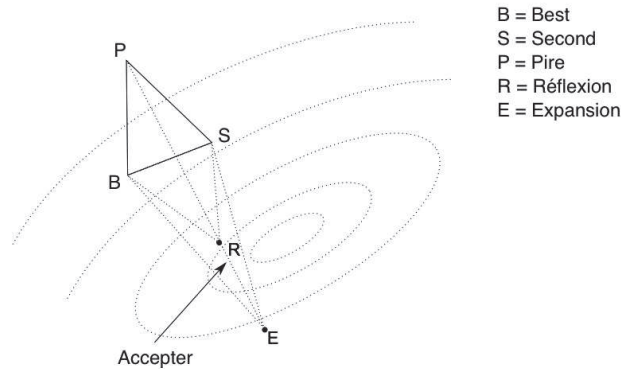


FIGURE 3.13 – Illustration de l'opération de réflexion [86]

Si le point que l'on obtiendrait par réflexion n'est pas meilleur que le BEST, mais quand même meilleur que le deuxième plus mauvais point $F_{obj}(\mathbf{p}_{n_p})$, soit $F_{obj}(\mathbf{p}_{BEST}) < F_{obj}(\mathbf{p}_r) < F_{obj}(\mathbf{p}_{n_p})$ alors il faut calculer le point que l'on obtiendrait par contraction externe :

$$\mathbf{p}_{ce} = (1 + \rho_{NM} \cdot \gamma_{NM}) \bar{\mathbf{p}} - \rho_{NM} \cdot \gamma_{NM} \cdot \mathbf{p}_{n_p+1} \quad (3.45)$$

Dans ce cas si ce nouveau point est meilleur que celui de la réflexion, soit $F_{obj}(\mathbf{p}_{ce}) < F_{obj}(\mathbf{p}_r)$, alors on effectue la contraction externe :

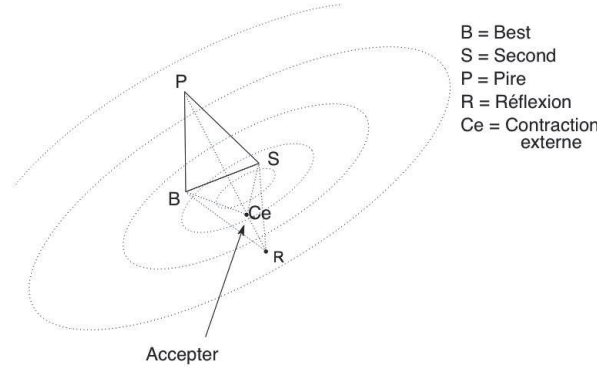


FIGURE 3.14 – Illustration de l'opération de contraction externe [86]

Sinon on effectue une réduction de la taille du simplexe centrée sur le BEST. Donc pour tous les sommets i du simplexe sauf le BEST nous faisons :

$$\mathbf{p}_i = \mathbf{p}_{BEST} + \sigma_{NM}(\mathbf{p}_i - \mathbf{p}_{BEST}) \quad (3.46)$$

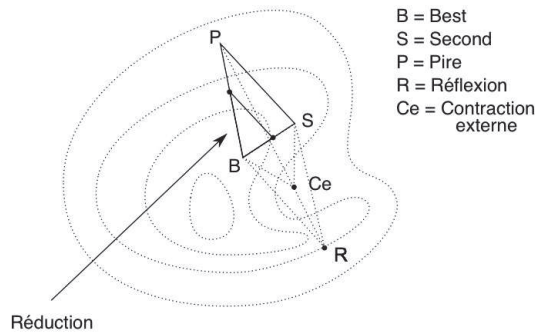


FIGURE 3.15 – Illustration de l'opération de réduction fautive de contraction externe [86]

Si le critère du point qui est obtenu par l'opération de réflexion n'est pas meilleur que le critère du second plus mauvais point $\mathbf{p}_{n_p}$, alors on calcule le point que l'on obtiendrait par une contraction interne :

$$\mathbf{p}_{ci} = (1 + \rho_{NM} \cdot \gamma_{NM}) \bar{\mathbf{p}} - \gamma_{NM} \cdot \mathbf{p}_{n_p+1} \quad (3.47)$$

Puis dans le cas où ce point est meilleur que le plus mauvais point, $F_{obj}(\mathbf{p}_{ci}) < F_{obj}(\mathbf{p}_{n_p+1})$, on effectue l'opération de contraction interne.

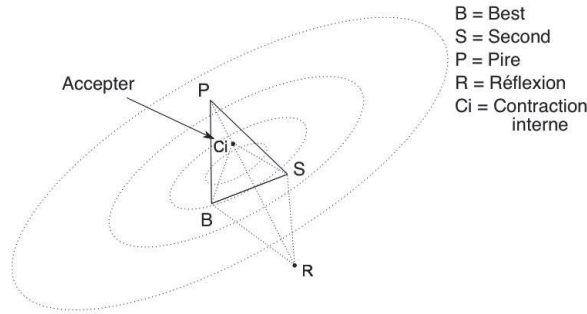


FIGURE 3.16 – Illustration de l'opération de contraction interne [86]

Si on ne tombe dans aucune des catégories précédentes, on effectue une réduction de la taille du simplexe centrée sur le BEST. Il s'agit de la même opération de réduction que la précédente, figure 3.15, seule la configuration des autres points change.

Donc pour tous les sommets i du simplexe sauf le BEST l'opération suivante est réalisée :

$$\mathbf{p}_i = \mathbf{p}_{BEST} + \sigma_{NM}(\mathbf{p}_i - \mathbf{p}_{BEST}) \quad (3.48)$$

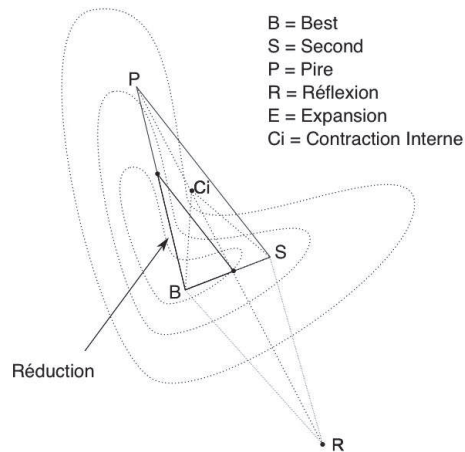


FIGURE 3.17 – Illustration de l'opération de réduction faute de contraction interne [86]

Puis l'algorithme recommence.

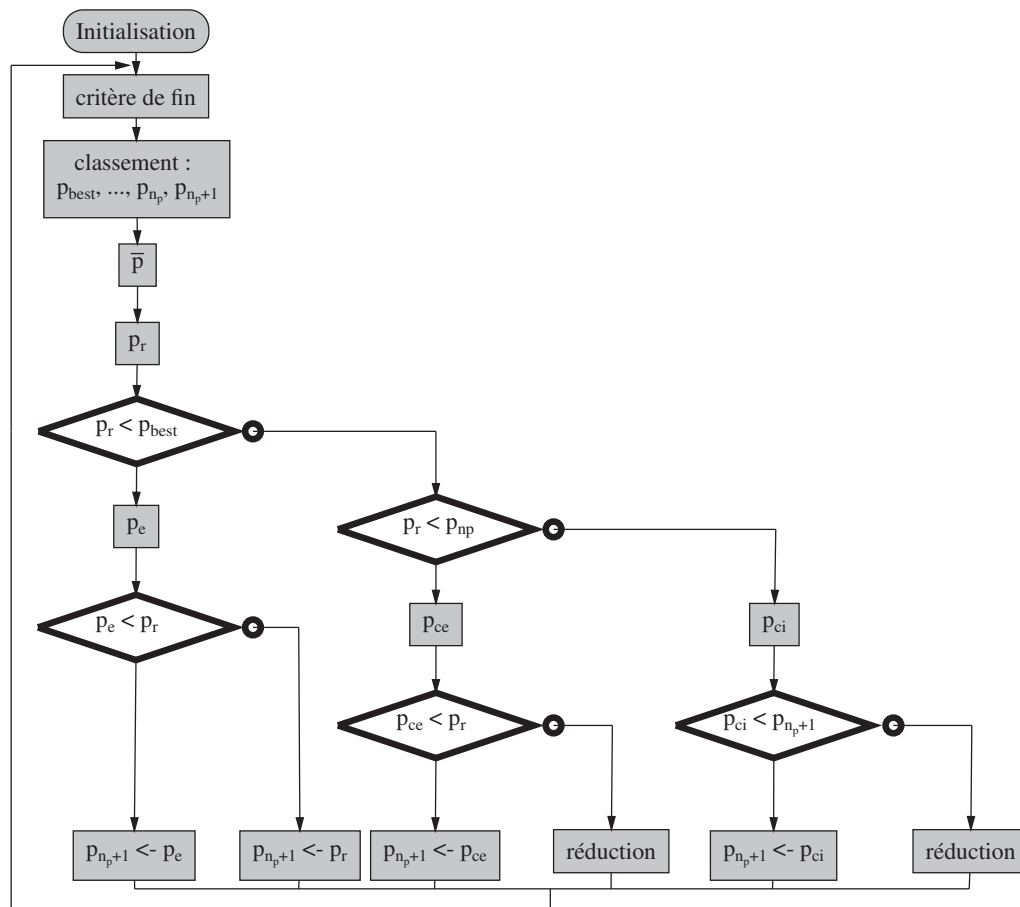


FIGURE 3.18 – Algorithme du simplexe de Nedler-Mead

3.3.3.2 Interprétation géométrique des différentes opérations mathématiques

- Le simplexe effectue une expansion quand il est encore loin de la solution. Il déplace le plus mauvais point selon une direction de descente.
- La réflexion correspond à une situation où le simplexe aborde une vallée du domaine de recherche. Dans une telle situation l'opération d'expansion ne permet pas d'améliorer la fonction objectif.
- Proche du minimum, les contractions internes et externes permettent de réduire la taille du simplexe tout en le déplaçant dans la direction de descente. Si trop d'opérations de ce type s'enchaînent à la suite loin de la solution, on obtient un simplexe dégénéré qui pourrait converger prématurément. Dans ce cas, il faut réinitialiser le simplexe sur sa solution courante.
- Les étapes de réductions diminuent la taille du simplexe en rapprochant les sommets du *BEST*. Dans la pratique ces étapes sont rares.

3.3.4 Critère d'arrêt de l'algorithme

Pour chacun de ces algorithmes, la convergence est repérée quand un certain critère d'arrêt est atteint. Il existe plusieurs critères possibles parmi lesquels :

- On fixe un grand nombre d'itérations et on espère que la solution soit atteinte avant.
- Un critère de non-évolution de la fonction objectif peut être utilisé, $\|\vec{\nabla}(Fobj)(\mathbf{p}^k)\| < \varepsilon$. Mais on peut alors s'arrêter sur un point d'inflexion ou un point de selle.
- La non-évolution peut aussi être constatée sur les paramètres à identifier, $\|\Delta \mathbf{p}^{k+1} - \mathbf{p}^k\| \approx 0$, mais dans le cas des algorithmes évolutionnaires il peut se produire de nombreuses itérations sans que le meilleur candidat n'évolue. Il faut donc attendre plusieurs itérations après que cette condition soit atteinte afin de confirmer la convergence.
- Pour les algorithmes génétiques on peut fixer le critère par rapport à l'écart entre candidats, la convergence est atteinte quand les candidats se sont regroupés.

$$\max_{i=1 \dots n_p+1} \|\mathbf{p}_i - \mathbf{p}_{BEST}\| < \varepsilon \text{ [86]}.$$

Dans notre cas nous avons choisi de considérer le critère de non-évolution des paramètres après 30 itérations pour tous nos algorithmes.

3.4 Conclusion : adaptation des méthodes inverses à la DSC

Nous voulons utiliser une méthode inverse pour identifier des propriétés d'un corps à partir d'un thermogramme obtenu à la suite d'une expérience de DSC. Par rapport aux notations que nous avons définies dans ce chapitre, $\mathbf{y}_{exp}$ sera un vecteur colonne constitué des ordonnées du thermogramme expérimental, soit des valeurs du flux thermique en watt. L'abscisse est constitué des temps correspondant aux valeurs du flux pris en ordonnées : $\mathbf{y}_{exp} = \boldsymbol{\phi}_{exp} = {}^t(\phi_{exp}(t_1), \dots, \phi_{exp}(t_{n_m}))$.

Ne disposant d'informations suggérant que certains points ou ensembles de points que l'on pourrait positionner *a priori* par rapport aux abscisses sont plus importants ou moins importants que d'autres, la matrice $\mathbb{W}$ sera égale à la matrice identité, $\mathbb{I}$. $\mathbb{M}$ sera le modèle direct du transfert thermique qui incorpore une des équations d'état du chapitre 2.1.

$\mathbf{y}$ sera un vecteur colonne constitué des valeurs du flux du thermogramme numérique simulé par le modèle direct avec : $\mathbf{y} = \boldsymbol{\phi}_{num} = {}^t(\phi_{num}(t_1), \dots, \phi_{num}(t_{n_m}))$. Les temps ${}^t(t_1, \dots, t_{n_m})$ du thermogramme numérique seront les mêmes que ceux des observations expérimentales.

Quant au vecteur $\mathbf{p}$, il sera constitué par les propriétés thermodynamiques que l'on veut identifier des modèles présentés au chapitre 2, et les paramètres de transfert que l'on identifie ou adapte.

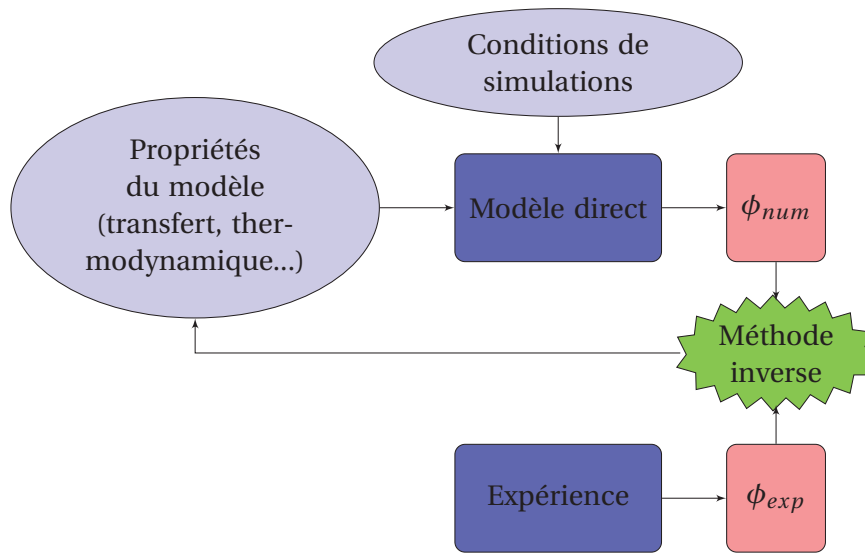


FIGURE 3.19 – Schéma de l'inversion appliqué à la D.S.C.

On en arrive à la fonction objectif. Les données expérimentales et numériques appartiennent toutes deux à un ensemble discret. Il faut donc discretiser la relation (3.4) :

$$F_{obj}(\mathbf{p}) = \sum_{i=1}^{n_m} (\phi_{num}(t_i, \mathbf{p}) - \phi_{exp}(t_i))^2 \quad (3.49)$$

Chapitre 4

Mise en œuvre des méthodes inverses pour la calorimétrie

Dans le chapitre 3 nous avons présenté ce qu'est une méthode inverse et nous avons détaillé les méthodes que nous avons envisagées. Ce chapitre 4 contient les résultats d'identification de plusieurs corps réalisés avec ces méthodes, c'est-à-dire comment nous avons identifié les paramètres des modèles directs présentés dans le chapitre 2, paramètres qui permettent de reconstruire les enthalpies des échantillons.

Si certains paramètres sont corrélés, par exemple ceux qui ne se trouvent que sous la forme d'un produit entre eux, alors il existe une infinité de solutions possibles pour ces deux paramètres. On peut savoir si certains des paramètres sont corrélés en étudiant les coefficients de sensibilités issus du modèle direct, qui sont définis dans le chapitre sur les méthodes inverses, équation (3.17). Après avoir déterminé un jeu de paramètres identifiables, nous allons tester notre méthode d'inversion sur différents thermogrammes en appliquant la méthode inverse sur un thermogramme numérique dont on connaît les paramètres, et vérifier que les paramètres sont retrouvés par cette méthode. Enfin, à partir de ces coefficients de sensibilités, on va définir un intervalle de confiance pour les paramètres que l'on identifie, avant d'appliquer notre méthode à des cas expérimentaux.

4.1 Etude des coefficients de sensibilité

Ces coefficients correspondent à l'étude paramétrique menée sur le modèle direct pour nos différentes équations d'état, ils reflètent donc les mêmes résultats que la section 2.3. Cette représentation est orientée pour l'inversion, on y constate mieux l'importance de chaque paramètre c'est-à-dire la facilité de les identifier par une méthode inverse. Le cas du binaire est particulièrement intéressant car il dépend d'un nombre relativement élevé de paramètres. On peut donc s'attendre à ce qu'il y ait une certaine corrélation entre eux, ce qui rendrait impossible l'identification simultanée de ces paramètres. L'étude des coefficients de sensibilité permet aussi de distinguer quels sont les paramètres qui seront facilement identifiés et ceux qui le seront difficilement. Comme nous le verrons dans la section 4.3.3, ils permettent de définir les intervalles de confiance pour les paramètres identifiés.

Nous présenterons dans cette section les coefficients de sensibilités pour deux fractions

massiques en soluté 0,57% et 10% de la solution de $H_2O - NH_4Cl$, afin d'étudier les cas de pics de thermogramme séparés ou fusionnés. Le cas du corps pur n'est pas présenté car c'est un cas particulier du modèle binaire. Nous ne présenterons les valeurs des sensibilités que pour une vitesse β de $10 K \cdot min^{-1}$, afin d'amplifier l'ordre de grandeur des sensibilités et aussi parce que nous voulons une méthode opérationnelle pour des vitesses de réchauffements relativement élevées.

Ces coefficients de sensibilité sont définis dans la sous-partie traitant de la méthode de Gauss, équation (3.17). Pour les calculer, nous avons considéré un jeu de paramètres de référence qui correspondent au thermogramme de référence. Ensuite, pour chaque coefficient, nous avons modifié la valeur du paramètre i de Δp_i , calculé un nouveau thermogramme, puis calculé la dérivée discrète de la sortie du modèle par rapport au paramètre i , équation (4.1). Cette opération est réalisée pour chacun des paramètres que l'on voudrait identifier.

$$\mathbb{S}_{p_i} = \left(\frac{\phi(\mathbf{p} + \Delta \mathbf{p}_i) - \phi(\mathbf{p})}{\Delta p_i} \right)_{p_j \neq i} \quad (4.1)$$

avec $\mathbb{S}_{p_i}$ de la colonne de la matrice $\mathbb{S}$ de l'équation (3.17) qui correspond au paramètre p_i .

Les figures 4.9 à 4.7 correspondent à des décalages d'un des paramètres de $\pm 10\%$ ou ± 1 et $0,5 K$ afin de rendre visibles les déformations des thermogrammes. Dans la suite de l'étude, les coefficients de sensibilités seront calculés avec un décalage de 1% ou $0,1 K$, les formes et les rapports d'ordres de grandeurs de ces coefficients restant inchangés. Pour que l'on puisse comparer les formes et les ordres de grandeurs de ces paramètres entre eux, on multiplie les coefficients par les paramètres qui leurs correspondent. On obtient des coefficients de sensibilité réduits qui ont tous la dimension d'un flux, équation (4.2).

$$p \mathbb{S}_{p_i} = p_i \times \mathbb{S}_{p_i} \quad (4.2)$$

4.1.1 Coefficient de sensibilité réduit relatif à T_E

Le modèle direct présente une forte sensibilité à la température eutectique car T_E positionne le début du premier pic de fusion du thermogramme. Ce pic représente une transition à température constante, et a donc une pente initiale dont la valeur est fixée par la résistance thermique entourant l'échantillon (voir (1.16)). Une modification de T_E va décaler cette pente à l'origine du premier pic. Avec ce décalage, les thermogrammes ne vont se recouvrir qu'en partie. Les courbes de flux seront distinctes sur les zones d'influences du paramètre, figure 4.1. Avec l'augmentation de la fraction massique de soluté, une influence sur le deuxième pic apparaît, augmentant la sensibilité du modèle à T_E , car du fait de la diminution de T_M (voir paragraphe 2.3.2.6) les deux pics du thermogramme vont se rapprocher.

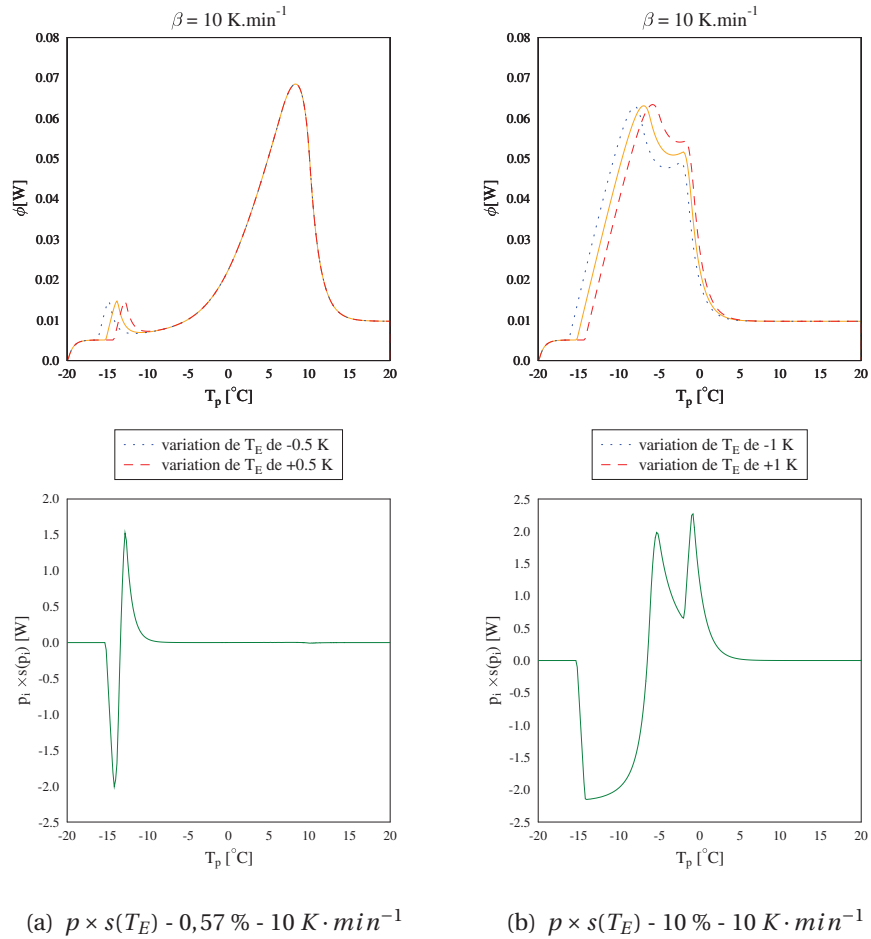
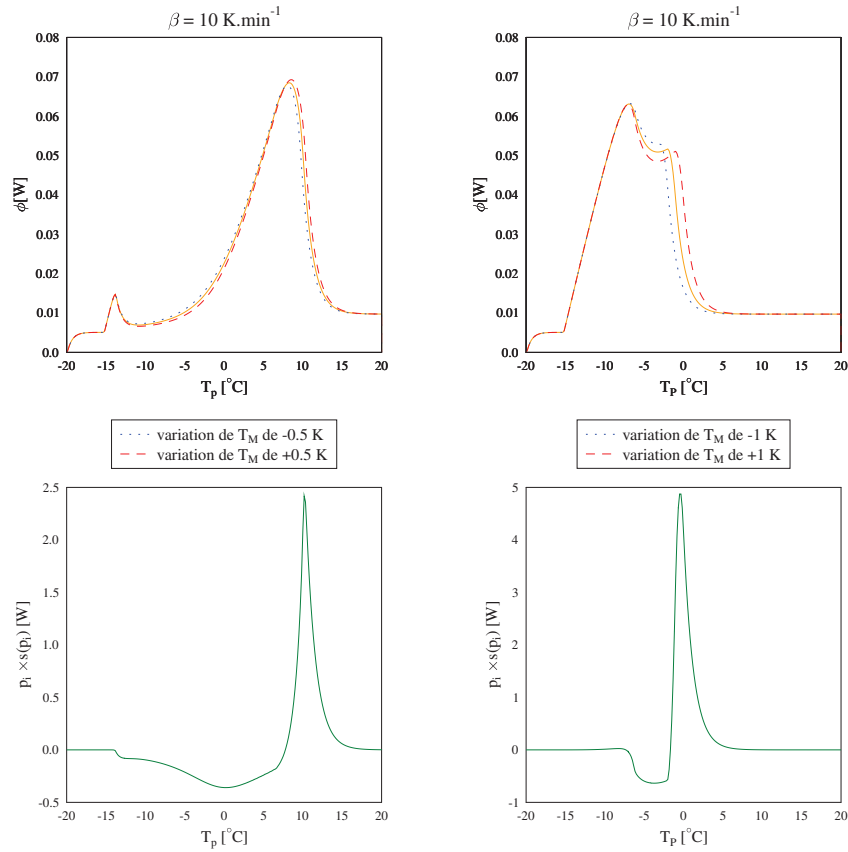


FIGURE 4.1 – Coefficient de sensibilité réduit relatif à T_E

4.1.2 Coefficient de sensibilité réduit relatif à T_M

Une variation de T_M , toutes choses égales par ailleurs, va décaler la température de fin de fusion et donc décale la fin du deuxième pic dans le même sens. D'après l'équation (2.51), une modification de T_M , toutes choses égales par ailleurs va modifier la fraction massique en soluté x_0 , donc la forme de $h(T)$ et donc tout le thermogramme de fusion.



(a) $p \times s(T_M) - 0,57 \% - 10 K \cdot min^{-1}$

(b) $p \times s(T_M) - 10 \% - 10 K \cdot min^{-1}$

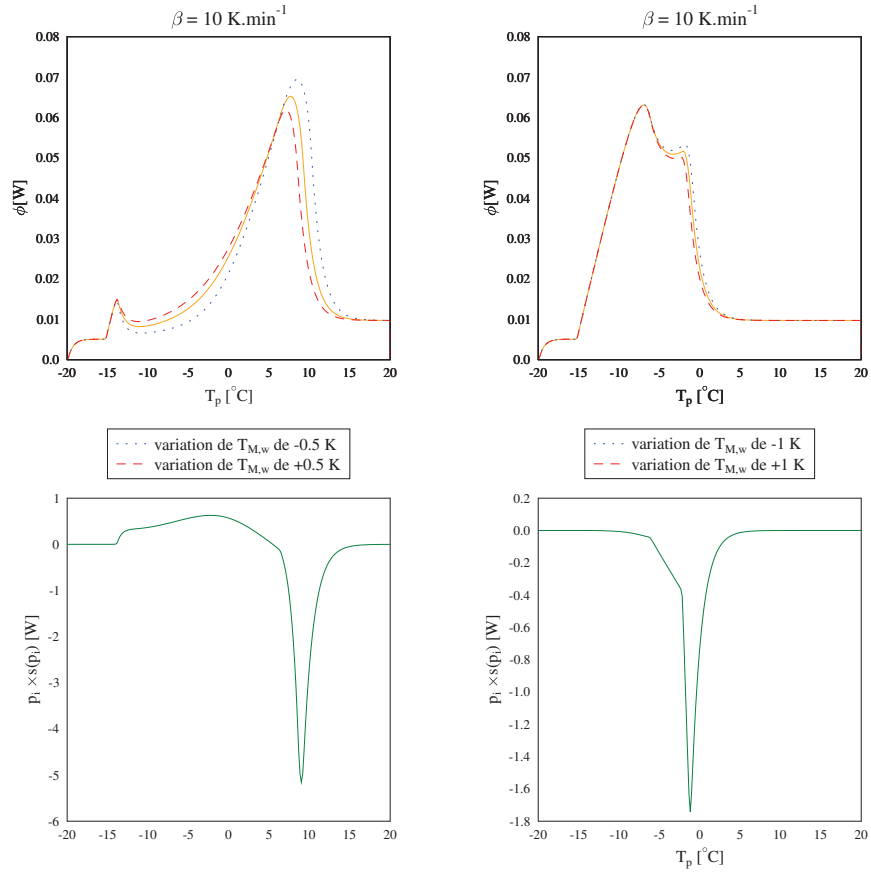
FIGURE 4.2 – Coefficient de sensibilité réduit relatif à T_M

4.1.3 Coefficient de sensibilité réduit relatif à $T_{M,w}$

La sensibilité à ce paramètre se porte principalement sur la fin du pic du solvant, figure 4.3. Dans les équations du modèle thermodynamique du binaire (2.53) et (2.54), donnant l'enthalpie à travers le changement de phase, les températures T_M et $T_{M,w}$ interviennent fréquemment sous la forme d'une différence. L'enthalpie va donc varier globalement de manière contraire selon ces deux paramètres. On retrouve cette différence de signe dans les

sensibilités. Néanmoins, ces deux températures interviennent aussi indépendamment l'une de l'autre dans ces équations. C'est pourquoi les coefficients de sensibilité ne sont pas les mêmes, permettant leur identification simultanée. Pour une faible fraction massique de soluté, $T_M \approx T_{M,w}$ et $\widetilde{L}_E \approx 0 \text{ J} \cdot \text{kg}^{-1}$, si on inclue cette relation dans les équations donnant l'enthalpie pour un binaire (2.52) (2.53) (2.54), la forme de $h(T)$ tend vers celle d'un corps pur.

Effectivement, quand la fraction massique de soluté x_0 , tend vers 0, le comportement de la solution tend vers le comportement d'un corps pur. Ceci explique pourquoi ces deux paramètres ont des coefficients de sensibilités très similaires pour de très faibles fractions massiques de soluté. Pour des concentrations supérieures, les formes des deux sensibilités se distinguent de plus en plus, notamment sur l'intervalle de temps précédant la fin de fusion où le paramètre $T_{M,w}$ influe de moins en moins. Cette diminution locale de sensibilité peut s'expliquer par l'importance croissante du paramètre T_E .



(a) $p \times s(T_{M,w})$ - 0,57 % - $10 \text{ K} \cdot \text{min}^{-1}$

(b) $p \times s(T_{M,w})$ - 10 % - $10 \text{ K} \cdot \text{min}^{-1}$

FIGURE 4.3 – Coefficient de sensibilité réduit relatif à $T_{M,w}$

4.1.4 Coefficient de sensibilité réduit relatif à $\widetilde{L}_E$

Pour une fraction massique de soluté, la sensibilité de la variation d'enthalpie à l'eutectique $\widetilde{L}_E$ se retrouve que sur le pic eutectique. Avec l'augmentation de la fraction massique de soluté, $\widetilde{L}_E$ augmente, il y a d'avantage d'énergie à échanger lors de la fusion eutectique, et comme la pente au début du pic eutectique est fixée (1.16), le pic sera plus large. De plus comme T_M diminue avec l'augmentation de la fraction massique de soluté si celle-ci reste inférieure à x_E , (2.51), les deux pics seront rapprochés et vont donc fusionner. Ce faisant, des parties de l'échantillon seront encore à la température T_E pendant le deuxième pic. Ceci explique que l'on observe un deuxième pic sur la sensibilité quand la fraction massique de soluté est plus élevée.

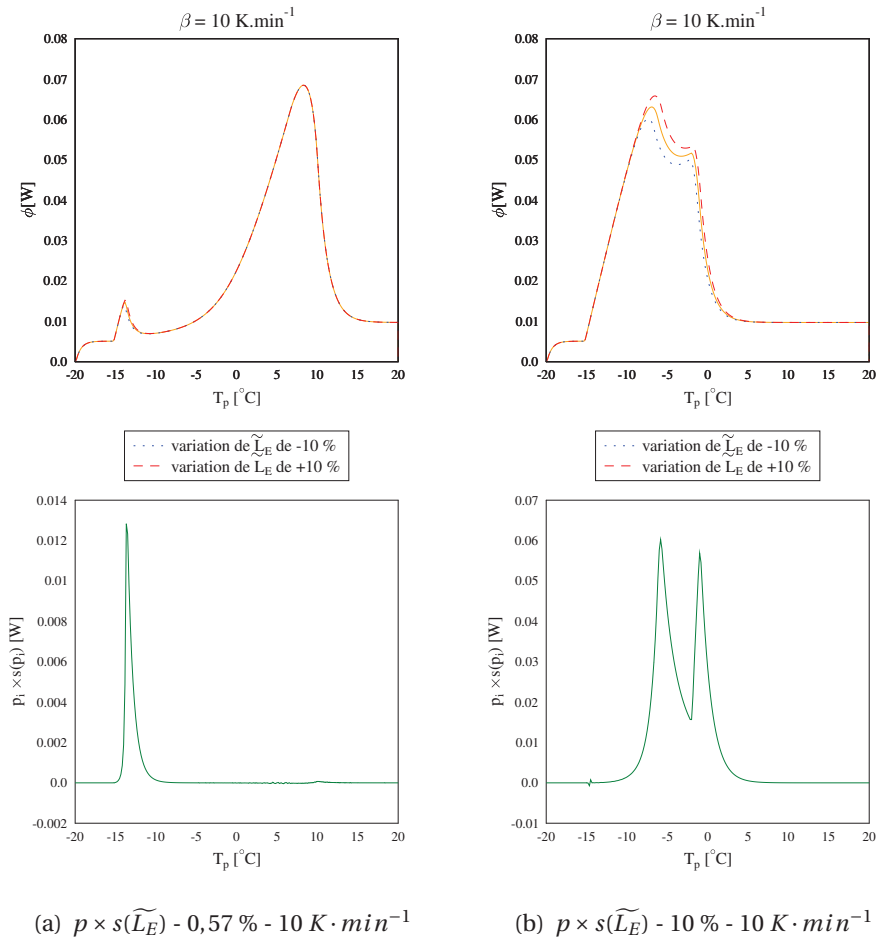
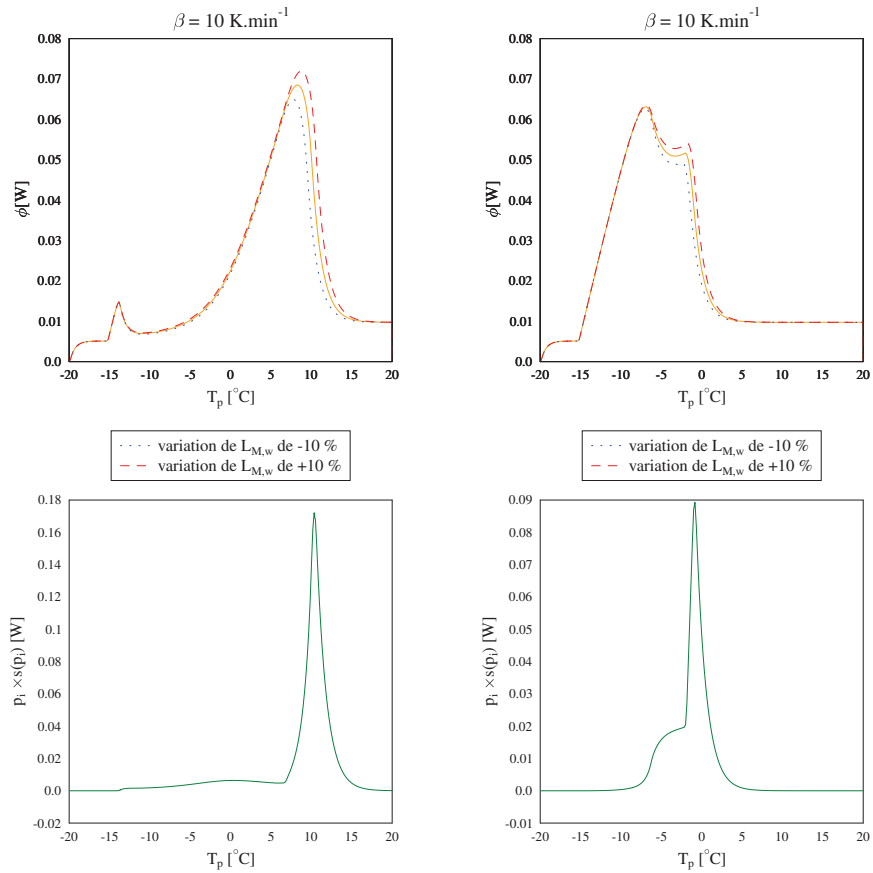


FIGURE 4.4 – Coefficient de sensibilité réduit relatif à $\widetilde{L}_E$

4.1.5 Coefficient de sensibilité réduit relatif à $L_{M,w}$

L'augmentation d'une chaleur latente, augmente la quantité d'énergie nécessaire à la fusion. L'énergie à échanger est plus importante toutes choses égales par ailleurs, notamment pour un même intervalle de température. La fusion consommant plus d'énergie, elle se déroulera sur un intervalle de temps plus long. Ce qui explique pourquoi la sensibilité à $L_{M,w}$ du modèle, n'est majoritairement présente qu'en fin de pic. On remarque que, pour de faibles fractions massiques de soluté, quand les pics sont séparés, la pente du début du deuxième pic est inchangée par la modification de $L_{M,w}$. On peut expliquer cela par le fait que pendant la fusion progressive se termine à une température proche de celle du corps pur. La pente du début de ce pic dépend donc principalement de la résistance thermique entre l'échantillon et le creuset, (1.16).



(a) $p \times s(L_{M,w}) - 0,57 \% - 10 K \cdot min^{-1}$

(b) $p \times s(L_{M,w}) - 10 \% - 10 K \cdot min^{-1}$

FIGURE 4.5 – Coefficient de sensibilité réduit relatif à $L_{M,w}$

4.1.6 Coefficient de sensibilité réduit relatif à c_L

Le coefficient de sensibilité de la capacité thermique liquide est nul avant la fusion eutectique car l'échantillon est solide à ce moment, quasi-nul au court de la fusion de l'échantillon car le transfert d'énergie y est majoritairement latent. Après la fusion la courbe de sensibilité augmente rapidement car le retour au régime stationnaire se fait par transfert de chaleur sensible.

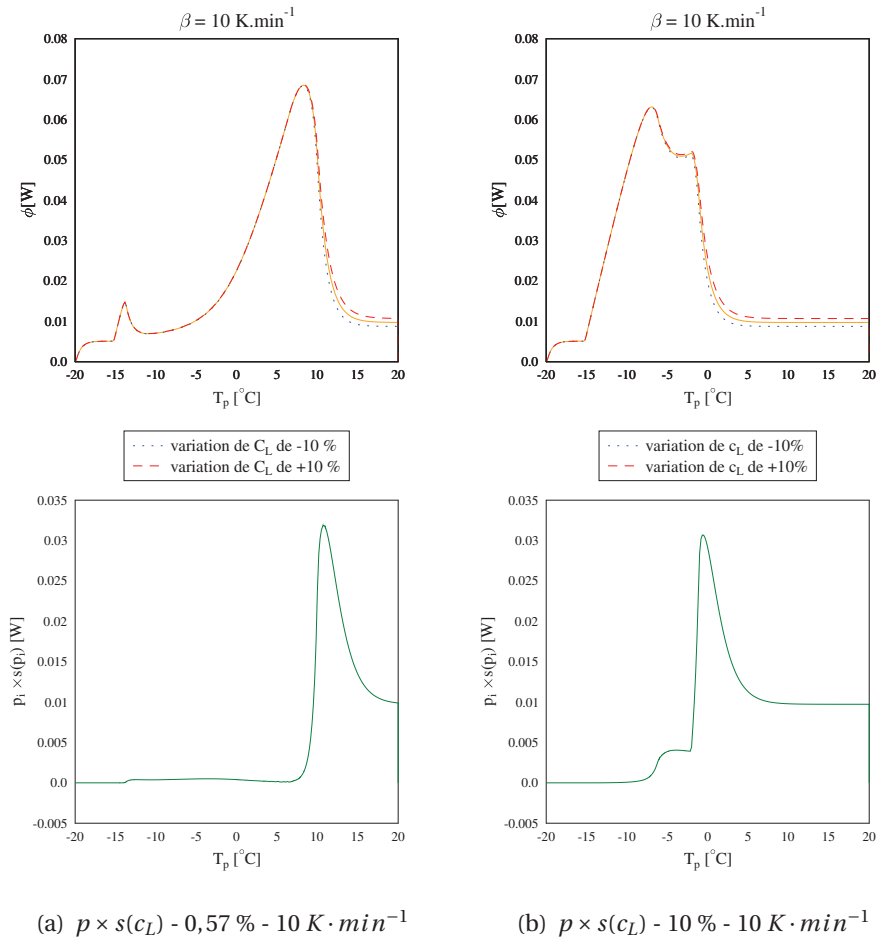


FIGURE 4.6 – Coefficient de sensibilité réduit relatif à c_L

4.1.7 Coefficient de sensibilité réduit relatif à c_S

Le coefficient de sensibilité de la capacité thermique solide est nul à la température eutectique car on a une transition à température constante. Ensuite il diminue au cours de la disparition de la phase solide avec une augmentation rapide aux points d'inflexion du thermogramme, à chaque changement dans la dynamique du transfert thermique.

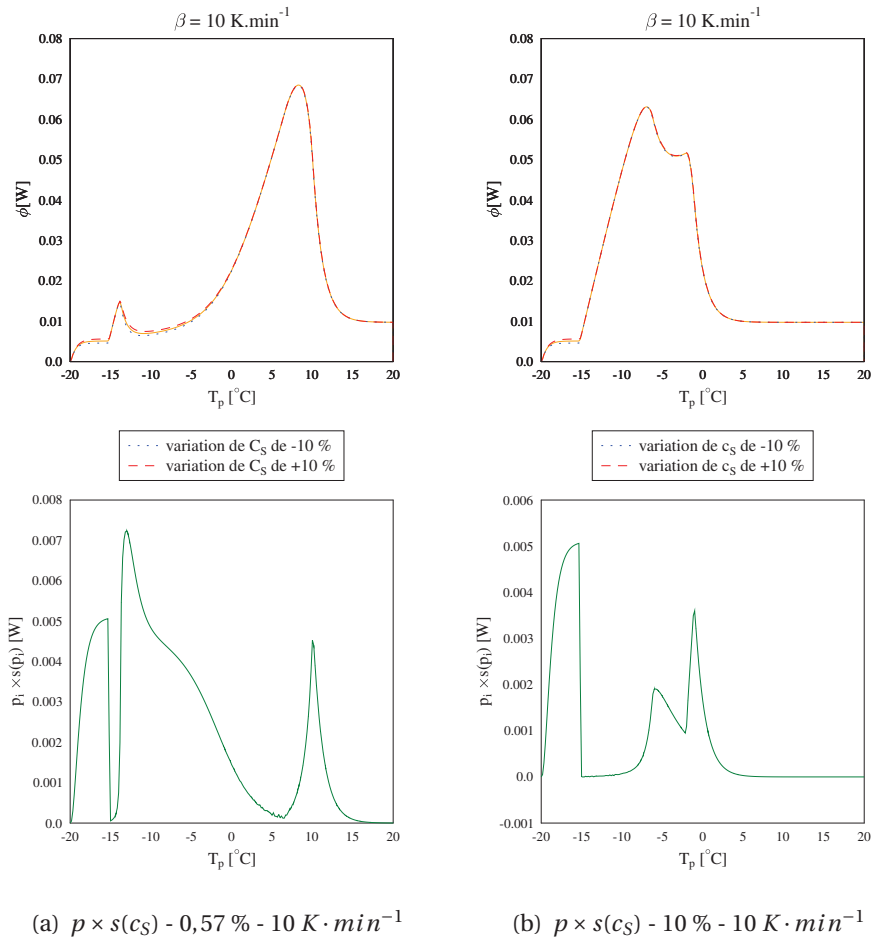


FIGURE 4.7 – Coefficient de sensibilité réduit relatif à c_S

4.1.8 Coefficient de sensibilité réduit relatif à ρ

Une variation de la masse volumique modifie l'ensemble du thermogramme mise à part la région correspondant à la pente eutectique, car la pente étant fixée, (1.16), la déformation du pic eutectique est reportée à la fin de ce pic. L'augmentation de ρ se traduit par une augmentation de l'inertie thermique du système. Au regard des quatre sensibilités précédentes, figures 4.4, 4.5, 4.6 et 4.7, il apparaît que la sensibilité du modèle au paramètre ρ est similaire à la superposition de ces courbes de sensibilité, on le vérifiera dans conclusion de cette étude, voir figure 4.13. Il en résulte un fort couplage entre les paramètres $\tilde{L}_E$, $L_{M,w}$, c_S , c_L et ρ par l'intermédiaire de ce dernier. On vérifiera cette hypothèse dans la conclusion de l'étude de sensibilité, voir figure 4.13.

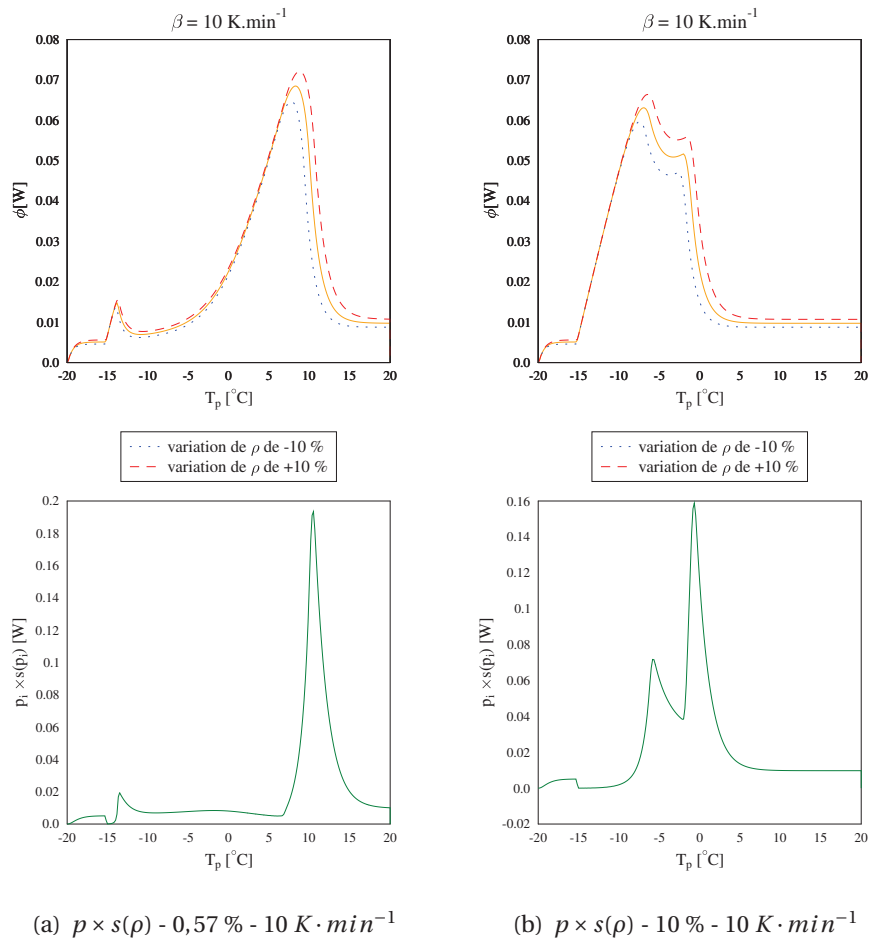


FIGURE 4.8 – Coefficient de sensibilité réduit relatif à ρ

4.1.9 Coefficient de sensibilité réduit relatif à λ_L

La sensibilité du modèle au paramètre λ_L s'observe principalement sur le deuxième pic du thermogramme, celui du solvant, sur le sommet. Avec l'augmentation de la fraction massique de soluté x_0 , on voit sur la courbe de sensibilité qu'une influence sur le premier pic du thermogramme apparaît. Cela s'explique par l'augmentation de la quantité de liquide qui apparaît au palier eutectique avec la fraction massique en soluté. Les valeurs de cette sensibilité sont très faibles comme le laisse présager la presque non-déformation du thermogramme. Ce paramètre sera donc difficile à identifier et ne pourra l'être qu'une fois que les autres paramètres seront identifiés. Il est possible que la méthode inverse paraisse avoir convergé sur tous les paramètres et qu'elle s'arrête selon le critère de convergence, alors qu'il lui faudrait plus d'itérations pour identifier la conductivité liquide. L'enthalpie que l'on cherche est indépendante de la conductivité en utilisant un modèle réduit. On identifie une conductivité apparente relative à la réduction de modèle. Elle ne présente pas d'intérêt en dehors de ce problème. On peut se satisfaire d'une valeur approchée.

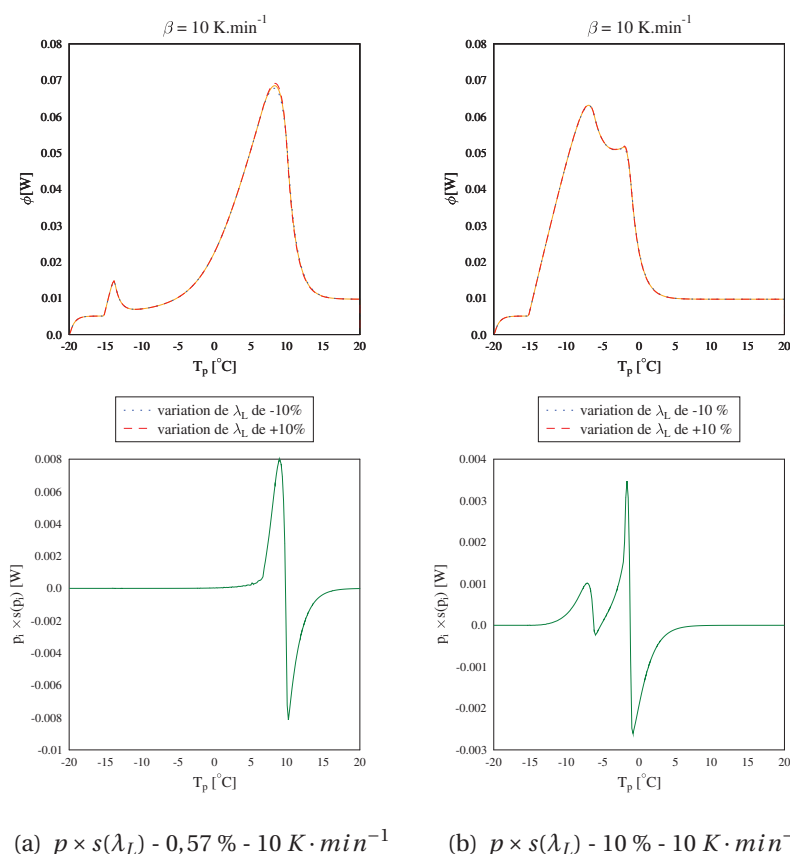


FIGURE 4.9 – Coefficient de sensibilité réduit relatif à λ_L

4.1.10 Coefficient de sensibilité réduit relatif à λ_S

Comme pour la conductivité liquide, la sensibilité du modèle à la conductivité solide est très faible. Mais dans le cas solide on observe une influence de la conductivité sur les courbes de sensibilité, correspondant aux deux pics du thermogramme pour toutes fractions massiques de soluté. Car les parties solides de l'échantillon sont soumises à un gradient thermique tout au long de la fusion contrairement au cas du corps pur où la phase solide est à une température uniforme pendant la fusion 2.3.1.1.

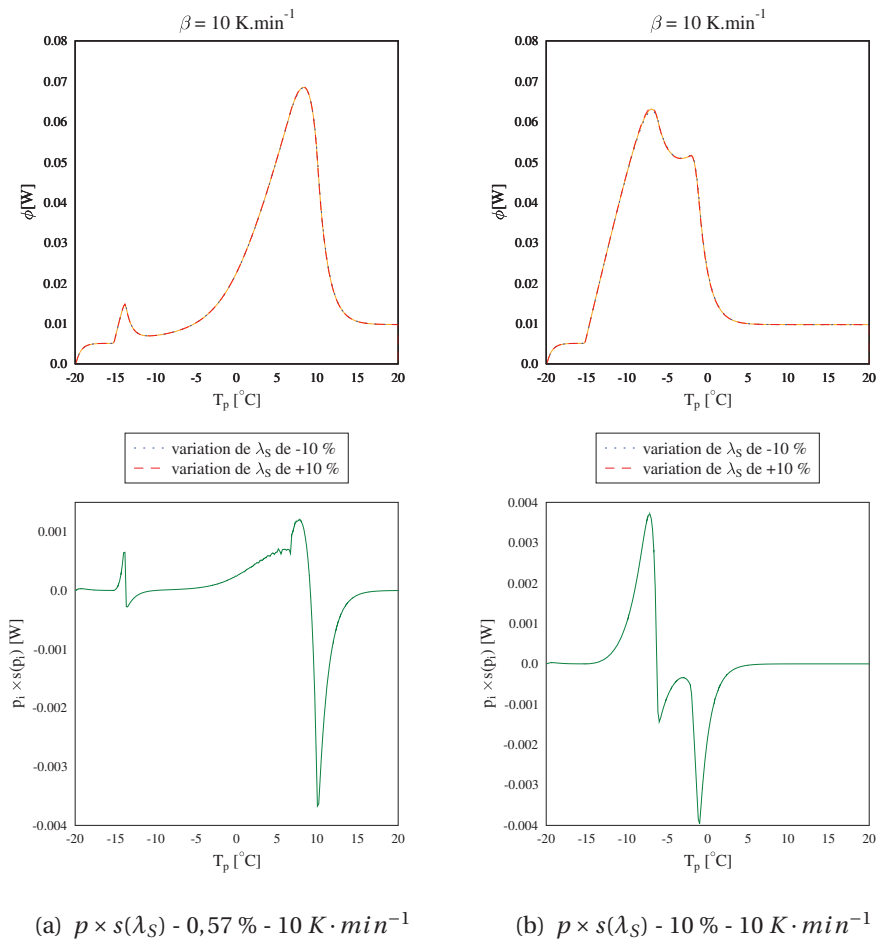


FIGURE 4.10 – Coefficient de sensibilité réduit relatif à λ_S

4.1.11 Coefficient de sensibilité réduit relatif à α

L'influence de α sur le thermogramme est de modifier la pente à l'origine des transitions types corps pur (1.16) et d'accroître ou de réduire l'échange d'énergie entre l'échantillon et le creuset. Cela se traduit par un étirement des pics vers le haut ou d'un étalement de ces pics. Cette influence est d'autant plus marquée sur les thermogrammes que l'on se rapproche des sommets, comme on peut le voir sur les sensibilités.

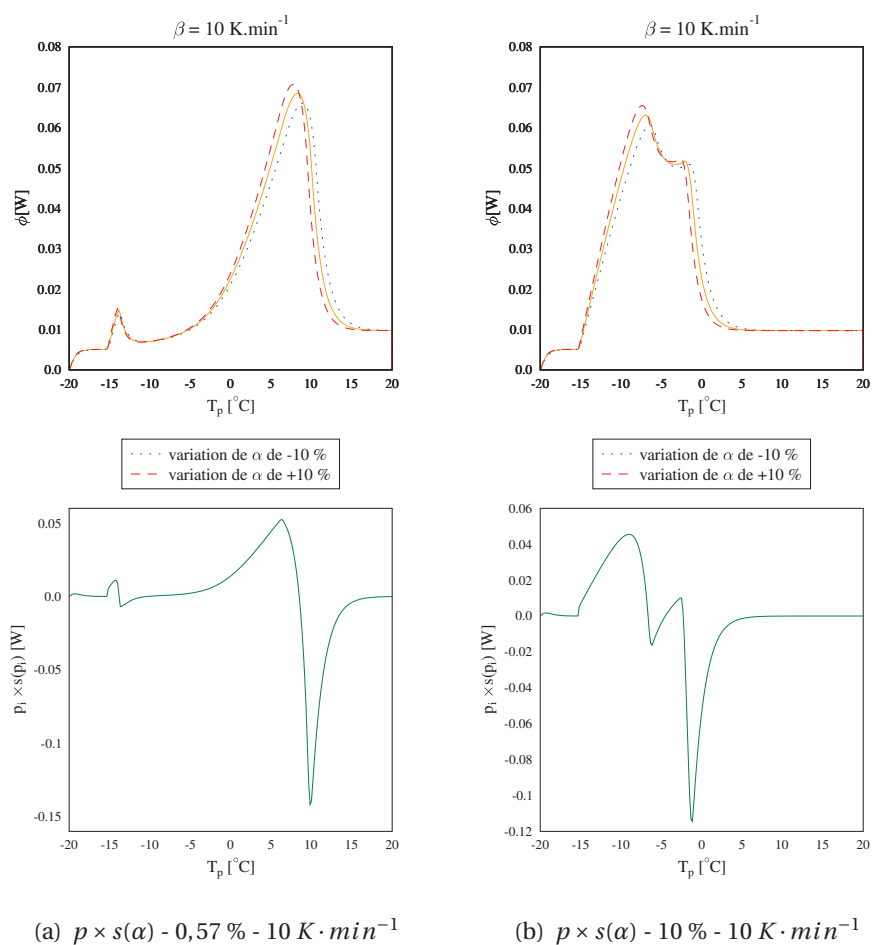


FIGURE 4.11 – Coefficient de sensibilité réduit relatif à α

4.1.12 Superposition des coefficients de sensibilité réduites

Cette étude nous a permis d'une part de constater que les paramètres les plus influents et donc les plus facilement identifiables sont les trois températures T_E , T_M et $T_{M,w}$, voir figure 4.12. On montre également que les conductivités sont les paramètres les moins influents et seront donc les plus difficiles à identifier, mais leurs valeurs exactes n'a pas un grand intérêt car il s'agit de conductivité modifiées pour la réduction de modèle. D'autre part, nous pouvons tirer de cette étude que la masse volumique va poser problème lors de l'identification. En effet, la courbe de sensibilité au paramètre ρ est très similaire à la somme des courbes de sensibilité de c_S , c_L , $L_{M,w}$ et $\tilde{L}_E$, voir figure 4.13. Cela caractérise la forte corrélation qui existe entre ces paramètres. En effet, elle est reliée à de nombreux paramètres par un produit, ce qui permet aux écarts entre les valeurs courantes des paramètres et leurs valeurs recherchées de se compenser. Par exemple si la chaleur latente du solvant, $L_{M,w}$, est supérieure à sa valeur recherchée et que la masse volumique est, elle, inférieure à sa valeur recherchée. La méthode inverse risque de préférer donner aux autres paramètres liés à ρ des valeurs supérieures à leurs valeurs recherchées, si la sensibilité à $L_{M,w}$ est la plus élevée du lot. On a expliqué que la masse volumique était une approximation nécessaire pour conserver un maillage fixe. La valeur que l'on identifie n'est pas plus pertinente qu'une autre approximation, nous la considérerons donc comme un paramètre d'entrée qui est connue. Dans la section suivante nous allons comparer les trois algorithmes d'optimisation que nous avons envisagés pour l'inversion.

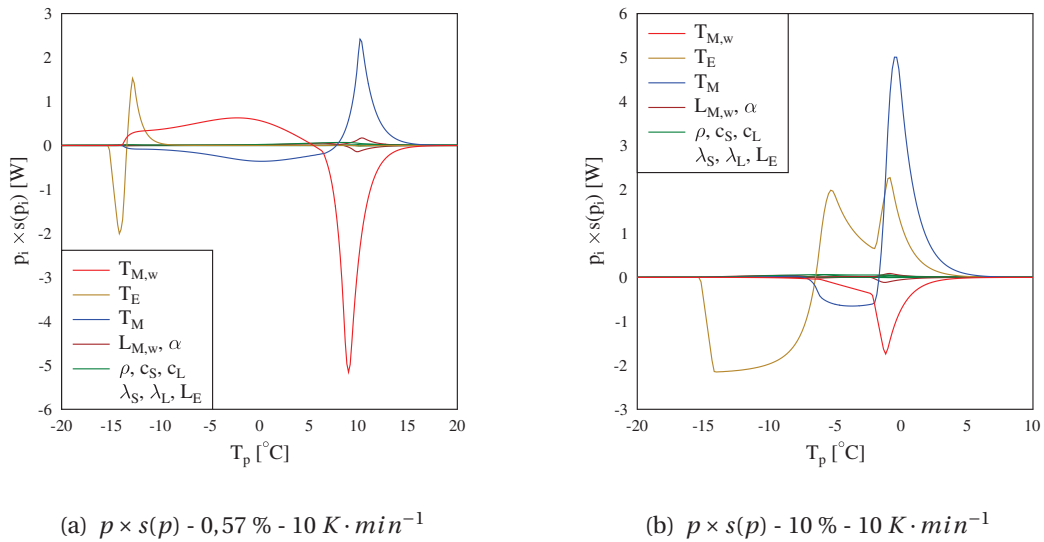
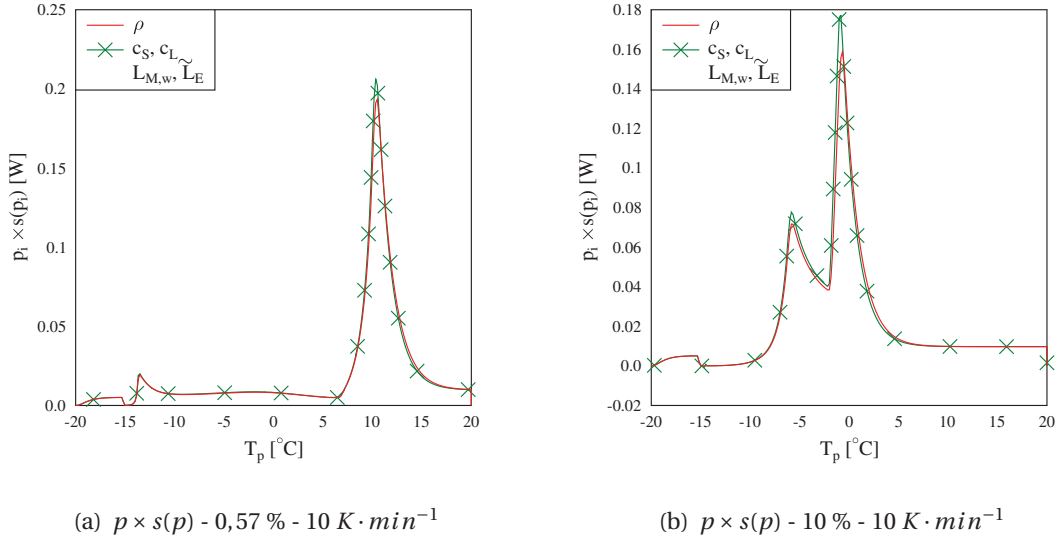


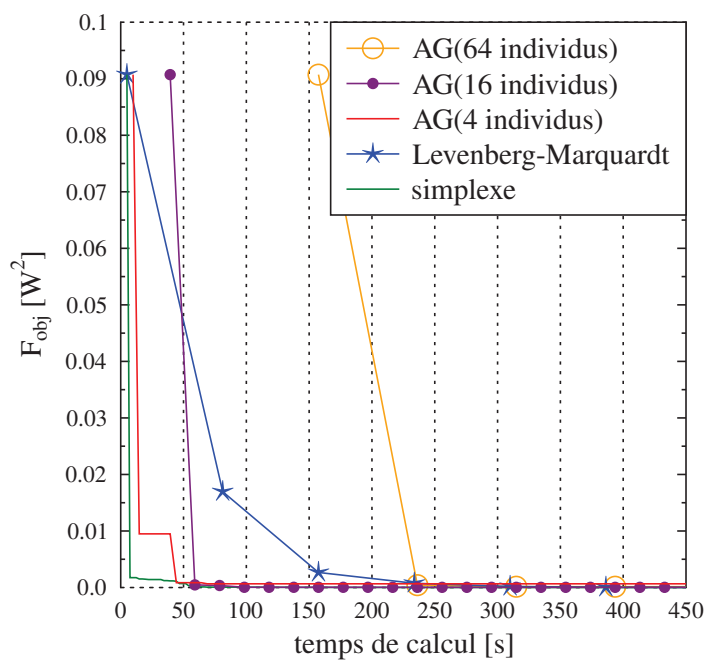
FIGURE 4.12 – Superposition des coefficients de sensibilité réduites

FIGURE 4.13 – Comparaison de la courbe de sensibilité à ρ à la somme de celles à c_S , c_L , $L_{M,w}$ et $\tilde{L}_E$

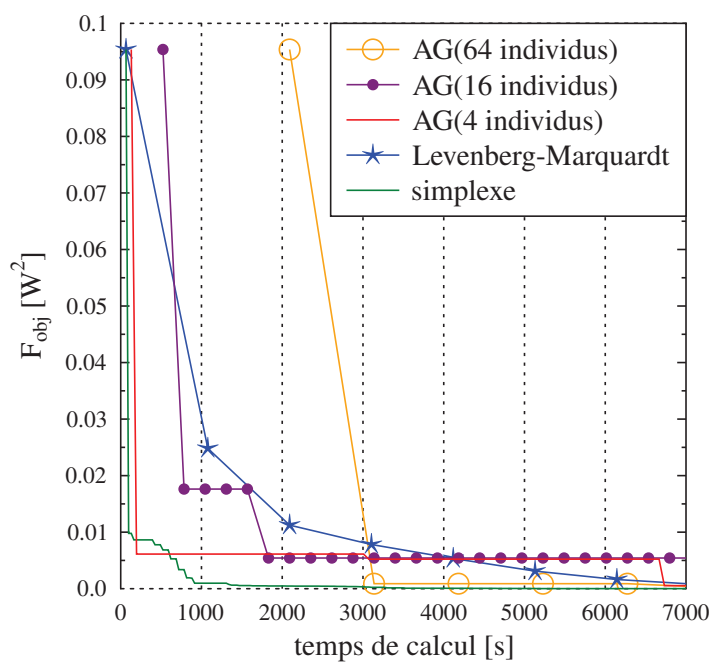
4.2 Comparaison de trois types d'algorithmes d'inversion

Pour comparer les trois méthodes que nous avons présentées au chapitre 3 que sont le simplexe, l'algorithme génétique et notre version modifiée de Levenberg-Marquardt. Nous allons tester les méthodes d'inversion sur le modèle 1D d'un corps pur, l'eau, et d'une solution binaire, $H_2O - NH_4Cl$ à 2,5 %, par des échantillons de 10,3 mg réchauffés à $5 K \cdot \min^{-1}$. La figure 4.14 donne la diminution de la fonction objectif au long du processus d'identification afin de s'assurer que les algorithmes ne rencontrent pas de problème. Le modèle direct du corps pur est résolu en 2,46 s et celui de la solution binaire est résolue en 32,68 s, pour un maillage 1D de 20 nœuds et un pas de temps constant de 10^{-4} s. Les algorithmes de Levenberg-Marquardt et du simplexe sont initialisés avec des paramètres 10% ou 0,5 K en-dessous de leur valeur à la solution, à partir de là le simplexe est construit avec des sommets à 1% ou 0,1 K. L'algorithme génétique est initialisé dans un cercle de 10% ou 0.5 K autour de la solution.

Sur la figure 4.14 on voit que c'est le simplexe qui identifie les paramètres le plus rapidement.



(a) modèle du corps pur



(b) modèle de la solution binaire

FIGURE 4.14 – Evolution de la fonction objectif au cours de l'identification globale.

Sur les figures 4.14(a) et 4.14(b) on voit que l'algorithme génétique identifie la solution en peu d'itérations. Ceci est dû au fait que l'initialisation se fait dans un espace restreint proche de la solution, l'algorithme a donc de grandes chances d'évaluer un bon candidat du premier coup. On remarque aussi que plus la population de l'algorithme génétique est faible, au plus les itérations seront rapides car la population à évaluer est moins importante. Mais la population étant alors plus faible, ses estimations de la solution seront de moindres qualités à chaque itération et il en faudra d'avantage pour identifier une meilleure solution que si la population était plus élevée.

Les avantages principaux de l'algorithme de descente de Levenberg-Marquardt est que sa convergence peut être démontrée et que la convergence est sur la solution contrairement à l'algorithme génétique qui lui, doit rassembler sa population autour de la solution afin de pouvoir l'approcher au mieux, voir figure 4.15.

Au cours des premières itérations, l'algorithme génétique peut améliorer grandement la solution parce qu'elle est dispersée sur tout le domaine de recherche et il se peut qu'aucun individu ne soit proche de la solution. Les individus enfants ont donc des chances d'avoir une bien meilleure valeur pour leur fonction objectif que leurs parents. Alors que, à la fin du processus, la population est concentrée autour de la solution et donc tous les individus ont un même ordre de grandeur pour la fonction objectif.

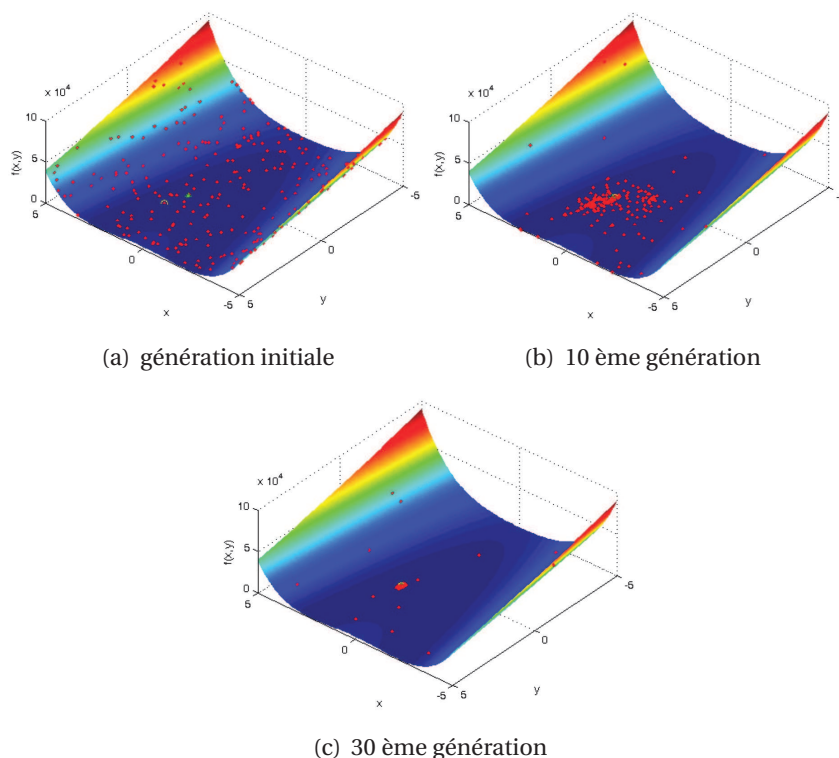


FIGURE 4.15 – Evolution de la population de l'algorithme génétique pour la résolution de la fonction de Rosenbrock. [77]

Voici nos conclusions sur ces algorithmes :

4.2.1 Algorithme de Levenberg-Marquardt mis en œuvre : (voir 3.3.1.6)

Mathématiquement très complet et dont la convergence a été démontrée, il cherche à emprunter toujours le meilleur chemin de convergence, mais cette recherche prend beaucoup de temps et peut ne pas aboutir. Si la fonction objectif peut être exprimée directement en fonction des paramètres, alors cette méthode est bien adaptée car les calculs sont exacts. Mais dans notre cas, on ne peut exprimer directement la fonction objectif et on est obligé d'utiliser des approximations pour le calcul des dérivées. De ce fait le calcul des gradients de la fonction objectif est inexact et donc la matrice Hessienne approximée nécessite d'être régularisée pour en améliorer le conditionnement, avant d'être inversée, les paramètres de régularisation n'ayant pas de règles quant à leurs valeurs. Un autre problème de cet algorithme est son incapacité à sortir des minimums locaux, le calcul du gradient ne fournissant dans ce cas aucune direction de descente.

4.2.2 Algorithme Génétique : (voir 3.3.2.1)

Il est robuste, on est assuré d'avoir la convergence, mais il consomme beaucoup de temps de calcul. Les minimums locaux n'ont pas d'influences particulières. Cela s'explique par son caractère aléatoire, l'algorithme combinant les solutions selon un procédé que l'on a choisi à l'avance. Le problème de cet algorithme est qu'il utilise une grande population, ce qui nécessite donc d'évaluer de nombreuses fois la fonction objectif, i.e. résoudre de nombreuses fois le modèle de la fusion, ce qui nécessite d'autant plus de temps de calcul que la population est importante. Il présente néanmoins l'avantage d'avoir une initialisation relativement facile à choisir, car on initialise un domaine et pas un point.

4.2.3 Algorithme du Simplexe : (voir 3.3.3.1)

Cet algorithme n'a pas d'étapes mathématiquement lourdes comme le calcul d'un gradient ou l'inversion d'une matrice et peut se révéler très rapide si il a été initialisé judicieusement et de taille adéquate. Cet algorithme est à la base une méthode de descente selon l'hyperplan tangent. Mais c'est aussi une méthode exploratrice comme les algorithmes évolutionnaires. Il teste plusieurs solutions potentiels appartenant à l'hyperplan pour évaluer le gradient au lieu de le calculer, afin de déplacer l'un de ses sommets sans nécessairement améliorer la meilleure solution courante comme avec les méthodes déterministes. Nous ne l'avons pas observé sur les résultats que nous présentons mais certains auteurs ont trouvé que selon la taille et l'initialisation du simplexe, il peut être pris dans un minimum local, différent du minimum global. Mais une réinitialisation du simplexe sur la solution courante quand ça arrive permet d'atteindre le minimum global [87] [88].

Nous avons choisi de nous servir de l'algorithme du simplexe pour notre méthode inverse de part sa simplicité et sa vitesse.

4.3 Applications de l'inversion à des cas théoriques

4.3.1 Identification sur des thermogrammes simulés

Avant de tester notre méthode sur des cas réels où il est parfois difficile de trouver les valeurs des paramètres dans la littérature, et d'évaluer l'impact des différentes sources d'erreurs, nous allons d'abord inverser avec le modèle 1D des thermogrammes simulés avec le modèle 2D, pour voir si on retrouve les paramètres énergétiques qui ont servi lors des simulations. Nous avons considéré des échantillons de rayon 2,125 mm et de hauteur 0,78 mm.

Les paramètres énergétiques qui ont servi à construire ces thermogrammes sont donnés dans les tableaux suivants :

Tableau 4.1 – Propriétés de corps purs [63] [64]

Matériaux	ρ $kg \cdot m^{-3}$	c_S $J \cdot kg^{-1} \cdot K^{-1}$	c_L $J \cdot kg^{-1} \cdot K^{-1}$	T_M $^{\circ}C$	L_M $kJ \cdot kg^{-1}$
eau	1000	2060	4180	0	333
hexadécane	780	1680	2310	18,3	228,9

Tableau 4.2 – Propriétés de solutions binaires. [52]

Matériaux	ρ $kg \cdot m^{-3}$	$c_L - c_S$ $J \cdot kg^{-1} \cdot K^{-1}$	$T_{M,w}$ $^{\circ}C$	$L_{M,w}$ $kJ \cdot kg^{-1}$	T_E $^{\circ}C$	x %	$\tilde{L}_E^1$ $kJ \cdot kg^{-1}$	T_M^2 $^{\circ}C$
$H_2O - NH_4Cl$	1000	2120	0	333	-15,9	0,57	8,2	-0,46
						2,5	35,9	-2,02
						5	74,4	-4,03
						10	144	-8,07

Les identifications sur ces thermogrammes de références ont donc été faites avec l'algorithme du simplexe. Les initialisations ont été faites à 10% ou 0,5 K en dessous des valeurs des paramètres de référence, comme pour la comparaison entre les trois algorithmes d'inversions, section 4.2. Les thermogrammes correspondant à ces identifications sont disponibles dans les figures 4.16 et 4.17. Les paramètres correspondants à ces identifications sont disponibles dans les tableaux 4.3 et 4.4. Une excellente concordance est obtenue sur les thermogrammes et les propriétés énergétiques correspondent à celles qui ont servi pour générer les thermogrammes de références. Ces identifications ont été obtenues en moins de 200 itérations.

1. Ces valeurs ont été déterminées à partir de nos identifications sur des thermogrammes expérimentaux

2. calculée avec l'équation (2.51)

Tableau 4.3 – Propriétés énergétiques des corps purs identifiées

Matériaux	ρ $kg \cdot m^{-3}$	c_S $J \cdot kg^{-1} \cdot K^{-1}$	c_L $J \cdot kg^{-1} \cdot K^{-1}$	T_M $^{\circ}C$	L_M $kJ \cdot kg^{-1}$
eau	1000	2059,9	4180,03	-3×10^{-5}	333,2
hexadécane	780	1679,9	2308	18,3	229,1

Tableau 4.4 – Propriétés énergétiques des solutions binaires identifiées

Matériaux	ρ $kg \cdot m^{-3}$	$c_L - c_S$ $J \cdot kg^{-1} \cdot K^{-1}$	$T_{M,w}$ $^{\circ}C$	$L_{M,w}$ $kJ \cdot kg^{-1}$	T_E $^{\circ}C$	x %	$\widetilde{L}_E$ $kJ \cdot kg^{-1}$	T_M $^{\circ}C$
$H_2O - NH_4Cl$	1000	2118,5	9×10^{-3}	332,8	-15,86	0,57	8,24	-0,46
		2119,6	$1,3 \times 10^{-2}$	332,8	-15,9	2,5	35,93	-2,02
		2120,6	4×10^{-2}	333,5	-15,9	5	74,44	-4,03
		210,2	5×10^{-4}	333,0	-15,9	10	143,9	-8,07

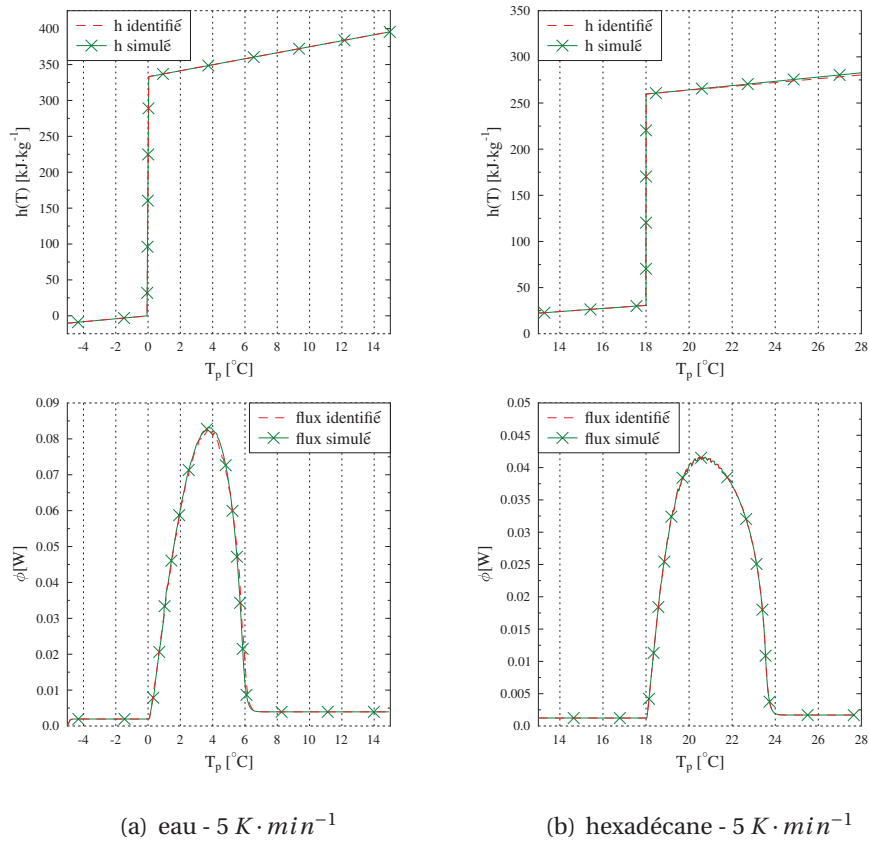
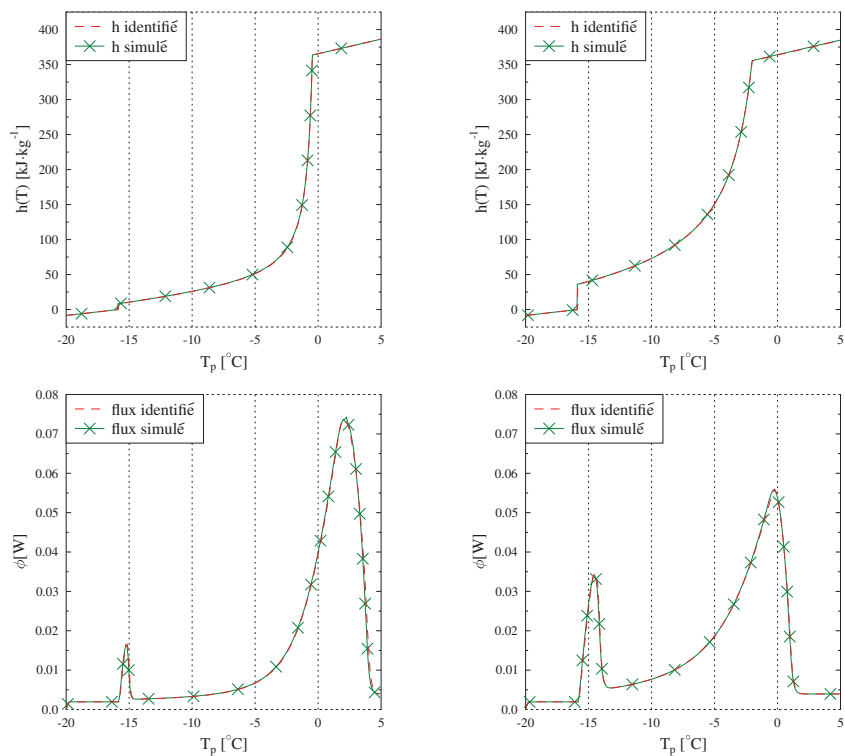
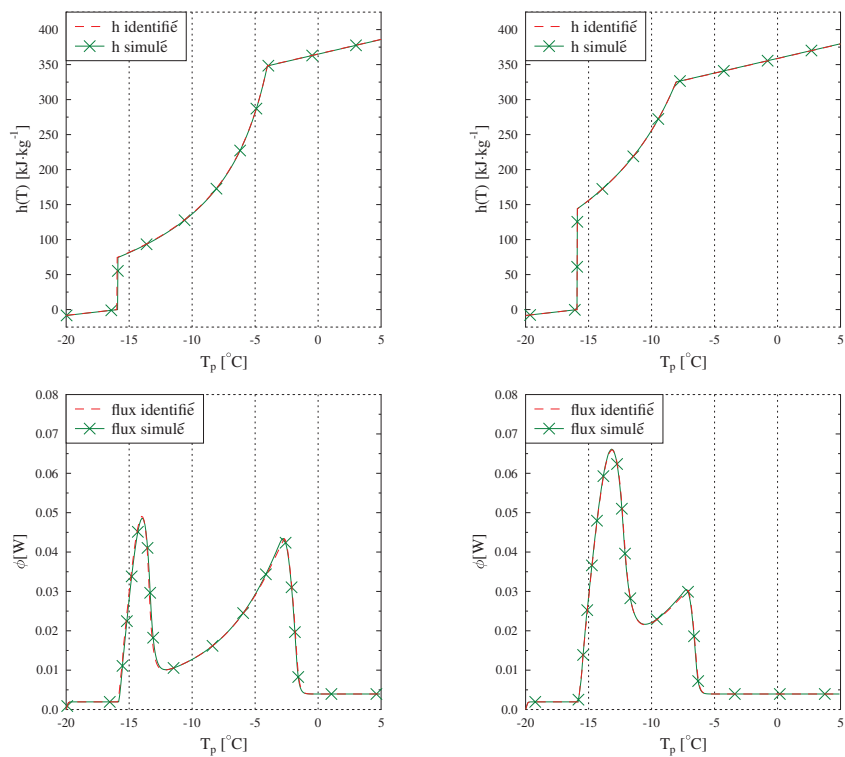


FIGURE 4.16 – Identifications sur des thermogrammes simulés de corps purs

(a) $\text{H}_2\text{O}-\text{NH}_4\text{Cl}$ - 0,57 % - $5\text{ K}\cdot\text{min}^{-1}$ (b) $\text{H}_2\text{O}-\text{NH}_4\text{Cl}$ - 2,5 % - $5\text{ K}\cdot\text{min}^{-1}$ (c) $\text{H}_2\text{O}-\text{NH}_4\text{Cl}$ - 5,0 % - $5\text{ K}\cdot\text{min}^{-1}$ (d) $\text{H}_2\text{O}-\text{NH}_4\text{Cl}$ - 10 % - $5\text{ K}\cdot\text{min}^{-1}$ FIGURE 4.17 – Identifications sur des thermogrammes simulés de solution binaire $\text{H}_2\text{O}-\text{NH}_4\text{Cl}$

4.3.2 Influence de l'erreur sur l'inversion : incertitudes sur les résultats

Une information importante pour juger de la qualité des résultats d'une méthode inverse est l'influence sur les paramètres identifiés, des erreurs sur les paramètres d'entrée du système. Ces erreurs sont de six types [72] :

- Le calcul numérique est source d'erreur. D'abord par la discrétisation du milieu, de phénomènes continus. Souvent on ne peut calculer les vraies valeurs des dérivées, on utilise alors des développements limités. L'erreur peut également venir du schéma numérique, le numéricien doit stabiliser le schéma avec les paramètres numériques de son code.
- Il peut y avoir des erreurs dans la modélisation des phénomènes. Par exemple nous avons négligé les variations de volume, de masse volumique et nous avons fait des approximations sur les capacités thermiques.
- Les résultats expérimentaux sur lesquels on va appliquer la méthode inverse présentent une incertitude expérimentale ou bruit. Aucune mesure n'est parfaite. Il s'agit là d'une erreur aléatoire sur chaque point expérimental. C'est l'erreur stochastique.
- Les systèmes d'acquisition de la mesure peuvent aussi donner des erreurs. Par exemple s'ils sont mal étalonnés, s'il y a un défaut sur les signaux électriques ou si la mauvaise utilisation d'un filtre introduit un biais.
- Il arrive que sur le jeu de paramètres du modèle, certains soient déjà connus. Ces valeurs sont connues à une incertitude près. Or si on fixe un paramètre du modèle, ce dernier est considéré comme ayant sa valeur vraie et donc l'incertitude sur sa valeur devient une erreur sur l'entrée du modèle. C'est l'erreur déterministe.
- Enfin, comme on l'a évoqué précédemment, les méthodes de régularisation peuvent induire un biais dans la résolution du problème inverse.

Nous considérons que les calculs numériques se sont bien déroulés avec un maillage à 20 nœuds, un pas de temps constant de 10^{-4} s, et que l'effort de modélisation a été suffisant afin que l'on puisse négliger les erreurs de modélisation et de calcul numérique. Nous allons aussi considérer que la calorimétrie est suffisamment mature pour que les systèmes d'acquisitions ne provoquent pas d'erreurs supplémentaires. Néanmoins nous ne sommes pas sûr de comment la régulation est effectuée. Le fichier de résultats donne deux variations de température : « Program Temperature » et « Sample Temperature ». La variation de « Program Temperature » a une dépendance linéaire en fonction du temps, nous avons donc considéré qu'elle était la température de plateau T_p . Par précaution nous allons suivre l'avis du constructeur donnant une incertitude d'offset sur T_p de $\Delta_T p_i = 0,1 \text{ K}$ [32]. Nous allons rajouter cette incertitude dans le calcul de l'intervalle de confiance des températures identifiées par le processus d'inversion.

Nous avons présenté trois méthodes d'inversion : une version modifiée de l'algorithme de Levenberg-Marquardt, l'algorithme génétique et le simplexe. Les deux dernières sont des méthodes stochastiques et ne requièrent pas de régularisation. Pour la première, la modification proposée par [61] diminue le biais de la régularisation à mesure que l'on s'approche de la solution. Nous pouvons donc considérer cette source d'erreur comme nulle.

Dans ces conditions, il ne reste plus que deux sources d'erreurs à étudier : l'erreur déterministe sur les paramètres « connus », et l'erreur stochastique ou « bruit ». Nous allons donc les déterminer puis les combiner pour obtenir un intervalle de confiance sur le paramètre p_i ,

$\Delta_c p_i$, causé par les paramètres d'entrées que l'on définit par [61] [89] [90] :

$$\Delta_c p_i = \sqrt{\sum_{\substack{j=1 \\ j \neq i}}^{n_{p,e}} (\Delta_{\pi_j} p_i)^2 + (\Delta_{\sigma} p_i)^2 + (\Delta_T p_i)^2} \quad (4.3)$$

Avec $\Delta_{\pi_j} p_i$ l'erreur résultante sur le paramètre j d'une erreur sur le paramètre d'entrée π , $\Delta_{\sigma} p_i$ l'influence d'un bruit d'écart-type σ sur le paramètre i et $\Delta_T p_i = 0$ sauf si p_i est une température.

Les méthodes de descente permettent un calcul théorique des erreurs dues à l'inversion. Nous validerons ces calculs en comparant les résultats avec ceux que l'on obtient en identifiant sur les thermogrammes numériques que l'on a générés avec une erreur connue. Ces expériences numériques seront détaillées à chaque fois.

4.3.2.1 Erreur déterministe

Ici nous allons étudier l'impact d'une erreur sur chacun des paramètres énergétiques et de transfert si il étaient pris comme paramètre d'entrée. Afin de conforter notre choix de paramètres d'entrée pour notre méthode d'inversion.

Considérons l'itération finale de la méthode de Levenberg-Marquardt. Proche de la solution le terme de régularisation est quasiment nul, si on considère de plus que $\omega^k = 1$ dans l'équation (3.36) pour que la norme de l'erreur soit contenue dans $\Delta \mathbf{p}$, celle-ci devient :

$$\Delta \mathbf{p} = \left[{}^t \mathbb{S}(\mathbf{p}^k) \cdot \mathbb{S}(\mathbf{p}^k) \right]^{-1} \cdot {}^t \mathbb{S}(\mathbf{p}^k) \cdot (\mathbf{y}^{k+1} - \mathbf{y}^k) \approx 0 \text{ car } \mathbf{y}^{k+1} \approx \mathbf{y}^k \quad (4.4)$$

Supposons que l'on introduise une variation $\delta \pi$ sur le paramètre d'entrée π . Il va en résulter une variation $\Delta \mathbf{y}$ sur la sortie du modèle. En admettant que cette variation corresponde à la dernière itération d'une identification dans le sens opposé. L'équation (4.4) devient :

$$\Delta_{\pi} \mathbf{p} = \left[{}^t \mathbb{S}(\mathbf{p}^k) \cdot \mathbb{S}(\mathbf{p}^k) \right]^{-1} \cdot {}^t \mathbb{S}(\mathbf{p}^k) \cdot (-\Delta \mathbf{y}) \quad (4.5)$$

Si la variation $\Delta \pi$ est faible par rapport au paramètre π , et que la variation est nulle sur les autres, on peut linéariser le modèle direct autour du paramètre π avec un développement limité :

$$\mathbf{y}(\pi + \Delta \pi) = \mathbf{y}(\pi) + \mathbb{S}(\pi) \cdot \Delta \pi \quad (4.6)$$

Or $\Delta \pi$ est un vecteur constitué de zéro et du scalaire $\Delta \pi$, on peut donc réécrire l'équation (4.6) avec ce scalaire et remplacer la matrice $\mathbb{S}(\pi)$ par sa colonne qui correspond au paramètre π , $\mathbb{S}_{\pi}$.

$$\Delta \mathbf{y} = \mathbb{S}_{\pi}(\pi) \times \Delta \pi \quad (4.7)$$

En injectant ce résultat dans l'équation précédente on a :

$$\Delta_{\pi} \mathbf{p} = - \left[{}^t \mathbb{S}(\mathbf{p}^k) \cdot \mathbb{S}(\mathbf{p}^k) \right]^{-1} \cdot {}^t \mathbb{S} \cdot \mathbb{S}_{\pi}(\pi) \times \Delta \pi \quad (4.8)$$

Connaissant les matrices de sensibilité du modèle on peut prédire l'erreur sur l'inversion, qui a été engendrée une erreur sur un paramètre supposé connu.

Nous allons maintenant comparer les résultats de ce calcul théorique avec une détermination « expérimentale ». On obtient ces dernières valeurs en simulant un thermogramme en ayant faussé la valeur de l'un des paramètres de 1% ou 0,1 K par rapport à sa valeur dite de référence, en fixant ce paramètre et en identifiant les autres paramètres avec la méthode du simplexe. L'écart entre les valeurs que l'on obtient par cette identification et les valeurs de référence sont les erreurs engendrées par une erreur $\Delta\pi$ sur le paramètre d'entrée π .

Pour la masse volumique ρ nous pouvons calculer une incertitude sur sa valeur en considérant un échantillon fictif qui remplit parfaitement la cellule. La cellule cylindrique est de dimension $R=2,15 \text{ mm}$, $Z=0,78 \text{ mm}$, ce qui pour un échantillon de masse volumique $\rho = 1000 \text{ kg} \cdot \text{m}^{-3}$ donne une masse de 11,33 mg. En considérant que l'incertitude sur le rayon, et la hauteur est de 0,01 mm et que la masse est connue à $\pm 0,1 \text{ mg}$ près. Il vient :

$$\begin{aligned} \frac{\Delta\rho}{\rho} &= \frac{\Delta m}{m} + \frac{2\Delta R}{R} + \frac{\Delta Z}{Z} \\ \Delta\rho &= 27,5 \text{ kg} \cdot \text{m}^{-3} \approx 3\% \end{aligned} \quad (4.9)$$

On considèrera cette dernière valeur pour l'erreur déterministe sur ρ .

On voit sur ces tableaux que des erreurs sur les températures, $T_{M,w}$ et T_E , entraînent des erreurs très importantes sur le résultat de l'inversion dans le cas du corps pur et de la solution binaire, les autres températures caractéristiques ont également une influence importante mais dans une moindre mesure. Si l'étalonnage en température du calorimètre n'est pas parfait et si on considère une de ces températures comme paramètres d'entrée du modèle, on va, de ce fait, introduire une erreur déterministe sur ces températures, ce qui va fausser le résultat.

Les paramètres α et λ_L ont également une influence forte dans le cas du corps pur, mais cela ne nous pose pas de problème car ces propriétés étant modifiées par la réduction de modèle on ne peut considérer ces paramètres comme des paramètres d'entrée.

Il est intéressant de noter qu'une erreur sur ρ , c_S , c_L , $L_{M,w}$ ou $\tilde{L}_E$ sera reportée de la même quantité en valeur absolue sur les autres mêmes paramètres. Ceci est dû à ce que ρ relie tous ces paramètres ensemble, une erreur déterministe sur l'un d'entre eux va se reporter sur ces autres paramètres lors de l'identification, comme nous l'avons prévu au cours de l'analyse des coefficients de sensibilité 4.1.

Le tableau 4.6(b) indique qu'une erreur sur les paramètres d'entrée n'a pratiquement pas d'influence sur les températures T_E , T_M et $T_{M,w}$ que l'on identifie. Cela pour les deux modèles thermodynamiques. Alors que, en retour, une erreur sur ces températures influence de manière plus importante les autres paramètres du modèle. C'est pourquoi on se propose de rajouter l'incertitude des capteurs de températures, $\Delta T p_i$, dans le calcul de l'intervalle de confiance de T_E , T_M et $T_{M,w}$.

Tableau 4.5 – Influence théorique et expérimentale d’une erreur déterministe de 1% ou 0,1 K sur le résultat de l’inversion mise à part ρ où l’erreur est de 3% - cas d’un corps pur

(a) eau - propriétés énergétiques

X	$\varepsilon(c_S)$		$\varepsilon(c_L)$		$\varepsilon(T_{M,w})$		$\varepsilon(L_{M,w})$	
	exp %	th %	exp %	th %	exp %	th %	exp %	th %
ρ	3,2	3	2,9	3	0,001	8×10^{-4}	3,1	3
c_S	X	X	0,89	0,9	$8,1 \times 10^{-5}$	8×10^{-5}	0,89	0,9
c_L	0,97	1	X	X	$6,5 \times 10^{-5}$	$6,5 \times 10^{-5}$	0,98	1
λ_S	$2,4 \times 10^{-2}$	2×10^{-2}	0,03	0,04	$4,5 \times 10^{-5}$	$4,8 \times 10^{-5}$	$2,1 \times 10^{-2}$	2×10^{-2}
λ_L	1,4	1,5	1,7	1,5	1×10^{-4}	$9,5 \times 10^{-4}$	1,5	1,5
$T_{M,w}$	4,5	3,3	4,6	4,5	X	X	4,8	4,8
$L_{M,w}$	0,99	1	0,96	1	1×10^{-4}	$9,5 \times 10^{-5}$	X	X
α	5,1	5	3,9	3,7	0,04	0,03	5,5	5,5

(b) eau - propriétés pour le modèle

X	$\varepsilon(\rho)$		$\varepsilon(\lambda_S^{1D})$		$\varepsilon(\lambda_L^{1D})$		$\varepsilon(\alpha)$	
	exp %	th %	exp %	th %	exp %	th %	exp %	th %
ρ	X	X	$0,9 \times 10^{-3}$	2×10^{-3}	3×10^{-4}	1×10^{-4}	0,06	0,04
c_S	0,89	0,9	0,03	0,03	$5,2 \times 10^{-3}$	5×10^{-3}	$0,9 \times 10^{-2}$	1×10^{-2}
c_L	0,98	1	$5,6 \times 10^{-2}$	6×10^{-2}	$6,3 \times 10^{-3}$	6×10^{-3}	$9,3 \times 10^{-3}$	9×10^{-3}
λ_S	$2,4 \times 10^{-2}$	3×10^{-2}	X	X	$4,2 \times 10^{-2}$	4×10^{-2}	$2,3 \times 10^{-2}$	$2,4 \times 10^{-2}$
λ_L	1,4	1,5	16	15	X	X	0,25	0,24
$T_{M,w}$	4,5	4,5	3,1	3	0,2	0,17	0,6	0,5
$L_{M,w}$	0,99	1	$2,9 \times 10^{-2}$	3×10^{-2}	$5,8 \times 10^{-3}$	6×10^{-3}	2×10^{-2}	$1,4 \times 10^{-2}$
α	5,1	5	12	12,5	0,35	0,4	X	X

Tableau 4.6 – Influence théorique et expérimentale d'une erreur déterministe de 1% ou 0,1 K sur le résultat de l'inversion mise à part ρ ou l'erreur est de 3% - cas d'un binaire(a) $H_2O - NH_4Cl$, 2,5% - propriétés pour le modèle

X	$\varepsilon(\rho)$		$\varepsilon(\lambda_S^{1D})$		$\varepsilon(\lambda_L^{1D})$		$\varepsilon(\alpha)$	
	exp %	th %	exp %	th %	exp %	th %	exp %	th %
ρ	X	X	0,3	0,09	0,1	0,06	4×10^{-3}	7×10^{-3}
c_S	0,94	1	$2,6 \times 10^{-2}$	3×10^{-2}	$2,3 \times 10^{-3}$	2×10^{-3}	3×10^{-3}	3×10^{-3}
c_L	0,96	1	$2,6 \times 10^{-2}$	3×10^{-2}	3×10^{-3}	3×10^{-3}	$2,1 \times 10^{-3}$	2×10^{-3}
λ_S^{1D}	0,33	0,3	X	X	$3,2 \times 10^{-2}$	4×10^{-2}	0,19	0,2
λ_L^{1D}	0,6	0,6	0,3	0,2	X	X	0,2	0,18
$T_{M,w}$	0,5	0,6	0,6	0,5	$6,7 \times 10^{-2}$	7×10^{-2}	0,15	0,14
$L_{M,w}$	0,9	1	0,07	0,04	7×10^{-3}	4×10^{-3}	8×10^{-3}	6×10^{-3}
T_E	3,6	5,5	60	86	1,2	2,2	19	25
$\widetilde{L}_E$	1	1	0,05	0,03	0,01	0,02	3×10^{-3}	3×10^{-3}
T_M	0,5	0,7	3,7	4,2	0,2	0,3	1,3	1,2
α	0,5	0,5	3	3,7	0,1	0,1	X	X

(b) $H_2O - NH_4Cl$, 2,5% - températures

X	$\varepsilon(T_E)$		$\varepsilon(T_M)$		$\varepsilon(T_{M,w})$	
	exp %	th %	exp %	th %	exp %	th %
ρ	4×10^{-6}	7×10^{-6}	9×10^{-5}	$6,6 \times 10^{-5}$	$0,83 \times 10^{-3}$	1×10^{-3}
c_S	7×10^{-7}	6×10^{-7}	2×10^{-5}	$2,5 \times 10^{-5}$	$4,2 \times 10^{-4}$	4×10^{-4}
c_L	$2,7 \times 10^{-6}$	2×10^{-6}	$1,6 \times 10^{-5}$	2×10^{-5}	$3,2 \times 10^{-4}$	3×10^{-4}
λ_S^{1D}	4×10^{-4}	$3,7 \times 10^{-4}$	$9,5 \times 10^{-4}$	1×10^{-3}	$2,4 \times 10^{-3}$	3×10^{-3}
λ_L^{1D}	4×10^{-4}	$2,6 \times 10^{-4}$	2×10^{-3}	3×10^{-3}	0,01	0,01
$T_{M,w}$	3×10^{-4}	$2,3 \times 10^{-4}$	2×10^{-3}	1×10^{-3}	X	X
$L_{M,w}$	4×10^{-6}	2×10^{-6}	4×10^{-5}	$4,4 \times 10^{-5}$	$5,3 \times 10^{-4}$	5×10^{-4}
T_E	X	X	0,1	0,16	0,4	0,35
$\widetilde{L}_E$	3×10^{-6}	4×10^{-6}	$1,8 \times 10^{-5}$	2×10^{-5}	4×10^{-4}	$3,5 \times 10^{-4}$
T_M	$2,1 \times 10^{-3}$	2×10^{-3}	X	X	0,02	0,02
α	2×10^{-3}	$1,8 \times 10^{-3}$	6×10^{-3}	6×10^{-3}	$1,2 \times 10^{-2}$	1×10^{-2}

(c) $H_2O - NH_4Cl$, 2,5% - propriétés énergétiques

X	$\varepsilon(c_S)$		$\varepsilon(c_L)$		$\varepsilon(L_{M,w})$		$\varepsilon(\widetilde{L}_E)$	
	exp %	th %	exp %	th %	exp %	th %	exp %	th %
ρ	3,16	3	3,14	3	3,06	3	2,89	3
c_S	X	X	0,94	1	1	1	0,9	1
c_L	1,01	1	X	X	1,04	1	0,95	1
λ_S^{1D}	0,37	0,3	0,34	0,3	0,41	0,45	0,36	0,3
λ_L^{1D}	0,8	0,7	1	0,8	1,3	1,35	0,7	0,6
$T_{M,w}$	0,5	0,6	0,5	0,6	0,8	0,8	0,6	0,6
$L_{M,w}$	0,97	1	0,93	1	X	X	1,1	1
T_E	1,5	1,6	5	5	5	5,1	9	9,2
$\widetilde{L}_E$	1,06	1	1,04	1	1	1	X	X
T_M	1	0,9	0,7	0,7	0,2	0,16	0,8	0,75
α	0,4	0,5	0,4	0,4	1	0,9	0,3	0,4

4.3.2.2 Erreurs stochastiques

Ce sont les erreurs dues à la mesure. On considère qu'elles sont additives, de moyenne nulle et qu'elles appartiennent à une distribution qui suit une loi normale $N(0, \sigma^2)$ de moyenne nulle et de variance constante σ^2 .

Considérer que l'erreur est de moyenne nulle, c'est considérer que la mesure est non-biaisée, qu'elle ne présente pas d'erreur systématique.

On peut exprimer la matrice de covariance associée sous la forme [71] :

$$\text{cov}(\Delta \mathbf{p}) = \sigma^2 \cdot [{}^t \mathbb{S} \cdot \mathbb{S}]^{-1} \quad (4.10)$$

La diagonale principale de cette matrice est constituée par les variances correspondant à chaque paramètre.

$$\begin{pmatrix} \sigma_{p_1}^2 \\ \dots \\ \sigma_{p_{n_p}}^2 \end{pmatrix} = \text{diag}(\text{cov}(\Delta \mathbf{p})) = \text{diag}(\sigma^2 \cdot [{}^t \mathbb{S} \cdot \mathbb{S}]^{-1}) \quad (4.11)$$

Il y a une probabilité de 68 % que la vraie valeur soit dans un intervalle $\pm \sigma$, 95 % qu'elle soit dans un intervalle de largeur 2σ , 99 % qu'elle soit dans un intervalle de largeur $2,58\sigma$, ...

Nous allons poser :

$$\Delta_\sigma p_i = 2,58\sigma_{p_i} \quad (4.12)$$

Maintenant, comme pour l'erreur déterministe, nous allons comparer le calcul théorique avec des expériences simulées. Le bruit expérimental inhérent à l'appareillage et aux conditions opératoires peut-être modélisé en ajoutant au signal, ici le thermogramme, une erreur gaussienne de moyenne nulle et de variance σ_{bruit}^2 . Le livret de procédure standard du calorimètre nous informe que le bruit du signal est de $5 \mu W$ [32]. Vérifions sur un thermogramme expérimental cette valeur.

Prenons un thermogramme où nous avons ajusté le zéro et comparons le flux enregistré dans une zone de conduction pure (voir figure 4.18) où la théorie [35] dit que le flux est égal à une valeur constante, ϕ_{moyen} .

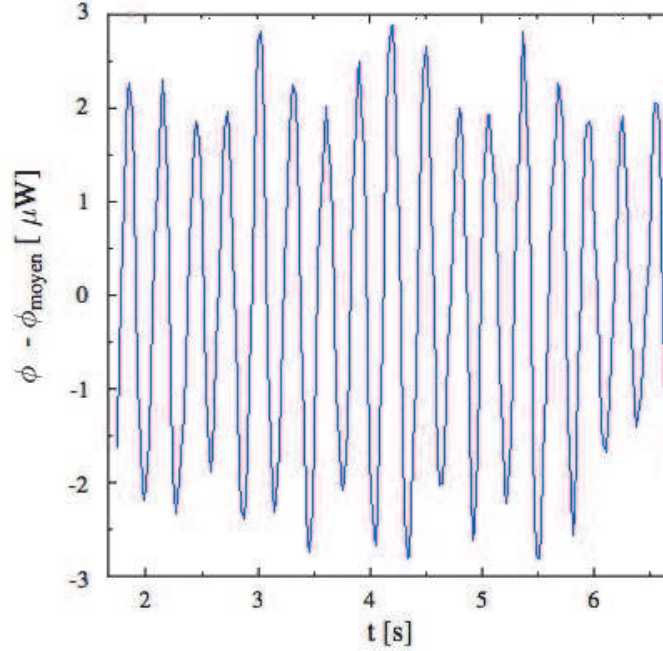


FIGURE 4.18 – Bruit expérimental du calorimètre

De ces mesures on obtient :

$$\sigma_{exp} = \sqrt{\frac{1}{n-1} \sum_{i=1}^n (\phi(t_i) - \phi_{moyen})^2} \approx 1,6 \mu W \quad (4.13)$$

Cette valeur est du même ordre de grandeur et inférieure à celle du livret technique [32] qui donne la valeur de $5 \mu W$, que nous pourrions retenir.

Nous allons donc bruiteur un thermogramme numérique de référence avec cet écart-type en ajoutant une distribution gaussienne de moyenne nulle et d'écart-type σ_{bruit} à cette référence. Puis identifier par inversion sur le thermogramme bruité les paramètres qui permettent de générer le thermogramme qui lui correspond le mieux, afin de les comparer à ceux qui ont servi pour générer le thermogramme numérique. Cette opération sera réalisée 200 fois afin de pouvoir mener une analyse statistique des erreurs dues au bruit sur les paramètres identifiés, pour déterminer les σ_{p_i} .

Comparons ces résultats expérimentaux avec nos prévisions du calcul théorique.

Tableau 4.7 – Influence théorique d’une erreur stochastique de $5 \mu W$ sur le résultat de l’inversion

(a) eau								
X	$\varepsilon(\rho)$ %	$\varepsilon(c_S)$ %	$\varepsilon(c_L)$ %	$\varepsilon(\lambda_S^{1D})$ %	$\varepsilon(\lambda_L^{1D})$ %	$\varepsilon(T_{M,w})$ %	$\varepsilon(L_{M,w})$ %	$\varepsilon(\alpha)$ %
th	1	1	1	2,7	0,05	1×10^{-4}	1	0,03
exp	0,8	1,1	0,9	2,4	0,03	$0,9 \times 10^{-4}$	1	0,06

(b) $H_2O - NH_4Cl$: 2,5% - propriétés énergétiques et températures							
X	$\varepsilon(c_S)$ %	$\varepsilon(c_L)$ %	$\varepsilon(T_{M,w})$ %	$\varepsilon(L_{M,w})$ %	$\varepsilon(T_E)$ %	$\varepsilon(\tilde{L}_E)$ %	$\varepsilon(T_M)$ %
th	1,6	1,6	2×10^{-3}	1,6	$2,5 \times 10^{-4}$	1,6	$7,6 \times 10^{-4}$
exp	1,4	1,5	3×10^{-3}	1,4	$1,1 \times 10^{-4}$	1,4	9×10^{-4}

(c) $H_2O - NH_4Cl$: 2,5% - propriétés pour le modèle				
X	$\varepsilon(\rho)$ %	$\varepsilon(\lambda_S^{1D})$ %	$\varepsilon(\lambda_L^{1D})$ %	$\varepsilon(\alpha)$ %
th	1,6	0,48	9×10^{-2}	0,12
exp	1,7	0,39	$9,7 \times 10^{-2}$	0,1

Pour l’erreur déterministe et l’erreur stochastique, nos prévisions théoriques concordent avec les expériences de simulations. Nous pouvons donc donner un intervalle de confiance, (4.3), pour nos paramètres identifiés en nous basant sur les prévisions du calcul théorique.

4.3.3 Choix des paramètres d’entrée et intervalle de confiance

L’étude de sensibilité, section 4.1, nous a renseigné sur la possibilité d’identifier les différents paramètres du modèle 1D. C’est-à-dire que les paramètres peuvent être identifiés simultanément dès lors qu’on considère la masse volumique comme un paramètre d’entrée.

Pour ce travail nous allons identifier tous les paramètres énergétiques pour comparer les résultats de la méthode à la littérature. Mais il est intéressant de noter qu’une erreur déterministe sur les température $T_{M,w}$, T_E et T_M engendre d’importantes erreurs sur le résultat de l’inversion, tableau 4.6(b). Il n’est donc pas intéressant de considérer une de ces températures comme paramètre d’entrée, car une petite erreur dans la littérature disponible ou une erreur de l’étalonnage en température du calorimètre fausserait notablement le résultat de l’identification.

Du fait que l’on utilise un modèle réduit, on est obligé d’identifier les propriétés de transferts car comme elles sont modifiées par la réduction de modèle, on ne peut prévoir cette modification au préalable.

On considère donc qu’un seul paramètre d’entrée, la masse volumique ρ . A partir des erreurs déterministes et stochastique des tableaux 4.5, 4.6 et 4.7, le domaine de confiance sur le

paramètres identifié i sera donc d'après l'équation (4.3) :

$$\Delta_c p_i = \sqrt{(\Delta_\rho p_i)^2 + (\Delta_\sigma p_i)^2 + (\Delta_T p_i)^2} \quad (4.14)$$

Tableau 4.8 – Intervalle de confiance sur les paramètres identifiés

(a) eau							
X	$\varepsilon(c_S)$ %	$\varepsilon(c_L)$ %	$\varepsilon(\lambda_S^{1D})$ %	$\varepsilon(\lambda_L^{1D})$ %	$\varepsilon(T_{M,w})$ %	$\varepsilon(L_{M,w})$ %	$\varepsilon(\alpha)$ %
$\Delta_c p_i$	3,16	3,16	2,7	0,086	$3,7 \times 10^{-2}$	3,16	0,05

(b) $H_2O - NH_4Cl$, 2.5% - propriétés énergétiques							
X	$\varepsilon(c_S)$ %	$\varepsilon(c_L)$ %	$\varepsilon(T_{M,w})$ %	$\varepsilon(L_{M,w})$ %	$\varepsilon(T_E)$ %	$\varepsilon(\widetilde{L}_E)$ %	$\varepsilon(T_M)$ %
$\Delta_c p_i$	3,4	3,4	$3,7 \times 10^{-2}$	3,4	$3,9 \times 10^{-2}$	3,4	$3,7 \times 10^{-2}$

(c) $H_2O - NH_4Cl$, 2.5% - propriétés pour le modèle			
X	$\varepsilon(\lambda_S^{1D})$ %	$\varepsilon(\lambda_L^{1D})$ %	$\varepsilon(\alpha)$ %
$\Delta_c p_i$	0,5	0,09	0,12

Les incertitudes sur les températures dues à l'inversion sont faibles devant l'incertitude de mesure des températures du calorimètre $\Delta_T p_i$. C'est pourquoi les intervalles de confiance des températures que nous donnons par la suite sont sensiblement égaux à cette incertitude de mesure.

4.3.4 Erreur sur l'enthalpie identifiée : analyse type « Monte Carlo »

Nous savons maintenant reconstruire le thermogramme et identifier les paramètres du modèle thermodynamique associé à l'échantillon, avec un intervalle de confiance. Maintenant il nous faut reconstruire l'enthalpie de l'échantillon à partir des paramètres que nous avons identifiés. Chacun de ces paramètres disposent d'un intervalle de confiance. Nous allons générer 100 courbes d'enthalpie avec les valeurs des paramètres prises aléatoirement dans leur intervalle de confiance, afin d'obtenir une distribution normale des courbes d'enthalpie. Ainsi on obtient un vecteur écart-type pour l'enthalpie que l'on identifie σ_h associé à une enthalpie moyenne h_{moyen} . Les résultats de cette analyse seront présentés en même temps que les résultats d'inversion, figures 4.19, 4.20, 4.22, 4.23, 4.24, 4.25, 4.29(a), 4.29(b) et 4.29(c).

4.4 Application à des cas expérimentaux

Nous avons vu qu'avec une méthode inverse on pouvait retrouver les paramètres du modèle thermodynamique qui ont servi à simuler un thermogramme et on connaît un intervalle de confiance pour ces résultats de l'inversion. A présent, nous allons vérifier que si on applique

ces mêmes méthodes à une expérience réelle, on retrouve les propriétés énergétiques de la littérature dans cet intervalle de confiance, en supposant que la valeur que nous avons prise dans la littérature soit la bonne et qu'il n'y a pas d'erreur d'étalonnage. Nous allons aussi présenter les enthalpies que nous avons reconstituées à partir des paramètres que nous avons identifiés. Parmi ces résultats nous ne présenterons pas les valeurs des capacités thermiques mais leur différence car leurs valeurs dépendent de la ligne de base, si le zéro a bien été calé par rapport au zéro du modèle numérique. Les expériences qui ont servi de base à ces identifications ont été réalisées avec un appareil DSC Pyris Diamond de Perkin-Elmer. Rappelons que cet appareil dispose de cellules hermétiques en aluminium de rayon $R = 2,15$ mm et de hauteur $Z = 0.78$ mm, suffisantes pour contenir des échantillons de volume avoisinant les $10 \mu\text{L}$. Les masses des échantillons ont été mesurées à l'aide d'une balance AP250D Ohaus, qui est précise à $0,1 \text{ mg}$.

4.4.1 Résultats pour un corps pur

D'abord nous avons travaillé sur le cas le plus connu, celui du corps pur dont nous avons développé le modèle dans le chapitre 2.1.2. Nous avons commencé cette étude par un échantillon d'eau de masse $m = 11,4 \text{ mg}$. Le processus d'identification a permis de reconstruire correctement le thermogramme expérimental, identifiant par là même les différents paramètres du modèle thermodynamique du corps pur. Par la suite nous avons testé d'autres corps, dont un échantillon d'acide benzoïque de masse $m = 3,7 \text{ mg}$. Les thermogrammes reconstitués pour ces deux corps à différentes vitesses sont représentés dans les figures 4.19 et 4.20.

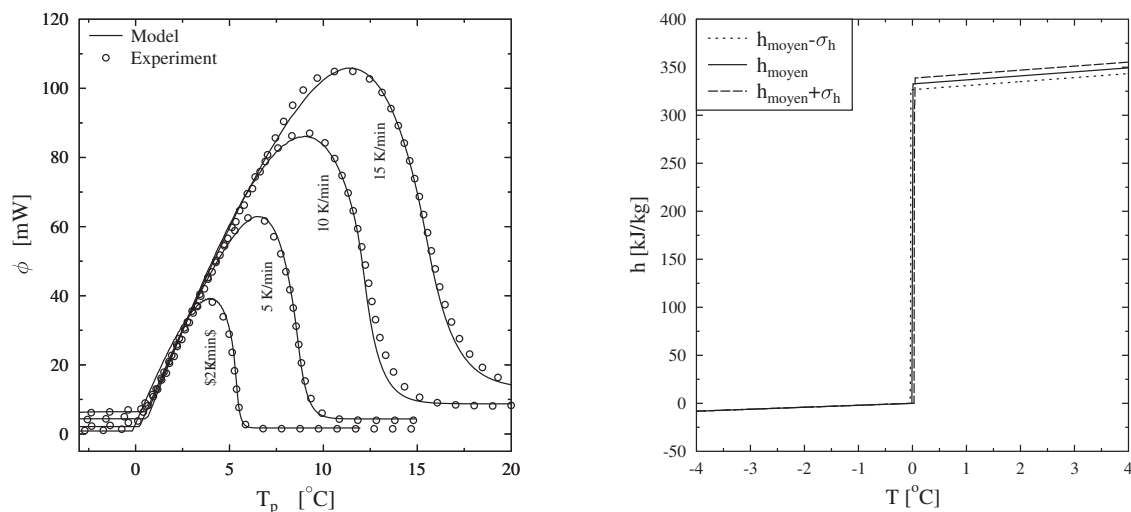


FIGURE 4.19 – Comparaison thermogramme expérimental/numérique et enthalpie reconstitués de l'eau

L'analyse statistique de l'enthalpie de l'eau, a été réalisée à partir des valeurs moyennes identifiées sur les thermogrammes expérimentaux réalisés pour plusieurs vitesses. Soit $T_{M,w} =$

0°C et $L_{M,w} = 330,75 \text{ kJ} \cdot \text{K}^{-1}$.

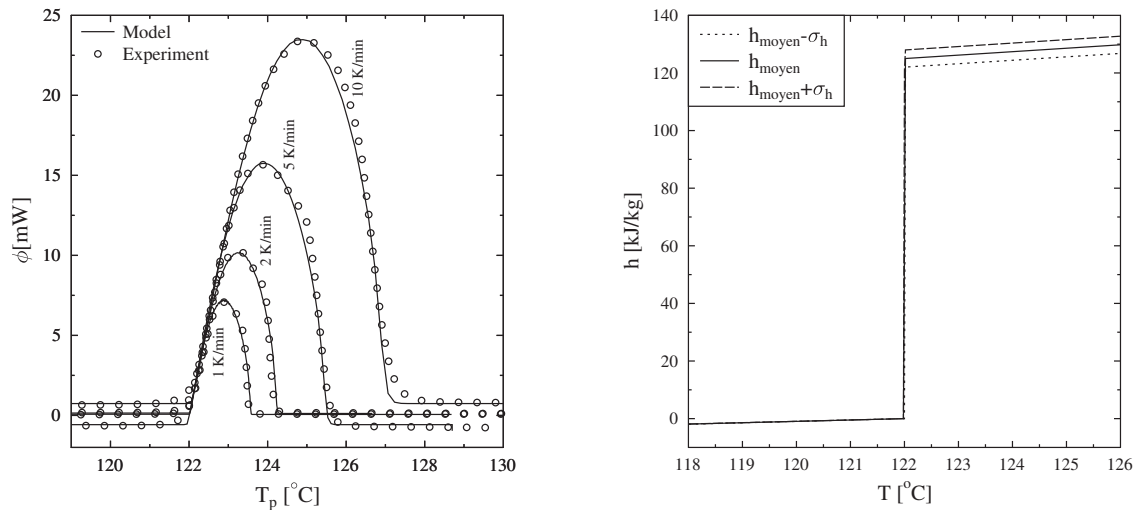


FIGURE 4.20 – Comparaison thermogramme expérimental/numérique et enthalpie reconstituée de l'acide benzoïque

L'analyse statistique de l'enthalpie de l'acide benzoïque, a été réalisée à partir des valeurs moyennes identifiées sur les thermogrammes expérimentaux réalisés pour plusieurs vitesses. Soit $T_{M,w} = 122,03^\circ\text{C}$ et $L_{M,w} = 123,75 \text{ kJ} \cdot \text{K}^{-1}$.

Les valeurs des paramètres identifiés pour les deux corps sont disponibles dans le tableau 4.9.

On présente également les valeurs déterminées par une approche directe pour comparaison.

Tableau 4.9 – Paramètres identifiés sur des thermogrammes expérimentales de corps pur et valeurs déterminées par mesures directes. Valeurs de la littérature [63] [91]

Material	β $\text{K} \cdot \text{min}^{-1}$	T_M $^\circ\text{C}$	$(T_M)_{\text{littérature}}$ $^\circ\text{C}$	L_M $\text{kJ} \cdot \text{kg}^{-1}$	$(L_M)_{\text{littérature}}$ $\text{kJ} \cdot \text{kg}^{-1}$	T_{onset} $^\circ\text{C}$	$S_{\text{surface pic}}$ $\text{kJ} \cdot \text{kg}^{-1}$
eau	2	$0 \pm 0,1$	0,02	330 ± 10	333	0,08	339
	5	$0,1 \pm 0,1$	0	325 ± 10	333	0,2	341
	10	$0,14 \pm 0,1$	0	335 ± 10	333	0,39	337
	15	$0,1 \pm 0,1$	0	338 ± 10	333	0,58	339
acide benzoïque	1	$122,2 \pm 0,1$	122	133 ± 5	132	122,3	130
	2	$122,0 \pm 0,1$	122	130 ± 5	132	122,3	130
	5	$121,8 \pm 0,1$	122	129 ± 5	132	122,3	129,9
	10	$122,1 \pm 0,1$	122	127 ± 5	132	122,3	129,8

Pour l'échantillon d'eau, l'étalonnage de l'appareil de DSC a été réalisé à une seule vitesse, $2 \text{ K} \cdot \text{min}^{-1}$. On a considéré pour T_{onset} , la définition disant qu'il s'agit de l'abscisse de l'intersection entre la pente à l'origine du pic et le plateau de conduction solide, voir figure 4.21(a).

On voit sur cette figure que, étalonner à une seule vitesse, conduit à une dépendance par rapport à la vitesse de T_{onset} . Par contre on remarque sur cette figure, qu'une autre définition de T_{onset} , l'intersection de la pente à l'origine avec l'axe des abscisses donne une valeur de T_{onset} , indépendante de la vitesse et égale à T_M du corps pur. On remarque sur cette figure que les débuts de pic de fusion ne sont pas des brusques ruptures de pente comme on pouvait s'y attendre pour un corps pur, mais que les thermogrammes présentent des arrondis ce qui s'explique par la présence d'impuretés dans les échantillons. Cela peut se comparer à un thermogramme d'une solution binaire de très faible fraction massique bien plus faible que celle des figures 4.22 à 4.26.

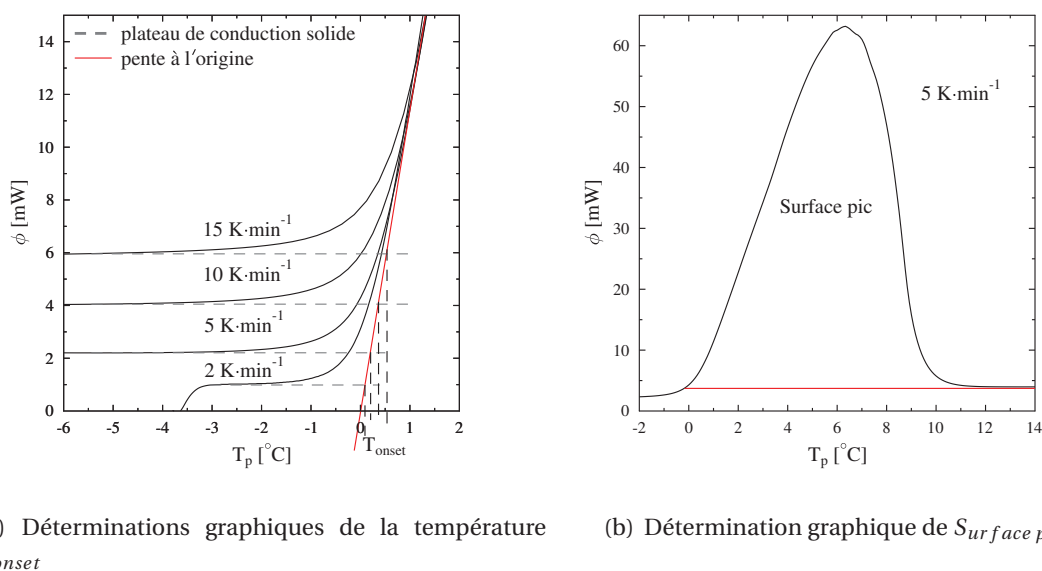


FIGURE 4.21 – Détermination par mesures directes

$S_{urface\ pic}$ est la surface exprimée en $kJ \cdot kg^{-1}$ comprise entre la courbe du thermogramme sur la durée du pic de fusion et la prolongation de la ligne de base liquide. Que l'on assimile à la chaleur latente du corps pur, voir figure 4.21(b) [33].

Dans ce tableau 4.9, l'étalonnage pour les expériences relatives à l'acide benzoïque a été réalisé à chaque vitesses. Dans ce cas on peut noter que T_{onset} est indépendante de la vitesse de réchauffement.

Dans ce tableau 4.1, on peut noter que les chaleurs latentes L_M identifiées pour l'eau et pour l'acide benzoïque, sont toutes contenues dans les intervalles de confiance pour les différentes vitesses. Et les valeurs de la littérature sont également comprises dans ces intervalles. On peut dire que l'on obtient par la méthode un résultat concordant qui est contenu dans l'intervalle de confiance, quelle que soit la vitesse de réchauffement comprise entre 1 et 15 $K \cdot min^{-1}$.

4.4.2 Résultats pour des solutions binaires

Les résultats que nous allons présenter correspondent à deux solutions binaires : $H_2O - NH_4Cl$ et $H_2O - KCl$.

4.4.2.1 solution de $H_2O - NH_4Cl$

Pour la première solution, nous avons testé quatre échantillons différents :

- un échantillon de 7 mg de fraction massique de soluté $x_0 = 0,480\%$
- un échantillon de 8,3 mg de fraction massique de soluté $x_0 = 5,00\%$
- un échantillon de 8,8 mg de fraction massique de soluté $x_0 = 9,02\%$
- un échantillon de 8,6 mg de fraction massique de soluté $x_0 = 10,00\%$

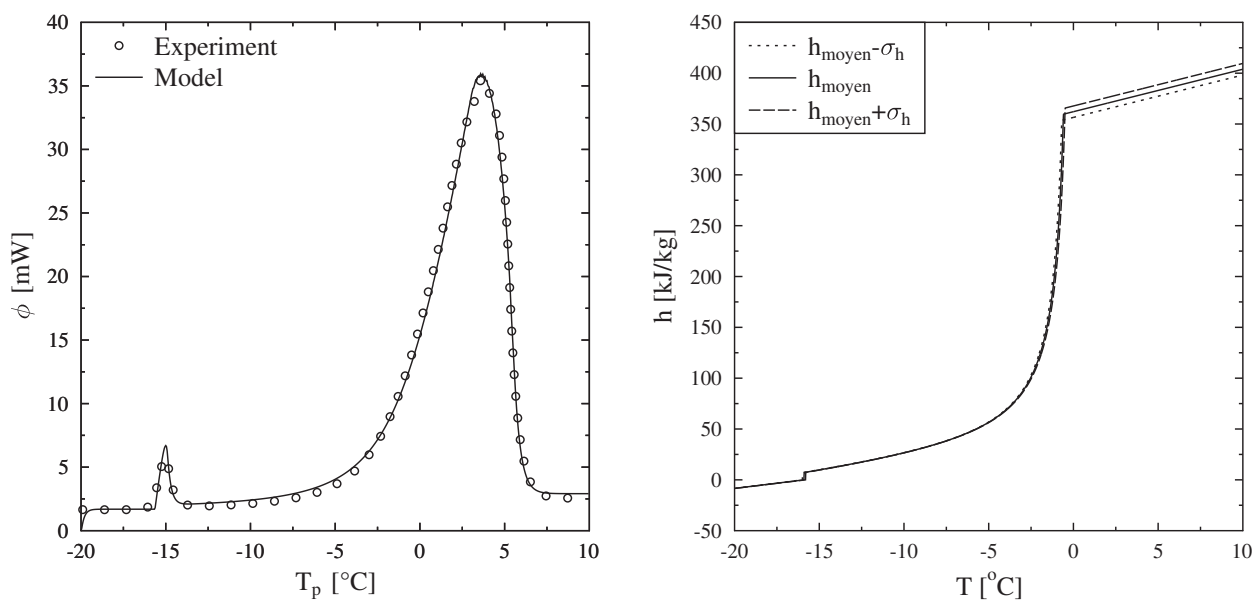
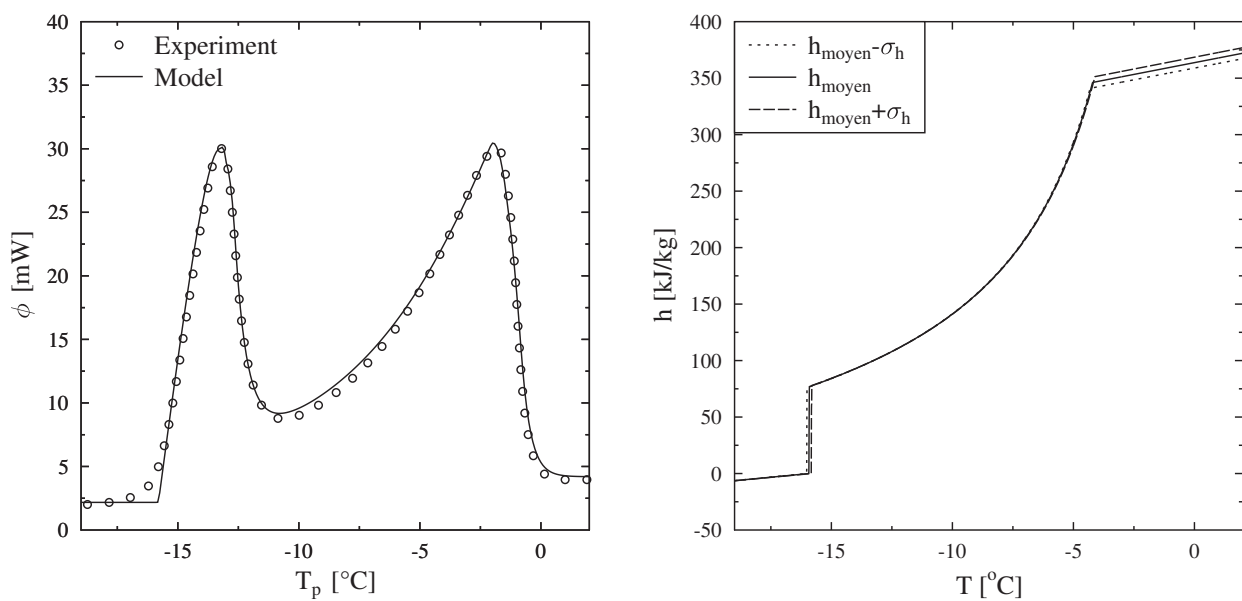
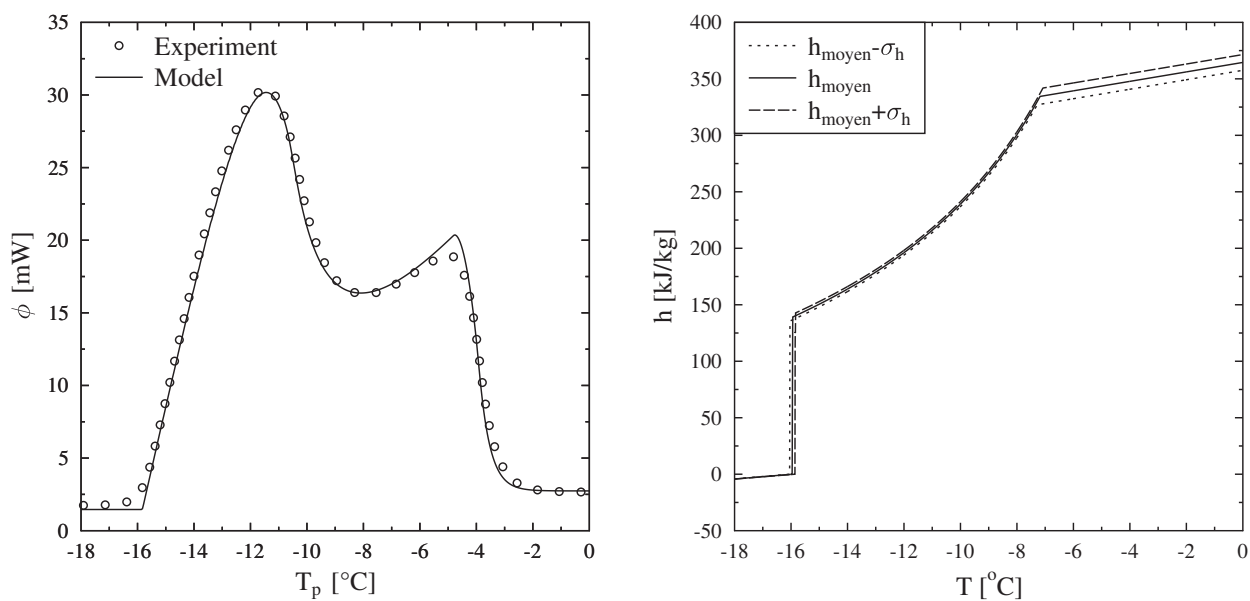
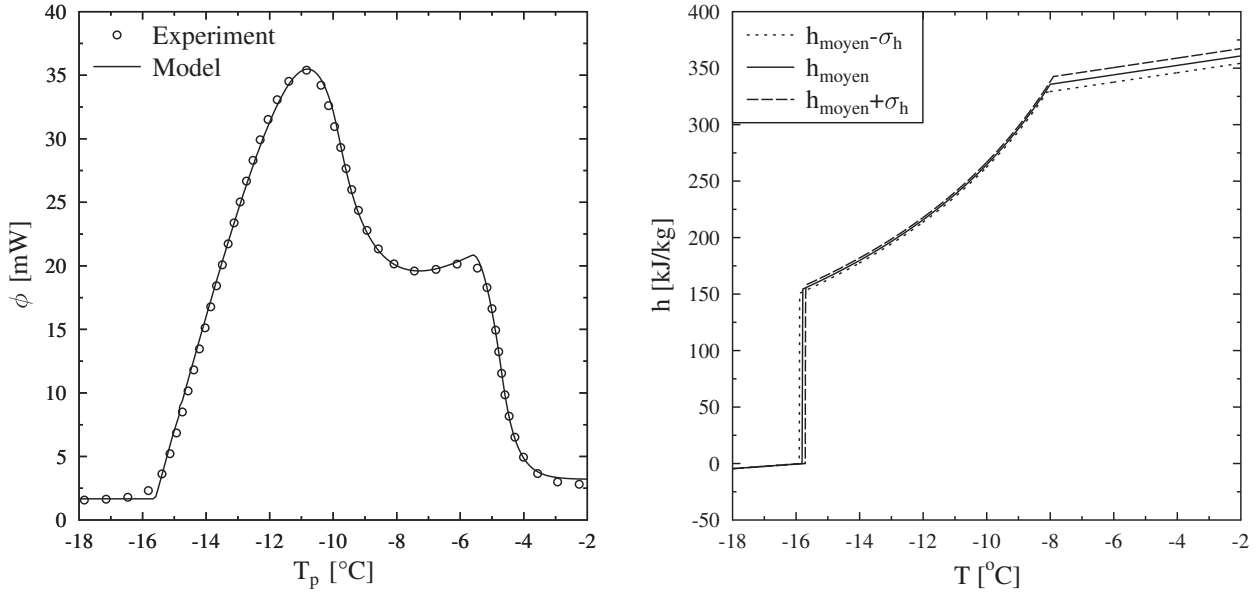


FIGURE 4.22 – $x_0 = 0,480\%$ et $\beta = 5 K \cdot min^{-1}$

FIGURE 4.23 – $x_0 = 5,00\%$ et $\beta = 5 \text{ K} \cdot \text{min}^{-1}$ FIGURE 4.24 – $x_0 = 9,02\%$ et $\beta = 5 \text{ K} \cdot \text{min}^{-1}$

FIGURE 4.25 – $x_0 = 10,00\%$ et $\beta = 5 \text{ K} \cdot \text{min}^{-1}$

Sur les figures 2.28, 2.30, 2.32, 2.34 et 2.36 et T_E , on ne distingue que l'influence des incertitudes sur les paramètres $\tilde{L}_E, L_{M,w}$. Ceci est dû d'abord à ce que les valeurs de ces deux paramètres sont élevées puisque leurs influences vont masquer celles des autres paramètres comme c_S et c_L . Ensuite, on peut remarquer dans la sous-section 2.3.2 que les paramètres $T_{M,w}$ et T_M ont des influences contraires sur l'enthalpie. L'influence des incertitudes sur ces deux températures vont se compenser.

Comme les intervalles de confiance sur les températures identifiées sont faibles (tableau 4.8) et que les variations d'enthalpies $\tilde{L}_E$ et $L_{M,w}$ sont élevées. On ne remarque sur les courbes d'enthalpies que l'influence des incertitudes sur $\tilde{L}_E$ et $L_{M,w}$. De plus on peut remarquer dans la sous-section 2.3.2 que les groupes de paramètres $T_{M,w}$, T_E et T_M , $\tilde{L}_E$, $L_{M,w}$ ont des influences qui se compensent sur l'enthalpie, voir figures 2.35, 2.29 et 2.33, 2.31, 2.37.

Les résultats des figures 4.22, 4.23, 4.24 et 4.25 correspondent à une vitesse de réchauffement de $5 \text{ K} \cdot \text{min}^{-1}$. Afin de montrer qu'ici aussi la vitesse de réchauffement n'a pas d'influence sur l'identification, que notre méthode est robuste vis-à-vis de ce paramètre, nous avons de nouveau effectué les tests à une vitesse de $2 \text{ K} \cdot \text{min}^{-1}$. Nous nous sommes contentés de les effectuer pour deux concentrations pour lesquelles les solutions ont des comportements thermiques très différents : $x_0 = 0,480\%$ et $x_0 = 9,02\%$. Les thermogrammes correspondants sont aux figures 4.26(a) 4.26(b).

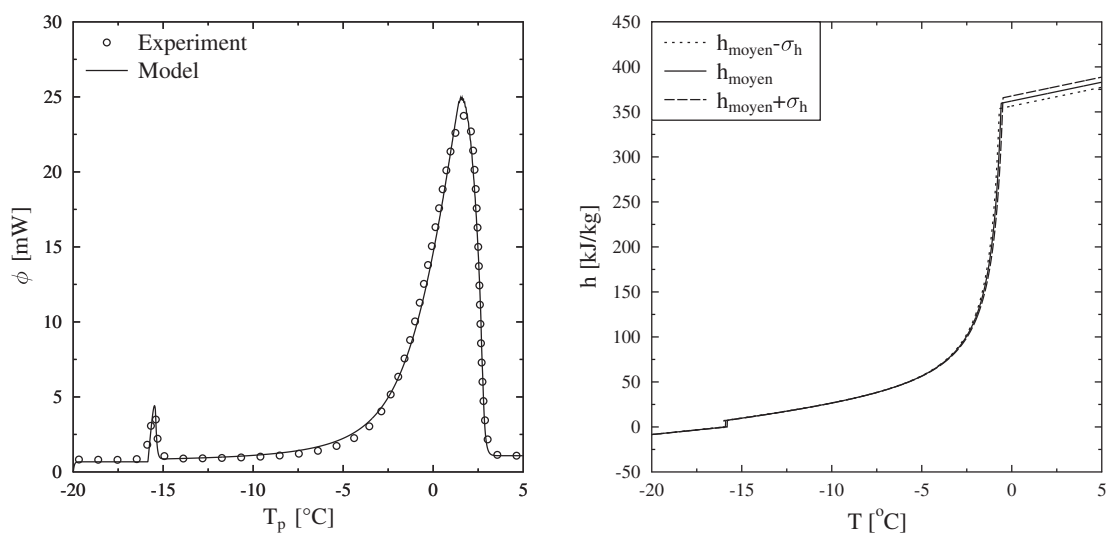
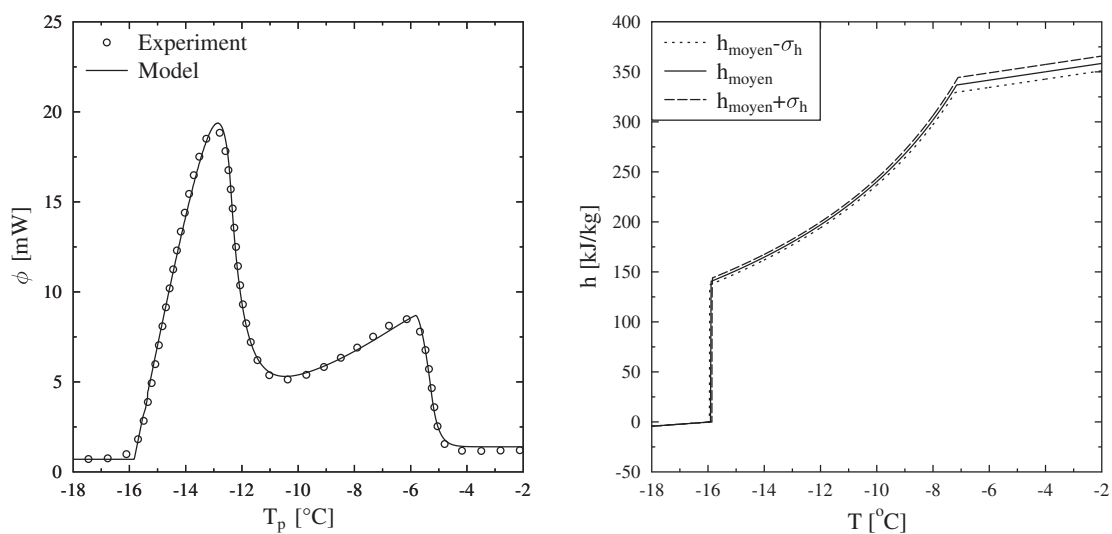
(a) $x_0 = 0,480\%$.(b) $x_0 = 9,02\%$.

FIGURE 4.26 – Comparaison des thermogrammes expérimentaux et numériques après identification (solution de $H_2O - NH_4Cl$) pour $\beta = 2\text{K} \cdot \text{min}^{-1}$.

Les valeurs des paramètres identifiés pour les solutions sont réparties dans les tableaux 4.10 et 4.11.

Tableau 4.10 – Paramètres identifiés sur des thermogrammes expérimentaux du binaire H_2O-NH_4Cl

(a) Paramètres du solvant [63] [57]

x_0 %	β $K \cdot \min^{-1}$	$c_L - c_S$ $J \cdot kg^{-1} \cdot K^{-1}$	$T_{M,w}$ $^{\circ}C$	$(T_{M,w})_{Littérature}$ $^{\circ}C$	$L_{M,w}$ $kJ \cdot kg^{-1}$	$(L_{M,w})_{Littérature}$ %
0.48	2	2071 ± 212	$0,0 \pm 0,1$	0	327 ± 11	333
	5	2004 ± 212	$0,01 \pm 0,1$	0	323 ± 11	333
5.00	5	2315 ± 230	$-0,07 \pm 0,1$	0	333 ± 12	333
9.02	2	2117 ± 280	$-0,39 \pm 0,08$	0	315 ± 20	333
	5	1843 ± 280	$-0,1 \pm 0,08$	0	323 ± 20	333
10.0	5	1740 ± 300	$-0,05 \pm 0,08$	0	332 ± 20	333

(b) Paramètres du soluté [92] [93]

x_0 %	β $K \cdot \min^{-1}$	T_E $^{\circ}C$	$(T_E)_{Littérature}$ $^{\circ}C$	$\widetilde{L}_E$ $kJ \cdot kg^{-1}$	T_M $^{\circ}C$
0,48	2	$-15,9 \pm 0,1$	-15,9	$7,2 \pm 0,2$	$-0,4 \pm 0,1$
	5	$-15,86 \pm 0,1$	-15,9	$7,39 \pm 0,2$	$-0,39 \pm 0,1$
5,00	5	$-15,91 \pm 0,1$	-15,9	$74,4 \pm 1,7$	$-4,1 \pm 0,1$
9,02	2	$-15,9 \pm 0,13$	-15,9	132 ± 6	$-7,18 \pm 0,1$
	5	$-15,95 \pm 0,13$	-15,9	128 ± 6	$-7,14 \pm 0,1$
10,0	5	$-15,8 \pm 0,13$	-15,9	144 ± 8	$-8,0 \pm 0,1$

Superposons les profils enthalpiques des solutions de H_2O-NH_4Cl à 0,48% et 9,02% qui ont été obtenus à $2\text{ K}\cdot\text{min}^{-1}$ et $5\text{ K}\cdot\text{min}^{-1}$ sur la figure 4.27. Encore une fois on remarque qu'il n'y a pas d'influence de la vitesse de réchauffement sur la détermination de l'enthalpie alors que les thermogrammes sont très différents. On note que la droite qui passe par les sommets des pics du solvant pour plusieurs vitesses coupe bien l'axe des abscisses à la température de liquidus T_M [94].

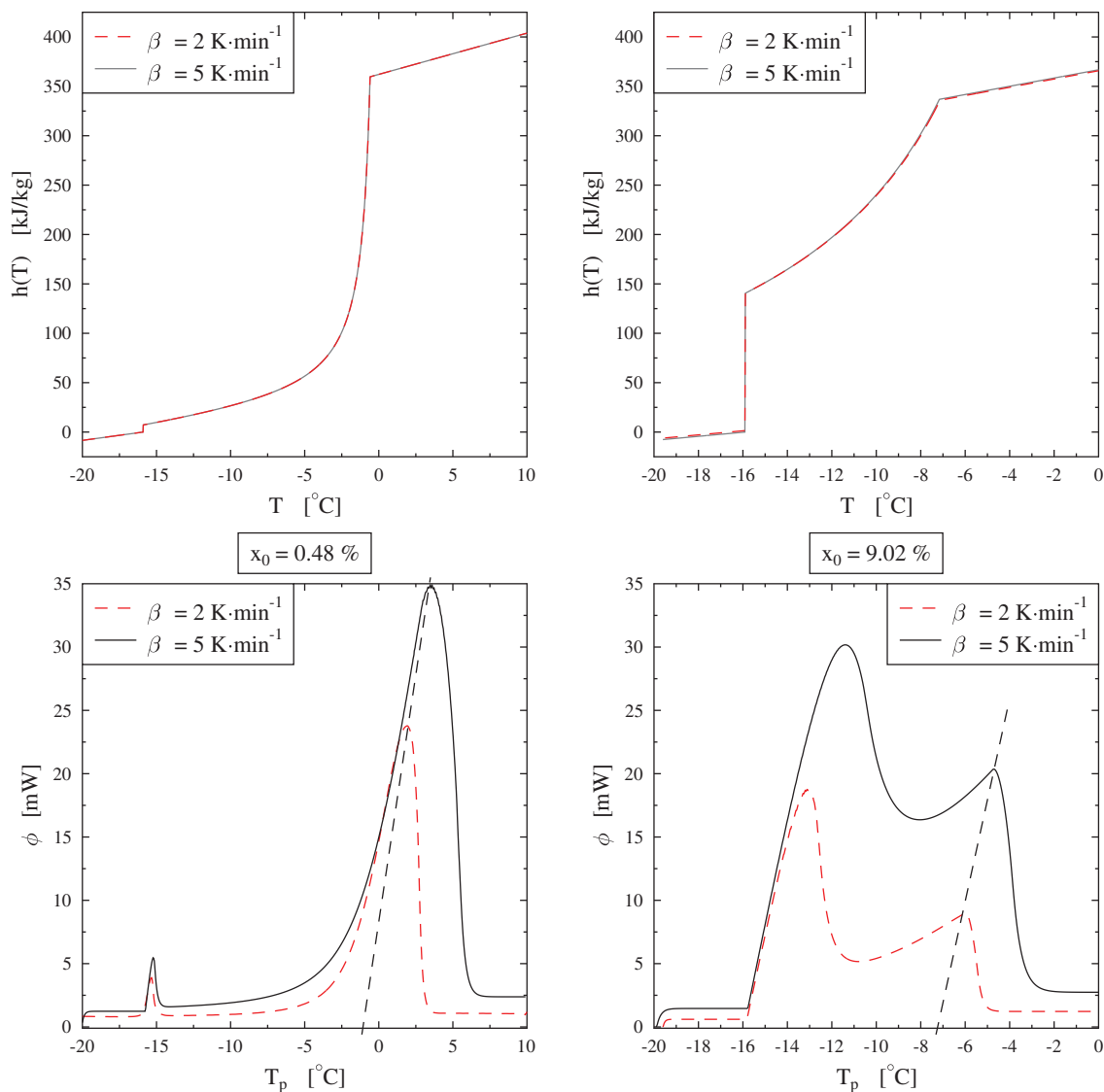


FIGURE 4.27 – Thermogrammes et enthalpies identifiés de solutions de H_2O-NH_4Cl pour deux vitesses

Remarquons que dans le cas d'une solution binaire, la température T_{onset} permet de déterminer la température eutectique. Mais définir une température T_{onset} sur le pic du solvant en prolongeant la pente droite de ce pic sur le thermogramme gradué en T_p , ne conduit pas à la température de liquidus T_M , voir figure 4.28.

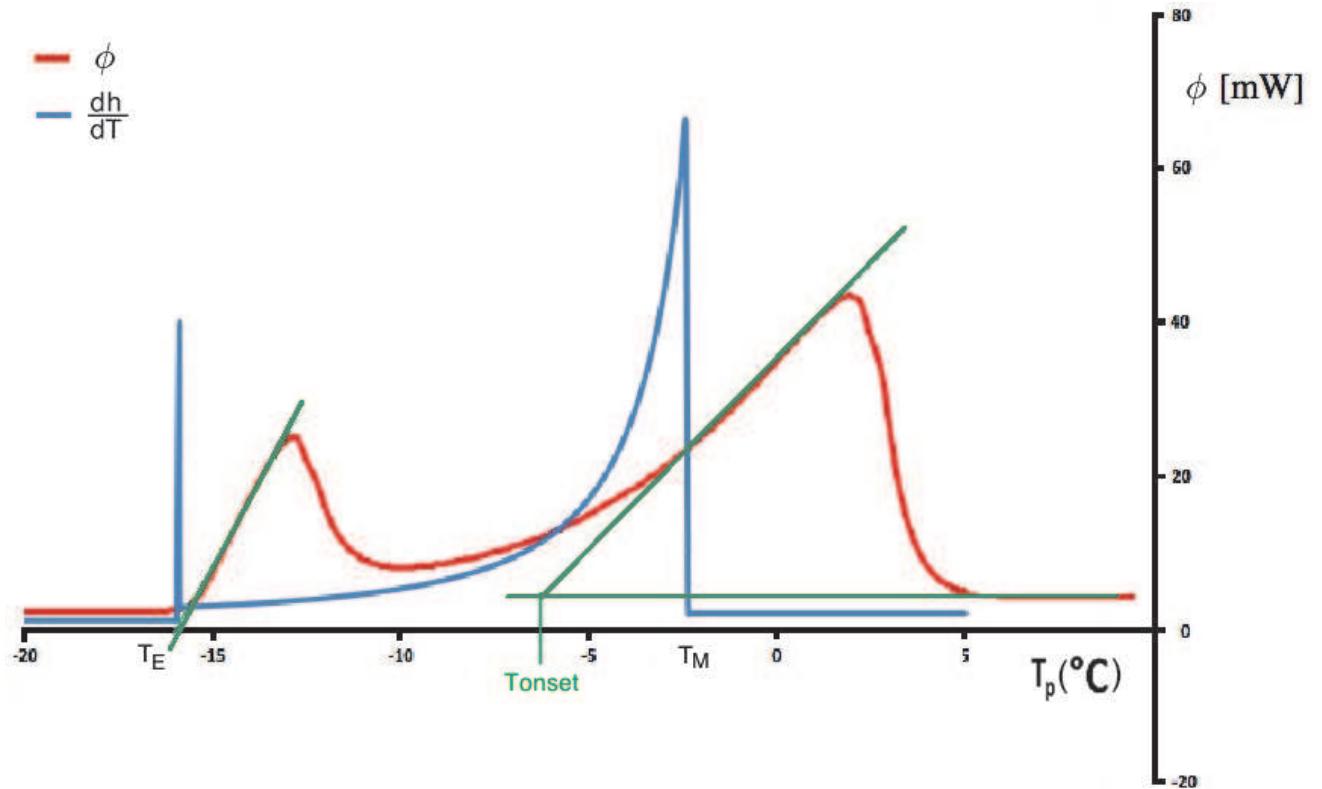


FIGURE 4.28 – T_{onset} de la solution binaire

4.4.2.2 Solution de $H_2O - KCl$

Quant à la deuxième solution, nous avons utilisé trois échantillons différents :

- un échantillon de 11,7 mg de fraction massique de soluté $x_0 = 0,615\%$
- un échantillon de 7,4 mg de fraction massique de soluté $x_0 = 1,130\%$
- un échantillon de 10,3 mg de fraction massique de soluté $x_0 = 2,007\%$

Les comparaisons des thermogrammes expérimentaux et numériques sont représentées dans la figure 4.29 ainsi que les enthalpies reconstituées. On observe toujours une excellente correspondance entre les thermogrammes expérimentaux et identifiés.

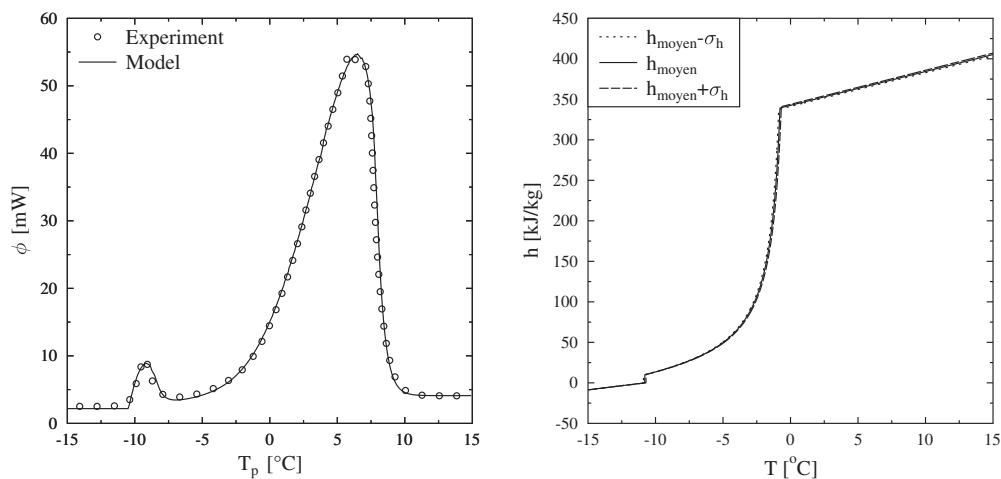
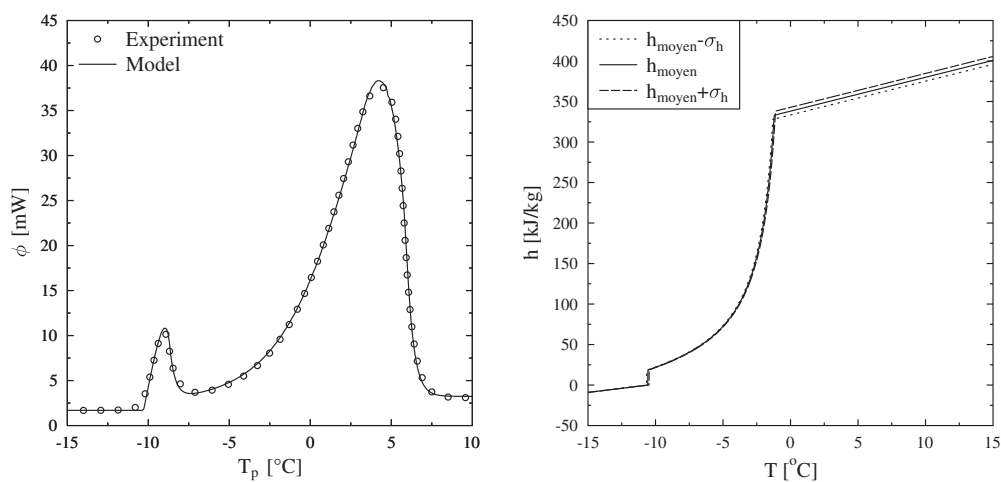
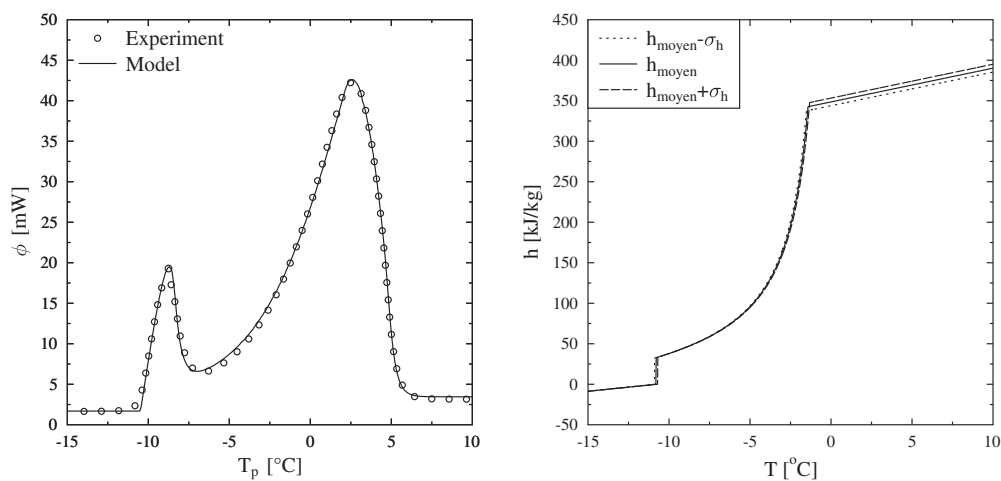
(a) $x_0 = 0,615\%$.(b) $x_0 = 1,130\%$.(c) $x_0 = 2,007\%$.

FIGURE 4.29 – Comparaison des thermogrammes expérimentaux et numériques après identification
(solution de $H_2O - KCl$)

Tableau 4.11 – Paramètres identifiés sur des thermogrammes expérimentaux de binaire H_2O-NH_4Cl

(a) Paramètres du solvant [63] [57]

x_0 %	β $K \cdot min^{-1}$	$c_L - c_S$ $J \cdot kg^{-1} \cdot K^{-1}$	$T_{M,w}$ $^{\circ}C$	$(T_{M,w})_{Littérature}$ $^{\circ}C$	$L_{M,w}$ $kJ \cdot kg^{-1}$	$(L_{M,w})_{Littérature}$ $J \cdot kg^{-1} \cdot K^{-1}$
0,615	5	1965 ± 211	$-0,3 \pm 0,1$	0	334 ± 11	333
1,13	5	2069 ± 211	$-0,8 \pm 0,1$	0	336 ± 14	333
2,007	5	2042 ± 212	$-0,8 \pm 0,1$	0	333 ± 4	333

(b) Paramètres du soluté [52]

x_0 %	β $K \cdot min^{-1}$	T_E $^{\circ}C$	$(T_E)_{Littérature}$ $^{\circ}C$	$\widetilde{L}_E$ $kJ \cdot kg^{-1}$	T_M $^{\circ}C$
0,615	5	$-10,87 \pm 0,1$	-10,6	$10,7 \pm 0,4$	$-0,63 \pm 0,11$
1,13	5	$-11,51 \pm 0,1$	-10,6	$16,1 \pm 0,6$	$-1,43 \pm 0,1$
2,007	5	$-11,45 \pm 0,1$	-10,6	$26,1 \pm 1$	$-1,92 \pm 0,1$

Dans les tableaux reftab :identBSKCl a et 4.11(b), on note que les valeurs identifiées pour chaque température caractéristique (T_E ou $T_{M,w}$) sont indépendantes de la composition, ce qui est attendu. On observe cependant un décalage systématique par rapport aux valeurs de la littérature.

Les résultats pour le binaire précédent ainsi que pour les corps purs étant satisfaisants, nous pensons ici qu'il s'agit d'un problème d'étalonnage en température lors de l'enregistrement du thermogramme. Selon cette hypothèse, les écarts entre les trois températures identifiées sont justes. Nous allons recalculer ces températures de sorte que la température eutectique prenne la valeur donnée dans la littérature. Les résultats corrigés montrent alors une température de fusion de l'eau pure conforme à la valeur attendue (aux incertitudes près), voir tableau 4.12.

Tableau 4.12 – Températures corrigées pour le binaire $H_2O - KCl$

x_0 %	β $K \cdot min^{-1}$	$T_{M,w}$ $^{\circ}C$	T_E $^{\circ}C$	T_M $^{\circ}C$
0,615	5	$-0,03 \pm 0,1$	-10,6	$-0,36 \pm 0,11$
1,13	5	$0,11 \pm 0,1$	-10,6	$-0,52 \pm 0,1$
2,007	5	$0,05 \pm 0,1$	-10,6	$-1,07 \pm 0,1$

4.4.3 Concurrence des résultats d'une même solution entre eux : détermination des chaleurs eutectiques L_E

On se propose maintenant de déterminer les chaleurs latentes eutectiques de ces deux solutions à partir de nos résultats d'identification et de la relation (2.50) de la sous-section 2.1.3

du chapitre traitant du modèle direct de la DSC, mais rappelons cette équation :

$$\widetilde{L}_E = \frac{x_0}{x_E} L_E \quad (4.15)$$

On obtient par régression linéaire des valeurs pour L_E de $283 \text{ kJ} \cdot \text{kg}^{-1}$ pour les solutions de $\text{H}_2\text{O} - \text{NH}_4\text{Cl}$ et $213 \text{ kJ} \cdot \text{kg}^{-1}$ pour les solutions de $\text{H}_2\text{O} - \text{KCl}$ d'après les régressions linéaires présentées aux figures 4.30(a) et 4.30(b). Nous n'avons pas trouvé ces valeurs dans la littérature pour les comparer, elles restent donc à confirmer.

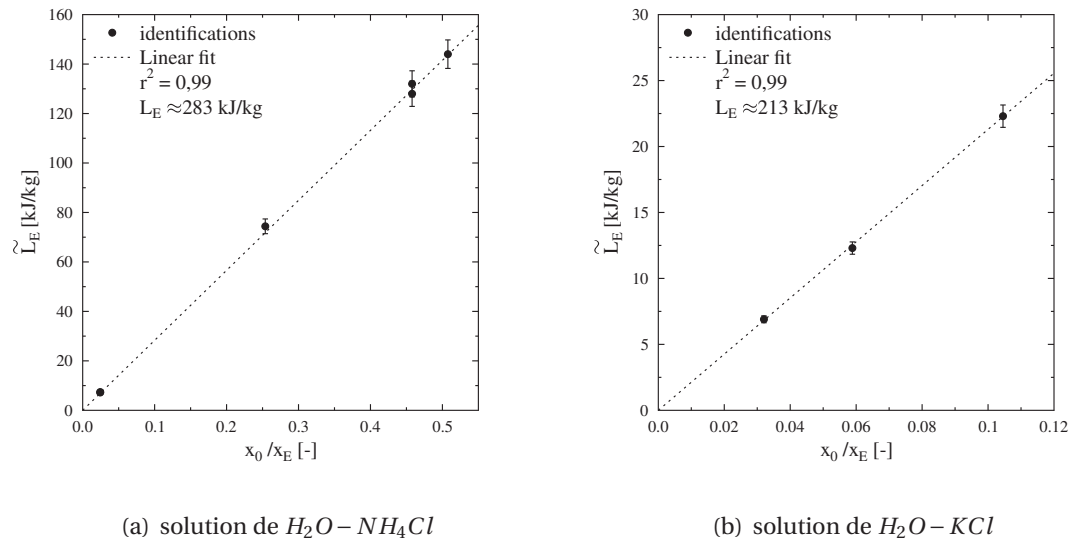


FIGURE 4.30 – Détermination des chaleurs latentes eutectiques L_E

Finalement, nous proposons de comparer les températures T_E et T_M identifiées en reconstituant la partie du diagramme de phase qui précède le point eutectique pour chaque solution binaire. Les résultats correspondants sont visibles sur les figures 4.31(a) et 4.31(b), ces figures ne présentent pas d'incertitude car celle-ci sont trop faibles pour être visibles.

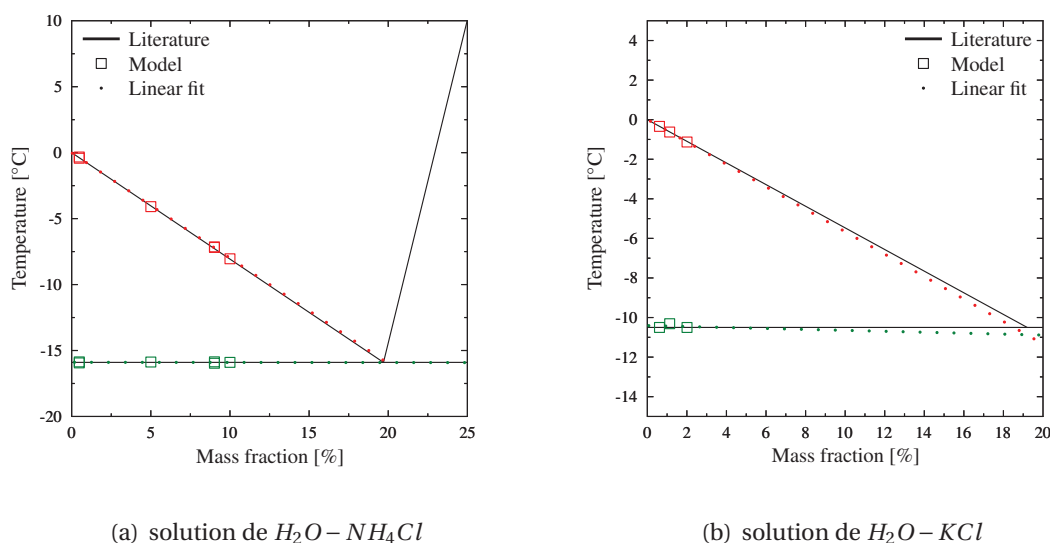


FIGURE 4.31 – Comparaison des paramètres identifiés avec leur valeurs exactes sur un diagramme de phase

4.4.4 Conclusion : incertitude sur l'enthalpie fonction de la température

Comme on peut le voir sur la figure 4.27 l'utilisation des méthodes inverses pour caractériser les matériaux à changement de phase conduit à des résultats qui sont indépendants de la vitesse de réchauffement du calorimètre dans une gamme comprise entre 2 et $15 \text{ K} \cdot \text{min}^{-1}$ et des échantillons d'une dizaine de mg. Une masse plus importante pourrait rendre non-négligeable le transport de matière par convection que l'on ne modélise pas dans cette étude. Des tests ont été effectués dans une gamme plus large entre 1 et $15 \text{ K} \cdot \text{min}^{-1}$ dans le cas du corps pur et on obtient des résultats identiques pour ces vitesses.

Le fait que l'on parvienne à reconstruire les diagrammes de phase de la figure 4.31 suppose que les températures caractéristiques aient été identifiées de manière correcte. De plus les valeurs des coefficients de régression linéaire de la figure 4.30 indiquent une concordance des résultats en énergie.

Enfin, nous avons reconstitué les enthalpies dans les figures 4.19, 4.20, 4.22, 4.23, 4.24, 4.25, 4.26 et 4.29 en considérant l'intervalle de confiance de chaque paramètre. Il apparaît sur ces figures que seules les incertitudes sur $L_{M,w}$ et $\widetilde{L}_E$ influent de manière notable sur l'enthalpie, celles sur les températures étant très faibles. Ceci est dû aux fortes valeurs de ces paramètres pour les produits utilisés et à leur prédominance sur le plan énergétique. D'ailleurs, l'influence de l'incertitude de L_E ne se remarque que pour les concentrations les plus élevées, quand $\widetilde{L}_E$ est du même ordre de grandeur que $L_{M,w}$. Ces relations enthalpie-température sont cependant identifiées avec des précisions meilleure que celles généralement admises pour les applications MCP [36], soit 10% sur l'enthalpie et 1 K sur la température.

Chapitre 5

Application à un mortier contenant des MCP

Dans les chapitres précédents, une approche permettant de caractériser les propriétés énergétiques des matériaux à changements de phases a été développée. Le projet ANR stock-E 2010 dans lequel s'inscrit cette thèse a pour cadre le domaine du bâtiment. La méthode qui a été développée va maintenant être appliquée à un cas pratique d'utilisation des MCP dans ce domaine, celui d'un mur contenant des inclusions de MCP. La DSC permet de faire de la caractérisation à l'échelle mésoscopique, mais pour une application dans l'industrie du bâtiment il serait intéressant de pouvoir caractériser à une échelle macroscopique le mélange constituant le mur. C'est pourquoi un autre dispositif expérimental est étudié.

Dans ce chapitre les méthodes inverses développées précédemment sont appliquées pour caractériser un échantillon de mortier contenant des MCP. Par manque de temps pour le séchage du mortier, nous n'avons pas réalisé notre propre expérience, le dispositif que nous avons construit n'est donc pas présenté. Nous avons donc utilisé les résultats expérimentaux d'une autre équipe du projet ANR stock-E 2010 issue du Laboratoire de Génie Civil et géo-Environnement (LGCgE) de l'université d'Artois à Béthune, qui ont accepté de nous les fournir. Nous allons d'abord brièvement présenter le dispositif expérimental de cette équipe, avant d'aborder le modèle direct que nous avons développé. Enfin, nous allons présenter les résultats de la caractérisation que nous avons obtenue par la méthode inverse et les comparer aux valeurs obtenues par mesures directes.

5.1 Dispositif expérimental du LGCgE [95] [96]

Le dispositif sert un échantillon de mortier de dimensions $250 \text{ mm} \times 250 \text{ mm} \times 40 \text{ mm}$ entre deux plaques échangeuses de dimensions $500 \text{ mm} \times 500 \text{ mm} \times 40 \text{ mm}$ posées sur leur tranche. Les échangeurs sont parcourus par une eau glycolée qui provient de deux bains thermo-régulés. Le dispositif formé par l'échantillon et les deux plaques est calorifugé de toutes parts par de l'isolant pour éviter les pertes thermiques, notamment latérales. Des schémas du dispositif expérimental sont donnés aux figures 5.3 et 5.4.

La forme aplatie, ainsi que le calorifugage permet de considérer le transfert thermique comme

unidimensionnel. Nous avons vérifié cette hypothèse en développant un modèle 3D de conduction, voir figure 5.1, afin de simuler le gradient interne de température.

Ce modèle est résolu par un schéma implicite où il n'y a pas de problème de divergence causé par un maillage trop important par rapport au pas de temps. Pour ce qui est des conditions aux limites, nous avons modélisé les contacts échantillons/échangeurs par des résistances, $R_{sup} = R_{inf} = RTC$, qui ont été évaluées pour des cas similaires au notre par plusieurs auteurs à $RTC = 10^{-4} \text{ K} \cdot \text{m}^2 \cdot \text{W}^{-1}$ [97] [98] [99] [100]. Les résistances thermiques latérales, R_{lat} , ont été évaluées à $20 \text{ K} \cdot \text{m}^2 \cdot \text{W}^{-1}$. Cette simulation qui utilise les propriétés que l'on a identifiées par la méthode inverse et qui seront présentées dans la section 5.5, permet de conforter l'hypothèse unidimensionnel du transfert thermique, voir figure 5.2. En effet on distingue une large zone au centre où les isothermes sont parallèles entre elles.

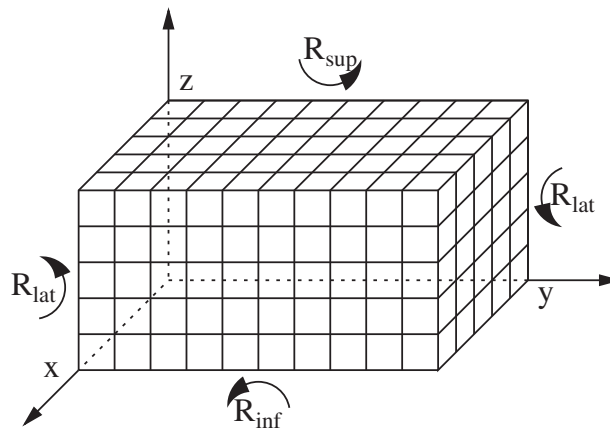


FIGURE 5.1 – Maillage 3D

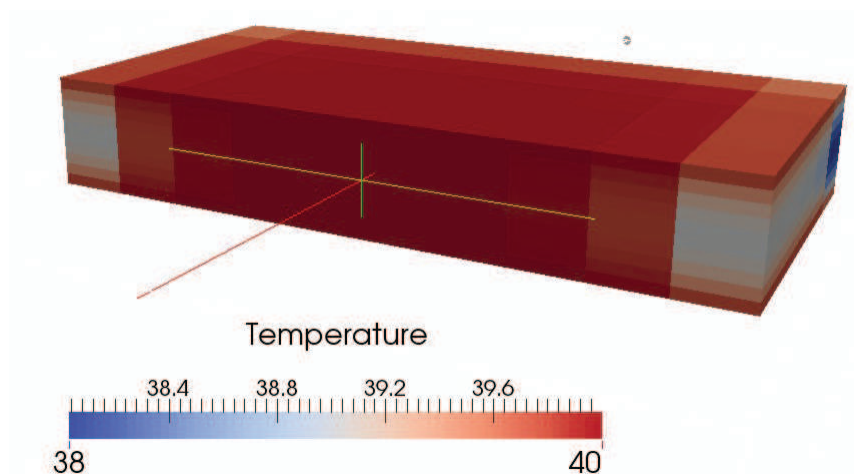


FIGURE 5.2 – Modélisation 3D de la conduction dans l'échantillon refroidi de manière symétrique - coupe interne

Avec ce dispositif, on mesure l'évolution du flux chaleur lors de cycles de réchauffements

et refroidissements pour évaluer le comportement thermique de l'échantillon disposés au centre des faces actives de l'échantillon. L'acquisition de ce flux se fait par l'intermédiaire de deux fluxmètres à gradient tangentiel d'épaisseur 0,2 mm, de surface $150 \times 150 \text{ mm}^2$ et de sensibilités $110 \mu\text{V} \cdot \text{W}^{-1} \cdot \text{m}^2$. La présence d'anneaux gardés autour des fluxmètres permet de considérer qu'il n'y a pas d'effet thermique transversal au niveau du fluxmètre et que le flux qui le traverse est unidirectionnel. Ces fluxmètres disposent également de thermocouples intégrés qui permettent de mesurer une température de surface de l'échantillon. Ces capteurs sont reliés à une centrale-multimètre multivoie Keithley 2700.

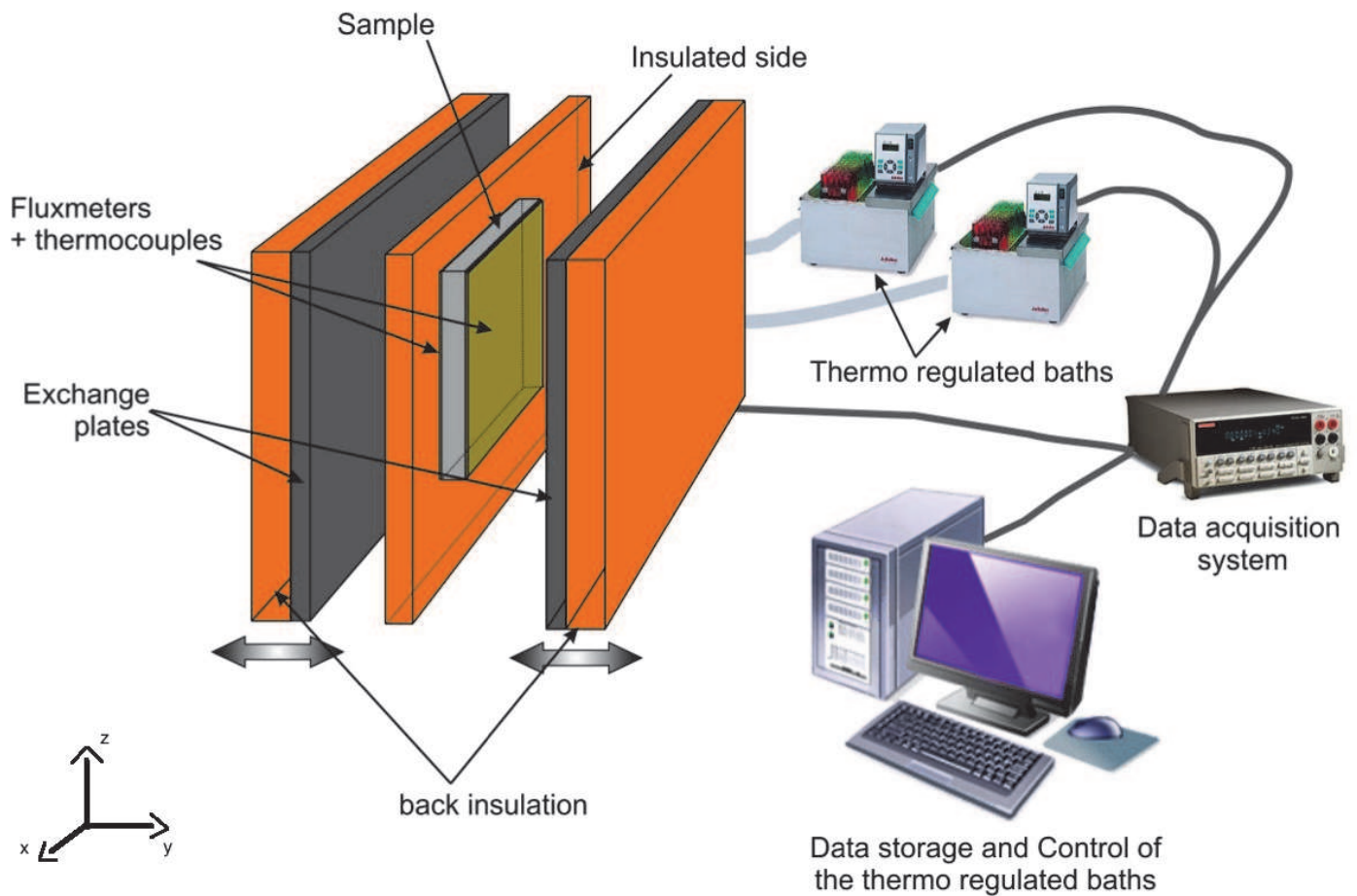


FIGURE 5.3 – Représentation globale du dispositif [101]

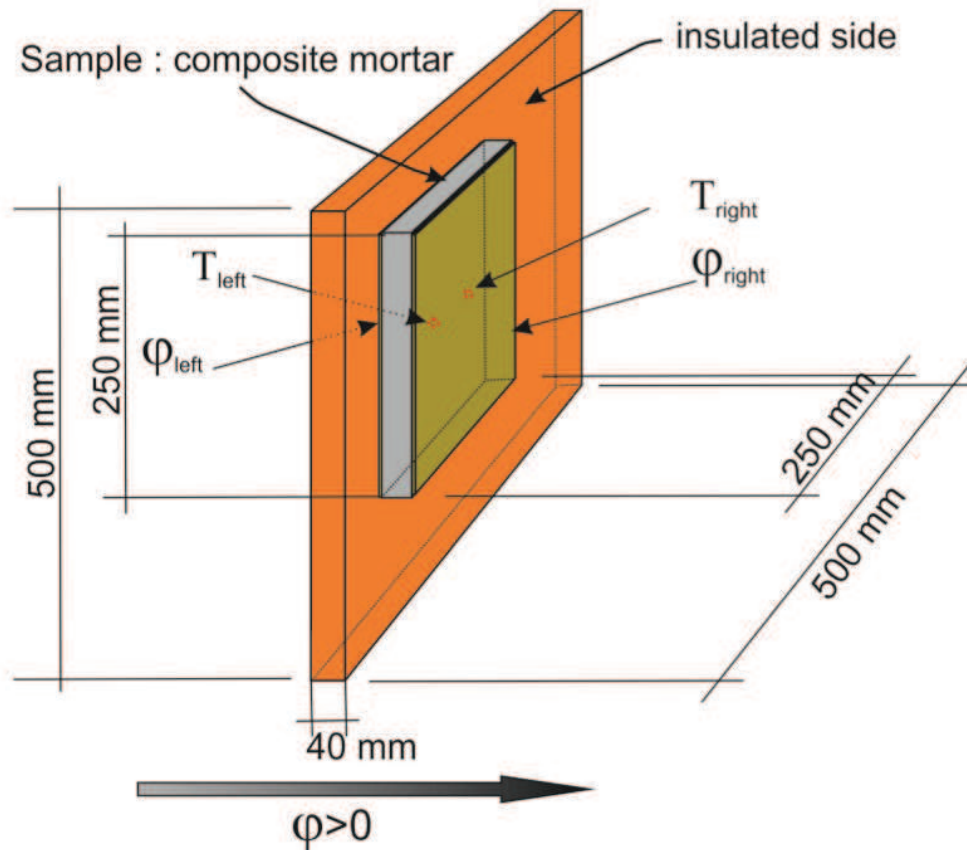


FIGURE 5.4 – Dimensions et instrumentation du dispositif [101]

Les températures T_{left} et T_{right} sont égales à la température de consigne T_p .

5.2 Résultats expérimentaux du LGCE [101]

Le mortier contient un MCP micro-encapsulé de l'industriel BASF (Micronal PCM ®DS 5001 X). Le mélange entre les différents matériaux a été fait en plusieurs étapes afin d'homogénéiser le mélange avec l'humidité ambiante et afin de soumettre les micro-capsules à un minimum de stress mécanique. La préparation de l'échantillon a été suivie d'une période de séchage d'environ deux mois. Les proportions de mélange eau/sable/ciment du mortier sont disponibles dans le tableau 5.1.

ratio massique entre le ciment et le sable	ratio entre le ciment et l'eau	ration MCP et (ciment + sable)
1 : 2,6	1 : 1,03	1 : 6,6

Tableau 5.1 – Ratio massique de mélange du mortier [96] [101]

L'échantillon de mortier a une masse de 3,12 kg et comporte 12,4% de MCP en masse. En considérant les dimensions de l'échantillon, la masse volumique est de $1248 \text{ kg} \cdot \text{m}^{-3}$. Dans

les figures 5.5, 5.6 et 5.7 nous présentons les courbes de flux expérimentales qui ont été obtenues avec ce dispositif.

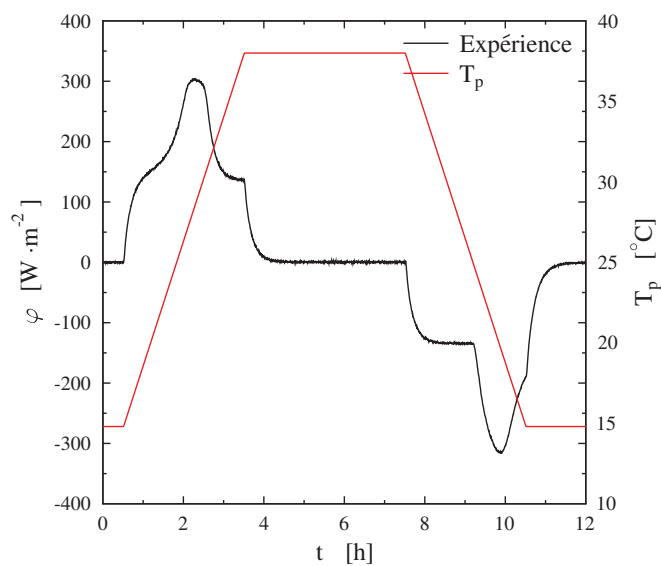


FIGURE 5.5 – Rampes de 3h - $7,73 \text{ K} \cdot \text{h}^{-1}$

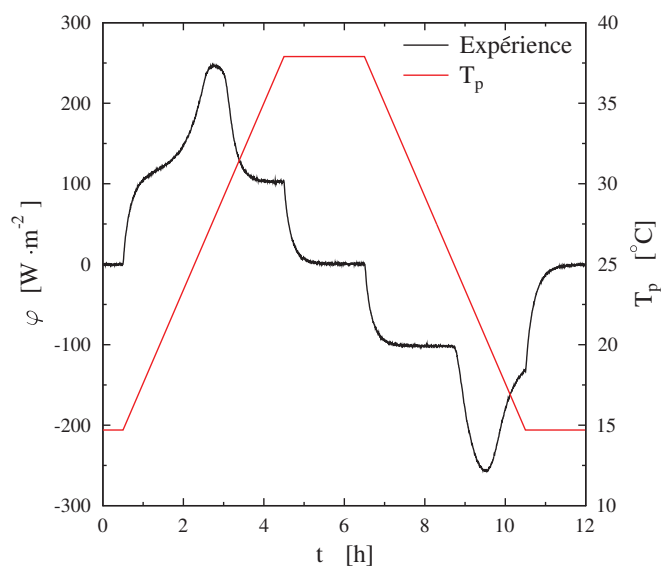
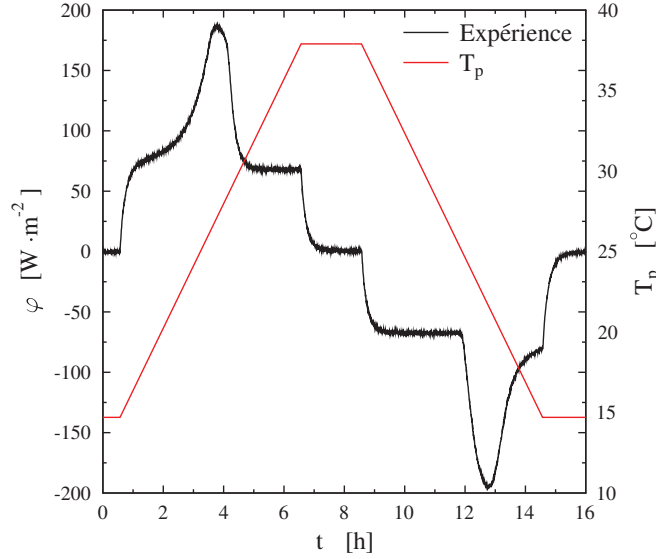


FIGURE 5.6 – Rampes de 4h - $5,8 \text{ K} \cdot \text{h}^{-1}$

FIGURE 5.7 – Rampes de 6h - $3,87 \text{ K} \cdot \text{h}^{-1}$

5.3 Modèle direct

Le modèle direct que nous allons décrire est résolu de la même manière que celui que nous avons développé pour la DSC (voir chapitre 2). La structure du code étant la même, les seuls éléments que l'on va décrire, sont la géométrie, l'équation d'état et les conditions aux limites. On considère un modèle unidimensionnel de part la présence d'isolant sur les faces latérales de l'échantillon et de sa géométrie aplatie.

5.3.1 Géométrie, maillage et conditions aux limites

On prend un maillage régulier à une dimension selon l'épaisseur e de l'échantillon, voir figure 5.8.

$$x(i) = \frac{i-1}{NI}e \quad \forall i \in [1; NI+1] \quad (5.1)$$

En considérant que les propriétés du mortier sont homogènes dans l'échantillon, le bilan (2.1) devient :

$$\rho \frac{\partial h}{\partial t} = \lambda \frac{\partial^2 T}{\partial x^2} \quad (5.2)$$

Pour ce modèle, comme seule une faible portion de l'échantillon est soumise au changement de phase, on considère de plus que la conductivité est constante quels que soient la température et l'état physique du MCP. Il s'agit d'une conductivité équivalente, définie dans la norme européenne EN-12667, dont on pourra comparer la valeur que l'on identifie à la

valeur déterminée par une mesure directe, voir section 5.6.

Les conditions aux limites entre l'échantillon et les plaques échangeuses à la température uniforme $T_p(t)$, sont des contacts solide/solide. Le phénomène de constriction thermique est usuellement représenté par des résistances de contact (RTC) [97] [98] [99] [102]. Nous avons choisi de considérer un coefficient d'échange équivalent égal à l'inverse de la résistance thermique de contact, $\alpha = \frac{1}{RTC}$. Ainsi en gardant les mêmes notations que dans le chapitre 2, les conditions aux limites s'écrivent :

$$\lambda \left(\frac{\partial T}{\partial x} \right)_{x=0} = \alpha (T(0, t) - T_p(t)) \quad (5.3)$$

$$-\lambda \left(\frac{\partial T}{\partial x} \right)_{x=e} = \alpha (T(L, t) - T_p(t)) \quad (5.4)$$

$$T_p(t) = T_0 + \beta(t) \times t \quad (5.5)$$

avec $\beta(t)$ une fonction continue par morceau qui prend successivement les valeurs $(0, \beta, 0, -\beta, 0)$ pour les évolutions de températures des expériences, voir figure 5.6.

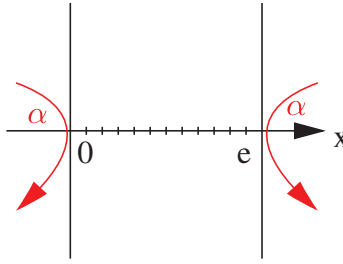


FIGURE 5.8 – Modélisation thermique du mortier

5.3.2 Equation d'état

Le MCP qui est incorporé dans le mortier a un thermogramme semblable à ceux d'un binaire sans eutectique [101]. Nous avons donc assimilé le comportement thermodynamique du MCP à celui d'une solution binaire sans eutectique faiblement concentré (voir figure 5.24 et 5.25).

Le mortier est un mélange d'agréats que l'on considère comme un mélange de phases distinctes. D'après le premier principe de la thermodynamique, l'enthalpie du mélange mortier et MCP s'écrit donc :

$$h(T) = (1 - \chi_{MCP})h_{b\acute{e}ton} + \chi_{MCP}h_{MCP} \quad (5.6)$$

χ_{MCP} est le taux massique de MCP dans le mortier.

En notant c la capacité thermique massique du mortier sans MCP, $c_{S,MCP}$, $c_{L,MCP}$ les capacités massiques solide et liquide du MCP, T_M la température de fin de fusion du MCP, T_{MCP}

une température caractéristique d'un corps pur que l'on associe au MCP et L_{MCP} une variation d'enthalpie associé au changement de phase du MCP pur. L'enthalpie devient :

$$T < T_M$$

$$\begin{aligned} h(T) = & (1 - \chi_{MCP})c \times (T - T_M) + \chi_{MCP}c_{S,MCP} \times (T - T_M) + \chi_{MCP} \left(1 - \frac{T_{MCP} - T_M}{T_{MCP} - T} \right) L_{MCP} \\ & + \chi_{MCP}(c_{S,MCP} - c_{L,MCP})(T_{MCP} - T_M) \ln \left(\frac{T_{MCP} - T}{T_{MCP} - T_M} \right) + G \end{aligned} \quad (5.7)$$

$$T \geq T_M$$

$$h(T) = (1 - \chi_{MCP})c \times (T - T_M) + \chi_{MCP}c_{L,MCP} \times (T - T_M) + G \quad (5.8)$$

$$\text{Avec } G = c_S \times (T_M - T_{ref}) + \left(\frac{T_{MCP} - T_M}{T_{MCP} - T_{ref}} - 1 \right) \tilde{L}_{MCP} + (c_S - c_L) \times (T_M - T_{MCP}) \ln \left(\frac{T_{MCP} - T_{ref}}{T_{MCP} - T_M} \right) + h_{ref}$$

la constante d'intégration de sorte que : $h_{ref} = 0 \text{ J} \cdot \text{kg}^{-1}$, $T_{ref} = 15^\circ \text{C}$.

En posant :

$$c_S = (1 - \chi_{MCP})c + \chi_{MCP}c_{S,MCP} \quad (5.9)$$

$$c_L = (1 - \chi_{MCP})c + \chi_{MCP}c_{L,MCP} \quad (5.10)$$

$$\tilde{L}_{MCP} = \chi_{MCP}L_{MCP} \quad (5.11)$$

On obtient les expressions suivantes pour l'enthalpie :

$$T < T_M$$

$$h = c_S \times (T - T_M) + \left(1 - \frac{T_{MCP} - T_M}{T_{MCP} - T} \right) \tilde{L}_{MCP} + (c_S - c_L) \times (T_{MCP} - T_M) \ln \left(\frac{T_{MCP} - T}{T_{MCP} - T_M} \right) + G \quad (5.12)$$

$$T \geq T_M$$

$$h = c_L \times (T - T_M) + G \quad (5.13)$$

L'expression que l'on obtient correspond à une enthalpie de solution binaire avec un liquide rectiligne, équations (2.52), (2.53) et (2.54), mais qui ne présente pas d'eutectique.

5.4 Caractérisation par méthode inverse

Comme pour l'application des méthodes inverses sur la DSC, chapitre 4, nous allons d'abord étudier les sensibilités du modèle direct aux divers paramètres. Puis nous déterminerons un intervalle de confiance pour chacun des paramètres que l'on identifie. Avant de présenter les résultats de l'identification et l'enthalpie que nous avons obtenue pour l'échantillon de mortier. Enfin nous comparerons les valeurs identifiées à des valeurs obtenues par mesure directe.

5.4.1 Etude de sensibilité

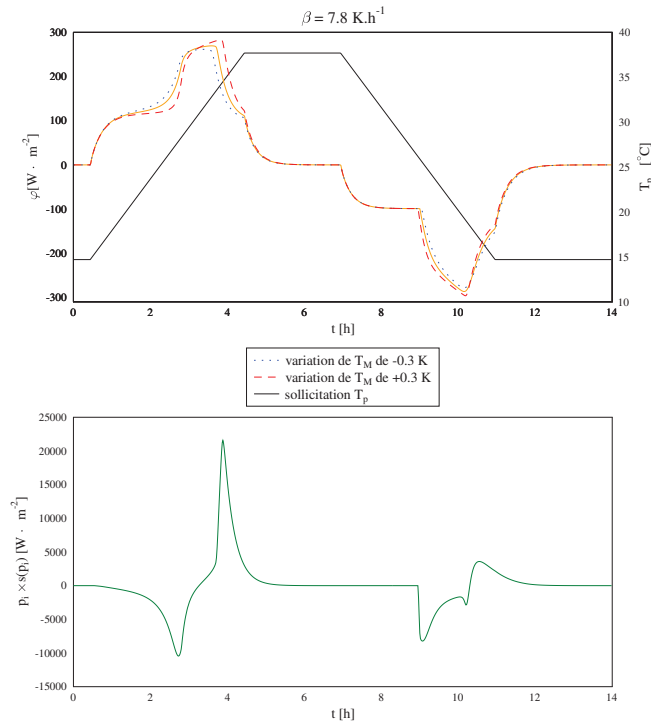
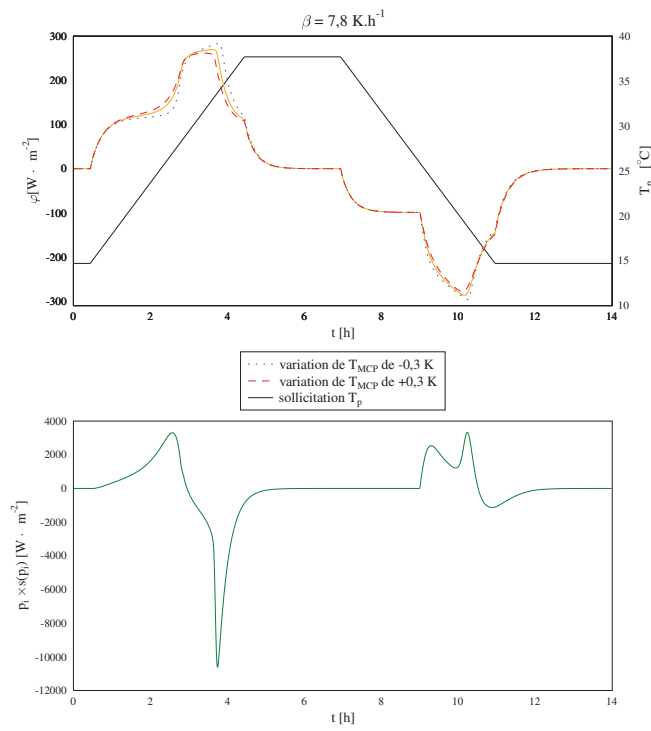
La forme de l'enthalpie du mortier étant la même que celle de la solution binaire si on annule la transition eutectique (voir sous-section 5.3.2 et 2.1.3), on va donc trouver un comportement similaire du modèle par rapport aux paramètres énergétiques, mais la géométrie étant différente le comportement du modèle selon les paramètres de transfert sera différent. Les courbes de sensibilité seront calculées pour des écarts de $\pm 10\%$ ou $0,3\text{ K}$ afin de rendre la déformation visible sur les courbes de flux.

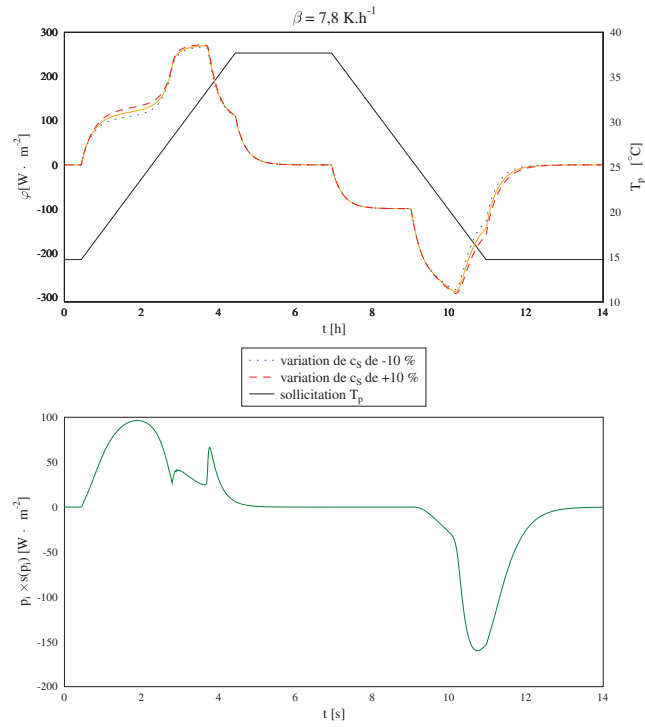
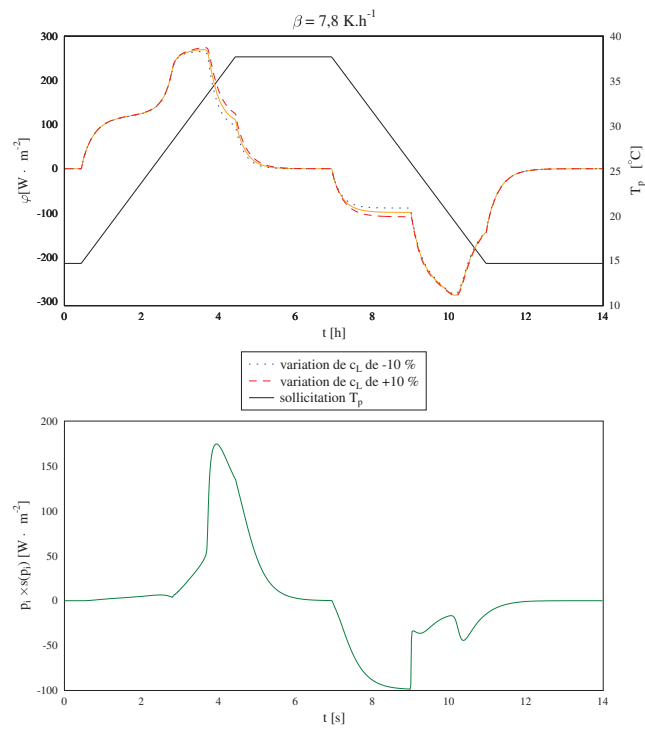
On remarque que les courbes de sensibilités de T_{MCP} et T_M , figures 5.9 et 5.10, sont de formes similaires mais de signes opposés, comme cela est le cas pour la solution binaire pour de très faibles fractions massiques en soluté x_0 , voir figure 4.2.

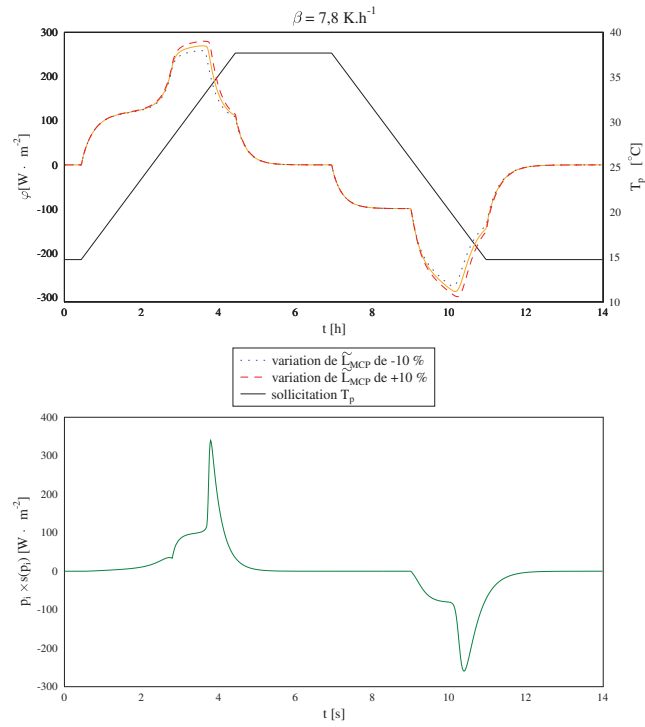
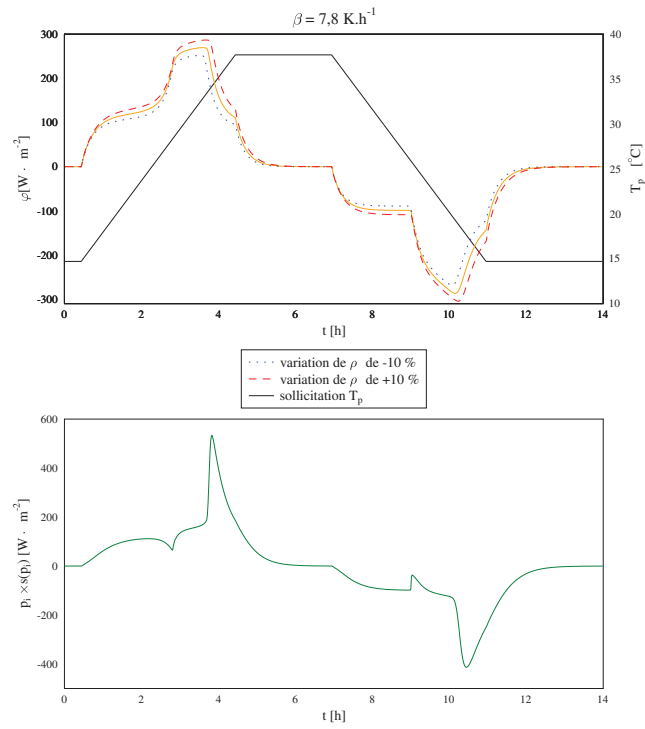
Les courbes de sensibilité du modèle aux paramètres c_S et c_L , figures 5.11 et 5.12, sont très différentes lors de la fusion et lors de la cristallisation. En effet, c_S intervient en amont de la fusion du MCP, son influence diminue fortement après le début du pic de fusion avant une rapide augmentation de l'influence pendant le retour à l'équilibre thermique de l'échantillon. Lors de la phase de cristallisation, la sensibilité du modèle à ce paramètre ne se retrouve qu'après le sommet du pic de cristallisation. Le rôle de c_L étant le même que c_S mais dans la phase liquide, et, ces deux paramètres ayant des valeurs très proches (voir les résultats de l'inversion dans le tableau 5.5), les courbes de sensibilités seront similaires au signe près si on intervertit les phases de fusion et de cristallisation.

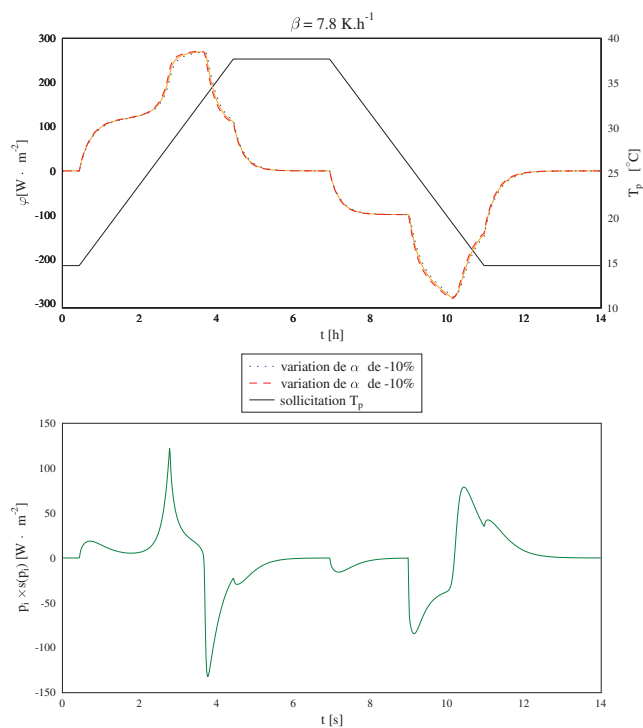
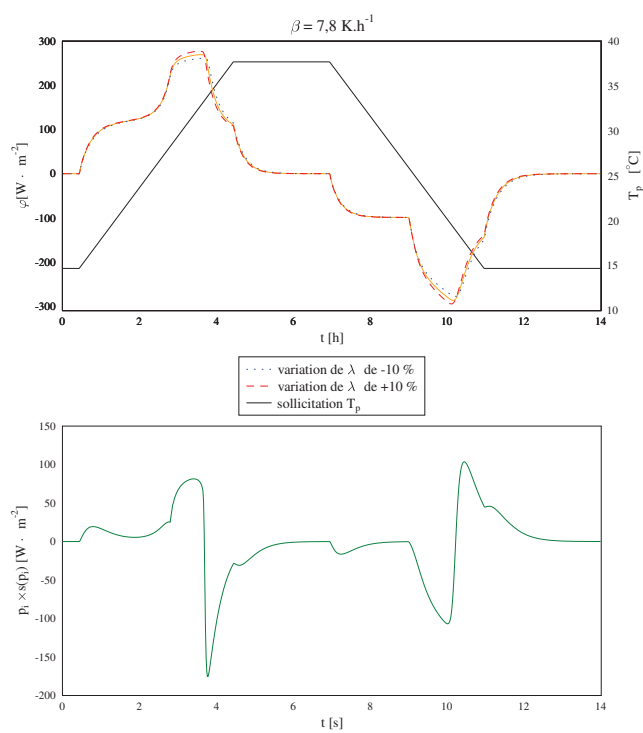
Sur la figure 5.13, est représentée la courbe de sensibilité du modèle au paramètre $\tilde{L}_{MCP}$. Celle-ci se retrouve principalement sur le sommet et la fin du pic de fusion car il s'agit d'une augmentation ou diminution de la quantité d'énergie qui est échangée au cours de la fusion. La courbe de sensibilité à ρ , figure 5.14, apparaît comme la superposition des courbes de sensibilité aux paramètres $\tilde{L}_{MCP}$, c_S et c_L , voir figure 5.18, comme c'est le cas pour les coefficients de sensibilité du modèle binaire, figure 4.8. On va donc fixer le paramètre ρ et ne pas chercher à l'identifier par une méthode inverse.

Pour ce modèle les coefficients de sensibilité à la conductivité équivalente λ et le coefficient d'échange équivalent α , figures 5.16 et 5.15, sont du même ordre de grandeur contrairement au modèle développé pour la DSC, figures 4.9, 4.10 et 4.11. Mais les courbes de sensibilités ayant des formes différentes, on peut identifier ces deux paramètres simultanément.

FIGURE 5.9 – Coefficient de sensibilité réduit relatif à T_M FIGURE 5.10 – Coefficient de sensibilité réduit relatif à T_{MCP}

FIGURE 5.11 – Coefficient de sensibilité réduit relatif à c_s FIGURE 5.12 – Coefficient de sensibilité réduit relatif à c_L

FIGURE 5.13 – Coefficient de sensibilité réduit relatif à $\tilde{L}_{MCP}$ FIGURE 5.14 – Coefficient de sensibilité réduit relatif à ρ

FIGURE 5.15 – Coefficient de sensibilité réduit relatif à α FIGURE 5.16 – Coefficient de sensibilité réduit relatif à λ

Cette étude nous a permis de vérifier que les paramètres étaient indépendants les uns des autres car leur courbe de sensibilité sont différentes, mise à part la masse volumique que l'on considère donc comme paramètre d'entrée de l'inversion. En effet, sa courbe de sensibilité est la somme des courbes de sensibilité aux paramètres c_S , c_L et $\tilde{L}_{M,w}$, caractérisant leur corrélation, voir figure 5.18. On voit aussi que les paramètres T_M et T_{MCP} sont les plus influents comme cela fut constaté dans une précédente étude [101]. On a remarqué que le modèle était plus sensible aux paramètres lors d'une fusion que lors d'une cristallisation. Les paramètres étant d'autant plus faciles à identifier que le modèle leur est sensible, la présence d'un faible degré de surfusion lors de la cristallisation ne devrait pas fausser l'identification.

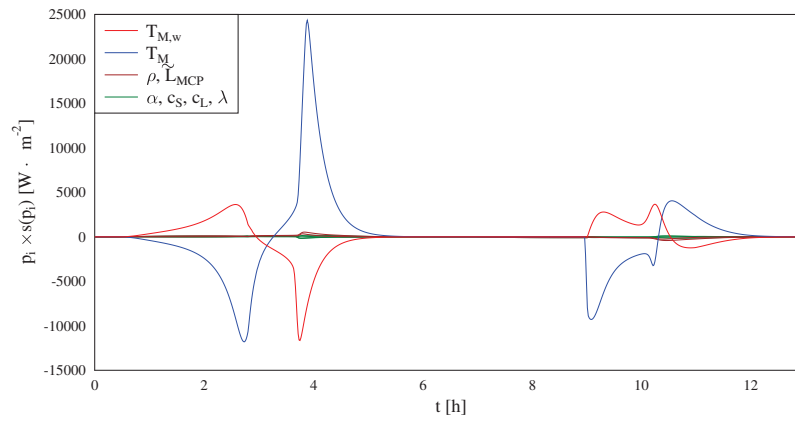


FIGURE 5.17 – Superposition des coefficients de sensibilité réduits

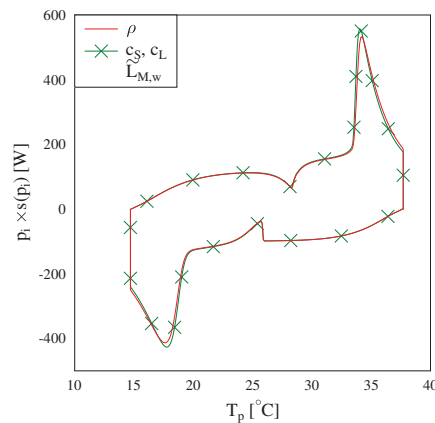


FIGURE 5.18 – Comparaison de la courbe de sensibilité au paramètre ρ à la somme de celles aux paramètres c_S , c_L et $\tilde{L}_{M,w}$

5.4.2 Etude de l'erreur

Comme pour le chapitre sur l'application des méthodes inverses sur les thermogrammes de DSC du chapitre 4, nous allons évaluer l'impact de diverses erreurs sur les résultats de l'identification pour avoir un intervalle de confiance, pour les paramètres que l'on identifie. L'impact de ces erreurs est évalué par un calcul théorique basé sur les matrices de sensibilité.

Il existe toujours six types d'erreurs pour le dispositif du LGCgE. Nous allons reprendre les mêmes considérations que dans le cas de la DSC. C'est-à-dire que nous nous limiterons aux erreurs déterministes ($\Delta_\pi p_j$) et stochastiques ($\Delta_\sigma p_i$) engendrées par la méthode inverse auquel on rajoute dans le cas des températures une incertitude de mesure $\Delta_T p_i = 0,1 \text{ K}$.

5.4.2.1 Erreur déterministe

Dans ce chapitre l'échantillon est majoritairement constitué de mortier dont les propriétés ne varient pas dans le domaine de température étudié et le MCP utilisé étant une paraffine, la masse volumique va très peu varier au cours des changements de phase. La dilatation volumique n'est pas ici une source d'erreur déterministe pour ρ . Par contre, contrairement au cas de la DSC, la géométrie de l'échantillon est constante et connue. On peut donc calculer l'incertitude sur ρ et la considérer comme étant une erreur déterministe.

En considérant que l'incertitude sur les longueur (L), largeur (L) et épaisseur (e) de l'échantillon est de 1 mm et que la masse est connue à $\pm 0,1 \text{ g}$ près. Il vient :

$$\begin{aligned} \frac{\Delta \rho}{\rho} &= \frac{\Delta m}{m} + \frac{2\Delta L}{L} + \frac{\Delta e}{e} \\ \Delta \rho &= 41,2 \text{ kg} \cdot \text{m}^{-3} \approx 3,3\% \end{aligned} \quad (5.14)$$

Pour les autres paramètres nous allons considérer une erreur déterministe de 1% ou 0,1 K.

Ces erreurs sur les paramètres identifiés qui sont engendrées par une erreur sur un paramètre d'entrée sont données par la relation (4.8) du chapitre 4. Leurs valeurs sont données dans le tableau 5.2. Dans ce tableau on fixe un paramètre de la colonne X, et on identifie les paramètres de la ligne correspondante.

Tableau 5.2 – Influence théorique d'une erreur déterministe de 3.3% sur ρ et de 1% ou 0.1 K sur les autres paramètres sur le résultat de l'inversion - cas du mortier

X	$\varepsilon(\rho)$ %	$\varepsilon(c_S)$ %	$\varepsilon(c_L)$ %	$\varepsilon(\lambda)$ %	$\varepsilon(T_{MCP})$ %	$\varepsilon(\tilde{L}_{MCP})$ %	$\varepsilon(T_M)$ %	$\varepsilon(\alpha)$ %
ρ	X	3,2	3,2	11×10^{-4}	$4,4 \times 10^{-4}$	3,3	$3,7 \times 10^{-4}$	22×10^{-3}
c_S	1	X	1	$1,9 \times 10^{-4}$	$1,3 \times 10^{-4}$	1	$1,1 \times 10^{-4}$	$6,7 \times 10^{-3}$
c_L	1	1	X	$2,3 \times 10^{-4}$	$1,3 \times 10^{-4}$	1	$1,1 \times 10^{-4}$	$6,7 \times 10^{-3}$
λ	0,3	0,2	0,2	X	$1,2 \times 10^{-3}$	0,45	1×10^{-3}	0,99
T_{MCP}	9,06	9,02	9,06	0,86	X	9,17	0,15	1,3
$\tilde{L}_{MCP}$	1	0,98	0,99	4,1	$1,3 \times 10^{-4}$	X	10^{-4}	$6,6 \times 10^{-3}$
T_M	9,75	9,71	9,73	2,17	0,17	9,8	X	1,17
α	4,2	4,1	4,2	0,57	$1,5 \times 10^{-3}$	4,1	10^{-3}	X

5.4.2.2 Erreur stochastique

Ces erreurs sur les paramètres identifiés qui ont été engendrées par le bruit expérimental sont données par les relations (4.11) et (4.12) du chapitre 4. Leurs valeurs sont données dans le tableau 5.3.

Prenons un thermogramme et comparons le flux enregistré dans une zone de flux stable ou le flux, φ_{moyen} , est constant, voir figure 5.19. Ce bruit correspond à la courbe de flux présentée en figure 5.22.

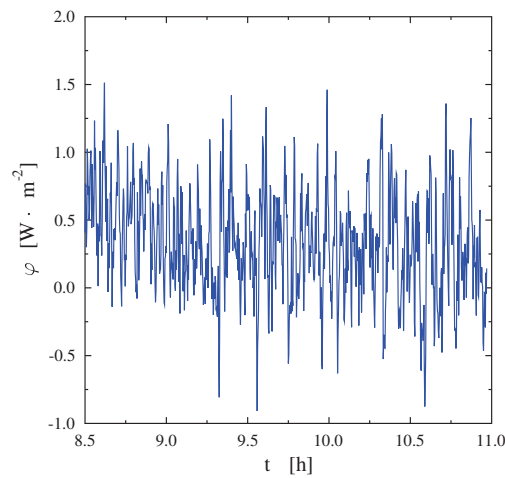


FIGURE 5.19 – Bruit expérimental sur une courbe de flux du mortier, figure 5.7

De ces mesures on obtient :

$$\sigma_{exp} = \sqrt{\frac{1}{n-1} \sum_{i=1}^n (\varphi(t_i) - \varphi_{moyen})^2} \approx 0,395 \text{ W} \cdot \text{m}^{-2} \quad (5.15)$$

Tableau 5.3 – Influence théorique d’une erreur stochastique de $0.395 \text{ W} \cdot \text{m}^{-2}$ sur le résultat de l’inversion - cas du mortier

X	$\varepsilon(\rho)$ %	$\varepsilon(c_S)$ %	$\varepsilon(c_L)$ %	$\varepsilon(\lambda)$ %	$\varepsilon(T_{MCP})$ %	$\varepsilon(\tilde{L}_{MCP})$ %	$\varepsilon(T_M)$ %	$\varepsilon(\alpha)$ %
$\Delta_\sigma p_i$	3,35	3,56	3,33	2,36	$2,3 \times 10^{-3}$	3,1	$7,6 \times 10^{-3}$	37,8

5.4.2.3 Intervalle de confiance

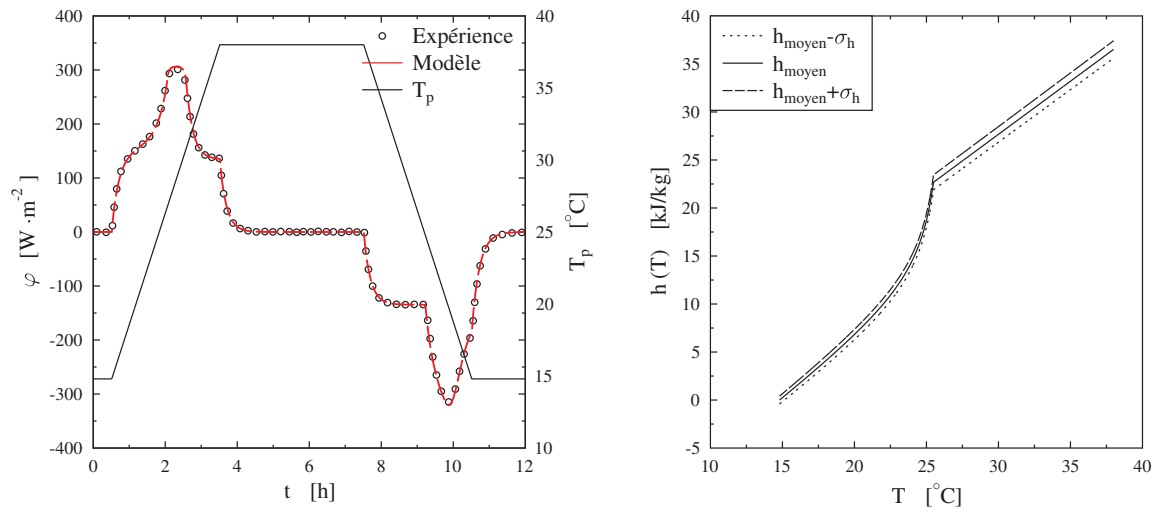
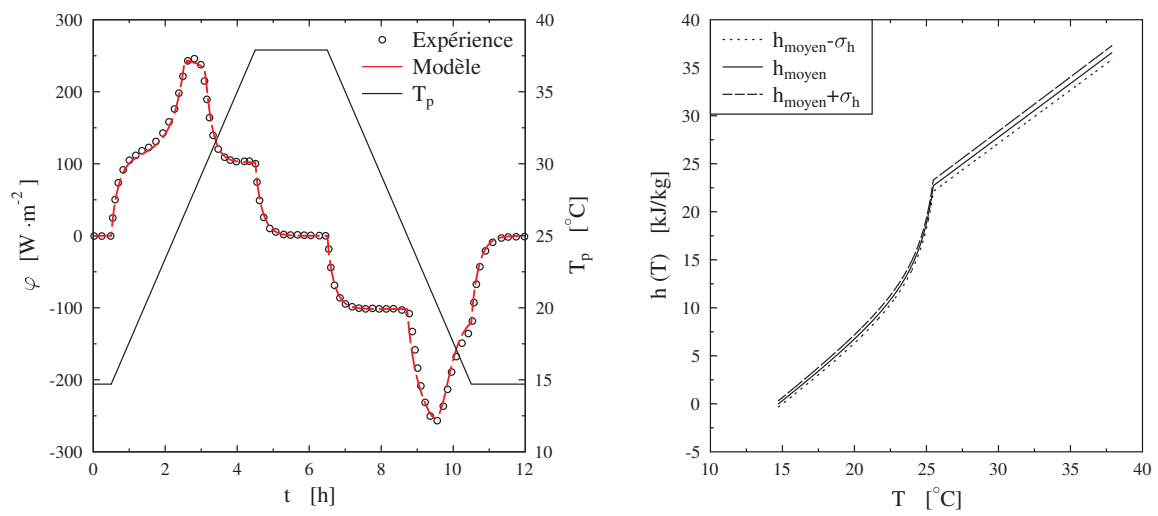
Comme pour le chapitre 4, un intervalle de confiance est associé à chacun des paramètres identifiés selon la relation (4.14), ces intervalles sont présentés dans le tableau 5.4. Ce qui permet de réaliser une étude statistique sur l’enthalpie que l’on a identifiée, les résultats de celle-ci seront présentés avec les résultats de l’inversion.

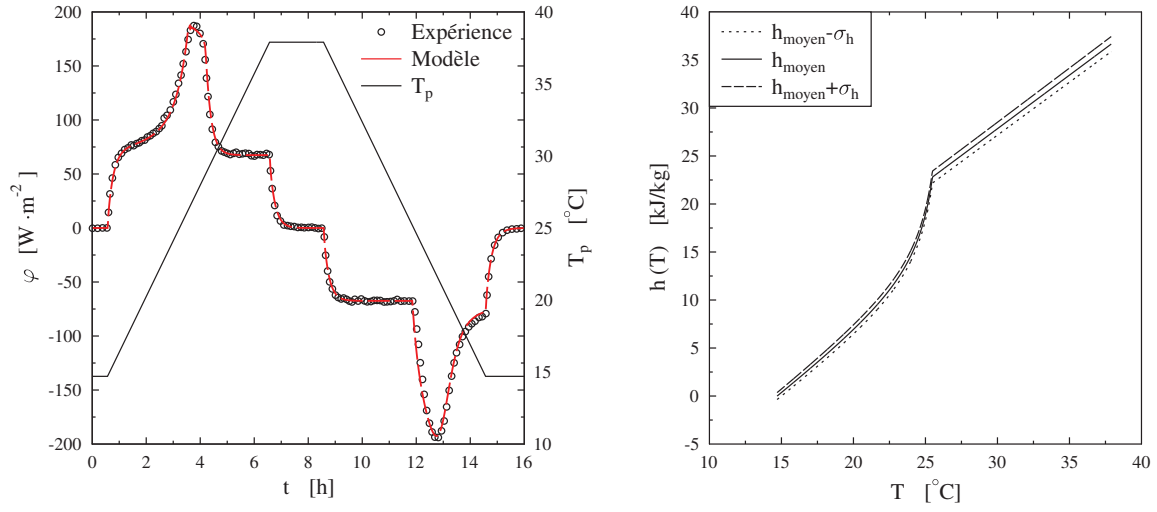
Tableau 5.4 – Intervalle de confiance des paramètres identifiés - cas du mortier

X	$\varepsilon(c_S)$ %	$\varepsilon(c_L)$ %	$\varepsilon(\lambda)$ %	$\varepsilon(T_{MCP})$ %	$\varepsilon(\tilde{L}_{MCP})$ %	$\varepsilon(T_M)$ %	$\varepsilon(\alpha)$ %
$\Delta_c p_i$	5	4,8	2,36	0,03	4,8	0,03	37,9

5.5 Résultat d’inversion

Nous allons maintenant présenter les résultats d’identifications par inversion réalisés sur cet échantillon. Ces résultats seront présentés avec leurs intervalles de confiance du tableau 5.4.

FIGURE 5.20 – Rampes de $3\text{h} - 7,73 \text{ K} \cdot \text{h}^{-1}$ FIGURE 5.21 – Rampes de $4\text{h} - 5,8 \text{ K} \cdot \text{h}^{-1}$

FIGURE 5.22 – Rampes de 6h - $3,87 \text{ K} \cdot \text{h}^{-1}$

Comme pour le cas de la DSC, chapitre 4, les courbes de flux expérimentales et identifiées se superposent parfaitement et les intervalles de confiance sur les paramètres identifiés sont suffisamment faibles pour que l'on puisse se fier à l'enthalpie que nous avons reconstruite. Cette enthalpie est indépendante de la vitesse de réchauffement car on obtient la même pour les trois vitesses, voir figure 5.23.

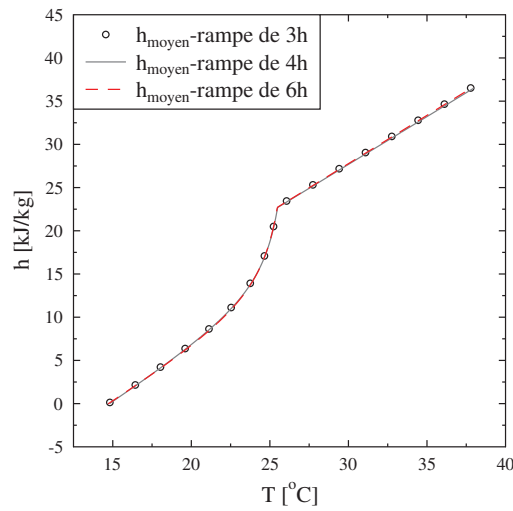
FIGURE 5.23 – Comparaison des enthalpies identifiées aux vitesses de $10,3-7,8-5,2 \text{ K} \cdot \text{h}^{-1}$

Tableau 5.5 – Paramètres identifiés sur les courbes de flux expérimentales du mortier à 12,4 % de MCP

β $K \cdot h^{-1}$	c_S $J \cdot kg^{-1} \cdot K^{-1}$	c_L $J \cdot kg^{-1} \cdot K^{-1}$	$\tilde{L}_{MCP}$ $J \cdot kg^{-1}$	T_M $^{\circ}C$	T_{MCP} $^{\circ}C$	λ $W \cdot m^{-1} \cdot K^{-1}$	α $W \cdot m^{-2} \cdot K^{-1}$
10,3	1145 ± 57	1109 ± 54	11,8 ± 0,6	25,5 ± 0,1	26,7 ± 0,1	0,52 ± 0,02	109 ± 42
7.8	1115 ± 57	1118 ± 54	12,1 ± 0,6	25,5 ± 0,1	26,7 ± 0,1	0,50 ± 0,02	136 ± 52
5.2	1130 ± 57	1112 ± 54	12,0 ± 0,6	25,5 ± 0,1	26,7 ± 0,1	0,51 ± 0,02	135 ± 52

5.6 Comparaison de la caractérisation par méthode inverse et des caractérisations directes

Le dispositif que nous utilisons, figure 5.3, permet en outre de déterminer une conductivité apparente du mortier. On détermine la conductivité équivalente en soumettant le mortier à un gradient de température, c'est-à-dire en fixant des consignes de température différentes mais constantes pour les deux échangeurs du dispositif, figure 5.3. Et en approximant le gradient de température par une forme linéaire :

$$\lambda = \frac{(\varphi_D + \varphi_G) \times e}{2|T_D - T_G|} \quad (5.16)$$

où T_D , T_G , φ_D et φ_G sont les températures mesurées par les thermocouples des fluxmètres et les flux surfaciques mesurés par les fluxmètres, voir figure 5.3.

Par une méthode similaire on peut évaluer les capacités solide et liquide en soumettant un échantillon homogène en température à un même créneau de température ΔT de la part des deux échangeurs, et en relevant l'énergie échangée par l'échantillon $E_{sensible}$ à travers les fluxmètres pour atteindre un nouvel état d'équilibre thermique. Un bilan énergétique permet ensuite de déterminer la capacité à l'état physique dans lequel se trouve le MCP au cours de l'expérience.

$$c = \frac{E_{sensible}}{\Delta T} \quad (5.17)$$

Enfin, comme $c_S \approx c_L$ on peut calculer la variation d'enthalpie $\tilde{L}_{MCP}$ due au changement de phase en résolvant un bilan énergétique entre deux états stationnaires aux températures T_1 et T_2 , avant et après le changement de phase. Donc en soustrayant la chaleur sensible échangée $E_{sensible}$ à la quantité d'énergie échangée au cours de la transformation $E_{sensible+latent}$. La chaleur sensible est approximé comme étant la moyenne des deux capacités qui sont presque de même valeurs, multipliée par l'écart maximal de température au cours de cette transformation [103].

$$\tilde{L}_{MCP} = E_{sensible+latent} - E_{sensible} \approx E_{sensible+latent} - \frac{c_S + c_L}{2} \times (T_2 - T_1) \quad (5.18)$$

Les valeurs qui ont été déterminées par l'équipe de Laurent Zalewski sont dans le tableau 5.6.

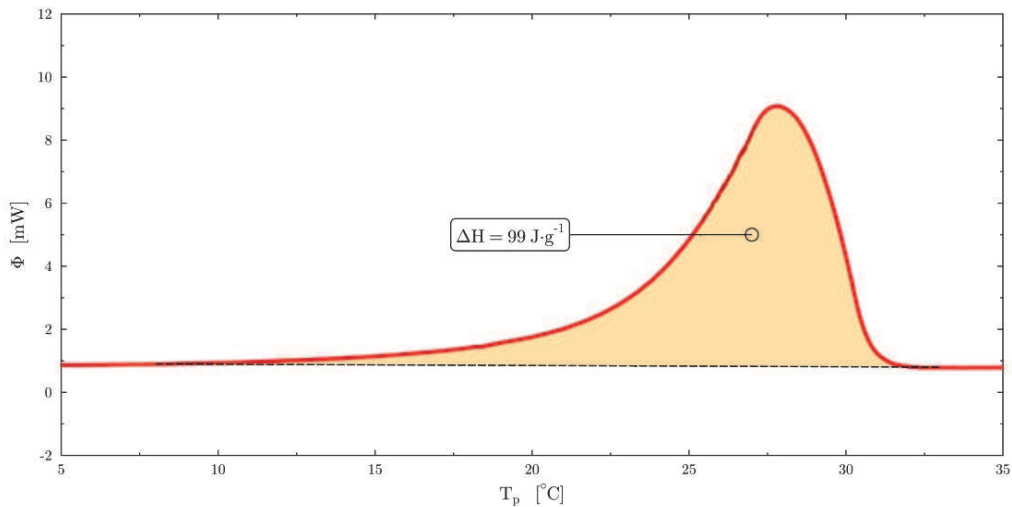
Tableau 5.6 – Propriétés de l'échantillon à 12.4% de MCP qui ont été déterminé avec le dispositif [101]

c_S ($J \cdot kg^{-1} \cdot K^{-1}$)	c_L ($J \cdot kg^{-1} \cdot K^{-1}$)	$\tilde{L}_{MCP}$ ($J \cdot g^{-1}$)	λ ($W \cdot m^{-1} \cdot K^{-1}$)
1119 ± 56	1080 ± 54	$11,589 \pm 1952$	$0,55 \pm 0,02$

En comparant les tableaux 5.5 et 5.6, on remarque que toutes les propriétés qui ont été identifiées par les deux méthodes concordent.

De plus, on a déterminé par DSC d'une part la variation d'enthalpie L_{MCP} et d'autre part la température T_M , voir figures 5.24 et 5.25, sur un échantillon de MCP pur, du Micronal PCM ®DS 5001 X de masse $m = 5,78 \text{ mg}$. On a estimé L_{MCP} à $99 J \cdot kg^{-1}$, figure 5.24, et T_M à $25,5^\circ C$, figure 5.25, on retrouve cette valeur de T_M dans le tableau 5.6. Cela nous permet de vérifier notre modélisation du mortier, en calculant $\tilde{L}_{MCP}$ par le biais de l'équation (5.11), et en approximant la surface du pic du thermogramme de DSC à $L_{MCP} \cdot 99 \times 12,4\% = 12,276 = 11,589 \pm 1952 J \cdot g^{-1}$. La variation d'enthalpie au cours de la fusion est concordante avec cette hypothèse aux précisions de mesures près.

La température T_M du mortier correspond à la température de liquidus de notre binaire dont il a la même forme d'enthalpie [101]. Cette température peut être déterminé en superposant plusieurs expériences de DSC sur le même échantillon mais à différentes vitesses. La droite qui relie les sommets des pics croise l'axe en T_p à la température T_M [94], voir figure 5.25. Car l'axe horizontal correspond au cas limite d'une expérience à vitesse nulle où la transition de phase se ferait à l'équilibre thermique.

FIGURE 5.24 – Caractérisation de L_{MCP} par calorimétrie DSC à $5 K \cdot min^{-1}$ [101]

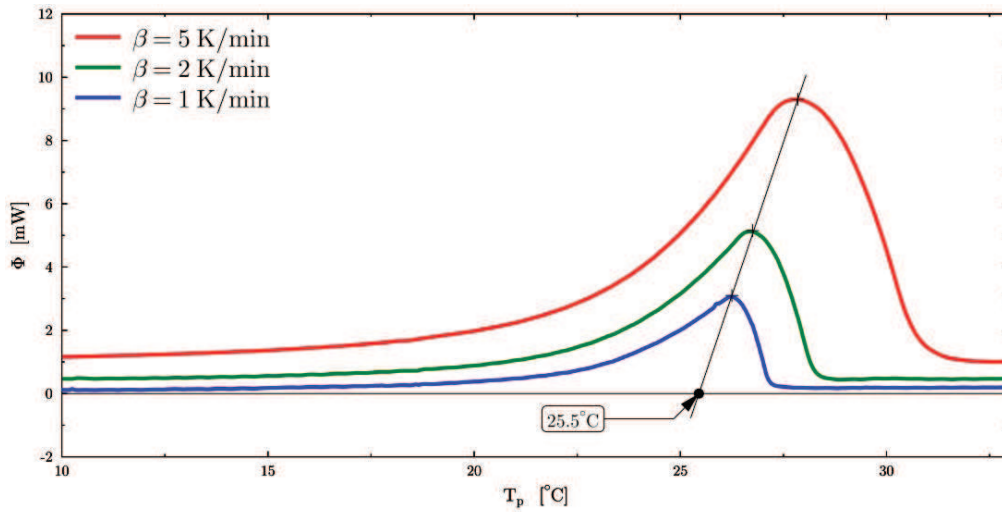


FIGURE 5.25 – Caractérisation de T_M par calorimétrie DSC [101]

5.7 Conclusion

Dans ce chapitre nous avons caractérisé un mortier contenant des MCP selon l'approche des méthodes inverses, et nous avons obtenu une même enthalpie pour représenter le comportement thermodynamique du mortier lors d'un cycle de fusion-cristallisation pour deux vitesses différentes. Nos résultats concordent avec ceux qui ont été obtenus par des méthodes directes. Mais la caractérisation n'est pas complète du fait de la non-répétabilité des échantillons. Dans la pratique il n'est pas réalisable de reproduire un mortier donné. Néanmoins par cette méthode une seule expérience est nécessaire pour caractériser les propriétés du mortier, alors que plusieurs expériences avec différents appareils étaient nécessaires jusqu'à présent.

Conclusion générale

L'objectif de cette thèse était de montrer que l'utilisation des méthodes inverses sur des thermogrammes de DSC bien étalonnés permet de déterminer la fonction enthalpie d'un matériau à changement de phase solide-liquide.

Dans le premier chapitre nous avons d'abord posé le problème du changement de phase et la modélisation que nous avons adoptée. Le changement de phase, solide-liquide en l'occurrence, est modélisé par un bilan énergétique pour lequel on a présenté les trois méthodes qui sont utilisées pour sa résolution. Nous avons choisi de résoudre ce bilan énergétique par la méthode enthalpique en considérant que la transformation suit localement l'équilibre thermodynamique. Puis nous avons décrit les principes et fonctionnement d'un calorimètre différentiel à balayage. Cet appareil nous sert de base pour développer notre méthode d'identification par méthode inverse. Enfin nous avons illustré les défauts de caractérisation de la méthode dite du « c_p équivalent » pour déterminer l'enthalpie, défaut qui n'est pas résolu même à faible vitesse. Sont également décrites les erreurs de dimensionnement des procédés énergétiques qui en résultent.

Dans le deuxième chapitre nous avons développé notre modélisation directe de la DSC afin de générer un thermogramme numérique identique au thermogramme expérimental que l'on obtient dans les mêmes conditions. Ce modèle se veut thermodynamiquement consistant dans le sens où il utilise l'équation d'état correspondant à l'échantillon. Pour cette étude nous avons considéré deux cas : celui d'un corps pur et celui d'une solution binaire présentant un liquidus rectiligne. Les paramètres du modèle direct contiendront les propriétés énergétiques qui serviront à reconstituer l'enthalpie massique de l'échantillon, les propriétés de transfert et celles de simulation. Une étude paramétrique de ce modèle est présentée dans ce chapitre 2 pour illustrer l'influence de chaque paramètre sur l'enthalpie correspondante et les conséquences sur la modification du thermogramme. Une validation expérimentale de notre modèle direct de DSC est également présentée. Pour terminer ce deuxième chapitre, nous proposons un modèle réduit à une dimension afin de réduire le temps de calcul. Ce modèle réduit modifie pour cela les propriétés de transfert sans altérer les propriétés énergétiques (température et variation d'enthalpie).

Dans le troisième chapitre, nous présentons le principe des méthodes inverses axées sur l'identification de paramètres par l'utilisation de trois méthodes d'optimisation différentes : la méthode de Levenberg-Marquardt, l'algorithme génétique et la méthode du simplexe de Nedler-Mead. Ce chapitre se terminera en spécifiant l'écriture de ces méthodes pour identifier les paramètres du modèle direct de DSC afin de reconstruire le même thermogramme.

Dans le quatrième chapitre, nous aborderons les résultats de l'identification dans le cas de la DSC. En premier lieu, une étude des courbes de sensibilités aux paramètres du modèle direct est conduite afin de vérifier la non-interdépendance des paramètres entre eux. Et le cas échéant nous retirerons ce/ces paramètres du processus d'identification. Ensuite une comparaison succincte des trois méthodes d'optimisation est faite sur un cas. On en a déduit le choix de la méthode du simplexe de Nedler-Mead.

La méthode étant en place, les outils étant choisis, nous montrons ensuite que notre approche par méthode inverse fonctionne, en re-déterminant les propriétés énergétiques correspondant à différents thermogrammes que l'on a simulés. Enfin, nous appliquons la méthode à des thermogrammes expérimentaux et vérifions la concordance des résultats entre eux ainsi qu'avec la littérature et la théorie thermodynamique. Il est bien montré que, bien que des thermogrammes puissent par exemple avec différentes des vitesses de réchauffement, avoir des formes très différentes, les variations d'enthalpie sont déterminées intrinsèquement vis à vis de la thermodynamique.

Pour clore cette thèse, nous présentons dans le chapitre 5 une utilisation de la méthode inverse pour caractériser un échantillon macroscopique de mortier contenant des inclusions de MCP. Pour ce faire un modèle a été développé pour simuler les courbes de flux que l'on enregistre sur un dispositif expérimental associé au mortier. Ce modèle utilise une équation d'état similaire à celle des solutions binaires présentant un liquidus rectiligne.

Le principal avantage de notre approche basée sur l'inversion est qu'elle permet de caractériser les propriétés énergétiques d'un échantillon à partir d'une seule expérience de calorimétrie. Cette détermination nécessite cependant de faire un choix *a priori* concernant la forme de l'équation d'état et donc de son expression mathématique.

Ainsi, la méthode développée ici pour les corps purs et les solutions binaires faiblement concentrées, peut être facilement étendue et adaptée à d'autres cas dont on est capable de fournir une expression mathématique pour $h(T)$. Des travaux sont d'ores et déjà en cours au laboratoire, notamment sur les solutions solides et les modèles de type NRTL.

Nous évaluons également la possibilité de généraliser la méthode en n'imposant plus d'expression mathématique pour $h(T)$, ce qui permettrait de ne plus avoir de choix *a priori* à poser.

Cette méthode devrait aussi contribuer à la mise au point d'une qualification des propriétés thermiques des MCP utilisés dans différentes applications comme dans le bâtiment objet de départ de cette étude, ou dans d'autres domaines de stockage thermique (centrales solaire, industrie du froid, protection électronique...).

Bibliographie

- [1] Commissariat général au développement durable. Chiffres clés de l'environnement - édition 2013, 2013.
- [2] GIEC. 5e Rapport du GIEC, résumé technique, 2013.
- [3] Key World Energy Statistics. iea, 2014.
- [4] Commissariat général au développement durable. Chiffres clés de l'énergie - édition 2013, 2014.
- [5] Les Technologies Clés 2005. ISBN : 2-11-091986-8. les Éditions de l'industrie, Télédor 536 Paris., 2000.
- [6] B. Multon, J. Aubry, P. Haessig, and H. Ben Ahmed. Systèmes de stockage d'énergie électrique. Techniques de l'ingénieur, 2013.
- [7] J. P. Bédécarrats. Etude des transformations des matériaux à changement de phase encapsulés destinés au stockage du froid. PhD thesis, Université de Pau et des Pays de l'Adour, 1993.
- [8] J. P. Dumas. Stockage du froid par chaleur latente. Techniques de l'ingénieur, 2002.
- [9] P. Reghem. Etude hydrodynamique de fluide diphasiques solide-liquide en conduite circulaire : Application au coulis de glace. PhD thesis, Université de Pau et des Pays de l'Adour, 2002.
- [10] D. N. Nkwette and F. Haghighat. Thermal energy storage with phase change material. A state of the art review. Sustainable Cities and Society, (10) :87–100, 2014.
- [11] L. Yang, H. Yan, and J. Lam. Thermal comfort and building energy consumption implications. A review. Applied Energy, (115) :164–173, 2014.
- [12] Y. Dutil, D. R. Rousse, S. Lassue, A. Joulin, J. Virgone, F. Kuznik, K. Johannes, J. P. Dumas, J. P. Bédécarrats, A. Castell, and L. F. Cabeza. Modeling phase change materials behavior in bulding applications : Comments on material characterization and model validation. Renewable Energy, 61 :132–135, 2014.
- [13] A. Sharma, V. Tyagi, C. Chen, and D. Buddy and. Review on thermal energy storage with phase change materials and applications. Renewable and Sustainable Energy Reviews, (13) :318–345, 2009.
- [14] N. R. Jankowski and F. P. McCluskey. A rview of phase change materials for vehicule component thermal buffering. Applied Energy, (113) :1525–1561, 2014.
- [15] C. Castellòn, E. Günther, H. Mehling, S. Hiebler, and L. F. Cabeza. Determination of the enthalpy of PCM as a function of temperature using a heat-flux DSC. A study of different measurement procedures and their accuracy. International Journal of Energy Research, 32(13) :1258–1265, 2008.

- [16] M. Pomianowski, P. Heiselberg, and Y. Zhang. Review of thermal energy storage technologies based on PCM application in buildings. Energy and Buildings, (67) :56–69, 2013.
- [17] Y. Dutil, D. R. Rousse, N.B. Salah, S. Lassue, and L. Zalewski. A review on phase-change materials : Matematical modeling and simulations. Renewable and Sustainable Energy Reviews, 15 :112–130, 2011.
- [18] V. Rakesh. Modeling Melting : Apparent Specific Heat. COMSOL, 2007.
- [19] M. Koschenz and B. Lehmann. Development of a thermally activated ceiling panel with PCM for application in lightweight and retrofitted buildings. Energy and Buildings, 28(11-12) :1291–1298, 2008.
- [20] M. Ibáñez, A. Lázaro, B. Zalba, and L. F. Cabeza. An approach to the simulation of PCMs in building applications using TRNSYS. Applied Thermal Engineering, 25(11-12) :1796–1807, 2005.
- [21] F. Kuznik, J. Virgone, and K. Johannes. Developement and validation of a new TRN-SYS type for the simulation of external building walls containing PCM. Energy and Buildings, 42(7) :1004–1009, 2012.
- [22] A. Bricard and D. Gobin. Transfert de chaleur avec changement d'état solide-liquide. Techniques de l'ingénieur, 2001.
- [23] L. Landau. Heat conduction in a melting solid. Applied Math, (8) :81–94, 1950.
- [24] C. Bérnard and D. Gobin and A. Zanolì. Moving boundary problem : heat conduction in the solid phase of a phase-change material during melting driven by natural convection in the liquid. International Journal of Heat and Mass Transfer, 29(11) :1669–1681, 1986.
- [25] S. Geoffroy and S. Mergui and D. Gobin. Heat and mass transfert during solidification of a binary solution from a horizontal plate. Experimental Thermal and Fluid Science, (29) :169–178, 2005.
- [26] V. Sobotka, A. Agazzi, N. Boyard, and D. Delaunay. Parametric model for the analytical determination of the solidification and cooling times of semi-crystalline polymers. Applied Thermal Engineering, 50 :416–421, 2013.
- [27] C.R. Swaminathan and V.R. Voller. Towards a general numerical scheme for solidification systems. International Journal of Heat and Mass Transfer, 40(12) :2859–2869, 1997.
- [28] V. R. Voller. Advances in Numerical Heat Transfert, volume 1, chapter 9. Taylor and Francis, 1996.
- [29] V. R. Voller. An enthalpy method for modeling dendritic growth in a binary alloy. International Journal of Heat and Mass Transfer, 51 :823–834, 2008.
- [30] C.-H. Chang and L.-S. Chao. Modeling analysis of melting and solidifying processes in excimer laser crystallization of a-Si films with effective specific heat-enthalpy method. International Journal of Heat and Mass Transfer, 35 :571–576, 2008.
- [31] T. Norton, A. Delgado, E. Hogan, P. Grace, and Da-Wen Sun. Simulation of high pressure freezing processes by enthalpy method. Journal of Food Engineering, 91 :260–268, 2009.
- [32] Perkin-Elmer. SOP (Standard Operating Procedure) for the Perkin-Elmer Pyris 1 DSC.

- [33] J. P. Dumas, D. Clausse, and F. Broto. A study of thermograms obtained through differential scanning calorimetry of an emulsion of a supercooled liquid. Thermochimica Acta, (13) :261–275, 1975.
- [34] Y. Zeraouli. Etude thermique des transformations des émulsions concentrées. Application à la calorimétrie à balayage. PhD thesis, Université de Pau et des Pays de l'Adour, 1991.
- [35] J. Schawe and C. Schick. Influence of the heat conductivity of the sample on DSC curves and its correction. Thermochimica Acta, (187) :335–349, 1991.
- [36] E. Günther, S. Hiebler, H. Medhling, and R. Redlich. Enthalpy of Phase Change Materials as a Function of Temperature : Required Accuracy and Suitable Measurement Methods. International Journal of Thermophysics, (30) :1257–1269, 2009.
- [37] W. Lin, D. Dalmazzone, W. Furst, A. Delahaye, L. Fournaison, and P. Clain. Accurate DSC measurement of the phase transition temperature in the TPBP-water system. Journal of Chemical Thermodynamics, (61) :132–137, 2013.
- [38] C. Barreneche, A. Solé, L. Miro, I. Martorell, A. I. Fernández, and L. F. Cabeza. Study on differential scanning calorimetry analysis with two operation modes and organic and inorganic phase change material (PCM). Thermochimica Acta, (553) :233–26, 2013.
- [39] M.J. Richardson. Quantitative aspects of differential scanning calorimetry. Thermochimica Acta, (300) :15–28, 1997.
- [40] S. Gibout, E. Franquet, J.-P. Bédécarrats, and J. P. Dumas. Comparison of different modelings of pure substances during melting in a DSC experiment. Thermochimica Acta, (528) :1–8, 2012.
- [41] GEMO. Détermination des caractéristiques physiques par D.S.C., number PPH-610 - 00, Avril 2008.
- [42] Dartmouth college. Chemistry 75 : Phase Transitions and Differential Scanning Calorimetry, 2010.
- [43] X. Tardif, B. Pignon, N. Boyard, J. W. P. Schmelzer, V. Sobotka, D. Delaunay, and C. Schick. Experimental study of crystallisation of PolyEtherEterKeton (PEEK) over a large temperature range using a nano-calorimeter. Polymer Testing, 36 :10–19, 2014.
- [44] G.W.H. Höhne, W.F. Hemminger, and H.J. Flammersheim. Differential Scanning Calorimetry 2nd edition. Springer, New York, 2003.
- [45] A. Lazaro, C. Penalosa, A. Solé, G. Diarce, T. Hausmann, M. Fois and B. EZalba, S. Gshwander, and L. F. Cabeza. Intercomparative tests on phase change materials characterisation with differential scanning calorimeter. Applied Energy, 2012.
- [46] J. P. Dumas, S. Gibout, L. Zalewski, K. Johannes, E. Franquet, S. Lassue, J. P. Bédécarrats, P. Tittlein, and F. Kuznik. Interpretation of calorimetry experiments to characterise phase change materials. International Journal of Thermal Sciences, 78 :48–55, 2014.
- [47] Proceedings Oklahoma Academy of Science. Limitation in the apparent heat capacity formulation for heat transfer with phase change, volume 67 :83-88, 1987.
- [48] E. Günther, S. Hiebler, and H. Mehling. Determination of the heat storage capacity of PCM and PCM objects as a function of temperature. In Proc. Ecstock - The Tenth Int. Conf. on Thermal Energy Storage (Stockton College, NJ, USA)., 2006.

- [49] J. P. Dumas, S. Gibout, L. Zalewski, K. Johannes, E. Franquet, S. Lassue, J.-P. Bedecarrats, and P. Tittlein. Nécessité de l'interprétation correcte de la calorimétrie pour l'utilisation des Matériaux à Changement de Phases (MCP). In Congrès français de thermique SFT2012, 2012.
- [50] J. P. Dumas. Etude de la rupture de métastabilité et du polymorphisme des corps organiques. PhD thesis, Université de Pau et des Pays de l'Adour, 1976.
- [51] Y. Zeraouli, A. J. Ehmimed, and J. P. Dumas. Modèles de transferts thermiques lors de la fusion d'une solution binaire dispersée. International Journal of Thermal Science, (39) :780–796, 2000.
- [52] A. Jamil, T. Kousksou, K. El Omari, Y. Zeraouli, and Y. Le Guer. Heat transfer in salt solutions enclosed in DSC cells. Thermochimica Acta, 2010.
- [53] F. Agyenim, N. Hewitt, P. Eames, and M. Smyth. A review of materials, heat transfer and phase change problem formulation for latent heat thermal energy storage systems. Renewable and Sustainable Energy Reviews, (14) :615–628, 2010.
- [54] A.W. Date. Introduction to Computational Fluid Dynamics. Cambridge University Press, 2005.
- [55] A.B. Tayler. Mathematical formulation of Stefan problems, in : Moving Boundary Problems in Heat Flow and Diffusion. Univ. Oxford Conf., pages 120–137, 1975.
- [56] V. Voller and C. Swaminathan. General source-based method for solidification phase change. Numerical Heat Transfer, (B-19) :175–189, 1991.
- [57] P. Verma, S. Varun, and Singal. Review of mathematical modeling on latent heat thermal energy storage systems using phase-change material. Renewable and Sustainable Energy Reviews, (12) :999–1031, 2008.
- [58] J. Timmermans. The Physico-Chemical Constants of Binary Systems in Concentrated Solutions. New York, Interscience Publishers, 1960.
- [59] T. Kousksou, A. Jamil, Y. Zeraouli, and J. P. Dumas. Equilibrium liquidus temperatures of binary mixtures from differential scanning calorimetry. Chemical Engineering Science, (62) :6516–6523, 2007.
- [60] Cours Université Paris 6 (Jussieu) : lois de conservation et modélisation VIII : Volumes Finis [online]. 2007. URL : http://www.ann.jussieu.fr/~despres/BD_fichiers/vf.pdf.
- [61] S. Gibout. Méthodes inverses de calcul appliquées à l'étude des transferts thermiques lors de la cristallisation de liquides dispersés surfondus. PhD thesis, L'Université de Pau et des Pays de l'Adour, 2001.
- [62] W. Maréchal. Modélisation d'expérience de calorimétrie lors d'un changement d'état solide-liquide d'un corps pur. Master's thesis, école nationale supérieure en génie des technologies industrielles, 2011.
- [63] B. Zalba, J. M. Marin, L. F. Cabeza, and H. Melhling. Review on thermal energy storage with phase change : materials, heat transfer analysis and applications. Applied Thermal Engineering, (23) :251–283, 2003.
- [64] Q. Duan, F. L. Tan, and K. C. Leong. A numerical study of solidification of n-hexadecane based on the enthalpy formulation. Journal of Materials Processing Technology, (120) :249–258, 2002.

- [65] E. Moukhina. Enthalpy calibration for wide DSC peaks. Thermochimica Acta, (522) :96–99, 2011.
- [66] B. Han, J. H. Choi, J. A. Dantzig, and J. C. Bischof. A quantitative analysis on latent heat of an aqueous binary mixture. Cryobiology, (52) :146–151, 2006.
- [67] E. Franquet, S. Gibout, and J. P. Dumas. Simulation des transferts thermiques en calorimétrie différentielle à balayage : influence de la géométrie de la surface libre sur le thermogramme. In Actes du congrès annuel de la Société Française de Thermique (Société Française de Thermique), Valenciennes, 2010.
- [68] S. Gibout, E. Franquet, W. Maréchal, and J. P. Dumas. On the use of a reduced model for the simulation of melting of solutions in DSC experiments. Thermochimica Acta, 566(20) :118–123, 2013.
- [69] C. Vassard. Le problème du voyageur de commerce, Brochure APMEP n° 121.
- [70] Marc Bonnet. Problèmes inverses. Master's thesis, Ecole Centrale de Paris, octobre 2008.
- [71] Y. Jarny and D. Maillet. Thermal Measurements and Inverse Techniques. In Ecole de printemps Métrologie Thermique et Techniques Inverses, volume 5. imprimerie centrale de l'université de Nantes, 2011.
- [72] D. Petit and D. Maillet. Techniques inverses et estimation de paramètres. Partie 2. Techniques de l'ingénieur, 2012.
- [73] M. Cuer. Des questions bien posées dans des problèmes inverses de gravimétrie et géomagnétisme : une nouvelle application de la programmation linéaire. PhD thesis, USTL Montpellier, France, 1984.
- [74] P. C. Sabatier. Problèmes inverses et applications. In Ecole d'été d'analyse numérique. CEA/INRIA/EDE, 1985.
- [75] J.-M. Reneaume. Notes de cours d'optimisation. ENSGTI, Octobre 2009.
- [76] A. Donev. Numerical Methods I : Singular Value Decomposition. Courant Institute, NYU, October 2010.
- [77] S. Gibout. Méthodes d'optimisation. Université de Pau et des Pays de l'Adour - UFR de Sciences et Techniques, 2005.
- [78] J.V. Beck and K.J. Arnold. Parameter estimation in Engineering and Science. PhD thesis, Wiley New York, 1977.
- [79] A. Tikhonov and V. Arsenin. Méthodes de résolution de problèmes mal posés. Mir, 1976.
- [80] A. W. Westerbg and S. W. Director. A modified least squares algorithm for solving sparse nxn sets of nonlinear equations. Compute and Chemical Engineering, 2 :77–81, 1978.
- [81] A. Ranganathan. The Levenberg-Marquardt Algorithm [online]. august 2013. URL : <http://www.ananth.in/Notes.html>.
- [82] L. Gosselin, M. Tye-Gingras, and F. Mathieu-Potvin. Review of utilization of genetic algorithms in heat transfer problems. International Journal of Heat and Mass Transfer, (52) :2169–2188, 2009.
- [83] E. Lutton. Algorithmes génétiques et algorithmes évolutionnaires. Techniques de l'ingénieur, 2006.

- [84] I. Ayadi, L. Bouillaut, and P. Aknin. Optimisation par algorithme génétique de la maintenance préventive dans un contexte de modélisation par modèles graphiques probabilistes. In Lambda Mu 17, 17ème Congrès de Maîtrise des Risques et de Sécurité de Fonctionnement, La Rochelle : France, 2010.
- [85] H. Nabli. An overview on the simplex algorithm. Applied Mathematics and Computation, (210) :479–489, 2009.
- [86] M. Baudin. Nedler-Mead User’s Manual. Technical report, Scialb, April 2010.
- [87] M. A. Luersen and R. Le Riche. Globalisation de l’Algorithme de Nedler-Mead : Application aux Composites. Technical report, INSA de Rouen - LMR - Laboratoire de Mécanique, Décembre 2001.
- [88] M. A. Luersen. GBNM : Un Algorithme d’Optimisation par Recherche Directe - Application à la Conception de Monopalmes de Nage. PhD thesis, Ecole Doctorale SPMI - INSA de Rouen, Décembre 2004.
- [89] W. G. Steele, R. A. Ferguson, R. P. Taylor, and H. W. Coleman. Comparaison of ANSI/ASME and ISO models for calculation of uncertainty. ISA Transactions, 33 :339–352, 1994.
- [90] R. H. Dieck. Measurement uncertainty models. ISA Transactions, 36(1) :29–35, 1997.
- [91] [online]April 2014. URL : <http://webbook.nist.gov/cgi/cbook.cgi?ID=65-85-0&Units=SI>.
- [92] K. Kaneko and T. Koyaguchi. Simultaneous crystallization and melting at both the roof and floor of crustal magma chambers : Experimental study using NH₄Cl-H₂O binary eutectic system. Journal of volcanology and geothermal research, (96) :161–174, 2000.
- [93] A. Jamil, T. Kousksou, Y. Zeraouli, S. Gibout, and J. P. Dumas. Simulation of the thermal transfert during an eutectic melting of a binary solution. Thermochimica Acta, (441) :30–34, 2006.
- [94] A. Jamil, T. Kousksou, Y. Zeraouli, and J. P. Dumas. Liquidus temperatures determination of the dispersed binary system. Thermochimica Acta, 471 :1–6, 2008.
- [95] Z. Younsi, L. Zalewski, D. R. Rousse, A. Joulin, and S. Lassue. Thermophysical characterization of phase change materials with heat flux sensors. In 5th European Thermal-Science Conference, 2008.
- [96] A. Joulin, Z. Younsi, L. Zalewski, S. Lassue, D. R. Rousse, and J.-P. Cavrot. Experimental and numerical investigation of a phase change material : Thermal-energy storage and release. Applied Energy, (88) :2454–2462, 2011.
- [97] S. Garcia. Experimental design optimization and thermophysical parameter estimation of composite materials using genetic algorithms. PhD thesis, Virginia Polytechnic Institute and State University and Université de Nantes, juin 1999.
- [98] H. Belghazi. Modélisation analytique du transfert instationnaire de la chaleur dans un matériau bicouche en contact imparfait et soumis à une source de chaleur en mouvement. PhD thesis, Université de Limoges, 2008.
- [99] A. Fuentes, P. Rogeon, P. Carré, T. Loulou, and A. Morançais. Modélisation du contact électrothermique - Conséquences sur la température de contact. SFT 2011, juin 2011.
- [100] V. Sobotka, N. Lefevre, Y. Jarny, and D. Delaunay. Inverse methodology to determine mold set-point temperature in resin transfer molding process. International Journal of Thermal Sciences, 49 :2138–2147, 2010.

- [101] E. Franquet, S. Gibout, P. Tittlein, L. Zalewski, and J. P. Dumas. Experimental and theoretical analysis of a cement mortar containing microencapsulated PCM. Applied Energy, 73(1) :30–38, 2014.
- [102] V. Sobotka, N. Lefevre, Y. Jarny, and D. Delaunay. Inverse methodology to determine mold set-point temperature in resin transfer molding process. International Journal of Thermal Sciences, 49 :2138–2147, 2010.
- [103] A. Joulin, L. Zalewski, S. Lassue, and H. Naji. Experimental investigation of thermal characteristics of a mortar with or without a micro-encapsulated phase change material. Applied Thermal Engineering, 66 :171–180, 2014.

Annexe A

Recherche de l'optimum géométrique pour la réduction de modèle

A.1 Cas du rapport minimal des surfaces 1D et 2D

Il apparait dans les résultats des tableaux 2.6, 2.7, 2.8 et 2.9 qu'il existe une géométrie pour laquelle la réduction de modèle ne modifie pas ou de manière négligeable les conductivités. Cette géométrie dépend de l'hétérogénéité des coefficients d'échanges équivalents entre les différentes surfaces. Si le paramètre α de la face supérieure est assez proche de la valeur de α sur les autres faces (tableau 2.8) alors cette géométrie optimale correspond à une géométrie intermédiaire entre le cylindre haut et le cylindre plat, il pourrait s'agir de la géométrie qui correspond à un minimum du rapport des surfaces 1D et 2D. Mais si la valeur du paramètre α de la face supérieure est très inférieur à celle de α sur les autres faces (tableau 2.7), la géométrie optimum est un cylindre plat. En effet le modèle 1D suppose un échange uniforme sur toute la surface, et le modèle cylindrique suppose une inhomogénéité dans les coefficients α_i . Alors pour que les conductivités ne soient pas être modifiées afin que le thermogramme 1D coïncide avec le thermogramme 2D, la surface supérieure doit être augmentée de sorte que l'échange autour du cylindre se rapproche du cas homogène. Comme on l'a évoqué précédemment, cet écart dans les coefficients α_i s'explique par la présence d'une couche d'air entre le dessus de l'échantillon et le creuset. Donc si le creuset est rempli au maximum ou si l'échange est amélioré au niveau du couvercle du creuset et si ce creuset a cette géométrie optimale, qui correspondrait au minimum du rapport des surfaces, l'inversion utilisant le modèle réduit donnera accès aux conductivités. Repérons ce minimum et vérifions notre hypothèse. Nous considérerons une non-uniformité de 60%.

Modèle 2D

$$V_{2D} = \pi r_{2D}^2 z \quad (\text{A.1})$$

$$\begin{aligned} S_{2D} &= 2\pi(r_{2D}^2 + r_{2D}z) \\ &= 2\pi(r_{2D}^2 + \frac{V_{2D}}{\pi r_{2D}}) \\ &= 2\pi(\frac{V_{2D}}{z} + \sqrt{\frac{V_{2D}z}{\pi}}) \end{aligned} \quad (\text{A.2})$$

Modèle 1D

$$V_{1D} = \frac{4}{3}\pi r_{1D}^3 \quad (\text{A.3})$$

$$\begin{aligned}
S_{1D} &= 4\pi \left(\frac{3V_{1D}}{4\pi} \right)^{\frac{2}{3}} \\
&= 4\pi r_{1D}^2
\end{aligned} \tag{A.4}$$

Comme on travaille en isovolume, $V_{2D} = V_{1D} = V$, écrivons le rapport des surfaces et cherchons son extremum.

$$\frac{S_{2D}}{S_{1D}} = \frac{r_{2D}^2 + \frac{V}{\pi r_{2D}}}{2 \left(\frac{3V}{4\pi} \right)^{\frac{2}{3}}} = \frac{\frac{V}{\pi z} + \sqrt{\frac{Vz}{\pi}}}{2 \left(\frac{3V}{4\pi} \right)^{\frac{2}{3}}} \tag{A.5}$$

$$\begin{aligned}
\left(\frac{d \frac{S_{2D}}{S_{1D}}}{dr_{2D}} \right)_{V=c^st} &= 0 \quad (=) \quad 2r_{2D} - \frac{V}{\pi r_{2D}^2} = 0 \Rightarrow r_{ext} = \sqrt[3]{\frac{V}{2\pi}} \\
&\Leftrightarrow \frac{-V}{\pi z^2} + \sqrt{\frac{V}{\pi}} \frac{1}{2\sqrt{z}} = 0 \\
&\Leftrightarrow \frac{2V}{\pi} \sqrt{z} = \sqrt{\frac{V}{\pi}} z^2 \\
&\Leftrightarrow 2\sqrt{\frac{V}{\pi}} = z^{\frac{3}{2}} \Rightarrow z_{ext} = \sqrt[3]{\frac{4V}{\pi}}
\end{aligned} \tag{A.6}$$

$$\pi r_{extremum}^2 z_{extremum} = \pi \frac{V^{\frac{2}{3}}}{(2\pi)^{\frac{2}{3}}} \left(\frac{4}{\pi} \right)^{\frac{1}{3}} V^{\frac{1}{3}} = V \tag{A.7}$$

Confirmons qu'il s'agit d'un minimum. Pour se faire, calculons la dérivée seconde en prenant comme variable unique r_{2D} .

$$r_{1D}^2 \left(\frac{d^2 \frac{S_{2D}}{S_{1D}}}{dr_{2D}^2} \right)_{r_{extremum}, V=c^st} = r_{1D}^2 \frac{2 + \frac{2V}{\pi r_{extremum}^3}}{2 \left(\frac{3V}{4\pi} \right)^{\frac{2}{3}}} = \frac{6r_{1D}^2}{2 \left(\frac{3V}{4\pi} \right)^{\frac{2}{3}}} = 3 > 0 \tag{A.8}$$

Il s'agit bien d'un minimum.

Nous avons fait le calcul, et il apparait que pour un rapport de surfaces minimales, les modèles 2D-cylindrique et 1D-sphérique produisent le même thermogramme pour des conductivités pratiquement identiques aux erreurs numériques près, tableau A.1.

$$\begin{cases} r_{min} &= \sqrt[3]{\frac{V}{2\pi}} \\ z_{min} &= \sqrt[3]{\frac{4V}{\pi}} \\ \left(\frac{S_{2D}}{S_{1D}} \right)_{min} &= \frac{1+8^{\frac{1}{3}}}{2 \left(\frac{9}{4} \right)^{\frac{1}{3}}} = 1,144714243... \end{cases} \tag{A.9}$$

Vous noterez qu'un tel cylindre est celui où la hauteur est égale au diamètre.

Tableau A.1 – Rapport des conductivités pour une géométrie minimale.

Matériaux	$\left(\frac{\lambda_{1D}}{\lambda_{2D}} \right)_S$			$\left(\frac{\lambda_{1D}}{\lambda_{2D}} \right)_L$		
$\beta \text{ (K} \cdot \text{min}^{-1})$	2	5	10	2	5	10
eau	1,0039	1,0002	1,0010	1,0038	1,0003	1,0010
$H_2O/NH_4Cl : 0,57\%$	1,0002	0,9913	1,0000	1,0105	1,015	0,9999
$H_2O/NH_4Cl : 2,5\%$	1,003	1,0003	1,0000	0,9849	1,009	1,002
$H_2O/NH_4Cl : 5,0\%$	1,0006	1,038	1,010	1,010	0,9735	1,010
$H_2O/NH_4Cl : 10\%$	0,9997	1,010	1,0045	1,017	1,004	1,0087

A.2 Conclusion

On a montré qu'il existait une géométrie cylindrique pour laquelle la réduction au modèle sphérique n'aurait presque aucune influence sur les conductivités si l'échange était suffisamment uniforme sur toute la surface, ou qu'il n'y ait pas d'importante couche d'air au-dessus de l'échantillon. Dans ces conditions le modèle réduit identifierait quasiment la vraie conductivité. Le calorimètre fixant le rayon, d'après (A.9) on peut donc déterminer la quantité avec laquelle il faut remplir l'échantillon pour que l'expérience de calorimétrie soit faite avec un cylindre que l'on dit minimal. On se retrouve avec la difficulté de remplissage de l'échantillon, ensuite les cellules de calorimétrie ne sont pas forcément adaptées. Par exemple pour notre calorimètre de Perkin-Elmer, les cellules ont un rayon de 2,15 mm et sont faites pour contenir une dizaine de mg. Pour obtenir cette géométrie minimale avec ce rayon il faudrait 62,4 mg, on dépasse les capacités de la cellule. Par contre le calorimètre SETARAM DSC 131 qui, lui, peut contenir une telle masse devrait convenir. Malheureusement nous ne disposons pas de cellules « sur-mesures » qui nous auraient permis de vérifier expérimentalement notre idée.

Abstract

With the development of intermittent sources of energy and the depletion of fossil fuels, the subject of energy storage is becoming an important topic. One of the studied options is the latent thermal storage using of phase change materials (PCM). One application for this type of energy storage is to improve the thermal insulation in buildings. To make the best use of these materials it is necessary to be able to predict their energy behavior. This requires a precise knowledge of their thermophysical properties, first of all of the specific enthalpy function of the material $h(T)$.

Currently, it is often suggested to approximate the enthalpy by the direct integration of the thermograms obtained through calorimetry experiments (notion of « equivalent » calorific capacity). This approach is false because thermograms are only a time related representation of complex phenomena where thermal transfers arise in the cell of the calorimeter acting with the thermophysical properties. As a result, for example, the shape of thermograms depends on the heating rate and on the mass of the sample, which is not the case for the enthalpy of the PCM, which depends, at constant pressure, only on the temperature or on the concentration (for the solutions).

We propose to compare the results given by a numerical direct model with experimental thermograms. The main objective in this thesis is then to use this direct model in an inverse method in order to identify the parameters of the equation of state, which enables us to calculate the specific enthalpy $h(T)$.

First of all, the detail of an enthalpy model is presented, and then validated by comparison with experiments, allowing us to reconstruct the thermograms of pure substances or of salt solutions, of which the enthalpies are known. A study of the influence of the various parameters (c_p , λ , T_M , L_M ...) on the shape of thermograms is also undertaken in order to deduce their sensibilities. A reduced model is then developed in order to reduce the calculating time of the direct model. This optimized model allows the use of inverse methods with acceptable durations.

Several inverses algorithms are then presented : Levenberg-Marquardt, evolutionary and Simplex which has proved to be the fastest).

We shall then apply this algorithm to identify, from calorimetric experiments, the enthalpy function of pure substances or of salt solutions. The results that we obtain show that it is possible to identify a function $h(T)$ independent of the heating rate and of the mass, which validates the method. An analysis of the various sources of errors in the identification process and of their influences on the result allows us to estimate the quality of the enthalpy function that we identify.

Finally, the same approach has been used to analyze an experiment on a composite material that is used in buildings (cement with inclusion of PCM micro-encapsulated). Again in this case, our methods lead to a relevant energetic characterization.

Key words : phase change materials, melting, characterization, modelization, inverse method, experimental validation

Résumé

Avec le développement des énergies intermittentes et la raréfaction des énergies fossiles, le sujet du stockage de l'énergie prend de plus en plus d'ampleur. Une des voies étudiée est le stockage thermique par utilisation de matériaux à changement de phase (MCP). Cette voie est en outre développée pour améliorer l'inertie thermique dans le secteur du bâtiment. Pour utiliser au mieux ces matériaux il est nécessaire de pouvoir prévoir leur comportement énergétique. Cela nécessite une connaissance précise des propriétés thermophysiques, et en premier lieu de la fonction enthalpie massique $h(T)$.

Actuellement, il est souvent proposé d'approximer cette enthalpie par l'intégration directe des thermogrammes de la calorimétrie utilisant notamment la notion de capacité calorifique « équivalente ». Cette approche est cependant fautive car le thermogramme n'est qu'une représentation en fonction du temps de phénomènes complexes faisant intervenir non seulement les propriétés énergétique du matériaux mais également les transferts thermiques au sein de la cellule du calorimètre. Il en résulte, par exemple, que la forme des thermogrammes, et donc l'enthalpie apparente, dépend de la vitesse de réchauffement et de la masse de l'échantillon ce qui n'est pas le cas de l'enthalpie des MCP qui ne dépend, à pression fixe, que de la température ou de la concentration (pour les solutions).

On propose de comparer la sortie d'un modèle numérique direct avec des thermogrammes expérimentaux. L'objectif principal de cette thèse est alors d'utiliser ce modèle dans le cadre d'une méthode inverse permettant l'identification des paramètres de l'équation d'état permettant alors de calculer l'enthalpie massique $h(T)$.

Dans un premier temps, il est donc présenté le détail d'un modèle 2D dit enthalpique qui néglige la convection, validé par l'expérience, permettant de reconstituer les thermogrammes de corps purs ou de solutions binaires dont les enthalpies sont connues. Il en est déduit une étude de l'influence des différents paramètres (c_p , λ , T_M , L_M ...) sur la forme des thermogrammes pour en déduire leurs sensibilités. Une réduction de ce modèle est ensuite effectuée pour réduire le temps de calcul du modèle direct en vue de l'utilisation dans une méthode inverse.

Cette dernière est décrite ainsi que les algorithmes d'optimisation correspondants (de Levenberg-Marquardt, génétique ou du simplexe qui s'est avéré le plus rapide) dans un second temps.

Nous appliquerons ensuite cet algorithme pour identifier, à partir d'expériences, la fonction enthalpie de corps purs ou de solutions binaires. Les résultats obtenus montrent qu'il est possible d'identifier une fonction $h(T)$ indépendante de la vitesse de réchauffement et de la masse, ce qui valide la méthode. Une analyse des différentes sources d'erreurs dans le processus d'identification et leurs influences sur le résultat permet d'évaluer la qualité de la fonction enthalpie que l'on identifie.

Enfin, cette même approche a été utilisée pour analyser une expérience réalisée sur un échantillon d'un matériau composite utilisé dans le bâtiment (ciment avec inclusion de MCP micro-encapsulé). Dans ce cas encore, nos méthodes permettent une caractérisation énergétique pertinente.

Mots clés : matériaux à changement de phase, fusion, caractérisation, modélisation, méthode inverse, validation expérimentale