

Table des matières

Résumé	10
Abstract	11
Notations	12
I Motivations et plan détaillé de la thèse	14
I. 1 Motivations industrielles	14
I. 1.1 Contexte du nucléaire civil pour la production d'électricité	14
I. 1.2 Le contexte accident grave pour les réacteurs nucléaires	15
I. 1.3 Besoins numériques et méthodologiques pour la simulation des accidents graves	16
I. 2 Contenu de la thèse	17
I. 2.1 Chapitre contexte scientifique	17
I. 2.2 Chapitre modèle linéaire dans le contexte bayésien	18
I. 2.3 Chapitre reconstruction non-paramétrique par la maximisation de γ -entropie	19
I. 3 Perspectives dans le cadre des applications en thermodynamique chimique	20
II Contexte scientifique	22
II. 1 L'apprentissage statistique et l'analyse d'incertitude	22
II. 1.1 Généralités sur l'apprentissage statistique	22
II. 1.2 Analyse de l'incertitude	25
II. 1.2.1 Quantification de l'incertitude d'entrée	26
II. 1.2.2 Propagation de l'incertitude	27
II. 1.2.3 Analyse de sensibilité	28
II. 2 Description des systèmes chimiques	28
II. 3 La méthode Calphad	30
II. 4 Les problèmes du calcul thermodynamique	32
II. 4.1 Le problème d'équilibre	32
II. 4.2 Le problème de reconstruction fonctionnelle avec le modèle linéaire	33
II. 4.3 Les données d'assimilation	36
II. 4.3.1 Données de structure	37

II. 4.3.2 Données thermodynamiques	38
II. 4.3.3 Données de diagramme de phase	39
II. 5 Modélisation du système Cu-Mg dans la formulation classique de thermodynamique	40
II. 5.1 Description du système	41
II. 5.2 Modélisation classique en thermodynamique	43
II. 5.3 Données d'assimilation	45
III The linear model in the Bayesian framework	47
III. 1 General results with the linear model	47
III. 2 The Bayesian framework	49
III. 2.1 General definitions and properties in the Bayesian framework	49
III. 2.2 Conjugate prior distribution for the Gaussian linear model	50
III. 2.3 Sensitivity to the prior parameters	54
III. 2.4 Levelled linear models	54
III. 3 Bayesian point of view for the Calphad-based modelling	56
III. 3.1 Motivations	56
III. 3.2 Calphad-based modelling	57
III. 3.3 Literature review of the uncertainty analysis for the specific framework of computational thermodynamics	59
III. 3.4 Uncertainty for the thermodynamic quantities	63
III. 3.5 Uncertainty for the phase diagram	64
III. 4 Application to the binary system Cu-Mg case study	64
III. 4.1 Case study modelling	64
III. 4.2 Results	67
IV Non parametric reconstruction with linear constraints : the γ-entropy maxi- misation problem	72
IV. 1 Introduction	72
IV. 2 The specific γ -entropy maximisation problem for a single reconstruction	79
IV. 2.1 Transfer principle	80
IV. 2.2 γ -divergence minimisation problem	81
IV. 2.3 Maximum entropy problem	83
IV. 2.4 Maximum Entropy on the Mean	86
IV. 2.5 Connection with classical minimisation problems	89
IV. 3 The γ -entropy maximisation problem for the reconstruction of a multidimensional function	91
IV. 3.1 The γ -entropy maximisation problem for the multidimensional case in the convex analysis framework	92
IV. 3.2 Linear transfer principle for the multidimensional case	101
IV. 3.3 The embedding into the MEM framework for the multidimensional case .	103

IV. 4 Applications	105
IV. 4.1 Application in the case $p = 1$	105
IV. 4.1.1 Reconstruction of an univariate convex function	106
IV. 4.1.2 Reconstruction of a bivariate polynomial function	107
IV. 4.2 Applications in the case $p \neq 1$	110
IV. 4.2.1 Phase diagram description in Computational Thermodynamics .	110
IV. 4.2.2 Phase diagram with an ideal solution	112
 Annexes	 114
A Bibliographic references for the thermodynamic data	114
B Classical algorithms of Bayesian estimation	116
C Some general definitions and properties	118
D Multidimensional Bregman distance	120
D. 1 Bregman distance bounds	120
D. 2 Cauchy result for Bregman distance minimising sequence	122
 Bibliography	 124

Résumé

Cette thèse s'inscrit dans les domaines de l'apprentissage statistique et de l'analyse d'incertitude dans le cadre d'un problème de thermodynamique chimique. On s'intéresse à deux problèmes liés à la reconstruction d'une fonction f multidimensionnelle : en premier lieu, un problème de propagation de l'incertitude dans le cadre du modèle linéaire bayésien ; en second lieu, la reconstruction non paramétrique d'une fonction par maximisation d'un critère d'analyse convexe.

La spécificité du problème de thermodynamique chimique considéré dans cette thèse réside dans la structure compliquée des données : les données d'assimilation ne sont pas des observations directes de la quantité f à reconstruire. Dans le cadre du modèle linéaire bayésien, on propose une méthode permettant de prendre en compte la spécificité de ce modèle. On considère un cas de thermodynamique chimique réel pour l'application de la méthode.

Dans le cadre de la reconstruction non-paramétrique, on s'intéresse à la reconstruction d'une fonction f multidimensionnelle sur un compact U et telle que f satisfait un certain nombre de contraintes intégrales très générales. On se propose de résoudre ce problème de reconstruction fonctionnelle dans le cadre de la maximisation de la γ -entropie sous contraintes. A savoir, on se donne une fonction convexe γ de $\mathbb{R}^p$ dans $\mathbb{R}_+$, munie de bonnes propriétés et on s'intéresse à la maximisation sous contraintes de la quantité $I_\gamma(f) = - \int_U \gamma(f) dP_U$. On expliquera que ce problème peut être rapproché d'un autre problème connu portant cette fois-ci sur des mesures signées F . Ces problèmes ont été étudiés dans le cas d'une unique reconstruction, à savoir une unique fonction ou une unique mesure. Nous nous proposons d'étudier le cas plus général d'une fonction ou d'une mesure à valeurs dans $\mathbb{R}^p$.

Mots-clefs : Problèmes de maximisation de l'entropie, Statistique bayésienne, Modèle linéaire, Propagation de l'incertitude, Application en thermodynamique, Méthode Calphad.

Abstract

The topics addressed in this thesis lie in statistical learning and uncertainty analysis within the frame of a thermodynamic problem. We are interested in two problems linked with the reconstruction of a multidimensional function f : an uncertainty propagation problem within the Bayesian linear model in the first place ; a non parametric reconstruction problem using a convex analysis maximisation criterion in the second place.

The specificity of the thermodynamic problem considered in this manuscript consists in the complicated structure of the data : the assimilation data are not direct observations of the quantity f to reconstruct. In the framework of the Bayesian linear model, we propose a method allowing us to take into account the thermodynamic model specificity. We consider a real case study for the application of the method.

In the framework of non-parametric reconstruction problems, we are interested in the reconstruction of a multidimensional function f over a compact set U and such that f satisfies a certain amount of very general integral constraints. We propose to solve this reconstruction problem by setting it in the frame of the γ -entropy maximisation problem under constraints. That is, we define a convex function γ of $\mathbb{R}^p$ to $\mathbb{R}_+$ with some good properties and we are interested in the maximisation under constraints of the quantity $I_\gamma(f) = - \int_U \gamma(f) dP_U$. We explain that this problem can be linked to another one that deals with signed measures F . Such problems have been studied in the case of a single function or a single measure reconstruction. We propose to study the more general case of the reconstruction of a function or a measure with values in $\mathbb{R}^p$.

Key-words : Entropy maximisation problems, Bayesian statistics, Linear model, Uncertainty propagation, Application in thermodynamics, Calphad method.

Notations

Version française

I	Nombre d'éléments chimiques.
K	Nombre d'espèces chimiques.
P	Pression.
T	Température.
$\mathbf{N} = (N_1, \dots, N_I)$	Composition chimique, donnée en nombre de mole.
$\mathbf{E} = (P, T, \mathbf{N})$	Condition expérimentale, les valeurs sont dans un ensemble compact $\mathcal{K} \subset \mathbb{R}_+^{I+2}$.
$\boldsymbol{\lambda} = (\lambda^{(1)}, \dots, \lambda^{(r)})$	Proportions de phase.
$\boldsymbol{\gamma} = (\gamma^{(1)}, \dots, \gamma^{(r)})$	Occupation des sites cristallographiques. Chaque phase α contient un nombre s_α de sous-réseaux. $\boldsymbol{\gamma}^{(\alpha)} = (\gamma_1^{(\alpha,1)}, \dots, \gamma_K^{(\alpha,1)}, \dots, \gamma_1^{(\alpha,s_\alpha)}, \dots, \gamma_K^{(\alpha,s_\alpha)})$ $\gamma_k^{(\alpha,s)} \in [0; 1], \forall \alpha, s, k$
$G_{\mathbf{E}}$	Énergie totale pour une condition $\mathbf{E}$ donnée.
$G_{\mathbf{E}}^{(\alpha)}$	Fonction d'énergie libre de Gibbs de la phase α pour une condition $\mathbf{E}$ donnée.
Z	Valeur d'énergie, utilisée comme sortie dans le modèle linéaire.
$a^{\alpha,s}$	Nombre total de sites cristallographiques pour le sous-réseau s de la phase α , i.e. la quantité d'espèces chimiques que peut contenir une unité du sous-réseau s .
$b_{i,k}$	Correspondance entre l'espèce chimique k et l'élément chimique i , e.g. pour les molécules de dioxygène O_2 , $b_{O,O_2} = 2$.
$M_i^{(\alpha)}$	Nombre de mole de l'élément i en phase α .
ν_k	Nombre relatif de charges pour l'espèce chimique k , $\nu_k \in \mathbb{Z}$.

Remarque : Les matrices sont notées à l'aide d'une notation en gras (e.g. matrice $\mathbf{A}$) et les vecteurs à l'aide d'une notation en gras italique (e.g. vecteur $\boldsymbol{x}$).

Notations

English version

I	Number of chemical elements.
K	Number of chemical species.
P	Pressure.
T	Temperature.
$\mathbf{N} = (N_1, \dots, N_I)$	Chemical compositions, in number of mole.
$\mathbf{E} = (P, T, \mathbf{N})$	Experimental conditions, vector with values in a compact set $\mathcal{K} \subset \mathbb{R}_+^{I+2}$.
$\boldsymbol{\lambda} = (\lambda^{(1)}, \dots, \lambda^{(r)})$	Phase proportions.
$\boldsymbol{\gamma} = (\gamma^{(1)}, \dots, \gamma^{(r)})$	Site fraction occupancy. Each phase α contains s_α sublattices. $\boldsymbol{\gamma}^{(\alpha)} = (\gamma_1^{(\alpha,1)}, \dots, \gamma_K^{(\alpha,1)}, \dots, \gamma_1^{(\alpha,s_\alpha)}, \dots, \gamma_K^{(\alpha,s_\alpha)})$ $\gamma_k^{(\alpha,s)} \in [0; 1], \forall \alpha, s, k$
$G_{\mathbf{E}}$	Total energy of a system for a given $\mathbf{E}$.
$G_{\mathbf{E}}^{(\alpha)}$	Gibbs free energy function of the phase α for a given $\mathbf{E}$.
Z	Energy value, used as the output of the linear model.
$a^{\alpha,s}$	Total number of site fractions for the sublattice s of the phase α by formula unit, i.e. the number of chemical species that can be contained in one unit of sublattice s .
$b_{i,k}$	The correspondence between chemical specie and the chemical element, e.g. for dioxygen molecules O_2 , $b_{O,O_2} = 2$.
$M_i^{(\alpha)}$	Number of mole of element i in phase α .
ν_k	The relative number of charges for chemical specie k , $\nu_k \in \mathbb{Z}$.

Remark : Matrices are denoted by a bold notation (e.g. matrix $\mathbf{A}$) and vectors by an italic bold notation (e.g. vector $\mathbf{x}$).

Chapitre I

Motivations et plan détaillé de la thèse

I. 1 Motivations industrielles

I. 1.1 Contexte du nucléaire civil pour la production d'électricité

Dès le début des années 50, l'énergie nucléaire en France connaît un développement considérable. En effet, la nécessité d'importer le pétrole et les crises pétrolières des années 70 ont abouti au développement important de la filière nucléaire pour la production d'électricité en France, marquant un fort contraste avec ses voisins européens.

De manière simplifiée, une centrale nucléaire utilise la chaleur produite par la réaction de fission nucléaire pour produire de l'électricité. La centrale est découpée en trois circuits indépendants, voir la Figure I.1 :

- . Le circuit primaire est représenté en jaune. Son rôle est d'assurer le transfert de la chaleur produite par la réaction de fission jusqu'au circuit secondaire. Le fluide qui permet la transmission de la chaleur est appelé le caloporteur.
Ce premier circuit est en contact avec la cuve du réacteur, à l'intérieur duquel a lieu la réaction de fission des matériaux combustibles sous l'effet du bombardement de neutrons. Un modérateur est utilisé suivant les technologies du réacteur. Son rôle est de contrôler la réaction de fission en chaîne en absorbant les neutrons.
- . Le circuit secondaire (bleu) est celui en charge de la production d'électricité.
Au contact du circuit primaire, de la vapeur se crée. Cette vapeur met en marche une turbine et un alternateur qui permettent la production d'électricité.
- . Enfin le circuit tertiaire (vert) est en charge du refroidissement du circuit secondaire.

La variété des matériaux et technologies utilisés, entre autres, pour les barres de combustible, les modérateurs et le caloporteur fait la spécificité de chaque filière nucléaire.

Les accidents dramatiques de Tchernobyl (1986) et Fukushima (2011), tous deux survenus

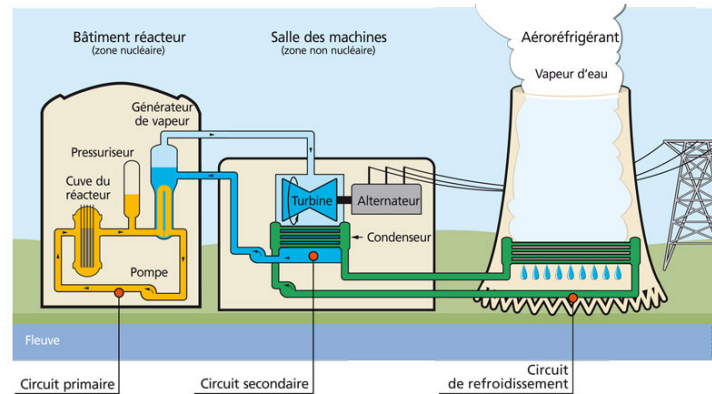


FIGURE I.1 – Principe de fonctionnement d’une centrale nucléaire. Référence figure [41].

à la suite d’une perte de refroidissement, motivent l’intégration de l’étude des données incertaines intervenant dans la simulation des accidents graves. Cette intégration est particulièrement essentielle pour la conception de nouveaux réacteurs.

I. 1.2 Le contexte accident grave pour les réacteurs nucléaires

La modélisation des accidents graves dans le contexte nucléaire s’intéresse au comportement des produits issus de la fission du combustible lorsque ceux-ci entrent en contact avec les matériaux environnants. Un tel contact survient après rupture de l’intégrité de la cuve du réacteur dont le rôle est de séparer le combustible des autres matériaux. Cette rupture a lieu suite à l’emballement de la réaction en chaîne se produisant au sein du réacteur, ce qui produit l’échauffement du cœur du réacteur.

Le contexte accident grave correspond aux conséquences associées à une perte prolongée du système de refroidissement (cf. Figure I.2) donnant lieu à une augmentation rapide et brutale des températures dans le cœur du réacteur. Cette hausse brutale conduit alors à la fonte de la gaine absorbante qui entoure et protège les matériaux combustibles. Si le système de refroidissement n’est pas rétabli au plus vite, les scénarios suivants sont à même de se produire :

1. Le cœur du réacteur nucléaire s’échauffe jusqu’à fonte de ses composants. Le corium, un magma métallique et minéral constitué des éléments fondus du cœur et des minéraux absorbés lors de son trajet, se forme alors.
2. Ce même corium peut sortir du cœur du réacteur, se relocaliser dans le fond de la cuve et interagir avec les structures métalliques la constituant.
3. Les interactions corium/métal peuvent amener à une rupture de la cuve ce qui amène le corium en contact du caloporteur du circuit primaire.
4. Le corium peut également s’étaler et interagir avec le béton formant le puit de la cuve.

Une fois encore, la spécificité de chaque filière nucléaire rend également spécifique l’étude des accidents graves. Nous renvoyons le lecteur intéressé vers [93, Chapitre 1] où l’auteur fournit une

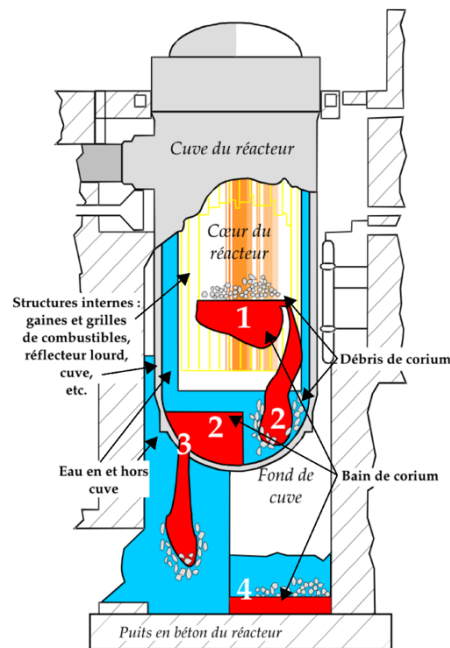


FIGURE I.2 – Les différents scénarios envisagés lors de la perte de refroidissement.

description détaillée du déroulement d'un accident grave pour un type de réacteur de deuxième génération, à savoir la génération actuelle de réacteurs dans le parc électrique français.

I. 1.3 Besoins numériques et méthodologiques pour la simulation des accidents graves

De manière évidente, l'étude des accidents graves ne peut se faire qu'au travers de la modélisation des phénomènes physiques, intervenant dans ce contexte, et de la simulation numérique. La problématique des accidents graves nécessite une simulation fiable et précise de résultats physiques.

Pour l'étude des accidents au sein des réacteurs nucléaires, le comportement thermodynamique des matériaux, c'est-à-dire les énergies des associations d'éléments chimiques constituant ces matériaux, a besoin d'être représenté par un modèle mathématique f . Ce modèle d'énergie est ensuite intégré dans un code multiphysique dont le rôle est de simuler le comportement du réacteur, ou d'une partie du réacteur, au cours du temps.

Le modèle mathématique des énergies f étant inconnu, le premier besoin dans ce contexte est la reconstruction d'un modèle approché, en adéquation avec une base de données, hors accident graves. L'approche courante de ce problème de reconstruction en thermodynamique est d'utiliser la méthode Calphad. La méthode Calphad consiste en la reconstruction des fonctions d'énergie grâce à des données de différentes natures. La description donnée par la méthode Calphad rentre dans le cadre de la thermodynamique et ne peut s'appliquer qu'aux systèmes dits *à l'équilibre*.

Les systèmes décrits sont donc à la fois en équilibre chimique, mécanique et thermique : les quantités des différentes espèces chimiques restent constantes, le système ne subit aucun effort et n'est pas le lieu d'un échange de chaleur avec l'extérieur.

La précision et la justesse de la reconstruction des fonctions d'énergie dépend de la quantité d'informations données, à savoir le nombre de points expérimentaux considérés pour réaliser l'estimation des fonctions d'énergie, ainsi que de la qualité des calculs et des mesures réalisés pour obtenir ces points expérimentaux, que l'on nommera indistinctement incertitude de mesure par la suite.

Le deuxième besoin est donc de proposer une méthode pour prendre en compte cette incertitude. Cette mesure de l'incertitude doit tout d'abord être possible au niveau de la sortie du modèle f , mais également être facile à propager aux éléments intervenant dans le couplage des fonctions d'énergie avec la simulation accident grave.

I. 2 Contenu de la thèse

Plus formellement d'un point de vue mathématique, on s'intéresse dans cette thèse à deux problèmes.

En premier lieu, on s'intéresse à la reconstruction d'une fonction f à valeurs dans $\mathbb{R}^r$ où $r > 1$ pour laquelle les données d'assimilation ne sont pas directement observées.

Le second problème abordé dans cette thèse est un problème de propagation de l'incertitude. On s'intéresse à quantifier la variabilité de la fonction f lorsque les données ayant servi à sa reconstruction sont bruitées.

Nous détaillons ci-après les éléments contenus dans chaque chapitre de la thèse.

I. 2.1 Chapitre contexte scientifique

Le Chapitre II propose dans un premier temps une introduction générale sur l'apprentissage statistique et les questions soulevées par le traitement de l'incertitude pour un modèle mathématique donné.

Nous spécifions, dans un second temps, la modélisation classique des systèmes chimiques pour l'étude thermodynamique des matériaux. Nous présentons la quantité d'intérêt à reconstruire, à savoir une fonction d'énergie f multidimensionnelle, et les données d'assimilation utilisées dans ce problème.

La spécificité du modèle thermodynamique est la diversité de natures des données d'assimilation utilisées pour la reconstruction de f . Nous formalisons les liens qu'il existe entre ces données d'assimilation et la quantité à reconstruire.

La méthode Calphad est une méthode de reconstruction fonctionnelle élaborée pour les problèmes de reconstructions de thermodynamique chimique. Cette méthode est conçue pour intégrer des données de différentes natures pour la reconstruction de la quantité d'intérêt f . Nous

présentons dans ce chapitre la méthode Calphad en spécifiant les problèmes d'optimisation mis en jeu dans ce contexte.

Enfin, nous présentons un cas concret de système chimique qui sera repris dans le Chapitre III pour l'application de la méthodologie de propagation des incertitudes.

I. 2.2 Chapitre modèle linéaire dans le contexte bayésien

Dans le Chapitre III, nous donnons dans un premier temps des rappels sur le modèle linéaire dans le cas d'utilisation de régresseurs très généraux, pris dans une base de fonctions bien choisies. Plus précisément, on s'intéresse au modèle linéaire

$$\mathbf{Z} = F(\mathbf{X})\boldsymbol{\theta} + \boldsymbol{\varepsilon}$$

où $\mathbf{Z}$ est un vecteur de N observations, $\boldsymbol{\varepsilon}$ est un vecteur gaussien centré, $F(\mathbf{X})$ est une base de p fonctions évaluées aux points de la matrice de design $\mathbf{X}$ et $\boldsymbol{\theta}$ est un paramètre inconnu de dimension p à déterminer. La matrice $\mathbf{X}$ de taille $N \times d$ contient pour chaque ligne les conditions pour l'observation de la quantité Z_j .

Dans le cadre du problème de propagation des incertitudes, on s'intéresse au modèle linéaire dans un contexte bayésien. Nous faisons des rappels généraux de définitions et de propriétés dans le cadre de la statistique bayésienne.

On s'intéresse ensuite au modèle Calphad dans le contexte du modèle linéaire bayésien. On cherche alors à reconstruire simultanément r fonctions d'énergie. On modélise chaque fonction d'énergie $f^{(\alpha)}$ par un modèle linéaire. Ces modèles sont de la forme suivante

$$f^{(\alpha)}(\mathbf{x}) = \sum_{m=1}^{p_\alpha} \theta_m^{(\alpha)} f_m^{(\alpha)}(\mathbf{x}) + B^{(\alpha)}(\mathbf{x}), \quad \forall \alpha = 1, \dots, r \quad (\text{I.1})$$

où $\mathbf{x} \in U \subset \mathbb{R}^d$ avec U compact, est une donnée expérimentale dans la modélisation de thermodynamique chimique. Nous détaillerons dans le Chapitre II les composantes du vecteur d'entrée $\mathbf{x}$, voir plus précisément la Section II. 2 et la Section II. 4. La fonction $B^{(\alpha)}$ représente une fonction connue de $\mathbf{x}$. Les fonctions $f_m^{(\alpha)}$ représentent les fonctions de régression.

Le choix du modèle linéaire pour la modélisation Calphad est conforté par plusieurs raisons.

Tout d'abord par rapport au choix de l'étude de l'incertitude par l'approche bayésienne, en choisissant une fonction de vraisemblance gaussienne et une loi *a priori* gaussienne sur le paramètre inconnu $\boldsymbol{\theta}$, on obtient une loi *a posteriori* gaussienne.

Enfin, une partie des données d'assimilation - qui seront par la suite nommées les *données thermodynamiques* - sont reliées à la quantité à reconstruire $f^{(\alpha)}$ par une transformation linéaire. C'est-à-dire, en notant à l'aide de $Z^{(\alpha)}$ une observation de type *données thermodynamiques* pour α donné et une valeur donnée de $(\mathbf{x})$, on peut noter pour l'observation j

$$Z_j^{(\alpha)} = \mathcal{T}_{j,\alpha} \left(f^{(\alpha)}(\mathbf{x}_j) \right) + \varepsilon_j$$

où ε_j représente l'incertitude de mesure sur la donnée $Z_j^{(\alpha)}$ et où $\mathcal{T}_{j,\alpha}$ est une transformation linéaire de la quantité d'intérêt $f^{(\alpha)}$.

La méthode de propagation de l'incertitude est appliquée au système chimique Cu-Mg. On s'intéresse au problème de la reconstruction des fonctions d'énergie des alliages formés à partir des éléments cuivre (Cu) et magnésium (Mg).

La partie application pour le modèle Calphad de ce Chapitre III fait l'objet d'un article en cours de rédaction pour le journal de thermodynamique chimique Calphad.

I. 2.3 Chapitre reconstruction non-paramétrique par la maximisation de γ -entropie

Dans le Chapitre IV, on s'intéresse à la reconstruction non paramétrique d'une fonction f à valeur dans $\mathbb{R}^r$, définie pour tout x dans un ensemble compact $U \subset \mathbb{R}^d$. On suppose que l'intérieur de U est non-vide et on note $f(x) = (f^1(x), \dots, f^r(x))$. On place notre problème de reconstruction dans un espace de probabilité $(U, \mathcal{B}(U), P_U)$ où $\mathcal{B}(U)$ sont les boréliens de U et P_U est une mesure de référence donnée. Dans ce cadre, on souhaite reconstruire f sur U telle que f satisfait les N contraintes intégrales suivantes

$$\int_U \sum_{i=1}^r \lambda^i(x) f^i(x) d\Phi_l(x) = z_l \quad 1 \leq l \leq N, \quad (\text{I.2})$$

où les Φ_l sont N mesures finies positives (connues) définies sur $(U, \mathcal{B}(U))$ et les λ^i sont des fonctions de poids continues, connues.

On se propose de résoudre ce problème de reconstruction fonctionnelle dans le cadre de la maximisation de la γ -entropie sous des contraintes linéaires. A savoir, on se donne une fonction convexe γ de $\mathbb{R}^r$ dans $\mathbb{R}_+$, munie de bonnes propriétés que nous détaillerons dans le Chapitre IV. On définit la γ -entropie par

$$I_\gamma(f) = - \int_U \gamma(f^1(x), \dots, f^r(x)) dP_U(x).$$

On appelle alors le problème de la maximisation de la γ -entropie I_γ sous contraintes linéaires le problème suivant

$$\begin{aligned} & \max I_\gamma(f) \\ & f : \int_U \sum_{i=1}^r \lambda^i(x) f^i(x) d\Phi(x) \in \mathcal{K}. \end{aligned}$$

On expliquera que ce problème peut être rapproché d'un autre problème connu portant cette fois-ci sur des mesures signées F , à savoir celui de la minimisation de la γ -divergence de F par rapport à une mesure de référence P_V (différente de P_U). On peut revenir au problème de départ en exprimant une régularisation bien choisie $\mathcal{T}$ de F et en s'intéressant à $f = \mathcal{T}(F)$.

Ces problèmes ont été étudiés dans le cas d'une unique reconstruction, à savoir une unique fonction ou une unique mesure. Nous nous proposons d'étudier le cas plus général d'une fonction ou d'une mesure à valeurs dans $\mathbb{R}^r$.

Ce chapitre remplit deux fonctions.

Il propose dans un premier temps un résumé détaillé des études portant sur les problèmes de minimisation de la γ divergence lorsque la quantité à reconstruire est une unique mesure. Par extension, on s'intéresse également aux travaux des auteurs ayant traité le cas où la quantité à reconstruire est une fonction dans $\mathbb{R}$.

Dans un second temps, on s'intéresse au cas de la reconstruction multidimensionnelle. On propose une méthode générale de reconstruction basée sur la maximisation de la γ -entropie pour des fonctions soumises à des contraintes intégrales très générales.

Ce chapitre a fait l'objet d'un article qui a été proposé au journal *RAIRO Operation Research*, [37].

I. 3 Perspectives dans le cadre des applications en thermodynamique chimique

Actuellement, le parc électrique français issu de la filière nucléaire est principalement constitué de réacteurs dits de deuxième génération. La troisième génération de réacteurs est en cours de déploiement en vue de remplacer ces derniers. Les besoins croissants en énergie et les retours d'expérience, notamment dans le cadre de la sûreté nucléaire, justifient l'investissement dans le développement d'améliorations majeures. A cette fin, les réacteurs nucléaires dits de quatrième générations sont à l'étude. Parmi eux, la technologie des réacteurs à neutrons rapides à caloporteur sodium (RNR-Na) constitue un développement envisagé. Au CEA, le projet *SFRAG* regroupe les travaux sur l'étude des accidents grave pour les réacteurs de quatrième génération.

Les principales différences de cette nouvelle filière nucléaire avec la filière actuelle résident en deux points principaux. Premièrement, les matériaux combustibles prévus pour la fission nucléaire sont plus lourds que ceux utilisés dans la filière actuelle. Le combustible de cette nouvelle filière consiste en de l'uranium très enrichi ou du plutonium - qui est actuellement un produit de fission des réacteurs de seconde génération et qui est jusqu'à présent considéré comme un déchet. L'utilisation des produits de fission en tant que matériaux combustibles constitue le principal atout offert par la technologie RNR-Na, permettant, en plus du recyclage des produits de fission, de diminuer la consommation d'uranium naturel. Deuxièmement, à la différence des réacteurs décrits précédemment, le sodium joue le rôle du fluide caloporteur dans les circuits primaires et secondaires. Le sodium est choisi pour ses bonnes propriétés et sa bonne compatibilité avec les matériaux mis en jeu et une grande disponibilité à faible coût. Cependant, il est très réactif avec l'air ou l'eau, ce qui demande un contrôle constant.

Une description plus poussée des réacteurs RNR-Na est disponible dans [41].

Un objectif postérieur aux travaux de cette thèse est d'appliquer la méthodologie développée dans le Chapitre III à des systèmes chimiques plus compliqués et intervenant dans le cadre applicatif des RNR-Na. Un intérêt tout particulier est porté aux systèmes chimiques mettant en jeu les produits de fission (principalement Uranium U et Plutonium Pu) dans leurs interactions

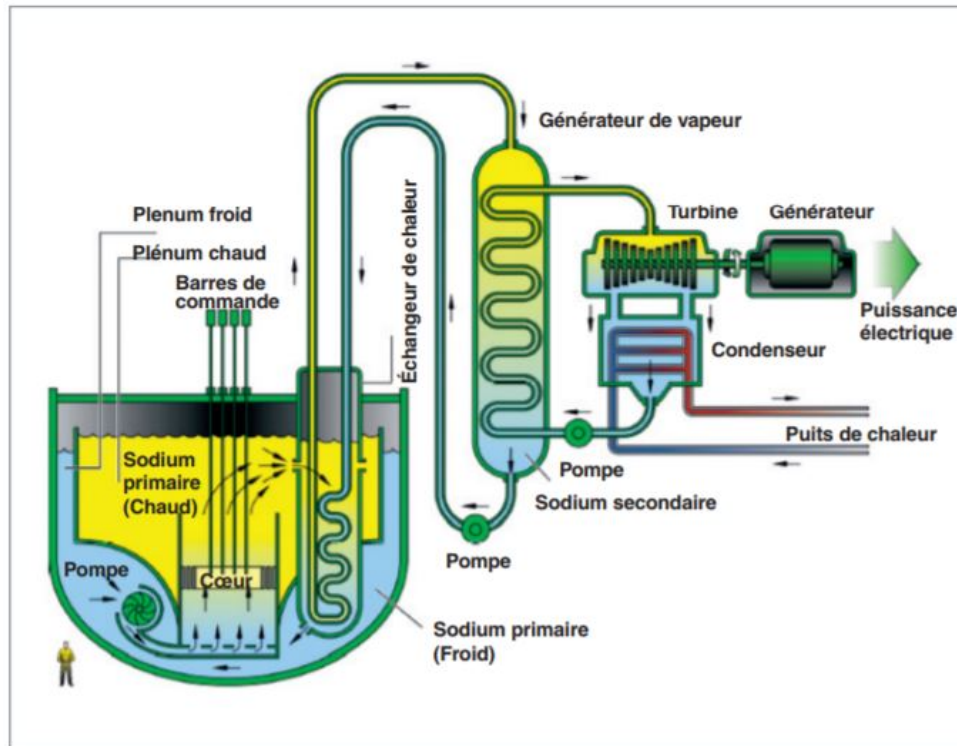


FIGURE I.3 – Schéma du réacteur RNR-Na. Référence figure [41]

avec les matériaux constituant les gaines entourant le combustible ainsi que les métaux de la cuve (principalement Chrome *Cr*, Fer *Fe* et Nickel *Ni*) et enfin le caloporteur (Sodium *Na*).

Dans le cadre des réacteurs de seconde génération, un couplage avec des logiciels de calcul thermodynamique est déjà disponible, voir les travaux de [93]. Le logiciel PROCOR (PROpagation of CORium) est une plateforme de simulation d'accidents graves dont le but est de simuler des phénomènes transitoires intervenant dans le cas de la propagation du corium. Il s'agit d'un code intégral : l'approche se fait à un niveau global. Le modèle physique est donc général et la simulation porte sur l'ensemble de la cuve.

La simulation touche le cadre de la rétention en fond de cuve : à la suite de l'emballement de la réaction en chaîne, la température du cœur augmente conduisant à la fusion de la gaine de combustible. Le corium formé vient alors se déposer en fond de cuve.

Dans le cadre de la simulation accidents grave pour les réacteurs RNR-Na, un logiciel de simulations des phénomènes transitoires intervenant dans le cas de la propagation du corium est en développement, voir les travaux [28] et [5] détaillant les implémentations sur le logiciel SIMMER. On peut également citer les travaux des auteurs de [60] ayant pour but le développement d'indices de sensibilité pour les paramètres de simulation d'un scénario d'accident grave, dans le cas de l'utilisation des réacteurs de type RNR-Na.

Les travaux futurs à cette thèse visent l'intégration de l'analyse d'incertitude pour les fonctions thermodynamiques dans la modélisation des accidents graves des RNR-Na.

Chapitre II

Contexte scientifique

Ce chapitre a pour but de présenter le cadre statistique dans lequel nous abordons les problèmes de thermodynamique étudiés dans ce manuscrit. Ce chapitre est articulé de la manière suivante. Nous proposons, dans un premier temps, une présentation générale de l'apprentissage statistique et de l'analyse des incertitudes. Nous détaillons la modélisation du type systèmes étudiés en thermodynamique chimique dans la Section II. 2. Nous présentons dans la Section II. 3 la méthode Calphad dont l'étude est au cœur de la thèse. Cette méthode nécessite la résolution de plusieurs problèmes. Ces problèmes sont décrits dans la Section II. 4. Enfin un exemple complet est détaillé dans la Section II. 5 à titre d'illustration.

II. 1 L'apprentissage statistique et l'analyse d'incertitude

II. 1.1 Généralités sur l'apprentissage statistique

L'apprentissage statistique joue un rôle important dans la compréhension de nombreux problèmes d'ingénierie. Le but recherché est d'*apprendre* une règle de décision ou un modèle de généralisation à partir d'exemples, ceux-ci constituant une base de d'apprentissage. Considérons une quantité de sortie $Y \in \mathbb{R}$ et un ensemble de grandeurs d'entrée $\mathbf{X} \in \mathbb{R}^d$. Considérons également que nous disposons d'une base d'apprentissage, c'est-à-dire un nombre N de paires entrées/sorties $(\mathbf{X}_i, Y_i)$, $i = 1, \dots, N$. On souhaite être en mesure de prévoir de nouvelles valeurs de la grandeur de sortie Y en fonction de nouvelles valeurs des grandeurs d'entrée $\mathbf{X}$.

Nous décrivons ci-après, de manière plus formelle, les éléments d'un problème de reconstruction d'une fonction reliant la quantité de sortie Y aux quantités d'entrée X . Nous nous plaçons de plus dans le cadre de l'apprentissage supervisé, c'est-à-dire que nous connaissons à la fois les entrées $\mathbf{X}_i$ et les sorties Y_i de la base d'apprentissage. Pour une approche plus complète des problèmes de l'apprentissage statistique, nous renvoyons le lecteur vers les ouvrages de référence suivants [92], [79] et [42].

Plus formellement, on se place dans un espace de probabilité $(\Omega, \mathcal{A}, \mu)$. On note par $\mathbf{X}$ une variable aléatoire à valeur dans $\mathbb{R}^d$ et par Y une variable aléatoire réelle. On suppose l'existence

d'une loi jointe $\mathbb{P}$ (inconnue) pour le couple $(\mathbf{X}, Y)$. On se donne un ensemble d'observations d'apprentissage $D = \{(\mathbf{x}_i, y_i), i = 1, \dots, N\}$, avec $\mathbf{x}_i \in \mathbb{R}^d$ et $y_i \in \mathbb{R}$ pour tout $i = 1, \dots, N$.

On s'intéresse alors au problème de régression suivant. On veut construire une fonction f de $\mathbb{R}^d$ à valeurs dans $\mathbb{R}$ permettant de généraliser la relation entre la sortie Y et les entrées $\mathbf{X}$.

Il peut exister plusieurs fonctions f candidates. Pour choisir entre elles, on utilise une fonction de perte $l(Y, f(\mathbf{X})) \in L^1(\mathbb{P})$. La fonction de perte mesure l'erreur de prédiction. L'erreur de prédiction moyenne est définie de la manière suivante

$$EPM(f) = \mathbb{E}_{\mathbf{X}, Y} [l(Y, f(\mathbf{X}))] = \int_{\mathbb{R}^{d+1}} l(y, f(x)) \mathbb{P}(dx, dy)$$

où l est une fonction de $\mathbb{R}^2$ dans $\mathbb{R}$. L'exemple le plus connu est la fonction de perte quadratique l_2 . La fonction l_2 associe alors à la paire $(y, f(x))$ le carré de la différence, à savoir

$$l_2(y, f(x)) = (y - f(x))^2.$$

Sous réserve que f soit de carré intégrable, l'erreur de prédiction moyenne pour la fonction de perte quadratique est alors le carré de la norme L^2 de l'erreur de prédiction

$$EPM_2(f) = \mathbb{E}_{\mathbf{X}, Y} [l_2(Y, f(\mathbf{X}))] = \int_{\mathbb{R}^{d+1}} (y - f(x))^2 \mathbb{P}(dx, dy).$$

Nous verrons par la suite, notamment dans le Chapitre IV, un plus large choix de fonctions de perte pouvant être considérées pour les problèmes de reconstruction fonctionnelle.

De manière générale, on peut conditionner l'erreur de prédiction moyenne par la loi de $\mathbf{X}$. L'expression de EPM devient

$$EPM(f) = \mathbb{E}_{\mathbf{X}} \left[\mathbb{E}_{Y|\mathbf{X}} [l(Y, f(\mathbf{X})) | \mathbf{X}] \right].$$

On peut alors minimiser EPM ponctuellement en $\mathbf{X}$. Dans le cas de la fonction de perte quadratique, on peut montrer que le meilleur choix de reconstruction pour la fonction de régression f est

$$f(x) = \mathbb{E}_{Y|\mathbf{X}} [Y | \mathbf{X} = x].$$

On considère maintenant le modèle statistique suivant pour la loi jointe de $(\mathbf{X}, Y)$

$$Y = f(\mathbf{X}) + \varepsilon \tag{II.1}$$

où ε modélise l'erreur aléatoire du modèle et f la relation (inconnue) entre $\mathbf{X}$ et Y . Cette erreur est centrée, i.e. $\mathbb{E}[\varepsilon] = 0$, et indépendante de $\mathbf{X}$. Pour ce modèle, le meilleur choix de reconstruction est donnée par

$$f(x) = \mathbb{E}_{Y|\mathbf{X}} [Y | \mathbf{X} = x].$$

Dans la suite manuscrit, on s'intéresse à deux types de reconstruction du modèle f inconnu.

Dans un premier temps, dans le Chapitre III, on regarde une reconstruction paramétrique à l'aide du modèle linéaire. L'hypothèse de reconstruction choisie dans le cas du modèle linéaire est

que l'approximation $\hat{f}$ peut être définie comme une combinaison linéaire d'une base de fonctions bien choisies $(f^i)_{i=1,\dots,p}$ où p est le nombre de fonctions de base. On définit alors

$$\hat{f}(\mathbf{x}) = \sum_{i=1}^p \theta^i f^i(\mathbf{x}) \quad (\text{II.2})$$

où les θ^i sont des poids à déterminer à l'aide de la base d'apprentissage D .

Les choix classiques de fonctions de bases $f^i, i = 1, \dots, p$ sont les polynômes, les séries de Fourier ou les bases d'ondelettes. Il existe un grand nombre de méthodes pour déterminer les valeurs de θ^i . La méthode la plus courante est la minimisation des moindres carrés où on cherche le minimum en $\boldsymbol{\theta} = (\theta^1, \dots, \theta^p)$ de la quantité

$$\sum_{j=1}^N \left(y_j - \sum_{i=1}^p \theta^i f^i(\mathbf{x}_j) \right)^2.$$

Dans le contexte qui nous intéresse au Chapitre III, on s'intéresse à la reconstruction simultanée d'un nombre r de fonctions suivant le modèle suivant

$$f^{(\alpha)}(\mathbf{x}) = \sum_{p=1}^{p_\alpha} \theta_p^{(\alpha)} F_p^{(\alpha)}(\mathbf{x}) + B^{(\alpha)}(\mathbf{x}), \quad \forall \alpha \in [1; r] \quad (\text{II.3})$$

où les fonctions de biais $B^{(\alpha)}$ sont des fonctions connues. L'estimation du vecteur inconnu $\boldsymbol{\theta}$ se fait par minimisation des moindres carrés en considérant des observations modélisées de la manière suivante

$$Z_j^{(\alpha)} = \mathcal{T}_{j,\alpha} \left(f^{(\alpha)}(\mathbf{x}_j) \right) + \varepsilon_j.$$

Dans le Chapitre IV, on se place dans l'espace de probabilité $(U, \mathcal{B}(U), P_U)$ et on s'intéresse à la reconstruction non paramétrique d'une fonction f à valeur dans $\mathbb{R}^r$, définie pour tout $x \in U \subset \mathbb{R}^d$. On suppose que l'intérieur de U est non-vidé et on note $f(x) = (f^1(x), \dots, f^r(x))$. Dans ce cadre, on souhaite reconstruire f sur U telle que f satisfait les N contraintes intégrales suivantes

$$\int_U \sum_{i=1}^r \lambda^i(x) f^i(x) d\Phi_l(x) = z_l \quad 1 \leq l \leq N, \quad (\text{II.4})$$

où les Φ_l sont N mesures finies positives (connues) définies sur $(U, \mathcal{B}(U))$ et les λ^i sont des fonctions de poids continues, connues.

Comme énoncé précédemment, on se propose de résoudre ce problème de reconstruction fonctionnelle dans le cadre de la maximisation de la γ -entropie sous des contraintes linéaires. A savoir, on se donne une fonction convexe γ de $\mathbb{R}^r$ dans $\mathbb{R}_+$. On définit la γ -entropie par

$$I_\gamma(f) = - \int_U \gamma \left(f^1(x), \dots, f^r(x) \right) dP_U(x).$$

On appelle alors le problème de la maximisation de la γ -entropie I_γ sous contraintes linéaires le problème suivant

$$\begin{aligned} & \max I_\gamma(f) \\ & f : \int_U \sum_{i=1}^r \lambda^i(x) f^i(x) d\Phi(x) \in \mathcal{K}. \end{aligned}$$

Nous verrons que la forme très générale de ce problème, nous permettra de considérer un très vaste choix quant à l'information utilisée pour reconstruire la quantité f .

Le choix de modélisation de l'erreur de modèle (II.1) par un bruit additif est une bonne approximation de la réalité pour plusieurs raisons. Pour la plupart des systèmes liés à des paires entrées/sorties $(\mathbf{X}, Y)$, il n'existe pas de relation déterministe $Y = f(\mathbf{X})$. De plus, il existe d'autres variables non mesurées qui contribuent à Y comme par exemple la précision de la méthode d'acquisition des valeurs de Y . Plus généralement on peut classer ces sources d'incertitudes sur la valeur de Y en différentes catégories :

- . l'incertitude survenant de phénomènes aléatoires,
- . l'incertitude survenant suite à un manque de connaissance vis-à-vis du système étudié (aussi appelée incertitude épistémiques),
- . enfin les erreurs de mesure liée à la précision de l'appareil d'acquisition utilisé.

Le reconstruction fonctionnelle décrite précédemment soulève plusieurs questions. Parmi elles, on s'intéresse à la qualité de la prédiction proposée par la fonction reconstruite lorsque les valeurs disponibles pour l'apprentissage de la donnée de sortie Y sont bruitées. On souhaite savoir comment l'incertitude de la base d'apprentissage se reflète dans la prédiction proposée par la reconstruction $\hat{f}$.

II. 1.2 Analyse de l'incertitude

Des grandeurs mal connues ou erronées peuvent parfois intervenir dans les modèles mathématiques. Ces incertitudes peuvent venir de l'existence de phénomènes aléatoires internes au modèle, d'un manque de connaissance ou d'un manque de précision de l'appareil ou du procédé servant à l'acquisition des données.

L'analyse de l'incertitude d'un modèle mathématique a plusieurs intérêts. Elle peut permettre de comprendre l'influence ou l'importance relative des différentes sources d'incertitudes, améliorer la qualité d'un modèle prédictif (pour permettre par exemple une prédiction plus précise), améliorer l'exploitation d'un code ou d'une simulation ou montrer la conformité d'un système avec un critère.

Plus précisément, on considère la sortie Y d'un modèle "boîte noire". On représente ce modèle "boîte noire" par une fonction f dépendant de variables incertaines $\mathbf{X}$ et de paramètres fixes $\mathbf{A}$. La sortie Y est l'image par f des variables d'entrées $\mathbf{X}$. Pour connaître parfaitement f , on pourrait calculer un grand nombre de fois les valeurs de Y associées à un jeu de valeurs de $\mathbf{X}$, ce qui peut s'avérer difficile. En effet, dans certains contextes, l'obtention d'une seule valeur de Y peut se révéler être très coûteuse.

Dans le contexte applicatif qui nous intéresse, les éléments de la base d'apprentissage D sont très souvent le résultat soit d'une expérience coûteuse d'étude des matériaux, soit le résultat d'une simulation numérique d'un modèle compliqué de physique. On veut donc tirer parti au maximum de l'information obtenue à l'aide de chaque point de la base d'apprentissage.

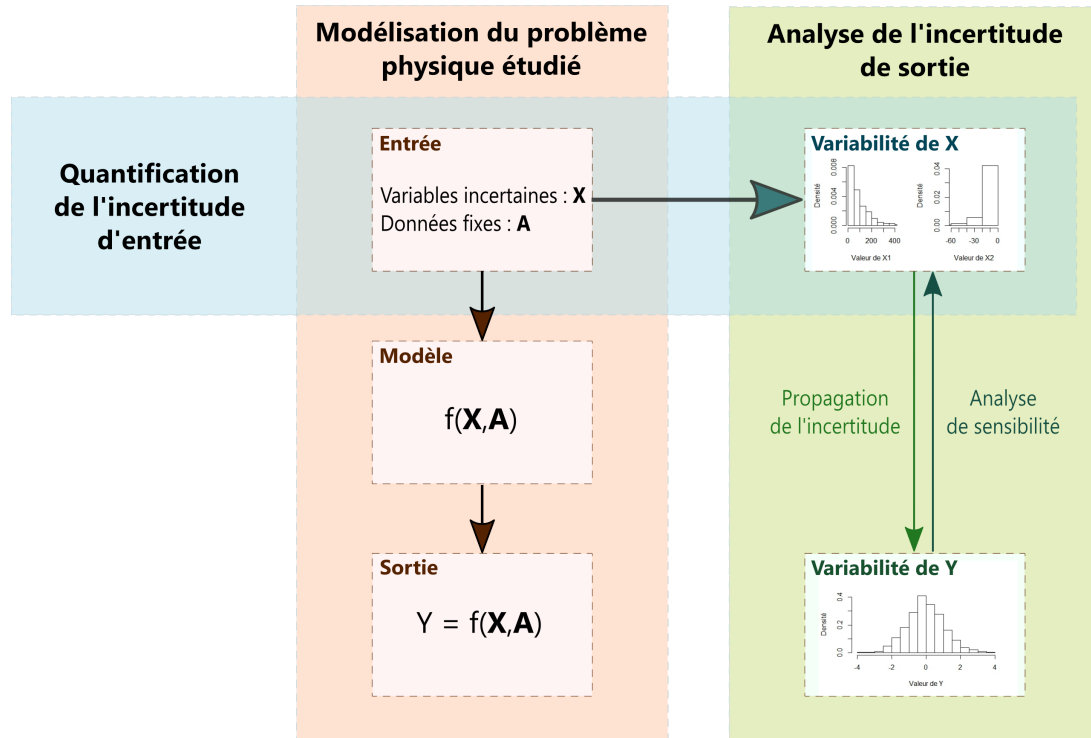


FIGURE II.1 – Modélisation et traitement des incertitudes associés à un problème de physique.

Dans le cadre du traitement de l'incertitude, on peut s'intéresser aux problèmes suivants, schématisés dans la Figure II.1 :

- . **La quantification de l'incertitude d'entrée** s'intéresse à la modélisation de l'incertitude liée aux variables incertaines, utilisées en entrée du code.
- . **La propagation de l'incertitude** permet d'estimer la dispersion de la sortie Y en fonction de la variabilité des grandeurs incertaines X .
- . **L'analyse de sensibilité** permet d'évaluer la sensibilité de la sortie de code en fonction d'une variable ou d'un groupe de variables.

Les problèmes de quantification et de propagation de l'incertitude visent la reconstruction d'une distribution de probabilité. Nous donnons quelques détails et exemples de méthodes sur les différentes étapes du traitement de l'incertitude dans les sections suivantes.

II. 1.2.1 Quantification de l'incertitude d'entrée

Les différentes sources d'incertitudes des grandeurs entrant en jeu dans un problème sont modélisées par des distributions de probabilités. Des informations très variées peuvent être utilisées pour construire ces distributions. Il peut s'agir de données expérimentales, de données numériques haute fidélité ou d'une information plus littérale comme l'avis d'un expert, un standard, une norme de sécurité ou une recommandation. On parlera indistinctement dans cette section de données d'assimilation pour la quantification d'incertitude.

Dans un problème de quantification de l'incertitude, on s'intéresse à la construction d'un modèle probabiliste P reproduisant l'information des données d'assimilation.

Dans le cas de données d'assimilation numériques, on suppose qu'il existe un ensemble empirique de N observations de la grandeur d'entrée, notées $\mathbf{x}_1, \dots, \mathbf{x}_N$. On suppose de plus que ces données sont des observations indépendantes et identiquement distribuées et on note

$$\hat{P}_N = \frac{1}{n} \sum_{i=1}^N \delta_{\mathbf{x}_i}.$$

On peut aborder de deux manières la reconstruction du modèle probabiliste.

Du point-de-vue de la statistique paramétrique, on fait l'hypothèse *a priori* que la distribution appartient à une famille paramétrique de lois $\{P_\vartheta : \vartheta \in \Theta\}$. La détermination de la distribution se résume alors à la détermination du paramètre ϑ de la loi.

Dans le cas du point-de-vue non paramétrique, l'estimation de la distribution ne repose pas sur la détermination de paramètres. Voir par exemple la méthode de reconstruction par noyaux pour l'estimation de la fonction de densité de [94].

Il est nécessaire de supposer que le choix de la famille paramétrique peut être validé. On peut citer par exemple les tests de Kolmogorov-Smirnov ou de Cramer-von Mises. Nous renvoyons le lecteur intéressé vers les ouvrages [26] et [25] pour une formulation plus précise du test d'adéquation d'une loi à des données d'assimilation et les méthodes employées.

La méthode du maximum d'entropie est un autre moyen de déterminer une distribution. Nous reviendrons plus longuement et pour un cadre plus général dans le Chapitre IV sur la méthode du maximum d'entropie. Les différentes méthodes évoquées précédemment supposent néanmoins qu'il existe suffisamment de données pour la détermination de la distribution.

II. 1.2.2 Propagation de l'incertitude

Dans le cadre de la propagation de l'incertitude, on cherche à comprendre la dispersion de la sortie Y en fonction de la variabilité des grandeurs incertaines $\mathbf{X}$. On se place dans le contexte d'un modèle

$$Y = f(\mathbf{X})$$

où $\mathbf{X}$ est une variable aléatoire dans $\mathbb{R}^d$ et Y une variable aléatoire réelle et f est un simulateur déterministe.

L'exemple le plus simple de propagation de l'incertitude est la simulation Monte Carlo par laquelle on cherche la détermination d'une quantité par la simulation aléatoire. On s'intéresse à la quantité I , qui peut être par exemple la valeur moyenne de f ou une valeur seuil $c \in \mathbb{R}$, que l'on définit par

$$I = \int_{\mathbb{R}^d} h(f(\mathbf{x})) \mathbb{P}_{\mathbf{X}}(d\mathbf{x}).$$

Pour les exemples cités précédemment, on peut choisir h comme étant la fonction identité pour le cas de la valeur moyenne ou h comme étant la fonction indicatrice des valeurs plus petites que la valeur c , à savoir $y \rightarrow \mathbf{1}_{y \leq c}$.

On suppose que l'on dispose de n copies de la grandeur d'entrée $\mathbf{X}$. On peut alors définir l'estimateur empirique de la quantité d'intérêt I par

$$\hat{I}_n = \frac{1}{n} \sum_{i=1}^n h(f(\mathbf{X}_i)).$$

Par la loi forte des grands nombres, l'estimateur $\hat{I}_n$ converge de manière presque sûre vers I . On connaît également le comportement asymptotique de l'estimateur par le théorème de la limite centrale, à savoir la quantité définie par

$$\left(\frac{n}{\text{Var}(h(f(\mathbf{X})))} \right)^{\frac{1}{2}} (I - \hat{I}_n)$$

converge en loi vers la loi gaussienne centrée réduite $\mathcal{N}(0, 1)$, permet de déterminer facilement des intervalles de confiance

II. 1.2.3 Analyse de sensibilité

L'analyse de sensibilité a pour objectif premier de déterminer comment la variabilité des grandeurs d'entrée $\mathbf{X}$ contribue à la variabilité de la quantité de sortie Y . Le second objectif de l'analyse de sensibilité est de déterminer, parmi les variables d'entrée, celles qui influent significativement sur la variable de sortie.

Il existe une grande variété de méthodes d'analyse de sensibilité. On distingue les méthodes locales déterministes des méthodes globales qui requiert une modélisation probabiliste satisfaisante des variables d'entrée et de sortie. Nous renvoyons le lecteur intéressé vers les revues suivantes qui proposent un état de l'art détaillé de l'analyse de sensibilité [77], [78], [11].

II. 2 Description des systèmes chimiques

Un système chimique est défini par un mélange de K espèces chimiques. Pour un environnement donné (ou une condition expérimentale donnée), ces espèces forment des micro-structures. Ces micro-structures déterminent les propriétés du matériau qu'elles composent. L'étude des systèmes chimique en thermodynamique est motivée par le besoin de pouvoir prévoir les propriétés et le comportement d'un matériau constitué de nombreuses espèces chimiques en fonction d'une condition expérimentale donnée.

Une *condition expérimentale*, notée $\mathbf{E}$, est un vecteur à composantes positives formé d'une valeur de température T , d'une valeur de pression P et d'une condition de composition chimique $\mathbf{N}$. La composition chimique $\mathbf{N}$ consiste en un vecteur avec I composantes qui représentent les I quantités des éléments chimiques étudiés. On distingue l'espèce chimique de l'élément chimique. L'espèce chimique peut être constituée de plusieurs éléments chimiques (comme par exemple les molécules), d'un seul élément chimique (comme les ions monoatomiques) ou d'aucun élément chimique (comme la lacune qui désigne un vide dans une structure).

L'énergie totale d'un système pour une condition expérimentale $\mathbf{E} \in (\mathbb{R}_+)^{I+2}$ est décrite en utilisant les contributions des r phases pouvant être formées. Une phase est définie par un ensemble d'éléments chimiques avec une composition et des propriétés physiques homogènes. Dans le cas des phases solides, les atomes sont arrangés de manière périodique. Dans le cas des liquides, cet arrangement à longue distance n'existe plus et la phase est désordonnée. Une phase est dite *stable* lorsqu'à sa configuration est associée à une énergie de Gibbs minimale. C'est-à-dire que le système ne dépense pas d'énergie pour maintenir une telle configuration.

Les arrangements des atomes dans l'espace dans les phases solides, ou structures cristallines, sont caractérisées par leurs symétries. Les structures cristallines qui sont équivalentes vis-à-vis de leurs symétries forment des groupes. Ces groupes suivent la classification donnée par les types de groupes d'espace en dimension trois. Nous renvoyons le lecteur intéressé à [7] pour leurs descriptions complètes. La phase solide est décrite avec un motif, c'est-à-dire un volume unitaire constitué d'atomes, qui est translaté à l'infini dans les différentes directions de l'espace. Ce motif translaté définit le réseau cristallin.

Un motif unitaire est composé de sites cristallins. Ces sites sont les localisations des espèces chimiques dans le motif répété. On distingue les différents sites cristallins par leurs caractéristiques : le nombre de plus proches voisins, leurs positions relatives par rapport aux autres sites et leurs dimensions. L'ensemble des sites cristallins d'une phase α ayant les mêmes caractéristiques définit un sous-réseau. Les sites interstitiels du réseau principal sont un exemple de sous-réseau, voir exemple de la Figure II.2. Ces sites sont les espaces vides où les espèces chimiques peuvent s'insérer. Le nombre de sous-réseaux pour une phase α est noté s_α par la suite.

Le lecteur intéressé est renvoyé aux livres [10] et [16] où davantage de détails sont donnés au sujet de la description cristallographique.

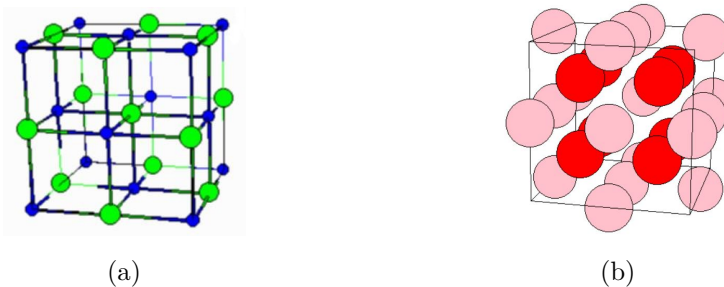


FIGURE II.2 – Deux exemples de sites interstitiels : (a) Site octaédrique : l'ensemble des sphères bleues définit le réseau principal. Entre six sites du réseau principal se trouvent les sites octaédriques, représentés ici par les sphères vertes. (b) Site tétraédrique : les sphères roses définissent le réseau principal. Entre quatre sites du réseau principal se trouvent les sites tétraédriques, représentés ici par les sphères rouges.

Par la suite, on note $G_{\mathbf{E}}(.,.)$ la fonction d'énergie totale du système, définie sur $[0; 1]^r \times [0; 1]^d$ et à valeurs dans $\mathbb{R}$. On note les variables de $G_{\mathbf{E}}$ à l'aide de $\boldsymbol{\lambda}$ et $\boldsymbol{\gamma}$. Le vecteur $\boldsymbol{\lambda} \in [0; 1]^r$ représente les proportions respectives des r différentes phases présentes à l'équilibre. Le vecteur

$\gamma \in [0; 1]^d$ où $d = K \times (\sum_{\alpha=1}^r s_\alpha)$ représente les proportions d'occupation des différents sites cristallins. Chaque composante $\gamma_k^{(\alpha, s_\alpha)}$ donne la proportion, comprise entre 0 et 1, de l'espèce chimique k présente dans la phase α et qui occupe le sous-réseau s_α . La fonction $G_E(., .)$ s'exprime alors comme une somme de r contributions énergétiques $G_E^{(\alpha)}$ pondérée par la quantité de la phase α notée $\lambda^{(\alpha)}$, c'est-à-dire

$$G_E(\boldsymbol{\lambda}, \boldsymbol{\gamma}) = \sum_{\alpha=1}^r \lambda^{(\alpha)} G_E^{(\alpha)}(\boldsymbol{\gamma}^{(\alpha)}), \quad \text{pour } (\boldsymbol{\lambda}, \boldsymbol{\gamma}) \in [0; 1]^r \times [0; 1]^d \quad (\text{II.5})$$

où $G_E^{(\alpha)}$ est la fonction d'énergie de la phase α , définie sur $[0; 1]^d$.

Dans certains cas plus simples, les fonctions d'énergie de Gibbs peuvent être exprimées à l'aide d'une fonction de $\boldsymbol{\gamma}$, à savoir la fraction molaire traditionnellement notée $\boldsymbol{x}$. Les fonctions d'énergie $G_E^{(\alpha)}(\boldsymbol{x})$ sont alors définies pour $\boldsymbol{x} \in [0; 1]^{\tilde{d}}$ où $\tilde{d} = I \times r$ et où chaque composante de $\boldsymbol{x}$ représente la fraction molaire de l'élément $i \in [1; I]$ présent dans la phase α .

La relation entre $\boldsymbol{x}$ et $\boldsymbol{\gamma}$ est définie ci-dessous

$$x_i^{(\alpha)} = \frac{\sum_{s=1}^{s_\alpha} a^{(\alpha, s)} \sum_{k=1}^K b_{i,k} \gamma_k^{(\alpha, s)}}{\sum_{j=1}^I \sum_{s=1}^{s_\alpha} a^{(\alpha, s)} \sum_{k=1}^K b_{j,k} \gamma_k^{(\alpha, s)}}, \quad \alpha = 1, \dots, r \text{ et } i = 1, \dots, I.$$

La quantité $a^{(\alpha, s)}$ représente le nombre total de sites cristallins pour le sous-réseau s de la phase α , c'est-à-dire le nombre d'espèces chimiques qui peut être contenu dans une unité de sous-réseau s . La quantité $b_{i,k}$ représente la correspondance entre l'espèce chimique k et l'élément chimique i , par exemple pour une molécule de dioxygène O_2 , $b_{O, O_2} = 2$. Par la suite, sauf si le contraire est mentionné, la formulation utilisant les occupations des sites cristallins $\boldsymbol{\gamma}$ sera utilisée, cette formulation étant celle permettant la modélisation la plus générale.

Un exemple complet et détaillé de la modélisation d'un système thermodynamique est proposé dans la Section II. 5 de ce chapitre à titre d'illustration. Cet exemple sera repris dans les différentes applications de la thèse.

II. 3 La méthode Calphad

La méthode CALPHAD (pour *CAL*culat*ion of PH*ase *D*iagram) est employée afin de déterminer les fonctions d'énergie libre, également appelée énergie de Gibbs, à partir de données de différentes natures. Nous renvoyons le lecteur intéressé vers les ouvrages [58], [46], [57], [87] sur la méthode Calphad. Les données d'assimilation utilisées par la méthode Calphad peuvent être soit des données de structure, soit des données thermodynamiques, soit des données de diagramme de phase, voir le tableau de la Figure II.3. Les données de structure sont des caractéristiques relatives aux différentes phases. Les données thermodynamiques sont liées aux fonctions d'énergie de Gibbs par des transformations linéaires (par exemple des dérivées ou des dérivées secondes des fonctions de Gibbs).

Natures de données d'assimilation

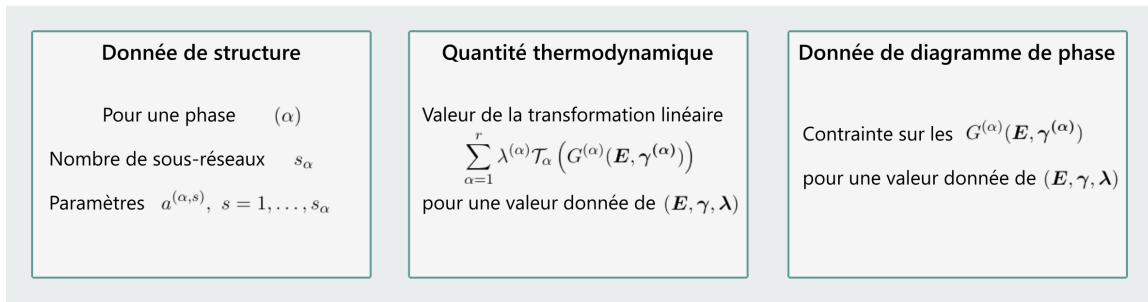


FIGURE II.3 – Les différentes catégories de données d'assimilation utilisées par la méthode Calphad.

Les données de diagramme de phase sont des conditions que doivent satisfaire les fonctions de Gibbs. C'est-à-dire que pour une condition expérimentale E donnée, est associée une contrainte sur les énergies $G^{(\alpha)}$ de plusieurs phases. L'information donnée par diagramme de phase est le résultat d'un problème de minimisation compliqué qui sera appelé par la suite *le problème d'équilibre*. Ce problème sera détaillé dans la Section II. 4.1 et nous donnons plus de détails sur les données de diagramme de phase dans la Section II. 4.3.3.

Plus généralement, la Section II. 4.3 de ce chapitre décrit plus précisément les différents types de données d'assimilation considérées dans les problèmes de calcul thermodynamique.

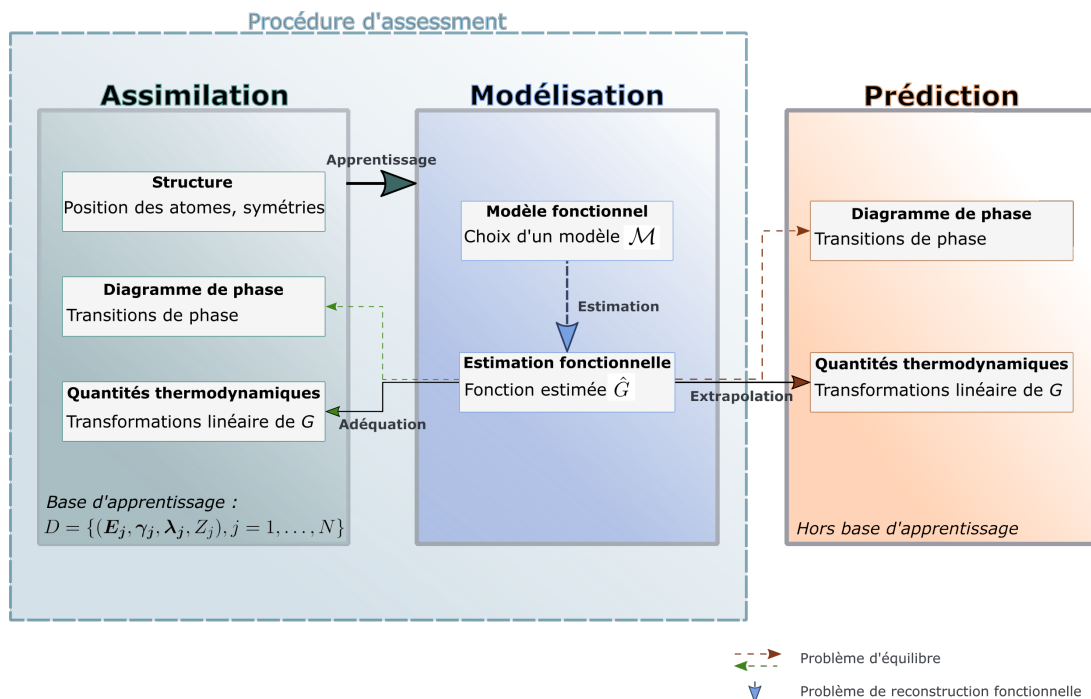


FIGURE II.4 – Le principe de la méthode Calphad et la procédure d'assessment.

L'assessment d'un système chimique à plusieurs éléments est la procédure de détermination des fonctions d'énergie libre de Gibbs des phases de ce système. La méthode de reconstruc-

tion fonctionnelle est basée sur une information partielle. Comme annoncé précédemment, cette information partielle peut être de plusieurs natures.

La méthode Calphad pour la procédure d'*assessment* est découpée en deux étapes qui sont répétées jusqu'à l'obtention d'un résultat satisfaisant pour l'utilisateur.

La première étape est une étape de modélisation. Cette étape vise à la reconstruction d'une ou plusieurs fonctions d'énergie de Gibbs, notée indistinctement G pour l'instant. Au vue des données d'assimilation disponibles, l'utilisateur fait le choix d'un modèle $\mathcal{M}$, c'est-à-dire une famille de fonctions à considérer pour reconstruire la quantité G voulue. On se sert alors des données d'assimilation pour ajuster la quantité G . Formellement on résout un *problème de reconstruction fonctionnelle*. A l'issue de cette étape de modélisation, l'utilisateur dispose d'un estimateur $\hat{G}$. Les modèles considérés dans cette thèse seront explicités dans la Section II. 4.2.

La deuxième étape est l'étape de validation de l'estimateur fonctionnel proposé $\hat{G}$. Muni de l'estimateur $\hat{G}$, l'utilisateur est en mesure de prédire de nouvelles données de diagramme de phase et de nouvelles données thermodynamiques, en dehors de la base d'apprentissage. Le diagramme de phase se détermine à partir des fonctions thermodynamiques en résolvant le *problème d'équilibre*.

Dans le cas général, la procédure d'*assessment* vise à la reconstruction de plusieurs fonctions. Les données d'assimilation peuvent être liées à une ou plusieurs des fonctions à reconstruire. Les utilisateurs de la méthode Calphad définissent alors une hiérarchie dans les fonctions à reconstruire.

II. 4 Les problèmes du calcul thermodynamique

Dans cette section sont décrits les deux problèmes cités précédemment : le *problème d'équilibre* et le *problème de reconstruction fonctionnelle*. Ces deux problèmes sont résolus alternativement lors de l'utilisation de la méthode Calphad pour la procédure d'*assessment*.

II. 4.1 Le problème d'équilibre

La thermodynamique permet la description des systèmes fermés, ce qui signifie qu'il n'y a pas de production ou disparition de matière. De plus les systèmes sont décrits dans leur état d'équilibre, ce qui signifie qu'aucune transformation chimique n'a lieu, qu'il n'y a pas de diffusion ou de perte de chaleur et que l'énergie de Gibbs du système global est la plus basse. Par conséquent le mélange optimal des phases à l'équilibre pour une condition expérimentale $\mathbf{E}$ est obtenu en résolvant un problème de minimisation sous contrainte.

L'utilisateur définit tout d'abord une liste de $r \in \mathbb{N}^*$ phases pouvant être considérées dans le mélange optimal des phases. Le problème de minimisation pour une condition expérimentale $\mathbf{E}$ donnée porte alors sur la détermination d'une paire optimale $(\boldsymbol{\lambda}, \boldsymbol{\gamma}) \in [0; 1]^r \times [0; 1]^d$ où $d = K \times (\sum_{\alpha=1}^r s_{\alpha})$, c'est-à-dire une paire formée par un vecteur de quantités de phase $\boldsymbol{\lambda} \in [0; 1]^r$

et un vecteur d'occupation de sites cristallins $\gamma \in [0; 1]^d$. L'ensemble des contraintes que le mélange optimal (λ^*, γ^*) doit respecter est noté $\mathcal{C}_F$. Cet ensemble est défini par les trois contraintes ci-dessous (a), (b) et (c).

L'équation (a) suivante assure que tous les sites cristallins soient occupés par une espèce chimique seulement

$$\sum_{k=1}^K \gamma_k^{(\alpha,s)} = 1, \quad \forall \alpha \in [1; r], s \in [1; s_\alpha]. \quad (\text{a})$$

Cette modélisation permet également de prendre en compte les défauts si l'ensemble des K espèces chimiques contient la *lacune*.

Les équations (b) et (c) ci-dessous sont des caractéristiques de l'état d'équilibre thermodynamique. La contrainte (b) est la contrainte de conservation de masse

$$\begin{aligned} \sum_{\alpha=1}^r \lambda^{(\alpha)} M_i^{(\alpha)} &= N_i, & \forall i \in [1; I] \\ \text{ou} \quad M_i^{(\alpha)} &= \sum_{s=1}^{s_\alpha} a^{(\alpha,s)} \sum_{k=1}^K b_{i,k} \gamma_k^{(\alpha,s)}, & \forall \alpha \in [1; r], i \in [1; I] \end{aligned} \quad (\text{b})$$

avec N_i la quantité de l'élément i et $M_i^{(\alpha)}$ la masse d'élément i présent en phase α , pour l'équilibre considéré (λ, γ) .

L'équation (c) assure la neutralité électronique du mélange de phase

$$\sum_{s=1}^{s_\alpha} a^{(\alpha,s)} \sum_{k=1}^K \nu_k \gamma_k^{(\alpha,s)} = 0, \quad \forall \alpha \in [1; r], \quad (\text{c})$$

où ν_k est la charge de l'espèce chimique k .

Finalement, le problème de minimisation associé au *problème d'équilibre* s'écrit

$$\min_{(\lambda, \gamma) \in \mathcal{C}_F} G_E(\lambda, \gamma) = \min_{(\lambda, \gamma) \in \mathcal{C}_F} \sum_{\alpha=1}^r \lambda^{(\alpha)} G_E^{(\alpha)}(\gamma^{(\alpha)}), \quad (\text{II.6})$$

où $\mathcal{C}_F$ est défini par

$$\mathcal{C}_F = \left\{ (\lambda, \gamma) \in [0; 1]^r \times [0; 1]^d : \begin{aligned} \sum_{k=1}^K \gamma_k^{(\alpha,s)} &= 1, & \forall \alpha \in [1; r], s \in [1; s_\alpha], \\ \sum_{\alpha=1}^r \lambda^{(\alpha)} M_i^{(\alpha)} &= N_i, & \forall i \in [1; I], \\ \sum_{s=1}^{s_\alpha} a^{(\alpha,s)} \sum_{k=1}^K \nu_k \gamma_k^{(\alpha,s)} &= 0, & \forall \alpha \in [1; r] \end{aligned} \right\}.$$

II. 4.2 Le problème de reconstruction fonctionnelle avec le modèle linéaire

Nous détaillons ici le problème de reconstruction fonctionnelle tel qu'il est classiquement abordé en thermodynamique.

Au cours de la procédure d'*assessment*, l'utilisateur cherche à reconstruire les r fonctions impliquées dans l'équilibre thermodynamique. Dans le cas de la formulation classique en thermodynamique, on modélise à l'aide d'un modèle linéaire les quantités d'intérêts, c'est-à-dire les

fonctions thermodynamiques. Les r fonctions associées aux phases stables intervenant dans le système à l'équilibre sont décrites à l'aide du modèle linéaire avec régresseurs fonctionnels. Un paramètre inconnu multidimensionnel, noté $\boldsymbol{\theta}^{(\alpha)}$ et de taille p_α , est à déterminer pour chaque phase stable. Nous renvoyons le lecteur intéressé vers [58, Chapitre 5] pour une description exhaustive des modèles de fonctions de base utilisés pour la formulation classique en thermodynamique des fonctions d'énergie de Gibbs.

Plus précisément, avec la formulation classique, la fonction d'énergie de Gibbs associée à la phase α s'écrit pour tout $\boldsymbol{\gamma}^{(\alpha)} \in [0; 1]^{d_\alpha}$ où $d_\alpha = K \times s_\alpha$

$$G_{\mathbf{E}}^{(\alpha)}(\boldsymbol{\gamma}^{(\alpha)}) = \sum_{p=1}^{p_\alpha} \theta_p^{(\alpha)} F_{p,\mathbf{E}}^{(\alpha)}(\boldsymbol{\gamma}^{(\alpha)}) + B_{\mathbf{E}}^{(\alpha)}(\boldsymbol{\gamma}^{(\alpha)}) \quad \forall \alpha \in [1; r] \quad (\text{II.7})$$

où $F_{p,\mathbf{E}}^{(\alpha)}$ est une fonction qui appartient à une base bien choisie et $B_{\mathbf{E}}^{(\alpha)}$ est une fonction connue.

On détaille par la suite trois exemples de modèles pour les fonctions d'énergie de Gibbs associées à trois types de phase différents. Sont détaillés pour chaque phase, les fonctions de biais $B_{\mathbf{E}}^{(\alpha)}$ et les régresseurs $F_{p,\mathbf{E}}^{(\alpha)}$.

En premier lieu, est décrit le cas du composé stœchiométrique. Cet exemple constitue le cas le plus simple de modélisation des fonctions de Gibbs. Il s'agit d'une phase solide qui n'existe que pour une composition donnée, c'est-à-dire que le domaine de la fonction d'énergie associée $G_{\mathbf{E}}^{(\alpha)}$ est réduite au singleton $\{\boldsymbol{\gamma}_0^{(\alpha)}\}$, à savoir la composition pour laquelle la phase existe. Du point de vue du cristal, le composé stœchiométrique est caractérisé par le fait que chaque sous-réseau ne peut être occupé que par une seule espèce chimique. La partie connue de la fonction d'énergie $B_{\mathbf{E}}^{(\alpha)}$ est modélisée par

$$B_{\mathbf{E}}^{(\alpha)}(\boldsymbol{\gamma}^{(\alpha)}) = \begin{cases} \sum_{i=1}^I b_i H_{\mathbf{E},i}^o & \text{si } \boldsymbol{\gamma}^{(\alpha)} = \boldsymbol{\gamma}_0^{(\alpha)} \\ \infty & \text{sinon,} \end{cases} \quad (\text{II.8})$$

où les $b_i \in \mathbb{N}$ sont les coefficients stœchiométriques pour le composé solide α considéré et $H_{\mathbf{E},i}^o$ est l'enthalpie des éléments chimiques purs i . L'enthalpie $H_{\mathbf{E},i}^o$ est une fonction de $\mathbf{E}$ connue. Les termes de la partie inconnue dans ce cas sont donnés par

$$G_{\mathbf{E}}^{(\alpha)}(\boldsymbol{\gamma}^{(\alpha)}) - B_{\mathbf{E}}^{(\alpha)}(\boldsymbol{\gamma}^{(\alpha)}) = \theta_0^{(\alpha)} T \log(T) + \theta_1^{(\alpha)} T^{-1} + \sum_{p=2}^{p_{max}} \theta_p^{(\alpha)} T^{p-2}. \quad (\text{II.9})$$

Nous détaillons dans un second cas l'exemple de la phase liquide. Dans le cas d'une phase liquide, le nombre de sous-réseau associé s_α est toujours égal à 1. Les espèces chimiques considérées sont alors les éléments chimiques. La partie connue $B_{\mathbf{E}}^{(\alpha)}$ est définie comme la somme d'un terme dit *de référence* et un terme dit *de configuration*. Elle est donnée par

$$\begin{aligned} B_{\mathbf{E}}^{(\alpha)}(\boldsymbol{\gamma}^{(\alpha)}) &= G_{\mathbf{E}}^{(\alpha,ref)}(\boldsymbol{\gamma}^{(\alpha)}) + G_{\mathbf{E}}^{(\alpha,cf)}(\boldsymbol{\gamma}^{(\alpha)}) \\ &= \sum_{i=1}^I \left(\gamma_i^{(\alpha)} G_{\mathbf{E},i}^o + RT \gamma_i^{(\alpha)} \log(\gamma_i^{(\alpha)}) \right) \end{aligned} \quad (\text{II.10})$$

où R est la constante des gaz parfaits et $G_{E,i}^o$ est l'énergie de Gibbs pour l'élément pur i . La fonction $G_{E,i}^o$ est une fonction connue des conditions expérimentales $\mathbf{E}$. Dans le cas d'un système binaire à deux éléments 1 and 2, la forme de l'énergie d'excès, à savoir le terme inconnu de la fonction d'énergie de Gibbs, est défini selon le modèle dit *modèle binaire d'énergie d'excès de Redlich-Kister*

$$G_{\mathbf{E}}^{(\alpha,ex)}(\gamma^{(\alpha)}) = \gamma_1^{(\alpha)} \gamma_2^{(\alpha)} \left\{ \sum_{p=0}^{p_{max}} \left(\theta_{p,1}^{(\alpha)} + \theta_{p,2}^{(\alpha)} T \right) \left(\gamma_1^{(\alpha)} - \gamma_2^{(\alpha)} \right)^p \right\}. \quad (\text{II.11})$$

Enfin le modèle le plus complet est celui du composé solide ayant un domaine de composition, c'est-à-dire une phase solide pour laquelle le domaine de la fonction d'énergie de Gibbs n'est pas réduit à un singleton. Soit le composé solide $(A, B)_a(C, D)_b$ avec un domaine de composition. La notation $(A, B)_a(C, D)_b$ permet de définir les espèces chimiques occupant le premier (respectivement le second) sous-réseau. Dans l'exemple donné, le premier sous-réseau peut être occupé par l'espèce chimique A ou l'espèce chimique B (et respectivement par l'espèce chimique C ou l'espèce chimique D pour le second sous-réseau). Les indices a et b définissent les coefficients stœchiométriques de chaque sous-réseau, à savoir les proportions de chaque sous-réseau par rapport au total.

La partie connue $B_{\mathbf{E}}^{(\alpha)}$ est définie comme suit

$$\begin{aligned} B_{\mathbf{E}}^{(\alpha)}(\gamma^{(\alpha)}) = & \gamma_A^{(\alpha,1)} \gamma_C^{(\alpha,2)} G_{\mathbf{E},A:C} + \gamma_A^{(\alpha,1)} \gamma_D^{(\alpha,2)} G_{\mathbf{E},A:D} \\ & + \gamma_B^{(\alpha,1)} \gamma_C^{(\alpha,2)} G_{\mathbf{E},B:C} + \gamma_B^{(\alpha,1)} \gamma_D^{(\alpha,2)} G_{\mathbf{E},B:D} \\ & + G_{\mathbf{E}}^{(\alpha,cf)}(\gamma^{(\alpha)}). \end{aligned} \quad (\text{II.12})$$

La fonction $G_{\mathbf{E},A:C}$ est une fonction connue de la condition expérimentale $\mathbf{E}$. Elle correspond à l'énergie qu'un composé solide $A_a C_b$ aurait eue pour une configuration donnée. Ces fonctions sont souvent exprimées à l'aide des fonctions des éléments purs dans leurs phases stables ou métastables. De telles fonctions sont définies dans l'article de Dinsdale [27], utilisé comme référence pour les éléments purs.

Le terme de configuration est donné par

$$\begin{aligned} G_{\mathbf{E}}^{(\alpha,cf)}(\gamma^{(\alpha)}) = RT \left[a \left(\gamma_A^{(\alpha,1)} \log \left(\gamma_A^{(\alpha,1)} \right) + \gamma_B^{(\alpha,1)} \log \left(\gamma_B^{(\alpha,1)} \right) \right) \right. \\ \left. + b \left(\gamma_C^{(\alpha,2)} \log \left(\gamma_C^{(\alpha,2)} \right) + \gamma_D^{(\alpha,2)} \log \left(\gamma_D^{(\alpha,2)} \right) \right) \right]. \end{aligned} \quad (\text{II.13})$$

La fonction $G^{(\alpha)}$ considérée comporte deux parties inconnues que l'on note $G_{\mathbf{E}}^{(\alpha,c)}$ et $G_{\mathbf{E}}^{(\alpha,ex)}$. La première partie de la fonction inconnue est donnée par

$$G_{\mathbf{E}}^{(\alpha,c)}(\gamma^{(\alpha)}) = \theta_0^{(\alpha)} T \log(T) + \theta_1^{(\alpha)} T^{-1} + \sum_{p=2}^{p_{max}} \theta_p^{(\alpha)} T^{p-2}. \quad (\text{II.14})$$

La deuxième partie inconnue $G_E^{(\alpha, ex)}$, comme dans le cas précédent de la phase liquide, est la partie relative à la modélisation des interactions appelé le *terme d'excès*. Celui-ci traduit les interactions entre les différentes espèces chimiques intervenant dans le composé $(A, B)_a(C, D)_b$. Dans le cas considéré ici en exemple, cinq configurations sont possibles. L'ensemble de ces configurations sera noté $\mathcal{C}$. Pour le composé solide $(A, B)_a(C, D)_b$, l'ensemble $\mathcal{C}$ est le suivant

$$\mathcal{C} = \{(A, B : C), (A, B : D), (A : C, D), (B : C, D), (A, B : C, D)\} \quad (\text{II.15})$$

où $(A, B : C)$ signifie la configuration dans laquelle le premier sous-réseau est occupé en partie par l'espèce A et en partie par l'espèce B , le second sous-réseau étant entièrement occupé par l'espèce C . Il en est de même pour les quatre autres configurations de $\mathcal{C}$. Le *terme d'excès* de l'énergie de Gibbs, contient alors autant de termes que d'éléments définis dans l'ensemble $\mathcal{C}$.

Pour chaque configuration $c \in \mathcal{C}$, on définit l'ensemble $\mathcal{K}_c$ des espèces impliquées dans la configuration c . Le *terme d'excès* de la fonction de Gibbs s'exprime alors de la manière suivante

$$G_E^{(\alpha, ex)}(\gamma^{(\alpha)}) = \sum_{c \in \mathcal{C}} \left(\prod_{k \in \mathcal{K}_c} \gamma_k^{(\alpha, s_k)} \right) L_{E,c}(\gamma^{(\alpha)}, \theta_c^{(\alpha)}), \quad (\text{II.16})$$

où la fonction $L_{E,c}$ est une fonction d'interaction, dépendant de la configuration considérée c .

Il existe trois formes de fonction $L_{E,c}$. Dans le cas d'une configuration $(A, B : C)$, la fonction d'interaction est donnée par

$$L_{E,(A,B:C)}(\gamma^{(\alpha)}, \theta_{(A,B:C)}^{(\alpha)}) = \sum_{p=0}^{p_{max}} \left(\theta_{(A,B:C),p,1}^{(\alpha)} + \theta_{(A,B:C),p,2}^{(\alpha)} T \right) \left(\gamma_A^{(\alpha,1)} - \gamma_B^{(\alpha,1)} \right)^p.$$

Dans le cas d'une configuration $(A : C, D)$, elle est donnée par

$$L_{E,(A:C,D)}(\gamma^{(\alpha)}, \theta_{(A:C,D)}^{(\alpha)}) = \sum_{p=0}^{p_{max}} \left(\theta_{(A:C,D),p,1}^{(\alpha)} + \theta_{(A:C,D),p,2}^{(\alpha)} T \right) \left(\gamma_C^{(\alpha,2)} - \gamma_D^{(\alpha,2)} \right)^p.$$

Dans le cas d'une configuration $(A, B : C, D)$, elle est donnée par

$$\begin{aligned} L_{E,(A,B:C,D)}(\gamma^{(\alpha)}, \theta_{(A,B:C,D)}^{(\alpha)}) &= \sum_{p=0}^{p_{max}} \left(\theta_{(A,B:C,D),p,1}^{(\alpha,1)} + \theta_{(A,B:C,D),p,2}^{(\alpha,1)} T \right) \left(\gamma_A^{(\alpha,1)} - \gamma_B^{(\alpha,1)} \right)^p \\ &+ \sum_{p=0}^{p_{max}} \left(\theta_{(A,B:C,D),p,1}^{(\alpha,2)} + \theta_{(A,B:C,D),p,2}^{(\alpha,2)} T \right) \left(\gamma_C^{(\alpha,2)} - \gamma_D^{(\alpha,2)} \right)^p. \end{aligned}$$

II. 4.3 Les données d'assimilation

Pour la reconstruction des fonctions d'énergie de Gibbs, la méthode Calphad se sert d'une information partielle obtenue grâce à différents types de données, à savoir des données de structure, des données thermodynamiques et des données de diagramme de phase. Ces différents types de données seront décrites ici.

Par la suite, un résultat d'expérience sera noté par le vecteur $(\mathbf{E}_j, \gamma_j, \lambda_j, Z_j)$ pour $j = 1, \dots, N$. La quantité Z_j sera utilisée pour donner la valeur d'énergie associée à l'expérience j .

II. 4.3.1 Données de structure

Nous donnons dans cette section quelques détails sur l'obtention des données de structure. Ces données sont très importantes dans la mesure où elles permettent de faire une hypothèse forte sur la nature de la phase présente pour une condition expérimentale donnée.

Nous expliquons brièvement le principe de la diffraction à rayons X (DRX) utilisée pour l'étude des solides, voir le schéma de principe de la Figure II.5. Cette méthode expérimentale permet de valider ou de rejeter l'hypothèse émise sur les structures présentes dans l'échantillon.

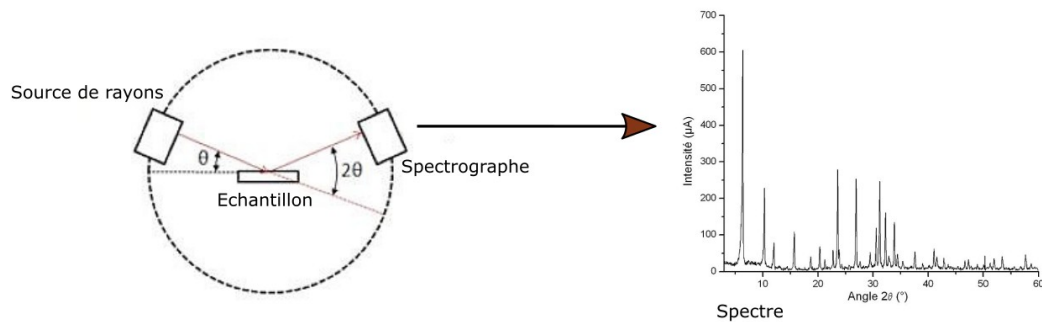


FIGURE II.5 – Principe de la diffraction à rayons X (DRX) pour l'étude des matériaux.

Un échantillon de la matière à analyser est sélectionné et est placé de manière à ce que le faisceau de rayons X touche la surface de l'échantillon avec un angle d'incidence connu θ , c'est-à-dire l'angle avec lequel le rayon frappe l'échantillon. Une partie des rayons est absorbée par l'échantillon, l'autre est réfléchi. Le fait que tous les rayons ne se comportent pas de la même manière vis-à-vis de l'échantillon produit des interférences dont l'intensité est mesurée par un capteur.

L'intensité et la forme des interférences varient en fonction de l'angle d'incidence θ et des arrangements spatiaux des espèces présentes. En faisant varier θ , il est possible de déterminer la structure de la phase présente dans l'échantillon.

L'analyse se déroule de la manière suivante :

- . L'utilisateur définit un ensemble de structures $E_0 = \{\alpha_0, \dots, \alpha_{r_0}\}$ pouvant être celles de la phase de l'échantillon observé. Chaque structure possédant un motif propre, un spectre théorique S_0^α peut être associé à chaque élément de E_0 .
- . L'utilisateur teste l'hypothèse : "la phase de l'échantillon est la phase $\alpha \in E_0$ ".
- . L'hypothèse est acceptée si le spectre théorique S_0^α est suffisamment proche du spectre expérimental obtenu et rejetée sinon.

L'analyse de la diffraction des rayons X permet également la détermination des paramètres des structures cristallographiques (géométrie de la maille, dimensions de la maille, distances entre les centres des atomes, positions et occupations par les différentes espèces chimiques). Nous renvoyons le lecteur intéressé vers [30, Chapitre 4] et plus généralement vers [43] où une description exhaustive des différentes méthodes de spectroscopie est proposée.

II. 4.3.2 Données thermodynamiques

Le terme *données thermodynamiques* concerne les valeurs des grandeurs thermodynamiques associées à l'énergie de Gibbs ou à une transformation linéaire de l'énergie de Gibbs.

Considérons dans un premier temps, un cas de reconstruction plus simple que le cadre classique d'assessment. On considère dans ce cas que pour les N expériences $(\mathbf{E}_j, \boldsymbol{\gamma}_j, \boldsymbol{\lambda}_j, Z_j)$ où $j = 1, \dots, N$, les valeurs Z_j sont les valeurs des mélanges des énergies de Gibbs pour la composition associée. La quantité observée Z_j peut alors être modélisée par

$$Z_j = \sum_{\alpha=1}^r \lambda_j^{(\alpha)} G_{\mathbf{E}_j}^{(\alpha)}(\boldsymbol{\gamma}_j^{(\alpha)}) + \varepsilon_j \quad (\text{II.17})$$

où ε_j représente l'incertitude de la procédure d'acquisition pour la valeur d'énergie Z_j .

Dans le cadre réel de l'assessment, les grandeurs thermodynamiques associées à des dérivées ou des transformations linéaires de l'énergie de Gibbs peuvent être utilisées pour le problème de détermination fonctionnelle. Les exemples de transformations linéaires les plus courants sont donnés ci-dessous

- l'entropie chimique $S_{\mathbf{E}}^{(\alpha)}(\boldsymbol{\gamma}^{(\alpha)}) = -\frac{\partial G_{\mathbf{E}}^{(\alpha)}}{\partial T}(\boldsymbol{\gamma}^{(\alpha)})$
- l'enthalpie $H_{\mathbf{E}}^{(\alpha)}(\boldsymbol{\gamma}^{(\alpha)}) = G_{\mathbf{E}}^{(\alpha)}(\boldsymbol{\gamma}^{(\alpha)}) + T \times S_{\mathbf{E}}^{(\alpha)}(\boldsymbol{\gamma}^{(\alpha)})$
- le potentiel chimique $\mu_{\mathbf{E},i}(\boldsymbol{\gamma}) = \frac{\partial G_{\mathbf{E}}^{(\alpha)}}{\partial N_i}(\boldsymbol{\gamma}^{(\alpha)})$ ¹
- la capacité calorifique $C_{p\mathbf{E}}^{(\alpha)}(\boldsymbol{\gamma}^{(\alpha)}) = -T \frac{\partial^2 G_{\mathbf{E}}^{(\alpha)}}{\partial T^2}(\boldsymbol{\gamma}^{(\alpha)})$.

Dans le cas où ce type de données plus générales sont utilisées pour l'assimilation, on peut noter les N résultats d'expériences $(\mathbf{E}_j, \boldsymbol{\gamma}_j, \boldsymbol{\lambda}_j, U_j)$. La quantité observée U_j est alors modélisée par

$$U_j = \sum_{\alpha=1}^r \lambda_j^{(\alpha)} \mathcal{T}_{j,\alpha} \left(G_{\mathbf{E}_j}^{(\alpha)}(\boldsymbol{\gamma}_j^{(\alpha)}) \right) + \eta_j \quad (\text{II.18})$$

où η_j représente l'incertitude de la procédure d'acquisition pour l'observation U_j et où $\mathcal{T}_{j,\alpha}$ est l'opérateur de transformation linéaire appliqué aux fonctions d'énergie de Gibbs $G_{\mathbf{E}_j}^{(\alpha)}$ correspondant à la donnée observée.

On peut également remarquer que le choix du modèle linéaire de l'approche classique en thermodynamique présente un certain avantage dans la mesure où toutes les données thermodynamiques s'écrivent à l'aide d'un modèle linéaire. En effet, les observations du modèle linéaire pour les fonctions d'énergie s'écrivent alors

$$Z_j = \sum_{\alpha=1}^r \lambda_j^{(\alpha)} \left(\sum_{p=1}^{p_\alpha} \theta_p^{(\alpha)} F_{p,\mathbf{E}_j}^{(\alpha)}(\boldsymbol{\gamma}_j^{(\alpha)}) + B_{\mathbf{E}_j}^{(\alpha)}(\boldsymbol{\gamma}_j^{(\alpha)}) \right) + \varepsilon_j \quad (\text{II.19})$$

1. On note que le potentiel chimique est propre à un élément chimique et ne dépend pas de la phase dans laquelle se trouve l'élément chimique.

et pour les autres quantités thermodynamiques

$$U_j = \sum_{\alpha=1}^r \lambda_j^{(\alpha)} \left(\sum_{p=1}^{p_\alpha} \theta_p^{(\alpha)} \tilde{F}_{p, \mathbf{E}_j}^{(\alpha)}(\gamma_j^{(\alpha)}) + \tilde{B}_{\mathbf{E}_j}^{(\alpha)}(\gamma_j^{(\alpha)}) \right) + \eta_j \quad (\text{II.20})$$

où $\tilde{F}_{p, \mathbf{E}_j}^{(\alpha)}$ (respectivement $\tilde{B}_{\mathbf{E}_j}^{(\alpha)}$) est la fonction $F_{p, \mathbf{E}_j}^{(\alpha)}$ (respectivement $B_{\mathbf{E}_j}^{(\alpha)}$) à laquelle a été appliquée la transformation linéaire $\mathcal{T}_{j, \alpha}$ correspondante. Le modèle reste linéaire pour le paramètre inconnu $\boldsymbol{\theta}$.

Par la suite, sauf mention particulière, on notera toujours les données d'assimilation de type données thermodynamiques par le vecteur $(\mathbf{E}_j, \gamma_j, \lambda_j, Z_j)$.

II. 4.3.3 Données de diagramme de phase

Le diagramme de phase se présente comme une carte des phases stables pour les différentes espèces chimiques du système étudié. L'énergie interne d'un arrangement varie avec les différentes variables d'état, à savoir la composition chimique, la température et la pression. Établir le diagramme de phase d'un système chimique signifie établir une partition du domaine des variables d'états accessibles en un certain nombre de zones, chaque zone pouvant être associée à une ou plusieurs phases.

Une phase stable est une phase avec l'énergie de Gibbs la plus basse. Pour certaines conditions expérimentales, un mélange de plusieurs phases peut avoir une énergie totale plus basse que les phases prises séparément. Les phases intervenant dans le dit-mélange sont alors les phases stables.

Les données de diagramme de phases consistent en des localisations dans la carte des phases stables. Elles consistent

- . soit en des valeurs de conditions expérimentales pour lesquelles apparaissent des changements de phases stables ;
- . soit en une information sur les phases stables présentes à une condition expérimentale donnée.

La Figure II.6 montre un exemple de diagramme de phase. Le diagramme de phase de la Figure II.6 est un diagramme de phase d'un *système binaire* A-B, i.e. à deux éléments. Dans cet exemple, les variables d'états pouvant varier sont la température et la composition relative de la quantité d'espèce par rapport au total des espèces A et B. Le bord gauche (respectivement droit) correspond aux zones où seule l'espèce A (respectivement B) est présente. La pression est fixée.

Afin de déterminer les différentes sous-division du diagramme, il est nécessaire de déterminer l'enveloppe convexe minimale de la liste des r fonctions d'énergie impliquées dans le système étudié. Ce problème de détermination se pose pour chaque valeur de température. Pour un domaine de composition où l'enveloppe convexe est confondue avec une unique fonction d'énergie,

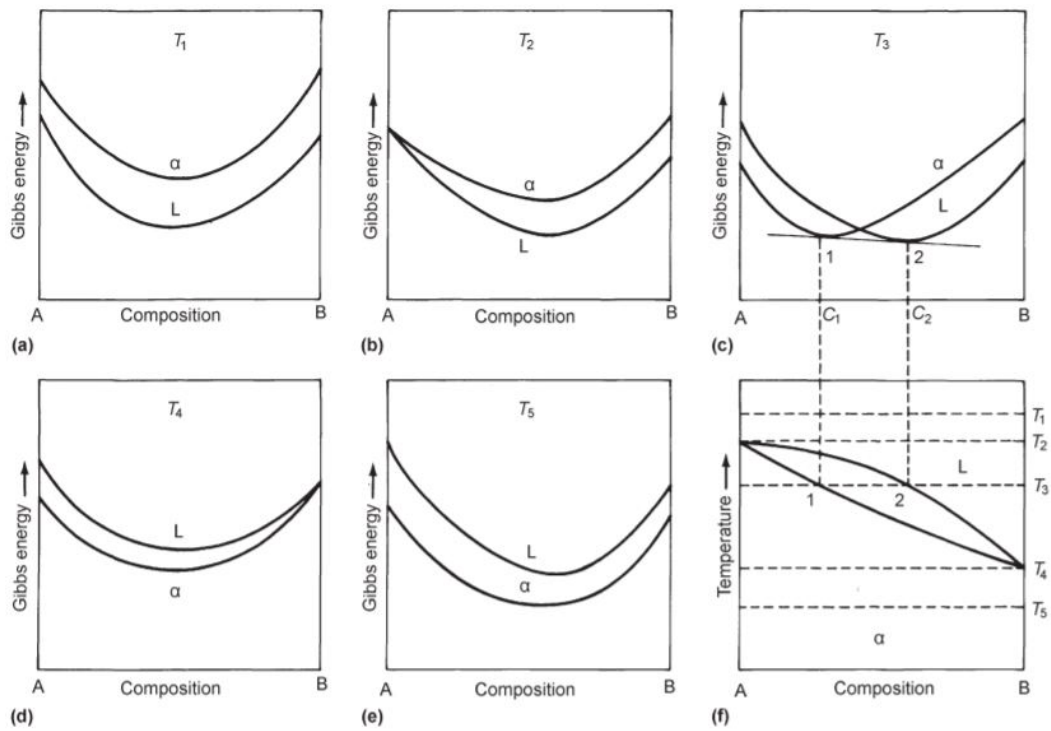


FIGURE II.6 – *Connexions entre le diagramme de phase (f) d'un système A-B à deux phases (la phase α est la phase principale pour les faibles température, la phase L pour les hautes température) et (a)-(e) les énergies propres à chaque phase pour différentes valeurs de température. Les figures (a)-(e) représentent les fonctions d'énergie de Gibbs pour les phases α et L en fonction de la composition relative, pour une température donnée. La composition relative correspond à la quantité de l'espèce chimique B par rapport à la somme des espèces A et B. Pour les températures T_1 et T_2 , la phase L est celle de plus basse énergie tandis que pour les température T_4 et T_5 , c'est la phase α . A la température T_3 , il existe une tangente commune aux deux fonctions d'énergie. Par conséquent, entre les compositions C_1 and C_2 , un mélange des phases α et L est la combinaison la plus stable.*

Référence de la figure [64].

la région correspondante dans le diagramme de phase est labellisée par le nom de la phase associée. Pour un domaine de composition où l'enveloppe convexe est confondue avec une tangente commune à plusieurs fonctions d'énergie, la région correspondante dans le diagramme de phase est labellisée par le nom des phases associées à ces fonctions d'énergie.

II. 5 Modélisation du système Cu-Mg dans la formulation classique de thermodynamique

A titre d'exemple, nous détaillons dans cette section un système complet avec le formalisme classique de la thermodynamique. Nous nous intéressons à l'ensemble des phases formées à partir

de cuivre Cu et de magnésium Mg. Nous donnons la description du système (espèces chimiques existantes, phases stables et méta-stables). Nous rappelons le modèle de fonctions utilisées pour le problème de reconstruction fonctionnelle. Enfin nous listons le type de données utilisées au cours de la procédure d’assessment.

Le système Cu-Mg est un cas d’étude proposé par la SATA². Ses organisateurs ont fourni l’avis d’expert et ont sélectionné les données d’assimilation nécessaires à l’étude du système. Cet exemple sera repris dans le manuscrit.

II. 5.1 Description du système

Dans cette section, nous détaillons de manière brève le système cuivre magnésium (système Cu-Mg) d’un point de vue thermochimique. Nous nous intéressons aux différentes structures cristallographiques que peuvent prendre les alliages formé de cuivre et de magnésium.

Le système étudié comporte 3 espèces chimique ($K = 3$) : l’atome de cuivre Cu, l’atome de magnésium Mg et la lacune notée Va. La lacune est un élément sans volume ni masse utilisé pour la modélisation des défauts. Elle correspond à un site cristallographique vide. Le cuivre est l’élément chimique de numéro atomique 29³ dont la forme native est le cuivre solide dans une structure cubique à faces centrées FCC_A1. La nomenclature des phases correspond à l’abréviation du nom de la structure en anglais (FCC : faces centered cubic) suivie de la notation *Strukturbericht* des structures cristallines. Cette notation est celle établie par le journal *Zeitschrift für Kristallographie – Crystalline Materials*. Le magnésium est l’élément chimique de numéro atomique 12⁴. La structure du magnésium pur est la structure hexagonale compacte HCP_A3. Nous renvoyons le lecteur aux représentations de la Figure II.7.

L’étude se fait à des températures comprises entre 300 et 1400K et à la pression atmosphérique 1 bar. Par la suite, on note l’intervalle de température de l’étude $D_T = [300; 1400]$. Dans cet intervalle de température, il existe cinq phases stables, à savoir une phase liquide et quatre phases solides (notée FCC_A1, CUMG2, LAVES_C15 et HCP_A3). Ces phases peuvent être présentes seules ou dans un mélange de phases stables.

Nous listons ci-dessous les caractéristiques de ces différentes structures :

- . La phase FCC_A1 contient 2 sous-réseau ($s_{FCC_A1} = 2$). Le réseau principal comporte 4 sites cristallographiques ($a^{(FCC_A1,1)} = 4$)⁵ pouvant être occupés par des atomes Cu ou

2. School for Advanced Thermodynamic Assessment

3. C’est à dire qu’il comporte 29 protons.

4. C’est à dire qu’il comporte 12 protons. L’atome de magnésium est donc plus petit et plus léger que l’atome de cuivre décrit précédemment.

5. Le nombre de sites est obtenu en comptant le nombre d’atomes par maille unitaire. Pour cet exemple, la maille unitaire est cubique. Elle comporte un atome à chaque sommet et un atome à chaque centre de face du cube. Le motif de la maille unitaire est translaté à l’infini dans les trois directions de l’espace. L’atome placé sur un sommet est alors partagé par 8 copies de la maille unitaire. Chaque atome ne devant être compté qu’une seule fois dans le dénombrement des atomes de la maille unitaire, on applique un coefficient $\frac{1}{8}$ devant chaque atome se trouvant sur un sommet du cube. De même les atomes se trouvant au centre des faces du cubes sont partagés

Mg représentés par les sphères jaunes Figure II.7 (a). Le deuxième sous-réseau est occupé par 4 sites interstitiels vacants ($a^{(FCC_A1,2)} = 4$).

- La phase HCP_A3 contient 2 sous-réseau ($s_{HCP_A3} = 2$). Le réseau principal comporte 6 sites cristallographiques ($a^{(HCP_A3,1)} = 6$) pouvant être occupés par des atomes Cu ou Mg représenté par les sphères vertes Figure II.7 (b). Le deuxième sous-réseau est occupé par 3 sites interstitiels vacants ($a^{(HCP_A3,2)} = 3$).
- La phase LAVES_C15 contient 2 sous-réseau ($s_{LAVES_C15} = 2$). Le réseau principal comporte 16 sites cristallographiques ($a^{(LAVES_C15,1)} = 16$) pouvant être occupés par des atomes Cu ou Mg. Le deuxième sous-réseau comporte 8 sites ($a^{(LAVES_C15,2)} = 8$) pouvant être occupés par des atomes Cu ou Mg. La configuration idéale est celle représentée Figure II.7 (c). Les atomes Cu remplissent alors le premier sous-réseau et les atomes Mg le second sous-réseau.
- La phase CUMG2 contient 2 sous-réseau ($s_{CUMG2} = 2$). Le réseau principal comporte 16 sites cristallographiques ($a^{(CUMG2,1)} = 16$) occupés par des atomes Cu représenté par les sphères blanches Figure II.7 (d). Le deuxième sous-réseau comporte 32 sites ($a^{(CUMG2,2)} = 32$) occupés par des atomes Mg.

Les caractéristiques de ces phases ont été définies par l'Union Internationale de Cristallographie. Elles sont référencées dans les *Tables Internationales de Cristallographie*, voir [3].

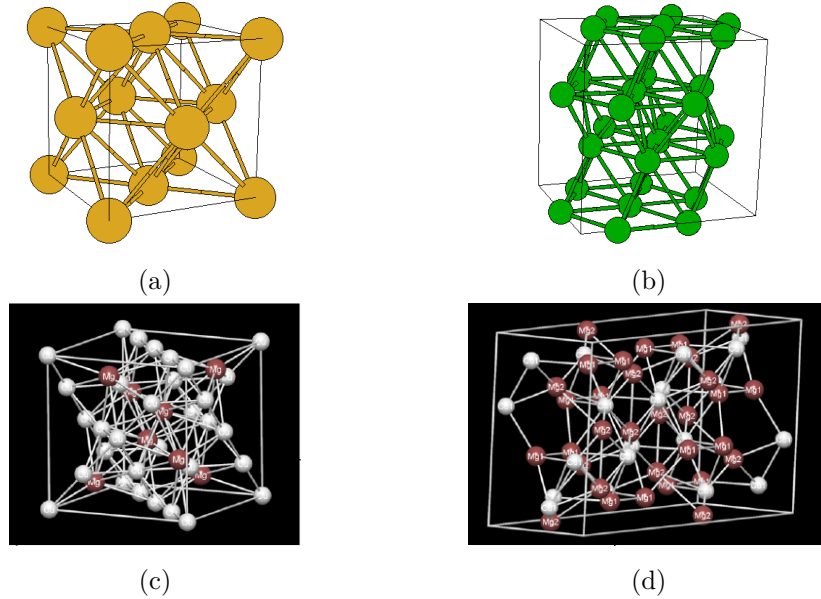


FIGURE II.7 – Les structures solides du système Cu-Mg : (a) FCC_A1, (b) HCP_A3, (c) LAVES_C15 où les atomes de cuivre sont en blanc et les atomes de magnésium en rouge, (d) CUMG2 où les atomes de cuivre sont représentées par des sphères blanches et les atomes de magnésium par des sphères rouges.

par 2 copies de la maille unitaire. On applique donc un coefficient $\frac{1}{2}$. Le dénombrement des atomes de la maille unitaire est donc le suivant : $8 \times \frac{1}{8} + 6 \times \frac{1}{2} = 4$.

II. 5.2 Modélisation classique en thermodynamique

Pour le problème d'assessment associé au système Cu-Mg, on se propose de reconstruire cinq fonctions d'énergie de Gibbs associées aux cinq arrangements de mélange du cuivre et du magnésium. Ces arrangements sont une phase liquide (notée LIQUID) et quatre phases solides notées par FCC_A1, CUMG2, LAVES_C15 et HCP_A3. Ces cinq énergies de Gibbs dans le formalisme usuel de thermodynamique dépendent d'un paramètre θ avec 18 composantes. La description de ces cinq fonctions est donnée ci-après.

La fonction de la phase LIQUID est totalement connue sur les bords du domaine de composition. Elle correspond à la fonction de phase liquide de Cu sur le bord gauche et à la fonction de phase liquide de Mg sur le bord droit. Ces fonctions sont estimées dans la publication de Dinsdale [27] qui sert de référence pour les éléments purs. Le terme de biais (connu) a la forme suivante

$$B_E^{(liq)}(\gamma^{(liq)}) = \gamma_{Cu}^{(liq)} G_{E,Cu}^{liq} + \gamma_{Mg}^{(liq)} G_{E,Mg}^{liq} + RT \left(\gamma_{Cu}^{(liq)} \log(\gamma_{Cu}^{(liq)}) + \gamma_{Mg}^{(liq)} \log(\gamma_{Mg}^{(liq)}) \right), \quad \gamma^{(liq)} \in [0; 1]^2$$

où $G_{E,Cu}^{liq}$ et $G_{E,Mg}^{liq}$ sont les fonctions de Gibbs des éléments purs Cu et Mg dans leur phase liquide, estimées par [27], et R est la constante des gaz parfaits.

Dans le cas d'un système binaire avec deux éléments chimiques, le terme inconnu de la fonction d'énergie, qui correspond au terme dit d'excès, est construit sur le modèle d'excès de Redlich-Kister

$$G_E^{(liq,ex)}(\gamma^{(liq)}) = \gamma_{Cu}^{(liq)} \gamma_{Mg}^{(liq)} \left\{ \left(\theta_{0,1}^{(liq)} + \theta_{0,2}^{(liq)} T \right) + \theta_{1,1}^{(liq)} \left(\gamma_{Cu}^{(liq)} - \gamma_{Mg}^{(liq)} \right) \right\}, \quad \gamma^{(liq)} \in [0; 1]^2. \quad (II.21)$$

La phase FCC_A1 est la configuration de référence des atomes Cu solides. Le terme connu $B_E^{(A1)}$ s'exprime comme la somme du terme dit de référence et d'un terme de configuration. Le terme $B_E^{(A1)}$ est donné par

$$B_E^{(A1)}(\gamma^{(A1)}) = G_E^{(A1,ref)}(\gamma^{(A1)}) + G_E^{(A1,cf)}(\gamma^{(A1)}), \quad \gamma^{(A1)} \in [0; 1]^2$$

avec

$$G_E^{(A1,ref)}(\gamma^{(A1)}) = \gamma_{Cu}^{(A1,1)} G_{E,Cu}^{FCC} + \gamma_{Mg}^{(A1,1)} G_{E,Mg}^{FCC}$$

et

$$G_E^{(A1,cf)}(\gamma^{(A1)}) = RT \left(\gamma_{Cu}^{(A1,1)} \log(\gamma_{Cu}^{(A1,1)}) + \gamma_{Mg}^{(A1,1)} \log(\gamma_{Mg}^{(A1,1)}) \right).$$

Les fonctions $G_{E,Cu}^{FCC}$ et $G_{E,Mg}^{FCC}$ sont des fonctions de Gibbs pour les éléments purs, données par [27]. La partie inconnue de la fonction correspond au terme d'interaction, donné par

$$G_E^{(A1,ex)}(\gamma^{(A1)}) = \gamma_{Cu}^{(A1,1)} \gamma_{Mg}^{(A1,1)} \theta^{(A1)}.$$

La phase HCP_A3 est la configuration de référence des atomes Mg. Le terme connu $B_E^{(A3)}$ est donné par

$$B_E^{(A3)}(\gamma^{(A3)}) = G_E^{(A3,ref)}(\gamma^{(A3)}) + G_E^{(A3,cf)}(\gamma^{(A3)}), \quad \gamma^{(A3)} \in [0; 1]^2$$

avec

$$G_E^{(A3,ref)}(\gamma^{(A3)}) = \gamma_{Cu}^{(A3,1)} G_{E,Cu}^{HCP} + \gamma_{Mg}^{(A3,1)} G_{E,Mg}^{HCP} \quad (II.22)$$

et

$$G_E^{(A3,cf)}(\gamma^{(A3)}) = RT \left(\gamma_{Cu}^{(A3,1)} \log \left(\gamma_{Cu}^{(A3,1)} \right) + \gamma_{Mg}^{(A3,1)} \log \left(\gamma_{Mg}^{(A3,1)} \right) \right). \quad (II.23)$$

Les fonctions $G_{E,Cu}^{HCP}$ et $G_{E,Mg}^{HCP}$ sont des fonctions des éléments purs, estimées par [27]. La partie inconnue de la fonction correspond au terme d'interaction, donné par

$$G_E^{(A3,ex)}(\gamma^{(A3)}) = \gamma_{Cu}^{(A3,1)} \gamma_{Mg}^{(A3,1)} \theta^{(A3)}. \quad (II.24)$$

La phase CUMG2 est un composé stœchiométrique $(Cu)_1(Mg)_2$ existant pour une seule composition, à savoir un tiers de Cu et deux tiers de Mg. Les sites cristallographiques du premier sous-réseau sont alors tous égaux à 1 pour Cu et 0 pour Mg et inversement pour le second sous-réseau. La partie connue $B_E^{(CuMg_2)}$ est nulle et la partie inconnue est donnée par

$$G_{E,Cu:Mg}^{o,CuMg_2} = \begin{cases} \left[\theta_1^{(CuMg_2)} + \theta_2^{(CuMg_2)} T + \theta_3^{(CuMg_2)} T \log(T) \right. \\ \quad \left. + \theta_4^{(CuMg_2)} T^2 + \theta_5^{(CuMg_2)} \frac{1}{T} \right] & \text{si } \gamma_{Cu}^{(CuMg_2,1)} = 1, \gamma_{Mg}^{(CuMg_2,2)} = 1, \\ +\infty & \text{sinon.} \end{cases}$$

La phase LAVES_C15 est un composé solide $(Cu, Mg)_2(Cu, Mg)_1$ disposant d'un domaine de composition. La partie connue de la fonction d'énergie $B_E^{(C15)}$ est donnée par

$$\begin{aligned} B_E^{(C15)}(\gamma^{(C15)}) = & \gamma_{Cu}^{(C15,1)} \gamma_{Cu}^{(C15,2)} G_{E,Cu:Cu}^o \\ & + \gamma_{Mg}^{(C15,1)} \gamma_{Cu}^{(C15,2)} G_{E,Mg:Cu}^o + \gamma_{Mg}^{(C15,1)} \gamma_{Mg}^{(C15,2)} G_{E,Mg:Mg}^o \\ & + RT \left[2 \left(\gamma_{Cu}^{(C15,1)} \log \left(\gamma_{Cu}^{(C15,1)} \right) + \gamma_{Mg}^{(C15,1)} \log \left(\gamma_{Mg}^{(C15,1)} \right) \right) \right. \\ & \left. + \left(\gamma_{Cu}^{(C15,2)} \log \left(\gamma_{Cu}^{(C15,2)} \right) + \gamma_{Mg}^{(C15,2)} \log \left(\gamma_{Mg}^{(C15,2)} \right) \right) \right], \\ & \text{pour } \gamma^{(C15)} \in [0; 1]^4. \end{aligned}$$

Deux termes sont inconnus. Le premier est le terme de référence des composés solides idéaux. Dans ce modèle tous les sites cristallographiques du premier sous-réseau sont occupés par les

atomes de Cu et ceux du deuxième sous-réseau par les atomes de Mg. On a alors un modèle similaire à celui de la phase CUMG2, à savoir

$$\gamma_{Cu}^{(C15,1)} \gamma_{Mg}^{(C15,2)} G_{E,Cu:Mg}^o = \gamma_{Cu}^{(C15,1)} \gamma_{Mg}^{(C15,2)} \left[\theta_{(Cu:Mg),1}^{(C15)} + \theta_{(Cu:Mg),2}^{(C15)} T + \theta_{(Cu:Mg),3}^{(C15)} T \log(T) \right. \\ \left. + \theta_{(Cu:Mg),4}^{(C15)} T^2 + \theta_{(Cu:Mg),5}^{(C15)} \frac{1}{T} + \theta_{(Cu:Mg),6}^{(C15)} T^3 \right].$$

Le deuxième terme est le terme d'excès. On note par $\mathcal{C}_{C15}$ l'ensemble des quatre configurations possible d'atomes dans le cas de la phase LAVES_C15

$$\mathcal{C}_{C15} = \{(Cu, Mg : Cu), (Cu, Mg : Mg), (Cu : Cu, Mg), (Mg : Cu, Mg)\}.$$

Le terme d'excès est la somme des contributions énergétiques de ces quatre configuration :

$$G_E^{(C15,ex)}(\gamma^{(C15)}) = \sum_{c \in \mathcal{C}_{C15}} \left(\prod_{k \in \mathcal{K}_c} \gamma_k^{(C15,s_k)} \right) L_{E,c}(\gamma^{(C15)}, \theta_c^{(C15)}).$$

Dans l'exemple de la première configuration $(Cu, Mg : Cu)$ de l'ensemble $\mathcal{C}_{C15}$, la contribution énergétique est donnée par

$$\gamma_{Cu}^{(C15,1)} \gamma_{Mg}^{(C15,1)} \gamma_{Cu}^{(C15,2)} L_{E,c}(\gamma^{(C15)}, \theta_c^{(C15)}). \quad (\text{II.25})$$

On réduit le nombre de paramètres de modèle en faisant l'hypothèse que les contributions dépendent seulement du sous-réseau considéré. On a alors pour les interactions du premier sous-réseau

$$L_{E,(Cu,Mg:Cu)}(\gamma^{(C15)}, \theta_{(Cu,Mg:.)}) = L_{E,(Cu,Mg:Mg)}(\gamma^{(C15)}, \theta_{(Cu,Mg:.)}) \\ = \theta_{(Cu,Mg:.)}^{(C15)}$$

et pour le second sous-réseau

$$L_{E,(Cu:Cu,Mg)}(\gamma^{(C15)}, \theta_{(.:Cu,Mg)}) = L_{E,(Mg:Cu,Mg)}(\gamma^{(C15)}, \theta_{(.:Cu,Mg)}) \\ = \theta_{(.:Cu,Mg)}^{(C15)}.$$

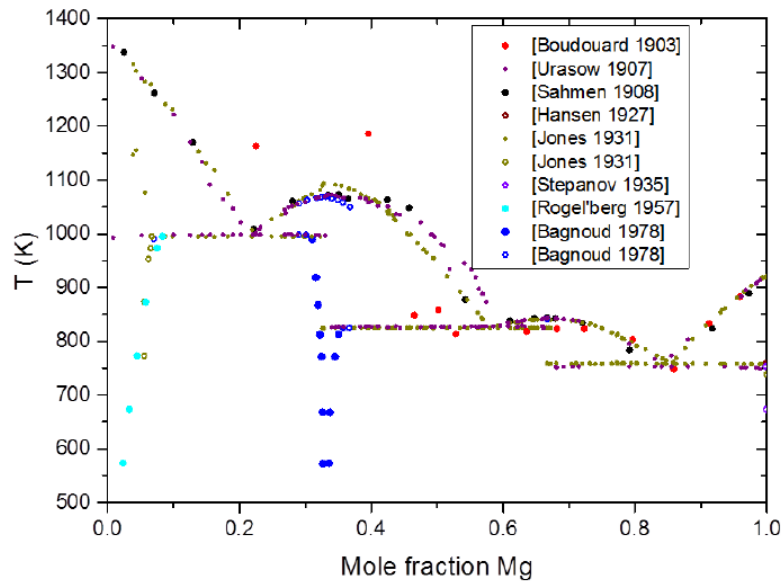
Le terme d'excès s'écrit alors

$$G_E^{(C15,ex)}(\gamma^{(C15)}) = \gamma_{Cu}^{(C15,1)} \gamma_{Mg}^{(C15,1)} \theta_{(Cu,Mg:.)}^{(C15)} + \gamma_{Cu}^{(C15,2)} \gamma_{Mg}^{(C15,2)} \theta_{(.:Cu,Mg)}^{(C15)}.$$

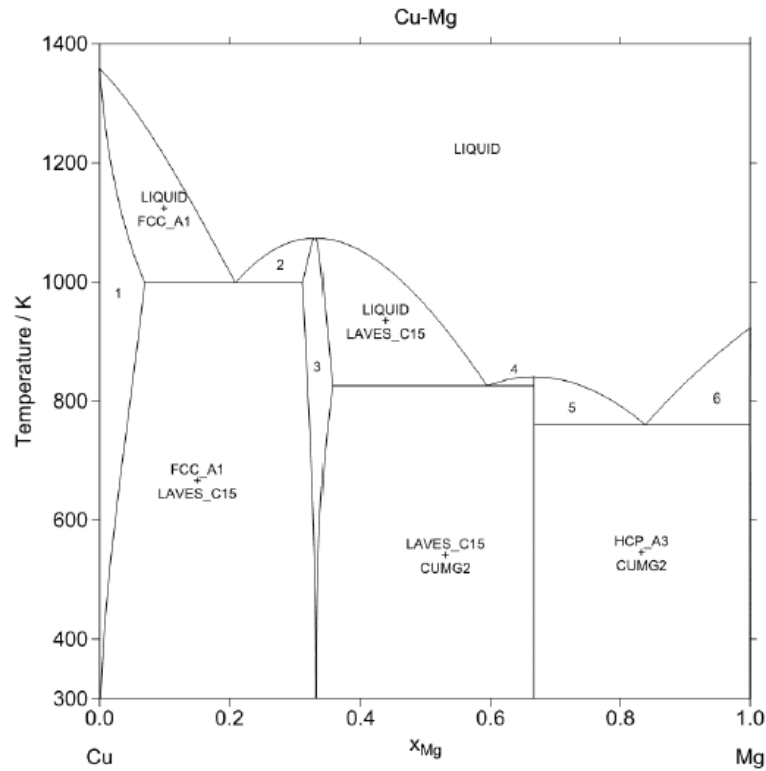
II. 5.3 Données d'assimilation

La sélection des données utilisées pour l'assessment du système Cu-Mg a été faite par les organisateurs de la SATA School. La liste des références est donnée dans le tableau A.1 de l'Annexe A.

La Figure II.8 (a) place les données de diagramme de phase utilisées pour l'assessment dans un graphe température/composition. La Figure II.8 (b) montre le diagramme de phase idéal du système Cu-Mg.



(a)



(b)

FIGURE II.8 – (a) : Les données de diagramme de phase pour le système binaire Cu-Mg. L'abscisse donne la composition relative en Mg ($x_{Mg} = 0$ signifie que seul l'élément Cu est présent et inversement) et l'ordonnée donne la température du changement de phase.

(b) : Le diagramme de phase idéal du système Cu-Mg. L'abscisse donne la composition relative en Mg et l'ordonnée donne la température.

Labels manquants du diagramme : 1. FCC_A1. 2. LIQUID + LAVES_C15. 3. LAVES_C15. 4. LIQUID + CUMG2. 5. LIQUID + CUMG2. 6. LIQUID + HCP_A3.

Chapitre III

The linear model in the Bayesian framework

III. 1 General results with the linear model

We recall first in this section some general results with the linear model. We refer the interested reader to the books of Azais and Bardet [4], Saporta [80] and Hastie, Tibshirani and Friedman [42] for more details.

We consider the quantities $Z_j, j = 1, \dots, N$ to be N realisations of the following linear model

$$Z_j = F(\mathbf{X}_j)\boldsymbol{\theta} + \varepsilon_j$$

where ε_j are independent centred gaussian variable $\varepsilon_j \sim \mathcal{N}(0, \sigma_\varepsilon^2)$ with variance $\sigma_\varepsilon^2 > 0$. Random variables ε_j can represent measurement uncertainty or computation code uncertainty. The column vector $\boldsymbol{\theta}$ of size p is the unknown parameter to be estimated. The line vector $F(\mathbf{X}_j) = (f_1(\mathbf{X}_j) \ \cdots \ f_p(\mathbf{X}_j))$ consists in the regressor basis evaluated at design point $\mathbf{X}_j$. Design point $\mathbf{X}_j$ is a multidimensional quantity with d components describing the experimental conditions.

If quantity Z_j is a vector of dimension n_Z , ε_j is a vector of the same dimension than the observed quantity Z_j and is characterised by a covariance matrix $n_Z \times n_Z$. $F(\mathbf{X}_j)$ corresponds to $n_Z \times p$ matrix. In the following, the case of $n_Z = 1$ will be considered to explain the approach.

We use the following notations

- the N observations

$$\mathbf{Z} = (Z_1, \dots, Z_N)^T$$

- the design matrix with N lines and d columns with d the number of features for the experimental conditions

$$\mathbf{X} = \begin{pmatrix} \mathbf{X}_1 \\ \vdots \\ \mathbf{X}_N \end{pmatrix} = \begin{pmatrix} X_{1,1} & \cdots & X_{1,d} \\ \vdots & & \vdots \\ X_{N,1} & \cdots & X_{N,d} \end{pmatrix}$$

- the matrix of the regressor basis evaluated at the points of the design matrix with N lines and p columns, we suppose $\mathbf{F}$ of full rank,

$$\mathbf{F} = \begin{pmatrix} f_1(\mathbf{X}_1) & \cdots & f_p(\mathbf{X}_1) \\ \vdots & & \vdots \\ f_1(\mathbf{X}_N) & \cdots & f_p(\mathbf{X}_N) \end{pmatrix}.$$

One needs to estimate the p components of the unknown parameter $\boldsymbol{\theta}$ in the light of the vector of observations $\mathbf{Z}$. Additionally one may also provides an estimation of noise variance σ_ε^2 when it is unknown.

The frequentist approach consists in defining a likelihood function which relies on the values of the observations $\mathbf{Z}$ and the value of parameters $\boldsymbol{\theta}$ and σ_ε^2 . Because of the assumption of a gaussian noise ε_j , the likelihood function for an observation Z_j is given by

$$L(Z_j | \boldsymbol{\theta}, \sigma_\varepsilon^2) = \frac{1}{\sqrt{2\pi\sigma_\varepsilon^2}} \exp\left(-\frac{1}{2\sigma_\varepsilon^2} (Z_j - F(X_j)\boldsymbol{\theta})^2\right).$$

As observations are supposed to be independent measurements, the likelihood function for the whole vector of observations $\mathbf{Z}$ is given by

$$\begin{aligned} L(\mathbf{Z} | \boldsymbol{\theta}, \sigma_\varepsilon^2) &= \prod_{j=1}^N L(Z_j | \boldsymbol{\theta}, \sigma_\varepsilon^2) \\ &= \left(\frac{1}{2\pi\sigma_\varepsilon^2}\right)^{\frac{N}{2}} \exp\left(-\frac{1}{2\sigma_\varepsilon^2} \sum_{j=1}^N (Z_j - F(X_j)\boldsymbol{\theta})^2\right). \end{aligned} \quad (\text{III.1})$$

The estimation of unknown quantities $\boldsymbol{\theta}$ and σ_ε^2 consists in choosing the values $\hat{\boldsymbol{\theta}}$ and $\hat{\sigma}_\varepsilon^2$ that maximises the likelihood function (III.1). This leads to

$$\begin{aligned} \hat{\boldsymbol{\theta}} &= (\mathbf{F}^T \mathbf{F})^{-1} \mathbf{F}^T \mathbf{Z} \\ \hat{\sigma}_\varepsilon^2 &= \frac{1}{N} \|\mathbf{Z} - \mathbf{F} \hat{\boldsymbol{\theta}}\|^2. \end{aligned}$$

Confidence in the estimator can be evaluated by computing the confidence ellipsoid of one estimator $\hat{\boldsymbol{\theta}}$ for a level α . One need to solve the quadratic problem of finding all vectors $\mathbf{c} \in \mathbb{R}^p$ such that

$$\frac{1}{\hat{\sigma}_\varepsilon^2} (\mathbf{c} - \hat{\boldsymbol{\theta}})^T (\mathbf{F}^T \mathbf{F})^{-1} (\mathbf{c} - \hat{\boldsymbol{\theta}}) \leq pq \quad (\text{III.2})$$

with q the $1 - \alpha$ quantile of the Fisher's law $\mathcal{F}(p, N - p)$ with degrees of freedom $(p, N - p)$, see [33]. The subset of all vectors $\mathbf{c}$ that are solutions of (III.2) defines the confidence ellipsoid.

In the case of an heteroscedastic noise, that is the linear model

$$Z_j = F(\mathbf{X}_j)\boldsymbol{\theta} + \varepsilon_j$$

where ε_j are independent centred gaussian variable $\varepsilon_j \sim \mathcal{N}(0, \sigma_j^2)$ with variance $\sigma_j^2 > 0$, the likelihood function for the whole vector of observations $\mathbf{Z}$ is given by

$$\begin{aligned} L(\mathbf{Z} \mid \boldsymbol{\theta}, \sigma_1^2, \dots, \sigma_N^2) &= \prod_{j=1}^N L(Z_j \mid \boldsymbol{\theta}, \sigma_j^2) \\ &= \frac{1}{(2\pi)^{\frac{N}{2}} |\boldsymbol{\Sigma}|^{\frac{1}{2}}} \exp \left(-\frac{1}{2} (\mathbf{Z} - \mathbf{F}\boldsymbol{\theta})^T \boldsymbol{\Sigma}^{-1} (\mathbf{Z} - \mathbf{F}\boldsymbol{\theta}) \right) \end{aligned} \quad (\text{III.3})$$

where $\boldsymbol{\Sigma}$ is the diagonal matrix of experiments uncertainties σ_j^2

$$\boldsymbol{\Sigma} = \begin{pmatrix} \sigma_1^2 & 0 & \cdots & 0 \\ 0 & \sigma_2^2 & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ 0 & \cdots & 0 & \sigma_N^2 \end{pmatrix}.$$

In that case, the estimation of unknown quantity $\boldsymbol{\theta}$ consists in choosing the values $\hat{\boldsymbol{\theta}}$ that maximises the likelihood function (III.3). This leads to

$$\hat{\boldsymbol{\theta}} = (\mathbf{F}^T \boldsymbol{\Sigma}^{-1} \mathbf{F})^{-1} \mathbf{F}^T \boldsymbol{\Sigma}^{-1} \mathbf{Z}.$$

III. 2 The Bayesian framework

III. 2.1 General definitions and properties in the Bayesian framework

In this Section, we briefly set the Bayesian approach in a very general framework. We refer the interested reader to [72] for a comprehensive approach of the Bayesian framework. We describe how to get the *posterior* distribution of some random variable $U \in \mathbb{R}$, that is the probability distribution for such random variable considering an other random variable $V \in \mathbb{R}$ to be known. We recall the basic definitions and properties of the Bayesian approach. We give the classical algorithms of the Bayesian estimation and Bayesian sampling.

Let us now consider that there exists a *prior* knowledge on the random variable U , that is there exists an expert advice about the values that U can take. The Bayesian approach allows to integrate such anterior knowledge on U , that is the so-called *prior* distribution of such random variable we will further denote by π . *Prior* knowledge can either come from previous studies or expert advice.

We first recall the Bayes' theorem on which is based the Bayesian approach.

Theorem 1. *Consider a probability space $\Omega = (A, \mathcal{P}(A), \mathbb{P})$ and two random variables U and V defined from Ω to $\mathbb{R}$. Bayes' theorem applied to the events $E_1 = \{U = u\}$ and $E_2 = \{V = v\}$ gives the conditional probability of event E_1 knowing that event E_2 occurs. Provides $\mathbb{P}(V = v) \neq 0$, this conditional probability is equal to*

$$\mathbb{P}(U = u \mid V = v) = \frac{\mathbb{P}(V = v \mid U = u) \mathbb{P}(U = u)}{\mathbb{P}(V = v)}.$$

We consider now a parametric model where the quantity of interest Z relies on a multidimensional parameter $\boldsymbol{\theta}$. As in the *frequentist* approach, the *Bayesian* approach features a likelihood function which is denoted by $L(Z | \boldsymbol{\theta})$. In addition, is required a *prior* distribution which will be denoted by $\pi(\boldsymbol{\theta})$. Based on Bayes' theorem, one can compute the *posterior* distribution of $\{\boldsymbol{\theta} | Z\}$, that is

$$\mathbb{P}(\boldsymbol{\theta} = \boldsymbol{\theta}_0 | Z = z_0) = \frac{L(Z = z_0 | \boldsymbol{\theta} = \boldsymbol{\theta}_0)\pi(\boldsymbol{\theta} = \boldsymbol{\theta}_0)}{m(z_0)}$$

where $m(z_0)$ is the marginal law of the observation z_0 . Denoting Θ the set of possible values of parameter $\boldsymbol{\theta}$, the marginal $m(z_0)$ is defined by

$$m(z_0) = \int_{\Theta} L(Z = z_0 | \boldsymbol{\theta} = \boldsymbol{\theta}_0)\pi(\boldsymbol{\theta} = \boldsymbol{\theta}_0)d\boldsymbol{\theta}_0.$$

Once having at hand the *posterior* distribution, one can also be interested in the prediction of an unknown quantity $\psi(\boldsymbol{\theta})$ based on the vector of N observations $\mathbf{Z} = (Z_1, \dots, Z_N)^T$. This corresponds to the *posterior* prediction of $\psi(\boldsymbol{\theta})$ knowing $\mathbf{Z} = \mathbf{z}$ which is given by

$$p(\psi(\boldsymbol{\theta}) = y | \mathbf{Z} = \mathbf{z}) = \int_{\Theta} L(\psi(\boldsymbol{\theta}) = y | \boldsymbol{\theta} = \boldsymbol{\theta}_0)p(\boldsymbol{\theta} = \boldsymbol{\theta}_0 | \mathbf{Z} = \mathbf{z})d\boldsymbol{\theta}_0.$$

We will only consider linear application of $\boldsymbol{\theta}$ as function $\psi(\boldsymbol{\theta})$ in the sequel.

As a parallel of the confidence interval in the *frequentist* approach, one can compute the credible set of the parameter $\boldsymbol{\theta}$. A credible set of level $100(1 - \alpha)\%$ of parameter $\boldsymbol{\theta}$ is a subset $\mathcal{S}$ of Θ such that

$$\begin{aligned} \mathbb{P}(\boldsymbol{\theta} \in \mathcal{S} | Z = z_0) &= \int_{\mathcal{S}} \mathbb{P}(\boldsymbol{\theta} = \boldsymbol{\theta}_0 | Z = z_0)d\boldsymbol{\theta}_0 \\ &\geq 1 - \alpha. \end{aligned}$$

There exists several possibilities to compute such credible sets. One possibility is the high *posterior* density rule will be followed. Such rule gives a set with the most likely values of parameter $\boldsymbol{\theta}$. One wants to find a constant $k(\alpha)$ which depends on the selected credible level α . The credible set $\mathcal{S}$ is then defined by

$$\mathcal{S} = \left\{ \boldsymbol{\theta}_0 \in \Theta : \mathbb{P}(\boldsymbol{\theta} = \boldsymbol{\theta}_0 | Z = z_0) \geq k(\alpha) \right\}.$$

That is that the value $k(\alpha)$ is the constant value such that $\mathbb{P}(\boldsymbol{\theta} \in \mathcal{S} | Z = z_0) \geq 1 - \alpha$.

III. 2.2 Conjugate prior distribution for the Gaussian linear model

In this Section, we explain the specificities of a bayesian model using conjugate *prior* distributions as it is proposed in our work. This approach make a specific use of the regression model used for the Gibbs energy function.

We consider the following linear model

$$\mathbf{Z} = F(\mathbf{X})\boldsymbol{\theta} + \boldsymbol{\varepsilon}$$

where $\mathbf{Z}$ is a vector of N observations, $F(\mathbf{X})$ is a basis of p functions located at the points of the design matrix $\mathbf{X}$ and $\boldsymbol{\theta}$ is an unknown p dimensional parameter. $\mathbf{X}$ is a $N \times d$ matrix for which each line sums up the experimental conditions for observation Z_i . In the frame of the Calphad-based modelling, the vector $\mathbf{Z}$ represents thermodynamic quantities and the $N \times p$ matrix $F(\mathbf{X})$ is a basis of p functions applied at the N experimental conditions. The vector $\boldsymbol{\varepsilon}$ models the experiment uncertainty and is a centred Gaussian with covariance matrix $\boldsymbol{\Sigma}$.

In our framework, we propose to use a Gaussian likelihood function with a Gaussian *prior* distribution. The likelihood model for the observation vector $\mathbf{Z}$ is given by

$$L(\mathbf{Z} = \mathbf{z} \mid \boldsymbol{\theta} = \boldsymbol{\theta}_1) = \frac{1}{(2\pi)^{N/2} |\boldsymbol{\Sigma}|^{1/2}} \exp \left(-\frac{1}{2} (\mathbf{z} - F(\mathbf{X})\boldsymbol{\theta}_1)^T \boldsymbol{\Sigma}^{-1} (\mathbf{z} - F(\mathbf{X})\boldsymbol{\theta}_1) \right) \quad (\text{III.4})$$

where $|\boldsymbol{\Sigma}|$ is the determinant of the covariance matrix $\boldsymbol{\Sigma}$.

We choose a Gaussian *prior* distribution for the parameter $\boldsymbol{\theta}_0$ centred in the value $\boldsymbol{\theta}_0$ and with a *prior* covariance matrix $\boldsymbol{\Sigma}_0$. The density function for *prior* distribution is given by

$$\pi(\boldsymbol{\theta} = \boldsymbol{\theta}_1) = \frac{1}{(2\pi)^{p/2} |\boldsymbol{\Sigma}_0|^{1/2}} \exp \left(-\frac{1}{2} (\boldsymbol{\theta}_1 - \boldsymbol{\theta}_0)^T \boldsymbol{\Sigma}_0^{-1} (\boldsymbol{\theta}_1 - \boldsymbol{\theta}_0) \right).$$

In that case, the *posterior* distribution for parameter $\boldsymbol{\theta}$ is a Gaussian distribution, see the following result.

Proposition 2. *Consider a multivariate Gaussian prior distribution $\mathcal{N}_p(\boldsymbol{\theta}_0, \boldsymbol{\Sigma}_0)$ for the multi-dimensional parameter in the linear model $\mathbf{Z} = F(\mathbf{X})\boldsymbol{\theta} + \boldsymbol{\varepsilon}$. The posterior distribution, for the likelihood function (III.4), is then the Gaussian distribution $\mathcal{N}_p(\boldsymbol{\theta}_P, \boldsymbol{\Sigma}_P)$ with mean vector*

$$\boldsymbol{\theta}_P = \boldsymbol{\Sigma}_P \left(F(\mathbf{X})^T \boldsymbol{\Sigma}^{-1} \mathbf{z} + \boldsymbol{\Sigma}_0^{-1} \boldsymbol{\theta}_0 \right) \quad (\text{III.5})$$

and covariance matrix

$$\boldsymbol{\Sigma}_P = \left(F(\mathbf{X})^T \boldsymbol{\Sigma}^{-1} F(\mathbf{X}) + \boldsymbol{\Sigma}_0^{-1} \right)^{-1}. \quad (\text{III.6})$$

Proof. Using Bayes' theorem, we have

$$\begin{aligned}
 & \mathbb{P}(\boldsymbol{\theta} = \boldsymbol{\theta}_1 \mid \mathbf{Z} = \mathbf{z}) \\
 & \propto L(\mathbf{Z} = \mathbf{z} \mid \boldsymbol{\theta} = \boldsymbol{\theta}_1) \pi(\boldsymbol{\theta} = \boldsymbol{\theta}_1) \\
 & \propto \exp\left(-\frac{1}{2}(\mathbf{z} - F(\mathbf{X})\boldsymbol{\theta}_1)^T \boldsymbol{\Sigma}^{-1}(\mathbf{z} - F(\mathbf{X})\boldsymbol{\theta}_1)\right) \exp\left(-\frac{1}{2}(\boldsymbol{\theta}_1 - \boldsymbol{\theta}_0)^T \boldsymbol{\Sigma}_0^{-1}(\boldsymbol{\theta}_1 - \boldsymbol{\theta}_0)\right) \\
 & \propto \exp\left(-\frac{1}{2}\left\{ \boldsymbol{\theta}_1^T F(\mathbf{X})^T \boldsymbol{\Sigma}^{-1} F(\mathbf{X}) \boldsymbol{\theta}_1 + \boldsymbol{\theta}_1^T \boldsymbol{\Sigma}_0^{-1} \boldsymbol{\theta}_1 \right. \right. \\
 & \quad \left. \left. - \boldsymbol{\theta}_1^T \left(F(\mathbf{X})^T \boldsymbol{\Sigma}^{-1} \mathbf{z} + \boldsymbol{\Sigma}_0^{-1} \boldsymbol{\theta}_0\right) - \left(F(\mathbf{X})^T \boldsymbol{\Sigma}^{-1} \mathbf{z} + \boldsymbol{\Sigma}_0^{-1} \boldsymbol{\theta}_0\right)^T \boldsymbol{\theta}_1 \right\}\right) \\
 & \propto \exp\left(-\frac{1}{2}\left\{ \boldsymbol{\theta}_1^T \left(F(\mathbf{X})^T \boldsymbol{\Sigma}^{-1} F(\mathbf{X}) + \boldsymbol{\Sigma}_0^{-1}\right) \boldsymbol{\theta}_1 \right. \right. \\
 & \quad \left. \left. - \boldsymbol{\theta}_1^T \left(F(\mathbf{X})^T \boldsymbol{\Sigma}^{-1} \mathbf{z} + \boldsymbol{\Sigma}_0^{-1} \boldsymbol{\theta}_0\right) - \left(F(\mathbf{X})^T \boldsymbol{\Sigma}^{-1} \mathbf{z} + \boldsymbol{\Sigma}_0^{-1} \boldsymbol{\theta}_0\right)^T \boldsymbol{\theta}_1 \right\}\right) \\
 & \propto \exp\left(-\frac{1}{2}\left\{ \boldsymbol{\theta}_1^T \boldsymbol{\Sigma}_P^{-1} \boldsymbol{\theta}_1 - \boldsymbol{\theta}_1^T \boldsymbol{\Sigma}_P^{-1} \boldsymbol{\Sigma}_P \left(F(\mathbf{X})^T \boldsymbol{\Sigma}^{-1} \mathbf{z} + \boldsymbol{\Sigma}_0^{-1} \boldsymbol{\theta}_0\right) \right. \right. \\
 & \quad \left. \left. - \left(F(\mathbf{X})^T \boldsymbol{\Sigma}^{-1} \mathbf{z} + \boldsymbol{\Sigma}_0^{-1} \boldsymbol{\theta}_0\right)^T \boldsymbol{\Sigma}_P \boldsymbol{\Sigma}_P^{-1} \boldsymbol{\theta}_1 \right\}\right) \\
 & \propto \exp\left(-\frac{1}{2}(\boldsymbol{\theta}_1 - \boldsymbol{\theta}_P)^T \boldsymbol{\Sigma}_P^{-1}(\boldsymbol{\theta}_1 - \boldsymbol{\theta}_P)\right)
 \end{aligned}$$

by using the relations (III.5) and (III.6). The associated proper posterior distribution is then

$$\mathbb{P}(\boldsymbol{\theta} = \boldsymbol{\theta}_1 \mid \mathbf{Z} = \mathbf{z}) = \frac{1}{(2\pi)^{d/2} |\boldsymbol{\Sigma}_P|^{1/2}} \exp\left(-\frac{1}{2}(\boldsymbol{\theta}_1 - \boldsymbol{\theta}_P)^T \boldsymbol{\Sigma}_P^{-1}(\boldsymbol{\theta}_1 - \boldsymbol{\theta}_P)\right).$$

□

The strength of our method is that it provides an analytical *posterior* distribution. The prediction of thermodynamic quantity for a new experiment is possible and we provide a closed form expression for the variance of such prediction, see the next proposition.

Proposition 3. Consider a multivariate Gaussian posterior distribution $\mathcal{N}_d(\boldsymbol{\theta}_P, \boldsymbol{\Sigma}_P)$ in the linear model $\mathbf{Z} = F(\mathbf{X})\boldsymbol{\theta} + \boldsymbol{\varepsilon}$ and the likelihood function (III.4), the posterior predictive distribution for the estimated value of $F(x_{N+1})\boldsymbol{\theta}$ is the Gaussian distribution $\mathcal{N}(F(x_{N+1})\boldsymbol{\theta}_P, \sigma_{N+1}^2 + \sigma_P^2)$, with

$$\sigma_P^2 = F(x_{N+1})\boldsymbol{\Sigma}_P F(x_{N+1})^T. \quad (\text{III.7})$$

Proof. We first set the following definitions

$$\begin{aligned}
 \mu_n &= \Sigma_N \left(\frac{z_{N+1}}{\sigma_{N+1}^2} + \frac{F(x_{N+1})\boldsymbol{\theta}_P}{\sigma_P^2} \right) \\
 \Sigma_n &= \frac{\sigma_{N+1}^2 \sigma_P^2}{\sigma_{N+1}^2 + \sigma_P^2}.
 \end{aligned} \quad (\text{III.8})$$

The *posterior* prediction for a value z_{N+1} outside of the set of observations is given by

$$\begin{aligned}
& \mathbb{P}(F(x_{N+1})\boldsymbol{\theta} = z_{N+1} \mid \mathbf{Z} = \mathbf{z}) \\
&= \int_{\mathbb{R}^d} L(F(x_{N+1})\boldsymbol{\theta} = z_{N+1} \mid \boldsymbol{\theta} = \boldsymbol{\theta}_1) \mathbb{P}(F(x_{N+1})\boldsymbol{\theta} = F(x_{N+1})\boldsymbol{\theta}_1 \mid \mathbf{Z} = \mathbf{z}) d\boldsymbol{\theta}_1 \\
&= \int_{\mathbb{R}^d} \frac{1}{(2\pi\sigma_{N+1}^2)^{1/2}} \exp\left(-\frac{1}{2} \frac{(z_{N+1} - F(x_{N+1})\boldsymbol{\theta}_1)^2}{\sigma_{N+1}^2}\right) \\
&\quad \frac{1}{(2\pi\sigma_P^2)^{1/2}} \exp\left(-\frac{1}{2} \frac{(F(x_{N+1})\boldsymbol{\theta}_1 - F(x_{N+1})\boldsymbol{\theta}_P)^2}{\sigma_P^2}\right) d\boldsymbol{\theta}_1 \\
&= \int_{\mathbb{R}^d} \frac{1}{2\pi(\sigma_{N+1}^2\sigma_P^2)^{1/2}} \exp\left(-\frac{1}{2} A(z_{N+1}, F(x_{N+1}), \boldsymbol{\theta}_1, \sigma_{N+1}^2, \boldsymbol{\theta}_P, \sigma_P^2)\right) d\boldsymbol{\theta}_1.
\end{aligned}$$

Focusing on the term in the exponential multiplied by -2 ,

$$\begin{aligned}
& A(z_{N+1}, F(x_{N+1}), \boldsymbol{\theta}_1, \sigma_{N+1}^2, \boldsymbol{\theta}_P, \sigma_P^2) \\
&= \frac{(z_{N+1} - F(x_{N+1})\boldsymbol{\theta}_1)^2}{\sigma_{N+1}^2} + \frac{(F(x_{N+1})\boldsymbol{\theta}_1 - F(x_{N+1})\boldsymbol{\theta}_P)^2}{\sigma_P^2} \\
&= \frac{z_{N+1}^2}{\sigma_{N+1}^2} + \frac{(F(x_{N+1})\boldsymbol{\theta}_P)^2}{\sigma_P^2} + (F(x_{N+1})\boldsymbol{\theta}_1)^2 \left(\frac{1}{\sigma_{N+1}^2} + \frac{1}{\sigma_P^2}\right) \\
&\quad - 2(F(x_{N+1})\boldsymbol{\theta}_1) \left(\frac{z_{N+1}}{\sigma_{N+1}^2} + \frac{F(x_{N+1})\boldsymbol{\theta}_P}{\sigma_P^2}\right) \\
&= B(z_{N+1}, F(x_{N+1}), \sigma_{N+1}^2, \boldsymbol{\theta}_P, \sigma_P^2) + \frac{(F(x_{N+1})\boldsymbol{\theta}_1 - \mu_n)^2}{\Sigma_n}
\end{aligned} \tag{III.9}$$

by using the definitions set in the expressions (III.8). The last term in equation (III.9) is the only one depending on $\boldsymbol{\theta}_1$ and is dealt in the integration.

The first term $B(z_{N+1}, F(x_{N+1}), \sigma_{N+1}^2, \boldsymbol{\theta}_P, \sigma_P^2)$ can be rewritten as follows

$$\begin{aligned}
B(z_{N+1}, F(x_{N+1}), \sigma_{N+1}^2, \boldsymbol{\theta}_P, \sigma_P^2) &= \frac{z_{N+1}^2}{\sigma_{N+1}^2} + \frac{(F(x_{N+1})\boldsymbol{\theta}_P)^2}{\sigma_P^2} - \frac{\mu_n^2}{\Sigma_n} \\
&= \frac{z_{N+1}^2}{\sigma_{N+1}^2} + \frac{(F(x_{N+1})\boldsymbol{\theta}_P)^2}{\sigma_P^2} - \Sigma_n \left(\frac{z_{N+1}}{\sigma_{N+1}^2} + \frac{F(x_{N+1})\boldsymbol{\theta}_P}{\sigma_P^2}\right)^2 \\
&= \frac{(z_{N+1} - F(x_{N+1})\boldsymbol{\theta}_P)^2}{\sigma_{N+1}^2 + \sigma_P^2} \\
&\quad + z_{N+1}^2 \left(\frac{1}{\sigma_{N+1}^2} - \frac{1}{\sigma_{N+1}^2 + \sigma_P^2} - \frac{\Sigma_n}{\sigma_{N+1}^4}\right) \\
&\quad + (F(x_{N+1})\boldsymbol{\theta}_P)^2 \left(\frac{1}{\sigma_P^2} - \frac{1}{\sigma_{N+1}^2 + \sigma_P^2} - \frac{\Sigma_n}{\sigma_P^4}\right) \\
&= \frac{(z_{N+1} - F(x_{N+1})\boldsymbol{\theta}_P)^2}{\sigma_{N+1}^2 + \sigma_P^2}.
\end{aligned}$$

The last expression comes by noticing that

$$\frac{1}{\sigma_{N+1}^2 + \sigma_P^2} + \frac{\Sigma_n}{\sigma_{N+1}^4} = \frac{1}{\sigma_{N+1}^2 + \sigma_P^2} + \frac{\sigma_P^2}{\sigma_{N+1}^2(\sigma_{N+1}^2 + \sigma_P^2)} = \frac{1}{\sigma_{N+1}^2}$$

and

$$\frac{1}{\sigma_{N+1}^2 + \sigma_P^2} + \frac{\Sigma_n}{\sigma_P^4} = \frac{1}{\sigma_{N+1}^2 + \sigma_P^2} + \frac{\sigma_{N+1}^2}{\sigma_P^2 (\sigma_{N+1}^2 + \sigma_P^2)} = \frac{1}{\sigma_P^2}.$$

□

We give further in the Section III. 3 some examples and the advantages of the closed formulas obtained in the Proposition 2 and Proposition 3 for our thermodynamic framework.

III. 2.3 Sensitivity to the prior parameters

From Section III. 2.2, the *posterior* mean and variance in the case of the linear model and recalled below are function of the *prior* mean and variance, where *posterior* mean θ_p is given by

$$\theta_p = \Sigma_p \left(F(\mathbf{X})^T \Sigma^{-1} \mathbf{Z} + \Sigma_0^{-1} \theta_0 \right) \quad (\text{III.10})$$

and posterior covariance matrix Σ_p is given by

$$\Sigma_p = \left(F(\mathbf{X})^T \Sigma^{-1} F(\mathbf{X}) + \Sigma_0^{-1} \right)^{-1}. \quad (\text{III.11})$$

Establishing the sensitivity to the *prior* parameters consists in determining how much the values of the *posterior* distribution varies when the *prior* parameters values are modified.

III. 2.4 Levelled linear models

In the present section, we consider to deal with two sets of observations of different kinds. We denote these sets by vectors $\mathbf{Z}$ and $\mathbf{Y}$. We propose to model these observations using a system of linear models.

More formally written, we consider the following system of linear models

$$\begin{cases} \mathbf{Z} = F^{(1)}(\mathbf{X})\theta^{(1)} + F^{(2)}(\mathbf{X})\theta^{(2)} + \varepsilon_Z \\ \mathbf{Y} = G^{(1)}(\mathbf{X})\theta^{(1)} + \varepsilon_Y \end{cases} \quad (\text{III.12})$$

where $\theta^{(1)} \in \mathbb{R}^{d_1}$ and $\theta^{(2)} \in \mathbb{R}^{d_2}$ are two multidimensional unknown where ε_Z and ε_Y are independent centred Gaussian variables with covariance matrix Σ_Z and Σ_Y . The covariance matrix Σ_Z (respectively Σ_Y) represents the experiment uncertainty for observations $\mathbf{Z}$ (respectively $\mathbf{Y}$).

For the following proposition, we assume that the inverses of the covariance matrices Σ_Z and Σ_Y exist and that $F^{(1)}(\mathbf{X})$, $F^{(2)}(\mathbf{X})$ and $G^{(1)}(\mathbf{X})$ are full rank matrices.

Proposition 4. *Let us consider the linear model (III.12).*

Let the posterior distribution of $\theta^{(1)}$ knowing $\mathbf{Y}$ be $\mathcal{N}_{d_1}(\theta_P^{(1)}, \Sigma_P^{(1)})$.

Let Π^1 and Π^2 be defined by

$$\Pi^1 = \left((F^{(2)}(\mathbf{X}))^T F^{(2)}(\mathbf{X}) \right)^{-1} (F^{(2)}(\mathbf{X}))^T$$

and

$$\mathbf{\Pi}^2 = \mathbf{\Pi}^1 F^{(1)}(\mathbf{X}).$$

Then, the joint posterior distribution of $(\boldsymbol{\theta}^{(1)}, \boldsymbol{\theta}^{(2)})$ knowing $(\mathbf{Y}, \mathbf{Z})$ is $\mathcal{N}_{d_1+d_2}(\boldsymbol{\theta}_P, \boldsymbol{\Sigma}_P)$ where

$$\boldsymbol{\theta}_P = \begin{pmatrix} \boldsymbol{\theta}_P^{(1)} \\ \mathbf{\Pi}^1 \mathbf{Z} - \mathbf{\Pi}^2 \boldsymbol{\theta}_P^{(1)} \end{pmatrix}$$

and

$$\boldsymbol{\Sigma}_P = \begin{pmatrix} \boldsymbol{\Sigma}_P^{(1)} & -\boldsymbol{\Sigma}_P^{(1)}(\mathbf{\Pi}^2)^T \\ -\mathbf{\Pi}^2 \boldsymbol{\Sigma}_P^{(1)} & \boldsymbol{\Sigma}_Z + \mathbf{\Pi}^2 \boldsymbol{\Sigma}_P^{(1)}(\mathbf{\Pi}^2)^T \end{pmatrix}.$$

Proof. Using Bayes' theorem, we have that

$$\begin{aligned} & \mathbb{P}(\boldsymbol{\theta}^{(1)} = \boldsymbol{\theta}_1, \boldsymbol{\theta}^{(2)} = \boldsymbol{\theta}_2 \mid \mathbf{Y} = \mathbf{y}, \mathbf{Z} = \mathbf{z}) \\ & \propto \mathbb{P}(\boldsymbol{\theta}^{(1)} = \boldsymbol{\theta}_1 \mid \mathbf{Y} = \mathbf{y}) \mathbb{P}(\boldsymbol{\theta}^{(2)} = \boldsymbol{\theta}_2 \mid \boldsymbol{\theta}^{(1)} = \boldsymbol{\theta}_1, \mathbf{Z} = \mathbf{z}) \\ & \propto \exp\left(-\frac{1}{2}(\boldsymbol{\theta}_1 - \boldsymbol{\theta}_P^{(1)})^T \boldsymbol{\Sigma}_P^{(1)-1}(\boldsymbol{\theta}_1 - \boldsymbol{\theta}_P^{(1)})\right) \\ & \quad \exp\left(-\frac{1}{2}(\mathbf{z} - F^{(1)}(\mathbf{X})\boldsymbol{\theta}_1 - F^{(2)}(\mathbf{X})\boldsymbol{\theta}_2)^T \boldsymbol{\Sigma}_Z^{-1}(\mathbf{z} - F^{(1)}(\mathbf{X})\boldsymbol{\theta}_1 - F^{(2)}(\mathbf{X})\boldsymbol{\theta}_2)\right) \\ & \propto \exp\left(-\frac{1}{2}(\boldsymbol{\theta}_1 - \boldsymbol{\theta}_P^{(1)})^T \left((\mathbf{\Pi}^2)^T \boldsymbol{\Sigma}_Z^{-1} \mathbf{\Pi}^2 + \boldsymbol{\Sigma}_P^{(1)-1}\right)(\boldsymbol{\theta}_1 - \boldsymbol{\theta}_P^{(1)})\right) \\ & \quad \exp\left(-\frac{1}{2}(\boldsymbol{\theta}_2 - \mathbf{\Pi}^1 \mathbf{z} + \mathbf{\Pi}^2 \boldsymbol{\theta}_P^{(1)})^T \boldsymbol{\Sigma}_Z^{-1}(\boldsymbol{\theta}_2 - \mathbf{\Pi}^1 \mathbf{z} + \mathbf{\Pi}^2 \boldsymbol{\theta}_P^{(1)})\right) \\ & \quad \exp\left(\frac{1}{2}(\boldsymbol{\theta}_2 - \mathbf{\Pi}^1 \mathbf{z} + \mathbf{\Pi}^2 \boldsymbol{\theta}_P^{(1)})^T \boldsymbol{\Sigma}_Z^{-1} \mathbf{\Pi}^2 (\boldsymbol{\theta}_1 - \boldsymbol{\theta}_P^{(1)})\right) \\ & \quad \exp\left(\frac{1}{2}(\boldsymbol{\theta}_1 - \boldsymbol{\theta}_P^{(1)})^T (\mathbf{\Pi}^2)^T \boldsymbol{\Sigma}_Z^{-1}(\boldsymbol{\theta}_2 - \mathbf{\Pi}^1 \mathbf{z} + \mathbf{\Pi}^2 \boldsymbol{\theta}_P^{(1)})\right). \end{aligned}$$

We denote by $\boldsymbol{\Sigma}_P^{-1}$ the following matrix

$$\boldsymbol{\Sigma}_P^{-1} = \begin{pmatrix} (\mathbf{\Pi}^2)^T \boldsymbol{\Sigma}_Z^{-1} \mathbf{\Pi}^2 + \boldsymbol{\Sigma}_P^{(1)-1} & (\mathbf{\Pi}^2)^T \boldsymbol{\Sigma}_Z^{-1} \\ \boldsymbol{\Sigma}_Z^{-1} \mathbf{\Pi}^2 & \boldsymbol{\Sigma}_Z^{-1} \end{pmatrix}.$$

The matrix $\boldsymbol{\Sigma}_P$ is obtained by using the Schur complement for block matrix inversion, see details in [95], provided that the matrix $\left((\mathbf{\Pi}^2)^T \boldsymbol{\Sigma}_Z^{-1} \mathbf{\Pi}^2 + \boldsymbol{\Sigma}_P^{(1)-1}\right)$ can be inverted. In the end, we have

$$\boldsymbol{\Sigma}_P = \begin{pmatrix} \boldsymbol{\Sigma}_P^{(1)} & -\boldsymbol{\Sigma}_P^{(1)}(\mathbf{\Pi}^2)^T \\ -\mathbf{\Pi}^2 \boldsymbol{\Sigma}_P^{(1)} & \boldsymbol{\Sigma}_Z + \mathbf{\Pi}^2 \boldsymbol{\Sigma}_P^{(1)}(\mathbf{\Pi}^2)^T \end{pmatrix}.$$

□

III. 3 Bayesian point of view for the Calphad-based modelling

III. 3.1 Motivations

The Calphad method is a functional estimation process that uses assimilation data of different natures. In the Calphad-based modelling, the assessor targets the determination of a number r of Gibbs energy functions. Such functions are essential in the comprehension and the modelling of many problems from physics, thermodynamics and chemistry.

To be more precise, the chemical species can have different structural arrangements. These arrangements are called *phases*. Each phase α has its inner Gibbs energy function, we denote in the following by $G^{(\alpha)}$ for the energy function associated to the phase α .

These functions are modelled using a general regression model. We give below with **Example 1** an introducing example of a reconstruction problem for the Calphad method. We will generalise more in the next section.

Example 1. *As an example, let us consider a fictive case with $r = 2$ Gibbs energy functions, with a liquid phase which function is denoted by $G^{(\alpha_1)}$ and a compound phase which function is denoted by $G^{(\alpha_2)}$. We assume we known $G^{(\alpha_1)}$ but not $G^{(\alpha_2)}$. The model for the unknown function $G^{(\alpha_2)}$ is a function of the temperature $T > 0$, we can write as follows*

$$G^{(\alpha_2)}(T) = \theta_1^{(\alpha_2)} + \theta_2^{(\alpha_2)} T + \theta_3^{(\alpha_2)} T \log(T) + \theta_4^{(\alpha_2)} T^2 + \theta_5^{(\alpha_2)} \frac{1}{T}, \quad \text{for } T \in [T_{\min}; T_{\max}]$$

where $\theta_1^{(\alpha_2)}, \dots, \theta_5^{(\alpha_2)}$ are the components of an unknown vector $\boldsymbol{\theta}^{(\alpha_2)} \in \mathbb{R}^5$ to assess.

One specificity of the Calphad-based modelling is the complicated structure of the assimilation data. Indeed, the assimilation data are not values of the quantity of interest G but can be related to G either by a linear transformation or by more complicated transformations. Taking back the **Example 1**, we do not observe values of $G^{(\alpha_2)}$. So in order to estimate $\boldsymbol{\theta}^{(\alpha_2)}$, one may use heat capacities data, denoted by $C_p^{(\alpha_2)}$, mixing enthalpy data, denoted by $H_{mix}^{(\alpha_2)}$, and a condition about $G^{(\alpha_2)}$. The quantities $C_p^{(\alpha_2)}$ and $H_{mix}^{(\alpha_2)}$ are related to the quantity of interest by the following relations

$$C_p^{(\alpha_2)}(T) = -T \frac{\partial^2 G^{(\alpha_2)}}{\partial T^2}(T)$$

and

$$H_{mix}^{(\alpha_2)}(T) = G^{(\alpha_2)}(T) + T \frac{\partial G^{(\alpha_2)}}{\partial T}(T).$$

In those cases, the assimilation data are values of $C_p^{(\alpha_2)}$ and $H_{mix}^{(\alpha_2)}$ for given values of T .

Those quantities alone does not permit to determine all the components of $\boldsymbol{\theta}^{(\alpha_2)}$, that is why we also need a condition about $G^{(\alpha_2)}$. One may use the melting temperature T_0 to give the following condition to be satisfied by $G^{(\alpha_1)}$ and $G^{(\alpha_2)}$

$$G^{(\alpha_1)}(T_0) = G^{(\alpha_2)}(T_0).$$

In that case, the assimilation data is the value T_0 where the transition occurs, that is the value T_0 for which we have the above condition.

For more details about the assimilation data used in the Calphad method, we refer the reader to the Section II. 4.3 where we list the different relations that exist between G and the data used to determine G .

In the assimilation process - that is when fitting the data - some errors may be introduced in the model, either by the measurement uncertainty of the data or error in the modelling of the G^α . The uncertainty propagation aims to understand the effects on the computed data of these errors. Despite the importance of an uncertainty study, its integration in the Calphad-based modelling is still in its early stages. In our work, we propose a methodology to study the uncertainty propagation using a Bayesian approach within the linear model.

The Bayesian approach is a statistical paradigm which features both a likelihood function and a *prior* distribution. On one hand, the likelihood function quantifies the deviation of the data from the regression model. On the other hand, the *prior* distribution provides anterior knowledge on the parameters. These elements combined with Bayes' theorem enable to compute the so-called *posterior* distribution. The *posterior* distribution provides the probability distribution for the parameters values considering we have observed a given data base. The interested reader is referred to [72] for a comprehensive approach of the Bayesian paradigm and to the Section III. 2 of this chapter.

In our work, we propose a Bayesian approach which has not been applied yet in the Calphad-based modelling. We use Gaussian conjugate *prior* distributions which are well fitted for the Gibbs energy function framework when those functions are modelled with the linear model. Our approach benefits from the linear model as it allows having an analytical posterior distribution for the model parameters. Therefore an analytical formula to provide credible intervals for the model parameters and calculated thermodynamic data can be directly computed. The approach also permits to get a fast estimation of the uncertainty on the computed phase diagram data.

The method is applied to the Cu-Mg system.

III. 3.2 Calphad-based modelling

We remind to the reader some modelling elements of the Chapter II.

In our work, we consider that we can model the Gibbs energy function with the linear model. That is that the Gibbs functions $G^{(\alpha)}$, for $\alpha = 1, \dots, r$ can be expressed by

$$G^{(\alpha)}(\mathbf{E}, \boldsymbol{\gamma}^{(\alpha)}) = \sum_{p=1}^{p_\alpha} \theta_p^{(\alpha)} F_p^{(\alpha)}(\mathbf{E}, \boldsymbol{\gamma}^{(\alpha)}) + B^{(\alpha)}(\mathbf{E}, \boldsymbol{\gamma}^{(\alpha)}), \quad \forall \alpha \in [1; r] \quad (\text{III.13})$$

for the experimental condition $\mathbf{E} \in U \subset \mathbb{R}^d$ with U compact and the structure $\boldsymbol{\gamma}^{(\alpha)} \in [0; 1]^{d_\alpha}$, see Section II. 2 and Section II. 4 for the description of these quantities. The function $B^{(\alpha)}$ is a known function of $\mathbf{E}$ and $\boldsymbol{\gamma}^{(\alpha)}$. The regression functions $F_p^{(\alpha)}$ are basis functions chosen by the assessor.

For all $\alpha = 1, \dots, r$, the vector $\boldsymbol{\theta}^{(\alpha)} = (\theta_1^{(\alpha)}, \dots, \theta_{p_\alpha}^{(\alpha)})$ is the vector of unknown parameters associated to the Gibbs function $G^{(\alpha)}$. More generally, we will denote by $\boldsymbol{\theta} = (\boldsymbol{\theta}^{(1)T}, \dots, \boldsymbol{\theta}^{(r)T})^T$ the vector of unknown parameters for the whole system.

As mentioned before, one specificity of the Calphad-based modelling is the complicated structure of the data. Instead of direct observations of the quantity $G^{(\alpha)}$, one can use two kind of observations that are called *thermodynamic data* and *phase diagram data*.

The *thermodynamic data* are related to the $G^{(\alpha)}$ by a linear transformation. By denoting with Z a thermodynamic observation for a given α and a given value of $(\mathbf{E}, \boldsymbol{\gamma}^{(\alpha)})$, we can write for the j -th experiment

$$Z_j^{(\alpha)} = \mathcal{T}_{j,\alpha} \left(G^{(\alpha)}(\mathbf{E}_j, \boldsymbol{\gamma}_j^{(\alpha)}) \right) + \varepsilon_j$$

where ε_j represents the measure uncertainty associated to the observation $Z_j^{(\alpha)}$ and where $\mathcal{T}_{j,\alpha}$ is the linear transformation of the Gibbs function $G^{(\alpha)}$ corresponding to the observed data.

More generally, a thermodynamic observation can be related to a mix of several phases. We denote the mix values of phases with a vector $\boldsymbol{\lambda} \in [0; 1]^r$. Then, we can write for a given value of $(\boldsymbol{\lambda}_j, \mathbf{E}_j, \boldsymbol{\gamma}_j)$ the j -th observation

$$Z_j = \sum_{\alpha=1}^r \lambda_j^{(\alpha)} \mathcal{T}_{j,\alpha} \left(G^{(\alpha)}(\mathbf{E}_j, \boldsymbol{\gamma}_j^{(\alpha)}) \right) + \varepsilon_j$$

where ε_j represents the uncertainty associated to the observation Z_j and where $\mathcal{T}_{j,\alpha}$ is the linear transformation of the Gibbs function $G^{(\alpha)}$ corresponding to the observed data.

In the sequel, we will use this last very general formulation.

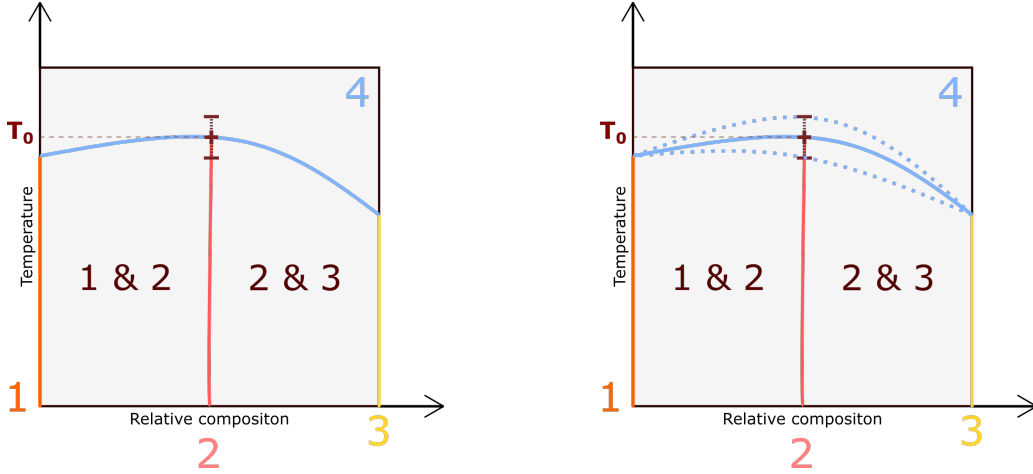
Notice that with the assumption of a linear model for $G^{(\alpha)}$, the thermodynamic data write

$$Z_j = \sum_{\alpha=1}^r \lambda_j^{(\alpha)} \left(\sum_{p=1}^{p_\alpha} \theta_p^{(\alpha)} \tilde{F}_p^{(\alpha)}(\mathbf{E}_j, \boldsymbol{\gamma}_j^{(\alpha)}) + \tilde{B}^{(\alpha)}(\mathbf{E}_j, \boldsymbol{\gamma}_j^{(\alpha)}) \right) + \varepsilon_j$$

where $\tilde{F}_p^{(\alpha)}$ (respectively $\tilde{B}^{(\alpha)}$) is the regression function $F_p^{(\alpha)}$ (respectively is the known bias function $\tilde{B}^{(\alpha)}$) to which the linear transformation $\mathcal{T}_{j,\alpha}$ is applied. The model remain a linear model with an unknown parameter $\boldsymbol{\theta}$.

The phase diagram is a map of stable phases, that is the phase - or group of phases - with the lowest Gibbs energy, see Section II. 4.1 and Section II. 4.3.3. The *phase diagram data* consists in a constraint on the $G^{(\alpha)}$ for a given value of $(\boldsymbol{\lambda}, \mathbf{E}, \boldsymbol{\gamma})$.

As an illustration, see the fictive example proposed in the Figure III.1. The example features 3 compounds (denoted 1, 2 and 3) and a liquid phase (denoted by 4). A melting temperature with an error bar is measured for compound 2 and is denoted by T_0 .

(a) Uncertainty on the melting temperature T_0 .

(b) Uncertainty on the transition curve.

FIGURE III.1 – Example of phase diagram uncertainties with 4 phases. The Figure III.1a displays the error bar on the melting temperature T_0 of phase 2. In the Figure III.1b, the uncertainty is reported on the whole phase transition curve associated to the apparition of phase 4 (in light blue). The plain line corresponds to the transition curve associated to the T_0 value. The dotted lines correspond to the transition curves associated to the extremal values of the error bar.

III. 3.3 Literature review of the uncertainty analysis for the specific framework of computational thermodynamics

Before getting to the details of the method proposed in our work, we first recall some early and more recent works that target the uncertainty integration in the Calphad-based modelling. In all these parametric methods, the unknown parameter to assess will be denoted by θ , the vector of thermodynamic observations by $\mathbf{Z} \in \mathbb{R}^N$, where its components are modelled as follows

$$\begin{aligned} Z_j &= \sum_{\alpha=1}^r \lambda_j^{(\alpha)} \mathcal{T}_{j,\alpha} \left(G^{(\alpha)}(\mathbf{E}_j, \gamma_j^{(\alpha)}) \right) + \varepsilon_j \\ &= z_j(\boldsymbol{\lambda}_j, \mathbf{E}_j, \boldsymbol{\gamma}_j) + \varepsilon_j \end{aligned}$$

for a given value of $(\boldsymbol{\lambda}_j, \mathbf{E}_j, \boldsymbol{\gamma}_j)$ and with ε_j the measure uncertainty associated to the observation Z_j , for $j = 1, \dots, N$. In the following, we denote by $z_j(\boldsymbol{\lambda}_j, \mathbf{E}_j, \boldsymbol{\gamma}_j)$ the thermodynamic quantity predicted by the Calphad model and by $\mathbf{z}$ the vector with components $z_j(\boldsymbol{\lambda}_j, \mathbf{E}_j, \boldsymbol{\gamma}_j)$, $j = 1, \dots, N$. Notice that some authors do not use the linear model for the Gibbs functions $G^{(\alpha)}$ but other general parametric models. In some of these works, the authors can also be interested in the phase transition obtained with the fitted model. For example see the work of [67] in which the authors propose a propagation of the parameter uncertainties to the phase diagram in order to obtain a probability of stability for the concerned phases. Such phase transition function is denoted by $\psi_E(\hat{\theta})$ in the following. We recall in the Appendix B some classical algorithms used for the bayesian estimation and that are mentioned in the following.

One of the earliest works in uncertainty estimation and propagation is presented in [59].

The aim is the estimation of confidence intervals for computed phase transitions. The author proposes an iterative method based on a Taylor development of the regression model residues, that is the deviation of the measured thermodynamic quantities Z_j from the Calphad model $z_j(\lambda_j, E_j, \gamma_j)$ of such quantity. The covariance matrix of the estimated value of θ is approximated by the Jacobian matrix of the residues with respect to the parameter and evaluated at the final iteration of the method. The uncertainty is then propagated to the phase diagram quantities using the Taylor development, that is, by denoting by $\psi_E(\cdot)$ the computed phase transition, the variance of quantity $\psi_E(\hat{\theta})$ is approximated by

$$\begin{pmatrix} \frac{\partial \psi_E(\hat{\theta})}{\partial \theta_1} & \dots & \frac{\partial \psi_E(\hat{\theta})}{\partial \theta_p} \end{pmatrix} \text{Cov}(\hat{\theta}) \begin{pmatrix} \frac{\partial \psi_E(\hat{\theta})}{\partial \theta_1} \\ \vdots \\ \frac{\partial \psi_E(\hat{\theta})}{\partial \theta_p} \end{pmatrix}.$$

One of the first Bayesian approach for the unknown parameter in the Gibbs energy modelling appears in [53]. The author considers a linear model in which the unknown parameter is fitted by a sequential Bayesian estimation. In order to estimate the values of the unknown parameters, one must define an objective function. Such objective function depending on the parameters values and consists in a criterion to minimise¹. The values of the parameters that realise the minimum of the objective function are chosen for the parameter estimation. In [53], the objective function considered consists in minimising a quadratic criterion penalised by a *prior* information. It is expressed as follows

$$(\mathbf{z} - \mathbf{Z})^T \mathbf{C}_{\mathbf{P},\mathbf{z}}^{-1} (\mathbf{z} - \mathbf{Z}) + (\boldsymbol{\theta} - \boldsymbol{\theta}_0)^T \mathbf{C}_0^{-1} (\boldsymbol{\theta} - \boldsymbol{\theta}_0)$$

where $\boldsymbol{\theta}_0$ and $\mathbf{C}_0$ are respectively the mean vector and the covariance matrix of the *prior* distribution, $\mathbf{Z}$ is the vector of observations used to fit the model, $\mathbf{z}$ is the vector of corresponding predicted quantities and $\mathbf{C}_{\mathbf{P},\mathbf{z}}$ is the estimation of the *posterior* covariance matrix associated to quantity $\mathbf{z}$. Quantities $\boldsymbol{\theta}$ and $\mathbf{C}_{\mathbf{P},\mathbf{z}}$ are updated by a sequential algorithm based on an algorithm coming from control theory and communication that is the so-called extended Kalman filter method, see [50] for more details about this method.

A common choice of objective function for the Calphad-based modelling is the weighted least squares minimisation, that is one wants to minimise on the parameter $\boldsymbol{\theta}$ the value of

$$\sum_{j=1}^N w_j (z_j - Z_j)^2$$

where Z_j is the j -th observation and z_j is the corresponding computed quantity for the fitted model. Usually, the assessors use observations that come from multiple experiments, that is several datasets. The issue when using weighted objective function is how to choose the weights

1. In our case, we consider objective function to minimise but depending on the context one may deal with objective function to maximise

in such case. The authors in [97] propose a new way to determine the value of weights for each datasets. The method is based on the computation of cross-validation scores that measure the quality of prediction of a fitted model. The authors propose to use equal weights for all observations coming from the same dataset. We denote by $Z_{i,j}$ the j -th experiment of dataset i , $\mathbf{Z}^i$ the i -th dataset and $z_{i,j}^l$ the computed quantity corresponding to observation $Z_{i,j}$ using the regression model fitted with database $\mathbf{Z}^l$. The cross validation score for the regression model based on $\mathbf{Z}^l$, is then

$$s^l = \frac{1}{\sum_{i \neq l} N_i} \sum_{i \neq l} \sum_{j=1}^{N_i} (Z_{i,j} - z_{i,j}^l)^2.$$

The authors finally set weights w_l with

$$w_l = 1 - \frac{s^l - \min(s^i, i = 1, \dots, k)}{\max(s^i, i = 1, \dots, k)}, \quad l = 1, \dots, k.$$

Starting in the 90s, iterative experimental methods have become popular for uncertainty estimation and propagation. Following the algorithms developments and improvements in statistics, the first algorithms to be used are algorithms developed with the Markov Chain Monte Carlo (MCMC) methods. MCMC methods are a class of algorithms used to sample random variables values from a probability distribution. We call state space of the Markov chain the set of all possible values that can be sampled from the probability distribution. In the bayesian framework, such algorithms can be used to determine the *posterior* distribution for $\boldsymbol{\theta}$ knowing the assimilation data. In these algorithms, the Markov chain property of the method comes from the fact that in the chain of the N realisations $\boldsymbol{\theta}^{(1)}, \dots, \boldsymbol{\theta}^{(N)}$, solely the quantity $\boldsymbol{\theta}^{(i-1)}$ is used to generate the quantity $\boldsymbol{\theta}^{(i)}$. One wishes that, as the number of MCMC samplings N grows, the algorithm tends to sample values of $\boldsymbol{\theta}$ from an invariant measure which is the *posterior* distribution.

We give below some details about the MCMC algorithms to generate the values of $\boldsymbol{\theta}$. One important definition in the following algorithms is the definition of a conditional density function, that is the probability density function of an event providing that another event occurs. We denote in the following the conditional density function by $q(\cdot | \boldsymbol{\theta}^{(i)})$, to sample a new parameter value $\boldsymbol{\theta}$ from the parameter i -th value $\boldsymbol{\theta}^{(i)}$.

The first MCMC algorithm to be known is the Metropolis-Hastings algorithm which has been used for the Calphad-based modelling in the works [29] and [44]. The advantage of Metropolis-Hastings method is that it requires to know the *posterior* distribution $\mathcal{P}$ up to a constant. The Metropolis-Hastings algorithm is briefly recalled in Algorithm 2. See [61] for a more detailed description and a convergence analysis.

The Metropolis-Hasting algorithm can be poorly suited for the exploration of the space state of the Markov chain when the latter is thin in one direction. One challenging case is for example when dealing with a multimodal distribution, that is when considering a probability distribution which has several local maxima and minima. The authors of [68] propose to deal with improved versions of the algorithm that remedy that issue.

An empirical model selection is considered in [44]. The authors make a choice of n linear models $M^{(i)}$, $i = 1, \dots, n$ for the Gibbs energy functions with different basis functions and different ranks. The authors also define a reference linear model $M^{(0)}$. Such model $M^{(0)}$ is used to compare the different models $M^{(i)}$. The model selection is based on the Bayes' factor between the two models $M^{(0)}$ and $M^{(i)}$. Such factor is defined as the ratio R of the data likelihood in model $M^{(i)}$ over the data likelihood in model $M^{(0)}$, that is

$$R = \frac{\mathbb{P}(D|M^{(i)})}{\mathbb{P}(D|M^{(0)})} = \frac{\int \mathbb{P}(\boldsymbol{\theta}^{(i)}|M^{(i)})\mathbb{P}(D|\boldsymbol{\theta}^{(i)}, M^{(i)}) d\boldsymbol{\theta}^{(i)}}{\int \mathbb{P}(\boldsymbol{\theta}^{(0)}|M^{(0)})\mathbb{P}(D|\boldsymbol{\theta}^{(0)}, M^{(0)}) d\boldsymbol{\theta}^{(0)}} = \frac{\mathbb{P}(M^{(i)}|D)}{\mathbb{P}(M^{(0)}|D)} \frac{\mathbb{P}(M^{(0)})}{\mathbb{P}(M^{(i)})}.$$

In the end, the model $M^{(i)}$ with the highest Bayes' factor with respect to reference model $M^{(0)}$ is selected.

In [66], the authors propose to use Akaike Information Criterion (AIC) (see [2]) for model selection in the frame of the linear model. In their work, the dimension of the parameter $\boldsymbol{\theta}$ for the linear model is no longer fixed. The idea is to choose the dimension of $\boldsymbol{\theta}$ in agreement with the number of observations at hand. The objective function considered is the following one

$$\min_{\boldsymbol{\theta}} \left\{ 2 \dim(\boldsymbol{\theta}) + N \log(\|\mathbf{Z} - \mathbf{z}\|^2) \right\}$$

where N is the number of observations used to fit the model and $\dim(\boldsymbol{\theta})$ is the number of features for parameter $\boldsymbol{\theta}$.

Authors in [65], [17] and then in [18] proposed to use another MCMC algorithm that is called the Gibbs sampler. Gibbs sampler algorithm (Algorithm 3) can be used in the case of full known conditional *posterior* distributions $\{p(\theta_j|\theta_{i \neq j}, \mathbf{Z})\}$.

The authors in [18] already pointed out the lack of ease when considering *posterior* distributions which are not analytical. In their study, they consider a Gaussian *prior* distribution with a uniform likelihood function. In this case, the *posterior* distribution is a truncated Gaussian distribution from which it can be challenging to sample random variables. They propose an approximation of the *posterior* uncertainty using another analytical distribution, namely the χ^2 distribution.

As an alternative of the use of MCMC methods, genetic algorithms have been used in [88], [70] and [71] to compute the posterior distribution of parameter $\boldsymbol{\theta}$. Genetic algorithms fall into the much broader category of evolutionary algorithms. Such algorithms have been named this way as they attempt to simulate the biological evolution processes within the optimisation framework. Genetic algorithms are used to optimise non-smooth multivariate objective function. Each feature of the parameter $\boldsymbol{\theta}$ to optimise is seen as a sequence of k binary variables which can take as values 0 or 1. The algorithm is shortly described below in Algorithm 4. For a comprehensive study, see [39].

The authors in [88] propose an estimation of phase diagram uncertainty for the two binary systems UO_2 - PuO_2 and UO_2 - BeO . In the cases considered by the authors, there exists an expli-

cit expression in parameter $\boldsymbol{\theta}$ of phase diagram curves, *liquidus*² and *solidus*³ boundary, which eases the propagation. All studies featuring the chemical species considered in the application are aggregated in order to define a range of feasible parameters. Such feasibility range is used to determine the *prior* distribution.

III. 3.4 Uncertainty for the thermodynamic quantities

The assumption that the Gibbs energy functions follow a linear model makes easy the uncertainty propagation to the thermodynamic quantities. As all thermodynamic quantities can be expressed as a linear transformation of parameter $\boldsymbol{\theta}_p$, the Proposition 3 can be directly used by replacing basis function for the Gibbs energy function F by its equivalent $\tilde{F}$ for the target thermodynamic quantity.

Taking back the **Example 1**, the following Gibbs energy function is considered for the stoichiometric compound

$$G^{\alpha_2}(T) = \theta_1^{\alpha_2} + \theta_2^{\alpha_2} T + \theta_3^{\alpha_2} T \log(T) + \theta_4^{\alpha_2} T^2 + \theta_5^{\alpha_2} \frac{1}{T}, \quad \text{for } T \in [T_{\min}; T_{\max}].$$

The associated heat capacity function is defined by

$$C_p^{\alpha_2}(T) = -T \frac{\partial^2 G^{\alpha_2}}{\partial T^2}(T),$$

that is

$$C_p^{\alpha_2}(T) = -\theta_3^{\alpha_2} - 2\theta_4^{\alpha_2} T - \theta_5^{\alpha_2} \frac{2}{T^2}.$$

In that example, the design matrix is a vector $\mathbf{T} = (T_1, \dots, T_N)$ of the N temperature values for each measure. We denote by $F(\mathbf{T})$ the following matrix

$$F(\mathbf{T}) = \begin{pmatrix} 1 & T_1 & T_1 \log(T_1) & T_1^2 & \frac{1}{T_1} \\ \vdots & \vdots & \vdots & \vdots & \vdots \\ 1 & T_N & T_N \log(T_N) & T_N^2 & \frac{1}{T_N} \end{pmatrix}.$$

Then, the *posterior* mean and covariance matrix for a heat capacity quantity for the experiments $\mathbf{T}$ are then given by

$$\tilde{F}(\mathbf{T})\boldsymbol{\theta}_p = \tilde{F}(\mathbf{T})\boldsymbol{\Sigma}_p \left(F(\mathbf{T})^T \boldsymbol{\Sigma}^{-1} \mathbf{Z} + \boldsymbol{\Sigma}_0^{-1} \boldsymbol{\theta}_0 \right)$$

and

$$\tilde{F}(\mathbf{T})\boldsymbol{\Sigma}_p \tilde{F}(\mathbf{T})^T = \tilde{F}(\mathbf{T}) \left(F(\mathbf{T})^T \boldsymbol{\Sigma}^{-1} F(\mathbf{T}) + \boldsymbol{\Sigma}_0^{-1} \right)^{-1} \tilde{F}(\mathbf{T})^T$$

2. It consists in the curve of the phase diagram above which all chemical species are in a liquid phase.

3. At the contrary, it consists in the curve of the phase diagram under which all chemical species are in a solid phase.

where matrix $\tilde{F}(\mathbf{T})$ is given by

$$\tilde{F}(\mathbf{T}) = (-1) \times \begin{pmatrix} 0 & 0 & 1 & 2T_1 & \frac{2}{T_1^2} \\ \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 1 & 2T_N & \frac{2}{T_N^2} \end{pmatrix}.$$

III. 3.5 Uncertainty for the phase diagram

The uncertainty for phase diagram data is estimated through Monte-Carlo simulations. The method consists in sampling a great number N of parameter $\boldsymbol{\theta}$ values from the posterior distribution. For each parameter value $\boldsymbol{\theta}_i$ the associated phase diagram is computed. The boundaries of each phase diagrams obtained for each sampled parameter $\boldsymbol{\theta}_i$ are denoted by $\mathcal{S}_i$, where $\mathcal{S}_i$ is the set of coordinates of the phase diagram where a phase transition occurs for the i -th draw. The algorithm presented below provides an estimation of credible area of level $1 - \alpha$ for phase diagram boundaries $\mathcal{S}$ values and denoted by $\mathcal{C}_{\mathcal{S},\alpha}$. Note that $\mathcal{C}_{\mathcal{S},\alpha}$ is only an estimation and is determined experimentally.

Algorithm 1 Phase diagram uncertainty by Monte-Carlo simulations.

Compute $\boldsymbol{\theta}_p$ from (III.5) and $\boldsymbol{\Sigma}_p$ from (III.6).

For i in $[1; N]$:

- . Draw $\boldsymbol{\theta}_{p_i}$ from $\mathcal{N}_d(\boldsymbol{\theta}_p, \boldsymbol{\Sigma}_p)$.
- . Compute phase diagram associated to $\boldsymbol{\theta}_{p_i} \rightarrow$ evaluate $\mathcal{S}_i$.

Determine $\mathcal{C}_{\mathcal{S},\alpha}$ the set such that $\mathbb{P}(\mathcal{S}_i \in \mathcal{C}_{\mathcal{S},\alpha}) \geq 1 - \alpha$.

III. 4 Application to the binary system Cu-Mg case study

III. 4.1 Case study modelling

This section recalls the modelling of the binary system Cu-Mg studied here. The full description can be found in the Section II. 5. The system features five stable phases over the range of temperature $[300K; 1400K]$. There are one liquid phase and four solid phases : the FCC phase for pure Cu, the HCP phase for pure Mg, the Laves C15 phase for the solid compound Cu_2Mg and the stoichiometric compound CuMg_2 .

The Cu-Mg system is a textbook case as it provides a high number of experiments and data. The data selection has been made by the organizers of the SATA⁴ on thermodynamic assessments. The complete list of references for the data used in the case study can be found in the Table A.1 of the Appendix A.

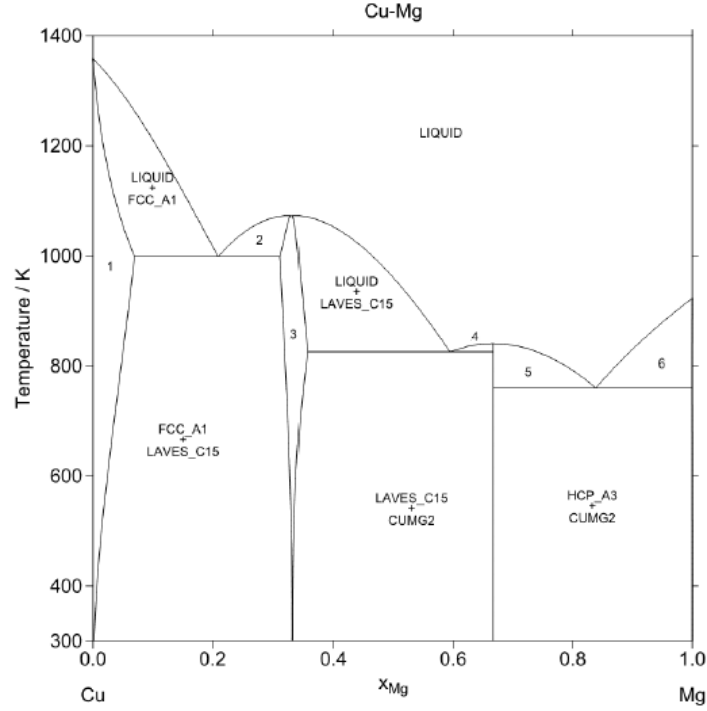


FIGURE III.2 – Phase diagram of the Cu-Mg binary system obtained at the SATA School. Unnamed phase diagram areas : 1. FCC. 2. Liquid and + Laves. 3. Laves. 4. Liquid and CuMg₂. 5. Liquid and CuMg₂. 6. Liquid and HCP.

Let us first recall the general model used for the Gibbs functions $G^{(\alpha)}$. For $\alpha = 1, \dots, r$, $G^{(\alpha)}$ is modelled by

$$G^{(\alpha)}(\mathbf{E}, \gamma^{(\alpha)}) = \sum_{p=1}^{p_\alpha} \theta_p^{(\alpha)} F_p^{(\alpha)}(\mathbf{E}, \gamma^{(\alpha)}) + B^{(\alpha)}(\mathbf{E}, \gamma^{(\alpha)}), \quad \forall \alpha \in [1; r] \quad (\text{III.14})$$

with $\mathbf{E} \in U \subset \mathbb{R}^d$ is the experimental condition, with U compact, and $\gamma^{(\alpha)} \in [0; 1]^{d_\alpha}$ the chemical structure. We detail further the modelling choice for each phase α of the system. In all the following expressions R is the gas constant, approximately equal to $8.315 \text{ J.K}^{-1}.\text{mol}^{-1}$.

The chosen model for the liquid phase (denoted by $\alpha = liq$), the FCC phase (denoted by $\alpha = A1$) and the HCP phase (denoted by $\alpha = A3$) is the sub-regular solid solution model. The known bias term has the following shape

$$B^{(\alpha)}(\mathbf{E}, \gamma^{(\alpha)}) = \gamma_{Cu}^{(\alpha)} G_{Cu}^{o,(\alpha)}(\mathbf{E}) + \gamma_{Mg}^{(\alpha)} G_{Mg}^{o,(\alpha)}(\mathbf{E}) + RT \left(\gamma_{Cu}^{(\alpha)} \log \left(\gamma_{Cu}^{(\alpha)} \right) + \gamma_{Mg}^{(\alpha)} \log \left(\gamma_{Mg}^{(\alpha)} \right) \right), \quad \gamma^{(\alpha)} \in [0; 1]^2$$

where $G_{Cu}^{o,(\alpha)}$ and $G_{Mg}^{o,(\alpha)}$, for $\alpha \in \{A1, A3, liq\}$, are the Gibbs energy functions used for the pure elements and provided by [27].

For the liquid phase, the unknown part is modelled by

$$\sum_{p=1}^{p_{liq}} \theta_p^{(liq)} F_p^{(liq)}(\mathbf{E}, \boldsymbol{\gamma}^{(liq)}) = \gamma_{Cu}^{(liq)} \gamma_{Mg}^{(liq)} \left\{ \theta_1^{(liq,0)} + \theta_2^{(liq,0)} T + \theta_1^{(liq,1)} \left(\gamma_{Cu}^{(liq)} - \gamma_{Mg}^{(liq)} \right) \right\}, \quad \boldsymbol{\gamma}^{(liq)} \in [0; 1]^2.$$

For the FCC phase, the unknown part is modelled by

$$\sum_{p=1}^{p_{A1}} \theta_p^{(A1)} F_p^{(A1)}(\mathbf{E}, \boldsymbol{\gamma}^{(A1)}) = \gamma_{Cu}^{(A1)} \gamma_{Mg}^{(A1)} \theta_1^{(A1)}, \quad \boldsymbol{\gamma}^{(A1)} \in [0; 1]^2.$$

For the HCP phase, the unknown part is modelled by

$$\sum_{p=1}^{p_{A3}} \theta_p^{(A3)} F_p^{(A3)}(\mathbf{E}, \boldsymbol{\gamma}^{(A3)}) = \gamma_{Cu}^{(A3)} \gamma_{Mg}^{(A3)} \theta_1^{(A3)}, \quad \boldsymbol{\gamma}^{(A3)} \in [0; 1]^2.$$

The Gibbs energy function for the stoichiometric compound $\alpha = CuMg_2$ only exists at a fixed composition with one third of Cu and two thirds of Mg, that is when site fraction of the first sublattice is equal to 1 for Cu and 0 for Mg and inversely for the second sublattice. The bias term is null. The model for the Gibbs energy function associated is

$$G^{(CuMg_2)}(\mathbf{E}) = \begin{cases} \left[\theta_1^{(CuMg_2)} + \theta_2^{(CuMg_2)} T + \theta_3^{(CuMg_2)} T \log(T) \right. \\ \quad \left. + \theta_4^{(CuMg_2)} T^2 + \theta_5^{(CuMg_2)} \frac{1}{T} \right] & \text{if } \gamma_{Cu}^{(CuMg_2,1)} = \gamma_{Mg}^{(CuMg_2,2)} = 1, \\ +\infty & \text{else.} \end{cases}$$

The Laves phase is a solid compound $(Cu, Mg)_2(Cu, Mg)_1$ with a composition domain. The compound-energy formalism (CEF) is used to describe the Gibbs energy function. So the known bias term is given by

$$\begin{aligned} B^{(C15)}(\mathbf{E}, \boldsymbol{\gamma}^{(C15)}) = & \gamma_{Cu}^{(C15,1)} \gamma_{Cu}^{(C15,2)} G_{Cu:Cu}(\mathbf{E}) + \gamma_{Mg}^{(C15,1)} \gamma_{Cu}^{(C15,2)} G_{Mg:Cu}(\mathbf{E}) \\ & + \gamma_{Mg}^{(C15,1)} \gamma_{Mg}^{(C15,2)} G_{Mg:Mg}(\mathbf{E}) \\ & + RT \left[2 \left(\gamma_{Cu}^{(C15,1)} \log \left(\gamma_{Cu}^{(C15,1)} \right) + \gamma_{Mg}^{(C15,1)} \log \left(\gamma_{Mg}^{(C15,1)} \right) \right) \right. \\ & \left. + \left(\gamma_{Cu}^{(C15,2)} \log \left(\gamma_{Cu}^{(C15,2)} \right) + \gamma_{Mg}^{(C15,2)} \log \left(\gamma_{Mg}^{(C15,2)} \right) \right) \right] \end{aligned}$$

for $\boldsymbol{\gamma}^{(C15)} \in [0; 1]^4$. The functions $G_{Cu:Cu}$, $G_{Mg:Cu}$ and $G_{Mg:Mg}$ are defined as

$$\begin{aligned} G_{Cu:Cu}(\mathbf{E}) &= 3G_{Cu}^{o,(FCC)}(\mathbf{E}) + 45000 \\ G_{Mg:Mg}(\mathbf{E}) &= 3G_{Mg}^{o,(FCC)}(\mathbf{E}) + 17700 \\ G_{Mg:Cu}(\mathbf{E}) &= 2G_{Mg}^{o,(HCP)}(\mathbf{E}) + G_{Cu}^{o,(FCC)}(\mathbf{E}) + 34720. \end{aligned}$$

The unknown part of the function is modelled by

$$\begin{aligned} \sum_{p=1}^{p_{C15}} \theta_p^{(C15)} F_p^{(C15)}(\mathbf{E}, \boldsymbol{\gamma}^{(C15)}) &= \gamma_{Cu}^{(C15,1)} \gamma_{Mg}^{(C15,1)} \theta_{(Cu,Mg:)}^{(C15)} + \gamma_{Cu}^{(C15,2)} \gamma_{Mg}^{(C15,2)} \theta_{(:,Cu,Mg)}^{(C15)} \\ &+ \gamma_{Cu}^{(C15,1)} \gamma_{Mg}^{(C15,2)} \left[\theta_1^{(C15)} + \theta_2^{(C15)} T + \theta_3^{(C15)} T \log(T) \right. \\ &\left. + \theta_4^{(C15)} T^2 + \theta_5^{(C15)} \frac{1}{T} + \theta_6^{(C15)} T^3 \right], \quad \boldsymbol{\gamma}^{(C15)} \in [0; 1]^4. \end{aligned}$$

We denote by $\boldsymbol{\theta} = (\boldsymbol{\theta}^{(liq)})^T, \theta^{(A1)}, \boldsymbol{\theta}^{(C15)T}, \boldsymbol{\theta}^{(CuMg_2)T}, \theta^{(A3)})^T \in \mathbb{R}^{18}$ the vector of all unknown parameters to assess, with $\boldsymbol{\theta}^{(liq)} \in \mathbb{R}^3$, $\theta^{(A1)} \in \mathbb{R}$, $\boldsymbol{\theta}^{(C15)} \in \mathbb{R}^8$, $\boldsymbol{\theta}^{(CuMg_2)} \in \mathbb{R}^5$ and $\theta^{(A3)} \in \mathbb{R}$. Figure III.2 displays the ideal phase diagram obtained at the end of the assessment proposed by the SATA School for the Cu-Mg system.

III. 4.2 Results

This Section displays the results obtained by the application of the Bayesian approach for the linear model. The parameter *prior* mean $\boldsymbol{\theta}_0$ was chosen equal to the result of the assessment of the Cu-Mg system made for the SATA school. The parameter *prior* covariance matrix $\boldsymbol{\Sigma}_0$ was chosen equal to $(F(\mathbf{X})^T F(\mathbf{X}))^{-1}$.

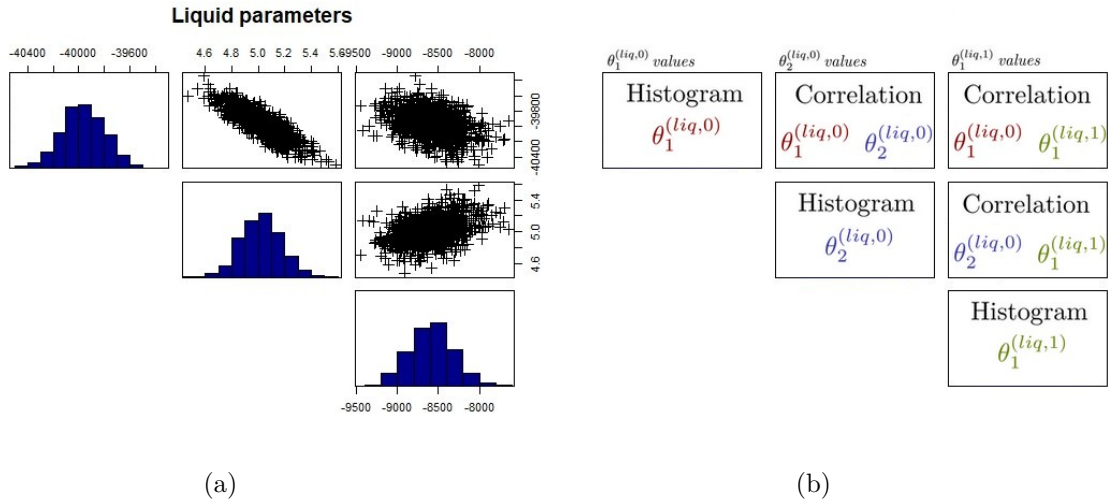


FIGURE III.3 – (a) Distributions and correlations of the interaction parameters associated to the Gibbs energy function for the liquid phase. (b) Definitions of each plots.

Figures III.3, III.4 and III.5 display plot matrices of $\boldsymbol{\theta}$ posterior distribution for the liquid phase, the Laves phase and the CuMg₂ compound. We briefly recall how to read such graphs with the example of the liquid phase parameters (Figure III.3). Notice that it is possible to construct such plot matrices for the whole parameter. However, we choose for clarity to group the

parameters by their related Gibbs energy functions. Diagonal plots are histograms of each liquid phase parameters, that are the order 0 interaction parameters $\theta_1^{(liq,0)}$ and $\theta_2^{(liq,0)}$ and the first order interaction parameter $\theta_1^{(liq,1)}$. The scatterplots above the diagonal display the correlations between the different parameters. The first row features correlation between parameters $\theta_1^{(liq,0)}$ and $\theta_2^{(liq,0)}$ for the middle scatterplot and between parameters $\theta_1^{(liq,0)}$ and $\theta_1^{(liq,1)}$ for the right scatterplot. The second row right scatterplot displays the correlation between parameters $\theta_2^{(liq,0)}$ and $\theta_1^{(liq,1)}$. The correlation scatterplots show elliptic points clouds which are specific of the Gaussian distribution. The correlation scatterplot for the order 0 interaction parameters is a rather elongated ellipse which means that parameters $\theta_1^{(liq,0)}$ and $\theta_2^{(liq,0)}$ values are strongly related. At the contrary, a circular scatterplot is the result of a low dependency.

For the particular value of $\gamma^{(C15)}$ where $\gamma_{Cu}^{(C15,1)} = \gamma_{Mg}^{(C15,2)} = 1$ and $\gamma_{Cu}^{(C15,2)} = \gamma_{Mg}^{(C15,1)} = 0$, the function $G^{(C15)}$ is expressed by

$$\theta_1^{(C15)} + \theta_2^{(C15)}T + \theta_3^{(C15)}T \log(T) + \theta_4^{(C15)}T^2 + \theta_5^{(C15)}\frac{1}{T} + \theta_6^{(C15)}T^3$$

all the others contributions being 0. The Figure III.4 displays the plot matrix obtained for the above parameters associated with the Laves phase.

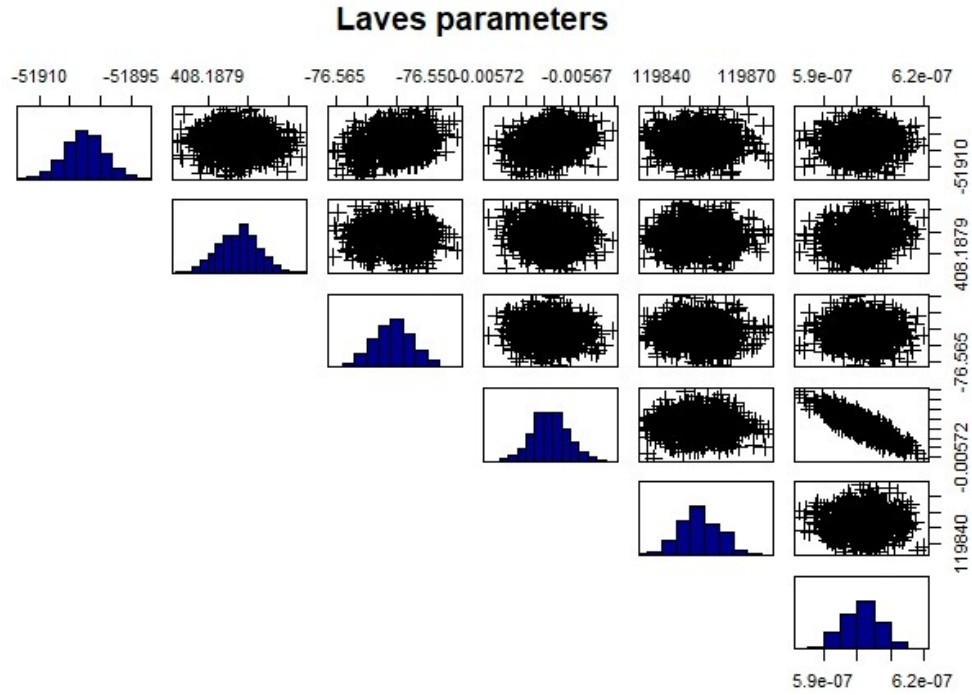


FIGURE III.4 – Distributions and correlations of the six parameters involved in the reference energy function of the Laves phase. The diagonal histograms are, from left to right, $\theta_1^{(C15)}$, $\theta_2^{(C15)}$, $\theta_3^{(C15)}$, $\theta_4^{(C15)}$, $\theta_5^{(C15)}$ and $\theta_6^{(C15)}$.

Finally the Figure III.5 displays the distribution and the correlations for the parameters

involved in the the CuMg₂ Gibbs energy function which is recalled below

$$G^{(CuMg_2)}(\mathbf{E}) = \begin{cases} \left[\theta_1^{(CuMg_2)} + \theta_2^{(CuMg_2)}T + \theta_3^{(CuMg_2)}T \log(T) \right. \\ \quad \left. + \theta_4^{(CuMg_2)}T^2 + \theta_5^{(CuMg_2)}\frac{1}{T} \right] & \text{if } \gamma_{Cu}^{(CuMg_2,1)} = \gamma_{Mg}^{(CuMg_2,2)} = 1, \\ +\infty & \text{else.} \end{cases}$$

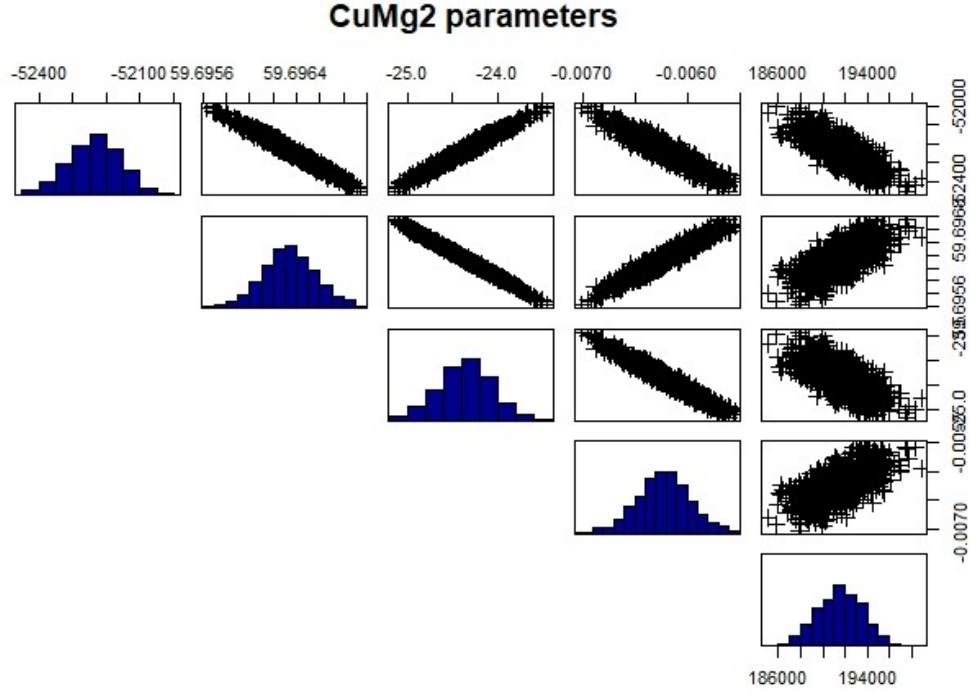


FIGURE III.5 – Distributions and correlations of the five parameters involved in the Gibbs energy function for the CuMg₂ compound. The diagonal histograms are, from left to right, $\theta_1^{(CuMg_2)}$, $\theta_2^{(CuMg_2)}$, $\theta_3^{(CuMg_2)}$, $\theta_4^{(CuMg_2)}$ and $\theta_5^{(CuMg_2)}$.

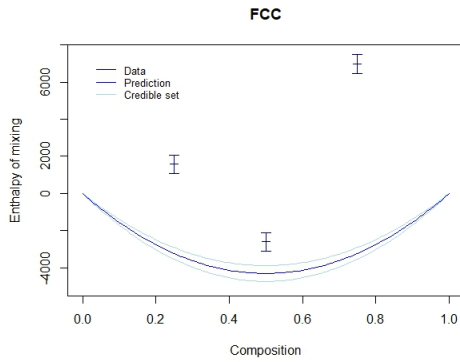
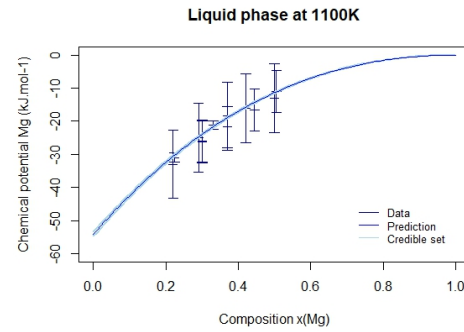
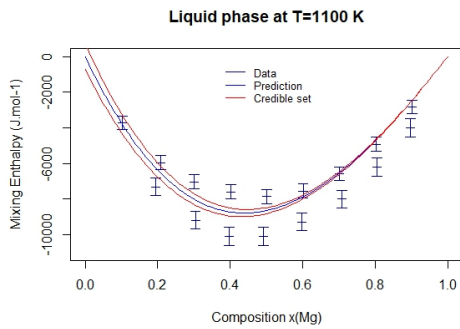
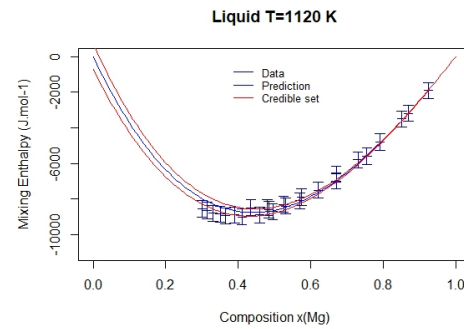
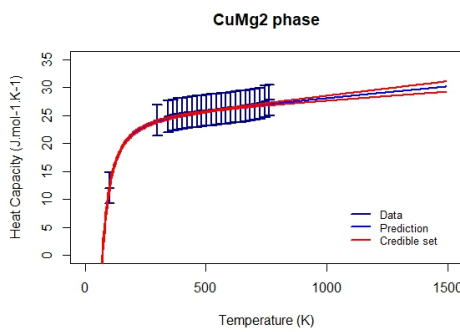
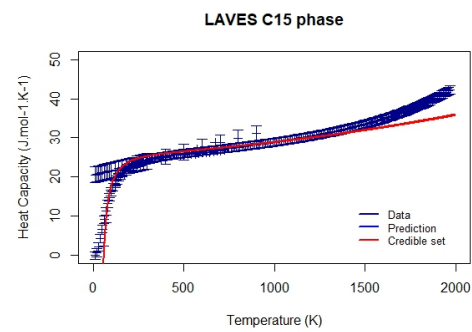
(a) *Enthalpy of formation function.*(b) *Chemical potential.*(c) *Enthalpy of mixing function.*(d) *Enthalpy of mixing function.*(e) *Heat capacity function.*(f) *Heat capacity function.*

FIGURE III.6 – Uncertainty propagation for different thermodynamic quantities for the FCC phase (III.6a), the liquid phase (III.6b, III.6c and III.6d), the CuMg_2 phase (III.6e) and the Laves phase (III.6f).

We then use the results of the Section III. 3.4 to propagate the uncertainty on the parameters to some thermodynamic quantities. The Figure III.6 displays the credible intervals obtained with our method and the agreement with the thermodynamic data. The agreement with the thermodynamic data is rather good except for the formation enthalpy function of the FCC phase (Figure III.6a). One can see that the credible intervals obtained are relatively thin around the prediction function even when the thermodynamic data are far from the model chosen. One can also notice that these credible intervals get thinner where a lot of assimilation data are available.

Finally, the Figure III.7 display the results obtained with the methodology proposed in the Section III. 3.5. The graph is obtained by sampling 100 values of θ from the posterior distribution. For each sample, the equilibrium is computed with the free software Open Calphad, see [89] for the Software description and algorithm explanation.

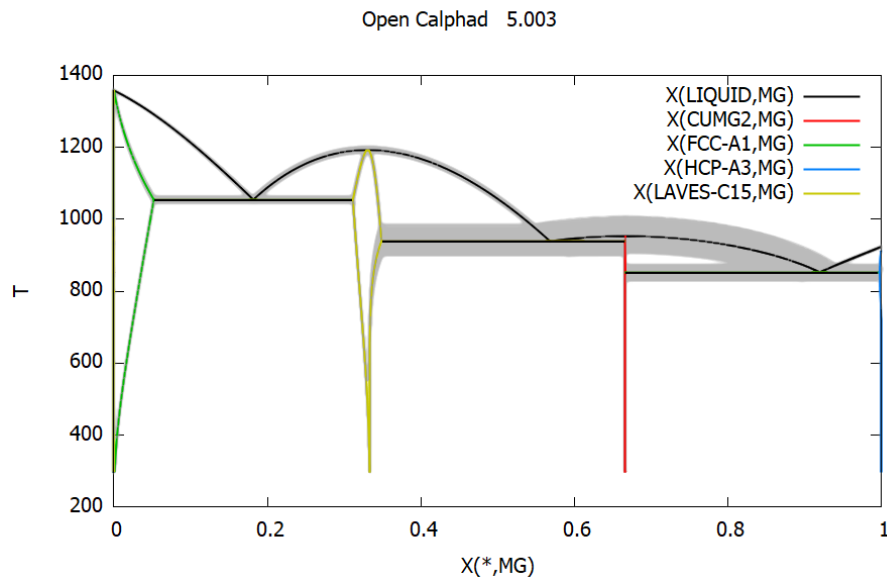


FIGURE III.7 – Monte-Carlo simulations of the phase diagram for 100 samples of the *posterior* distribution. Bold lines correspond to the phase diagram obtained for the parameters mean values. The equilibrium computations are made thanks to the free software Open Calphad, [89]. The Open Calphad source code can be compiled from the Github repository of Prof. Bo Sundman, see [1].

Chapitre IV

Non parametric reconstruction with linear constraints : the γ -entropy maximisation problem

IV. 1 Introduction

In some problems coming from applied physics, a multidimensional function f taking values in $\mathbb{R}^p$ ought to be reconstructed given a set of observations. In thermodynamics, information on the function of interest, namely the p components of the function f we wish to reconstruct, are indirectly available. In general the available information consists in the value of integrals that involves the unknown function f and some known weights $(\lambda^i)_{i=1,\dots,p}$. For example, one can consider an interpolation problem when the integration measure consists in Dirac masses. In this case we give at known locations the value of a scalar product between f and λ , see expression (IV.2). In the present work, we need to consider more general constraints. Therefore we study a reconstruction problem in which constraints are defined as integrals involving the unknown function f and the weight function λ against suitable measures Φ , see expressions (IV.1) and (IV.3).

In the sequel we provide a general method for the reconstruction of a p -real valued function from partial knowledge submitted to the general constraints previously discussed. We refer the interested reader to [58, Chap. 2] for the basic rules of thermodynamics and [58, Chap. 5] for the description of functions which are ordinarily considered for the reconstruction of thermodynamic quantities.

To be more precise, we consider a $\mathbb{R}^p$ valued function $f(x) = (f^1(x), \dots, f^p(x))$ defined for all x in the compact set $U \subset \mathbb{R}^d$ (we assume the interior of U to be non-empty). We set our work in a probability space $(U, \mathcal{B}(U), P_U)$ where $\mathcal{B}(U)$ is the Borel σ -algebra and P_U is the given reference measure. In such framework, we wish to reconstruct f over U and such that the

reconstructed quantity satisfies the N following integral constraints

$$\int_U \sum_{i=1}^p \lambda^i(x) f^i(x) d\Phi_l(x) = z_l \quad 1 \leq l \leq N, \quad (\text{IV.1})$$

where Φ_l are N positive (known) finite measures on $(U, \mathcal{B}(U))$ and λ^i are known continuous weight functions.

The expression of integral constraints as in (IV.1) allows to express a wide range of problems. For example, one can consider some pairs $(x_l, z_l) \in U \times \mathbb{R}$ for $l = 1, \dots, N$, and wish to solve the following interpolation equations

$$\sum_{i=1}^p \lambda^i(x_l) f^i(x_l) = z_l \quad 1 \leq l \leq N. \quad (\text{IV.2})$$

Expression (IV.2) can be obtained from (IV.1) by choosing $d\Phi_l(x) = \delta_{x_l}(dx)$ the Dirac measure located at x_l for all $l = 1, \dots, N$. Therefore, the integral constraints (IV.1) become interpolation constraints.

When U is a subset of $\mathbb{R}$, one can also involves the l first moments of f^i by taking $dP_U(x) = dx$ and $d\Phi_l(x) = x^l dP_U(x)$. Namely, integral constraints (IV.1) become in this case

$$\int_U \sum_{i=1}^p \lambda^i(x) f^i(x) x^l dP_U(x) = z_l \quad 1 \leq l \leq N.$$

In our work, the z_l represent N ideal real-valued measurements. In the case of noisy observations, a relaxed version of problem (IV.1) can be considered. The aim is then to reconstruct p real-valued functions on U such that

$$\int_U \sum_{i=1}^p \lambda^i(x) f^i(x) d\Phi_l(x) \in \mathcal{K}_l \quad 1 \leq l \leq N, \quad (\text{IV.3})$$

where $\mathcal{K}_l$ is an interval in $\mathbb{R}$. In the sequel $\mathcal{K}$ will denote the product of the N intervals $\mathcal{K}_l$.

In the general case, problem (IV.3) is ill-posed and has many solutions. In our work, we propose to choose among the solutions the function f that maximises I_γ the γ -entropy of the function f defined by

$$I_\gamma(f) = - \int_U \gamma(f^1(x), \dots, f^p(x)) dP_U(x) \quad (\text{IV.4})$$

where γ is a strictly convex function from $\mathbb{R}^p$ to $\mathbb{R}$. In this framework, the reconstruction problem we consider can be rephrased as

$$\begin{aligned} & \max I_\gamma(f) \\ & \text{s.t.} \quad \int_U \sum_{i=1}^p \lambda^i(x) f^i(x) d\Phi(x) \in \mathcal{K}. \end{aligned} \quad (\mathcal{F}_{\Phi, \gamma}^p)$$

The resolution of problem $(\mathcal{F}_{\Phi, \gamma}^p)$ is conducted in two steps. We consider a dual problem on (signed) measures as a first step. The second step consists in solving a discrete approximation of the dual problem. This approach is summarised in Figure IV.1 and Figure IV.2 (see on the next

page). The Figure IV.1 presents the approach in the case $p = 1$ and $\lambda(x) = 1, \forall x \in U$, which has already been treated in [21]. Figure IV.2 presents the case $p \neq 1$, which is the extension of [21] treated in this work. The resolution we propose involves the embedding into a more complicated framework. Let us sketch the description of this framework. Let V be a Polish space and P_V be the reference measure on V . Unless it is specified, V is a compact space. The resolution we propose involves a transfer between U and V . A more precise description of such transfer will be given in Section IV. 2 (unidimensional case) and Section IV. 3 (multidimensional case).

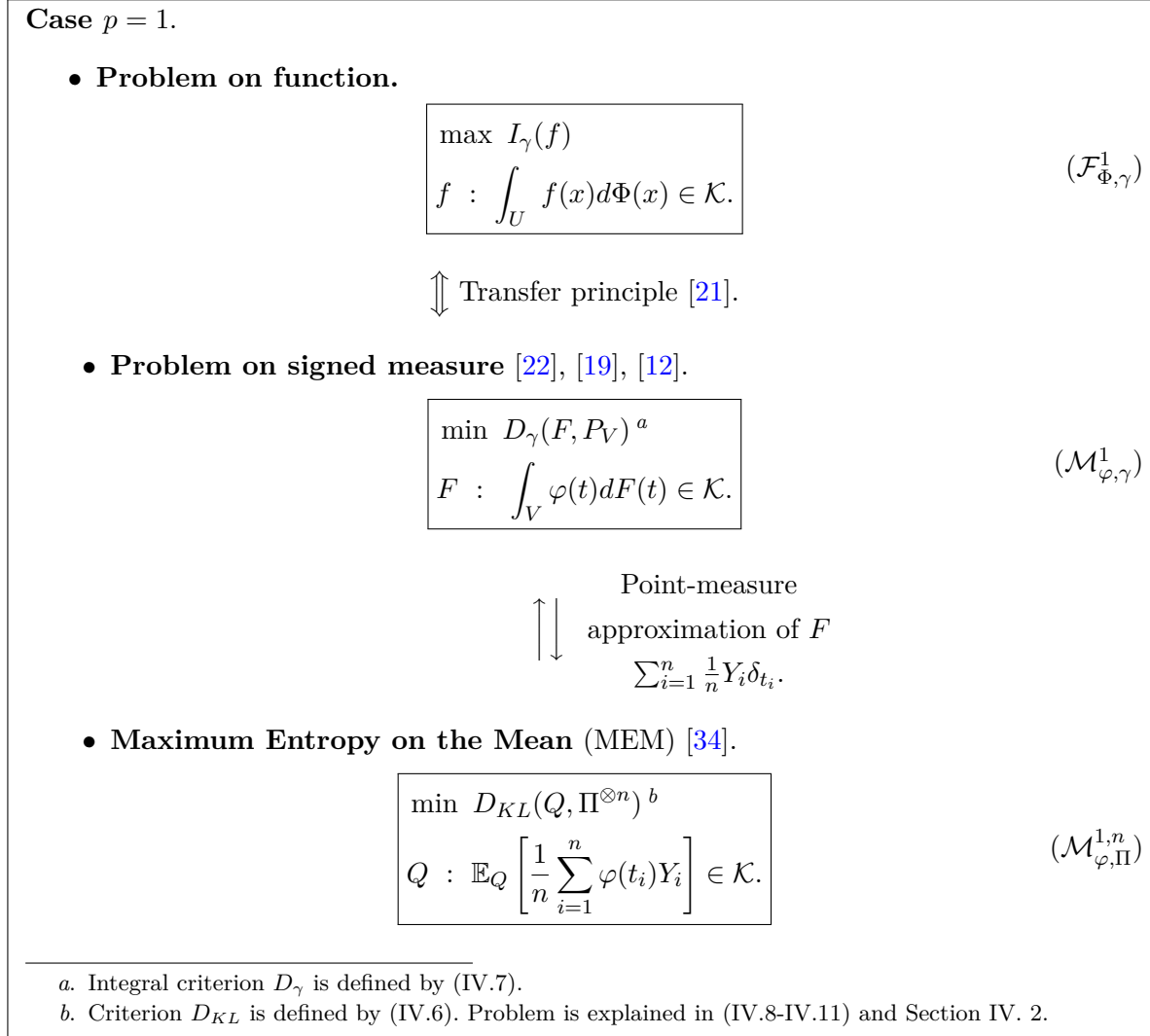


FIGURE IV.1 – Problems raised in the sequel, case $p = 1$. Such problems have already been studied, see Section IV. 2 for more details.

We first recall how the Maximum Entropy (ME) method is put in action. Originally, the ME method aims at the reconstruction of a probability measure P when dealing with information on the expectation under P of some random variables. We give below a first example.

Case $p \neq 1$.

- **Problem on functions.**

$$\begin{aligned} & \max I_\gamma(f) \\ & f : \int_U \sum_{i=1}^p \lambda^i(x) f^i(x) d\Phi(x) \in \mathcal{K}. \end{aligned} \quad (\mathcal{F}_{\Phi, \gamma}^p)$$

$\Updownarrow$ Transfer principle.

- **Problem on signed measures.**

$$\begin{aligned} & \min D_\gamma(F, P_V) \\ & F : \sum_{i=1}^p \int_V \varphi^i(t) dF^i(t) \in \mathcal{K}. \end{aligned} \quad (\mathcal{M}_{\varphi, \gamma}^p)$$

$\uparrow \downarrow$ Point-measure
approximations
of the F^i .

- **Maximum Entropy on the Mean.**

$$\begin{aligned} & \min D_{KL}(Q, \Pi^{\otimes n}) \\ & Q : \mathbb{E}_Q \left[\frac{1}{n} \sum_{i=1}^p \sum_{j=1}^n \varphi^i(t_j) Y_j^i \right] \in \mathcal{K}. \end{aligned} \quad (\mathcal{M}_{\varphi, \Pi}^{p, n})$$

FIGURE IV.2 – Problems raised in the sequel, $p \neq 1$. See Section IV. 3 for the extension of the method provided by [21] to solve inverse problems.

Example 2. When $V = \mathbb{R}$, one may want to reconstruct a probability measure P such that the quantity $\int t^k dP(t)$ for some $k \in \mathbb{N}^*$ is equal to given values m_k .

More precisely, define the entropy of a probability measure P with respect to the measure P_V as

$$S(P) = \begin{cases} - \int_V \log \left(\frac{dP}{dP_V} \right) dP & \text{if } P \ll P_V \text{ and } \log \left(\frac{dP}{dP_V} \right) \in L^1(P) \\ -\infty & \text{otherwise,} \end{cases} \quad (\text{IV.5})$$

where $P \ll P_V$ means that P is absolutely continuous with respect to P_V . ME method derives as solution the probability P^{ME} which maximises the entropy, provided that the information for the reconstructed probability measure meets the information asked.

In information theory and statistics, one usually considers the opposite of the entropy, that

is the so-called Kullback-Leibler divergence of P with respect to P_V which is defined by

$$D_{KL}(P, P_V) = \begin{cases} \int_V \log \left(\frac{dP}{dP_V} \right) dP & \text{if } P \ll P_V \\ & \text{and } \log \left(\frac{dP}{dP_V} \right) \in L^1(P) \\ +\infty & \text{otherwise.} \end{cases} \quad (\text{IV.6})$$

Equivalently the ME method derives as a solution the probability measure P^{ME} which minimises the Kullback-Leibler divergence from the reference measure P_V under the constraints. Reference measure P_V can be interpreted as a *prior* measure.

The Kullback-Leibler divergence defined in (IV.6) is called the I -divergence in [22] and [19]. The author also calls I -projection the probability measure that maximises the entropy (IV.5) on a convex set of probability measures. Further in [23] an axiomatic justification for the use of the ME method is provided.

In a more general case, the entropy problem can target a reconstruction of a signed measure. In this case, the γ -divergence D_γ , defined in the expression (IV.7), is considered instead of the Kullback-Leibler divergence D_{KL} (IV.6). The authors in [12] and [13] have studied the minimisation of the γ -divergence D_γ under linear constraints. Let F be a signed measure defined on V . The classical Lebesgue decomposition of F with respect to P_V is

$$F = F^a + F^s$$

with $F^a \ll P_V$ the absolutely continuous part and F^s the singular part. F^s is singular with respect to P_V means that it is concentrated on a set $\tilde{V}$ such that $P_V(\tilde{V}) = 0$. We recall as well the Jordan decomposition of measure F^s

$$F^s = F^{s,+} - F^{s,-}$$

with $F^{s,+}$ and $F^{s,-}$ two positive measures mutually singular. The γ -divergence D_γ is then defined as follows

$$D_\gamma(F, P_V) = \int_V \gamma \left(\frac{dF^a}{dP_V} \right) dP_V + b_\psi F^{s,+}(V) - a_\psi F^{s,-}(V) \quad (\text{IV.7})$$

where the integrand γ is a convex function and F^a , $F^{s,+}$ and $F^{s,-}$ are as defined previously. The scalar quantities b_ψ and a_ψ , with $a_\psi < b_\psi$, are the endpoints of ψ domain with ψ the convex conjugate of γ defined by

$$\psi(t) = \sup_y \{ty - \gamma(y)\}.$$

Taking back the Example 2, the reconstruction problem can be put in the frame of an optimisation problem as $(\mathcal{M}_{\varphi, \gamma}^1)$ defined in Figure IV.1. In this example, the objective function is the Kullback-Leibler divergence D_{KL} , that is the criterion D_γ when γ is the convex function defined on $\mathbb{R}_+^*$ by $y \rightarrow y \log(y) - y - 1$. The scalar quantities a_ψ and b_ψ are respectively equal to $-\infty$ and $+\infty$ which leads to a reconstruction with no singular parts. The moment constraint can be written as $\int_V \varphi(t) dF(t) \in \mathcal{K}$ by taking $\mathcal{K} = \{m_k\}$ and $\varphi : t \rightarrow t^k$. Finally, one has to add the constraint $\int_V dF(t) = 1$ to ensure that the reconstructed measure is a probability measure.

Notice that the expression in (IV.7) contains terms depending on the singular part F^s of measure F . Those terms may not be considered in the γ -divergence (IV.7) depending on the convex function γ used, see [15] and Example 2.

More generally the author in [54] and [55] studies the characterization of the optimal signed measure which minimises (IV.7) under linear constraints. Integral functionals with normal convex integrand, i.e. integrals for which the integrated function is strictly convex with respect to one of its variable, are studied. See more comments on normal convex integrand in [74, Chapter 14]. See also [56] for a systematic treatment of convex distance minimisation.

We now recall that many usual optimisation problems $(\mathcal{M}_{\varphi, \gamma}^1)$ can be set in an entropy maximisation problem frame, as proposed in the early paper [62]. This general embedding in ME is called the Maximum Entropy on the Mean (MEM) method and has been developed in [34] and [24]. The method is based on a suitable discretisation of the working set V . The reference measure P_V is approximated by a point-measure supported by n pixels, $t_1, \dots, t_n$ which are deterministic points in V such that $P_n = \frac{1}{n} \sum_{i=1}^n \delta_{t_i} \rightarrow P_V$. By associated to each pixel t_i a real random amplitude Y_i , we defined the random point-measure F_n by

$$F_n = \frac{1}{n} \sum_{i=1}^n Y_i \delta_{t_i}. \quad (\text{IV.8})$$

By construction $F_n \ll P_n$. Notice that we choose to present the simple case of real random amplitudes Y_i for this introduction, but one can consider more complicated constructions. We will do so in the Section IV. 3.3 where the amplitudes Y_i will be vectors in $\mathbb{R}^p$.

In the MEM problem, one wants to determine the "optimal" distribution Q to generate vectors of n real random amplitudes $(Y_1, \dots, Y_n)$. Optimal distribution Q must be such that the constraints considered in problem $(\mathcal{M}_{\varphi, \gamma}^1)$ applied to the random point-measure F_n is met on average, that is that

$$\mathbb{E}_Q \left[\int_V \varphi(t) dF_n(t) \right] = \mathbb{E}_Q \left[\frac{1}{n} \sum_{i=1}^n \varphi(t_i) Y_i \right] \in \mathcal{K}. \quad (\text{IV.9})$$

Taking back Example 2, the constraints applied to the point-measure approximation become

$$\begin{aligned} \mathbb{E}_Q \left[\frac{1}{n} \sum_{i=1}^n Y_i \right] &= 1 \quad \text{and} \\ \mathbb{E}_Q \left[\frac{1}{n} \sum_{i=1}^n (t_i)^k Y_i \right] &= m_k. \end{aligned} \quad (\text{IV.10})$$

We assume that the random amplitudes are independent. Let Π denote their *prior* distribution so that the reference distribution for $(Y_1, \dots, Y_n)$ is the tensor measure $\Pi^{\otimes n}$. In order to build the optimal distribution, we minimise the Kullback-Leibler divergence (IV.6) with respect to the prior $\Pi^{\otimes n}$ under the constraints (IV.9). When such "optimal" distribution exists, we will denote the solution by Q_n^{MEM} . Then let F_n^{MEM} be defined as follows

$$F_n^{MEM} = \mathbb{E}_{Q_n^{MEM}} \left[\frac{1}{n} \sum_{i=1}^n Y_i \delta_{t_i} \right]. \quad (\text{IV.11})$$

Notice that unlike F_n , the quantity F_n^{MEM} is no longer random. Let the log Laplace transform of probability measure Π be denoted by ψ_Π

$$\psi_\Pi(\tau) = \log \left(\int_{\mathbb{R}} \exp(\tau y) d\Pi(y) \right) \quad \text{for all } \tau \in \mathcal{D}_{\psi_\Pi}, \quad (\text{IV.12})$$

with $\mathcal{D}_{\psi_\Pi}$ the domain of ψ_Π . We denote by γ_Π the convex conjugate of the log-Laplace transform ψ_Π ¹. We hope that the reconstruction F_n^{MEM} is a good approximation (in a sense we will further specify) of the solution of the corresponding continuous problem $(\mathcal{M}_{\varphi,\gamma}^1)$, for which convex function γ is the function γ_Π . The properties of the minimising sequence $(F_n^{MEM})_n$ have been studied in [34]. The authors in [35] deal with a multidimensional case, that is estimating a vector of reconstructions when dealing with information on generalised moments of each components. In [36], Bayesian and MEM methods are proposed to solve inverse problems on measures. More details about the MEM method will be provided in Section IV. 2 for the case $p = 1$ and Section IV. 3.3 for the case $p \neq 1$. In particular we explain how to choose reference probability measure Π so that the criterion $D_{KL}(\cdot, \Pi^{\otimes n})$ for discrete problem $(\mathcal{M}_{\varphi,\Pi}^{1,n})$ is a good alternative to criterion $D_\gamma(\cdot, P_V)$ of continuous problem $(\mathcal{M}_{\varphi,\gamma}^1)$.

Back to the function reconstruction problem, an extension of the MEM method to solve generalised moment problems for function reconstruction as in $(\mathcal{F}_{\Phi,\gamma}^1)$ is proposed by [21] in the case $p = 1$ and $\lambda(x) = 1, \forall x \in U$. The method uses a transfer principle which links the function to reconstruct to a corresponding measure. The transfer relies on the use of suitable kernels. Such transfer is particularly useful when considering measures Φ in the constraints equations (IV.3) that might not all be absolutely continuous with respect to the reference measure P_U .

Our work has a twofold purpose. We first make a summary of entropy maximisation problems in order to reconstruct a single measure and, by extension with the linear transfer, entropy maximisation problems in order to reconstruct a single function. We propose a global synthesis of the entropy maximisation methods for such reconstruction problems from the convex analysis point of view, as well as in the framework of the embedding into the MEM setting. We then propose an extension of the entropy methods for a multidimensional case. Such extension is the main contribution of this work. We study the MEM embedding for the function reconstruction problem that is proposed in [21] in the extended case of inverse problems, that is when $p \neq 1$ and the λ^i are any known bounded continuous functions. We provide a general method of reconstruction based on the γ -entropy maximisation for functions submitted to generalised moment and interpolation constraints as in (IV.3).

This paper is organized as follows.

In Section IV. 2, we recall in a global synthesis of entropy maximisation methods some results for the specific case of a single function reconstruction problem $(\mathcal{F}_{\Phi,\gamma}^1)$ or a single measure reconstruction problem $(\mathcal{M}_{\varphi,\gamma}^1)$. First we describe how the transfer principle works, that is how a function problem $(\mathcal{F}_{\Phi,\gamma}^1)$ can be linked to the measure problem $(\mathcal{M}_{\varphi,\gamma}^1)$ in Section IV. 2.1. Then

1. Notice that as ψ_Π is a convex function, ψ_Π is also the convex conjugate of γ_Π .

we recall some results about the resolution of the γ -divergence minimisation problem $(\mathcal{M}_{\varphi,\gamma}^1)$ in Section IV. 2.2. In Section IV. 2.3, we take a specific look at problem $(\mathcal{M}_{\varphi,\gamma}^1)$ when the convex function γ is the function $y \rightarrow y \log(y) - y - 1$. This specific problem is the ME problem which will be denoted by $(\mathcal{M}_{\varphi,ME}^1)$. We then extend the class of studied optimisation problems by giving the construction of the MEM problem setting and we provide some properties of the MEM reconstruction in Section IV. 2.4. Finally in Section IV. 2.5, the setting of some usual optimisation problems into the entropy maximisation frame is reminded.

The main contributions of this paper are the results presented in Section IV. 3. It consists in the study of the entropy methods for the reconstruction of a multidimensional function submitted to very general constraints such as an integral inverse problem as in (IV.3). We extend the approach of [21] for the case $p \neq 1$ and any known bounded continuous functions λ^i . We study the embedding of the functions reconstruction problem $(\mathcal{F}_{\Phi,\gamma}^p)$ into the MEM problem $(\mathcal{M}_{\varphi,\Pi}^{p,n})$ framework. Such study is stepped in three independent parts. First in Section IV. 3.1, we study the problem $(\mathcal{F}_{\Phi,\gamma}^p)$ in a convex analysis framework. We express the optimal solution thanks to Fenchel duality theorem. This first approach lacks to give a suitable reconstruction when constraint measures Φ are not absolutely continuous with respect to P_U . To remedy this issue, we propose in Section IV. 3.2 to transfer the functions reconstruction problem $(\mathcal{F}_{\Phi,\gamma}^p)$ to a corresponding measures reconstruction problem $(\mathcal{M}_{\varphi,\gamma}^p)$. The transfer is performed using suitable continuous kernels. Finally in Section IV. 3.3, we set problem $(\mathcal{M}_{\varphi,\gamma}^p)$ obtained by the transfer into a MEM problem framework and we study the reconstructions given by problem $(\mathcal{M}_{\varphi,\Pi}^{p,n})$.

Applications will be presented in Section IV. 4. We consider first some simple examples of single function reconstructions and then a two-functions case study inspired by computational thermodynamics.

IV. 2 The specific γ -entropy maximisation problem for a single reconstruction

We give in this section some details about the γ -entropy maximisation problem in the case of a single reconstruction, that is that we are interested in a single function or a single measure reconstruction. Those problems have already been studied, see for example in [21] for problem $(\mathcal{F}_{\Phi,\gamma}^1)$, in [13] for problem $(\mathcal{M}_{\varphi,\gamma}^1)$ and in [34] for problem $(\mathcal{M}_{\varphi,\Pi}^{1,n})$. Let us recall some results of these authors.

In Section IV. 2.1 we are interested in the link between a function reconstruction problem $(\mathcal{F}_{\Phi,\gamma}^1)$ and a measure reconstruction problem $(\mathcal{M}_{\varphi,\gamma}^1)$. The function reconstruction problem is set as

$$\begin{aligned} \max \quad & I_\gamma(f) \\ f : \quad & \int_U f(x) d\Phi(x) \in \mathcal{K} \end{aligned} \quad (\mathcal{F}_{\Phi,\gamma}^1)$$

and the measure reconstruction problem as

$$\begin{aligned} \min \quad & D_\gamma(F, P_V) \\ F : \quad & \int_V \varphi(t) dF(t) \in \mathcal{K}. \end{aligned} \tag{\mathcal{M}_{\varphi, \gamma}^1}$$

The idea is to set a transformation from measures on V to functions on U . Such transformation is the linear transfer we will further describe in IV. 2.1.

We remind the reader that we consider a number N of constraints. Therefore Φ and φ take values in $\mathbb{R}^N$. In addition we will consider that functions φ are continuous.

In Section IV. 2.2 we study the γ -divergence minimisation problem under constraints, that is problem $(\mathcal{M}_{\varphi, \gamma}^1)$. We recall the results of [13] for the existence of an optimal solution to problem $(\mathcal{M}_{\varphi, \gamma}^1)$.

We take a closer look in Section IV. 2.3 at the Maximum Entropy (ME) problem. This problem corresponds to problem $(\mathcal{M}_{\varphi, \gamma}^1)$ in the special setting when γ is the function $y \rightarrow y \log(y) - y - 1$ and the γ -divergence coincides with the Kullback-Leibler divergence. We recall results on the existence of the optimum and its expression when such minimiser exists.

In Section IV. 2.4, we recall the setting of the n -th MEM problem $(\mathcal{M}_{\varphi, \Pi}^{1, n})$

$$\begin{aligned} \min \quad & D_{KL}(Q, \Pi^{\otimes n}) \\ Q : \quad & \mathbb{E}_Q \left[\frac{1}{n} \sum_{i=1}^n \varphi(t_i) Y_i \right] \in \mathcal{K}. \end{aligned} \tag{\mathcal{M}_{\varphi, \Pi}^{1, n}}$$

We give, when it exists, an expression for the minimiser Q_n^{ME} of the n -th problem. A function g_n^{ME} , related to the expectation of the random amplitudes under Q_n^{ME} , can be defined. We will see that under some assumptions, it exists one particular $v \in \mathbb{R}^N$ such that the sequence $(g_n^{ME})_n$ converges to a function $g_\infty^{ME}(t) = \psi'_\Pi(\langle v, \varphi(t) \rangle)$. This limit will be expressed.

Finally in Section IV. 2.5, we give examples of some classical optimisation problems embedding into the MEM framework.

IV. 2.1 Transfer principle

We briefly recall the idea of the transfer principle developed in [21], in order to put to work the γ -entropy methods in the case of a function reconstruction problem. Let us denote by $(U, \mathcal{B}(U), P_U)$ the probability space where $U \subset \mathbb{R}^d$ is compact (non empty), $\mathcal{B}(U)$ is the associated Borel set and P_U is the reference measure. In the case of the reconstruction problem for a single function, one wants to reconstruct over U a function f taking values in $\mathbb{R}$ such that f satisfies the integral constraints

$$\int_U f(x) d\Phi(x) \in \mathcal{K}. \tag{IV.13}$$

Let γ be a given convex function taking values in $\mathbb{R}$ and $\mathcal{D}_\gamma \subset \mathbb{R}$ its domain. Then the γ -entropy is defined by

$$I_\gamma(f) = - \int_U \gamma(f(x)) dP_U(x),$$

for a function f defined on U and taking values in $\mathcal{D}_\gamma$.

In order to chose among the functions that satisfy (IV.13), we propose as a selection criterion to maximise the γ -entropy. This means that we consider the optimisation problem $(\mathcal{F}_{\Phi, \gamma}^1)$

$$\begin{aligned} \max \quad & I_\gamma(f) \\ f : \quad & \int_U f(x) d\Phi(x) \in \mathcal{K}. \end{aligned} \tag{IV.13}$$

The method proposed by [21] is to transfer the function reconstruction problem to a measure reconstruction problem thanks to some continuous kernel K . Such kernel links the function to reconstruct to a corresponding signed measure. Recall that V is a Polish space and P_V the reference probability measure on V . The idea is that if one can reconstruct over V a signed measure F such that

$$\int_V \varphi(t) dF(t) \in \mathcal{K}, \tag{IV.14}$$

then a regularized function f_K can be reconstructed linked to the measure F . To do so, we proceed as follows.

We denote by K a continuous kernel defined over $U \times V$ taking values in $\mathbb{R}$. The kernel K is such that measure Φ of the integral constraint (IV.13) is linked to a regularized function φ_K involved in an integral constraint as in (IV.14). The relation linking Φ to φ_K is given by

$$\varphi_K(t) = \int_U K(x, t) d\Phi(x), \quad t \in V.$$

Therefore, for any continuous kernel K , one can reconstruct the regularized function f_K associated to F by defining

$$f_K(x) = \int_V K(x, t) dF(t), \quad x \in U.$$

Hence as a consequence of Fubini theorem, if the measure F satisfies (IV.14), the regularized function f_K defined above satisfies (IV.13).

IV. 2.2 γ -divergence minimisation problem

In this section, we recall the results provided by [13] for the γ -divergence minimisation problem for signed measure reconstruction. We set our work on V . Let F be a signed measure defined on V . The classical Lebesgue decomposition of F with respect to P_V is

$$F = F^a + F^s$$

with $F^a \ll P_V$ the absolutely continuous part and F^s the singular part. F^s singular with respect to P_V means that it is concentrated on a set $\tilde{V}$ such that $P_V(\tilde{V}) = 0$. Notice that F^a and F^s are still signed measures. Recall the Jordan decomposition of measure F^s

$$F^s = F^{s,+} - F^{s,-}$$

with $F^{s,+}$ and $F^{s,-}$ mutually singular. Let γ be essentially strictly convex. We denote by ψ its convex conjugate, that is

$$\psi(t) = \sup_{y \in \mathcal{D}_\gamma} \{ty - \gamma(y)\}.$$

Considering a_ψ and b_ψ , with $a_\psi < b_\psi$, the endpoints of ψ domain, we define the γ -divergence D_γ by

$$D_\gamma(F, P_V) = \int_V \gamma \left(\frac{dF^a}{dP_V} \right) dP_V + b_\psi F^{s,+}(V) - a_\psi F^{s,-}(V).$$

The problem we consider is the following

$$\begin{aligned} \min \quad & D_\gamma(F, P_V) \\ F : \quad & \int_V \varphi(t) dF(t) \in \mathcal{K}. \end{aligned} \tag{\mathcal{M}_{\varphi,\gamma}^1}$$

The authors in [13] study the existence conditions of an optimal solution using convex analysis tools. Their first result is to consider the following dual problem of $(\mathcal{M}_{\varphi,\gamma}^1)$ which relies on a Lagrange multiplier $v \in \mathbb{R}^N$

$$\sup_{v \in \mathbb{R}^N} \left\{ \inf_{c \in \mathcal{K}} \langle v, c \rangle - \int_V \psi(\langle v, \varphi(t) \rangle) dP_V(t) \right\}. \tag{\mathcal{M}_{\varphi,\gamma}^{1,*}}$$

In addition, they give conditions for problems $(\mathcal{M}_{\varphi,\gamma}^1)$ and $(\mathcal{M}_{\varphi,\gamma}^{1,*})$ to have solutions. These results are recalled in Theorem 5 below.

Theorem 5. [13, Theorem 3.4]

1. If there exists $v' \in \mathbb{R}^N$ such that $\langle v', \varphi(t) \rangle \in \mathcal{D}_\psi$ then

$$F : \int_V \varphi(t) dF(t) \in \mathcal{K} \quad D_\gamma(F, P_V) = \sup_{v \in \mathbb{R}^N} \left\{ \inf_{c \in \mathcal{K}} \langle v, c \rangle - \int_V \psi(\langle v, \varphi(t) \rangle) dP_V(t) \right\}.$$

2. If, in addition, $\int_V \psi(\langle v', \varphi(t) \rangle) dP_V(t)$ is finite and if it exists a signed measure F' such that $\int_V \varphi(t) dF'(t) \in \mathcal{K}$, $\frac{dF'}{dP_V}$ is in the relative interior of $\mathcal{D}_\gamma$ and $D_\gamma(F', P_V)$ is finite, then

$$F : \int_V \varphi(t) dF(t) \in \mathcal{K} \quad D_\gamma(F, P_V) = \min_{F : \int_V \varphi(t) dF(t) \in \mathcal{K}} D_\gamma(F, P_V)$$

and

$$\begin{aligned} & \sup_{v \in \mathbb{R}^N} \left\{ \inf_{c \in \mathcal{K}} \langle v, c \rangle - \int_V \psi(\langle v, \varphi(t) \rangle) dP_V(t) \right\} \\ &= \max_{v \in \mathbb{R}^N} \left\{ \inf_{c \in \mathcal{K}} \langle v, c \rangle - \int_V \psi(\langle v, \varphi(t) \rangle) dP_V(t) \right\}. \end{aligned}$$

The next theorem proposes a more precise characterisation of the solution for problem $(\mathcal{M}_{\varphi,\gamma}^1)$ under the same conditions specified in Theorem 5.

Theorem 6. [13, Theorem 4.1] Under the assumptions of Theorem 5.

1. The absolutely continuous part with respect to P_V of solution F^o of $(\mathcal{M}_{\varphi,\gamma}^1)$ is given by

$$\frac{dF^o}{dP_V}(t) = \psi'(\langle v^*, \varphi(t) \rangle)$$

where v^* is the solution of

$$\max_{v \in \mathbb{R}^N} \left\{ \inf_{c \in \mathcal{K}} \langle v, c \rangle - \int_V \psi(\langle v, \varphi(t) \rangle) dP_V(t) \right\}.$$

2. If, in addition, for all $v \in \mathbb{R}^N$ and for all $t \in V$, $\langle v, \varphi(t) \rangle$ is in the interior of $\mathcal{D}_\psi$, the singular part vanishes.

It can be noted that when $\mathcal{D}_\psi = \mathbb{R}$, Theorem 6 always gives solutions that are absolutely continuous with respect to P_V . The condition to have $\mathcal{D}_\psi = \mathbb{R}$, is to consider a function γ that is such that the ratio $|\frac{\gamma(y)}{y}|$ is equal to ∞ on the edges of $\mathcal{D}_\gamma$, see [13, Lemma 2.1].

We will see in Section IV. 2.4 that the approach proposed by the embedding in the MEM framework boils down to the same results provided by Theorem 6.

IV. 2.3 Maximum entropy problem

In this section we take a better look at the problem of maximising the entropy of a probability measure under generalised moments constraints, that is the ME problem

$$\begin{aligned} \min \quad & D_{KL}(P, P_V) \\ P : \quad & \int_V \varphi(t) dP(t) \in \mathcal{K} \end{aligned} \tag{\mathcal{M}_{\varphi,ME}^1}$$

with φ defined on V and taking values in $\mathbb{R}^N$. We remind the reader that in this section, the first component of function φ is the constant 1 and that $\mathcal{K}$ is the product $\{1\} \times \mathcal{K}_1 \times \dots \times \mathcal{K}_{N-1}$. Notice that for the results recalled in this section only, V does not need to be compact and φ has not to be continuous. The definition of the Kullback-Leibler divergence D_{KL} is recalled below

$$D_{KL}(P, P_V) = \begin{cases} \int_V \log \left(\frac{dP}{dP_V} \right) dP & \text{if } P \ll P_V \text{ and } \log \left(\frac{dP}{dP_V} \right) \in L^1(P) \\ +\infty & \text{otherwise.} \end{cases}$$

The problem proposed in Section IV. 2.2 is a more general setting of the original ME problem. The reconstruction provided by the ME method satisfies the following properties. Those are showed by Shore and Johnson in [84].

Proposition 7.

1. **Uniqueness** : If the solution of ME problem exists, it is unique.
2. **Coordinate independence** : The reconstruction is independent of coordinate system choice.
3. **System independence** : If the probability space $(V, \mathcal{B}(V), P_V)$ consists in the product space of m probability spaces, the reconstruction over the whole probability space is the tensor product of the reconstructions on each probability space.
 In other words, if $P_V = \otimes_{i=1}^m P_{V_i}$ where P_{V_i} is a reference probability measure on $(V_i, \mathcal{B}(V_i))$ and $(V, \mathcal{B}(V), P_V)$ is the product space of all $(V_i, \mathcal{B}(V_i), P_{V_i})$, then the reconstruction P on $(V, \mathcal{B}(V))$ is given by $\otimes_{i=1}^m P_i$ where P_i is the reconstruction on $(V_i, \mathcal{B}(V_i))$.
4. **Subset independence** : If probability space $(V, \mathcal{B}(V), P_V)$ consists in the union of m probability spaces, the reconstruction over the global space leads to the same measure than the reconstruction problem conditioning on each probability space.
 In other words, it does not matter whether one treats the information as a subset V_j of whole set V in a conditional constraint or in the full system.

We recall in this section some results of [20] and [34] in order to solve problem $(\mathcal{M}_{\varphi, ME}^1)$. Results are stated without any proof for the setting of our example.

We give first a definition of the generalised solution for problem $(\mathcal{M}_{\varphi, ME}^1)$. This definition requires the use of a minimising sequence of $D_{KL}(\cdot, P_V)$, defined as follows. Let (P_n) be a sequence of probability measures. (P_n) is a minimising sequence of $D_{KL}(\cdot, P_V)$ if we have

$$\lim_{n \rightarrow \infty} D_{KL}(P_n, P_V) = \inf_P D_{KL}(P, P_V). \quad (\text{IV.15})$$

Definition 1. [20] Let us consider a sequence of probability measures $(P_n)_{n \in \mathbb{N}}$ defined on $(V, \mathcal{B}(V))$ such that (P_n) converges and is a minimising sequence of $D_{KL}(\cdot, P_V)$ and such that for all n , probability measure P_n satisfies

$$\int_V \varphi(t) dP_n(t) \in \mathcal{K}.$$

Then we call generalised maximal entropy solution the measure P^{MEG} that is such that

$$P^{MEG} = \lim_{n \rightarrow \infty} P_n, \quad \text{in total variation.}$$

If P^{MEG} also meets the constraint, then it is called the maximal entropy solution of problem $(\mathcal{M}_{\varphi, ME}^1)$ denoted by P^{ME} .

We will now recall some results on the existence of a generalised solution for problem $(\mathcal{M}_{\varphi, ME}^1)$ and the shape of the solution when it exists.

Let us first define $\mathcal{P}_{\mathcal{K}}$ the subset of probability measures that satisfy the constraints of the ME problem $(\mathcal{M}_{\varphi, ME}^1)$, that is

$$\mathcal{P}_{\mathcal{K}} = \left\{ P \text{ probability measure on } V, \text{ such that } \int_V \varphi(t) dP(t) \in \mathcal{K} \right\}.$$

With the previous definition of P^{MEG} , we recall a result of [20] for the existence in our framework of the generalised solution of problem $(\mathcal{M}_{\varphi, ME}^1)$.

Lemma 8. [20] *If $\inf_{P \in \mathcal{P}_{\mathcal{K}}} D_{KL}(P, P_V) < \infty$, then P^{MEG} exists.*

Let us now introduce several definitions involved in the characterisation of the problem $(\mathcal{M}_{\varphi, ME}^1)$ solution. For all $v \in \mathbb{R}^N$, we define the quantity

$$Z_{P_V, \varphi}(v) = \int_V \exp(\langle v, \varphi(t) \rangle) dP_V(t) \quad (\text{IV.16})$$

where $\langle \cdot, \cdot \rangle$ is the usual scalar product in $\mathbb{R}^N$. We denote the domain of $Z_{P_V, \varphi}$ by $D_{P_V, \varphi}$, that is the subset of vectors in $\mathbb{R}^N$ that are such that $Z_{P_V, \varphi}(\cdot)$ is finite

$$D_{P_V, \varphi} = \left\{ v \in \mathbb{R}^N, Z_{P_V, \varphi}(v) < \infty \right\}. \quad (\text{IV.17})$$

The following definition describes the so-called exponential family with respect to probability measure P_V . The interested reader may be referred to [8] for more details about exponential models.

Definition 2. *The φ -Hellinger arc of P_V is a family of measures P_v that are defined by*

$$dP_v(t) = (Z_{P_V, \varphi}(v))^{-1} \exp(\langle v, \varphi(t) \rangle) dP_V(t) \quad (\text{IV.18})$$

for all $v \in D_{P_V, \varphi}$.

The family of measures P_v defined as in (IV.18) for all $v \in D_{P_V, \varphi}$ may also be called the exponential model with respect to P_V .

We recall below the Theorem 4 from [19] that characterises the generalised reconstruction P^{MEG} . This theorem describes the reconstruction P^{MEG} as an element of the φ -Hellinger arc of P_V . More important, the reconstruction problem $(\mathcal{M}_{\varphi, ME}^1)$, which is infinite dimensional, is transformed into the finite dimension problem (IV.19) that considers the vectors v in $D_{P_V, \varphi}$, see expression (IV.17).

Theorem 9. [19, Theorem 4] *The reconstruction P^{MEG} belongs to the φ -Hellinger arc of P_V if and only if it exists one measure P in $\mathcal{P}_{\mathcal{K}}$ such that $P \ll P_V$.*

Then, defines

$$H_{P_V, \varphi}(v, \mathcal{K}) = \inf_{c \in \mathcal{K}} \langle v, c \rangle - \log(Z_{P_V, \varphi}(v)), \quad \forall v \in D_{P_V, \varphi}$$

the reconstruction P^{MEG} is obtained by determining $v^ \in D_{P_V, \varphi}$ such that*

$$H_{P_V, \varphi}(v^*, \mathcal{K}) = \sup_{v \in D_{P_V, \varphi}} H_{P_V, \varphi}(v, \mathcal{K}). \quad (\text{IV.19})$$

The previous theorem does not ensure that the reconstruction P^{MEG} will satisfy the constraints of problem $(\mathcal{M}_{\varphi, ME}^1)$. Of course, the reconstruction P^{MEG} is more interesting when P^{MEG} belongs to $\mathcal{P}_K$. We then have the following corollary for a reconstruction that satisfies the constraints. Such reconstruction is then denoted by P^{ME} .

Corollary 10. [34] *If there exists a measure $P \in \mathcal{P}_K$ such that $P \ll P_V$ and if $D_{P_V, \varphi}$ is an open set, then P^{MEG} is the reconstruction P^{ME} in $\mathcal{P}_K$ and P^{ME} belongs to the φ -Hellinger arc of P_V .*

As an illustration, we propose the following simple example of a probability measure reconstruction which maximises the entropy with respect to the standard Gaussian distribution $\mathcal{N}(0, 1)$. The added constraints are fixed valued for the first and second order moments. By giving first order moment equal to m and second order moment equal to $m^2 + 1$, the reconstruction we obtain is, as one can expect, the Gaussian distribution centred in m and with unit variance.

Example 3. *The working probability space is $(\mathbb{R}, \mathcal{B}(\mathbb{R}), \mathcal{N})$ where $\mathcal{N}$ is the standard Gaussian distribution $\mathcal{N}(0, 1)$. We wish to reconstruct the probability measure with given first and second order moments which minimises the Kullback-Leibler divergence. Our problem is rewritten*

$$\begin{aligned} \min \quad & D_{KL}(P, \mathcal{N}) \\ \text{s.t.} \quad & \int_{\mathbb{R}} dP(t) = 1 \\ & \int_{\mathbb{R}} t dP(t) = m \\ & \int_{\mathbb{R}} t^2 dP(t) = m^2 + 1. \end{aligned}$$

Using the Theorem 9 recalled previously the problem becomes

$$\max_{v_1 \in \mathbb{R}, v_2 < \frac{1}{2}} v_1 m + v_2 (m^2 + 1) - \log \left(\int_{\mathbb{R}} \exp(v_1 t + v_2 t^2) d\mathcal{N}(t) \right). \quad (\text{IV.20})$$

Classical optimality criterion applied to the maximisation problem (IV.20) gives $v_1 = m$ and $v_2 = 0$. Then, the Radon-Nikodym derivative, see [69] and [63], with respect to $\mathcal{N}$ for the reconstructed probability measure P^{ME} is equal to

$$\frac{dP^{ME}}{d\mathcal{N}}(t) = \exp \left(-\frac{m^2}{2} + mt \right).$$

P^{ME} is therefore the Gaussian distribution $\mathcal{N}(m, 1)$.

IV. 2.4 Maximum Entropy on the Mean

We recall in this section how the MEM method works. Let us first recall the problem studied. We assume here that V is compact. V is discretised with a suitable deterministic sequence $t_1, \dots, t_n$ such that point probability measure $P_n = \frac{1}{n} \sum_{i=1}^n \delta_{t_i}$ approximates well the probability

measure P_V . We denote by Q a distribution that generates a vector $(Y_1, \dots, Y_n)$ of n real random amplitudes and by F_n the random point measure defined by

$$F_n = \frac{1}{n} \sum_{i=1}^n Y_i \delta_{t_i}. \quad (\text{IV.8})$$

One wants to determine the "optimal" distribution Q to generate $Y_1, \dots, Y_n$ such that point-measure F_n meets on average the constraints defined as follows

$$\mathbb{E}_Q \left[\int_V \varphi(t) dF_n(t) \right] = \mathbb{E}_Q \left[\frac{1}{n} \sum_{i=1}^n \varphi(t_i) Y_i \right] \in \mathcal{K},$$

with φ a continuous function taking values in $\mathbb{R}^N$.

We set Π a given reference distribution on $\mathbb{R}$ and we study the Kullback-Leibler divergence between the joint distribution Q and the tensor distribution $\Pi^{\otimes n}$ under some constraints. This is summed up in a more concise way by the following problem

$$\begin{aligned} \min \quad & D_{KL}(Q, \Pi^{\otimes n}) \\ Q : \quad & \mathbb{E}_Q \left[\frac{1}{n} \sum_{i=1}^n \varphi(t_i) Y_i \right] \in \mathcal{K}. \end{aligned} \quad (\mathcal{M}_{\varphi, \Pi}^{1,n})$$

We assume the support of Π to be $]a; b[$ with $-\infty \leq a < b \leq \infty$. The domain of the moment generating function of Π is denoted D_Π

$$D_\Pi = \left\{ \tau \in \mathbb{R} : \int_a^b \exp(\tau y) d\Pi(y) < \infty \right\}.$$

We denote by ψ_Π the logarithm of the moment generating function of Π

$$\psi_\Pi(\tau) = \log \int_a^b \exp(\tau y) d\Pi(y), \quad \forall \tau \in D_\Pi. \quad (\text{IV.21})$$

We recall below some results of [34] for the resolution of the n -th MEM problem and the convergence of the obtained solution. First [34, III.3, Lemma 3.1] recalled below gives sufficient conditions for the existence of a solution to the MEM problem $(\mathcal{M}_{\varphi, \Pi}^{1,n})$.

Lemma 11. [34, III.3, Lemma 3.1] *Let us assume the following*

(H1) : It exists at least one $c \in \mathcal{K}$ for which one can find $y_1, \dots, y_n$ with $y_i \in]a; b[$ for $i = 1, \dots, n$, that are such that $\frac{1}{n} \sum_{i=1}^n \varphi(t_i) y_i = c$.

Then, for n sufficiently large, it always exists a solution to the MEM problem $(\mathcal{M}_{\varphi, \Pi}^{1,n})$.

The Corollary 3.1 of [34, Section III.3] describes the solution of the MEM problem $(\mathcal{M}_{\varphi, \Pi}^{1,n})$ when it exists. As for problem $(\mathcal{M}_{\varphi, ME}^1)$, solving the reconstruction problem $(\mathcal{M}_{\varphi, \Pi}^{1,n})$ consists in solving a finite dimension problem that considers the vectors v in $D_{\Pi, \varphi}$, with $D_{\Pi, \varphi}$ the subset of $\mathbb{R}^N$ defined in (H3).

Corollary 12. [34, III.3, Corollary 3.1] *Let us assume that assumption (H1) is satisfied and the following*

(H2) : *The closed convex hull of Π support is $[a; b]$.*

(H3) : *Let the set $D_{\Pi, \varphi}$ be a non empty open set where $D_{\Pi, \varphi}$ is defined by*

$$D_{\Pi, \varphi} = \left\{ v \in \mathbb{R}^N, Z_{\Pi, \varphi}(v) = \int_{[a; b]^n} \prod_{i=1}^n \exp(\langle v, \varphi(t_i) \rangle y_i) d\Pi^{\otimes n}(y_1, \dots, y_n) < \infty \right\}$$

Then, for n sufficiently large, the solution Q_n^{ME} to the MEM problem $(\mathcal{M}_{\varphi, \Pi}^{1, n})$ is

$$Q_n^{ME}(y_1, \dots, y_n) = (Z_{\Pi, \varphi}(v^*))^{-1} \exp \left(\sum_{i=1}^n \langle v^*, \varphi(t_i) \rangle y_i \right) \Pi^{\otimes n}(y_1, \dots, y_n)$$

where $v^ \in D_{\Pi, \varphi} \subset \mathbb{R}^N$ is the unique maximiser of*

$$H_{\Pi, \varphi}^n(v, \mathcal{K}) = \inf_{c \in \mathcal{K}} \langle v, c \rangle - \log(Z_{\Pi, \varphi}(v)).$$

For the convergence result of the MEM reconstruction, we define the function g_n^{ME} by

$$g_n^{ME}(t) = \frac{1}{\sharp M_n(t)} \sum_{i \in M_n(t)} \mathbb{E}_{Q_n^{ME}}[Y_i]$$

with $M_n(t)$ the subset of indices in $[1; n]$ defined by

$$M_n(t) = \left\{ j \in [1; n] : \|t - t_j\| = \min_{i \in [1; n]} \|t - t_i\| \right\}$$

and $\sharp M_n(t)$ is the number of elements in $M_n(t)$.

The next theorem requires the strong assumption denoted by (H6). The notation ∂ of assumption (H6) refers to the edge of the set.

Theorem 13. [34, III.4, Theorem 3.1] *Under assumptions (H1), (H2) and*

(H4) : *D_{Π} is a non empty open set and it exists $v \in \mathbb{R}^k$ such that for all $t \in V$, $\langle v, \varphi(t) \rangle \in D_{\Pi}$.*

(H5) : *$D(\Pi, \varphi) = \{v \in \mathbb{R}^N, \forall t \in V, \langle v, \varphi(t) \rangle \in D_{\Pi}\}$ is non empty.*

(H6) : *$\forall v \in \partial D(\Pi, \varphi)$, we have $\lim_{u \in D(\Pi, \varphi), u \rightarrow v} \left| \int_V \varphi(t) \psi'_{\Pi}(\langle v, \varphi(t) \rangle) dP_V(t) \right| = +\infty$.*

Then g_n^{ME} converges to g_{∞}^{ME}

$$g_{\infty}^{ME}(t) = \psi'_{\Pi}(\langle v^*, \varphi(t) \rangle)$$

where v^ maximises*

$$H_{\Pi, \varphi}^{\infty}(v, \mathcal{K}) = \inf_{c \in \mathcal{K}} \langle v, c \rangle - \int_V \psi_{\Pi}(\langle v, \varphi(t) \rangle) dP_V(t).$$

Remark 1. Let the real random sequence (X_n) be defined by $X_n = \frac{1}{n} \sum_{i=1}^n Y_i \varphi(t_i)$ and let us denote by Q_n the law of X_n . Under the assumptions (H1) and (H2) and provided that ψ_Π is sufficiently regular, one can characterise the asymptotic behaviour of Q_n . As n tends to infinity, Q_n tends to concentrate on the events that belong to the compact set $\mathcal{K}$. That is, let $x \in \mathcal{K}$, we have that

$$Q_n(X_{n,l} > x_l, l = 1, \dots, N) \approx \exp(-nI(x)) \quad (\text{IV.22})$$

where the rate function is $I(x) = \sup_{v \in \mathbb{R}^N} \{ \langle v, x \rangle - \int_V \psi_\Pi(\langle v, \varphi(t) \rangle) dP_V(t) \}$.

This remark is more formally given in [34, III.4, Corollary 3.5] as the large deviation property of the sequence (Q_n) .

Thereafter in the Section IV. 3, we aim to reconstruct the p components of a vectorial measure F subject to the following integral constraints

$$\sum_{i=1}^p \int_V \varphi_l^i(t) dF^i(t) \in \mathcal{K}_l, \quad 1 \leq l \leq N.$$

We will follow the MEM construction given in the present section. For the multidimensional case, the random measure F_n will then be vectorial and the sequence $(Y_i)_{i=1, \dots, n}$ will be a sequence of vectorial amplitudes in $\mathbb{R}^p$.

IV. 2.5 Connection with classical minimisation problems

One can notice that the link between the γ -entropy maximisation problem $(\mathcal{M}_{\varphi, \gamma}^1)$ and MEM problem $(\mathcal{M}_{\varphi, \Pi}^{1,n})$. Indeed if one chooses the convex function γ involved in problem $(\mathcal{M}_{\varphi, \gamma}^1)$ to be the convex conjugate of ψ_Π , we have that

$$\max_{v \in D_{\Pi, \varphi, \psi_\Pi}} H_{\Pi, \varphi}^\infty(v, \mathcal{K}) = \inf_{F \in \mathcal{F}_\mathcal{K}} D_{\gamma_\Pi}(F, P_V)$$

where $H_{\Pi, \varphi}^\infty$ is the function defined in the Theorem 13.

In this section we detail some classical minimisation problems set in the MEM embedding. We set our work under the assumptions of Theorem 13.

Poisson distribution and Kullback-Leibler divergence minimisation

Let the reference distribution Π be a Poisson distribution $\mathcal{P}(\theta)$ with parameter $\theta \in]0; +\infty[$. The support of Π is $\mathbb{N}$. The log-Laplace transform ψ of a Poisson distribution $\mathcal{P}(\theta)$ is

$$\psi_{\mathcal{P}(\theta)}(\tau) = \theta(e^\tau - 1)$$

with domain $D_\Pi = \mathbb{R}$.

Its convex conjugate is the following function

$$\gamma_{\mathcal{P}(\theta)}(y) = y \log(y) - y(1 + \log(\theta)) + \theta, \quad \text{for } y \in \mathbb{R}_+.$$

The associated convex criterion to minimise for problem $(\mathcal{M}_{\varphi,\gamma}^1)$ becomes

$$D_{\gamma_{\mathcal{P}(\theta)}}(F, P_V) = \begin{cases} \int_V \log\left(\frac{dF}{dP_V}\right) dF - (1 + \log(\theta)) F(V) + \theta, & \text{if } F \ll P_V \\ +\infty & \text{and } \frac{dF}{dP_V} \geq 0 \\ & \text{otherwise.} \end{cases}$$

Such criterion gives the Kullback-Leibler divergence when $\theta = 1$

$$D_{KL}(F, P_V) = \begin{cases} \int_V \log\left(\frac{dF}{dP_V}\right) dF - F(V) + 1, & \text{if } F \ll P_V \text{ and } \frac{dF}{dP_V} \geq 0 \\ +\infty & \text{otherwise.} \end{cases}$$

Gaussian distribution and least squares minimisation

Let the reference distribution Π be a Gaussian distribution $\mathcal{N}(m, \sigma^2)$, $\sigma^2 > 0$. The support of Π is $]a; b[=]-\infty; +\infty[$. The log-Laplace transform ψ of a Gaussian distribution $\mathcal{N}(m, \sigma^2)$ is

$$\psi_{\mathcal{N}(m, \sigma^2)}(\tau) = \frac{\tau^2 \sigma^2}{2} + \tau m$$

with domain $D_\Pi = \mathbb{R}$.

Its convex conjugate is the function

$$\gamma_{\mathcal{N}(m, \sigma^2)}(y) = \frac{(y - m)^2}{2\sigma^2}, \quad \text{for } y \in \mathbb{R}.$$

The associated convex criterion to minimise for problem $(\mathcal{M}_{\varphi,\gamma}^1)$ becomes

$$D_{\gamma_{\mathcal{P}(\theta)}}(F, P_V) = \begin{cases} \int_V \frac{1}{2\sigma^2} \left(\frac{dF}{dP_V} - m\right)^2 dP_V, & \text{if } F \ll P_V \\ +\infty & \text{otherwise} \end{cases}$$

which gives the minimisation problem consisting in finding the least squares deviation of Radon-Nikodym derivative $\frac{dF}{dP_V}$ from constant m .

Exponential distribution and Burg entropy minimisation

Let the reference distribution Π be an exponential distribution $\mathcal{E}(\theta)$ with mean θ , $\theta > 0$. The support of Π is $[0; +\infty[$. The log-Laplace transform ψ of the exponential distribution $\mathcal{E}(\theta)$ is

$$\psi_{\mathcal{E}(\theta)}(\tau) = -\log(1 - \tau\theta)$$

with domain $D_\Pi = \left\{\tau < \frac{1}{\theta}\right\}$.

Its convex conjugate is the function

$$\gamma_{\mathcal{E}(\theta)}(y) = \log\left(\frac{\theta}{y}\right) + \frac{y}{\theta} - 1, \quad \text{for } y \in \mathbb{R}_+^*.$$

The associated convex criterion to minimise for problem $(\mathcal{M}_{\varphi,\gamma}^1)$ becomes

$$D_{\gamma_{\mathcal{P}(\theta)}}(F, P_V) = \begin{cases} \log(\theta) - \int_V \log\left(\frac{dF}{dP_V}\right) dP_V + \frac{F(V)}{\theta} - 1, & \text{if } F \ll P_V \\ & \text{and } F^s \geq 0 \\ +\infty & \text{otherwise.} \end{cases}$$

Such criterion gives the Burg-entropy of F when $\theta = 1$, which is the reverse Kullback-Leibler divergence of P_V with respect to F

$$D_{KL}(P_V, F) = \begin{cases} - \int_V \log\left(\frac{dF}{dP_V}\right) dP_V + F(V) - 1, & \text{if } F \ll P_V \text{ and } F^s \geq 0 \\ +\infty & \text{otherwise.} \end{cases}$$

IV. 3 The γ -entropy maximisation problem for the reconstruction of a multidimensional function

In this section we propose to study the γ -entropy maximisation method for the reconstruction of p real-valued functions with domain U (with U compact and non-empty) when they are subjected to very general constraints such as

$$\int_U \sum_{i=1}^p \lambda^i(x) f^i(x) d\Phi_l(x) \in \mathcal{K}_l \quad 1 \leq l \leq N. \quad (\text{IV.23})$$

This study is performed in three independent parts.

First in Section IV. 3.1, we study in the convex analysis framework the problem $(\mathcal{F}_{\Phi,\gamma}^p)$ recalled below

$$\begin{aligned} \max \quad & - \int_U \gamma(f^1, \dots, f^p) dP_U \\ f : \quad & \int_U \sum_{i=1}^p \lambda^i(x) f^i(x) d\Phi_l(x) \in \mathcal{K}_l, \quad l = 1, \dots, N. \end{aligned} \quad (\mathcal{F}_{\Phi,\gamma}^p)$$

We detail the construction of a dual problem of finite dimension. We are able to express an optimal solution thanks to Fenchel duality theorem when the constraint measures Φ_l are absolutely continuous with respect to P_U . However, we show that this first approach does not give a suitable reconstruction when the constraint measures Φ_l are not absolutely continuous with respect to P_U .

To remedy this issue, we propose in Section IV. 3.2 to transfer the functions reconstruction problem $(\mathcal{F}_{\Phi,\gamma}^p)$ to a corresponding reconstruction problem $(\mathcal{M}_{\varphi,\gamma}^p)$ on signed measures. Such problem is recalled below

$$\begin{aligned} \min \quad & D_\gamma(F, P_V) \\ F : \quad & \sum_{i=1}^p \int_V \varphi^i(t) dF^i(t) \in \mathcal{K}. \end{aligned} \quad (\mathcal{M}_{\varphi,\gamma}^p)$$

The linear transfer is performed by means of suitable continuous kernels.

Finally in Section IV. 3.3, we set the problem $(\mathcal{M}_{\varphi,\gamma}^p)$ obtained with the linear transfer into a MEM problem framework. We detail the construction of a sequence of random point-measures for the multidimensional framework. We study the reconstructions given by problem $(\mathcal{M}_{\varphi,\Pi}^{p,n})$.

IV. 3.1 The γ -entropy maximisation problem for the multidimensional case in the convex analysis framework

In this section we study the way to reconstruct a multidimensional function subject to an inverse problems by a γ -entropy maximisation approach. We propose to study such approach within the framework of convex analysis. We recall some general definitions and properties further used in the Appendix C.

We frame our work in a probability space $(U, \mathcal{B}(U), P_U)$ where U is compact (non empty) and P_U is the reference probability measure. We aim at the reconstruction of p real-valued functions such that

$$\int_U \sum_{i=1}^p \lambda^i(x) f^i(x) d\Phi_l(x) \in \mathcal{K}_l \quad 1 \leq l \leq N. \quad (\text{IV.24})$$

As there exists many solutions fitting the previous constraints, we decide to choose among them the solution to a convex problem under constraints. Given a closed convex function $\gamma : \mathbb{R}^p \rightarrow [0; +\infty]$, we decide to characterize the optimal solution $f^o = (f^{1,o}, \dots, f^{p,o})$ that maximises the γ -entropy defined by

$$- \int_U \gamma(f^1(x), \dots, f^p(x)) dP_U(x)$$

provided that f^o satisfies the constraints stated in (IV.24). We denote by $I_\gamma(\cdot)$ the opposite of the γ -entropy that is

$$I_\gamma(f) = \int_U \gamma(f^1(x), \dots, f^p(x)) dP_U(x).$$

To put it in a more concise way, we study problem $(\mathcal{F}_{\Phi,\gamma}^p)$

$$\begin{aligned} & \max -I_\gamma(f) \\ f : & \int_U \sum_{i=1}^p \lambda^i(x) f^i(x) d\Phi_l(x) \in \mathcal{K}_l, \quad l = 1, \dots, N \end{aligned} \quad (\mathcal{F}_{\Phi,\gamma}^p)$$

or equivalently

$$\begin{aligned} & \min I_\gamma(f) \\ f : & \int_U \sum_{i=1}^p \lambda^i(x) f^i(x) d\Phi_l(x) \in \mathcal{K}_l, \quad l = 1, \dots, N. \end{aligned} \quad (\mathcal{F}_{\Phi,\gamma}^p)$$

We denote by ψ the convex conjugate of γ , which is defined by

$$\psi(t) = \sup_y \{ty - \gamma(y)\}.$$

The domain of γ (respectively of ψ) is denoted by $\mathcal{D}_\gamma$ (respectively $\mathcal{D}_\psi$). We will further make the following assumption denoted by $\mathcal{H}_1$.

$\mathcal{H}_1$ γ (respectively its convex conjugate ψ) is a differentiable, closed, essentially strictly convex function for all interior points on its domain $y \in \text{int } \mathcal{D}_\gamma$ (respectively in $\mathcal{D}_\gamma$).

The minimum of $\gamma(y)$ is 0 and is attained at some $y_0 = (y_0^1, \dots, y_0^p)$ such that $y_0 \in \text{int } \mathcal{D}_\gamma$.

The convex function γ is such that the ratio $\frac{\gamma(y)}{\|y\|}$ tends to infinity on the edges of D_γ .

As in the case of the single function reconstruction (recalled in the Section IV. 2), we will need to define the multidimensional analogue of the γ -divergence of signed measures. Such γ -divergence features some terms that depend on the singular parts. As in the case in one dimension, the singular part vanishes when $\frac{\gamma(y)}{\|y\|}$ tends to infinity on the edges of D_γ . The assumption on the ratio $\frac{\gamma(y)}{\|y\|}$ is then taken for the sake of simplicity.

Let E be a convex set of measurable $\mathbb{R}^p$ -valued functions, the minimum of $I_\gamma(\cdot)$ on the convex E will be denoted as

$$I_\gamma(E) := \min_{f \in E} I_\gamma(f).$$

The first result we have is a characterisation of the minimum of $I_\gamma(\cdot)$ over E with respect to a specific convex functional. Define f_1 and f_2 two functions. Provide $\gamma(f_1)$ and $\gamma(f_2)$ are finite P_U a.s., we define γ -Bregman distance (see [14]) of function f_1 and f_2 on U as

$$B_\gamma(f_1, f_2) := \int_U [\gamma(f_1) - \gamma(f_2) - (\nabla \gamma(f_2))^T (f_1 - f_2)] dP_U. \quad (\text{IV.25})$$

One can remark that $B_\gamma(f_1, f_2) \geq 0$, $\forall f_1, f_2$ with finite I_γ values by the convexity of γ . The next theorem characterises the minimum of $I_\gamma(\cdot)$ over a convex set E of functions with respect to the Bregman distance.

Theorem 14. *Given any convex set E of measurable functions on U , such that $I_\gamma(E)$ is finite, there exists a differentiable function f° not necessarily in E such that for all $f \in E$ with $I_\gamma(f) < \infty$:*

$$I_\gamma(f) \geq I_\gamma(E) + B_\gamma(f, f^\circ). \quad (\text{IV.26})$$

In addition, f° is unique P_U -a.s. and any sequence of functions $f_n \in E$, for which $I_\gamma(f_n) \rightarrow I_\gamma(E)$, converges to f° in P_U .

Proof. We adapt the proof of [20] in the multidimensional case proposed in this section. The proof relies on an identity that exists for all function $f \in E$ with finite $I_\gamma(\cdot)$ value and that a I_γ -minimising sequence of function is in some weak sense a Cauchy sequence. In the following, $\|\cdot\|$ will denote the Euclidean norm of $\mathbb{R}^p$ -valued vector.

First notice that for all $\alpha \in]0; 1[$ and all $f, f_1 \in E$ such that $I_\gamma(f)$ and $I_\gamma(f_1)$ are finite, the following equality holds

$$\begin{aligned} \alpha I_\gamma(f) + (1 - \alpha) I_\gamma(f_1) &= I_\gamma(\alpha f + (1 - \alpha) f_1) + \alpha B_\gamma(f, \alpha f + (1 - \alpha) f_1) \\ &\quad + (1 - \alpha) B_\gamma(f_1, \alpha f + (1 - \alpha) f_1). \end{aligned} \quad (\text{IV.27})$$

Let denote by f_2 the function $f_2 = \alpha f + (1 - \alpha)f_1$. One can first notice that by developing and rearranging those terms, we have that

$$\alpha (\nabla \gamma(f_2))^T (f - f_2) + (1 - \alpha) (\nabla \gamma(f_2))^T (f_1 - f_2) = 0.$$

Therefore, only the $I_\gamma(\cdot)$ parts remain in the sum of $\alpha B_\gamma(f, f_2)$ and $(1 - \alpha)B_\gamma(f_1, f_2)$, that is

$$\alpha B_\gamma(f, f_2) + (1 - \alpha)B_\gamma(f_1, f_2) = \alpha I_\gamma(f) + (1 - \alpha)I_\gamma(f_1) + I_\gamma(f_2)$$

and we have identity (IV.27).

Let $(f_n) \subset E$ be a minimising sequence of I_γ and such that for all n , $I_\gamma(f_n) < \infty$. Then (f_n) is a Cauchy sequence in probability, that is

$$\lim_{n, m \rightarrow \infty} P_U(\{x : \|f_n(x) - f_m(x)\| > \varepsilon\}) = 0 \quad \forall \varepsilon > 0.$$

See Lemma 24 in Section D. 2 for the proof.

Then, there exists a subsequence (f_{n_k}) of (f_n) in $\mathbb{R}^p$ such that f_{n_k} converges a.s. to a function f^o in P_U measure and such f^o satisfies the inequality in (IV.26). Indeed by replacing f_1 by f_{n_k} in (IV.27), it becomes

$$\begin{aligned} I_\gamma(f) &= \frac{1}{\alpha} I_\gamma(\alpha f + (1 - \alpha)f_{n_k}) - \frac{1 - \alpha}{\alpha} I_\gamma(f_{n_k}) + B_\gamma(f, \alpha f + (1 - \alpha)f_{n_k}) \\ &\quad + \frac{1 - \alpha}{\alpha} B_\gamma(f_{n_k}, \alpha f + (1 - \alpha)f_{n_k}) \\ &\geq I_\gamma(E) + B_\gamma(f, \alpha f + (1 - \alpha)f_{n_k}) + \frac{1 - \alpha}{\alpha} (I_\gamma(E) - I_\gamma(f_{n_k})), \end{aligned} \quad (\text{IV.28})$$

the last inequality coming from the positivity of Bregman distance and from the fact that $\alpha f + (1 - \alpha)f_{n_k} \in E$ by the convexity of E and therefore $I_\gamma(\alpha f + (1 - \alpha)f_{n_k}) \geq I_\gamma(E)$.

As f_{n_k} is a I_γ -minimising sequence, the last term in (IV.28) tends to 0 when k goes to infinity.

Finally, taking a sequence (α_k) converging slower than f_{n_k} to 0, by Fatou's lemma

$$\liminf_{k \rightarrow \infty} B_\gamma(f, \alpha_k f + (1 - \alpha_k)f_{n_k}) \geq B_\gamma(f, f^o).$$

This proves the existence of f^o in inequality (IV.26). $\square$

Given N real-valued positive measures $\Phi_1, \dots, \Phi_N$, their Lebesgue decompositions are given by $\Phi_l = \phi_l P_U + \Sigma_l$, for $l = 1, \dots, N$ where their Radon-Nikodym derivatives with respects to P_U are denoted by ϕ_l . Measures Σ_l are singular with respect to P_U which means that they are concentrated on a set $\tilde{U}$ such that $P_U(\tilde{U}) = 0$. We denote $E_{\phi, \mathcal{K}}$ the subset of functions as follows

$$\begin{aligned} E_{\phi, \mathcal{K}} &:= \left\{ f = (f^1, \dots, f^p), \right. \\ &\quad \left. \int_U \sum_{i=1}^p \lambda^i(x) f^i(x) \phi_l(x) dP_U(x) \in \mathcal{K}_l, \quad l = 1, \dots, N \right\}. \end{aligned} \quad (\text{IV.29})$$

The next theorem leads to a description of the optimal function $(f^{1,o}, \dots, f^{p,o})$ which is solution on $E_{\phi, \mathcal{K}}$ of $I_\gamma(\cdot)$ minimisation problem. The result is obtained thanks to Fenchel's duality theorem for convex functions and is a generalisation in higher dimension of [21, Th. II.2].

Theorem 15. Suppose there exists a $\mathbb{R}^p$ -valued function $f \in E_{\phi, \mathcal{K}}$ and such that it satisfies

$$f(x) = (f^1, \dots, f^p) \in \text{int } \mathcal{D}_\gamma, \quad P_U - a.s.$$

Let L be the subspace of $\mathbb{R}^N$ such that for given $(\phi_1, \dots, \phi_N)$

$$L = \left\{ v \in \mathbb{R}^N : v_l = \int_U \sum_{i=1}^p \lambda^i(x) f^i(x) \phi_l(x) dP_U(x), l = 1, \dots, N, \right. \\ \left. \text{for some } f : U \rightarrow \mathbb{R}^p \right\}$$

and m be the following application

$$m(v) := \int_U \psi \left(\tau^1(x, v), \dots, \tau^p(x, v) \right) dP_U(x), \quad (\text{IV.30})$$

with $\tau^i(x, v) = \sum_{l=1}^N \lambda^i(x) v_l \phi_l(x)$. The minimum of $I_\gamma(\cdot)$ over the set $E_{\phi, \mathcal{K}}$ can be expressed by

$$I_\gamma(E_{\phi, \mathcal{K}}) = \max_{v \in \mathbb{R}^N} \left\{ \inf_{c \in \mathcal{K} \cap L} \langle v, c \rangle - m(v) \right\}. \quad (\text{IV.31})$$

Then, for $v^o \in \mathbb{R}^N$ which maximises (IV.31), the minimiser of $I_\gamma(\cdot)$ over $E_{\phi, \mathcal{K}}$ is

$$f^{i,o}(x) = \frac{\partial \psi}{\partial \tau_i}(\tau^1(x, v^o), \dots, \tau^p(x, v^o)), \quad \forall i = 1, \dots, p \quad (\text{IV.32})$$

and such $f^o = (f^{o,1}, \dots, f^{o,p})$ satisfies (IV.26).

Proof. Outline of the proof is as in [21] with a multidimensional approach :

1. Expression of the minimum in (IV.31) is given thanks to Fenchel's duality theorem.
2. Letting $\tau^i(x, v) = \sum_{l=1}^N \lambda^i(x) v_l \phi_l(x)$ and v^o the vector of $\mathbb{R}^N$ which maximises (IV.31), one must verify that $(\tau^1(x, v^o) \dots \tau^p(x, v^o))^T$ belongs to the interior of $\mathcal{D}_\psi$.
3. Candidate function f^o must be such that it satisfies the Bregman inequality (IV.26) of Theorem 14.

Let $k(z) := \inf \{ I_\gamma(f) : \int_U \sum_{i=1}^p \lambda^i(x) f^i(x) \phi(x) dP_U(x) = z \}$. Function k is a convex function with $\text{dom } k \subset L$. Let h be such that :

$$h(z) = \begin{cases} 0 & \text{if } z \in \mathcal{K} \cap L \\ +\infty & \text{else.} \end{cases} \quad (\text{IV.33})$$

Then $I_\gamma(E_{\phi, \mathcal{K}}) = \inf_{z \in \mathcal{K}} k(z) = \inf_{z \in \mathbb{R}^N} \{ k(z) + h(z) \}$.

In order to apply Fenchel duality theorem, we need that $\text{dom } k \cap \text{dom } h \neq \emptyset$. Equivalently, that means that there exists $z_0 \in \text{dom } k$ such that $z_0 \in \mathcal{K} \cap L$. From Theorem 15 assumptions, there exists $f \in \mathcal{D}_\gamma \cap E_{\phi, \mathcal{K}}$. By applying Lemma 16 there exists a closed convex set $\tilde{\mathcal{D}}$ included

in $\mathcal{D}_\gamma$ and a function $\tilde{f}$ such that it belongs to $\tilde{\mathcal{D}} \cap E_{\phi, \mathcal{K}}$. As $\tilde{f} \in \tilde{\mathcal{D}}$, $\gamma(\tilde{f})$ is well-defined and $I_\gamma(\tilde{f}) < \infty$. Let $z_0 \in \mathbb{R}^N$ be defined by

$$z_{0,l} = \int_U \sum_{i=1}^p \lambda^i(x) \tilde{f}^i(x) \phi_l(x) dP_U(x), \quad l = 1, \dots, N. \quad (\text{IV.34})$$

By definition $z_0 \in L$, as $I_\gamma(\tilde{f}) < \infty$, $z_0 \in \text{dom } k$ and as $\tilde{f} \in E_{\phi, \mathcal{K}}$, $z_0 \in \mathcal{K}$. Therefore $\text{dom } k \cap \text{dom } h \neq \emptyset$ and Fenchel-Moreau duality theorem can then be applied, see [73] and [74].

With the superscript $*$ denoting the convex conjugate, using Fenchel-Moreau duality theorem, we have the following equality

$$\inf_{z \in \mathbb{R}^N} \{k(z) + h(z)\} = \max_{v \in \mathbb{R}^N} \{-k^*(v) - h^*(-v)\} \quad (\text{IV.35})$$

and then $I_\gamma(E_{\phi, \mathcal{K}}) = \max_{v \in \mathbb{R}^N} \{-k^*(v) - h^*(-v)\}$.

The convex conjugate of k can be expressed as follow :

$$\begin{aligned} k^*(v) &= \sup_{z \in \mathbb{R}^N} \{\langle v, z \rangle - k(z)\} \\ &= \sup_{z \in \mathbb{R}^N} \sup_{f \in E_{\phi, \mathcal{K}}} \left\{ \langle v, z \rangle - \int_U \gamma(f) dP_U \right\} \\ &= \sup \left\{ \langle v, \int_U \sum_{i=1}^p \lambda^i(x) f^i(x) \phi(x) dP_U(x) \rangle - \int_U \gamma(f) dP_U \right\} \\ &= \sup \left\{ \int_U \sum_{i=1}^p f^i(x) \left\{ \sum_{l=1}^N v_l \lambda^i(x) \phi_l(x) \right\} - \gamma(f(x)) dP_U(x) \right\} \\ &= \int_U \psi \left(\sum_{l=1}^N v_l \lambda^1(x) \phi_l(x), \dots, \sum_{l=1}^N v_l \lambda^p(x) \phi_l(x) \right) dP_U(x) \\ &= \int_U \psi \left(\tau^1(x, v), \dots, \tau^p(x, v) \right) dP_U(x), \end{aligned}$$

with $\tau^i(x, v) = \sum_{l=1}^N \lambda^i(x) v_l \phi_l(x)$.

By definition of h , its convex conjugate is :

$$h^*(-v) = \sup_{c \in \mathcal{K} \cap L} \langle -v, c \rangle = - \inf_{c \in \mathcal{K} \cap L} \langle v, c \rangle.$$

Therefore we have equality (IV.31)

$$\begin{aligned} I_\gamma(E_{\phi, \mathcal{K}}) &= \min_{f \in E_{\phi, \mathcal{K}}} \int_U \gamma(f) dP_U \\ &= \max_{v \in \mathbb{R}^N} \{-k^*(v) - h^*(-v)\} \\ &= \max_{v \in \mathbb{R}^N} \left\{ \inf_{c \in \mathcal{K} \cap L} \langle v, c \rangle - m(v) \right\}. \end{aligned} \quad (\text{IV.36})$$

Let us denote by v^o , the vector in $\mathbb{R}^N$ which realises the maximum of $\left\{ \inf_{c \in \mathcal{K} \cap L} \langle v, c \rangle - m(v) \right\}$.

We remind that

$$\begin{aligned} m(v) &= \int_U \psi \left(\tau^1(x, v), \dots, \tau^p(x, v) \right) dP_U(x) \\ \text{with } \tau^i(x, v) &= \sum_{l=1}^N \lambda^i(x) v_l \phi_l(x). \end{aligned}$$

From the assumptions on ψ , $\tau^i(x, v^o)$ belongs to the interior of $\mathcal{D}_\psi$ and therefore one can differentiate ψ at $\tau^i(x, v^o)$.

Let us now demonstrate that $f^o = (f^{1,o}, \dots, f^{p,o})$ with

$$f^{i,o}(x) = \frac{\partial \psi}{\partial \tau_i} \left(\tau^1(x, v^o), \dots, \tau^p(x, v^o) \right), \quad \forall i = 1, \dots, p,$$

satisfies the Bregman inequality in (IV.26) for $E = E_{\phi, \mathcal{K}}$ and for any f such that $I_\gamma(f) < \infty$. Equivalently, it boils down to showing that for any f with $I_\gamma(f) < \infty$

$$\int_U \sum_{i=1}^p \left(\tau^i(x, v) - \frac{\partial \gamma}{\partial y_i}(f^o(x)) \right) (f^{o,i}(x) - f^i(x)) dP_U(x) \geq 0. \quad (\text{IV.37})$$

Inequality (IV.37) comes by using (IV.36) and noticing that for any $\tau \in \mathcal{D}_\psi$

$$\gamma(\nabla \psi(\tau)) = \tau^T \nabla \psi(\tau) - \psi(\tau) \quad (\text{IV.38})$$

and that for any f with $I_\gamma(f) < \infty$

$$\int_U \sum_{i=1}^p \tau^i(x, v^o) f^i(x) dP_U(x) \geq \inf_{z \in \mathcal{K} \cap L} \langle v^o, z \rangle.$$

Let us denote by $\mathcal{D}_\nabla$ the set of γ gradient

$$\mathcal{D}_\nabla = \{ \nabla \gamma(y), y \in \text{int } \mathcal{D}_\gamma \}.$$

For $\tau^o = (\tau^1(x, v^o), \dots, \tau^p(x, v^o)) \in \text{int } \mathcal{D}_\nabla$, using (IV.38) we have $\nabla \gamma(\nabla \psi(\tau^o)) = \tau^o$ and then (IV.37) stands with equality. When $\tau^o \in \mathbb{R}^p \setminus (\text{int } \mathcal{D}_\nabla)$, let $u \in \partial(\mathcal{D}_\gamma)$ such that $\nabla \gamma(u)$ defined below is the projection of τ^o on $\partial \mathcal{D}_\nabla$

$$\nabla \gamma(u) := \lim_{\substack{y \rightarrow u \\ y \in \text{int } \mathcal{D}_\gamma}} \nabla \gamma(y).$$

At such τ^o , we have $\psi(\tau^o) = u^T \tau^o - \gamma(u)$ and so $\nabla \psi(\tau^o) = u$. Therefore

$$\begin{aligned} \nabla \gamma(f^o) &= \nabla \gamma(\nabla \psi(\tau^o)) \\ &= \nabla \gamma(u) \end{aligned}$$

and

$$\begin{aligned} f^{o,i} &= (\nabla \psi(\tau^o))_i \\ &= u_i. \end{aligned}$$

If, for $i \in \{1, \dots, p\}$, we have $\tau^i(x, v^o) \leq (\nabla \gamma(u))_i$, then $u_i \leq f^i$ for any f such that $I_\gamma(f) < +\infty$. Therefore both terms of the product in (IV.37) are negative. Conversely, if we have $\tau^i(x, v^o) \geq (\nabla \gamma(u))_i$, then $u_i \geq f^i$ for any f such that $I_\gamma(f) < +\infty$. Then optimal solution f^o satisfies Bregman inequality. $\square$

To apply Fenchel duality Theorem in the previous Theorem 15, one needs to prove that the domains of the two convex conjugates functions k and h defined in the proof do not have an empty intersection. In order to prove so, the following lemma is required. Lemma 16 features a function $\tilde{f}$ which belongs to a closed convex set $\tilde{\mathcal{D}}$ in the interior of γ domain. Given a function f in the interior of $\mathcal{D}_\gamma$, we prove the existence of $\tilde{f}$ for which the constraint values $\int_U \sum_{i=1}^p \lambda^i(x) \tilde{f}^i(x) \phi_l(x) dP_U(x)$ are equal to the constraints values obtained with function f .

Lemma 16. *Let $\phi_1, \dots, \phi_N$ be given integrable functions. Assume there exists a measurable function $f : U \rightarrow \mathbb{R}^p$ such that*

$$\begin{aligned} f &\in \text{int } \mathcal{D}_\gamma \\ \sum_{i=1}^p \lambda^i(x) f^i(x) \phi_l(x) &\in L^1(P_U), \quad l = 1, \dots, N. \end{aligned}$$

Then, there exist $\tilde{f}$ and $\tilde{\mathcal{D}}$ such that $\tilde{f}$ is a function defined on U , $\tilde{f} \in \tilde{\mathcal{D}}$ and

$$\begin{aligned} &\int_U \sum_{i=1}^p \lambda^i(x) f^i(x) \phi_l(x) dP_U(x) \\ &= \int_U \sum_{i=1}^p \lambda^i(x) \tilde{f}^i(x) \phi_l(x) dP_U(x), \quad l = 1, \dots, N \end{aligned} \tag{IV.39}$$

with $\tilde{\mathcal{D}}$ a closed convex set of $\mathbb{R}^p$ such that $\tilde{\mathcal{D}} \subsetneq \text{int } \mathcal{D}_\gamma$.

Proof. Let L be the subset of vectors of $\mathbb{R}^N$ defined by

$$L = \left\{ v : v_l = \int_U \sum_{i=1}^p \lambda^i(x) h^i(x) \phi_l(x) dP_U(x), \quad l = 1, \dots, N, \quad h : U \rightarrow \mathbb{R}^p \right\}.$$

Let (D_n) be a sequence of closed convex sets such that for all n , $D_n \subsetneq \text{int } D_{n+1}$, meaning that D_n is strictly growing to its limit D_γ . Let $T_n = \{x \in U : f(x) \in D_n\}$.

Let L_n be the subset of vector in $\mathbb{R}^N$ such that

$$\begin{aligned} L_n = \left\{ v : v_l = \int_U \sum_{i=1}^p \lambda^i(x) h^i(x) \phi_l(x) dP_U(x), \quad l = 1, \dots, N, \right. \\ \left. h \text{ bounded on } T_n, h = 0 \text{ on } T_n^c \right\}. \end{aligned}$$

then there exists n_0 such that for all $n \geq n_0$, $L_n = L$. Indeed, if L_n is a proper subset of L , there exists v_n with $\|v_n\| \neq 0$ such that

$$\begin{aligned} \sum_{l=1}^N v_{n,l} \left(\int_U \sum_{i=1}^p \lambda^i(x) h^i(x) \phi_l(x) dP_U(x) \right) &= 0 \\ \int_U \sum_{i=1}^p \lambda^i(x) h^i(x) \left(\sum_{l=1}^N v_{n,l} \phi_l(x) \right) dP_U(x) &= 0. \end{aligned}$$

That implies that $\sum_{l=1}^N v_{n,l} \phi_l(x) = 0$ P_U -a.s. on T_n . Taking a convergent subsequence v_{n_k} with limits v such that $\|v\| \neq 0$, we have $\sum_{l=1}^N v_l \phi_l(x) = 0$ P_U -a.s. on U which goes against $\|v\| \neq 0$.

For $\delta > 0$, defines C^δ by

$$C^\delta = \left\{ v : v_l = \int_U \sum_{i=1}^p \lambda^i(x) h^i(x) \phi_l(x) dP_U(x), \quad l = 1, \dots, N, \right. \\ \left. \|h\| \leq \delta \text{ on } T_{n_0}, \quad h = 0 \text{ on } T_{n_0}^c \right\}.$$

The affine hull of C^δ equals L and $0 \in \text{int } C^\delta$ in the relative topology of L .

We denote by f_n the projection onto D_n of f at x . Therefore $\int_U \sum_{i=1}^p \lambda^i(x) (f_n^i(x) - f^i(x)) \phi_l(x) dP_U(x)$ approaches 0 as n tends to infinity. For $\delta > 0$ there exists then $n_1(\delta)$ such that for all $n \geq n_1(\delta)$, $\int_U \sum_{i=1}^p \lambda^i(x) (f_n^i(x) - f^i(x)) \phi_l(x) dP_U(x)$ belongs to C^δ . Therefore, it exists h defined on U such that $\|h\| \leq \delta$ on T_{n_0} and $h = 0$ on $T_{n_0}^c$ and such that

$$\int_U \sum_{i=1}^p \lambda^i(x) (f_n^i(x) - f^i(x)) \phi_l(x) dP_U(x) = \int_U \sum_{i=1}^p \lambda^i(x) h^i(x) \phi_l(x) dP_U(x).$$

For such h , we set $\tilde{f} = f_n + h$. $\tilde{f}$ belongs to $\tilde{\mathcal{D}}$ with $\tilde{\mathcal{D}} = D_n \cap D_{n_0}^\delta$ and

$$D_{n_0}^\delta = \{f(x) : \|f(x) - f_{n_0}(x)\| \leq \delta\}.$$

□

Next proposition describes under which conditions the infimum of the Theorem 15 is reached for an optimal function in $E_{\phi, \mathcal{K}}$.

Proposition 17. *For a function m defined as in (IV.30) and for v^o which maximises (IV.31). If v^o is an interior point of m domain, then optimal solution f^o , with components $f^{i,o}$ as in (IV.32), belongs to $E_{\phi, \mathcal{K}}$.*

Proof. Recall that

$$I_\gamma(E_{\phi, \mathcal{K}}) = \max_{v \in \mathbb{R}^N} (-k^*(v) - h^*(-v))$$

with $k^*(v) = \int_U \psi(\tau) dP_U$ and $h^*(-v) = - \inf_{z \in \mathcal{K} \cap L} \langle v, z \rangle$.

v^o belongs to the interior of $\text{dom } m$ implies that k^* is differentiable in v^o with its gradient d with components d_l defined for all $l = 1, \dots, N$ by

$$d_l = \int_U \sum_{i=1}^p \lambda^i(x) \frac{\partial \psi}{\partial \tau_i} \left(\tau^1(x, v^o), \dots, \tau^p(x, v^o) \right) \phi_l(x) dP_U(x).$$

By [73, Cor. 23.5.3], the subgradient of $h^*(-v)$ at $v = v^o$, denoted by ∂h^* is included in $(-\mathcal{K})$. As relative interiors of h^* and k^* have a non-empty intersection set, [73, Th. 23.8] implies on the sum of convex function subgradients that $\partial g = \{d\} + \partial h^*$. As g reaches its maximum at v^o , ∂g is a subset which contains 0, which implies that $d \in \mathcal{K}$.

Then, with $f^{i,o}(x) = \frac{d\psi}{d\tau_i}(\tau^1(x, v^o), \dots, \tau^p(x, v^o))$ for all $l = 1, \dots, N$,

$$\int_U \sum_{i=1}^p \lambda^i(x) f^{i,o}(x) \phi_l(x) dP_U(x) \in \mathcal{K}_l.$$

Therefore, $f^o \in E_{\phi, \mathcal{K}}$. □

The following Corollary is deduced from the Theorem 15 when dealing with measures Φ that are not absolutely continuous with respect to P_U .

Let us define the analogue of $E_{\phi, \mathcal{K}}$ for measures Φ that are not absolutely continuous with respect to P_U . $E_{\Phi, \mathcal{K}}$ is the set of function

$$E_{\Phi, \mathcal{K}} := \left\{ f = (f^1, \dots, f^p) : \int_U \sum_{i=1}^p \lambda^i(x) f^i(x) d\Phi_l(x) \in \mathcal{K}_l, \ l = 1, \dots, N \right\}.$$

As a result, we see that when considering measures Φ with a singular part, the optimum function defined in the Theorem 15 does not meet in fact the constraints (IV.24). The corollary stresses on the fact that the problem is ill-posed when dealing with measures Φ which are not absolutely continuous with respect to the reference measure P_U .

Corollary 18. *Suppose there exists a function $f \in E_{\Phi, \mathcal{K}}$ and such that it satisfies*

$$f(x) = (f^1(x), \dots, f^p(x)) \in \text{int } \mathcal{D}_\gamma, \quad P_U - \text{a.s.}$$

Let L, m be as defined in the Theorem 15 and $\Phi = \phi P_U + \Sigma$. The minimum of $I_\gamma(\cdot)$ over $E_{\Phi, \mathcal{K}}$ can be expressed by

$$I_\gamma(E_{\Phi, \mathcal{K}}) = \max_{v \in \mathbb{R}^N} \left\{ \inf_{c \in \tilde{\mathcal{K}} \cap L} \langle v, c \rangle - m(v) \right\} \quad (\text{IV.40})$$

for some $\tilde{\mathcal{K}}$ different of $\mathcal{K}$.

Therefore the optimal solution f^o defined in (IV.32) no longer meets the constraints in (IV.24).

Proof. Let $k_\Phi(c) = \inf \left\{ I_\gamma(f) : \int_U \sum_{i=1}^p \lambda^i(x) f^i(x) d\Phi(x) = c \right\}$ and $\Phi = \phi P_U + \Sigma$. Then,

$$\begin{aligned} I_\gamma(E_{\Phi, \mathcal{K}}) &= \inf_{f \in E_{\Phi, \mathcal{K}}} \{ I_\gamma(f) \} = \inf_{c \in \mathcal{K}} k_\Phi(c) \\ &= \inf \left\{ I_\gamma(f) : \int_U \sum_{i=1}^p \lambda^i(x) f^i(x) \phi(x) dP_U(x) \right. \\ &\quad \left. + \int_U \sum_{i=1}^p \lambda^i(x) f^i(x) d\Sigma(x) = c, \ c \in \mathcal{K} \right\} \\ &= \inf \left\{ I_\gamma(f) : \int_U \sum_{i=1}^p \lambda^i(x) f^i(x) \phi(x) dP_U(x) \right. \\ &\quad \left. = c - \int_U \sum_{i=1}^p \lambda^i(x) f^i(x) d\Sigma(x), \ c \in \mathcal{K} \right\} \\ &= \inf \left\{ I_\gamma(f) : \int_U \sum_{i=1}^p \lambda^i(x) f^i(x) \phi(x) dP_U(x) \tilde{c}, \ \tilde{c} \in \tilde{\mathcal{K}} \right\} \\ &= I_\gamma(E_{\phi, \tilde{\mathcal{K}}}), \end{aligned}$$

where $\tilde{\mathcal{K}} = \mathcal{K} + \left\{ c_\Sigma \in \mathbb{R}^N : c_\Sigma = \int_U \sum_{i=1}^p \lambda^i(x) f^i(x) d\Sigma(x) \right\}$

We apply the Theorem 15 to get expression (IV.40). $\square$

Corollary 18 points out the necessity of a different approach for solving problem described in (IV.24), particularly when dealing with Φ not absolutely continuous with respect to the reference measure on U , as it is the case for solving interpolation problems.

IV. 3.2 Linear transfer principle for the multidimensional case

Following equivalence of problem solutions introduced in [21], the inverse problem on functions described in (IV.24) can be treated as an inverse problem on measures. Sets and generic elements will be distinguished with V and t when discussing the measure reconstruction problem. Measures considered thereafter are always finite real-valued measures. The set of all finite measures on a set V will be denoted by $\mathcal{M}(V)$. The aim is then to reconstruct p real-valued measures $F^i \in \mathcal{M}(V)$ such that

$$\sum_{i=1}^p \int_V \varphi_l^i(t) dF^i(t) \in \mathcal{K}_l, \quad l = 1, \dots, N \quad (\text{IV.41})$$

with φ_l^i being given real-valued functions on V , $\mathcal{K}_l \subset \mathbb{R}$, for all $i = 1, \dots, p$ and $l = 1, \dots, N$.

Let us first denote the following assumption

$\mathcal{H}_2$: V is a compact metric space, P_V is a probability measure having full support, all φ_l^i are continuous and for each $i = 1, \dots, p$, $(\varphi_l^i)_{l=1, \dots, N}$ are linearly independent.

Given the assumption $\mathcal{H}_1$, a solution to problem (IV.41) can then be chosen by taking as optimal solution $(F^{1,o}, \dots, F^{p,o})$ the p -real valued measure which minimises the γ -divergence with respect to the reference measure P_V provided that F^o meets the constraints (IV.41). The opposite of the γ -entropy I_γ and the γ -divergence are linked by the following relation [13, theorem 2.7]

$$D_\gamma(F, P_V) = \int_V \gamma \left(\frac{dF_a^1}{dP_V}, \dots, \frac{dF_a^p}{dP_V} \right) dP_V$$

where F_a^i are measures absolutely continuous with respect to P_V .

Let us denote the set of p -real valued measure with $F^1, \dots, F^p \in \mathcal{M}(V)$ meeting the constraints described in (IV.41) by the set

$$S_{\varphi, \mathcal{K}} = \left\{ (F^1, \dots, F^p) : \sum_{i=1}^p \int_V \varphi_l^i(t) dF^i(t) \in \mathcal{K}_l, \quad l = 1, \dots, N \right\}. \quad (\text{IV.42})$$

The Theorem 19 below describes the γ -divergence minimiser $(F^{o,1}, \dots, F^{o,p})$ such that $(F^{o,1}, \dots, F^{o,p})$ belongs to $S_{\varphi, \mathcal{K}}$. The Theorem 19 is the analogue of the Theorem 15 when dealing with measure reconstruction.

Theorem 19. *Under assumption $\mathcal{H}_2$, suppose we have p measures absolutely continuous with respect to P_V such that $(F^1, \dots, F^p)^T \in S_{\varphi, \mathcal{K}}$ and such that they satisfy*

$$\left(\frac{dF^1}{dP_V}(t), \dots, \frac{dF^p}{dP_V}(t) \right) \in \text{int } \mathcal{D}_\psi, \quad P_V - a.s.$$

Then there exist p real-valued measures $F^{1,o}, \dots, F^{p,o}$, absolutely continuous with respect to P_V and such that $(F^{1,o}, \dots, F^{p,o})$ minimises $D_\gamma(\cdot, P_V)$ over $S_{\varphi, \mathcal{K}}$. Their Radon-Nikodym derivatives $f^{i,o}$ with respect to P_V are defined by

$$f^{i,o}(t) = \frac{\partial \psi}{\partial \tau^i} \left(\tau^1(t, v^o), \dots, \tau^p(t, v^o) \right), \quad \forall i = 1, \dots, p$$

with $\tau^i(t, v) = \sum_{l=1}^N v_l \varphi_l^i(t)$ and with v^o such that it maximises

$$D_\gamma(S_{\varphi, \mathcal{K}}, P_V) := \max_{v \in \mathbb{R}^N} \left\{ \inf_{z \in \mathcal{K}} \langle v, z \rangle - \int_V \psi \left(\tau^1(t, v), \dots, \tau^p(t, v) \right) dP_V(t) \right\}.$$

Proof. Direct from the Theorem 15. □

Having recovered the p measures $F^{i,o}$ described in Theorem 19, a regularized reconstruction f^o is possible via the linear transfer principle. As a matter of fact, one can linearly transfer the constraints (IV.24) to the constraints (IV.41) by using suitable kernel K . Let $K^i(\cdot, \cdot)$ be measurable bounded real-valued functions on $U \times V$ for $i = 1, \dots, p$ such that

$$f_K^i(x) = \int_V K^i(x, t) dF^i(t) \quad \forall i = 1, \dots, p \tag{IV.43}$$

$$\varphi_l^i(t) = \int_U \lambda^i(x) K^i(x, t) d\Phi_l(x) \quad \forall i = 1, \dots, p, \quad \forall l = 1, \dots, N.$$

Then the Fubini theorem links the two sets of constraints as it follows

$$\begin{aligned} \int_U \sum_{i=1}^p \lambda^i(x) f_K^i(x) d\Phi_l(x) &= \sum_{i=1}^p \int_U \lambda^i(x) \left(\int_V K^i(x, t) dF^i(t) \right) d\Phi_l(x) \\ &= \sum_{i=1}^p \int_V \left(\int_U \lambda^i(x) K^i(x, t) d\Phi_l(x) \right) dF^i(t) \\ &= \sum_{i=1}^p \int_V \varphi_l^i(t) dF^i(t). \end{aligned}$$

It is then clear that if $(F^{1,o}, \dots, F^{p,o})$ is a solution for problem $(\mathcal{M}_{\varphi, \gamma}^p)$, then $f_K^o = (f_K^{1,o}, \dots, f_K^{p,o})$ is a regularized solution to problem $(\mathcal{F}_{\Phi, \gamma}^p)$. The components of f_K^o are defined by

$$f_K^{i,o}(x) = \int_V K^i(x, t) dF^{i,o}(t) \quad \forall i = 1, \dots, p$$

with $(F^{1,o}, \dots, F^{p,o})$ defined as in Theorem 19.

The advantage of the kernel transfer method is that if it occurs that some measures Φ are not absolutely continuous with respect to the reference measure -as it is the case when considering

Dirac measures for example-, the linear transfer principle provides continuous functions φ for problem (IV.41) by choosing kernel K to be continuous.

Choice of K is influenced by the prior knowledge on expected properties for the regularized solution f_K , see the applications in Section IV. 4 for some examples.

IV. 3.3 The embedding into the MEM framework for the multidimensional case

In this section we study the MEM method in the multidimensional case. We first detail the construction of the random point-measure involved in the reconstruction problem.

As before in Section IV. 2, define a sequence of discrete probability measures $(P_n)_n$ as follows

$$P_n = \frac{1}{n} \sum_{j=1}^n \delta_{t_j} \quad (\text{IV.44})$$

with $(t_j)_{j=1,\dots,n}$ a deterministic sequence of points in V such that the limit of probability measures P_n is the reference probability measure P_V . For each $i = 1, \dots, p$, a real valued random variable Y_j^i is associated to t_j . The random variable Y_j^i can be seen as a random amplitude for a signed measure F^i at location t_j . Then let F_n^i , for $i = 1, \dots, p$ be p random measures such that $F_n^i \ll P_n$ and such that they are defined for all $t \in V$ by

$$F_n^i(t) = \frac{1}{n} \sum_{j=1}^n Y_j^i \delta_{t_j}(t).$$

For all $j = 1, \dots, n$ the p real-valued vector of random variables $\mathbf{Y}_j = (Y_j^1, \dots, Y_j^p)^T$ is sampled from a reference distribution Π and we denote by $\Pi^{\otimes n}$ the joint distribution of the n independent, identically sampled vectors $\mathbf{Y}_n$ of dimension p . Replacing F^i by the discretised measure F_n^i , the measure constraint (IV.41) can be rewritten in the following way

$$\frac{1}{n} \sum_{j=1}^n \sum_{i=1}^p \varphi_l^i(t_j) Y_j^i \in \mathcal{K}_l \quad 1 \leq l \leq N. \quad (\text{IV.45})$$

MEM method consists then in finding the optimal joint distribution Q_n^{MEM} such that it minimises the divergence from the reference distribution $\Pi^{\otimes n}$ and such that the discrete constraints (IV.45) are met on average. The MEM problem can be rewritten as

$$\min_{Q \in \mathcal{Q}_{\varphi, \mathcal{K}}^n} D_{KL}(Q, \Pi^{\otimes n}). \quad (\text{IV.46})$$

where $\mathcal{Q}_{\varphi, \mathcal{K}}^n$ defines the set of distributions Q which generates $n \times p$ random amplitudes Y_j^i such that the expectation under Q of the constraints is satisfied, that is

$$\mathcal{Q}_{\varphi, \mathcal{K}}^n = \left\{ (\mathbf{Y}_1, \dots, \mathbf{Y}_n)^T \sim Q : \mathbb{E}_Q \left[\frac{1}{n} \sum_{j=1}^n \sum_{i=1}^p \varphi_l^i(t_j) Y_j^i \right] \in \mathcal{K}_l, \quad l = 1, \dots, N \right\},$$

where $\mathbf{Y}_j$ is the j -th sample $(Y_j^1, \dots, Y_j^p)$ of amplitudes for the p random measures $F_n^1, \dots, F_n^p$.

Let us denote the following assumption

$\mathcal{H}_3$: function γ considered in the γ -divergence problem $(\mathcal{F}_{\Phi, \gamma}^p)$ is the function γ_Π that is such that its conjugate function ψ_Π has its domain equal to $\mathbb{R}^p$ and corresponds to the logarithm of the moment generating function of Π

$$\psi_\Pi(\tau_1, \dots, \tau_p) = \log \int_{\mathbb{R}^p} \exp(\tau^T y) d\Pi(y). \quad (\text{IV.47})$$

Notice that the components of $\nabla \psi_\Pi$ are then

$$\frac{\partial \psi_\Pi}{\partial \tau_i}(\tau_1, \dots, \tau_p) = \int y_i \exp(\tau^T y - \psi_\Pi(\tau_1, \dots, \tau_p)) d\Pi(y). \quad (\text{IV.48})$$

Provided that it exists $y_j = (y_j^1 \dots y_j^p)^T$ in the interior of Π domain for $j = 1, \dots, n$ such that

$$\frac{1}{n} \sum_{j=1}^n \sum_{i=1}^p \varphi_l^i(t_j) y_j^i \in \mathcal{K}_l, \quad l = 1, \dots, N, \quad (\text{IV.49})$$

by standard theory of the ME method the minimiser Q_n^{MEM} of $K(Q, \Pi^{\otimes n})$ exists and it belongs to the exponential family through $\Pi^{\otimes n}$ spanned by the statistics $\frac{1}{n} \sum_{j=1}^n \sum_{i=1}^p \varphi_l^i(t_j) y_j^i$ for $l = 1, \dots, N$. Its expression is given by

$$Q_n^{MEM}(y_1, \dots, y_n) = \exp \left(\sum_{j=1}^n \tau_j^T y_j - \psi_\Pi(\tau_j) \right) \Pi^{\otimes n}(y_1, \dots, y_n) \quad (\text{IV.50})$$

where $\tau_j = (\tau^1(t_j, v_n^o), \dots, \tau^p(t_j, v_n^o))$ and $\tau^i(t, v) = \sum_{l=1}^N v_l \varphi_l^i(t)$ and v_n^o is the maximiser of the discrete dual problem

$$H_n(v) = \inf_{c \in \mathcal{K}} \langle v, c \rangle - \frac{1}{n} \sum_{j=1}^n \psi_\Pi(\tau^1(t_j, v), \dots, \tau^p(t_j, v)). \quad (\text{IV.51})$$

We define the vector of measures $F_n^{MEM} = (F_n^{1, MEM}, \dots, F_n^{p, MEM})^T$ with each component defined by

$$F_n^{i, MEM} = \mathbb{E}_{Q_n^{MEM}} [F_n^i]. \quad (\text{IV.52})$$

The next theorem describes the convergence of $(F_n^{MEM})_n$ sequence to the solution of the γ -divergence minimisation problem on signed measures $(\mathcal{M}_{\varphi, \gamma}^p)$.

Theorem 20. *Under assumptions $\mathcal{H}_2$ and $\mathcal{H}_3$, suppose there exists p measures $(F^1, \dots, F^p)^T \in S_{\varphi, \mathcal{K}}$ such that they satisfy*

$$\left(\frac{dF^1}{dP_V}(t), \dots, \frac{dF^p}{dP_V}(t) \right) \in \text{int } \mathcal{D}_\Pi \quad P_V - a.s.$$

Then the sequence (F_n^{MEM}) converges to $(F^{1, o}, \dots, F^{p, o})$ the minimiser of $D_\gamma(S_{\varphi, \mathcal{K}}, P_V)$ which Radon-Nikodym derivatives $f^{i, o}$ with respect to P_V are defined by

$$f^{i, o}(t) = \frac{\partial \psi}{\partial \tau^i}(\tau^1(t, v^o), \dots, \tau^p(t, v^o)), \quad \forall i = 1, \dots, p$$

with $\tau^i(t, v) = \sum_{l=1}^N v_l \varphi_l^i(t)$ and with v^o such that it maximises

$$D_\gamma(S_{\varphi, \mathcal{K}}, P_V) := \max_{v \in \mathbb{R}^N} \left\{ \inf_{z \in \mathcal{K}} \langle v, z \rangle - \int_V \psi(\tau^1(t, v), \dots, \tau^p(t, v)) dP_V(t) \right\}.$$

To link the problem studied previously with the constraint of function reconstruction problem $(\mathcal{F}_{\Phi, \gamma}^p)$ recalled below

$$\int_U \sum_{i=1}^p \lambda^i(x) f^i(x) d\Phi_l(x) \in \mathcal{K}_l \quad 1 \leq l \leq N,$$

one can consider the analogue of constraint (IV.45) for the function to reconstruct.

As problem $(\mathcal{F}_{\Phi, \gamma}^p)$ and problem $(\mathcal{M}_{\varphi, \gamma}^p)$ can be linked by the linear transfer studied in the Section IV. 3.2, there exists a regularized function f_K^i , associated to the reconstructed measure F^i and chosen kernel K^i , that is defined by

$$f_K^i(x) = \int_V K^i(x, t) dF^i(t) \quad \forall i = 1, \dots, p. \quad (\text{IV.53})$$

Then the function of interest f_K^i can be approximated by a random function f_n^i defined for all $x \in U$ by

$$\begin{aligned} f_n^i(x) &= \int_V K^i(x, t) F_n^i(dt) \\ &= \frac{1}{n} \sum_{j=1}^n Y_j^i K^i(x, t_j). \end{aligned} \quad (\text{IV.54})$$

For Q_n^{MEM} the solution of problem (IV.46), the regularized function $f_{n,K}^{MEM}$ has its components $f_{n,K}^{i, MEM}$ defined by

$$f_{n,K}^{i, MEM}(x) = \mathbb{E}_{Q_n^{MEM}} [f_n^i(x)] = \mathbb{E}_{Q_n^{MEM}} \left[\frac{1}{n} \sum_{j=1}^n Y_j^i K^i(x, t_j) \right]. \quad (\text{IV.55})$$

Then if distribution Q_n^{MEM} is a solution of problem $(\mathcal{M}_{\varphi, \Pi}^{p, n})$, the regularized solution $f_{n,K}^{MEM}$ defined above meets an approximation of the constraints of problem (IV.3). Such approximation is given by

$$\mathbb{E}_Q \left[\frac{1}{n} \sum_{j=1}^n \sum_{i=1}^p \varphi_l^i(t_j) Y_j^i \right] \in \mathcal{K}_l, \quad l = 1, \dots, N \quad (\text{IV.56})$$

From previous expression, it is easy to see that properties of the solution $f_{n,K}^{i, MEM}$ directly depends on the random amplitudes $(Y_j^i)_{j \in \mathbb{N}}$ properties and on the kernel K^i properties.

IV. 4 Applications

IV. 4.1 Application in the case $p = 1$

We consider in this section simple examples of functions reconstruction of one or two variables. The first example considered is the reconstruction of a real-valued convex function of one variable, $f : [-1; 1] \rightarrow \mathbb{R}$ which is solution of an interpolation problem. The second example considered is the reconstruction of a polynomial function of two variables, $f : [0; 1] \times [0; 1] \rightarrow \mathbb{R}$ which is solution of an interpolation problem.

IV. 4.1.1 Reconstruction of an univariate convex function

The first example considered is the reconstruction of a real-valued convex function of one variable, $f : U \rightarrow \mathbb{R}$ which is solution of an interpolation problem with $U = [-1; 1]$. We give $N = 3$ interpolation constraints. The function has a minimal value y_0 which is reached for a value $x_0 \in]-1; 1[$. The pair (x_0, y_0) consists in the first interpolation constraint. The two other points will be denoted by (x_1, y_1) and (x_2, y_2) where x_1 and x_2 belong respectively to the interval $] - 1; x_0[$ and $]x_0; 1[$, where the reconstructed function will be respectively decreasing and increasing. The set of interpolation values will be denoted by $\mathbf{z} = (y_1 - y_0, y_2 - y_0)$. For a reason explained in the following, the interpolation constraints are expressed as the increment from the minimum value.

In this example, we consider the log Laplace transform associated with the Poisson distribution

$$\psi_{\mathcal{P}(1)}(\tau) = (e^\tau - 1)$$

for the convex criterion to minimise. The objective function (IV.51) associated with the MEM problem becomes

$$\begin{aligned} H_n(v) &= \langle v, \mathbf{z} \rangle - \frac{1}{n} \sum_{j=1}^n \psi_{\mathcal{P}(1)}(\langle v, \boldsymbol{\varphi}(t_j) \rangle) \\ &= \langle v, \mathbf{z} \rangle - \frac{1}{n} \sum_{j=1}^n (\exp(\langle v, \boldsymbol{\varphi}(t_j) \rangle) - 1). \end{aligned} \quad (\text{IV.57})$$

The objective function (IV.57) has an analytic minimum when the number of discretisation points n is equal to 1. The analytic solution is

$$v_{o,l} = \frac{1}{\varphi_l(t)} \log \left(\frac{z_l}{\varphi_l(t)} \right), \quad l \in \{1; 2\}.$$

Otherwise, one can use the polynomial approximation of the exponential function. Solving the MEM problem is then reduced to finding the root of a polynomial.

In [51] the authors proved that

$$K_m^+(x, t) = \frac{(x - t)^{m-1}}{(m-1)!} \mathbf{1}_{[x_0; x]}(t), \quad t \in [-1; 1]$$

is a kernel which leads by the linear transfer to an increasing convex function with $m - 1$ derivatives and which is equal to 0 in x_0 . In our frame, the kernel used for the linear transfer is

$$K(x, t) = K_m^+(x, t) \mathbf{1}_{x > x_0}(x) + K_m^-(x, t) \mathbf{1}_{x < x_0}(x), \quad t \in V = [-1; 1]$$

where

$$K_m^-(x, t) = \frac{(t - x)^{m-1}}{(m-1)!} \mathbf{1}_{[x; x_0]}(t).$$

This will lead to a reconstructed function which reaches its minimum value 0 in x_0 , which is decreasing on the interval $[-1; x_0]$ and increasing on the interval $[x_0; 1]$.

In the end, the reconstructed function is given by

$$f_{n,K}(x) = y_0 + \frac{1}{n} \sum_{j=1}^n K(x, t_j) \exp \left(\sum_{l=1}^N v_{o,l} \varphi_l(t_j) \right).$$

The reconstructed function we obtain is displayed in Figure IV.3.

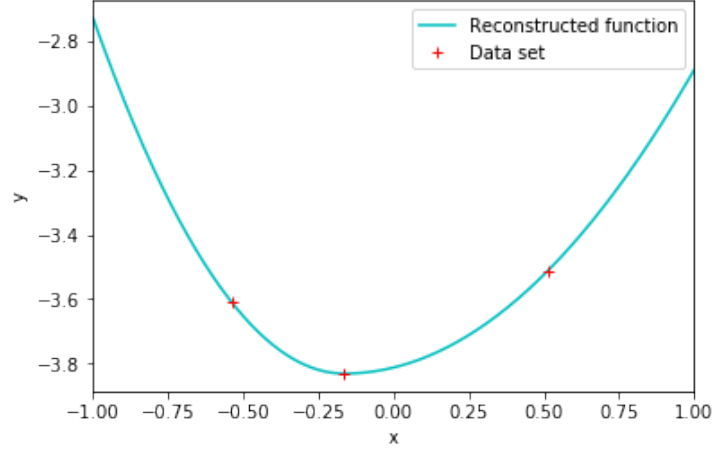


FIGURE IV.3 – *Poisson Log Laplace. Reconstruction of a convex function $f : [-1; 1] \rightarrow \mathbb{R}$ twice differentiable everywhere with 3 interpolation constraints. The minimum of the function is assumed to be known.*

IV. 4.1.2 Reconstruction of a bivariate polynomial function

The second example considered is the reconstruction of a polynomial function of two variables, $f : U \rightarrow \mathbb{R}$ which is solution of an interpolation problem with $U = [0; 1] \times [0; 1]$. We will choose an increasing number N of interpolation points thanks to a Latin Hypercube Sampler in $[0; 1] \times [0; 1]$, the domain of f . We denote by $\mathbf{z}$ the values of function f to interpolate at the design points.

In this example, we consider the log Laplace transform associated with the standard Gaussian distribution $\mathcal{N}(0, 1)$

$$\psi_{\mathcal{N}(0,1)}(\tau) = \frac{\tau^2}{2}$$

for the convex criterion to minimise. The objective function (IV.51) associated with the MEM problem becomes

$$\begin{aligned} H_n(v) &= \langle v, \mathbf{z} \rangle - \frac{1}{n} \sum_{j=1}^n \psi_{\mathcal{N}(0,1)}(\langle v, \varphi(t_j) \rangle) \\ &= \langle v, \mathbf{z} \rangle - \frac{1}{2n} \sum_{j=1}^n (\langle v, \varphi(t_j) \rangle)^2. \end{aligned} \tag{IV.58}$$

We choose $n = 100$ and the n discretising points are sampled from a Latin Hypercube Sampler in $V = [0; 1] \times [0; 1]$. The objective function (IV.58) has an analytic solution provided that the

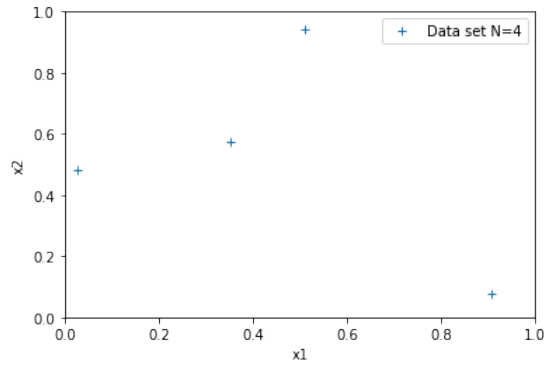
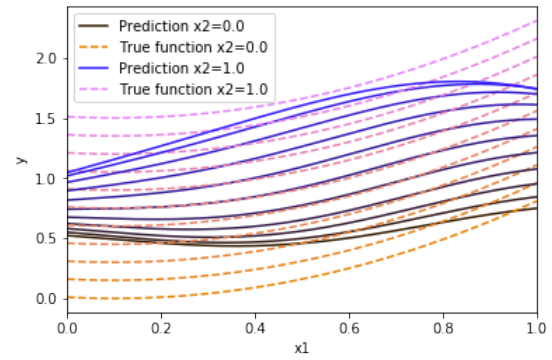
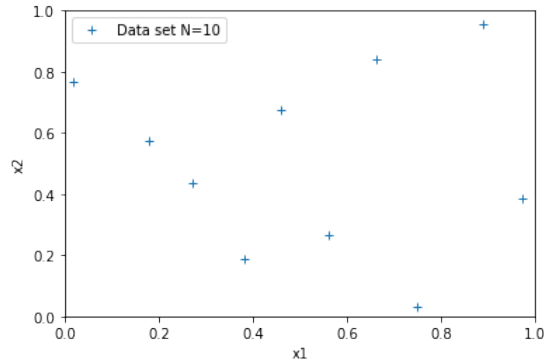
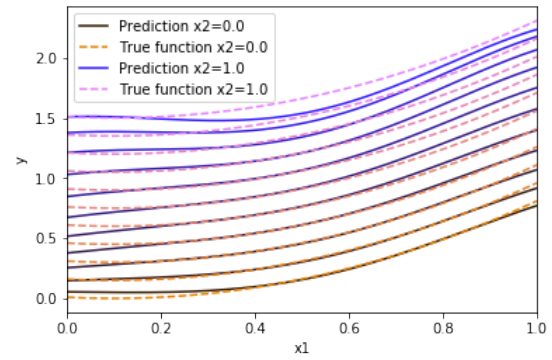
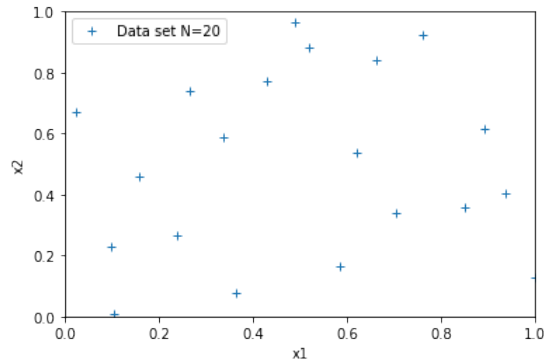
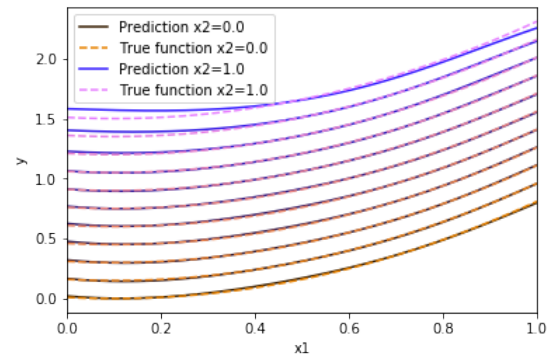

 (a) $N = 4$. Location of points used for training.

 (b) $N = 4$. Reconstructed function.

 (c) $N = 10$. Location of points used for training.

 (d) $N = 10$. Reconstructed function.

 (e) $N = 20$. Location of points used for training.

 (f) $N = 20$. Reconstructed function.

FIGURE IV.4 – *Gaussian Log Laplace*. Reconstructed functions $f : [0; 1] \times [0; 1] \rightarrow \mathbb{R}$ with the gaussian kernel for an increasing number of interpolation constraints. Figures (IV.4a), (IV.4c) and (IV.4e) display the design points used for the interpolation problem. Figures (IV.4b), (IV.4d) and (IV.4f) display the evolution of the reconstructed function with respect to its first component x_1 for several values of the second component x_2 . x_2 varies from 0 to 1 with a 0.1 step.

number of discretising points n is much larger than the number of constraints N . Notice that it is also required that the N components of the moment function $\boldsymbol{\varphi}$ are linearly independent. However this condition is already an assumption in $\mathcal{H}_2$ used for the Theorem 19. Optimal $\mathbf{v}$ is solution of the linear problem

$$\mathbf{A}\mathbf{v} = \mathbf{z}$$

where matrix $\mathbf{A}$ is

$$\mathbf{A} = \frac{1}{n} \sum_{j=1}^n \boldsymbol{\varphi}(t_j) \cdot \boldsymbol{\varphi}(t_j)^T.$$

We consider the symmetric gaussian kernel

$$K_{G,\theta}(u, w) = \exp\left(-\frac{(u-w)^2}{2\theta}\right), \quad (u, w) \in [0; 1] \times [0; 1]$$

which is largely used as a covariance kernel in kriging problems which leads to solutions infinitely differentiable. The kernel we use for a pair of points x and t where $(x, t) \in U^2 \times V^2 = [0; 1]^2 \times [0; 1]^2$ is the following product of gaussian kernels

$$\begin{aligned} K_\theta(x, t) &= K_{G,\theta_1}(x_1, t_1) K_{G,\theta_2}(x_2, t_2) \\ &= \exp\left(-\frac{(x_1 - t_1)^2}{2\theta_1}\right) \exp\left(-\frac{(x_2 - t_2)^2}{2\theta_2}\right) \end{aligned}$$

The best parameter θ is chosen using cross-validation, namely we choose the value of θ which minimises

$$\sum_{k=1}^N \left(f_{K_\theta}^{(k)}(x_k) - z_k \right)^2$$

with z_k the k -th value in the interpolation problem located at point x_k and $f_{K_\theta}^{(k)}$ is the reconstructed function obtained in removing (x_k, z_k) from the data set. This leads to a two step optimisation problem as matrix $\mathbf{A}$ depends on the value of θ and so does the optimal multiplier $\mathbf{v}$. Therefore in the following matrix $\mathbf{A}$ will be denoted with a subscript $\mathbf{A}_\theta$.

In order to determine the optimal parameter θ , we solve iteratively the following two stage problem

$$\begin{aligned} \text{Step 1 : } & \mathbf{v}_{m+1}^{(k)} \\ & \text{solve } \mathbf{A}_{\theta_m}^{(k)} \mathbf{v} = \mathbf{z}^{(k)} \text{ for all } k = 1, \dots, N. \\ \text{Step 2 : } & \theta_{m+1} \end{aligned}$$

$$\text{minimise } \sum_{k=1}^N \left(f_{m+1, K_\theta}^{(k)}(x_k) - z_k \right)^2.$$

The superscript $^{(k)}$ is used to note that the k -th experiment has been removed, that is that the k -th line has been removed from $\mathbf{A}_{\theta_m}$ for the matrix $\mathbf{A}_{\theta_m}^{(k)}$ and the k -th value from $\mathbf{z}$ for the vector of observation $\mathbf{z}^{(k)}$. The notation $f_{m+1, K_\theta}^{(k)}$ corresponds to the reconstructed function at the m -th iteration which is solution of the interpolation problem in which the k -th experiment has been removed.

The reconstructed function we obtain is displayed in Figure IV.4.

IV. 4.2 Applications in the case $p \neq 1$

We consider thereafter some toy models inspired from computational thermodynamics. At first we explain how to compute the so-called phase diagram. Then two toy models are solved using the method described in Section IV. 3.

IV. 4.2.1 Phase diagram description in Computational Thermodynamics

The application considered here derived from the assessment problem in Thermodynamics and the CALPHAD (CALculation of PHase Diagram) method [58], [46], [57], [87].

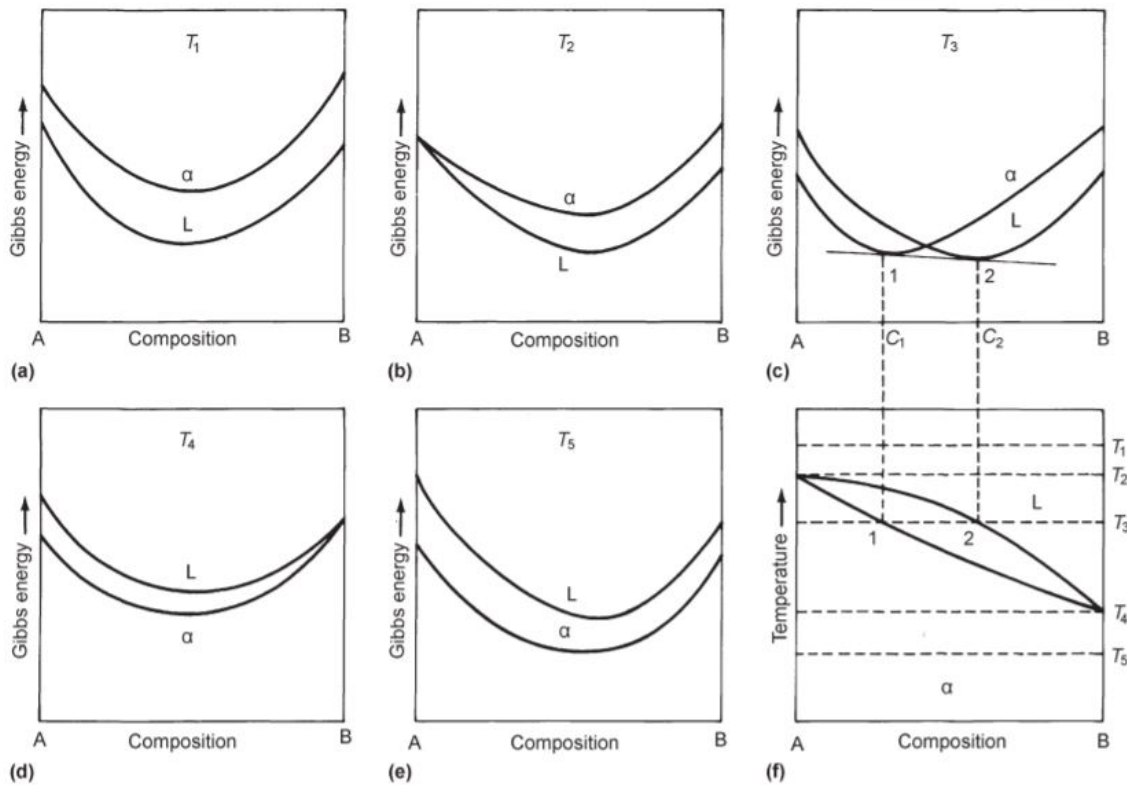


FIGURE IV.5 – Connections between the phase diagram (f) of a system A-B with two phases (the phase α for low temperature and the phase L for high temperature) and the inner energy (a)-(e) for each phase for different temperature values.

Figures (a)-(e) display the Gibbs energy for phases α and L at a given temperature with respect to the relative composition of chemical element B over the sum of chemical elements A and B. For temperature T_1 and T_2 , phase L has the lowest energy whereas for temperature T_4 and T_5 phase α has the lowest energy. At temperature T_3 , there exists a common tangent to the two energy functions. Therefore between composition C_1 and C_2 , a combination of phases α and L is the most stable.

Figure reference [64].

The CALPHAD method consists in the parametric reconstruction from partial information

of thermodynamic quantities, which are Gibbs energy functions and their derivatives. Data at hand for the reconstruction are thermodynamic quantities and phase diagram data. The thermodynamic quantities consist in linear transformations (e.g. first or second derivative) of the function which is ought to be reconstructed.

The phase diagram is a map of the chemical species spatial arrangements. The internal energy of such arrangements varies with state variables, which are the chemical composition, the temperature and the pressure. Establishing the phase diagram of a chemical system means to partition the domain of permissible state variables in several areas, each area featuring one or several stable phases. A stable phase is the phase with the lowest energy.

Figure IV.5 displays an example of a phase diagram. In order to determine the different divided areas of the diagram, one must compute the minimising convex hull of the list of the p functions involved in the system. This problem is performed for all temperature range. For the composition range where the minimising convex hull is confounded with a single energy function f^i , the corresponding area is the stability area of phase i . Such area is labelled by i in the phase diagram.

The phase diagram in Figure IV.5 is for a *binary system* A-B. Such phase diagram is called *binary phase diagram* as it features two chemical elements : A and B. In this example, the state variables which can vary are the temperature and the relative composition of B with respect to the total composition (that is of A and B), from solely element A at the left edge to solely element B at the right edge. The pressure is fixed.

The phase diagram data consist in locations in the map of stable phases, namely either temperatures or compositions where a change in the set of stable phases occurs or information on the stable phases for a given composition and a given temperature.

The novelty of the work presented here is that the reconstruction occurs in a non-parametric frame, contrarily as the usual reconstruction frame in thermodynamics.

Thereafter, the reconstruction of thermodynamic quantities is expressed as the inverse problem $(\mathcal{F}_{\Phi, \gamma}^p)$. The constraints that the solution $(f^1, \dots, f^p)$ must satisfy is recalled below

$$\int_U \sum_{i=1}^p \lambda^i(x) f^i(x) d\Phi_l(x) \in \mathcal{K}_l \quad 1 \leq l \leq N.$$

The functions to recover f^i , $i = 1, \dots, p$ are the Gibbs energy functions of the p stable phases which occur in the phase diagram. The functions λ^i represent the phase diagram data for each phase $i = 1, \dots, p$ and are supposed to be known all over the domain of permissible state variables. In the example considered thereafter, there are bounded, continuous functions but not necessary differentiable.

In this frame, x is a $(d + 1)$ -tuple of positive bounded quantities. First component is the temperature in Kelvin with values in $[T_0; T_{\max}]$, transposed to $[0; 1]$ without loss of generality. The d following components are composition data with values in $[0; 1]^d$. Kernels K^i , for all

$i = 1, \dots, p$, considered for the linear transfer, will be chosen as products of $d + 1$ component-wise kernels. In the examples considered in the following $d = 1$ or 2 .

IV. 4.2.2 Phase diagram with an ideal solution

In this section we study the case of a phase diagram with an *ideal solution*. It consists in the reconstruction of two functions, viz. the Gibbs energy functions associated with the two phases, outside of the domain where they are known.

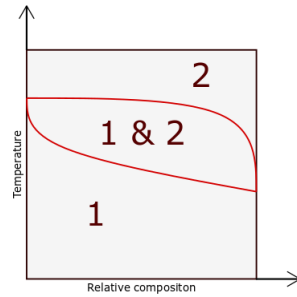


FIGURE IV.6 – Phase diagram corresponding to an ideal solution. There are two phases : the phase 1 is the stable phase for low temperature and the phase 2 for the high temperature. Behaviour with respect to the temperature of phases 1 and 2 are known when the composition is 0 or 1 for the temperature range they are respectively the stable phase.

The inverse problem associated to the functions to reconstruct is built in accordance with the phase diagram. The phase diagram with an ideal solution is usually displayed as in Figure IV.6. It features two phases and three divided area. The first area at the bottom of the diagram is the area associated with phase 1 which is the stable phase at low temperature. The second area at the top of the diagram is the area associated with phase 2 which is the stable phase at high temperature. The last area between the two first areas features a mix of phase 1 and phase 2.

Such problem may not be well-posed as there can exist an infinity of pairs of functions leading to the same phase diagram. In the following the univariate case is treated, which corresponds to the case of a given temperature value. Then functions solely depend on the relative composition of element B with respect to the total composition.

In the univariate case, the pair of functions $(f^1(x), f^2(x))$ is ought to be reconstructed for all $x \in [0; 1]$. In this case, the phase diagram constraints reduce to the following definition. Given $x_1, x_2 \in [0; 1]$ with $x_1 < x_2$. The interval $[x_1; x_2]$ consists in the compositions for which phase 1 and phase 2 coexist. Therefore functions λ^1 and λ^2 are defined as follows

$$\lambda^1(x) = \begin{cases} 1 & \text{if } x < x_1 \\ \frac{x_2 - x}{x_2 - x_1} & \text{if } x \in [x_1; x_2] \\ 0 & \text{else} \end{cases}$$

and

$$\lambda^2(x) = \begin{cases} 0 & \text{if } x < x_1 \\ 1 - \lambda^1(x) & \text{if } x \in [x_1; x_2] \\ 1 & \text{else.} \end{cases}$$

Observations for $x < x_1$ give values of solely f^1 and for $x > x_2$ give values of solely f^2 . In the interval $[x_1; x_2]$, one can observe

$$\lambda^1(x)f^1(x_1) + \lambda^2(x)f^2(x_2) \quad (\text{IV.59})$$

which does not provide direct information about the behaviour of f^1 nor f^2 on this interval.

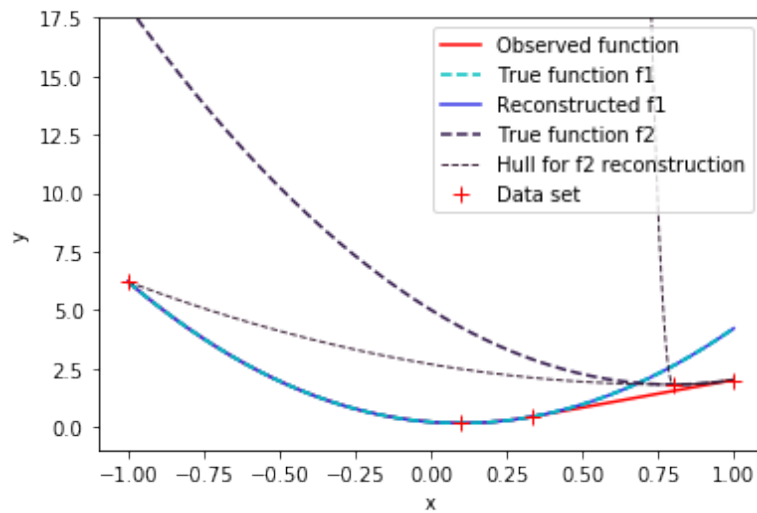


FIGURE IV.7 – *Bivariate Poisson Log Laplace. Reconstruction of two convex functions $f^1 : [-1; 1] \rightarrow \mathbb{R}$ and $f^2 : [-1; 1] \rightarrow \mathbb{R}$, twice differentiable everywhere. The minimum is assumed to be known for both functions.*

The two functions to be reconstructed ought to be convex twice differentiable. Therefore the same assumptions are taken than in the case treated in Section IV. 4.1.1 and the same method is applied. Figure IV.7 displays the results obtained for the reconstruction of the target functions.

Function f^1 is well reconstructed in this example but there exists a set of possible reconstructions for function f^2 . An extra constraint is required to have a unique solution for f^2 .

Annexe A

Bibliographic references for the thermodynamic data

Data type	Related phases	Bibliographic references
Heat capacities	CuMg ₂	[32]
	C15	[32], [96]
Mg partial pressure	Liquid	[81], [49], [38]
	C15 + CuMg ₂	[85]
	FCC + C15	[85]
Mixing entropies (EMF)	C15 + CuMg ₂	[31]
	FCC + C15	[31]
Mixing enthalpies	Liquid	[9], [45], [86]
Formation enthalpies	CuMg ₂	[52], [32]
	C15	[52]
Melting enthalpies	Liquid + CuMg ₂	[32]
	Liquid + C15	[32]
Formation enthalpies	C15	[90], [96]
	CuMg ₂	[90], [96]
	FCC	[83], [47]
	HCP	[83]
Liquidus	Liquid + FCC	[76], [48], [91]
	Liquid + C15	[76], [48], [91]
	Liquid + CuMg ₂	[76], [48], [91]
	Liquid + HCP	[76], [48], [91]
Eutectics	Liquid + FCC + C15	[48], [6]
	Liquid + C15 + CuMg ₂	[48], [6]
	Liquid + CuMg ₂ + HCP	[48]
Congruent meltings	Liquid + C15	[48], [6]
	Liquid + CuMg ₂	[48], [6]
Homogeneity range	FCC + C15	[6]
	C15 + CuMg ₂	[6]
Solvus	Liquid + FCC	[6], [75]
	Liquid + C15	[6], [75]
	Liquid + CuMg ₂	[48], [40], [82]
	Liquid + HCP	[6]

TABLE A.1 – Complete list of data and experiments used for the uncertainty study of the Cu-Mg system.

Annexe B

Classical algorithms of Bayesian estimation

We recall in this section some classical algorithms used in Bayesian estimation and that are mentioned in the Section III. 3.3.

Metropolis-Hastings algorithm

The first MCMC algorithm to be known is the Metropolis-Hastings algorithm. The advantage of Metropolis-Hastings method is that it requires to know the *posterior* distribution $\mathcal{P}$ up to a constant. The Metropolis-Hastings algorithm is briefly recalled in Algorithm 2. See [61] for a more detailed description and a convergence analysis.

Algorithm 2 Metropolis-Hastings algorithm.

Start with $\boldsymbol{\theta}^{(0)}$.

For $i \in [1; N]$:

 Draw $\boldsymbol{\theta}^{new}$ from $q(\cdot|\boldsymbol{\theta}^{(i-1)})$.

 Compute ratio r with

$$r = \frac{\mathcal{P}(\boldsymbol{\theta}^{new})q(\boldsymbol{\theta}^{(i-1)}|\boldsymbol{\theta}^{new})}{\mathcal{P}(\boldsymbol{\theta}^{(i-1)})q(\boldsymbol{\theta}^{new}|\boldsymbol{\theta}^{(i-1)})}.$$

 Set $\boldsymbol{\theta}^{(i)} = \boldsymbol{\theta}^{new}$ with probability equal to $\min(r, 1)$.

Gibbs sampler algorithm

Gibbs sampler algorithm is a MCMC algorithm that can be used in the case of full known conditional *posterior* distributions $\{p(\theta_j|\theta_{i \neq j}, \mathbf{Z})\}$.

Algorithm 3 Gibbs sampler algorithm.

Start with $(\theta_2^{(0)}, \dots, \theta_p^{(0)})$.

For $i = 1, \dots, N$

- . Draw $\theta_1^{(i)}$ from $p(\theta_1 | \theta_2^{(i-1)}, \dots, \theta_p^{(i-1)}, \mathbf{Z})$
 - $\vdots$
 - . Draw $\theta_j^{(i)}$ from $p(\theta_j | \theta_1^{(i)}, \dots, \theta_{j-1}^{(i)}, \theta_{j+1}^{(i-1)}, \dots, \theta_p^{(i-1)}, \mathbf{Z})$
 - $\vdots$
 - . Draw $\theta_p^{(i)}$ from $p(\theta_p | \theta_1^{(i)}, \dots, \theta_{p-1}^{(i)}, \mathbf{Z})$
-

Genetic algorithm

Genetic algorithms fall into the much broader category of evolutionary algorithms. Such algorithms have been named this way as they attempt to simulate the biological evolution processes within the optimisation framework. Genetic algorithms are used to optimise non-smooth multivariate objective function. Each feature of the parameter θ to optimise is seen as a sequence of k binary variables which can take as values 0 or 1. The algorithm is shortly described below in Algorithm 4. For a comprehensive study, see [39].

Algorithm 4 Genetic algorithm.

Start with an initial population of parameters $Pop^{(0)} = (\theta_1^{(0)}, \dots, \theta_{n^{(0)}}^{(0)})$.

For $i = 1, \dots, N$

- . For each element in $Pop^{(i)}$, evaluate the objectif function.
 - . **Selection.**
 - Determine $\tilde{Pop}^{(i)} \subset Pop^{(i)}$, the subset of the best evaluations.
 - . **Cross-over.**
 - Determine $\frac{n^{(i)}}{2}$ pairs.
 - For each pair, select one feature.
 - For each pair, cross feature values.
 - . **Mutation.**
 - Pick one feature binary variable.
 - Change its value with a given probability.
-

Annexe C

Some general definitions and properties

Definition 3. *General definitions.*

1. We call domain of a function γ the subset on which γ is finite, i.e.

$$\text{dom } \gamma = \{x : \gamma(x) < \infty\}.$$

2. A function γ is said to be closed (lower semicontinuous) if for each $\alpha \in \mathbb{R}$, the sublevel set $\{x \in \text{dom } \gamma : \gamma(x) \leq \alpha\}$ is a closed set.

Definition 4. *Convex analysis definitions. See for example [73] for more details.*

1. A convex function γ from $\mathbb{R}^p$ to $\mathbb{R}$ satisfies

$$\forall x, y \in \mathbb{R}^p, \eta \in [0, 1], \gamma(\eta x + (1 - \eta)y) \leq \eta\gamma(x) + (1 - \eta)\gamma(y). \quad (\text{C.1})$$

γ is called strictly convex when equality in (C.1) occurs only when $x = y$.

2. We call subgradient at $x \in \mathbb{R}^p$ of a convex function γ -and write $\partial\gamma_x$ - all vectors s in $\mathbb{R}^p$ which satisfy

$$\forall y \in \mathbb{R}^p, \gamma(y) \geq \gamma(x) + \langle s, y - x \rangle.$$

3. A function γ is essentially strictly convex if γ is strictly convex on all convex subset of its subgradient domain.
4. A function γ is essentially smooth if the interior of γ domain $\mathcal{D}_\gamma$ is not empty, γ is differentiable on the interior of $\mathcal{D}_\gamma$ and its gradient $\|\nabla\gamma\|$ tends to infinity when approaching to the edge of $\mathcal{D}_\gamma$.
5. The convex conjugate ψ of a function γ is defined by

$$\psi(\tau) := \sup_{y \in \mathbb{R}^p} \{\langle y, \tau \rangle - \gamma(y)\}.$$

Proposition 21. *Convex analysis property.*

1. *If γ is closed convex, then its biconjugate, that is the conjugate of convex conjugate ψ , is γ itself.*
2. *A function γ being essentially strictly convex is equivalent to its convex conjugate being essentially smooth.*
3. *Let γ be a convex function and let its domain be not empty. If $\gamma : \mathbb{R}^p \rightarrow \mathbb{R}$ is such that*

$$\frac{\gamma(y)}{\|y\|} \xrightarrow{y \in \partial D_\gamma} +\infty$$

then its convex conjugate has full domain, that is $\mathcal{D}_\psi = \mathbb{R}^p$.

Annexe D

Multidimensional Bregman distance

D. 1 Bregman distance bounds

This section generalises some results of [20] on the bounds of Bregman distance between two functions f and f_1 in the case of $\mathbb{R}^p$ -valued functions.

Lemma 22. *Let $(U, \mathcal{B}(U), P_U)$ be a probability space and E be a convex set of functions with values in $\mathbb{R}^p$. Recall the expression of $I_\gamma(f) = \int_U \gamma(f) dP_U$ and one of the Bregman distance for f_1, f_2 having finite I_γ values*

$$B_\gamma(f_1, f_2) := \int_U \left[\gamma(f_1) - \gamma(f_2) - (\nabla \gamma(f_2))^T (f_1 - f_2) \right] dP_U.$$

$\forall \varepsilon > 0, K > 0$, it exists $\iota > 0$ such that for all $f_1, f_2 \in E$ and $C \in \mathcal{B}(U)$ for which $\min_{x \in C} (\|f_1(x)\|, \|f_2(x)\|) \leq K$, then

$$P_U(C \cap \{x : \|f_1(x) - f_2(x)\| \geq \varepsilon\}) \leq \iota B_\gamma(f_1, f_2). \quad (\text{D.1})$$

Proof. Let f_1 and f_2 be two functions in E with finite I_γ values. Let denote by b_γ the integrand of B_γ that is

$$b_\gamma(f_1(x), f_2(x)) = \gamma(f_1(x)) - \gamma(f_2(x)) - (\nabla \gamma(f_2(x)))^T (f_1(x) - f_2(x)).$$

Let $C \in \mathcal{B}(U)$ be such that $\min_{x \in C} (\|f_1(x)\|, \|f_2(x)\|) \leq K$. For $x \in C$, if $\|f_1(x) - f_2(x)\| \geq \varepsilon$, then

$$b_\gamma(f_1(x), f_2(x)) \geq \min_{x \in C} \left(\min_{\text{if } \|f_1(x)\| \leq K} b_\gamma(f_1(x), f_1(x) + \varepsilon), \min_{\text{if } \|f_2(x)\| \leq K} b_\gamma(f_2(x) + \varepsilon, f_2(x)) \right) =: \frac{1}{\iota}$$

with $\iota > 0$. Therefore, for $C \in \mathcal{B}(U)$ for which $\min_{x \in C} (\|f_1(x)\|, \|f_2(x)\|) \leq K$

$$\begin{aligned} B_\gamma(f_1, f_2) &= \int_U b_\gamma(f_1(z), f_2(z)) dP_U(z) \\ &\geq \int_{C \cap \{x: \|f_1(x) - f_2(x)\| \geq \varepsilon\}} b_\gamma(f_1(z), f_2(z)) dP_U(z) \\ &\geq \int_{C \cap \{x: \|f_1(x) - f_2(x)\| \geq \varepsilon\}} \frac{1}{\iota} dP_U(z) \\ &\geq \frac{1}{\iota} P_U(C \cap \{x: \|f_1(x) - f_2(x)\| \geq \varepsilon\}) \end{aligned}$$

which gives the result. $\square$

Lemma 23. *Let $(U, \mathcal{B}(U), P_U)$ be a probability space and E be a convex set of functions with values in $\mathbb{R}^p$. γ is an essentially strictly convex twice differentiable function on $\mathbb{R}^p$ and its Hessian matrix is strictly definite positive on $\mathbb{R}^p$. Recall the Bregman distance for f_1, f_2 having finite I_γ values*

$$B_\gamma(f_1, f_2) := \int_U [\gamma(f_1) - \gamma(f_2) - (\nabla \gamma(f_2))^T (f_1 - f_2)] dP_U.$$

$\forall K > 0$, it exists $\beta > 0$ such that for all $f_1, f_2 \in \mathbb{R}^p$ and $C \in \mathcal{B}(U)$ with $C \subset \{x: \|f_2(x)\| \leq K\}$, then

$$P_U(C \cap \{x: \|f_1(x)\| \geq L\}) \leq \frac{\beta}{L^2} B_\gamma(f_1, f_2), \quad \text{if } L \geq 2K. \quad (\text{D.2})$$

Proof. Let denote by b_γ the integrand of B_γ , that is $b_\gamma(f_1, f_2) = \gamma(f_1) - \gamma(f_2) - (\nabla \gamma(f_2))^T (f_1 - f_2)$. Let $C \in \mathcal{B}(U)$ be $C \subset \{x: \|f_2(x)\| \leq K\}$.

Let u_K and v , two vector from $\mathbb{R}^p$ such that $\|u_K\| = K$ and v belongs to the p -ball of center u_K and of radius K , denoted by $\mathcal{B}(u_K, K)$. Then, using Taylor's theorem for the decomposition of $\gamma(v)$ centered in u_K , we have

$$\gamma(v) = \gamma(u_K) + (\nabla \gamma(u_K))^T (v - u_K) + R_\gamma(v, u_K)$$

where $R_\gamma(v, u_K)$ is the remainder term of order $o(K)$.

Then for all $x \in C$

$$\begin{aligned} b_\gamma(f_1(x), f_2(x)) &\geq b_\gamma(2u_K, u_K) \\ &\geq R_\gamma(2u_K, u_K) \\ &\geq \frac{1}{2} \|u_K\|^2 \varepsilon_{\gamma, K} \\ &\geq \frac{1}{8} L^2 \varepsilon_{\gamma, K} \end{aligned}$$

where $\varepsilon_{\gamma, K}$ is the smallest eigen value of γ Hessian matrix located at any point of $\mathcal{B}(u_K, K)$. Given the assumptions on γ , $\varepsilon_{\gamma, K}$ is strictly positive.

Therefore, for all $x \in C$ for which $\|f_2(x)\| \leq K$

$$\begin{aligned}
B_\gamma(f_1, f_2) &= \int_U b_\gamma(f_1(z), f_2(z)) dP_U(z) \\
&\geq \int_{C \cap \{x: \|f_1(x)\| \geq L\}} b_\gamma(f_1(z), f_2(z)) dP_U(z) \\
&\geq \int_{C \cap \{x: \|f_1(x)\| \geq L\}} \frac{8L^2}{\varepsilon_{\gamma, K}} dP_U(z) \\
&\geq \frac{8L^2}{\varepsilon_{\gamma, K}} P_U(C \cap \{x: \|f_1(x)\| \geq L\}) \\
&\geq \frac{L^2}{\beta} P_U(C \cap \{x: \|f_1(x)\| \geq L\})
\end{aligned}$$

which gives the result. $\square$

D. 2 Cauchy result for Bregman distance minimising sequence

Lemma 24. *Let $(U, \mathcal{B}(U), P_U)$ be a probability space and E be a convex set of functions with values in $\mathbb{R}^p$. Recall the expression of $I_\gamma(f) = \int_U \gamma(f) dP_U$ and the Bregman distance for f_1, f_2 two functions having finite I_γ values*

$$B_\gamma(f_1, f_2) := \int_U [\gamma(f_1) - \gamma(f_2) - (\nabla \gamma(f_2))^T (f_1 - f_2)] dP_U.$$

Let $(f_n) \subset E$ be a I_γ -minimising sequence, then (f_n) is a Cauchy sequence in probability P_U , meaning that

$$\lim_{n, m \rightarrow \infty} P_U(\{x: \|f_n(x) - f_m(x)\| > \varepsilon\}) = 0 \quad \forall \varepsilon > 0. \quad (\text{D.3})$$

Proof. Let $\eta > 0$ and K such that for $f \in \mathbb{R}^p$

$$P_U(\{x: \|f(x)\| > K\}) < \eta.$$

By applying Lemma 23 with $f_1 = \frac{f_n + f_m}{2}$, $f_2 = f$ and $C = \{x: \|f(x)\| \leq K\}$, we have

$$P_U\left(\{x: \|f(x)\| \leq K\} \cap \left\{x: \left\|\frac{f_n(x) + f_m(x)}{2}\right\| \geq L\right\}\right) \leq \frac{\beta}{L^2} B_\gamma\left(\frac{f_n + f_m}{2}, f\right),$$

if $L \geq 2K$. Choosing β and L such that

$$\frac{\beta}{L^2} B_\gamma\left(\frac{f_n + f_m}{2}, f\right) \leq \eta,$$

it follows that

$$P_U\left(\left\{x: \left\|\frac{f_n(x) + f_m(x)}{2}\right\| \geq L\right\}\right) \leq 2\eta.$$

Now applying Lemma 22 with $f_1 = f_m$, $f_2 = \frac{f_n + f_m}{2}$ and $C = \left\{x: \left\|\frac{f_n(x) + f_m(x)}{2}\right\| < L\right\}$, there exists ι such that

$$\begin{aligned}
P_U\left(\left\{x: \left\|\frac{f_n(x) + f_m(x)}{2}\right\| < L\right\} \cap \left\{x: \left\|f_m(x) - \frac{f_n(x) + f_m(x)}{2}\right\| \geq \varepsilon\right\}\right) \\
\leq \iota B_\gamma\left(f_m, \frac{f_n + f_m}{2}\right)
\end{aligned}$$

meaning that

$$\begin{aligned} P_U \left(\left\{ x : \left\| f_m(x) - \frac{f_n(x) + f_m(x)}{2} \right\| \geq \varepsilon \right\} \right) &\leq 2\eta + \iota B_\gamma \left(f_m, \frac{f_n + f_m}{2} \right) \\ P_U (\{x : \|f_m(x) - f_n(x)\| \geq 2\varepsilon\}) &\leq 2\eta + \iota B_\gamma \left(f_m, \frac{f_n + f_m}{2} \right) \end{aligned}$$

By taking η as small as possible, when n and m go to infinity, it ends up that

$$\lim_{n,m \rightarrow \infty} P_U (\{x : \|f_m(x) - f_n(x)\| \geq 2\varepsilon\}) = \lim_{n,m \rightarrow \infty} \iota B_\gamma \left(f_m, \frac{f_n + f_m}{2} \right).$$

Let $(f_n) \subset E$ be a I_γ -minimising sequence with f_n having finite I_γ values. By the positivity of B_γ , we have

$$\begin{aligned} I_\gamma \left(\frac{f_n + f_m}{2} \right) &\leq \frac{1}{2} (I_\gamma(f_n) + I_\gamma(f_m)) \\ &\leq \max_k I_\gamma(f_k). \end{aligned}$$

By the convexity of E , $\frac{f_n + f_m}{2}$ belongs to E and therefore $I_\gamma \left(\frac{f_n + f_m}{2} \right) \geq I_\gamma(E)$.

Using the following identity when n, m tend to infinity

$$I_\gamma(f_n) + I_\gamma(f_m) = 2I_\gamma \left(\frac{f_n + f_m}{2} \right) + B_\gamma \left(f_n, \frac{f_n + f_m}{2} \right) + B_\gamma \left(f_m, \frac{f_n + f_m}{2} \right), \quad (\text{D.4})$$

it implies that

$$\begin{aligned} \lim_{n,m \rightarrow \infty} B_\gamma \left(f_m, \frac{f_n + f_m}{2} \right) &= 0 \\ \lim_{n,m \rightarrow \infty} B_\gamma \left(f_n, \frac{f_n + f_m}{2} \right) &= 0. \end{aligned}$$

which proves that (f_n) is a Cauchy sequence in probability P_U . □

Bibliographie

- [1] The OpenCalphad Repository. <http://github.com/sundmanbo/opencalphad>, 2018.
- [2] AKAIKE, H. Information theory and an extension of the maximum likelihood principle. In *Selected Papers of Hirotugu Akaike*, E. Parzen, K. Tanabe, and G. Kitagawa, Eds., Springer Series in Statistics. Springer, 1998, pp. 199–213.
- [3] ARROYO, M., Ed. *International Tables for Crystallography*, 6 ed., vol. A. Wiley, 2016.
- [4] AZAÏS, J.-M., AND BARDET, J.-M. *Le Modèle Linéaire par l'exemple Régression, Analyse de la Variance et Plans d'Expériences Illustrations numériques avec les logiciels R, SAS et Splus*. Dunod, 2006.
- [5] BACHRATA, A., MARIE, N., BERTRAND, F., AND DROIN, J. B. Improvement of model for SIMMER code for SFR corium relocation studies. *International Journal of Physical and Mathematical Sciences* 8, 3 (2014), 6.
- [6] BAGNOUD, P., AND FESCHOTTE, P. Les systèmes binaires magnésium-cuivre et magnésium-nickel en particulier non-stoechiométrie des phases de laves MgCu_2 et MgNi_2 . *Zeitschrift fuer Metallkunde* (1978).
- [7] BARLOW, W. Über die Geometrischen Eigenschaften homogener starrer Strukturen und ihre Anwendung auf Krystalle. *Zeitschrift für Krystallographie und Mineralogie* 23 (1894), 1–63.
- [8] BARNDORFF-NIELSEN, O. *Information and Exponential Families : In Statistical Theory*. John Wiley & Sons, 1978.
- [9] BATALIN, G., SUDAVTSOVA, V., AND MIKHAILOVSKAYA, M. Thermodynamic properties of liquid alloys of the cu-mg systems. *Izv. V. U. Z. Tsvetn. Metall.* 2 (1987).
- [10] BORCHARDT-OTT, W. *Crystallography*. Springer Berlin Heidelberg, 1995.
- [11] BORGONOVO, E., AND PLISCHKE, E. Sensitivity analysis : A review of recent advances. *European Journal of Operational Research* 248, 3 (2016), 869–887.
- [12] BORWEIN, J., AND LEWIS, A. Duality relationships for entropy-like minimization problems. *SIAM Journal on Control and Optimization* (1991).
- [13] BORWEIN, J., AND LEWIS, A. Partially-finite programming in L_1 and the existence of maximum entropy estimates. *SIAM Journal on Optimization* (1993).
- [14] BREGMAN, L. M. The relaxation method of finding the common point of convex sets and its application to the solution of problems in convex programming. *USSR Computational Mathematics and Mathematical Physics* (1967).
- [15] BRONIATOWSKI, M., AND KEZIOU, A. Minimization of divergences on sets of signed measures. *Studia Scientiarum Mathematicarum Hungarica* (2006).

- [16] CALLISTER, W. D., AND RETHWISCH, D. G. *Materials Science and Engineering 7th Edition*. Wiley, 2001.
- [17] CHATTERJEE, N., MILLER, K., AND OLBRICHT, W. Bayes estimation : A novel approach to derivation of internally consistent thermodynamic data for minerals, their uncertainties, and correlations. part II : Application. *Physics and Chemistry of Minerals* (1994).
- [18] CHATTERJEE, N. D., KRÜGER, R., HALLER, G., AND OLBRICHT, W. The bayesian approach to an internally consistent thermodynamic database : theory, database, and generation of phase diagrams. *Contributions to Mineralogy and Petrology* (1998).
- [19] CSISZÁR, I. Sanov property, generalized i-projection and a conditional limit theorem. *The Annals of Probability* (1984).
- [20] CSISZÁR, I. Generalized projections for non-negative functions. *Acta Mathematica Hungarica* (1995).
- [21] CSISZÁR, I., GAMBOA, F., AND GASSIAT, E. MEM pixel correlated solutions for generalized moment and interpolation problems. *IEEE Transactions on Information Theory* (1998).
- [22] CSISZÁR, I. I-divergence geometry of probability distributions and minimization problems. *Ann. Probability* (1975).
- [23] CSISZÁR, I. Why least squares and maximum entropy? an axiomatic approach to inference for linear inverse problems. *The Annals of Statistics* (1991).
- [24] DACUNHA-CASTELLE, D., AND GAMBOA, F. Maximum d'entropie et problème des moments. *Annales de l'I.H.P. Probabilités et statistiques* (1990).
- [25] D'AGOSTINO, R. B., AND STEPHENS, M. A., Eds. *Goodness-of-fit techniques*, dekker ; 1st edition ed. CRP press, 1986.
- [26] DICKHAUS, T. *Theory of Nonparametric Tests*. Springer, 2018.
- [27] DINSDALE, A. T. SGTE data for pure elements. *Calphad* 15, 4 (1991), 317–425.
- [28] DROIN, J.-B. *Modélisation d'un transitoire de perte de débit primaire non protégé dans un RNR-Na*. PhD thesis, Université Grenoble Alpes, 2016.
- [29] DUONG, T. C., HACKENBERG, R. E., LANDA, A., HONARMANDI, P., TALAPATRA, A., VOLZ, H. M., LLOBET, A., SMITH, A. I., KING, G., BAJAJ, S., RUBAN, A., VITOS, L., TURCHI, P. E. A., AND ARRÓYAVE, R. Revisiting thermodynamics and kinetic diffusivities of uranium–niobium with bayesian uncertainty analysis. *Calphad* 55 (2016), 219–230.
- [30] EPP, J. *Materials Characterization Using Nondestructive Evaluation (NDE) Methods*. Woodhead Publishing, 2016.
- [31] EREMENKO, V., LUKASHEN, G., AND POLOTSKA, R. Some thermodynamic properties of magnesium-copper compounds. *RUSSIAN METALLURGY-METALLY-USSR* 1 (1968).
- [32] FEUFEL, H., AND SOMMER, F. Thermodynamic investigations of binary liquid and solid cu-mg and mg-ni alloys and ternary liquid cu-mg-ni alloys. *Journal of Alloys and Compounds* 224, 1 (1995-06), 42–54.
- [33] FISHER, R. A. *Statistical methods and scientific inference*. Statistical methods and scientific inference. Hafner Publishing Co., 1956.

- [34] GAMBOA, F. *Maximum d'entropie sur la moyenne*. PhD thesis, Université d'Orsay, 1989.
- [35] GAMBOA, F., AND GASSIAT, E. Maximum d'entropie et problème des moments : cas multidimensionnel. *Probability and Mathematical Statistics* (1991).
- [36] GAMBOA, F., AND GASSIAT, E. Bayesian methods and maximum entropy for ill-posed inverse problems. *The Annals of Statistics* (1997).
- [37] GAMBOA, F., GUENEAU, C., KLEIN, T., AND LAWRENCE, E. Maximum entropy on the mean approach to solve generalized inverse problems with an application in computational thermodynamics. <https://hal.archives-ouvertes.fr/hal-02903231>, 2020.
- [38] GARG, S. P., BHATT, Y. J., AND SUNDARAM, C. V. Thermodynamic study of liquid cu-mg alloys by vapor pressure measurements. *Metallurgical Transactions* 4, 1 (1973-01), 283–289.
- [39] GOLDBERG, D. E. *Genetic Algorithms in Search, Optimization and Machine Learning*, 1st ed. Addison Wesley, 1989.
- [40] GRIME, G., AND MORRIS-JONES, W. CXXVI. an x-ray investigation of the copper-magnesium alloys. *The London, Edinburgh, and Dublin Philosophical Magazine and Journal of Science* 7, 47 (1929-06), 1113–1134.
- [41] GUIDEZ, J., AND BONIN, B. *Réacteurs nucléaires à caloporteur sodium*. E-den, Une monographie de la Direction de l'énergie nucléaire. CEA Saclay ; Groupe Moniteur, 2014.
- [42] HASTIE, T., TIBSHIRANI, R., AND FRIEDMAN, J. *The Elements of Statistical Learning : Data Mining, Inference, and Prediction, Second Edition*, 2 ed. Springer Series in Statistics. Springer-Verlag, 2009.
- [43] HOLLAS, J. M. *Spectroscopie*. Sciences Sup, Dunod, 2003.
- [44] HONARMANDI, P., DUONG, T. C., GHOREISHI, S. F., ALLAIRE, D., AND ARROYAVE, R. Bayesian uncertainty quantification and information fusion in CALPHAD-based thermodynamic modeling. *arXiv :1806.05769 [cond-mat, stat]* (2018).
- [45] HULTGREN, R., DESAI, P. D., HAWKINS, D. T., GLEISER, M., AND KELLEY, K. K. Selected values of the thermodynamic properties of binary alloys. Tech. rep., National Standard Reference Data System, 1973.
- [46] JANSSON, B. *Computer operated methods for equilibrium calculations and evaluation of thermochemical model parameters*. PhD thesis, Royal Institute of Technology, 1984.
- [47] JIANG, C., WOLVERTON, C., SOFO, J., CHEN, L.-Q., AND LIU, Z.-K. First-principles study of binary bcc alloys using special quasirandom structures. *Physical Review B* 69, 21 (2004), 214202.
- [48] JONES, W. R. D. The copper-magnesium alloys. part 4. the equilibrium diagram. In *Annual Autumn Meeting* (1931).
- [49] JUNEJA, J. M., IYENGAR, G. N. K., AND ABRAHAM, K. Thermodynamic properties of liquid (magnesium + copper) alloys by vapour-pressure measurements made by a boiling-temperature method. *Journal of Chemical Thermodynamics* (1986), 11.
- [50] KALMAN, R. E. A new approach to linear filtering and prediction problems. *Journal of Basic Engineering* (1960). Publisher : American Society of Mechanical Engineers Digital Collection.

- [51] KEMPERMAN, J. H. B. Moment problems with convexity conditions. In *Optimizing Methods in Statistics* (1971), J. S. Rustagi, Ed., Academic Press.
- [52] KING, R., AND KLEPPA, O. A thermochemical study of some selected laves phases. *Acta Metallurgica* 12 (1964).
- [53] KÖNIGSBERGER, E. Improvement of excess parameters from thermodynamic and phase diagram data by a sequential bayes algorithm. *Calphad* (1991).
- [54] LEONARD, C. Minimizers of energy functionals. *Acta Mathematica Hungarica* (2001).
- [55] LEONARD, C. Minimizers of energy functional under not very integrable constraints. *Journal of Convex Analysis* (2003).
- [56] LIESE, F., AND VAJDA, I. *Convex statistical distances*. Teubner, 1987.
- [57] LIU, Z.-K. First-principles calculations and CALPHAD modeling of thermodynamics. *Journal of Phase Equilibria and Diffusion* 30 (2009).
- [58] LUKAS, H. L., FRIES, S. G., AND SUNDMAN, B. *Computational thermodynamics : the CAL-PHAD method*. Cambridge University Press, 2007. OCLC : 663969016.
- [59] MALAKHOV, D. V. Confidence intervals of calculated phase boundaries. *Calphad* (1997).
- [60] MARREL, A., MARIE, N., AND DE LOZZO, M. Advanced surrogate model and sensitivity analysis methods for sodium fast reactor accident assessment. *Reliability Engineering & System Safety* 138 (2015), 232–241.
- [61] METROPOLIS, N., ROSENBLUTH, A. W., ROSENBLUTH, M. N., TELLER, A. H., AND TELLER, E. Equation of state calculations by fast computing machines. *The Journal of Chemical Physics* (1953).
- [62] NAVAZA, J. On the maximum-entropy estimate of the electron density function. *Acta Crystallographica Section A* (1985).
- [63] NIKODYM, O. Sur une généralisation des intégrales de m. j. radon. *Fundamenta Mathematicae* 15, 1 (1930), 131–179.
- [64] OKAMOTO, H., SCHLESINGER, M. E., AND MUELLER, E. M., Eds. *ASM handbook, Volume 3 Alloy Phase Diagrams*, 10th edition ed. ASM International, 2016.
- [65] OLBRICHT, W., CHATTERJEE, N., AND MILLER, K. Bayes estimation : A novel approach to derivation of internally consistent thermodynamic data for minerals, their uncertainties, and correlations. part i : Theory. *Physics and Chemistry of Minerals* (1994).
- [66] OTIS, R. A., AND LIU, Z.-K. High-throughput thermodynamic modeling and uncertainty quantification for ICME. *JOM* 69 (2017).
- [67] PAULSON, N. H., BOCKLUND, B. J., OTIS, R. A., LIU, Z.-K., AND STAN, M. Quantified uncertainty in thermodynamic modeling for materials design. *arXiv :1901.10510 [cond-mat, physics :physics]* (2019).
- [68] PAULSON, N. H., JENNINGS, E., AND STAN, M. Bayesian strategies for uncertainty quantification of the thermodynamic properties of materials. *arXiv :1809.07365 [cond-mat]* (2018).
- [69] RADON, J. *Theorie und anwendungen der absolut additiven mengenfunktionen ...* Hölder, 1913. OCLC : 39514488.

- [70] REARDON, B. Fuzzy logic versus niched pareto multiobjective genetic algorithm optimization. *Modelling and Simulation in Materials Science and Engineering* (1998).
- [71] REARDON, B. J. Inversion of micromechanical powder consolidation and sintering models using bayesian inference and genetic algorithms. *Modelling and Simulation in Materials Science and Engineering*, 6 (1999).
- [72] ROBERT, C. P. *The Bayesian Choice : A Decision-Theoretic Motivation*. Springer Texts in Statistics. Springer, 1994.
- [73] ROCKAFELLAR, R. T. *Convex analysis*. Princeton university press, 1970.
- [74] ROCKAFELLAR, R. T., AND WETS, R. J.-B. *Variational Analysis*. Grundlehren der mathematischen Wissenschaften. Springer-Verlag, 1998.
- [75] ROGELBERG, I. Cu-mg phase diagram. *Tr. Gos. Nauchn.-Issled.* 16 (1957), 82–89.
- [76] SAHMEN, R. Über die legierungen des kupfers mit kobalt, eisen, mangan und magnesium. *Zeitschrift für anorganische Chemie* 57, 1 (1908-01-25), 1–33.
- [77] SALTELLI, A., RATTO, M., TARANTOLA, S., AND CAMPOLONGO, F. Sensitivity analysis for chemical models. *Chemical Reviews* 105, 7 (2005), 2811–2828. Publisher : American Chemical Society.
- [78] SALTELLI, A., RATTO, M., TARANTOLA, S., AND CAMPOLONGO, F. Update 1 of : Sensitivity analysis for chemical models. *Chemical Reviews* 112, 5 (2012), PR1–PR21. Publisher : American Chemical Society.
- [79] SANTNER, T. J., WILLIAMS, B. J., AND NOTZ, W. I. *The Design and Analysis of Computer Experiments*. Springer Series in Statistics. Springer-Verlag, 2003.
- [80] SAPORTA, G. *Probabilités, analyse des données et statistique*. TECHNIP, 2011.
- [81] SCHMAHL, N. G., AND SIEBEN, P. Vapour pressures CF magnesium in its binary alloys with copper, nickel and lead and their thermodynamical evaluation. *Symposium on Physical Chemistry of Metallic Solutions and Intermetallic Compounds*. (1959), 17.
- [82] SEDERMAN, V. XXIX. the cu₂mg phase in the copper-magnesium system. *The London, Edinburgh, and Dublin Philosophical Magazine and Journal of Science* 18, 118 (1934), 343–352.
- [83] SHIN, D., VAN DE WALLE, A., WANG, Y., AND LIU, Z.-K. First-principles study of ternary fcc solution phases from special quasirandom structures. *Physical Review B* 76, 14 (2007), 144204.
- [84] SHORE, J., AND JOHNSON, R. Axiomatic derivation of the principle of maximum entropy and the principle of minimum cross-entropy. *IEEE Transactions on Information Theory* (1980).
- [85] SMITH, J. F., AND CHRISTIAN, J. L. Thermodynamics of formation of copper-magnesium and nickel-magnesium compounds from vapor pressure measurements. *Acta Metallurgica* 8 (1960), 7.
- [86] SOMMER, F., LEE, J. J., AND PREDEL, B. Calorimetric investigations of liquid alkaline earth metal alloys. *Ber. Bunsenges. Phys. Chem.* (1983), 7.
- [87] SPENCER, P. J. A brief history of CALPHAD. *Calphad* 32 (2008).
- [88] STAN, M., AND REARDON, B. A bayesian approach to evaluating the uncertainty of thermodynamic data and phase diagrams. *Calphad* 27, 3 (2003), 319–323.

- [89] SUNDMAN, B., LU, X.-G., AND OHTANI, H. The implementation of an algorithm to calculate thermodynamic equilibria for multi-component systems with non-ideal phases in a free software. *Computational Materials Science* (2015).
- [90] TAYLOR, R. H., CURTAROLO, S., AND HART, G. L. Guiding the experimental discovery of magnesium alloys. *Physical Review B* 84, 8 (2011), 084101.
- [91] URAZOV, G. Alloys of copper and magnesium. *Zhurnal Russkogo Fiziko-Khimicheskogo Obshchestva* 39 (1907), 1556–1581.
- [92] VAPNIK, V. *The Nature of Statistical Learning Theory*, 2 ed. Information Science and Statistics. Springer-Verlag, 2000.
- [93] VIOT, L. *Couplage et synchronisation de modèles dans un code scénario d’accidents graves dans les réacteurs nucléaires*. phdthesis, Université Paris-Saclay, 2018.
- [94] WAND, M. P., AND JONES, M. C. *Kernel Smoothing*. CRC Press, 1994.
- [95] ZHANG, F., Ed. *The Schur Complement and Its Applications*. Numerical Methods and Algorithms. Springer US, 2005.
- [96] ZHOU, S., WANG, Y., SHI, F. G., SOMMER, F., CHEN, L.-Q., LIU, Z.-K., AND NAPOLITANO, R. E. Modeling of thermodynamic properties and phase equilibria for the cu-mg binary system. *Journal of Phase Equilibria and Diffusion* 28, 2 (2007-04-01), 158–166.
- [97] ZOMORODPOOSH, S., BOCKLUND, B., OBAIED, A., OTIS, R., LIU, Z. K., AND ROSLYAKOVA, I. Statistical approach for automated weighting of datasets : Application to heat capacity data. *Calphad* 71 (2020), 101994.