

UNIVERSITÉ D'ORLÉANS
*ÉCOLE DOCTORALE Mathématiques, Informatiques,
Physique Théorique et Ingénierie des Systèmes*

Institut Denis Poisson & BRGM

THÈSE présentée par :

Titouan SIMONNET

Soutenue le **06 novembre 2024**

pour obtenir le grade de : **Docteur de l'Université d'Orléans**

Discipline/ Spécialité : **Mathématiques, Mathématiques Appliquées**

**Apprentissage et réseaux de neurones en
tomographie par diffraction de rayons X.**

Application à l'identification minéralogique

THÈSE dirigée par :

M. GALERNE Bruno

Professeur, Université d'Orléans

RAPPORTEURS :

M. FERRAGE Eric

Directeur de recherches, Université de Poitiers

M. MOUSSAOUI Saïd

Professeur, Ecole Centrale de Nantes

JURY :

Mme CLAUSEL Marianne

Examinatrice, Professeure, Université de Lorraine

M. DEMORTIERE Arnaud

Examinateur, Chargé de recherches, Université de Picardie

Mme FALL Mame Diarra

Co-encadrante, Maître de conférences, Université d'Orléans

M. FERRAGE Eric

Rapporteur, Directeur de recherches, Université de Poitiers

M. GALERNE Bruno

Directeur de thèse, Professeur, Université d'Orléans

M. GRANGEON Sylvain

Co-encadrant, Docteur, BRGM

M. HARBA Rachid

Co-encadrant, Professeur émérite, Université d'Orléans

M. MOUSSAOUI Saïd

Président du jury, Rapporteur, Professeur, Ecole Centrale de Nantes

INVITE :

M. CLARET Francis

Co-encadrant, Docteur, BRGM

Remerciements

Ces remerciements sont sans doute parmi les mots les plus compliqués à écrire puisque ce sont les plus personnels. Mais il est évident que cette thèse n'aurait pas été possible tout seul. Je tiens donc à remercier l'ensemble des personnes qui m'ont aidé durant ces trois années de près ou de loin, en espérant n'oublier personne. J'en oublierais probablement et je m'excuse par avance pour celles et ceux que je ne cite pas.

Je commencerais par Bruno Galerne qui a construit ce projet et qui m'a dirigé pour ces travaux de thèse. Je tiens à te remercier pour ton encadrement, pour tes conseils, et ton expertise sur le *deep learning* qui m'ont permis d'avancer. Nos échanges m'ont beaucoup apporté et appris durant ces trois années.

Ensuite, je voudrais remercier les personnes qui ont co-encadré cette thèse et qui ont énormément contribué à ces travaux : Francis Claret, Diarra Fall, Rachid Harba et Sylvain Grangeon. Francis, merci pour ton investissement tout au long de la thèse. Tu as pris le temps chaque fois que cela a été nécessaire, ce qui m'a toujours permis d'avancer. Diarra, ta motivation et tes conseils durant ces trois années m'ont permis de rester investi tout au long de ce projet. Merci pour ta contribution et ton aide pour la rédaction des productions scientifiques. Rachid, merci d'avoir participé à l'encadrement, nos échanges lors de nos réunions m'ont toujours permis de clarifier mes idées et d'améliorer mes travaux. Sylvain, enfin, merci pour ton investissement, ta disponibilité, ton intérêt pour mes travaux et pour ces nombreuses discussions qui m'ont permis d'avancer. Tu as su m'apprendre ce dont j'avais besoin pour maîtriser la partie physique sur la diffraction de rayons X, et tu as toujours su répondre à mes questions sur le sujet. Ta curiosité pour le domaine mathématique de la thèse nous a aussi permis d'avoir des échanges constructifs et productifs.

Je tiens aussi à remercier Eric Ferrage et Saïd Moussaoui d'avoir rapporté cette thèse. Merci d'avoir pris le temps de lire attentivement ce manuscrit. Vos remarques m'ont permis d'éclaircir certains points et d'améliorer la qualité de mon manuscrit. Merci également à Marianne Clausel et Arnaud Demortière d'avoir accepté de faire partie du jury.

En dépendant de plusieurs laboratoires, j'ai eu la chance de diversifier les échanges, les points de vue et les approches scientifiques et méthodologiques. Je voudrais donc remercier l'ensemble des membres de l'Institut Denis Poisson et de l'unité MG2 du BRGM pour les échanges et les discussions que j'ai pu avoir ces trois dernières années. Je voudrais remercier tout particulièrement Christophe Chiaberge de m'avoir aidé à optimiser mes algorithmes de simulations, Nicolas Maubec pour les mesures des échantillons au laboratoire et à Romain Théron qui m'a aidé pour l'utilisation des serveurs de calculs de l'IDP. Merci également à Anne et Marie à l'IDP, ainsi qu'à Prescillia au BRGM pour leur aide sur les tâches administratives.

Je remercie également le projet ANR AI.iO qui a participé au financement de ma thèse.

Je complétera en remerciant les doctorants et post-doctorants (Merci Charles) de ces

laboratoires. A l'IDP, ceux que j'ai pu côtoyer : Maxime, Simon, Emile, Gaëtan, Sonia, Léo, Bruno, Emma, Paguiel, Cissé, ... Et j'en oublie sûrement. Au BRGM, les doctorants de l'unité : Anne-Cécile, Djouhar, Lucile, et surtout Raphaël avec qui j'ai eu la chance de partager le bureau. Une salutation particulière pour docteur Charles : *mens sana in corpore sano*.

Enfin, je tiens à remercier mes parents et ma petite sœur pour leur soutien. Un mot également pour mes amis depuis le lycée : Dmytro, Mehdi et Pierre et à ceux du club d'athlétisme : Aymeric et Serge notamment. Merci de m'avoir accompagné, supporté et encouragé toutes ces années. Pour finir, un merci particulier à Océane, tu m'aura beaucoup aidée dans la rédaction de ce manuscrit et soutenue sur la fin de ces trois années.

Table des matières

1	Introduction	1
1.1	Contexte général	1
1.2	Structure du manuscrit	4
1.3	Résumé des contributions	5
1.3.1	Simulation de diffractogrammes	5
1.3.2	Fonction de perte pour l'inférence de proportion	6
1.3.3	Applications aux diffractogrammes expérimentaux	7
1.3.4	L'approche des Transformers	8
1.4	Publications	8
2	La diffraction de rayons X et simulation de diffractogrammes	9
2.1	Les matériaux analysés	9
2.1.1	La structure des cristaux	9
2.2	La diffraction des rayons X	12
2.2.1	Acquisition de données en laboratoire	13
2.2.2	Acquisition au synchrotron	14
2.3	Calcul théorique de diffractogrammes	15
2.3.1	Facteur de diffusion atomique	15
2.3.2	Facteur de structure	17
2.3.3	Fonction d'interférence	17
2.4	Simulation numérique de diffractogrammes	19
2.4.1	Première approche : un calcul simplifié	19
2.4.2	Seconde approche : un simulation précise et complète	20
2.4.3	Amélioration des simulations	22
2.4.4	Résultats des simulations et constitution de bases de données	23
3	Les réseaux de neurones pour l'identification de phases minérales	27
3.1	Les réseaux de neurones artificiels	27
3.1.1	Le modèle	28
3.1.2	Les données	29
3.1.3	La fonction de coût	30
3.2	Identification de phases	32
4	Inférence de proportions	37
4.1	Une modélisation de Dirichlet pour l'inférence de proportions	37
4.1.1	Formulation du problème	37
4.1.2	La distribution de Dirichlet	37
4.1.3	Modélisation	38
4.1.4	Comparaison avec d'autres fonctions de pertes	40

4.1.5	Métriques d'évaluation pour l'inférence de proportions	41
4.2	Quantification de phases minérales sur des simulations	42
4.3	Imagerie Hyperspectrale	47
4.3.1	Jasper Ridge	48
4.3.2	Urban	51
4.4	Images MNIST	57
4.5	Conclusion	59
5	DRX : Analyse d'un jeu de données de laboratoire	61
5.1	Les données	61
5.1.1	Rietveld	62
5.1.2	Une base de test pour l'approche des réseaux de neurones	63
5.2	Méthode proposée	65
5.2.1	Comparaison de deux simulations	65
5.2.2	Génération des bases de données	67
5.3	Résultats sur les données expérimentales	69
5.4	Conclusion et discussion	70
6	DRX : Analyse d'un jeu de données tomographie	75
6.1	Description des données	75
6.1.1	Traitement des diffractogrammes	77
6.2	Analyse des données grâce aux réseaux de neurones	79
6.2.1	Construction de la base de données	79
6.2.2	Entraînement du réseau de neurones	81
6.3	Résultats	81
6.3.1	Un modèle à 11 classes	87
6.3.2	Reconstitution	88
6.4	Conclusion et discussion	89
7	Une nouvelle approche : les Transformers	91
7.1	Introduction	91
7.2	L'architecture des Transformers	92
7.2.1	Les réseaux de neurones récurrents	92
7.2.2	Attention	93
7.2.3	Encodage du signal et vecteurs de positions	93
7.2.4	Les Transformers pour l'analyse des signaux de diffraction	95
7.2.5	Des outils de visualisation	96
7.3	Classification de diffractogrammes simple phase	97
7.3.1	Entraînement et résultats de la classification	97
7.3.2	Visualisation des résultats	98
7.4	Inférence de proportions : Quantification des phases minérales des données de laboratoire	101
7.4.1	Entraînement et résultats pour la quantification des phases minérales	101
7.4.2	Visualisation des résultats	104
7.5	Conclusion	106

8	Conclusion et perspectives	107
8.1	Conclusion	107
8.2	Perspectives	108
8.2.1	Amélioration des résultats pour la quantification de phases minérales	108
8.2.2	Autres pistes exploratoires	108
A	Tableaux des amplitudes de variations pour les simulations	123

Chapitre 1

Introduction

1.1 Contexte général

Dans le cadre de cette thèse, nous nous sommes intéressés aux matériaux et notamment à leur composition. Pour ces travaux, ils sont définis comme des substances correspondant à un assemblage de différents composants, que l'on nomme phase. Ces matériaux sont omniprésents dans nos sociétés (minerais, construction, équipements, etc) et peuvent être d'origine naturelle, comme les sols ou le diamant, ou d'origine artificielle, comme les batteries, ou l'acier. Qu'ils soient de provenance naturelle ou synthétique les matériaux solides sont complexes, et constitués de nombreuses phases, notamment minérales, qui possèdent leurs propriétés propres. Elles confèrent aux minéraux un rôle clef dans les enjeux liés aux changements climatiques et aux technologies vertes (VIDAL *et al.* 2013), à l'assainissement de la pollution (GRANGEON *et al.* 2020) ou encore à la géothermie (BIRD & SPIELER 2004). La compréhension de ces propriétés et de la réactivité chimique ou physique associée à ces minéraux sont donc des enjeux majeurs. Leur évolution dans le temps ou en réponse à une contrainte (chimique ou physique) le sont également. Pour cela, une connaissance approfondie de la nature et de la quantité des phases minérales est nécessaire pour comprendre, modéliser et prédire les propriétés du solide. Par exemple, les propriétés d'absorption chimique ou les propriétés mécaniques des matériaux dépendent de la nature et de l'abondance des phases constituantes (S. V. KRIVOVICHEV *et al.* 2022 ; NDLOVU *et al.* 2011 ; WOOD & STRENS 1979). Ces phases sont généralement des solides sous forme cristalline, correspondant à des assemblages ordonnés d'atomes se répétant périodiquement dans les directions de l'espace.

Parmi les différentes techniques pour étudier les solides, la diffraction par rayons X (DRX) est une méthode de choix pour la détermination de la nature des phases minérales composant le solide. Des signaux, appelés diffractogrammes, sont les résultats de cette analyse et peuvent être décrits de manière schématique comme « l'empreinte digitale d'une phase cristalline ». Ils comportent de nombreux pics de diffraction d'intensité variable. Des variations sont observables à la fois sur les intensités relatives des pics d'un même signal, qui dépendent directement de l'angle de diffraction. Elles sont aussi observables sur les intensités absolues d'un diffractogramme, où l'un des principaux facteurs est la composition atomique du solide analysé. C'est finalement l'étude de ces signaux qui mène à la composition minérale du solide, permettant à la fois d'identifier et de quantifier les minéraux présents.

La méthode couramment utilisée pour déterminer la proportion des différentes phases est l'affinement Rietveld (DOEBELIN & KLEEGERG 2015 ; RIETVELD 1969). Cette mé-

thode est basée sur un algorithme qui calcule un modèle, i.e un diffractogramme synthétique. Pour minimiser l’erreur avec le signal original, et proposer le meilleur modèle possible, de nombreux paramètres sont affinés. Parmi ces paramètres, nous retrouvons notamment les abondances de chaque phase minérale, qui permettent de déterminer leur quantité dans le solide.

En général, l’affinement Rietveld nécessite un travail préalable d’identification des phases minérales. Les composants (en termes de phases minérales) du matériau doivent donc être connus avant de faire l’affinement. Cette tâche doit être accomplie en amont de l’affinement et peut requérir un temps conséquent. Mais, dès lors que l’identification des phases est bien réalisée, la méthode de Rietveld permet de quantifier les proportions des différentes phases constituant le solide. L’affinement est rapide (de l’ordre de quelques secondes par échantillon) et propose une quantification très précise, complétée par des mesures d’incertitudes. Néanmoins, dans certains cas, le nombre de données est très important, rendant le travail préalable d’identification difficile voire impossible car trop chronophage.

C’est le cas par exemple pour des jeux de données issus de la tomographie par diffraction de rayons X (DRX-CT). C’est une méthode d’analyse non destructive, pour caractériser *in-situ* (CLARET *et al.* 2018) un échantillon et sa réactivité (JACQUES *et al.* 2013; JENSEN *et al.* 2015). Cette technique de DRX-CT permet de ne pas se limiter à l’analyse d’échantillons sous forme de dépôt orienté (comme pour les argiles) ou de poudre (BISH & POST 1990), comme pour la DRX couramment utilisée en laboratoire. Les analyses peuvent donc être réalisées sur des échantillons de taille plus importante, qui conservent une forme naturelle. Elles permettent d’enregistrer des données en trois dimensions sous forme de voxels à l’échelle nanométrique (CLARET *et al.* 2018). Cela offre la capacité d’étudier un solide dans son ensemble, ainsi que l’organisation tridimensionnelle de ses minéraux. Sur ce type de données (CLARET *et al.* 2018), composée d’une dizaine de phases minérales, l’approche de Rietveld n’est plus possible, car elle implique d’identifier les phases dans chaque voxel avant de lancer la procédure de quantification.

Ces jeux de données nécessitent donc une analyse automatique des signaux, en utilisant par exemple les méthodes d’apprentissage automatique (i.e *machine learning*) décrites par MITCHELL & MITCHELL (1997). Une alternative déjà très populaire consiste à utiliser les réseaux de neurones (HAYKIN 2009) dont le développement est important ces dernières années. Différentes raisons expliquent l’essor des réseaux de neurones, à commencer par l’émergence de l’apprentissage profond (i.e *deep learning*) (GOODFELLOW *et al.* 2016). Le développement des outils informatiques et des GPU permettant des calculs toujours plus rapides est un autre facteur. De plus, les performances des réseaux de neurones surpassent des approches plus traditionnelles dans de nombreux domaines. C’est par exemple le cas des *Support Vector Machine* (SVM) ou de la régression logistique qui sont moins performants pour des tâches de classification. Les réseaux de neurones sont donc très largement utilisés pour différentes applications, comme la classification d’images (HE *et al.* 2016) ou le traitement du signal (GAMBOA 2017).

Les réseaux de neurones ont également été utilisés pour analyser des diffractogrammes. Déjà en 1999, GRIFFEN (1999) proposait d’utiliser un réseau de neurones pour quantifier les phases minérales dans des argilites. Même si ces travaux se limitaient à des mélanges simulés et composés seulement de deux ou trois phases, ils ont permis de valider le potentiel des réseaux de neurones pour l’analyse des données de diffraction. C’est plus récemment que de nombreuses études sur le sujet ont été produites. Une revue de ces travaux est proposée par SURDU & GYÖRGY (2023). Pour le moment, une majorité se concentre sur

des problèmes de classification. C’est le cas de l’approche de OVIEDO *et al.* (2019) basée sur des réseaux de neurones convolutifs (CNN), qui permet d’identifier le groupe d’espace (parmi 7 classes) ou la dimension du cristal. Une autre problématique traitée concerne l’identification de la symétrie du cristal (VECSEI *et al.* 2019; ZALOGA *et al.* 2020). La classification selon le système cristallin, et le groupe d’extinction (PARK *et al.* 2017) a aussi fait l’objet de recherches. D’autres travaux permettent une identification fiable et précise des matériaux (parmi plus de 1 000 classes) pour des réseaux métallo-organiques, à partir de signaux issus de la DRX (H. WANG *et al.* 2020). Les réseaux métallo-organiques sont des solides poreux hybrides cristallins, constitués d’ions métalliques reliés entre eux par des ligands organiques. Plus récemment encore, un Transformer, un nouveau type de modèle basé sur le principe de *self-attention* (VASWANI *et al.* 2017), propose une tâche similaire pour les réseaux métallo-organiques (CHEN *et al.* 2024). L’ensemble de ces études présente d’excellents résultats sur ces problèmes de classification et d’identification de phases.

Néanmoins, un problème apparaît lorsque l’on souhaite utiliser un réseau de neurones. Pour être performant, il a besoin d’un important nombre de données labélisées pour s’entraîner à réaliser la tâche qui lui est assignée. Mais, obtenir un nombre suffisant de diffractogrammes expérimentaux dont la phase minérale est connue est quasiment impossible. D’abord car le temps d’acquisition des données est trop important. Mais aussi parce qu’il faut obtenir une gamme de minéraux suffisamment large pour le minéral lui-même (par exemple différentes calcites) et pour constituer des mélanges polyphasiques. Ainsi, beaucoup de travaux utilisent des diffractogrammes synthétiques provenant de la *Inorganic crystal structure databases* (ICSD¹) pour entraîner les modèles, c’est notamment le cas des travaux suivants : (CHITTURI *et al.* 2021; LEE *et al.* 2021, 2020; PARK *et al.* 2017; SALGADO *et al.* 2023), etc. D’autres choisissent plutôt d’utiliser leurs algorithmes de simulation, comme c’est le cas dans les travaux de POLINE *et al.* (2024), où les diffractogrammes sont générés par un code Python.

On peut aussi augmenter artificiellement la taille de la base de données avec la *data augmentation*. Ce type d’approche a pour but de construire de nouvelles données d’apprentissage à partir de celles existantes. Par exemple, OVIEDO *et al.* (2019) proposent une approche de ce type appuyée par des principes physiques, pour les diffractogrammes. Malgré tout, des données synthétiques complètent la base d’entraînement. Une autre option consiste à combiner le bruit des données réelles avec la position théorique des pics (H. WANG *et al.* 2020) pour synthétiser des diffractogrammes. Ces approches sont complexes, d’une part car la mesure du bruit des données est complexe et d’autre part car les mailles non cubiques n’ont pas de dépendance linéaire entre la variation des paramètres de maille et la position des pics de diffraction.

Enfin, certains modèles sont entraînés uniquement sur des diffractogrammes expérimentaux (TATLIER 2011). Dans ce cas, la base de données est de taille très réduite. De plus, la création de bases de données réelles, qui est déjà complexe pour une tâche de classification, devient impossible lorsque l’on souhaite quantifier les phases minérales en se basant sur le signal de diffraction.

Ce problème d’inférence de proportions est donc une tâche très complexe à gérer avec les réseaux de neurones. Pour illustrer cette difficulté, on peut citer les travaux de LEE *et al.* (2021). Dans l’article, les auteurs écrivaient : “*Deep learning (DL) models trained with synthetic XRD data have never accomplished a satisfactory quantitative XRD analysis for the exact prediction of a constituent-phase fraction in unknown multiphase*

1. <http://www.fiz-karlsruhe.de/icsd.html>

inorganic compounds, although DL-based phase identification has been successful". Cet extrait souligne la complexité de la quantification de phases minérales, tout en insistant sur le problème de l'entraînement du modèle avec des données synthétiques. Les premiers travaux de LEE *et al.* (2020) proposaient de simplifier cette tâche de quantification en problème de classification. L'espace des proportions est partitionné en plusieurs sous-ensembles (0%-33%, 33%-66%, 66%-100%). Le réseau est alors capable de classer la phase minérale dans une plage de proportion plus réduite, mais toujours assez large. Une méthode alternative (LEE *et al.* 2021) consiste à combiner un CNN pour identifier les phases avec d'autres approches d'apprentissage automatique, pour réaliser la quantification dans un second temps. Ces deux études utilisent la base de données ICSD (dont l'accès nécessite un abonnement payant) pour récupérer des diffractogrammes théoriques afin d'entraîner les modèles. Les travaux concernant ce type de problème lié à l'analyse des signaux de diffraction sont pour le moment peu développés.

Mais d'autres domaines, tels que l'imagerie hyperspectrale, proposent déjà un grand nombre d'études pour des problématiques d'inférence de proportions. L'un des enjeux avec ces images est de quantifier les composants dans chaque pixel à partir d'un signal (BIOUCAS-DIAS *et al.* 2012). Ce type de données est très proche de celles obtenues avec la tomographie par DRX. Parmi les différentes solutions existantes, il y a l'algorithme SUnSAL (BIOUCAS-DIAS & FIGUEIREDO 2010) qui est très répandu, ou encore des approches bayésiennes (DOBIGEON *et al.* 2009). Ces dernières années, un grand nombre de travaux basés sur les réseaux de neurones proposent des alternatives. Notamment en utilisant des CNN pour prédire les abondances des composants (ZHANG *et al.* 2018) (WAN *et al.* 2021). D'autres plus récemment encore utilisent sur des Transformers (SCHEIBENREIF *et al.* 2023).

1.2 Structure du manuscrit

Le manuscrit est organisé en sept chapitres qui reprennent les différentes contributions de cette thèse.

Le premier chapitre introduit l'analyse des matériaux en utilisant la diffraction des rayons X. L'ensemble des outils et les algorithmes permettant le calcul théorique d'un signal (diffractogramme) issus de la DRX sont présentés. Pour illustrer les résultats, quelques simulations sont exposées.

Le deuxième chapitre introduit les réseaux de neurones. Un premier exemple simple est réalisé pour le traitement des diffractogrammes. Il s'agit de l'identification de phases minérales pour un problème avec 4 classes (mélanges de 4 phases dans des proportions variables).

Le troisième chapitre présente la méthodologie mise en place pour répondre à la problématique d'inférence de proportions en proposant des fonctions de perte adaptées pour les réseaux de neurones. Nous chercherons à déterminer le pourcentage de chaque composant dans un mélange. Divers exemples illustreront les résultats : les signaux obtenus par la DRX, les images hyperspectrales et des images synthétiques obtenues à partir de la base de données MNIST (LECUN 1998).

Le quatrième chapitre se concentre sur la quantification de phases minérales pour des données expérimentales. Elles sont issues d'une série de mesures effectuées au BRGM à l'aide d'un diffractomètre de laboratoire. Il sera notamment question de la qualité des simulations, et de création d'une base de données adaptée.

Le cinquième chapitre introduit le jeu de données issu de l'analyse tomographique par

DRX obtenu sur synchrotron. Ce type de jeu de données propose une analyse en trois dimensions du solide. En s'appuyant sur les travaux du chapitre précédent, l'enjeu est de construire une méthode pour quantifier les phases minérales.

Le sixième chapitre présente une approche novatrice des Transformers ou modèle auto-attentif (architecture d'apprentissage profond) pour la DRX, toujours pour quantifier les phases minérales à partir du diffractogramme. Cette approche est complétée par différents outils qui permettent la compréhension du signal et du modèle.

Le dernier chapitre conclut ce manuscrit et dresse un bilan des contributions. Diverses pistes d'améliorations ainsi que des perspectives sont également évoquées.

1.3 Résumé des contributions

L'objectif de cette thèse est d'effectuer une reconstruction tridimensionnelle complète de la minéralogie des solides, en proposant une méthode basée sur les réseaux de neurones pour l'analyse automatique de grands jeux de données, par exemple ceux issus d'études tomographiques par rayons X. L'identification et la quantification des phases minérales dans un échantillon à partir de diffractogrammes seront les enjeux majeurs. Pour répondre aux nombreuses problématiques liées à ces tâches, plusieurs contributions ont été mises en place.

1.3.1 Simulation de diffractogrammes

Premièrement, pour répondre à la problématique des bases de données très vastes nécessaires à l'entraînement des réseaux de neurones, des algorithmes de simulation de diffractogrammes ont été développés.

En effet, les méthodes neuronales nécessitent de larges bases de données afin que le réseau apprenne correctement la tâche qui lui est assignée. Dans le cas des signaux de diffraction, le réseau devra s'accoutumer aux pics de diffractions de chaque phase, et à la variabilité de ces derniers, que ce soit sur leur position, leur largeur ou leur intensité. Il faut donc constituer une base de données répondant à ces critères. Avoir une base assez grande n'est pas réalisable avec des diffractogrammes expérimentaux. Les plages de variation seraient limitées à une série d'échantillons, et le nombre de données serait contraint par l'aspect chronophage de l'ensemble du processus d'analyse, de la constitution des échantillons au post-traitement des données après la DRX.

Ainsi, il est nécessaire de pouvoir disposer d'une grande quantité de données pour répondre à ces problématiques. Deux manières de simuler ont été mises en place. Une première approche se distingue par sa vitesse de calcul. La seconde par un calcul qui décrit la physique du signal au mieux. Ce code permet notamment de modéliser des phases turbostratiques comme les argiles (FERRAGE *et al.* 2005a, 2010, 2005b). La première méthode est entièrement développée dans le cadre de cette thèse, alors que la seconde a été améliorée, notamment en automatisant et parallélisant un code existant développé par S. Grangeon au BRGM. Pour ces algorithmes, une description de la maille de la phase minérale est suffisante pour obtenir le signal de diffraction. La figure 1.1 illustre un exemple de simulation pour la phase minérale de la Gibbsite. Une comparaison avec une donnée réelle mesurée en laboratoire permet de voir la qualité de la donnée synthétique.

Des signaux multiphasés ou polyphasiques (mélange de plusieurs phases) sont générés à partir des diffractogrammes des minéraux simples en réalisant des combinaisons linéaires.

Cela permet de constituer des données contenant à la fois de la variabilité sur les diffractogrammes simple phase (représentant une unique phase minérale), mais aussi de couvrir l'ensemble des vecteurs de proportions possibles pour les signaux qui sont multi-phases (i.e des mélanges). Un grand nombre de données a pu être simulé pour constituer les bases de données nécessaires à l'entraînement de réseaux de neurones.

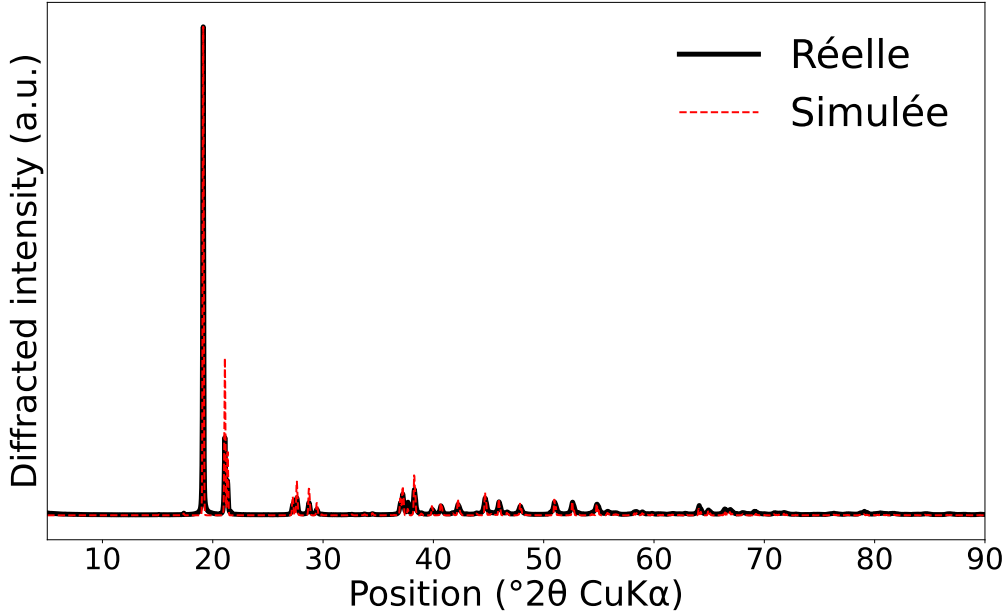


FIGURE 1.1 – Comparaison d’un diffractogramme de Gibbsite réel avec une simulation. Ce signal correspond à l’intensité des rayons X diffractés en fonction d’un angle θ (plus d’explications concernant la DRX sont détaillées dans la section 2.2).

1.3.2 Fonction de perte pour l’inférence de proportion

La modélisation d’une fonction de perte spécifique à l’inférence de proportion pour les réseaux de neurones a ensuite été proposée, principalement pour faire face au problème de quantification de phases minérales dans des mélanges.

Comme cela a été évoqué, les réseaux de neurones sont très performants pour des problèmes de classification. Un grand nombre de travaux étayent ces propos, pour le traitement du signal et plus spécialement dans le domaine médical (JAHANKHANI *et al.* 2006), (JIANG & KONG 2007), mais aussi pour la classification d’image avec des réseaux très performants comme ResNet (HE *et al.* 2016).

Dans le cadre de la classification, l’entraînement du réseau consiste à optimiser une fonction de coût. Cette fonction quantifie les erreurs effectuées sur les données d’entraînement. En l’optimisant, cela permet de réduire le nombre d’erreurs. Considérant un problème à K classes, le but est de minimiser l’erreur entre la sortie du réseau, $\mathbf{a} = (a_1, \dots, a_K) \in \mathbb{R}^K$, et un vecteur label (ou étiquette), $\mathbf{l} = (l_1, \dots, l_K) \in \{0, 1\}^K$, qui encode la classe de la donnée. Parmi les exemples les plus communs, il y a l’erreur des moindres carrés (MSE pour *Mean Square Error* en anglais), ou celle de l’entropie croisée (CE pour *Cross-Entropy* en anglais).

Cependant, lorsque les données ne correspondent pas à une classe mais à un mélange de

plusieurs composants, ce n'est plus un problème de classification. La finalité ne sera plus de prédire une classe mais plutôt d'obtenir en sortie du réseau un vecteur de proportions $\mathbf{p} = (p_1, \dots, p_K)$ dont chaque coordonnée sera associée à une classe. Ce problème d'inférence de proportions est plus complexe et les travaux le concernant sont peu commun dans la littérature. Malgré tout, quelques études en imagerie hyperspectrale (WAN *et al.* 2021; ZHANG *et al.* 2018) utilisent des réseaux de neurones, mais sans une modélisation adaptée respectant les contraintes du vecteur $\mathbf{p}$. La première contrainte décrite par l'équation (1.1), est celle de non-négativité pour chacune des coordonnées du vecteur.

$$p_i \geq 0, \quad \text{pour tout } i = 1, \dots, K \quad (1.1)$$

La seconde décrite par l'équation (1.2) contraint la somme de tous les éléments du vecteur à être égale à 1,

$$\sum_{i=1}^K p_i = 1 \quad (1.2)$$

Toutefois les fonctions classiques (MSE ou CE, par exemple) ne garantissent pas d'obtenir un vecteur de proportions. Elles ne sont donc pas la solution adaptée pour la quantification de phases minérales. La création d'une fonction de perte spécifique pour l'inférence de proportions est donc nécessaire.

Ainsi des fonctions de perte adaptées ont été développées, en s'appuyant sur une modélisation construite à partir de la loi de Dirichlet. En choisissant une architecture appropriée, ces fonctions peuvent être appliquées à différents types de données pour estimer un vecteur de proportions. Les résultats montrent les performances et la stabilité de ces fonctions à travers les différentes applications considérées, notamment en comparant à des fonctions usuelles.

1.3.3 Applications aux diffractogrammes expérimentaux

En s'appuyant sur les contributions évoquées dans les sections précédentes 1.3.2 et 1.3.1, une solution pour le traitement des diffractogrammes expérimentaux est proposée. Elle s'appuie sur un réseau de neurones et fournit une méthode générique pour le traitement de jeux de données issus de la DRX.

La méthode comporte plusieurs étapes, il faut d'abord construire une base de données synthétiques à partir des simulations décrites ci-dessus 1.3.1, ce qui présente plusieurs avantages. Cela permet de faire face au manque de données expérimentales tout en incluant une grande variabilité dans les diffractogrammes. Les données labellisées nécessaires à l'entraînement d'un réseau de neurones peuvent donc être générées. Lorsque la base d'entraînement synthétique est construite, le réseau peut être entraîné en utilisant la fonction de perte mentionnée au paragraphe précédent 1.3.2. Une fois le réseau entraîné, il peut traiter des données réelles, qu'elles proviennent de laboratoires ou d'un synchrotron. C'est grâce aux données obtenues avec le synchrotron que la reconstitution tridimensionnelle de la minéralogie est possible.

En suivant ces trois étapes, la méthode peut être confrontée à n'importe quel type de jeu de données, pour des solides naturels ou artificiels. En effet, les algorithmes de simulation permettent de générer une grande majorité des phases minérales connues. Cela permet de s'affranchir d'un problème de construction de base à partir de données réelles. Cependant, la méthode est dépendante de la qualité des simulations, qui doivent être proches des données expérimentales afin de garantir de bons résultats. Un avantage

important est la vitesse d’analyse du réseau une fois entraîné, capable d’analyser plusieurs centaines de données en quelques secondes. En effet, il est capable de traiter rapidement un grand nombre de signaux et de prédire la composition des échantillons associés en termes de phases minérales.

1.3.4 L’approche des Transformers

La dernière contribution concerne les Transformers, un nouveau type d’architecture neuronale conçue notamment pour le traitement et la génération de texte. Ces réseaux sont très performants et surpassent les architectures traditionnelles notamment pour le traitement automatique du langage (BROWN *et al.* 2020; DEVLIN 2018), mais aussi le traitement d’images (DOSOVITSKIY *et al.* 2021).

Un Vision Transformer (ViT), habituellement utilisé pour le traitement d’image, a été adapté aux signaux de diffraction. C’est la première fois que ce type de réseau est développé pour l’inférence de proportions et la quantification de phases minérales. Une autre nouveauté sera d’accompagner les Transformers d’outils de visualisation tels que l’*Attention Rollout* (ABNAR & ZUIDEMA 2020) ou le *positional encoding* (DOSOVITSKIY *et al.* 2021) au traitement de signaux. Ces outils permettent de comprendre le traitement des données par le modèle. Ils fournissent aussi des indications sur la propagation de l’information au sein du Transformer.

1.4 Publications

Au cours de la thèse, plusieurs articles ont permis la valorisation des différentes contributions évoquées ci-dessus.

- Les résultats du chapitre 3 concernant l’inférence de proportions ont abouti à un article pour la conférence *European Signal Processing Conference* en 2023 (SIMONNET *et al.* 2023).
- Les travaux réalisés présentant une méthode d’analyse pour des données de diffraction expérimentales ont été publiés dans le journal *IUCrJ* en Septembre 2024 (SIMONNET *et al.* 2024b). Ces travaux sont présentés dans le chapitre 4. Un préfaçage de l’article a été proposé par l’éditeur (PRASIANAKIS 2024).
- Un article a également été soumis à la conférence *ICASSP* en 2024 (SIMONNET *et al.* 2024a). Cet article concerne le travail présenté dans le dernier chapitre, présentant l’application d’un Transformer pour quantifier les phases minérales. Il propose également des outils pour visualiser la propagation de l’information dans le modèle.

Chapitre 2

La diffraction de rayons X et simulation de diffractogrammes

2.1 Les matériaux analysés

2.1.1 La structure des cristaux

Avant de présenter plus en détails la diffraction par rayons X (DRX), et pour mieux comprendre les calculs qui suivent, une brève introduction sur la structure des cristaux est proposée.

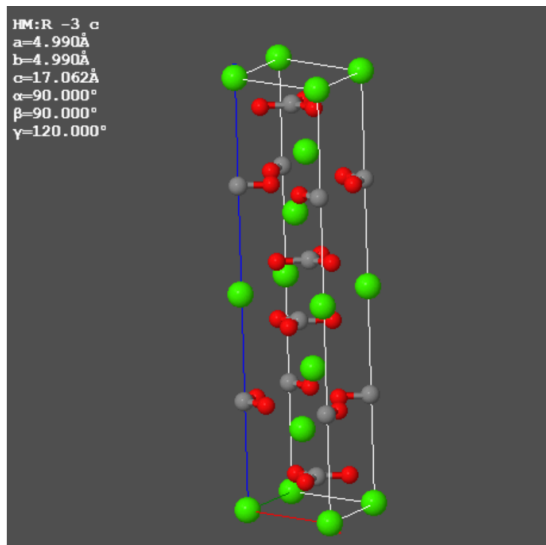
Un cristal est un solide dont les atomes sont assemblés à l'échelle moléculaire de manière régulière. Parmi eux, on peut par exemple citer, la neige, le sucre, le sel, etc.

Le cristal peut être décrit comme une répétition périodique d'un même motif dans les trois directions de l'espace. La plus petite structure qui se répète est appelée la maille, des exemples sont proposés sur la figure 2.1. C'est un parallélépipède contenant des atomes à des positions spécifiques. Le parallélépipède est décrit par trois vecteurs ($\vec{a}$, $\vec{b}$, $\vec{c}$) avec leurs normes et les angles entre chacun des vecteurs (α , β et γ). L'exemple le plus simple est $\|\vec{a}\| = \|\vec{b}\| = \|\vec{c}\|$, et $\alpha = \beta = \gamma = 90^\circ$. Le résultat est un cube comme sur l'exemple illustré de la figure 2.2a.

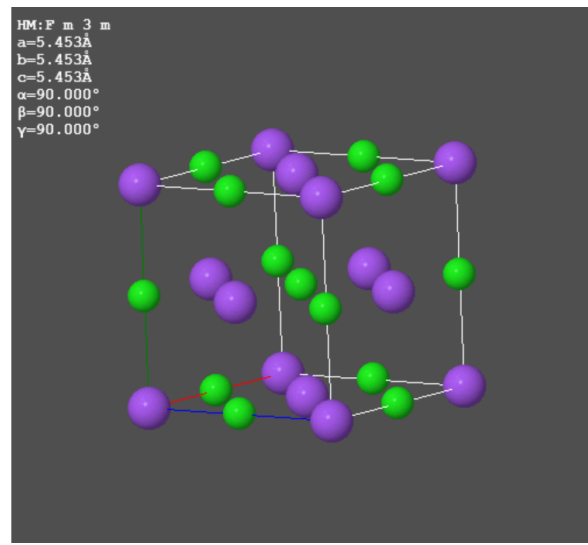
L'autre partie de la description est donnée par la position et la nature des atomes dans la maille. La position d'un atome est donnée par un vecteur $\vec{r}$ tel que $\vec{r} = x \times \vec{a} + y \times \vec{b} + z \times \vec{c}$, avec x , y et z les coordonnées fractionnelles de cet atome dans la maille (voir figure 2.2b). Des exemples de mailles sont donnés sur la figure 2.1. Le cristal est l'objet physique résultant de la répétition de mailles dans les trois directions de l'espace.

Cette maille est spécifique pour chaque **phase minérale**. La phase minérale fait référence à la composition chimique et structurale d'une substance cristalline. Lorsque la structure est connue, un fichier CIF (Crystallographic Information File) (HALL *et al.* 1991) y est associée. Plusieurs milliers de phases sont aujourd'hui connues avec fichier CIF spécifique pour chacune d'entre-elles. Il donne toutes les informations nécessaires pour décrire un cristal : paramètres de maille, groupe d'espace, symétrie, taux d'occupation, etc.

La description du cristal est complétée par le nombre de répétitions de maille dans chacune des directions de l'espace. On note N_1 le nombre de répétitions pour la direction du vecteur $\vec{a}$, N_2 pour la direction $\vec{b}$ et N_3 pour la direction $\vec{c}$. Dans le cas où $N_1 \times \vec{a} = N_2 \times \vec{b} = N_3 \times \vec{c}$, le cristal sera dit isotrope, sa dimension étant la même quelle que soit la direction de l'espace. Dans tous les autres cas, le cristal sera dit anisotrope. La structure est fréquemment décrite dans le **Réseau de Bravais**, l'espace contenant l'ensemble de

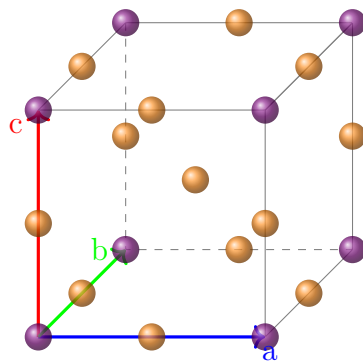


(a) Calcite

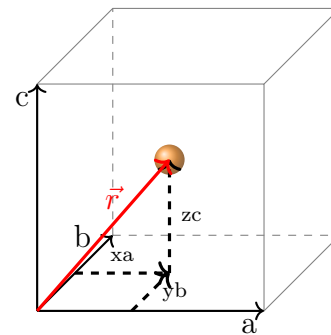


(b) Halite

FIGURE 2.1 – Exemples de mailles pour deux phases minérales différentes



(a)



(b)

FIGURE 2.2 – (Exemple d'une maille cubique (a) contenant des atomes de deux natures différentes (violet et orange) sur des positions données. Il est à noter que les sommets doivent être occupés par un atome. La figure (b) décrit la localisation d'un atome dans la maille.

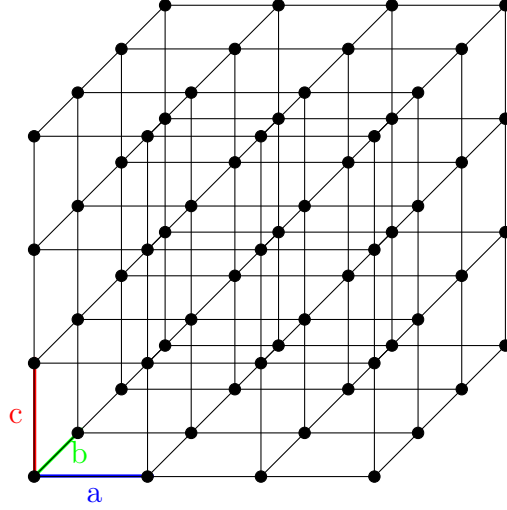


FIGURE 2.3 – Exemple d'un réseau de Bravais

ces mailles. La figure 2.3 donne un exemple de réseau de Bravais.

Réseau de Bravais C'est un quasi espace vectoriel, qui est utilisé notamment en cristallographie. La différence avec l'espace vectoriel est que les scalaires, notés m_i , définis dans ce type de réseau sont tels que $m_i \in \mathbb{Z}$ (non plus des nombres réels ou complexes). C'est donc une distribution régulière de points, que l'on appelle noeuds. Chacun de ces noeuds peut être vu comme le sommet d'une maille. Le réseau est donc une translation de mailles qui représente la périodicité de la distribution atomique d'un cristal. Un réseau de Bravais est donc un ensemble décrit à partir de plusieurs vecteurs :

$$m_1 \vec{a} + m_2 \vec{b} + m_3 \vec{c}, \quad (2.1)$$

où $m_1, m_2, m_3 \in \mathbb{Z}$, et $\vec{a}, \vec{b}, \vec{c}$ sont les vecteurs définissant la maille.

Le réseau est défini uniquement à partir des paramètres de maille ($\vec{a}, \vec{b}, \vec{c}, \alpha, \beta$, et γ). Mais, dans un souci de simplification, on utilisera par la suite un autre espace construit à partir du réseau de Bravais pour effectuer les calculs, l'espace réciproque.

Espace réciproque Il est important car il permet de simplifier les calculs théoriques liés à la DRX, qui seront donnés par la suite. La construction de cet espace se fait à partir du réseau de Bravais de la manière suivante :

$$\begin{cases} A^* = (\vec{b} \wedge \vec{c})/V \\ B^* = (\vec{c} \wedge \vec{a})/V, \\ C^* = (\vec{a} \wedge \vec{b})/V \end{cases}$$

où $\wedge$ représente le produit vectoriel et $V = a \times b \times c$ est le volume de la maille.

On peut aussi écrire :

$$a_i \cdot A_j^* = \delta_{ij} = \begin{cases} 1 \text{ si } i = j \\ 0 \text{ sinon} \end{cases}.$$

Par la suite, on notera (h, k, l) les coordonnées dans l'espace réciproque.

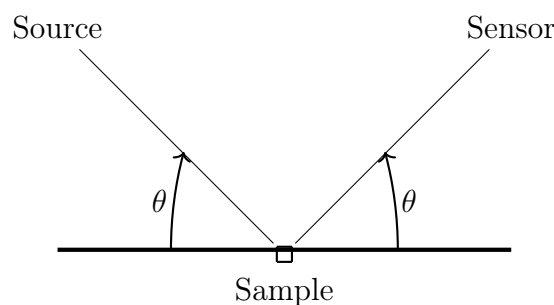


FIGURE 2.4 – Schéma de la diffraction par rayons X

Les matériaux sont, en réalité, composés de plusieurs phases minérales différentes, et correspondent à un assemblage de celles-ci. La détermination précise des abondances de chaque phase est possible, notamment grâce aux signaux obtenus par la DRX. Les notions et notations introduites dans cette section permettent d'exposer les calculs théoriques liés au phénomène de diffraction à l'échelle d'un cristal. Ces calculs seront nécessaires par la suite pour le calcul théorique des diffractogrammes.

2.2 La diffraction des rayons X

La diffraction des rayons X (*X-Ray Diffraction* en anglais) (DRX) est une méthode d'analyse des matériaux, notamment utilisée lorsqu'ils sont sous leur forme cristalline. Les matériaux analysés peuvent être, par exemple, sous la forme de lame ou de poudre. Le principe est d'envoyer des rayons X sur l'échantillon et de collecter les rayons X diffractés.

Les rayons X sont une forme de rayonnement électromagnétique constitués de photo-électrons, qui partagent des propriétés des ondes et des particules et interagissent avec la matière de différentes manières. Pour le principe de diffraction, on s'intéressera uniquement à la diffusion élastique des rayons X, c'est-à-dire aux rayons X qui conserveront la même longueur d'onde après interaction avec la matière. La diffusion non élastique, due notamment à l'absorption par les atomes de l'énergie des rayons X, génère une émission de longueur d'onde différente. Elle sera discutée plus tard dans la section 2.4.3.

La figure 2.4 illustre une des configurations géométriques possibles du procédé, qui est celle employée pour ce travail. Elle permet de voir que la source des rayons X et le capteur sont positionnés à un angle θ par rapport à un plan échantillon.

Avec une longueur d'onde des rayons X constante, notée λ , l'analyse se fait en balayant une plage angulaire en 2θ donnée. On considère un angle 2θ puisque les angles θ au niveau de l'émetteur et du récepteur varient en même temps. Au final, le résultat obtenu est un diffractogramme, un signal qui donne l'intensité des rayons X diffractés en fonction de l'angle 2θ . Plusieurs illustrations sont fournies à la fin du chapitre à la figure 2.14. Les variations d'intensité d'un angle de diffraction à un autre sont dues au **déphasage des ondes**. Le paragraphe qui suit donne quelques explications sur ce phénomène.

Déphasage d'onde Pour expliquer le déphasage d'onde, prenons l'exemple de deux ondes émises par la source de rayons X. Les ondes reçues par le capteur en sortie correspondent à la somme des ondes émergentes. Dans le cas où les ondes sont en phases (figure 2.5a), les interférences sont constructives et l'intensité des ondes diffractées est importante. Dans le cas où elles ne le sont pas, les interférences sont destructives, et

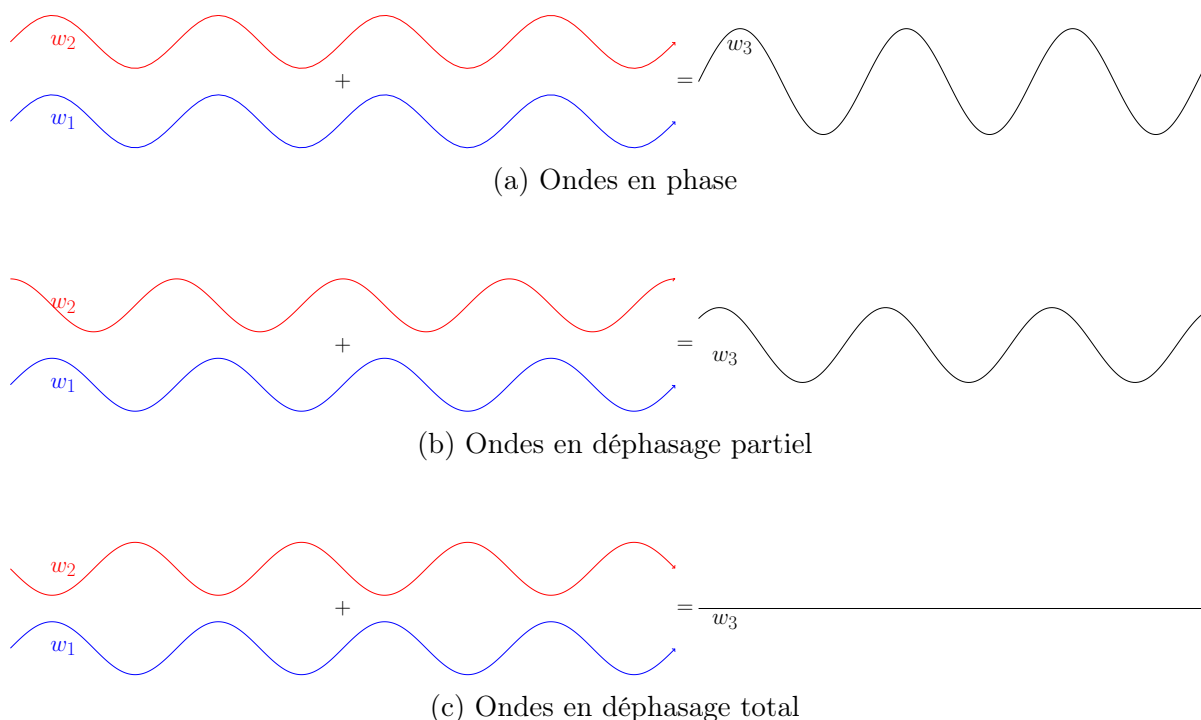


FIGURE 2.5 – Illustration du phénomène de déphasage d’onde à différents niveaux. Partant du cas (a) avec des ondes parfaitement en phase, jusqu’à un déphasage total (c).

l’intensité de diffraction est plus faible, voire nulle si le déphasage est total, comme sur l’exemple de la figure 2.5c.

Un déphasage des deux ondes se crée lorsque qu’il y a une différence de chemin parcouru, entre leur émission et leur réception par le capteur. Or, pour atteindre les différents éléments du cristal, les ondes ne vont pas toujours parcourir le même chemin. Ceci est vrai à différentes échelles, d’abord celle des électrons qui gravitent autour des noyaux des atomes, puis pour atteindre les différents atomes qui composent la maille, et enfin au niveau des plans d’atomes.

L’enjeu de la section sur le calcul théorique du diffractogramme (section 2.3) sera d’introduire les équations qui permettent le calcul des termes liés au déphasage à ces différentes échelles de la matière.

Plus de détails concernant les données acquises en laboratoire, ou au synchrotron, seront donnés dans les sous-sections qui suivent.

2.2.1 Acquisition de données en laboratoire

Les premières données que nous traiterons sont celles mesurées dans un laboratoire du BRGM. Néanmoins, le montage expérimental est celui couramment utilisé dans tous les laboratoires, rendant notre approche reproductible. Elles proviennent d’échantillons sous forme de poudre dont la taille du grain est inférieur à $10^{-6}m$. Ce sont des composites qui assemblent des pôles purs. La quantité de pôles purs est obtenue par pesage au milligramme près. Des rotations de l’échantillon au centre de l’analyse (voir figure 2.4) permettent de couvrir toutes les orientations du cristal possibles. Pour chaque pas en 2θ , une dizaine de rotations de l’échantillon sont effectuées. Le diffractomètre utilisé est un Bruker D8 Advance, avec un détecteur LynxEye XE-T. Une photo de l’appareil est présentée à la figure 2.6. La source de rayons X est une anode de cuivre de longueur d’onde

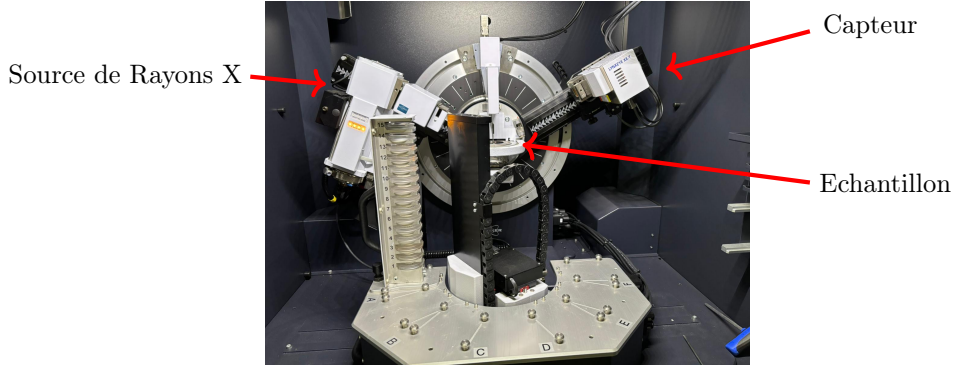


FIGURE 2.6 – Photo du diffractomètre : Bruker D8 Advance (au laboratoire du BRGM).

$\lambda = 1.5418 \text{ \AA}$. La taille du faisceau de rayons X est proche de $2\text{mm} \times 2\text{cm}$. Le nombre de photons envoyés sur l'échantillon est de l'ordre de 10^{-9} photons par seconde. La plage angulaire considérée en 2θ s'étend de 4.0001° à 89.9904° par pas de 0.029611° . Enfin, le temps d'acquisition d'un diffractogramme avec cet appareil est un paramètre réglable. Pour les expériences présentées ici, il a été fixé à deux heures.

2.2.2 Acquisition au synchrotron

Les données de tomographie diffraction ont été acquises sur la ligne ID11 du synchrotron ESRF dans le cadre de l'étude de CLARET *et al.* (2018). Cela permet notamment une analyse non destructive du solide. Une photo de l'appareil est donnée à la figure 2.7. Brièvement, les expériences ont été conduites en utilisant un faisceau monochromatique de longueur d'onde $\lambda=0.1897 \text{ \AA}$. La distance entre échantillon et détecteur, ainsi que l'énergie exacte du faisceau incident, ont été calibrées en utilisant un étalon NIST de CeO_2 . Le faisceau a été collimaté de manière à être pseudo-rectangulaire et de taille $10 \times 10\mu\text{m}$.

Les données ont été collectées en transmission, en utilisant un détecteur FReLoN, et de manière itérative en faisant translater l'échantillon selon un axe y perpendiculaire au faisceau et en le faisant tourner sur lui-même par pas angulaires ω .

L'ensemble des couples (y, ω) nécessaires à la constitution d'un sinogramme ont été acquis selon des pas de $10 \mu\text{m}$ et 1° , respectivement. Une transformée de Fourier inverse a ensuite été appliquée aux sinogrammes pour spatialiser les données dans l'espace réel.

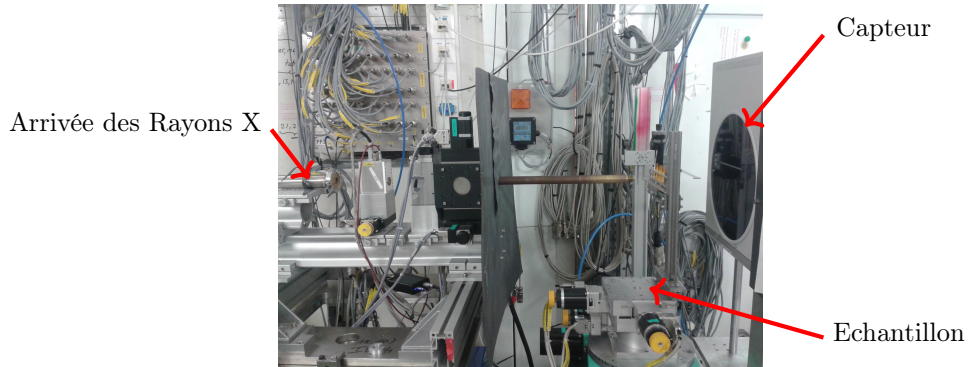


FIGURE 2.7 – Photo de l'appareil de mesure avec le synchrotron (Par S.Grangéon à l'ESRF, Grenoble).

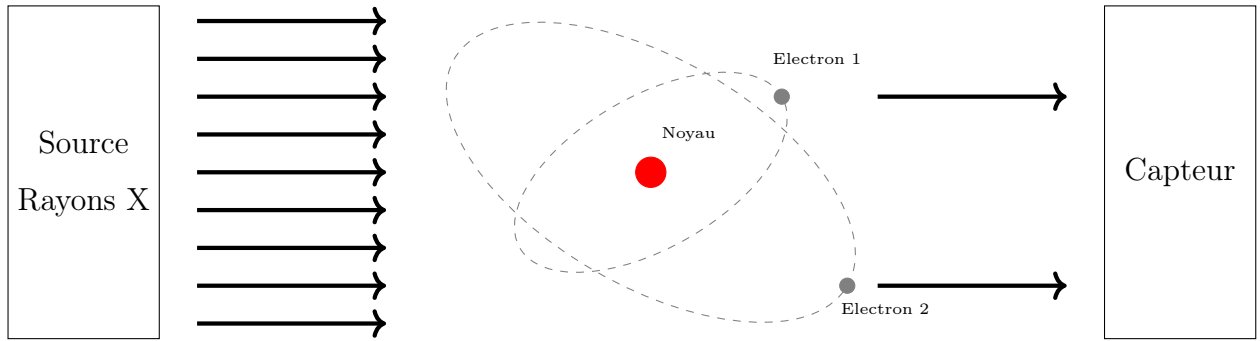


FIGURE 2.8 – Schéma simplifié de la diffraction au niveau d'un atome pour un angle $2\theta = 0$.

2.3 Calcul théorique de diffractogrammes

Le but de cette section n'est pas de présenter en détails le calcul théorique du diffractogramme, mais plutôt d'introduire les équations utiles pour l'implémentation d'un algorithme de simulation, ainsi que quelques brèves explications sur la provenance de celles-ci.

Tout d'abord, le diffractogramme est un signal résultant de la DRX. Il représente l'intensité des rayons X diffractés en fonction de l'angle de diffraction θ . Il est obtenu en calculant le déphasage des ondes émises par les rayons X à différentes échelles : l'atome, la maille et le cristal. Chacune de ces échelles étant associée aux trois principaux termes qui permettent un calcul théorique du signal : le facteur de diffusion atomique, la fonction d'interférence, et le facteur de structure.

2.3.1 Facteur de diffusion atomique

Ce premier facteur donne l'intensité diffractée au niveau atomique. Il résulte de l'interaction des rayons X avec les atomes. Les atomes sont constitués d'un noyau et d'un nuage d'électrons qui gravite autour du noyau. Et c'est plus précisément avec ces électrons que les rayons X vont interagir. Plus ils seront nombreux dans le nuage de l'atome, plus l'intensité diffractée sera importante, comme peut le montrer le schéma à visée pédagogique de la figure 2.8.

L'équation mise au point par CROMER & MANN (1968) permet de calculer le facteur de diffusion atomique pour tout angle θ de diffraction, et quel que soit l'atome. Son expression est la suivante :

$$f_0 = \sum_{i=1}^5 a_i \exp(-b_i s^2) + c, \quad (2.2)$$

où $(a_i, b_i, i = 1, \dots, 5)$ et c sont les coefficients spécifiques pour chaque atome, $s = \frac{1}{2d}$ est le terme dépendant de l'angle de diffraction θ . Dans ce terme, d est la distance interréticulaire (i.e la distance entre deux plans d'atomes), liée à θ par la loi de Bragg introduite par la suite (équation (2.5)).

La figure 2.9a montre les valeurs du facteur de diffusion atomique avec l'exemple de trois atomes, le carbone, l'oxygène et le calcium, de numéro atomique 6, 8 et 20 respectivement. Les facteurs sont calculés avec la formule décrite par l'équation (2.2). A

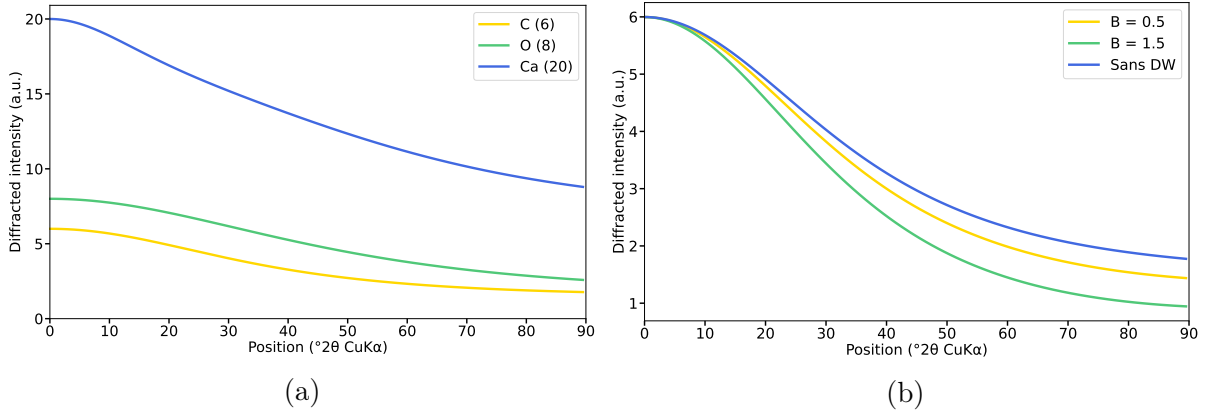


FIGURE 2.9 – (a) Exemples de facteurs de diffusion atomique calculés à partir de l’équation (2.2), pour trois atomes différents : le carbone, l’oxygène et le calcium dont les numéros atomiques sont 6, 8 et 20 respectivement. Ces numéros correspondent au nombre d’électrons qui gravitent dans le nuage autour de leur noyau. La figure (b) est une illustration de l’effet d’agitation thermique (Debye-Waller) avec des valeurs de B variables.

partir de cette figure, plusieurs observations peuvent être faites. A l’angle de diffraction $\theta = 0^\circ$, l’intensité diffractée correspond au numéro atomique de l’atome, donc au nombre d’électrons dans le nuage autour du noyau. A cet angle, toutes les ondes restent en phase (voir section 2.2), donc les interférences sont constructives et l’intensité est maximum.

On peut ensuite observer que l’intensité décroît au fur et à mesure que l’angle θ augmente jusqu’à ce que $\theta = 90^\circ$. Lorsque l’angle de diffraction augmente, une différence de chemin apparaît entre les ondes pour atteindre les différents électrons, menant à un déphasage et à des interférences destructives. A cette échelle, la différence de chemin ne peut pas être assez importante pour créer un déphasage complet. Ainsi l’intensité ne s’annule pas. On sera sur un déphasage partiel, comme illustré par la figure 2.5b.

Enfin, en observant le schéma de la diffraction introduit précédemment (figure 2.4), on comprend assez vite qu’une simple symétrie axiale permet d’obtenir l’intensité diffractée pour θ de 90° à 180° .

Facteur de Debye-Waller Le facteur de diffusion atomique peut être complété par l’effet de Debye-Waller (DEBYE 1915) qui correspond à l’agitation thermique des atomes. En pratique, plus cette agitation sera importante, plus l’intensité diffractée va diminuer lorsque l’angle θ augmente (voir figure 2.9b). On le calcule via,

$$f = f_0 \times \exp\left(-B \frac{\sin^2 \theta}{\lambda^2}\right), \quad (2.3)$$

où f_0 est le facteur de diffusion atomique et B est le paramètre qui décrit l’agitation thermique. Plus la valeur de B est élevée, plus l’agitation thermique sera importante et plus l’intensité sera impactée. Cette équation est valable si l’on suppose l’agitation isotrope. Une autre formule doit être considérée dans le cas d’une agitation anisotrope. Toutes ces équations sont détaillées dans le livre *Modern Powder Diffraction* (BISH & POST 1990).

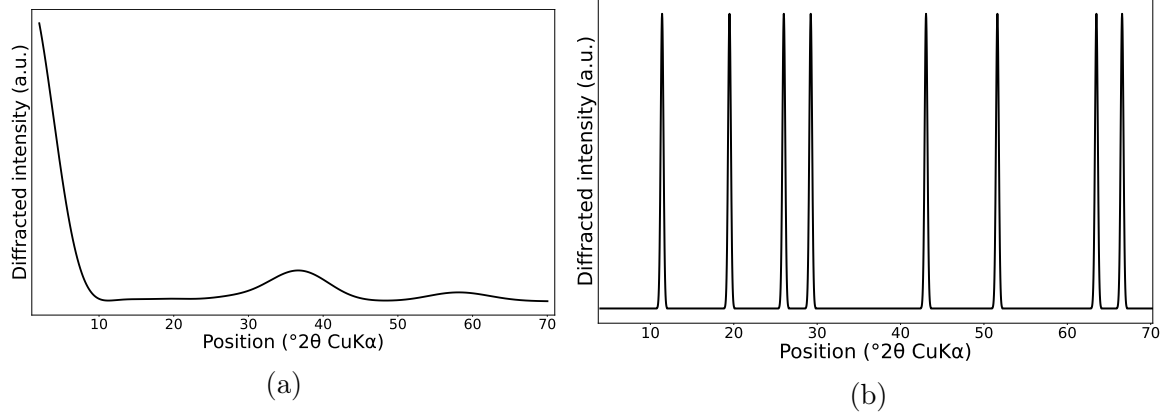


FIGURE 2.10 – Exemple d’un facteur de structure (a) et d’une fonction d’interférence (b).

2.3.2 Facteur de structure

Le facteur de diffusion atomique décrit l’interaction des rayons X au niveau atomique. Le facteur de structure décrit quant à lui l’intensité diffractée par une maille. Il prend en compte l’architecture de la maille, que ce soit sa forme (décrite par les paramètres de maille $a, b, c, \alpha, \beta, \gamma$) mais aussi la position et la nature des atomes qui la composent. Comme pour le facteur de diffusion atomique, le facteur de structure correspond au déphasage des ondes. A l’échelle de la maille, il permettra de calculer la différence de chemin des ondes pour atteindre les différents atomes de la maille. Comme il tient compte des atomes présents dans la maille, son calcul nécessite le facteur de diffusion atomique.

Dans le cas de maille centro-symétriques, pour avoir l’expression du facteur de structure, il faut se placer dans l’espace réciproque décrit dans la section 2.1.1, où l’on se place selon des coordonnées (h, k, l) . C’est une somme sur l’ensemble des N_{atom} atomes présents dans la maille, que l’on regroupe par nature d’atomes différents. Ainsi, l’expression utilise une somme sur le nombre d’atomes de natures différentes (somme sur P) et dans celle-ci, sur le nombre d’atomes de chaque type (somme sur N_P). Les coordonnées des atomes sont données par les triplets (x_n, y_n, z_n) . On obtient la formule,

$$\left(\left[\sum_P \sum_{N_P} f_{N_P} \cos 2\pi(hx_n + ky_n + lz_n) \right]^2 + \left[\sum_P \sum_{N_P} f_{N_P} \sin 2\pi(hx_n + ky_n + lz_n) \right]^2 \right)^{1/2}. \quad (2.4)$$

La figure 2.10a montre les valeurs de ce facteur pour un C-S-H.

2.3.3 Fonction d’interférence

Le dernier terme introduit est la fonction d’interférence. Il permet de calculer les interférences d’ondes liées à des interactions entre ondes issues des mailles du cristal. Lorsque deux ondes sont parfaitement en phase, les interférences sont constructives. Quand un déphasage apparaît, elles deviennent destructives, annulant l’intensité de diffraction. Une différence de chemin entre les ondes peut être à l’origine de ce déphasage.

Ce phénomène apparaît lorsque les ondes se reflètent sur des plans différents. Une illustration à visée pédagogique est montrée à la figure 2.11. Les ondes w_1 et w_2 sont envoyées selon un angle θ sur un plan, puis l’on collecte les ondes diffractées selon le même angle θ . Les ondes w_1 et w_2 se reflètent sur deux plans différents, P_1 et P_2 . D’après le schéma, l’onde w_1 parcourt une distance supplémentaire $AX_2 + X_2B$. Cette différence

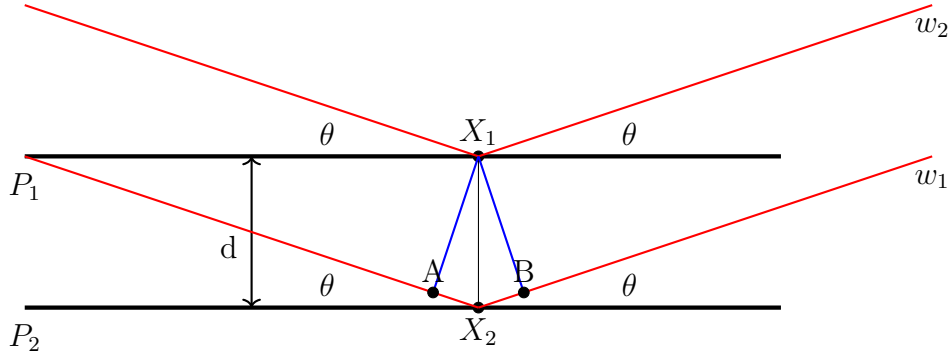


FIGURE 2.11 – Illustration simplifiée de la loi de Bragg entre deux plans P_1 et P_2 , séparés d'une distance d et soumis à deux ondes de rayons X w_1 et w_2 qui diffractent selon un angle θ .

de chemin entre les ondes, va amener un déphasage, et des interférences destructives dans le cas où la différence de chemin n'est pas multiple de la longueur d'onde λ des rayons X. Ceci aura pour conséquence d'annuler le phénomène de diffraction pour la majorité des angles 2θ considérés, amenant un signal parsemé uniquement de pics ponctuels. Ce phénomène d'interférence est régi par la **Loi de Bragg** qui stipule que :

$$2d \sin(\theta) = n\lambda, \quad (2.5)$$

où d est la distance entre deux plans d'atomes et $n \in \mathbb{N}$.

Preuve A l'aide du schéma 2.11, une idée de preuve peut être apportée. Pour que les interférences soient constructives, il est nécessaire que la longueur de la différence de chemin soit un multiple de la longueur d'onde des rayons X. Ainsi la distance $AX_2 + X_2B = 2 \times AX_2$ doit être telle que $2 \times AX_2 = n\lambda$, où $n \in \mathbb{N}$. Or, si l'on projette A sur le segment X_1X_2 , alors $AX_2 = d \sin(\theta)$. Donc $2 \times AX_2 = 2d \sin(\theta) = n\lambda$, ce qui conclut la démonstration.

Numériquement, la fonction Φ_h ci-dessous, permet de calculer la fonction d'interférence dans la direction du vecteur a (associée à la direction h de l'espace réciproque),

$$\Phi_h = \frac{\sin^2(\pi h N_1)}{N_1 \sin^2(\pi h)}, \quad (2.6)$$

où N_1 est le nombre de répétitions de maille dans la direction de a . Les calculs de Φ_k et Φ_l dans les directions de b et c respectivement sont également nécessaires pour le calcul de la fonction d'interférence.

La figure 2.10b montre une illustration de la fonction d'interférence. On constate qu'elle permet de localiser les pics, puisque l'on observe uniquement des Diracs. L'intensité sera déterminée par les deux autres termes : le facteur de diffusion atomique, et le facteur de structure.

Une fois toutes ces équations introduites, il est possible de calculer le diffractogramme théorique, à partir d'une description de la maille. La prochaine section introduit ces calculs.

2.4 Simulation numérique de diffractogrammes

Dans cette section, deux méthodes de simulation de diffractogrammes seront présentées. La première est assez répandue et reprend le fonctionnement de logiciels comme Mercury¹, ou encore celui de l'affinement Rietveld.

Les deux méthodes de simulation présentées dans cette section utilisent les équations (2.2),(2.3),(2.4),(2.6).

2.4.1 Première approche : un calcul simplifié

Cette approche utilise uniquement le facteur de diffusion atomique, et le facteur de structure dans les calculs. A partir de la loi de Bragg, et en choisissant des (h, k, l) entiers, on se limite à ceux qui respectent les conditions d'interférences constructives, i.e. là où il y aura des pics de diffraction. Puis, en se plaçant dans l'espace réciproque, il reste à calculer l'intensité à chacune des coordonnées $(h, k, l) \in \mathbb{N}^3$ que nous considérons. Chacun de ces points est alors associé à un angle θ . Finalement une distribution gaussienne autour des pics d'intensité est faite ce qui permet d'obtenir le diffractogramme.

Algorithme 1 : Simulation rapide de diffractogrammes

Data : Fichier CIF contenant les paramètres de maille, Vecteur de la plage angulaire de diffraction souhaitée θ

Result : Diffractogramme $D \in \mathbb{R}^d$

- 1 Calcul des facteurs de diffusion atomique F_{dif} (équation (2.2))
 - 2 Calcul des coordonnées de l'espace réciproque $(h, k, l) \in \mathbb{N}^3$ maximales, à partir de θ et des paramètres de maille donnés par *file*
 - 3 Création d'une matrice $M \in \mathcal{M}_{n \times 3}$ contenant tous les triplets (h, k, l) jusqu'à $h_{max}, k_{max}, l_{max}$
 - 4 **pour** $i \leftarrow 0$ **to** n **faire**
 - 5 Calcul du l'angle de diffraction θ'
 - 6 Interpolation du facteur de diffusion atomique pour θ' à partir de F_{dif}
 - 7 Calcul du facteur de structure F_s (équation (2.4))
 - 8 **fin**
 - 9 Somme des intensités pour les θ' qui se repètent
 - 10 Distribution gaussienne autour des pics pour les θ'
 - 11 Diffractogramme D dont l'intensité est nulle sauf aux pics de diffraction θ' , et autour avec les distributions gaussiennes.
-

Cette méthode présente l'avantage d'un calcul très rapide, permettant de simuler un grand nombre de diffractogrammes en un temps très court. Cependant, le calcul de l'intensité ne se fait pas pour tous les angles de diffractions, ce qui mène à des calculs approximatifs sur certains angles θ , notamment pour des cristaux dont la forme n'est pas la projection macroscopique de la maille. Ce qui est le cas pour les cristaux anisotropes, i.e lorsque l'égalité $N_1 \times a = N_2 \times b = N_3 \times c$ n'est pas vérifiée. Néanmoins, dans la majorité des cas, cette approximation est acceptée.

L'avantage principal de ce code est le temps de calcul, étant donné qu'il se limite aux calculs de deux facteurs aux angles où il y a des pics de diffraction. Cela évite notamment

1. <https://www.ccdc.cam.ac.uk/solutions/software/mercury>

tous les calculs aux angles où l'intensité diffractée sera nulle.

Finalement, l'ordre de grandeur du temps de calcul pour une simulation tourne autour de $10^{-1}s$ pour les mailles avec peu d'atomes comme le Quartz (LEVIEN *et al.* 1980). Les mailles comptant un nombre d'atomes plus important comme la Alite (ÁNGELES *et al.* 2008a) ou l'Ettringite (MOORE & TAYLOR 1970a) ont un temps de calcul de l'ordre de quelques secondes.

Tous les calculs sont faits avec Python, en utilisant le package Crystals (RENÉ DE COTRET *et al.* 2018) pour la lecture automatique des fichiers CIF. Le pseudo-code (algorithme 1) fournit plus de détails concernant cette méthode.

La seconde approche développée ne permet pas des calculs aussi rapides. Son objectif était plutôt de faire face à n'importe quels types de cristaux, y compris ceux possédant des défauts d'empilements, tout en se rapprochant le plus possible des mesures faites avec un diffractomètre.

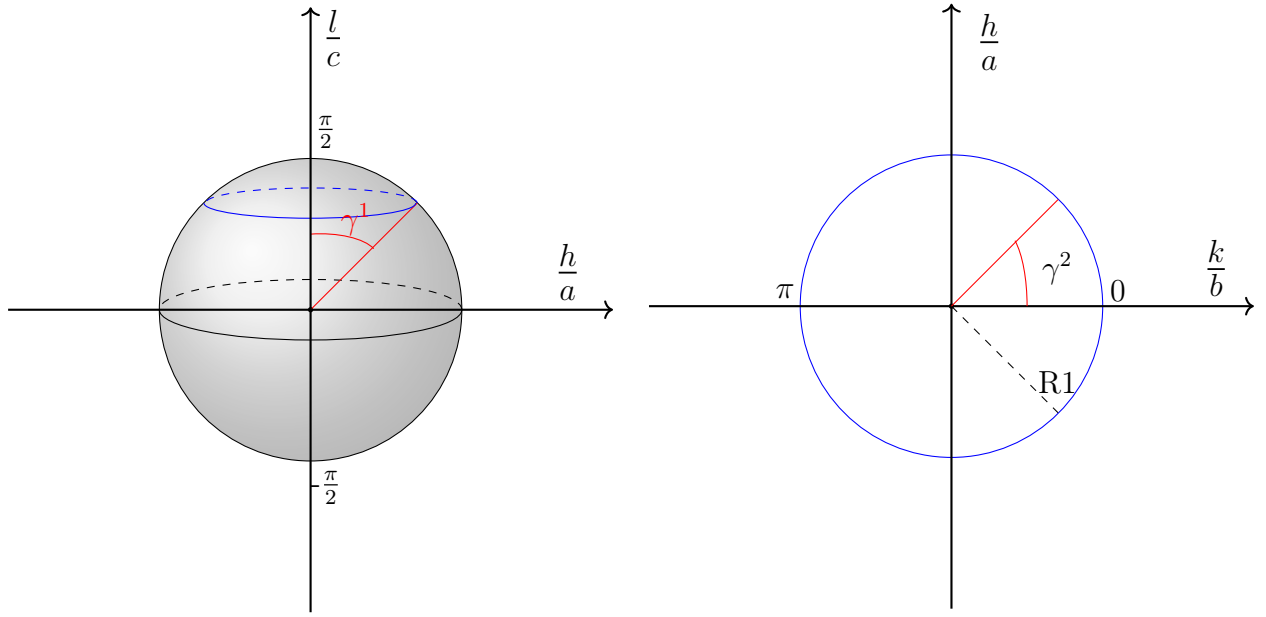
2.4.2 Seconde approche : une simulation précise et complète

La force de cette méthode est qu'elle permet de traiter n'importe quel cristal de manière précise. Elle prend en compte les orientations préférentielles des cristaux, et peut gérer des défauts d'empilement tels que le turbostratisme. C'est le cas, par exemple, des minéraux argileux qui peuvent présenter des empilements de feuillets aléatoires. Dans cette seconde approche, pour chaque angle de diffraction θ , l'intensité diffractée est calculée. Cependant les équations du facteur de structure (2.4), et de la fonction d'interférence (2.6) ne dépendent plus uniquement de l'angle de diffraction. Elles dépendent aussi des coordonnées dans l'espace réciproque (h, k, l) . C'est pourquoi il est nécessaire, pour chaque angle θ , de couvrir l'espace réciproque, en intégrant dans une sphère de cet espace. La sphère et la méthode d'intégration sont expliquées à la figure 2.12. Le pseudo code décrit à l'algorithme 2 expose le processus de calcul de manière simplifiée. Pour gérer des effets comme le turbostratisme, cela s'effectue lors du calcul de la fonction d'interférence (dans la direction $\vec{c}$). Les détails et équations sont données par BISH & POST (1990) dans le livre *Modern Powder diffraction*.

Cette approche aboutit évidemment à une plus grande précision, mais les temps de calculs sont plus conséquents. Pour 3 000 angles de diffraction, et en considérant un pas de 1 000 sur γ_1 et γ_2 pour l'intégration de la sphère, on a 3×10^9 boucles où il faut calculer les coordonnées de l'espace réciproque, le facteur de structure, et la fonction d'interférence. Une majorité de ces calculs est en plus effectuée dans le vide, car les pics de diffraction ne sont que ponctuels. Au final, dans des conditions équivalentes à l'algorithme précédent, les temps de calculs pour une simulation varient de quelques minutes à une heure selon la complexité de la maille.

La solution afin d'accélérer ces temps de calculs conséquents pour générer plusieurs milliers de données a été de combiner plusieurs langages de programmation. Tout d'abord Python pour la lecture du fichier CIF, et l'écriture des sorties. Puis, le Fortran est utilisé pour effectuer l'ensemble des boucles et bénéficier de la rapidité de calcul de ce langage plus bas niveau. Finalement, en optimisant et en parallélisant les calculs sur 128 cœurs en utilisant un serveur de calcul (serveur LETO de la région Centre²), les temps de calculs deviennent plus raisonnables et varient de quelques secondes à plusieurs minutes selon les cristaux. Par exemple pour le quartz, dont la maille contient un faible nombre d'atomes (une dizaine), le temps de calcul est d'environ cinq secondes. En comparaison,

2. <https://cascimodot.fr/doc/index.html>



(a) Sphère d'intégration de l'espace réciproque

(b) Découpe d'un plan horizontal de la sphère, en bleu sur la figure (a)

FIGURE 2.12 – Méthode d'intégration de l'espace réciproque. Une sphère est utilisée, elle est d'abord parcourue verticalement selon un pas γ^1 (a) qui permet d'intégrer ensuite selon des cercles (b) avec un pas γ^2 .

pour l'Ettringite, qui contient une centaine d'atomes, le temps de calcul est supérieur à deux minutes.

Algorithme 2 : Simulation complète de diffractogrammes

Data : Fichier CIF *file*, Vecteur de la plage angulaire de diffraction souhaitée θ , longueur d'onde λ des rayons X

Result : Diffractogramme $D \in \mathbb{R}^d$

1 Calcul des facteurs de diffusion atomique (équation (2.2))

2 **pour** $i \in \theta$ **faire**

3 Définition de la sphère de l'espace réciproque S

4 **pour** $j \in S$ **faire**

5 Calcul du facteur de structure (équation (2.4))

6 Calcul de la fonction d'interférence pour chaque direction a, b, c (équation (2.6))

7 I = Multiplication des deux termes

8 **fin**

9 Somme sur l'ensemble de la sphère des valeurs de I qui donne l'intensité diffractée pour l'angle θ

10 **fin**

Pour compléter les simulations, et notamment se rapprocher des données de laboratoire, des améliorations principalement liées à l'appareil de mesure ont été ajoutées.

2.4.3 Amélioration des simulations

Dans le cadre de nos travaux, deux termes principaux ont été ajoutés pour améliorer les simulations et reproduire plus finement les données expérimentales. D’abord, la fonction d’onde qui intègre les variations de longueur d’onde des rayons X émis, par exemple $K_{\alpha_1}, K_{\alpha_2}, \dots, K_{\beta_1}, K_{\beta_2}, \dots$. Puis un facteur d’absorption μ qui prend en compte l’absorption des rayons X par la matière.

La fonction d’onde Ce facteur est lié aux rayons X émis par la source. Cette source n’émet pas avec une seule longueur d’onde, mais toute une gamme liée aux émissions (par exemple α ou β) dû au matériau excité. Numériquement, cela va se traduire par le calcul de diffractogrammes à différentes longueurs d’ondes avec des contributions spécifiques. Pour obtenir le diffractogramme final, il restera à faire la combinaison linéaire de tous les signaux, les coefficients de la combinaison étant les contributions des longueurs d’onde. Intégrer cette fonction nécessite donc un temps conséquent, puisque le temps de calcul est quasiment multiplié par le nombre de longueurs d’ondes que l’on considère. En effet le facteur de structure et la fonction d’interférence sont dépendants de λ , la longueur d’onde considérée. Ce qui nécessite de les calculer à chaque itération de λ . Par contre, le facteur de diffusion atomique est indépendant de λ , et les calculs peuvent être effectués une fois seulement. Cependant comme le montre la figure 2.13a le signal expérimental est mieux reproduit si l’on compare via l’erreur des moindres carrés avec des données réelles. Ici, la comparaison se fait pour quatre phases minérales que nous présenterons plus en détails dans la section suivante 4.2.

C’est notamment pour les diffractogrammes de laboratoire du chapitre 5 que la fonction d’onde est importante. Grâce à une description de l’architecture de l’appareil de mesure, le logiciel Profex (DOEBELIN & KLEEGERG 2015) est capable de retrouver les longueurs d’ondes et les contributions associées. Le tableau 2.1 donne ces valeurs pour le diffractomètre de laboratoire décrit dans la section 2.2.1.

TABLE 2.1 – Fonction d’onde : Longueurs d’ondes et contributions calculés par Profex pour le diffractomètre décrit dans la section 2.2.1.

Longueurs d’ondes	1.5348	1.5406	1.5411	1.5444	1.5447	1.3923
Contributions	0.01586	0.56768	0.07601	0.25107	0.08688	0.00249

Le facteur d’absorption Ce facteur est lié à la diffusion non élastique des rayons X par les atomes. Ce paramètre n’influence pas directement la forme du diffractogramme, à savoir la position des pics et les rapports d’intensité entre eux. Seule l’intensité réelle des pics de diffraction change.

Cela va être important notamment lorsque l’on souhaitera faire des mélanges de plusieurs phases minérales pour générer des diffractogrammes polyphasiques. En effet, avec les calculs introduits précédemment, notamment le facteur de diffusion atomique et les atomes présents dans la maille, on s’attend à avoir des intensités maximales (ou des aires sous les pics) différentes d’un cristal à un autre. Donc, si les rapports d’intensité entre les phases minérales ne sont pas les bons, les mélanges obtenus ne seront pas conformes à la réalité. L’importance de l’influence de ce phénomène est illustré sur la figure 2.13b, où l’on

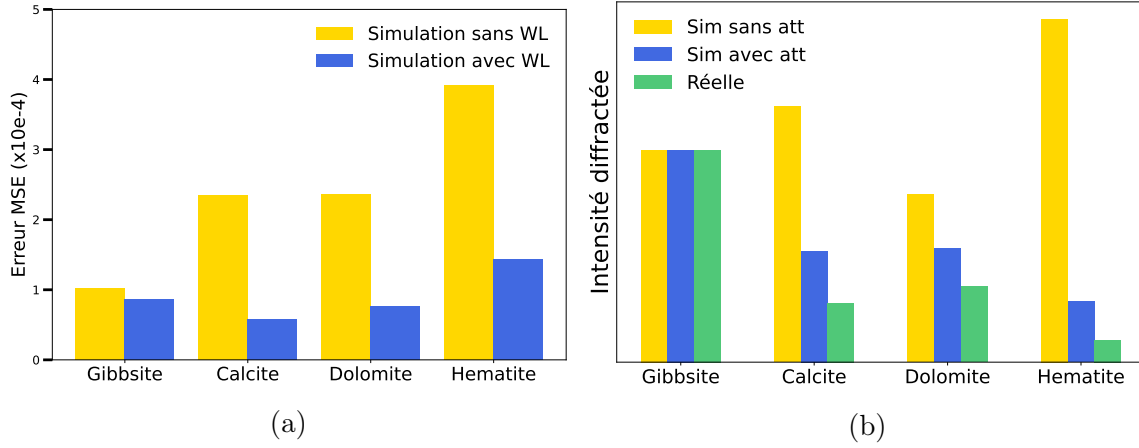


FIGURE 2.13 – (a) Erreur des moindres carrés entre les données de laboratoire et les simulations. Comparaison entre les simulations sans la fonction d’onde (WL, en jaune) et avec cette même fonction (en bleu). (b) Comparaison des intensités maximales avec l’effet du facteur d’absorption. En jaune sans ce facteur, en bleu avec, et en vert l’intensité maximale des données réelles.

observe les rapports des intensités maximales pour quatre phases minérales, en comparant avec des données de laboratoire, des simulations avec et sans ce facteur d’absorption.

D’autres facteurs tels que la fonction instrumentale (dépendante de l’architecture du diffractomètre) ou encore le bruit n’ont pas été pris en compte.

2.4.4 Résultats des simulations et constitution de bases de données

La figure 2.14 présente le résultat de ces simulations sur deux phases minérales (Calcite et Dolomite) et les compare avec des données acquises en laboratoire. Les deux méthodes de simulations sont illustrées.

La première ligne correspond à des données expérimentales collectées au BRGM, comme décrites dans la section 2.2.1. La deuxième ligne est associée au code simplifié, les meilleurs paramètres ayant été choisis pour faire des simulations les plus proches possibles de la réalité. C’est-à-dire que les paramètres de maille (a, b, c) ont été affinés pour correspondre à la localisation des pics, la largeur à mi hauteur et l’agitation thermique ajustées pour minimiser l’erreur des moindres carrés entre la simulation et la donnée expérimentale.

Enfin, la dernière ligne correspond au code complet. Pour la figure 2.14, les données ont été choisies parmi celles simulées pour construire la base de données. Nous nous limitons à ces données, même si la comparaison avec la donnée réelle n’est pas optimale. En effet les temps de calcul sont conséquents et cela requiert l’utilisation du serveur LETO de la région Centre. A noter que l’intensité à bas angle de diffraction (2θ entre 4° et 12°) est importante, dû à la diffusion des rayons X.

Grâce à ces deux codes de simulations, nous sommes capables de générer des bases de données pour alimenter des méthodes de *machine learning* et de *deep learning*. Ce sont ces méthodes que nous utiliserons par la suite, notamment les réseaux de neurones qui seront introduits dans le chapitre suivant. Les bases de données générées peuvent contenir un grand nombre de données d’une grande variété. En effet, selon les paramètres

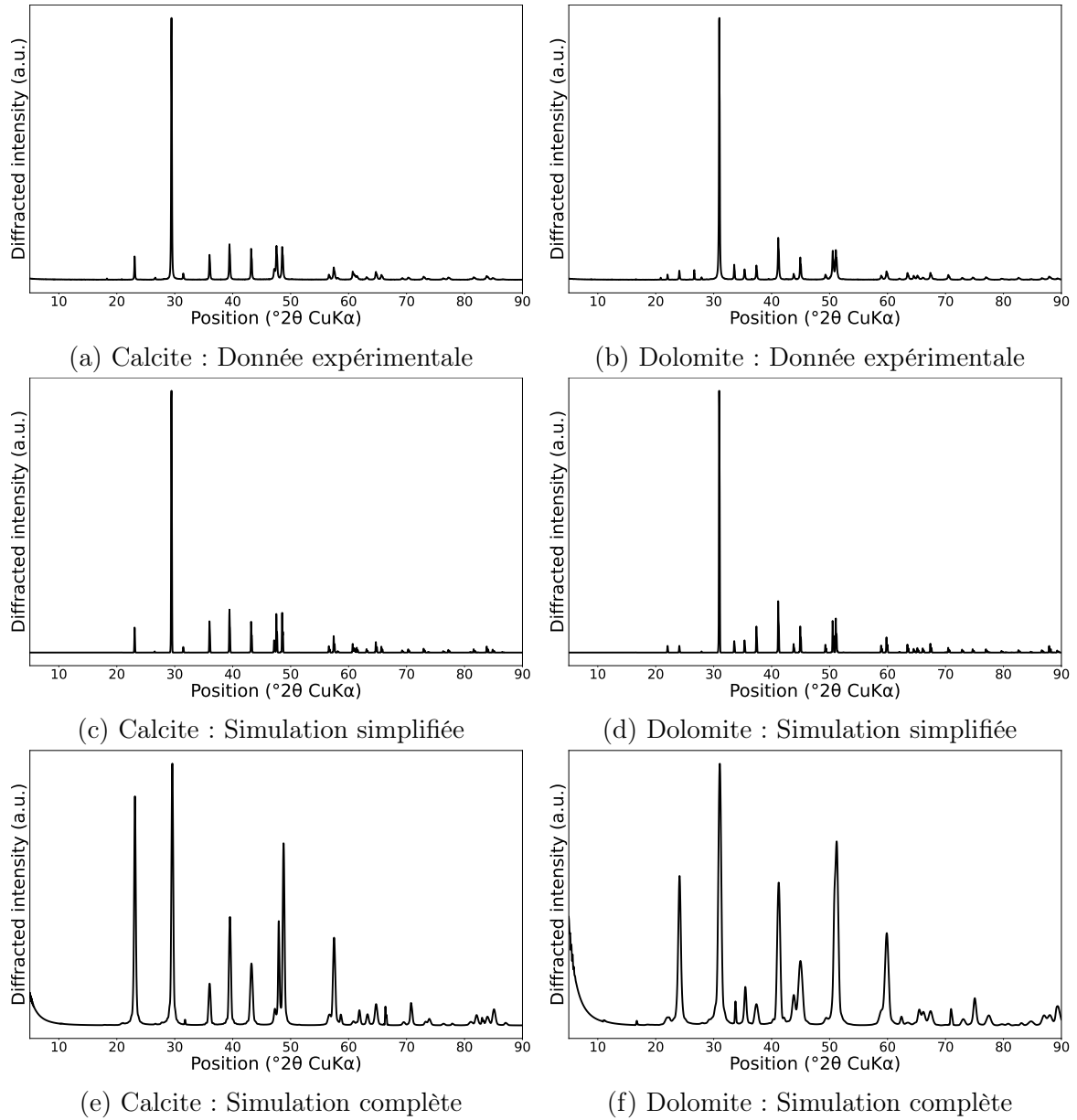


FIGURE 2.14 – Illustration des simulations pour deux phases minérales : la Calcite et la Dolomite. (a) et (b) correspondent à des données expérimentales collectées avec un diffractomètre au BRGM. (c) et (d) sont des simulations effectuées avec le code rapide (algorithme 2.4.1). Et (e) et (f) sont les simulations avec le code complet (algorithme 2.4.2)

de maille, l'agitation thermique, ou encore l'orientation du cristal, une grande gamme de diffractogrammes peut être créée. Cela permet de faire face à tout type de données pouvant être rencontrées sur le terrain, ainsi qu'aux différents appareils de mesure.

Chapitre 3

Les réseaux de neurones pour l'identification de phases minérales

3.1 Les réseaux de neurones artificiels

Dans le cadre de nos travaux pour le traitement des signaux de diffraction, le choix s'est porté sur les approches par réseaux de neurones. Ce type d'approche permet d'apprendre une tâche, en cherchant à imiter le fonctionnement du cerveau humain. Les tâches réalisables par un réseau de neurones sont diverses : segmentation (RONNEBERGER *et al.* 2015), transfert de style (GATYS *et al.* 2015), super-résolution (C. DONG *et al.* 2015), etc. Dans ce chapitre, nous traiterons principalement la tâche de classification, c'est-à-dire être capable, à partir d'une donnée, de retrouver la classe associée. Dans l'exemple de la DRX, l'objectif sera de retrouver la phase minérale à partir du diffractogramme.

Ces méthodes font partie des approches d'apprentissage automatique (BISHOP & NASRABADI 2006) (*Machine Learning*) et plus spécifiquement de l'apprentissage profond (GOODFELLOW *et al.* 2016). Leur utilisation est devenue commune puisqu'elles surpassent la majorité des méthodes traditionnelles pour un grand nombre de tâches, et leur développement ces dernières années a rendu leur emploi courant. Les applications sont diverses, de la détection de pathologie comme l'arthrose du genou (BRAHIM *et al.* 2019), à la génération d'images avec les *Generative Adversarial Networks* (GOODFELLOW *et al.* 2020) ou encore l'identification du genre musical d'une musique (CHOI *et al.* 2017). L'analyse des diffractogrammes n'a pas échappé à cette utilisation massive. En effet, un grand nombre de travaux a été réalisé pour, par exemple classer selon les phases minérales, les groupes d'espace ou encore la symétrie d'un échantillon (OVIEDO *et al.* 2019 ; PARK *et al.* 2017 ; VECSEI *et al.* 2019 ; ZALOGA *et al.* 2020).

Le but de ce chapitre est d'introduire les réseaux de neurones et de comprendre leur fonctionnement à travers l'exemple de la classification de diffractogrammes. Le fonctionnement simplifié d'un réseau de neurones peut se résumer en quelques lignes. L'objectif est de faire apprendre une tâche (de classification par exemple) à un **modèle** neuronal qui possède des paramètres (ou poids) Θ . C'est grâce à des **données** d'apprentissage labellisées que le modèle va apprendre la valeur des poids et sera capable d'apprendre à classer. Pour ces données labellisées, la classe de la donnée est connue, par exemple, pour un diffractogramme, la phase minérale associée est connue. Le modèle s'entraîne donc sur cette base de données en ajustant ses paramètres Θ pour améliorer ses performances lors de l'apprentissage. L'ensemble est réalisé en minimisant les erreurs entre la prédiction et la vérité grâce à une **fonction de perte** qui rétropropage cette erreur dans le modèle.

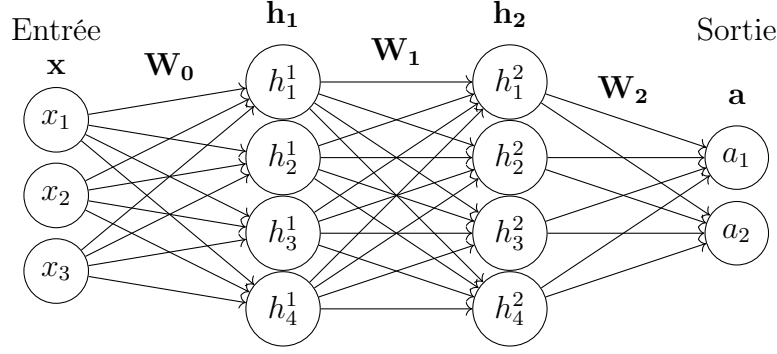


FIGURE 3.1 – Représentation schématique de l'architecture d'un réseau de neurones (FCNN), avec un vecteur d'entrée $\mathbf{x}$, deux couches cachées $\mathbf{h}_1$ et $\mathbf{h}_2$ et la sortie $\mathbf{a}$.

Pour résumer, ces approches reposent donc sur plusieurs éléments principaux : un modèle, des données, une fonction de coût (ou fonction de perte, ou encore *loss* en anglais) qui permet d'ajuster les poids au cours de l'apprentissage et un algorithme d'optimisation pour la minimisation de la fonction de coût.

3.1.1 Le modèle

On peut décrire le modèle en s'appuyant sur la figure 3.1 ci-dessus. Le type de modèle décrit par cette figure est un FCNN (pour *Fully Connected Neural Network*). Un modèle est une transformation f de la donnée d'entrée $\mathbf{x}$ pour obtenir en sortie un vecteur $\mathbf{a}$, que l'on transforme ensuite avec la fonction Softmax pour obtenir $\mathbf{y}$. De telle sorte que $\mathbf{y} = \text{Softmax}(\mathbf{a}) = \text{Softmax}(f(\mathbf{x}|\Theta))$ où Θ représente l'ensemble des paramètres du modèle. La fonction Softmax est définie pour chaque $i = 1, \dots, K$ comme :

$$\text{Softmax}(\mathbf{y}_i) = \frac{\exp(y_i)}{\sum_{j=1}^K \exp(y_j)}. \quad (3.1)$$

Cette fonction Softmax permet d'obtenir un vecteur de probabilité d'appartenance à chacune des classes. Au final, la classe prédite pour la donnée d'entrée $\mathbf{x}$ est obtenue en regardant la probabilité la plus importante, c'est-à-dire en regardant $\text{argmax}_k \mathbf{y}$. En prenant l'exemple de la figure 3.1, y_1 (resp. y_2) est la probabilité que $\mathbf{x}$ soit de la classe 1 (resp. 2). La classe choisie sera alors celle où la valeur de $\mathbf{y}$ est la plus grande.

La transformation f se fait généralement en plusieurs étapes, i.e couches cachées, notées $\mathbf{h}_1$ et $\mathbf{h}_2$ sur la figure 3.1. Plus le nombre de couches cachées est important, plus le réseau est dit profond, d'où le terme d'apprentissage profond.

Opération linéaire Les couches cachées peuvent être de différentes natures. Dans le cas le plus simple, celui d'un Perceptron multi-couches (MLP) (RUMELHART *et al.* 1986), l'opération pour passer d'une couche à la suivante est affine (mais généralement appelée couche linéaire). En notant W_i la matrice des poids associée et b_i la matrice de biais, on obtient l'équation suivante entre la couche i et $i + 1$:

$$h_{i+1} = g(W_i h_i + b_i). \quad (3.2)$$

Cette opération est complétée par une fonction d'activation non linéaire $g : \mathbb{R} \rightarrow \mathbb{R}$, qui est appliquée à chaque coordonnée de $\mathbf{h}_i$. Parmi les plus courantes on retrouve la

fonction ReLU ou encore la fonction sigmoid, qui sont respectivement définies par les équations (3.3) et (3.4). La fonction d'activation permet de casser la linéarité du modèle. Sans elle, le modèle se limiterait à une opération linéaire.

$$\text{ReLU}(h) = \max(0, h), \quad \text{avec } h \in \mathbb{R}. \quad (3.3)$$

$$\text{sigmoid}(h) = \frac{1}{1 + \exp^{-h}}, \quad \text{avec } h \in \mathbb{R}. \quad (3.4)$$

Convolution Dans nos travaux, nous utilisons des réseaux de neurones convolutifs (CNN) (LECUN *et al.* 1998). Les opérations entre les couches restent linéaires, mais restreintes à être des convolutions (voir figure 3.2). Pour une entrée $\mathbf{x}$ et un noyau de convolution $\mathbf{w}$ de taille M , la valeur de la sortie $\mathbf{y}$ à l'indice t est donnée par l'équation :

$$y(t) = (\mathbf{x} * \mathbf{w})(t) = \sum_{k=0}^M x(t - k)w(k) + b. \quad (3.5)$$

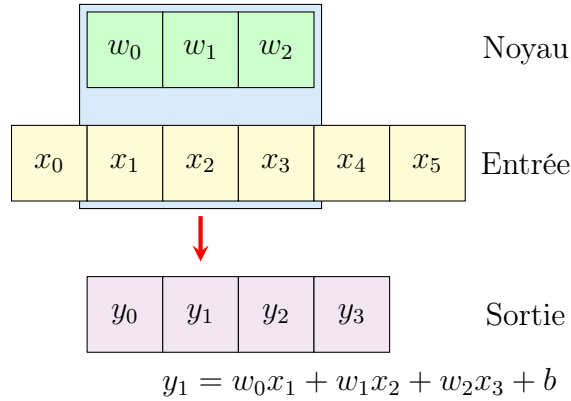


FIGURE 3.2 – Convolution pour un vecteur $\mathbf{x}$ par un vecteur de poids $\mathbf{w}$ de taille 3, le résultat est $\mathbf{y}$.

Ce type d'opération, très populaire en traitement d'image, permet d'extraire les caractéristiques principales des données, tout en réduisant le nombre de paramètres Θ du modèle. Plusieurs paramètres ajustables sont à prendre en compte pour une convolution, comme la taille du noyau de convolution, le pas entre chaque application (*stride*) et le nombre de canaux de la couche suivante. Tous ces paramètres feront varier de manière importante la taille de la sortie. Une fonction d'activation après l'opération convolutive est également toujours appliquée.

D'autres types d'opérations complètent ces couches comme les couches de pooling, de dropout ou de normalisation. La plupart d'entre-elles permettent d'ajuster la taille de la donnée à une couche donnée, ou encore de réduire la taille de certaines couches pour avoir un nombre de paramètres plus faible. Le vecteur Θ contient en général plusieurs millions de paramètres qui sont ajustés à partir des données d'entraînement.

3.1.2 Les données

Pour apprendre sa tâche de classification, le réseau va, généralement, utiliser un ensemble de données labellisées, noté $\mathcal{T}^r$. Cet ensemble contient $N_{\mathcal{T}^r}$ couples $(\mathbf{x}_i, \mathbf{l}_i)$,

$i = 1, \dots, N_{\mathcal{T}^r}$, où $\mathbf{x}_i$ est la donnée, et $\mathbf{l}_i$ la classe associée. Les vecteurs labels $\mathbf{l}_i$ sont des vecteurs indicateurs, dont tous les éléments valent 0, sauf celui qui indique la classe qui vaut 1. On parle alors d'apprentissage supervisé, car il est nécessaire que les données soient labellisées pour que le réseau apprenne à les trier. Cependant, l'apprentissage non supervisé est possible dans certains contextes, par exemple avec les *Generative Adversarial Networks* (GOODFELLOW *et al.* 2020), ou les Auto-Encodeur (AE) (HINTON & SALAKHUTDINOV 2006).

Afin de s'assurer que le réseau n'apprenne pas uniquement ces données, deux autres ensembles sont nécessaires. Un ensemble de validation $\mathcal{V}$ utilisé lors de l'entraînement du réseau pour tester les performances, on note $N_{\mathcal{V}}$ son cardinal. On considère également un ensemble de test $\mathcal{T}$ contenant $N_{\mathcal{T}}$ données. Il permet de tester le modèle à la fin de l'entraînement. Cet ensemble est uniquement utilisé une fois le modèle entraîné et ne doit pas être utilisé lors de l'entraînement. Pour éviter un effet de sur-apprentissage, il est essentiel que les trois ensembles soient indépendants. Dans le cas contraire, on aurait un réseau performant sur les données d'entraînement mais peu robuste sur les autres données. Il est donc important de construire une bonne base de données $\mathcal{T}^r$ pour que le modèle apprenne correctement.

Lors de l'entraînement, le réseau parcourt plusieurs fois l'ensemble des données pour apprendre au mieux. Lorsque le modèle parcourt une fois la base d'apprentissage cela correspond à ce que l'on appelle une **époque**. Généralement, les réseaux sont entraînés sur plusieurs centaines, voire milliers d'époques.

Le dernier élément à introduire est la fonction de coût. C'est grâce à elle que le réseau va pouvoir ajuster les paramètres Θ lors de la phase d'entraînement.

3.1.3 La fonction de coût

Le rôle de cette fonction, notée $\mathcal{L}$, pour “*loss*”, est de faire en sorte que le réseau apprenne la tâche qui lui est assignée au mieux à partir des données. La sortie du réseau doit permettre de désigner la classe associée à la donnée, en maximisant la vraisemblance d'une loi multinomiale, qui est paramétrée par la sortie du réseau $\mathbf{y}$. Cette fonction doit donc lier la sortie du réseau $\mathbf{y} = f(\mathbf{x}|\Theta)$ et le vecteur label $\mathbf{l}$, par exemple en minimisant la différence entre eux. La phase d'entraînement consiste alors à optimiser la fonction $\mathcal{L}$ pour trouver Θ' tel que :

$$\Theta' = \operatorname{argmin} \mathcal{L}(f(\mathbf{x}|\Theta), \mathbf{y}). \quad (3.6)$$

Dans la suite, nous noterons $\mathcal{L}$ la fonction de coût globale (sur l'ensemble des données de la base d'entraînement), $\mathcal{L}_i$ la fonction associée à une unique donnée i et $\mathcal{L}_b$ la fonction associée à un *batch* b .

Un exemple assez commun de fonction de coût est l'erreur des moindres carrés (MSE), définie pour une donnée i par :

$$\mathcal{L}_i^{MSE} = \|\mathbf{y}_i - \mathbf{l}_i\|^2. \quad (3.7)$$

La fonction de coût global est alors la somme sur l'ensemble d'apprentissage des $\mathcal{L}_i^{MSE}$:

$$\mathcal{L}^{MSE} = \sum_{i=1}^{N_{\mathcal{T}^r}} \mathcal{L}_i^{MSE}. \quad (3.8)$$

Un autre exemple de fonction de coût, que nous utiliserons par la suite est l'entropie croisée (CE). Elle est définie par :

$$\mathcal{L}_i^{CE} = - \sum_j l_{ij} \log(y_{ij}). \quad (3.9)$$

L'optimisation de ces fonctions se fait en utilisant des algorithmes de descente de gradient stochastique comme l'algorithme Adam (KINGMA & BA 2014) que l'on utilisera par la suite. Ils utilisent la rétropropagation (WERBOS 1990) pour mettre à jour les poids Θ du réseau en optimisant la fonction $\mathcal{L}$.

Descente de gradient stochastique Pour introduire le fonctionnement de ces algorithmes, nous commençons par présenter la descente de gradient standard. Cet algorithme permet un ajustement des poids du modèle en minimisant la fonction de perte. A chaque itération (mise à jour des poids Θ), il calcule le gradient de la fonction de perte $\mathcal{L}$ par rapport à tous les poids du modèle. Un taux d'apprentissage (*learning rate*) est utilisé comme paramètre de ces algorithmes, et donne le pas de la descente de gradient lors de l'optimisation de la fonction de coût. La mise à jour des poids est faite comme suit :

$$\Theta \leftarrow \Theta - \eta \nabla_{\Theta} \mathcal{L}(\Theta), \quad (3.10)$$

où η est le taux d'apprentissage (*learning rate*) et $\nabla_{\Theta} \mathcal{L}$ le gradient de la fonction de perte par rapport aux poids Θ .

Dans le cas de la descente de gradient standard, on utilise toutes les données de la base d'entraînement à chaque itération pour mettre à jour les poids. Cependant cela pose notamment des problèmes de mémoire lorsque que les bases d'entraînement contiennent un grand nombre de données, ce qui est généralement le cas.

La stratégie pour éviter ces limites est d'utiliser un algorithme de **descente de gradient stochastique**. L'ajustement des poids se fera alors sur des *batches* de données, c'est-à-dire en plusieurs petits sous-ensembles. La mise à jour des poids devient alors :

$$\Theta \leftarrow \Theta - \eta \nabla_{\Theta} \mathcal{L}_b(\Theta), \quad (3.11)$$

où $\mathcal{L}_b$ est la fonction de perte associée à un *batch* b .

Les *batches* sont constitués de sous-ensembles de données choisies aléatoirement. A chaque itération, les gradients sont donc légèrement différents car les *batches* considérés sont différents. Cette variabilité introduit la composante stochastique dans la mise à jour des paramètres. Cette stratégie présente plusieurs avantages, en commençant par la vitesse du temps de calcul. Le regroupement par *batch* permet de paralléliser les calculs sur GPU, et d'accélérer le temps d'entraînement. Ensuite, pour offrir un meilleur apprentissage et plus de stabilité, le traitement des données individuelles entraînerait du sur-apprentissage et des problèmes de convergence. Les poids Θ sont alors mis à jour après le passage de chaque *batch* dans le réseau, grâce aux algorithmes évoqués précédemment. Il n'est pas non plus possible de traiter l'ensemble de la base d'apprentissage car cela poserait des problèmes de mémoire.

L'algorithme Adam que l'on utilisera par la suite est une amélioration de la descente de gradient stochastique. Il utilise un fonctionnement similaire et intègre des concepts supplémentaires permettant une optimisation plus stable et plus rapide. Cela est possible notamment grâce à l'adaptation des taux d'apprentissage qui sont calculés à partir d'estimations des premiers et seconds moments des gradients (KINGMA & BA 2014).

3.2 Identification de phases

Nous allons dans un premier temps tester l'utilisation de réseaux de neurones pour le traitement et l'analyse des signaux de diffraction, sur un problème simple d'identification de phases. Ce problème avec des phases peu nombreuses permet une approche simplifiée. Ainsi, en plaçant en entrée du réseau un signal $\mathbf{x}$ à phase unique (i.e qui correspond une unique phase minérale), l'objectif est de retrouver sa phase minérale parmi K possibilités.

Nous avons fait le choix de traiter un problème avec $K = 6$ classes (i.e phases minérales) : Calcite (MARKGRAF & REEDER 1985) (CaCO_3), Dolomite (STEINFINK & SANS 1959) (CaMgC_2O_6), Gibbsite (BALAN *et al.* 2006) (AlO_3H_3), Hematite (BLAKE *et al.* 1966) (Fe_2O_3), Halite (WALKER *et al.* 2004) (NaCl) et Quartz (LEVIEN *et al.* 1980) (SiO_2). Ces phases ont été choisies car elles sont centro-symétriques, ce qui permet d'utiliser le code de simulation rapide (algorithme 1) présenté dans la section 2.4 du chapitre précédent. Aussi, la proximité en 2θ des pics de diffraction maximums de certaines phases permet de complexifier le problème auquel nous sommes confrontés (voir figure 3.3).

On introduit aussi de la variabilité au sein de chaque classe. D'abord en autorisant une variation de $\pm 5\%$ des paramètres de maille a , b et c , ensuite en jouant sur l'agitation thermique (i.e l'effet de Debye-Waller), où le coefficient B présenté précédemment dans la section 2.3.1 varie entre 0.1 et 2. Enfin on joue sur la largeur à mi-hauteur des pics, que l'on fait varier entre 0.01 et 0.2 grâce à la distribution gaussienne du code rapide décrit par l'algorithme 1. Tous les paramètres sont tirés de manière uniforme sur les intervalles décrits, les valeurs des amplitudes de variation sont données en annexe par le tableau A.1. Finalement, le réseau est entraîné sur 1800 données, regroupées en *batches* de 4 diffractogrammes.

Nous ajoutons que l'intensité réelle ne donne pas d'information pour notre problématique. La position et les ratios d'intensité des pics sont suffisants pour identifier et quantifier les phases minérales. Afin de considérer uniquement l'intensité relative, tous les signaux sont normalisés de manière à ce que la valeur maximale soit égale à 1.

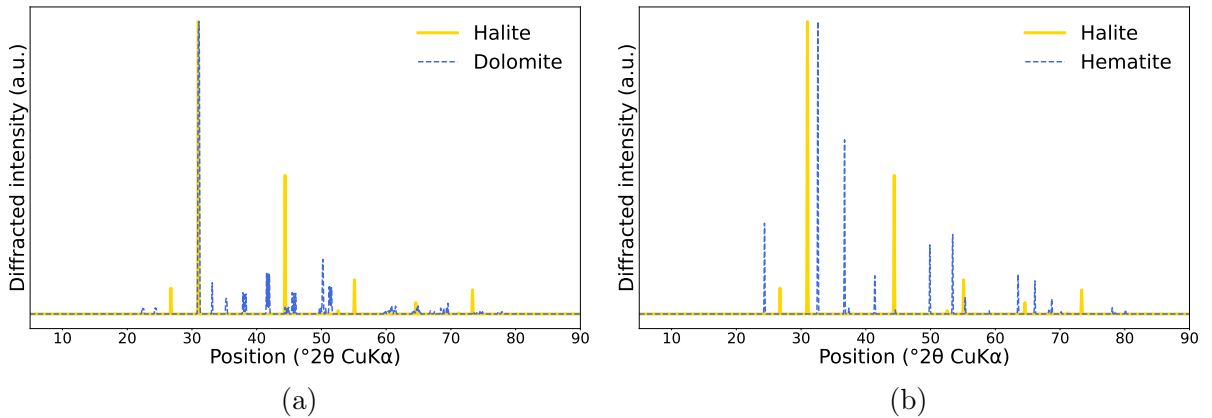


FIGURE 3.3 – Comparaison de simulations pour deux couples de phases différents (a) Halite et Dolomite et (b) Halite et Hématite.

Deux architectures de réseaux sont testées. D'abord un CNN qui est décrit par la figure 3.4, inspiré de travaux pour la classification de diffractogramme (OVIDEO *et al.* 2019). Il a été adapté à notre problème pour l'identification de phases minérales. Le réseau commence par trois couches de convolution, de noyau 8, 5 et 3 respectivement, chacune avec une valeur de *stride* identique à celle du noyau. Il compte ensuite deux

couches linéaires pour réduire la taille du vecteur de sortie, pour un total de 834 918 paramètres Θ à optimiser durant 100 époques d'entraînement. Le vecteur de sortie est de taille $K = 6$, qui est aussi le nombre de classes. La fonction d'activation est ReLU, et l'algorithme d'optimisation est Adam, avec un learning rate constant égal dont la valeur est 0.001.

Ensuite, un réseau de neurones entièrement composé de couches linéaires (FCNN) est aussi utilisé. Ce type de modèle est également un MLP, dont le nom a été évoqué précédemment dans ce chapitre. Le réseau est composé de six couches linéaires qui permettent de réduire la taille du vecteur d'entrée jusqu'à la dimension $K = 6$. La figure 3.5 décrit plus précisément l'architecture. Le désavantage de ce type de modèle est le nombre important de paramètres (ici 8 739 462, soit 10 fois plus que le CNN) nécessaires pour relier chacun des neurones entre deux couches successives, ce qui mène à des temps d'entraînement plus conséquents. Les paramètres d'entraînement (algorithme d'optimisation, taux d'apprentissage, nombre d'époques,) sont identiques à ceux du CNN.

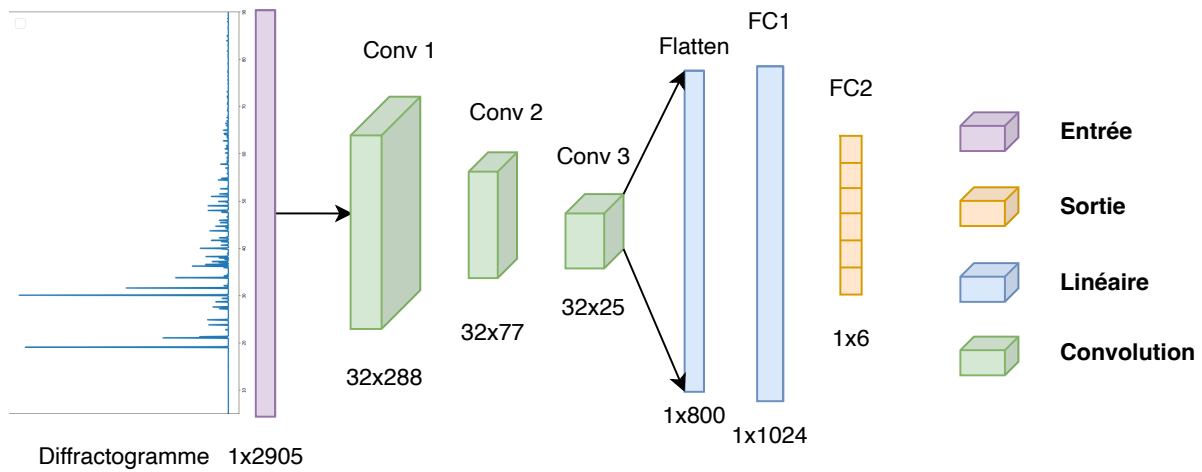


FIGURE 3.4 – Architecture du CNN pour la classification

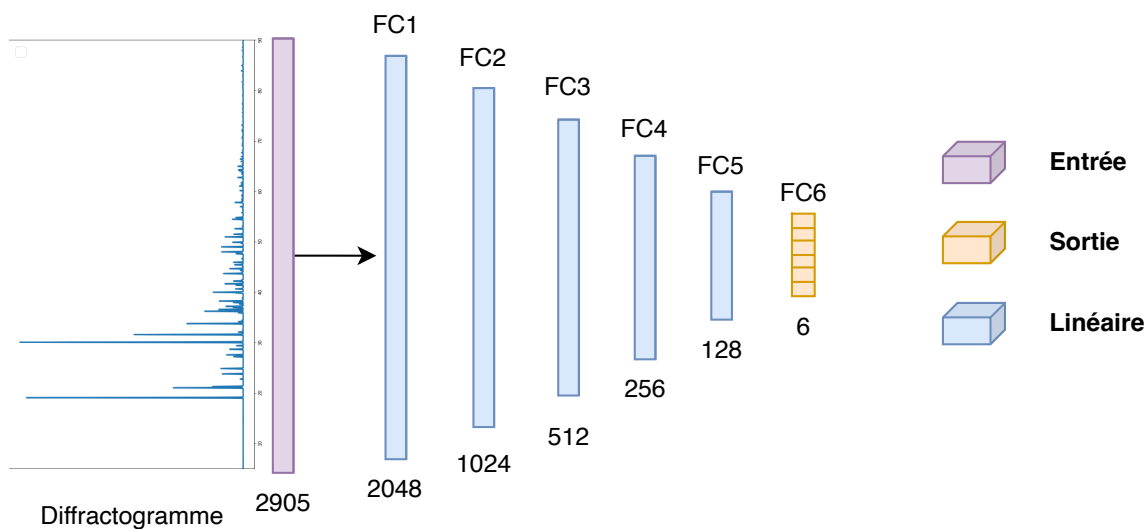


FIGURE 3.5 – Architecture du FCNN pour la classification

Pour cette première expérience, nous avons fait le choix de ne pas considérer de base de

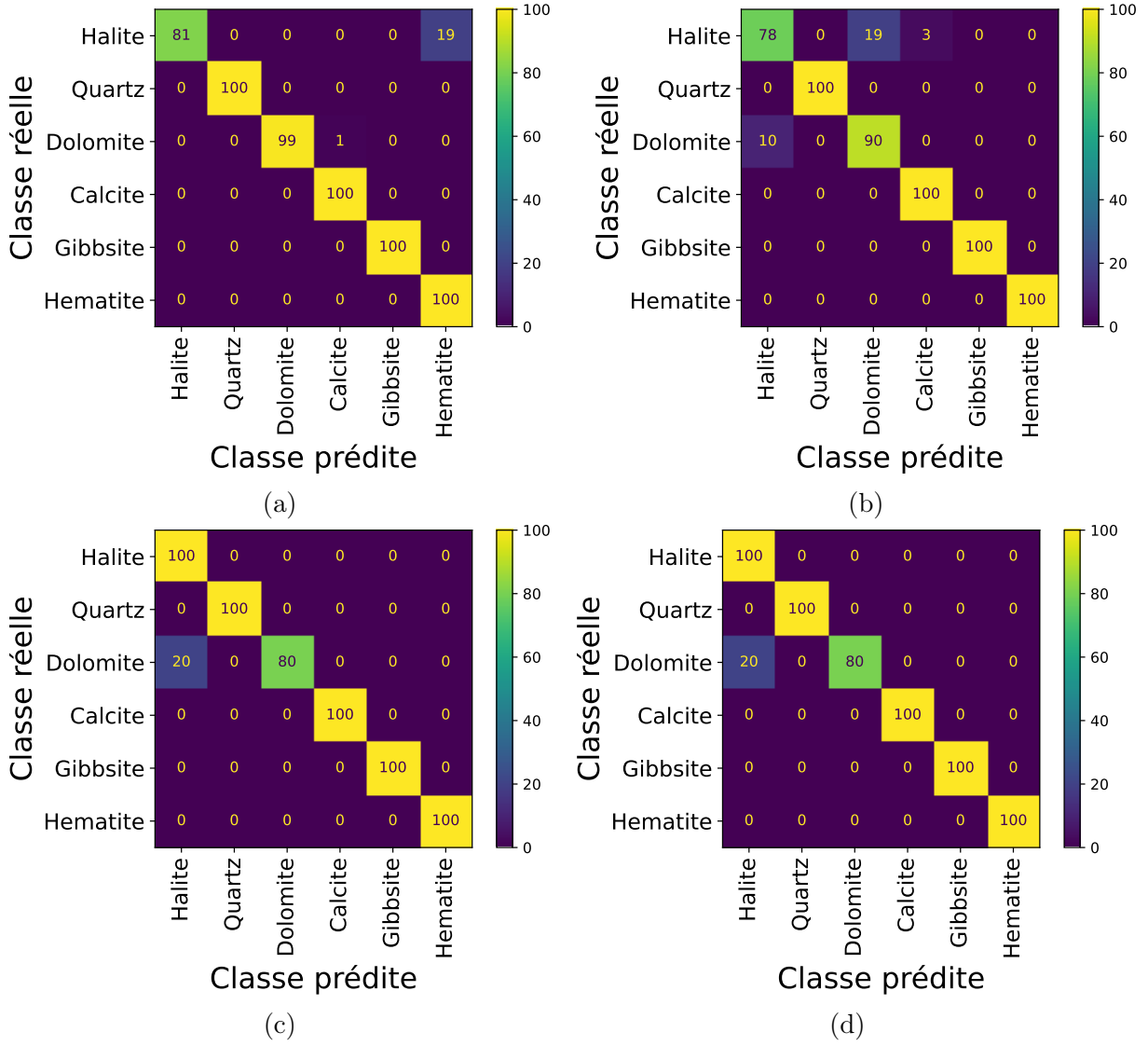


FIGURE 3.6 – Comparaison des matrices de confusion pour les différentes méthodes comparées : (a) Réseaux de neurones convolutifs (CNN), (b) Réseau de neurones avec des couches linéaires (FCNN), (c) Machine à vecteurs de support (SVM) et (d) Régression Logistique.

validation. Les performances sont directement évaluées sur l'ensemble de test contenant 100 diffractogrammes par phase, soit un total de 600 données. Les matrices de confusion des figures 3.6a et 3.6b fournissent un aperçu des résultats qui sont complétés par ceux du tableau 3.1.

Les résultats montrent que le CNN est efficace pour l'analyse des signaux issus de la diffraction, avec des mesures de l'ordre de 97%, et seulement 20 données parmi les 600 dont la classification a échoué. On remarque notamment avec la matrice de confusion que 19 signaux d'Halite ont été classés en Hématite. Ceci est probablement dû à la proximité de certains pics de diffraction communs à ces classes, observable sur la figure 3.3b.

Pour le FCNN, la classification est également très bonne, avec presque 95% des données bien classées. Les erreurs sont principalement dues à un mélange des phases de Dolomite et de Halite. La superposition des diffractogrammes de ces phases (voir figure 3.3a) montrent que les pics d'intensité maximales se confondent, ce qui explique les erreurs du FCNN.

Cette tâche de classification permet également de comparer les réseaux de neurones avec des méthodes de Machine Learning plus traditionnelles comme les SVM ou la régression logistique. Comme ces deux méthodes de classification sont répandues et leur efficacité a été prouvée, il est rassurant de voir que les approches neuronales présentent des erreurs similaires. Les performances de ces trois méthodes sont fournies par le tableau 3.1.

TABLE 3.1 – Résultats de la classification et comparaison des différentes méthodes testées pour l’identification de phases minérales.

Méthode	Paramètres	Accuracy ↑	F1 score ↑	Recall ↑
CNN	834 918	0.97	0.97	0.97
FCNN	8 739 462	0.94	0.95	0.95
SVM linéaire	X	0.97	0.97	0.97
Régression logistique	X	0.97	0.97	0.97

Dans le cadre de cette thèse, la classification n’est qu’une première étape pour se confronter aux données. L’objectif suivant est de traiter un problème plus complexe, celui de la quantification de phases minérales. C’est également la raison pour laquelle l’approche par réseaux de neurones a été choisie. Ses performances sont similaires aux modèles traditionnels de classification (SVM, régression logistique), et elle permet de s’adapter facilement à des problèmes plus complexes avec une modélisation adaptée. De plus, la comparaison des résultats avec les deux réseaux de neurones montre des meilleurs résultats avec le CNN. Pour le problème de quantification de phases, c’est donc cette architecture qui sera utilisée.

Chapitre 4

Inférence de proportions

Le problème de classification décrit dans le chapitre précédent est une première étape pour l'analyse des diffractogrammes. Cependant, en général, les diffractogrammes issus de l'analyse de solide ne correspondent pas à une unique phase minérale, mais plutôt à des mélanges de plusieurs phases qu'il est nécessaire de quantifier. Ce chapitre traite donc du problème d'inférence de proportions. Nous allons considérer trois applications, l'analyse de diffractogrammes, l'imagerie hyperspectrale et une base simulée d'images artificielles.

4.1 Une modélisation de Dirichlet pour l'inférence de proportions

4.1.1 Formulation du problème

Soit $\mathbf{x}$ une donnée (un signal ou une image) que l'on considère comme étant une combinaison linéaire de K composants,

$$\mathbf{x} = \sum_{j=1}^K p_j \mathbf{C}_j, \quad (4.1)$$

où $\mathbf{C}_j$ est le j ème composant ($j = 1, \dots, K$), et p_j sa proportion.

L'objectif de la méthode introduite dans ce chapitre est d'obtenir une estimation $\hat{\mathbf{p}}$ du vecteur de proportion $\mathbf{p}$, étant donné une entrée $\mathbf{x}$. Le vecteur $\mathbf{p}$ étant un vecteur de proportions, il appartient au simplexe de dimension K , noté Δ_K et défini par :

$$\Delta_K = \left\{ \mathbf{x} = (x_1, \dots, x_K) \in \mathbb{R}^K \mid x_j \geq 0, \ j = 1, \dots, K, \text{ and } \sum_{j=1}^K x_j = 1 \right\}. \quad (4.2)$$

Tout vecteur $\mathbf{p} \in \Delta_K$ doit donc respecter à la fois une contrainte de positivité sur ses composantes, mais aussi leur somme à 1.

4.1.2 La distribution de Dirichlet

La méthode que nous proposons pour l'inférence de proportions se base sur la distribution de Dirichlet, dont nous allons maintenant donner les principales caractéristiques.

Un vecteur aléatoire $\mathbf{P} = (P_1, \dots, P_K)$ est distribué suivant la loi de Dirichlet de paramètre $\boldsymbol{\alpha} = (\alpha_1, \dots, \alpha_K)$, $\alpha_i > 0 \forall i$, si sa densité de probabilité est donnée par :

$$f_{\mathbf{P}}^{\text{Dir}}(\mathbf{p}|\boldsymbol{\alpha}) = \frac{1}{\beta(\boldsymbol{\alpha})} \prod_{i=1}^K p_i^{\alpha_i-1}, \quad \text{où} \quad \beta(\boldsymbol{\alpha}) = \frac{\prod_{i=1}^K \Gamma(\alpha_i)}{\Gamma(S_{\boldsymbol{\alpha}})}. \quad (4.3)$$

On notera $\mathbf{P} \sim \text{Dir}(\boldsymbol{\alpha})$. $S_{\boldsymbol{\alpha}}$ est appelée la précision de la loi de Dirichlet et est définie par :

$$S_{\boldsymbol{\alpha}} = \sum_{i=1}^K \alpha_i. \quad (4.4)$$

La fonction gamma d'Euler est définie pour tout $z > 0$ par :

$$\Gamma(z) = \int_0^{+\infty} t^{z-1} \exp(-t) dt. \quad (4.5)$$

La distribution de Dirichlet est définie sur le simplexe Δ_K et sa densité est nulle en dehors. On a les expressions suivantes sur les marginales :

$$\mathbb{E}[P_i] = \frac{\alpha_i}{S_{\boldsymbol{\alpha}}} \quad \text{et} \quad \mathbb{E}[\log P_i] = \psi(\alpha_i) - \psi(S_{\boldsymbol{\alpha}}) \quad (4.6)$$

où ψ est la fonction digamma, la dérivée logarithmique de la fonction Γ :

$$\psi(z) = \frac{\Gamma'(z)}{\Gamma(z)}. \quad (4.7)$$

De même pour la variance on a :

$$\text{Var}(P_i) = \frac{\tilde{\alpha}_i(1 - \tilde{\alpha}_i)}{S_{\boldsymbol{\alpha}} + 1} \quad \text{et} \quad \text{Var}(\mathbf{P}) = \sum_{i=1}^K \text{Var}(P_i) \quad \text{avec} \quad \tilde{\alpha}_i = \mathbb{E}[P_i]. \quad (4.8)$$

L'influence des paramètres sur la variance est importante. Plus la précision $S_{\boldsymbol{\alpha}}$ est grande, plus la variance est faible, et inversement.

Le choix d'utiliser la distribution de Dirichlet pour modéliser les proportions est justifié par le fait que le vecteur $\mathbf{p}$ appartient au simplexe Δ_K . De plus, la modélisation de Dirichlet pour les réseaux de neurones a déjà été utilisée pour quantifier l'incertitude de classification (FALL 2024; SENSOY *et al.* 2018).

4.1.3 Modélisation

La modélisation que nous proposons est une adaptation des approches de classification. Pour rappel, la classification repose sur une maximisation de la vraisemblance d'une loi multinomiale (cf chapitre 3). Pour l'adapter à un problème d'estimation de proportions avec K classes, on voudra cette fois maximiser la vraisemblance d'une distribution de Dirichlet.

Soit $f(\cdot|\boldsymbol{\Theta}) : \mathbb{R}^d \rightarrow \mathbb{R}^K$ un réseau de neurones de paramètres $\boldsymbol{\Theta}$, $\mathbf{y}_i = f(\mathbf{x}_i|\boldsymbol{\Theta})$ la sortie du réseau étant donnée une entrée $\mathbf{x}_i$. Nous cherchons à établir une relation ϕ entre le vecteur de paramètres $\boldsymbol{\alpha}$ et la sortie du réseau $\mathbf{y}_i$, de telle sorte que :

$$\boldsymbol{\alpha} = \phi(\mathbf{y}_i). \quad (4.9)$$

La fonction ϕ doit respecter deux contraintes. La première est la contrainte de stricte positivité sur les α_i . La seconde est que la fonction doit être strictement croissante afin que les grandes valeurs de α_i aient une faible variance. Le choix s’est donc porté sur la fonction ϕ suivante :

$$\phi(\mathbf{y}) = \text{ReLU}(\mathbf{y}) + 1. \quad (4.10)$$

La fonction ReLU garantit d’avoir des valeurs supérieures ou égales à 0. Pour garantir la stricte positivité on rajoute 1. Cela permet aussi de garantir la non-dégénérescence de la loi de Dirichlet.

A chaque donnée d’entrée $\mathbf{x}_i$, on lui associe un vecteur aléatoire $\mathbf{P}_i$ tel que :

$$\mathbf{P}_i = (P_{i1}, \dots, P_{iK}) \sim \text{Dir}(\boldsymbol{\alpha}_i), \quad \text{où} \quad \boldsymbol{\alpha}_i = \phi(\mathbf{y}_i) = \text{ReLU}(f(\mathbf{x}_i|\boldsymbol{\Theta})) + 1. \quad (4.11)$$

La prediction finale $\hat{\mathbf{p}}_i$ pour la proportion $\mathbf{p}_i$ est obtenue en considérant l’espérance de la distribution de Dirichlet :

$$\hat{\mathbf{p}}_i = \frac{\boldsymbol{\alpha}_i}{S_{\boldsymbol{\alpha}_i}} \in \Delta_K. \quad (4.12)$$

Une fois cette modélisation introduite, il reste à entraîner le réseau selon cette modélisation afin de trouver le meilleur vecteur $\boldsymbol{\alpha}$ pour paramétrer la distribution de Dirichlet. On propose d’utiliser des fonctions de perte spécifiques.

Dirichlet et moindres carrés Pour cette première fonction, le terme de minimisation est l’erreur des moindres carrés. Par la suite, elle sera désignée comme la fonction “MSE & DIR”.

Nous considérons l’espérance de l’erreur des moindres carrés entre la variable aléatoire $\mathbf{P}_i$ et la proportion $\mathbf{p}_i$. Comme pour la classification, la fonction de perte totale est obtenue en sommant sur l’ensemble des données d’entraînement. L’entraînement s’appuie sur un ensemble de couples $(\mathbf{x}_i, \mathbf{p}_i)$, regroupant la donnée et le vecteur de proportion associé. Il s’agit donc d’apprentissage supervisé. La fonction de perte est donnée par :

$$\begin{aligned} \mathcal{L}_i^{SE}(\boldsymbol{\Theta}) &= \mathbb{E}(\|\mathbf{p}_i - \mathbf{P}_i\|^2) = \mathbb{E}(\|\mathbf{p}_i - \mathbb{E}[\mathbf{P}_i]\|^2 + \|\mathbf{P}_i - \mathbb{E}[\mathbf{P}_i]\|^2) \\ &= \|\mathbf{p}_i - \mathbb{E}[\mathbf{P}_i]\|^2 + \text{Var}(\mathbf{P}_i) \\ &= \|\mathbf{p}_i - \hat{\mathbf{p}}_i\|^2 + \text{Var}(\mathbf{P}_i) \end{aligned} \quad (4.13)$$

Dirichlet et entropie croisée Une alternative est d’utiliser l’entropie croisée à la place de l’erreur des moindres carrés :

$$\mathcal{L}_i^{CE}(\boldsymbol{\Theta}) = \mathbb{E} \left(\sum_{j=1}^K -p_{ij} \log(P_{ij}) \right) = \sum_{j=1}^K -p_{ij} [\psi(\alpha_i) - \psi(S_{\alpha_i})]. \quad (4.14)$$

Pour ces deux fonctions, il faut cependant noter l’élément suivant. Comme pour la classification, l’entraînement du réseau consiste à optimiser la fonction de perte, c’est-à-dire à chercher le $\boldsymbol{\Theta}'$ la minimisant. Or dans le cadre de cette modélisation, les deux fonctions considérées (équations (4.13) et (4.14)) ne sont pas convexes par rapport à la sortie du réseau et donc au vecteur $\hat{\mathbf{p}}_i$. La figure 4.1 en donne des représentations en 3D. Cependant, nous n’avons pas observé de problèmes de convergence de l’algorithme Adam (KINGMA & BA 2014) liés spécifiquement à ces fonctions.

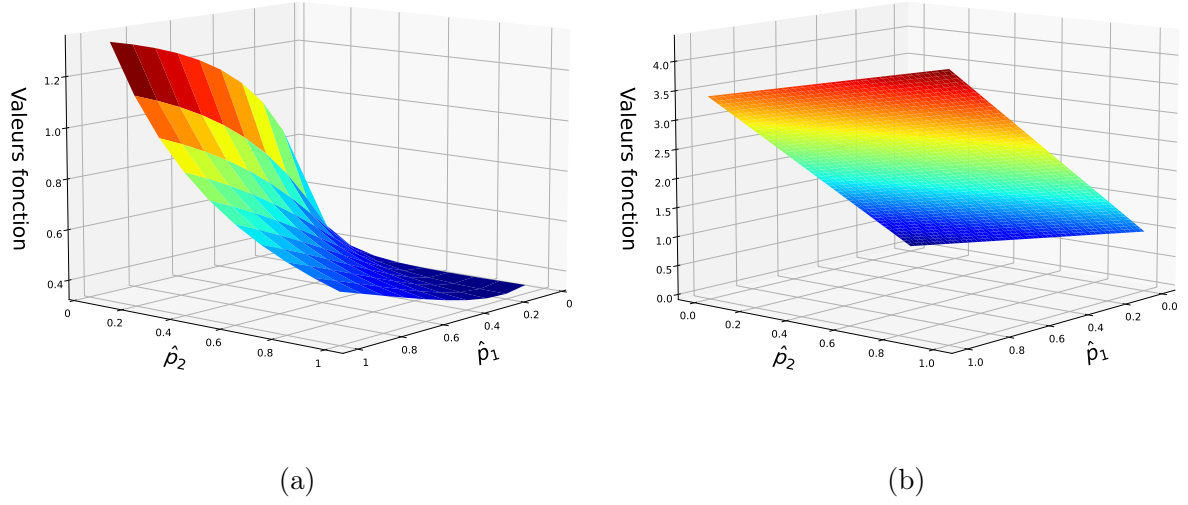


FIGURE 4.1 – Exemple de tracé des fonctions de pertes “MSE & DIR” (a) et “CE & DIR” (b). Pour les deux illustrations, $\mathbf{p} = (0.2, 0.8)$ et le tracé est effectué selon les variations de $\hat{p}_1$ et $\hat{p}_2$.

Pour conclure, la méthodologie proposée se base uniquement sur une adaptation de la fonction de coût. C’est la force principale de notre approche. Elle utilise des architectures de réseau déjà existantes traitant des données de même nature. Cela permet de s’adapter rapidement à d’autres problèmes d’inférence de proportions.

4.1.4 Comparaison avec d’autres fonctions de pertes

Les fonctions de coût introduites dans la section précédente seront comparées avec trois autres alternatives, sans une modélisation adaptée à ce problème d’inférence de proportions. Tout d’abord, nous considérons l’erreur des moindres carrés et l’entropie croisée dont les équations (4.15) et (4.16) du chapitre précédent pour la classification ont été légèrement adaptées. On considère pour la première

$$\mathcal{L}_i(\Theta) = \| \mathbf{p}_i - \text{ReLU}(\mathbf{y}_i) \|^2 \quad (4.15)$$

et pour la seconde :

$$\mathcal{L}_i(\Theta) = - \sum_{j=1}^K p_{ij} \log(\text{Softmax}(\mathbf{y}_{ij})) \quad (4.16)$$

Ces fonctions seront respectivement nommées par la suite “MSE” et “CE & Softmax”.

La troisième et dernière fonction de perte “MSE & proportion” est donnée par :

$$\mathcal{L}_i(\Theta) = \| \mathbf{p}_i - \hat{\mathbf{p}}_i \|^2, \quad \text{avec} \quad \hat{\mathbf{p}}_{ij} = \frac{\alpha_{ij}}{S_{\alpha_i}} \quad \text{et} \quad \alpha_i = \text{ReLU}(\mathbf{p}_i) + 1 \quad (4.17)$$

Elle reprend la fonction “MSE & DIR”, avec la modélisation de Dirichlet, mais sans le terme de variance. Cela permet de regarder l’erreur des moindres carrés entre la vérité $\mathbf{p}_i$ et le vecteur de proportions obtenu par la sortie du réseau $\mathbf{y}_i$.

Pour évaluer les capacités du modèle suivant les différentes fonctions de perte, on considère différentes mesures.

4.1.5 Métriques d'évaluation pour l'inférence de proportions

L'évaluation des performances pour chacune des fonctions de perte se fera avec trois mesures. Ces mesures sont définies sur l'ensemble de validation $\mathcal{V}$ défini précédemment, de cardinal $N_{\mathcal{V}}$. Elles seront également utilisées sur l'ensemble de test.

RMSE La RMSE (pour *Root Mean Square Error*), permet de quantifier l'habileté du réseau à trouver un vecteur de proportions $\hat{\mathbf{p}}_i$ proche du vecteur recherché $\mathbf{p}_i$ au sens des moindres carrés. Elle est définie par :

$$\text{RMSE} = \sqrt{\frac{1}{N_{\mathcal{V}}} \sum_{(\mathbf{x}_i, \mathbf{p}_i) \in \mathcal{V}} \frac{1}{K} \sum_{j=1}^K (p_{ij} - \hat{p}_{ij})^2}. \quad (4.18)$$

MMAE La MMAE (pour *Mean Maximal Absolute Error*) compare également $\hat{\mathbf{p}}_i$ et $\mathbf{p}_i$, mais en considérant la norme infinie, et donc l'erreur maximale (en valeur absolue) entre les deux vecteurs. On la définit par la formule suivante :

$$\text{MMAE} = \frac{1}{N_{\mathcal{V}}} \sum_{(\mathbf{x}_i, \mathbf{p}_i) \in \mathcal{V}} \max_{j \in \{1, \dots, K\}} |\hat{p}_{ij} - p_{ij}|. \quad (4.19)$$

Cette mesure est plus pénalisante que la RMSE car elle ne moyenne pas l'erreur sur l'ensemble du vecteur.

RRS La RRS (pour *Rate of Recovered Support*), permet de quantifier la capacité du réseau à retrouver les bons composants dans le mélange. Nous définissons d'abord les deux supports suivants pour un vecteur $\mathbf{p} = (p_1, \dots, p_K) \in \Delta_K$:

$$\text{supp}(\mathbf{p}) = \{j \in \{1, \dots, K\}, p_j > 0\} \quad \text{et} \quad \text{supp}_{\varepsilon}(\mathbf{p}) = \{j \in \{1, \dots, K\}, p_j > \varepsilon\}. \quad (4.20)$$

Le support permet d'identifier les éléments du vecteur de proportion dont la valeur est supérieure à 0. Pour la donnée associée au vecteur cela revient à trouver les composants présents. Pour le support du vecteur $\hat{\mathbf{p}}_i$ (la prédiction du modèle) on utilisera $\text{supp}_{\varepsilon}(\hat{\mathbf{p}}_i)$ plutôt que $\text{supp}(\hat{\mathbf{p}}_i)$, car le réseau de neurones prédit toujours une valeur strictement supérieure à 0 pour chaque coordonnée de $\hat{\mathbf{p}}_i$. Par conséquent, $\text{supp}(\hat{\mathbf{p}}_i)$ regrouperait tous les composants. En utilisant $\text{supp}_{\varepsilon}(\hat{\mathbf{p}}_i)$, cela permet d'associer un composant au support uniquement s'il est présent au delà d'un seuil $\varepsilon \in (0, 1)$ fixé (i.e si $\hat{p}_j > \varepsilon$). Par défaut, la valeur de ε est de 1%.

Finalement, la RRS permet de comparer le support du vecteur prédit $\hat{\mathbf{p}}_i$ et celui du vecteur recherché $\mathbf{p}_i$:

$$\text{RRS} = \frac{1}{N_{\mathcal{V}}} \sum_{(\mathbf{x}_i, \mathbf{p}_i) \in \mathcal{V}} \mathbb{1}(\text{supp}(\mathbf{p}_i) = \text{supp}_{\varepsilon}(\hat{\mathbf{p}}_i)). \quad (4.21)$$

Ces mesures vont permettre de comparer les différentes fonctions introduites dans cette section, en s'appuyant sur des applications variées que nous présentons dans les sections suivantes.

4.2 Quantification de phases minérales sur des simulations

Comme cela a été évoqué en introduction, l'utilisation des réseaux de neurones pour le traitement des diffractogrammes s'est beaucoup développée ces dernières années. On peut citer, entre autres, l'identification du groupe d'espace, la classification selon la dimension du cristal ou selon le système cristallin (PARK *et al.* 2017), (OVIEDO *et al.* 2019). L'identification de phases minérales comme dans le chapitre précédent 2 a aussi fait l'objet de recherches (H. WANG *et al.* 2020).

Mais la quantification des phases avec les réseaux de neurones reste encore une tâche difficile, même si des travaux sur le sujet commencent à émerger (LEE *et al.* 2021), (LEE *et al.* 2020). Récemment, une étude de POLINE *et al.* (2024) propose une quantification des phases minérales pour des données historiques obtenues par DRX par tomographie. La problématique principale pour la quantification avec les réseaux de neurones concerne la constitution d'une base d'apprentissage avec des données expérimentales. C'est la raison pour laquelle on considère ici des données synthétiques.

La première application présentée ici concerne la quantification de phases minérales à partir des diffractogrammes calculés. Dans un premier temps, nous limiterons les analyses à des diffractogrammes simulés, afin de se concentrer sur la tâche de quantification des phases, sans gérer la complexité de données réelles.

Nous considérons un ensemble de diffractogrammes où chacun est associé à un mélange de $K = 4$ composants, i.e. quatre phases minérales : Calcite (MARKGRAF & REEDER 1985) (CaCO_3), Dolomite (STEINFINK & SANS 1959) (CaMgC_2O_6), Gibbsite (BALAN *et al.* 2006) (AlO_3H_3) et Hématite (BLAKE *et al.* 1966) (Fe_2O_3). Ces phases ont une maille centro-symétrique, ce qui permet d'utiliser l'équation (2.4) pour le facteur de structure et donc l'utilisation du code de simulation rapide (cf. algorithme 1 du chapitre 2). Dans le cas contraire, le calcul du facteur de structure est plus complexe (BISH & POST 1990). La taille de chaque diffractogramme est de 2905, et permet de couvrir une plage angulaire de 4.0001° à $90.020\,055^\circ$ (en 2θ).

Chaque donnée est un mélange contenant entre une et quatre phases, que l'on peut décrire en reprenant l'équation (4.1) où $\mathbf{C}_j$ est le diffractogramme de la classe j . La figure 4.2 donne l'exemple d'un diffractogramme contenant 40% de Calcite ($\mathbf{C}_1$) et 60% de Gibbsite ($\mathbf{C}_2$).

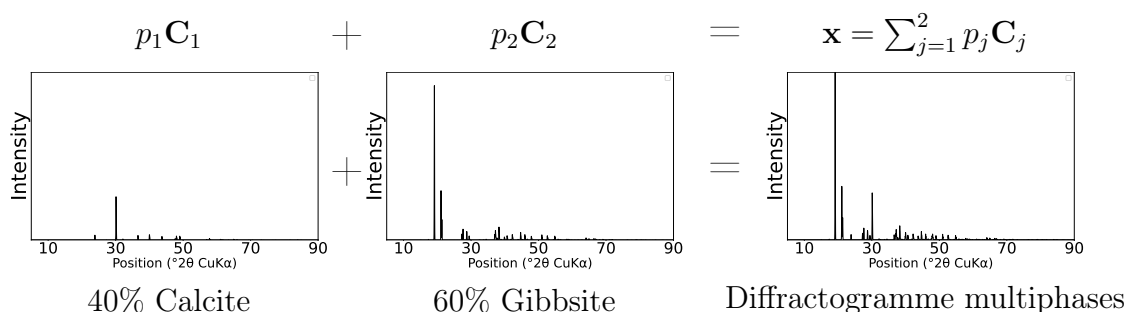


FIGURE 4.2 – Exemple d'un diffractogramme multiphase simulés avec 40% de Calcite et 60% de Gibbsite.

L'objectif sera de retrouver le vecteur de proportion $\mathbf{p}$. Pour cela, un total de 15 000 diffractogrammes mutiphases ont été générés à partir de $1\,500 \times 4$ diffractogrammes à phase

unique. La figure 4.3 illustre le processus de création de cette base. Ces 15 000 signaux ont été découpés en trois sous-ensembles (entraînement, validation et test). Chacun de ces sous-ensembles est créé à partir de signaux simple phase, c'est-à-dire à partir de diffractogrammes correspondant à une unique phase minérale. Les signaux simple phase assignés à la génération d'un sous-ensemble sont distincts de ceux assignés aux autres sous-ensembles. Cela permet de limiter tout biais sur les résultats des sous-ensembles de validation et de test.

Les 1 500 diffractogrammes de chacune des phases contiennent les mêmes fenêtres de variation pour les paramètres de maille, d'agitation thermique et de largeur à mi hauteur des pics que pour le problème de classification (voir section 3.2 et annexe A.1). Ces paramètres sont tirés de manière uniforme. Ces variations représentent la variation intra-classe, c'est-à-dire que pour une phase minérale j , il existe une infinité de diffractogrammes $\mathbf{C}_j$.

Les mélanges sont obtenus de la manière suivante. D'abord un tirage uniforme du nombre k de phases dans le mélange, k est donc compris entre 1 et $K = 4$. Ensuite, les k phases minérales du mélange sont tirées. Puis, un vecteur de proportion de taille k est construit, toujours de manière aléatoire. Enfin on effectue la combinaison linéaire du signal, les coefficients de la combinaison sont les éléments du vecteur de proportion tirés précédemment. Pour finir, les signaux sont normalisés de manière à ce que le maximum de chaque signal soit 1. La normalisation est effectuée après la combinaison linéaire des phases simples afin de respecter les rapports d'intensité entre les pics provenant de différentes phases minérales.

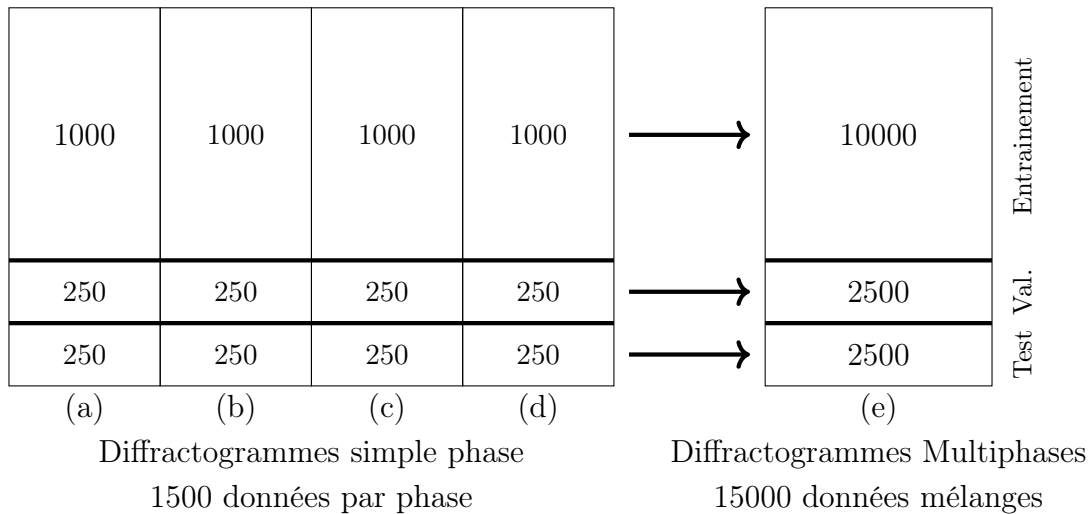


FIGURE 4.3 – Schéma de construction des différentes bases de données multiphases à partir de simulations simple phase. (a), (b), (c) et (d) correspondent aux 1 500 signaux de chaque phase, alors que (e) correspond aux 15 000 diffractogrammes multiphases.

Comme mentionné dans la section précédente, cette méthode pour l'inférence de proportion s'appuie sur un changement de fonction de perte. L'architecture du réseau est identique à celle du problème de classification 3.4, en adaptant simplement la dernière fonction linéaire pour obtenir une sortie de taille $K = 4$. Les 10 000 données de la base d'entraînement sont regroupées par *batches* de 4 durant les 100 époques de l'entraînement. L'optimizer est Adam avec un learning rate constant et égal à 0.001. Afin de s'assurer de la stabilité, et également pour offrir une meilleure comparaison, cinq entraînements du réseau (numérotés de 1 à 5) ont été effectués pour les cinq fonctions : "MSE & DIR", "CE

& DIR’, “CE & Softmax”, “MSE” et “MSE & Prop”. Cela permet d’obtenir une moyenne et un écart-type pour les mesures (RMSE, MMAE et RRS). Pour une comparaison équitable des fonctions, l’initialisation des paramètres du réseau est la même pour chacune des fonctions pour un numéro d’entraînement donné. Les cinq entraînements d’une même fonction ont des initialisations différentes. Mais, par exemple, l’entraînement numéro 1 de chaque fonction possède la même initialisation. Cela sera également le cas pour les deux autres applications qui seront présentées par la suite (imagerie hyperspectrale et MNIST).

La figure 4.4 montre l’évolution de la fonction de perte et celle des mesures (RMSE, MMAE et RRS) en fonction des époques. A l’issue de chaque époque, ces trois mesures sont calculées sur l’ensemble de validation, ce qui permet de suivre l’entraînement du réseau, et notamment d’éviter le sur-apprentissage. Seul le meilleur entraînement (celui avec la valeur de MMAE la plus faible sur l’ensemble de validation) est affiché pour cette figure.

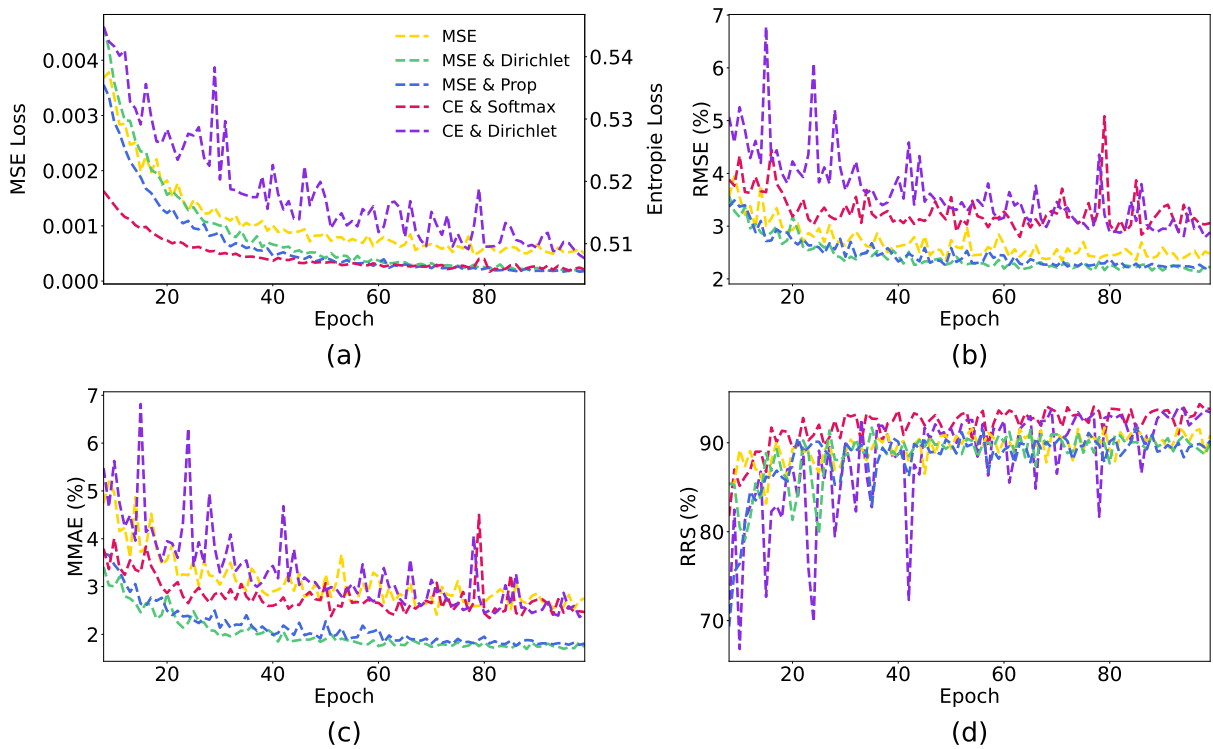


FIGURE 4.4 – Analyse de diffractogrammes simulés : (a) Evolution des valeurs des fonctions de perte, l’échelle de gauche correspond aux valeurs des fonctions utilisant l’erreur des moindres carrés (“MSE”, “MSE & DIR” et “MSE & Prop”). Les figures (b), (c) et (d) montrent l’évolution des valeurs des mesures (RMSE, MMAE et RRS respectivement) sur l’ensemble de validation au cours des époques d’entraînement.

Une première remarque concerne la figure 4.4 (a) représentant les fonctions de perte. Les échelles entre celles qui utilisent l’erreur de moindres carrés (sur la gauche) et celles utilisant l’entropie croisée (sur la droite) sont différentes. Cette différence est surtout marquée sur les valeurs de la *loss*, mais n’impacte en rien les valeurs sur les mesures d’évaluation. Nous remarquons également que les fonctions utilisant les mêmes termes de minimisation ont des valeurs très proches. La décroissance de chaque fonction est également assez nette lors des 70 premières époques, puis devient moins marquée. Il n’y a donc pas de souci d’optimisation, bien que les fonctions ne soient pas convexes.

Ensuite, sur l'ensemble de validation (figures 4.4 (b),(c) et (d)), l'évaluation des métriques MMAE et RMSE est regroupée car leur évolution est similaire. Les fonctions utilisant l'erreur des moindres carrés présentent des valeurs de MMAE et de RMSE plus faibles. Ces résultats ne sont pas surprenants car ces fonctions cherchent directement à minimiser ces erreurs contrairement à l'entropie croisée. Les valeurs obtenues sont proches de 2%. Ces résultats montrent l'efficacité de la méthode. En effet, cela signifie que le réseau est capable de trouver les bons pourcentages de chaque phase avec des erreurs de l'ordre de 2%. Les fonctions avec l'entropie croisée donnent également de bonnes performances, bien que moindres (entre 3% et 4% pour les deux mesures). Cependant, elles sont plus performantes sur la RRS, avec presque 95% des données dont le support retrouvé est correct ($\varepsilon = 1\%$), contre environ 90% pour celles utilisant la MSE.

Les résultats sur l'ensemble de test sont décrits par le tableau 4.1. Ils viennent confirmer les observations sur l'ensemble de validation. La faible valeur de l'écart-type sur les cinq entraînements est un indicateur de stabilité. Les deux fonctions dont les cinq entraînements ont fonctionné sont "MSE & DIR" et "CE & Softmax" et semblent donc plus stables. La meilleure performance est obtenue pour la fonction utilisant l'erreur des moindres carrés. Les trois autres fonctions présentent un ou plusieurs entraînements pour lesquels l'algorithme d'optimisation n'a pas convergé.

TABLE 4.1 – DRX : Résultats des différentes mesures (RMSE, MMAE et RRS) sur l'ensemble de test pour chacune des fonctions. Le tableau est séparé en 3 parties distinctes, avec tout d'abord les moyennes et écarts-types sur 5 entraînements. La deuxième partie ne retient que les entraînements performants (pour lesquels la fonction de perte a convergé), leur nombre étant indiqué entre parenthèses. Enfin, seul le meilleur entraînement est considéré, i.e celui avec la valeur de MMAE la plus faible sur l'ensemble de validation.

Fonction de perte	RMSE ↓	MMAE ↓	RRS _{1%} ↑
5 entraînements			
MSE & Dirichlet	02.22% ± 00.05	01.82% ± 0.03	<u>90.00% ± 00.54</u>
CE & Dirichlet	14.82% ± 14.67	21.49% ± 23.23	56.00% ± 45.72
<u>CE & Softmax</u>	<u>02.92% ± 00.18</u>	<u>02.49% ± 00.12</u>	93.42% ± 00.26
MSE	09.60% ± 08.85	11.78 ± 11.25%	73.41% ± 20.76
MSE & Prop	26.68% ± 12.20	40.33% ± 19.22	17.84% ± 35.68

C'est la raison pour laquelle dans la seconde partie du tableau, nous considérons la moyenne et l'écart-type uniquement sur les entraînements qui ont fonctionné, afin d'offrir une meilleure lecture. Encore une fois, la "MSE & DIR" fournit des meilleures valeurs avec des écarts types très faibles. Sur ces mêmes mesures d'erreur, la "MSE & Prop" offre des résultats assez proches (mais sur l'unique entraînement qui a fonctionné). Les fonctions d'entropie croisée se distinguent toujours par la RRS, avec des valeurs au-delà de 93%.

La dernière partie du tableau consiste à retenir le meilleur entraînement, c'est-à-dire celui avec la valeur de MMAE la plus faible sur l'ensemble de validation. Les constats sont identiques à ceux faits sur la deuxième partie du tableau. La meilleure performance est obtenue par la fonction "MSE & DIR" avec 2.20% sur la RMSE, 1.79% sur la MMAE et 90.0% pour la RRS.

Pour cette application il est important de rappeler, l'importance de la variation intra-classe qui complexifie le problème. En effet, le diffractogramme associé à une phase minérale n'est pas unique et peut prendre une infinité de formes. Les fluctuations des nombreux

paramètres (maille, agitation thermique, etc) peuvent notamment entraîner de légères variations sur la position, l'intensité et la forme des pics.

4.3 Imagerie Hyperspectrale

La seconde application pour le problème d'inférence de proportions concerne l'imagerie hyperspectrale. Ce sont des images pour lesquelles chaque pixel contient un spectre comme le montre la figure 4.5. Ces spectres couvrent une large gamme de longueurs d'ondes sur la bande électromagnétique. C'est grâce à eux que la composition de chaque pixel peut être identifiée pour l'image analysée.

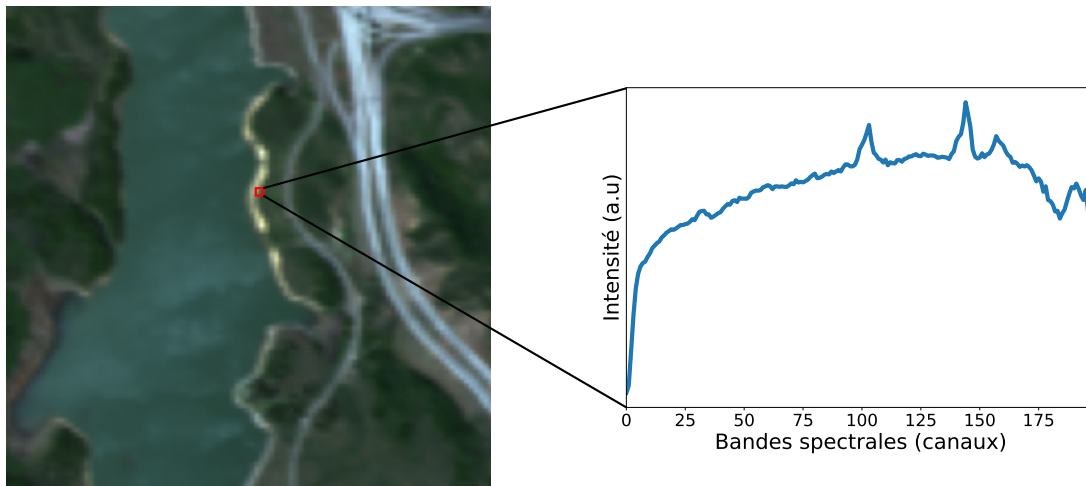


FIGURE 4.5 – Exemple de l'image hyperspectrale Jasper Ridge, avec l'illustration d'un spectre contenu dans un pixel.

Dans ce contexte, en reprenant l'équation (4.1), les vecteurs $\mathbf{C}_i$ correspondent aux composants. Dans la littérature sur les images hyperspectrales, les composants sont appelés *endmembers*, et chacun a une signature spectrale spécifique. Les *endmembers* sont variables d'une image à une autre, et peuvent éventuellement varier pour une même image (ZHU 2017). La tâche consiste à retrouver les abondances des *endmembers* dans chaque pixel, ce qui revient à chercher le vecteur $\mathbf{p}$.

L'avantage de l'imagerie hyperspectrale est que l'on peut aisément comparer notre approche avec d'autres méthodes de Machine Learning comme SUnSAL¹ (BIOUCAS-DIAS & FIGUEIREDO 2010) (un algorithme basé sur la méthode des directions alternées, qui permet notamment de trouver les abondances des *endmembers*). Il y a également des approches neuronales comme HyperAE² (PALSSON *et al.* 2018) (un auto-encodeur capable de trouver à la fois les abondances des *endmembers*, mais aussi de reconstituer leur signature spectrale). Enfin, l'algorithme UnDIP³ (RASTI *et al.* 2021) est capable d'extraire les *endmembers*, et également d'estimer leur proportion.

Nous allons considérer deux exemples d'images satellitaires, appelées respectivement Jasper Ridge et Urban et dont nous souhaitons retrouver la composition. Il est important de noter pour ces exemples que la vérité-terrain admise n'est pas la vérité absolue (due à la difficulté d'établir cette dernière) (ZHU 2017). Il est en effet impossible de déterminer l'abondance de chacun des composants pour chaque pixel d'une image satellite.

-
1. <https://github.com/Laadr/SUNSAL>
 2. <https://github.com/dv-fenix/HyperspecAE>
 3. <https://github.com/BehnoodRasti/UnDIP>

4.3.1 Jasper Ridge

La première image est une observation de la réserve biologique de Jasper Ridge en Californie (Etats-Unis). Elle est de taille $100 \times 100 \times 198$, et est représentée à la figure 4.5. C'est une sous-image extraite d'une image de dimension plus importante. La taille originale des spectres est de 224 (pour des longueurs d'ondes entre 380 et 2 500 nanomètres). Cependant à cause d'effets atmosphériques, seules 198 bandes spectrales sont retenues (ZHU 2017). Les *endmembers* sont l'eau, la terre, les arbres et la route. Nous avons donc un problème avec $K = 4$ classes. La figure 4.6 décrit les abondances de chacune sur l'ensemble de l'image.

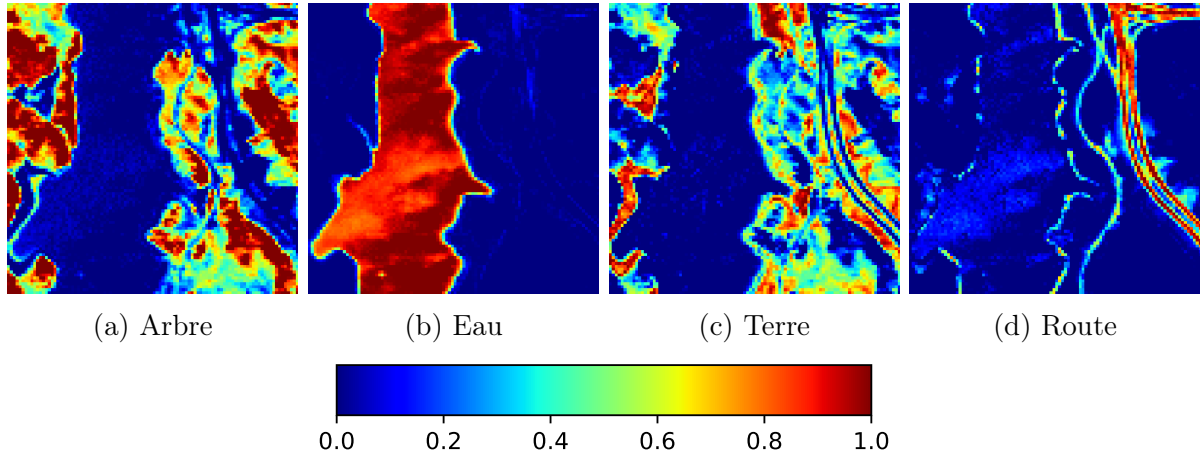


FIGURE 4.6 – Jasper Ridge : Abondances de chaque *endmember* pour l'ensemble de l'image hyperspectrale. Ces abondances correspondent à la vérité terrain établie.

Nous avons découpé l'image en trois sous-ensembles, comme le montre la figure 4.8. La partie du haut servira à entraîner le modèle, celle du milieu sert de base de validation, et celle du bas de base de test. Au final 50×100 pixels sont utilisés pour l'entraînement et 25×100 pour les deux autres bases. Ce découpage a été choisi car il permet de limiter la proximité entre chaque sous-ensemble, offrant des résultats plus fiables.

L'architecture du modèle utilisée est schématisée par la figure 4.7. Il s'agit d'un réseau de neurones convolutif introduit par ZHANG *et al.* (2018). Il est composé de quatre couches de convolution accompagnées d'un *MaxPooling*, puis de 2 couches linéaires et enfin la couche de sortie du modèle. L'opération de *MaxPooling* consiste à récupérer la plus grande valeur sur un support fixé, i.e sur une fenêtre du signal (après convolution). Elle est répétée tout au long du signal. Le réseau prend en entrée uniquement le spectre, ce qui permet au modèle de ne gérer que la dimension spectrale de l'image hyperspectrale. De récents travaux ont cependant montré que la dimension spatiale pouvait également être prise en compte, en utilisant des cubes de pixels (ZHANG *et al.* 2018) (i.e *patches hyperspectral*). Il est montré dans cette étude que ce type d'approche permet une amélioration significative des résultats.

Les spectres contenus dans les pixels ont été regroupés par *batches* de 10 sur les 100 époques de l'entraînement. Nous utilisons l'algorithme Adam comme optimiseur, avec un learning rate constant de 0.001.

Nous illustrons ces résultats à travers trois figures différentes. La première permet de suivre l'entraînement du réseau, avec d'abord l'évolution de la fonction de perte au cours

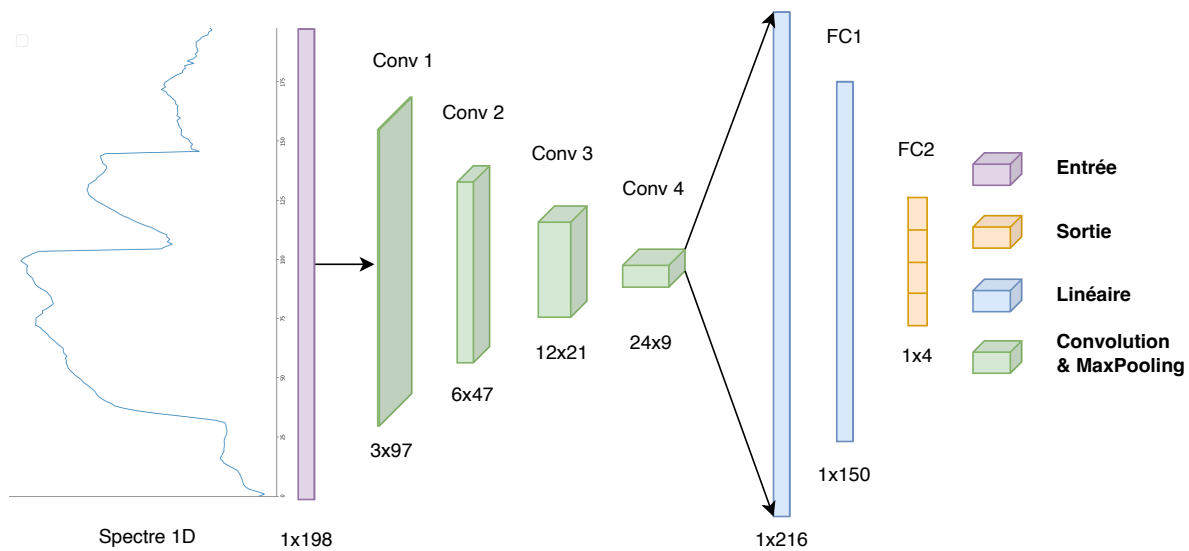


FIGURE 4.7 – Architecture du réseau de neurones convolutif utilisée pour l’inférence de proportions des images hyperspectrales

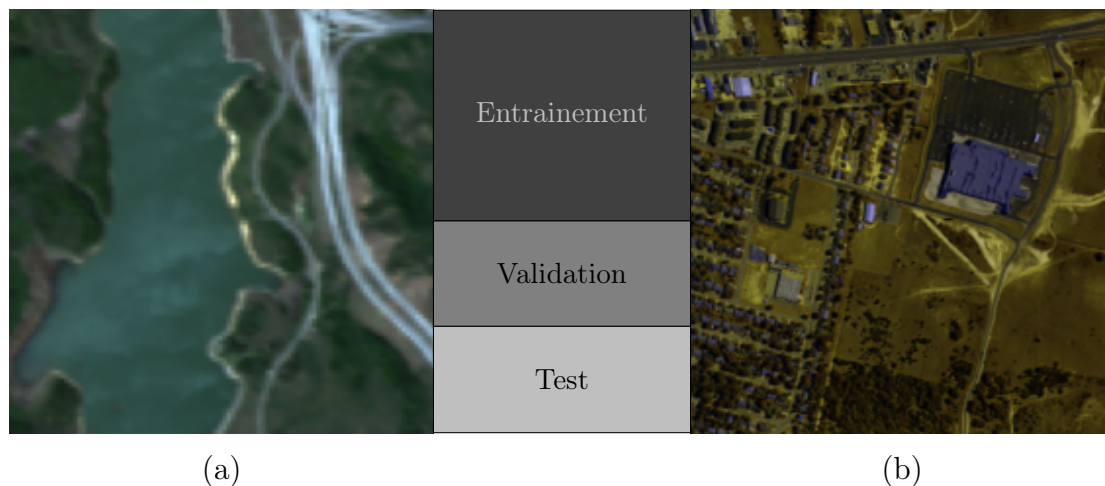


FIGURE 4.8 – Découpage des images hyperspectrales ((a) Jasper Ridge et (b) Urban) selon les trois sous-ensembles : Entrainement, Validation et Test.

des époques illustrée par la figure 4.9(a). Comme pour l’expérience précédente, les échelles de valeurs sont très différentes entre les fonctions qui utilisent l’erreur des moindres carrés et celles utilisant l’entropie croisée. Nous observons que toutes les fonctions décroissent au cours des 100 époques de l’entraînement. Toutefois la décroissance est moins marquée au fur et à mesure des époques. Cette analyse laisse penser qu’un nombre plus important d’époques serait nécessaire. L’évolution des mesures sur l’ensemble de validation disponible sur la figure 4.9 (b,c et d) permet de vérifier cette hypothèse. Cependant, sur ces mesures, la décroissance (ou croissance pour la RRS) n’est pas aussi nette. Sur la fin d’entraînement, la tendance est même plutôt à la stabilisation, ce qui laisse finalement penser que les 100 époques sont suffisantes pour ce problème d’inférence. Concernant les valeurs atteintes, la RMSE et la MMAE sont inférieures à 3% pour les meilleures fonctions “MSE & DIR” et “MSE & Prop”, les valeurs pour les fonctions “CE & DIR”, “CE & Softmax” et “MSE” sont elles plus élevées. Pour la RRS, les fonctions avec la modélisation de Dirichlet

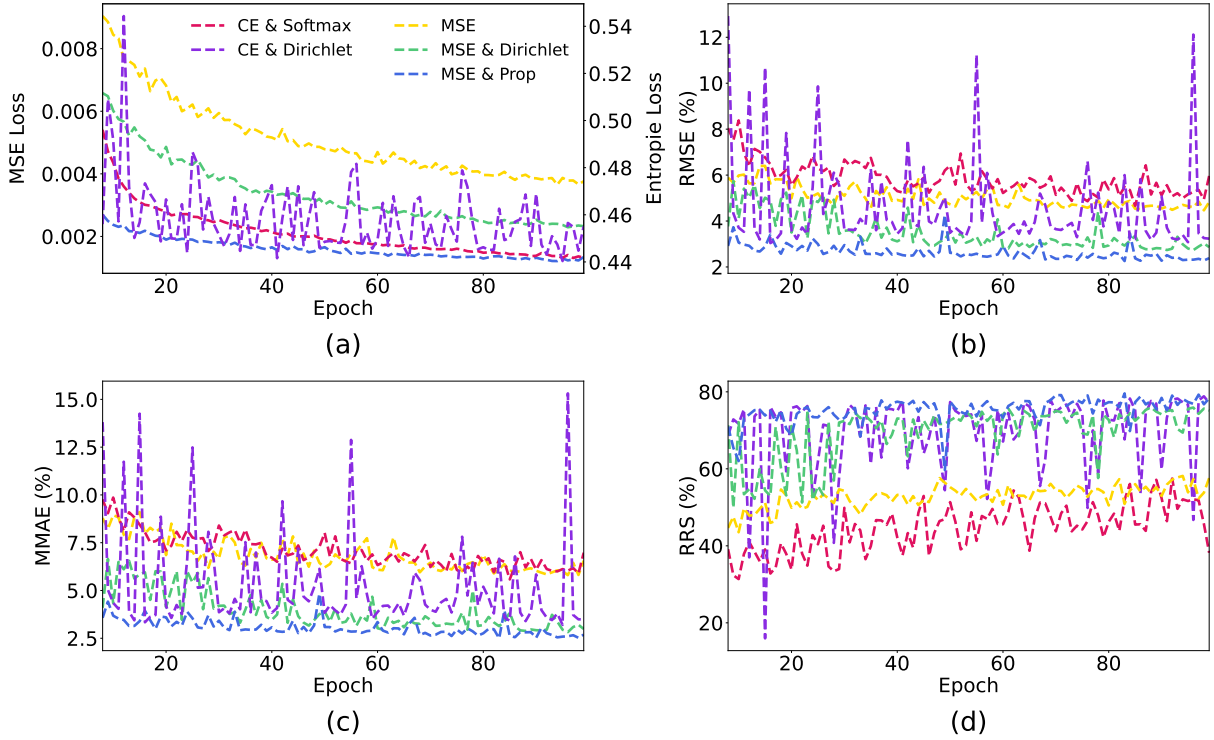


FIGURE 4.9 – Analyse de l’image Jasper Ridge : (a) Evolution des valeurs des fonctions de perte, l’échelle de gauche correspond aux valeurs des fonctions utilisant l’erreur des moindres carrés (MSE , $MSE \& Dirichlet$ et $MSE \& Prop$. Les figures (b), (c) et (d) montrent l’évolution des valeurs des mesures (RMSE, MMAE et RRS respectivement) sur l’ensemble de validation au cours des époques d’entraînement.

ont les valeurs les plus élevées, se rapprochant des 80%.

Ces valeurs sont confirmées sur l’ensemble de test dont les résultats sont présentés au tableau 4.2. Sur les entrainements performants, deux fonctions se distinguent particulièrement : “MSE & DIR” et “MSE & Prop”, offrant une excellente stabilité tout en étant capable de retrouver les bonnes proportions avec une MMAE de l’ordre de 2.5%, et de trouver les bons composants dans plus de 80% des cas. Contrairement à l’inférence de proportions sur les phases minérales, les fonctions utilisant l’entropie croisée, ici, ne sont pas nettement meilleures sur la RRS. Cette table permet également de se rendre compte que toutes les fonctions sont stables, avec un seul entrainement avec la fonction “MSE & Prop” qui n’a pas fonctionné. Enfin, la table offre une comparaison avec les trois méthodes alternatives présentées en début de section : SUnSAL, UnDIP et HyperAE. Parmi ces méthodes, c’est l’algorithme SUnSAL qui présente les meilleures performances sur les trois mesures, avec des valeurs inférieures à 8% pour la RMSE et la MMAE, et une RRS supérieure à 50%. On peut voir sur le tableau 4.2 que pour l’inférence de proportion, notre méthode, avec l’utilisation d’un réseau de neurones dédié et d’une fonction de perte spécifique, montre des résultats largement supérieurs à ces méthodes.

Pour terminer sur l’image Jasper Ridge, une illustration visuelle est proposée à la figure 4.10. C’est une comparaison entre la vérité et la prédiction des modèles sur l’ensemble de test. La première ligne est la vérité terrain, c’est-à-dire les abondances de chacun des composants. Les lignes suivantes correspondent à la valeur absolue de la différence entre la vérité et la prédiction des modèles. Les lignes sont associées à une fonction de perte, puis aux méthodes alternatives. Visuellement, trois lignes (2, 3 et 6) et donc trois fonctions se

TABLE 4.2 – Jasper Ridge : Résultats des différentes mesures (RMSE, MMAE et RRS) sur l’ensemble de test pour chacune des fonctions. La table est séparée en 4 parties distinctes, avec tout d’abord les moyennes et écarts-types sur 5 entraînements. La deuxième partie ne retient que les entraînements performants (pour lesquels la fonction de perte a convergé), leur nombre est indiqué entre parenthèses. Ensuite, seul le meilleur entraînement est considéré, i.e celui avec la valeur de MMAE la plus faible sur l’ensemble de validation. Enfin, la dernière partie concerne les mesures pour les méthodes alternatives.

Fonction de perte	RMSE ↓	MMAE ↓	RRS _{1%} ↑
5 entraînements			
MSE & Dirichlet	02.60% ± 00.07	02.44% ± 0.06	80.05% ± 00.67
<u>CE & Dirichlet</u>	<u>03.00% ± 00.06</u>	<u>03.24% ± 00.08</u>	<u>79.79% ± 01.39</u>
CE & Softmax	03.90% ± 0.09	04.14% ± 00.16	57.48% ± 01.76
MSE	03.60% ± 00.25	04.51% ± 00.26	67.50% ± 3.50
MSE & Prop	09.03% ± 13.20	13.25% ± 21.36	65.00% ± 30.77
Entraînements performants			
<u>MSE & Dirichlet</u> (5)	<u>02.60% ± 00.07</u>	02.44% ± 0.06	<u>80.05% ± 00.67</u>
<u>CE & Dirichlet</u> (5)	<u>03.00% ± 0.06</u>	<u>03.24% ± 00.08</u>	<u>79.79% ± 01.39</u>
<u>CE & Softmax</u> (5)	<u>03.90% ± 0.09</u>	<u>04.14% ± 00.16</u>	<u>57.48% ± 01.76</u>
<u>MSE</u> (5)	<u>03.60% ± 00.25</u>	<u>4.51% ± 00.26</u>	<u>67.50% ± 3.50</u>
MSE & Prop (4)	02.43% ± 00.14	<u>02.56% ± 00.22</u>	80.37% ± 01.20
Meilleur entraînement et comparaison			
MSE & Dirichlet	<u>02.48%</u>	02.34%	80.30%
CE & Dirichlet	02.90%	03.09%	77.80%
CE & Softmax	04.03%	03.87%	59.50%
MSE	03.29%	04.21%	74.00%
<u>MSE & Prop</u>	02.34%	<u>02.37%</u>	80.30%
SUnSAL	07.81%	07.83%	51.40%
UnDIP	25.91%	32.18%	37.50%
HyperAE	18.80%	22.70%	15.10%

distinguent : “MSE & DIR”, “CE & DIR” et “MSE & Prop”. Les erreurs restent faibles sur l’ensemble de l’image et leurs valeurs les plus élevées sont sur les zones frontières : la délimitation de la zone d’eau, les zones partageant plusieurs *endmembers*, comme la partie droite de l’image contenant à la fois des arbres, mais aussi de la terre. De manière plus générale, ces zones sont celles avec les erreurs les plus importantes, quelle que soit la fonction de perte, ou la méthode utilisée.

4.3.2 Urban

La seconde image hyperspectrale que nous traitons dans ce chapitre est l’image Urban (voir figure 4.8 (b)). Elle est très largement utilisée dans les études hyperspectrales (ZHU 2017), (ZHAO *et al.* 2022), (INCE & DOBIGEON 2022), (WAN *et al.* 2021). Il s’agit d’une vue satellite d’un lieu urbain. Elle comporte 307×307 pixels, contenant chacun des spectres de taille 210 qui couvrent des longueurs d’ondes entre 400 et 2 500 nanomètres. Cependant, comme pour Jasper Ridge, dû à certains effets atmosphériques, seules 162 bandes spectrales sont retenues (ZHU 2017). Plusieurs vérités existent pour l’image Urban avec un nombre de *endmembers* variable. Dans notre étude, nous utilisons la vérité avec $K = 6$ composants (ZHU 2017) : bitume, herbe, arbre, toit, métal et terre (voir figure 4.11).

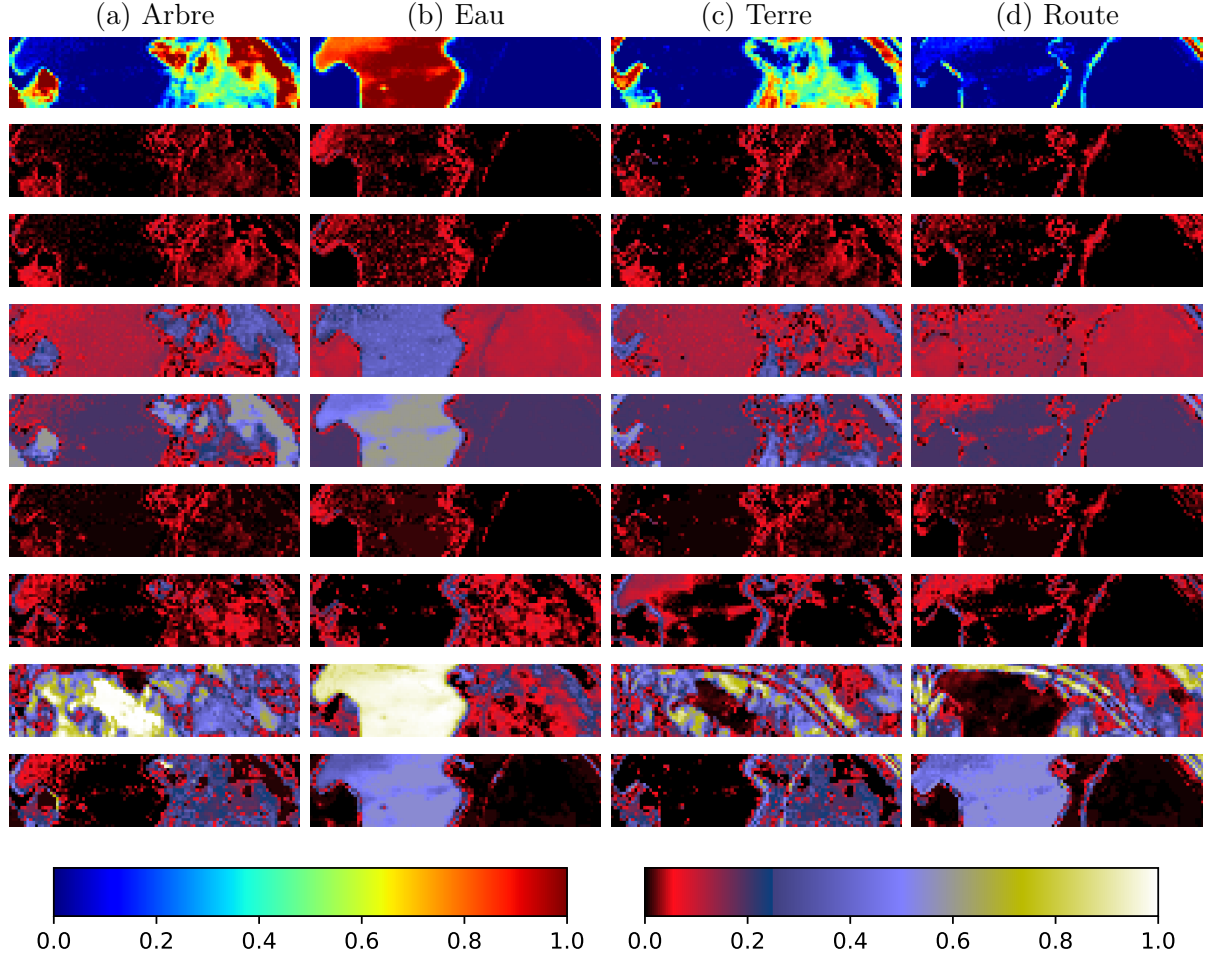


FIGURE 4.10 – Pour chaque classe de Jasper Ridge : Zone test (1^e ligne) et différence absolue entre la vérité et la prédiction respectivement pour les fonctions : “MSE & DIR” (2^e ligne), “CE & DIR” (3^e ligne), “CE & Softmax” (4^e ligne), “MSE” (5^e ligne), “MSE & Prop” (6^e ligne) et méthodes alternatives : SUNSAL (7^e ligne), UnDIP (8^e ligne) et HyperAE (9^e ligne).

Le processus d’entraînement est le même que celui décrit pour l’image Jasper Ridge. Comme illustré sur la figure 4.8, l’image est aussi découpée en trois parties pour limiter le biais entre les sous-ensembles. 50% des pixels sont utilisés pour l’entraînement, 25% pour la validation et 25% pour l’ensemble de test. L’architecture du réseau est presque identique, avec simplement une adaptation à la taille du spectre et au nombre de classes. La taille des *batches*, le nombre d’époques, ainsi que l’optimiseur restent également identiques.

Les résultats sont décrits avec trois figures similaires. La figure 4.12 montre l’évolution au cours des époques de la fonction de perte et des mesures (MMAE, RMSE et RRS) sur l’ensemble de validation. Les conclusions faites sont les mêmes que pour l’image Jasper Ridge. On peut aussi noter que la fonction “CE & Softmax” obtient une valeur de RRS très basse, avec tout juste 40% à la fin de l’entraînement du réseau. De plus, les valeurs de la RMSE et MMAE (resp. RRS) sont plus élevées (resp. basses) car l’image a deux composants supplémentaires.

Ces résultats sont une nouvelle fois confirmés sur l’ensemble de test (voir tableau 4.3) où les valeurs de chaque mesure sont indiquées. Sur cette expérience, tous les entraîne-

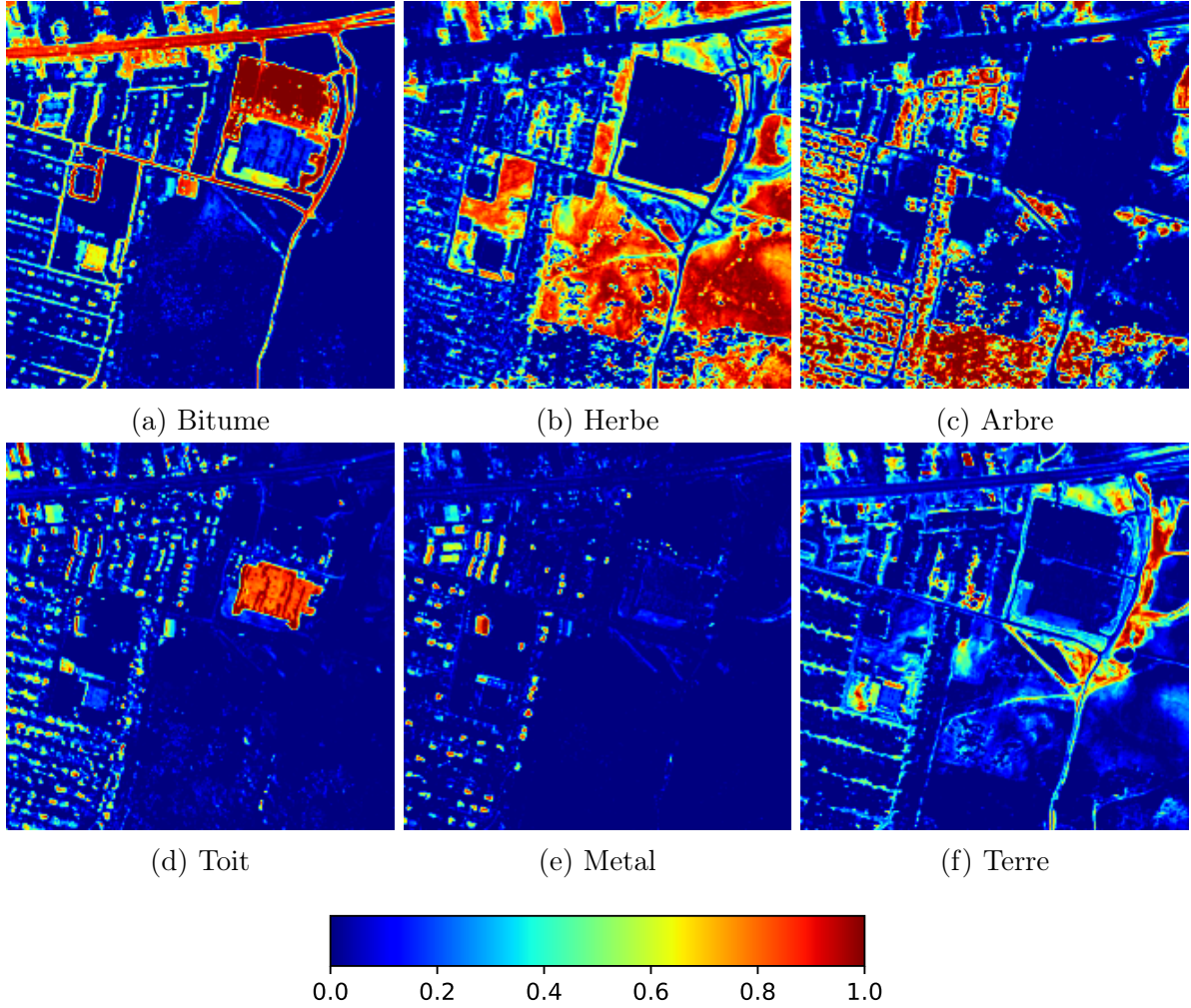


FIGURE 4.11 – Urban : Abondances de chaque *endmember* pour l’ensemble de l’image hyperspectrale. Ces abondances correspondent à la vérité terrain établie.

ments ont fonctionné. Le tableau ne comporte alors que deux parties. Les valeurs sont plus élevées que pour Jasper Ridge avec une erreur maximale moyenne de 5.45% pour la “MSE & Prop” et une RMSE de 4.17%, alors que la RRS vaut 56.10%. Sur ces images hyperspectrales, ce sont toujours les fonctions utilisant l’erreur des moindres carrés qui fournissent les valeurs de RRS les plus élevées. Enfin sur l’image Urban, l’avantage de combiner un réseau dédié à l’inférence de proportion avec une fonction dont la modélisation est adaptée est plus net. En effet les résultats sont nettement meilleurs que ceux obtenus avec les méthodes alternatives.

Les données du tableau 4.3 sont confirmées visuellement par la figure 4.14. Les erreurs les plus importantes sont pour les classes Herbe et Arbre qui sont confondues par les prédictions. Comme le montre la figure 4.13a, c’est la proximité entre ces deux classes et donc la difficulté à les distinguer qui mène à cette confusion. C’est également le cas pour les classes terre et route, leur proximité est illustrée sur la figure 4.13b. Les zones frontières entre deux *endmembers* différents comportent également les erreurs les plus importantes.

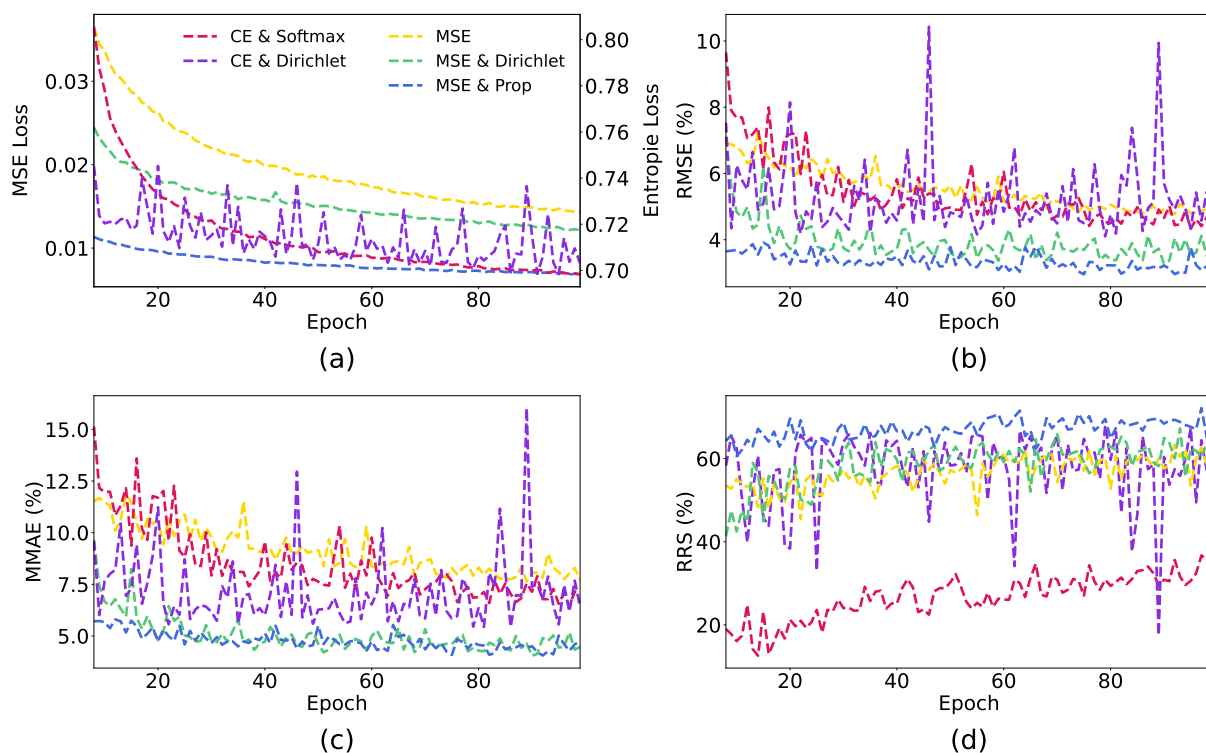


FIGURE 4.12 – Analyse de l'image Urban : (a) Evolution des valeurs des fonctions de perte, l'échelle de gauche correspond aux valeurs des fonctions utilisant l'erreur des moindres carrés (MSE , MSE & $Dirichlet$ et MSE & $Prop$). Les figures (b), (c) et (d) montrent l'évolution des valeurs des mesures (RMSE, MMAE et RRS respectivement) sur l'ensemble de validation au cours des époques d'entraînement.

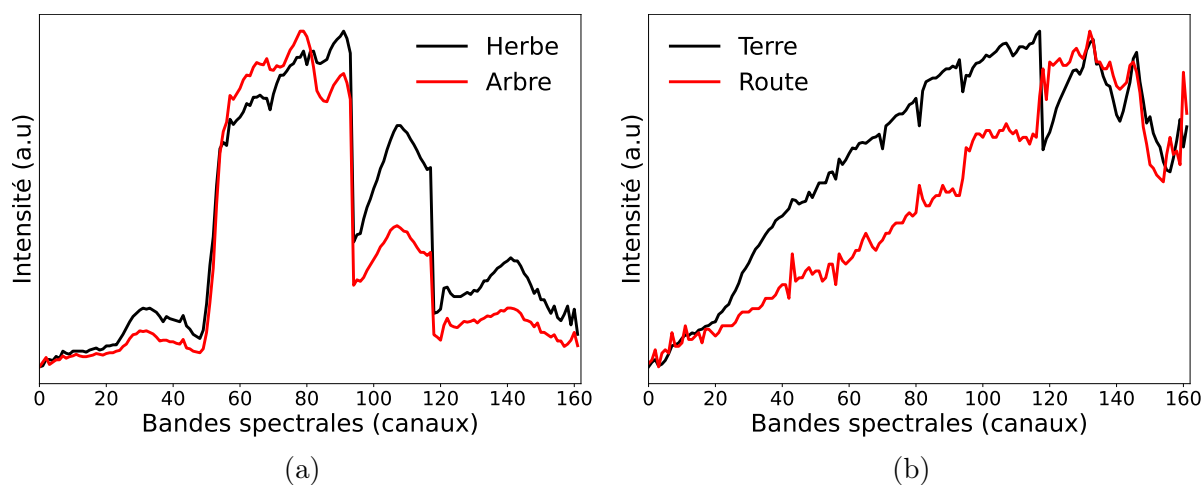


FIGURE 4.13 – Comparaison des spectres des *endmembers* de l'image Urban. La figure (a) compare Arbre et Herbe et la figure (b) Route et Terre.

TABLE 4.3 – Urban : Résultats des différentes mesures (RMSE, MMAE et RRS) sur l’ensemble de test pour chacune des fonctions. Le tableau est séparé en 4 parties, avec d’abord les moyennes et écarts-types sur 5 entraînements. La deuxième partie ne retient que les entraînements performants (pour lesquels la fonction de perte a convergé), leur nombre étant indiqué entre parenthèses. Ensuite, seul le meilleur entraînement est considéré, i.e celui avec la valeur de MMAE la plus faible sur l’ensemble de validation. Enfin, la dernière partie concerne les mesures pour les méthodes alternatives.

Fonction de perte	RMSE ↓	MMAE ↓	RRS _{1%} ↑
5 entraînements			
<u>MSE & Dirichlet</u>	<u>04.75% ± 00.19</u>	<u>06.13% ± 00.30</u>	53.90% ± 01.35
CE & Dirichlet	05.81% ± 00.36	08.32% ± 00.59	50.02% ± 01.12
CE & Softmax	05.04% ± 00.09	06.75% ± 00.26	41.59% ± 03.12
MSE	05.25% ± 00.17	08.62% ± 00.34	56.75% ± 02.70
MSE & Prop	04.25% ± 00.09	05.70% ± 00.16	<u>56.67% ± 00.73</u>
Meilleur entraînement et comparaison			
<u>MSE & Dirichlet</u>	<u>04.50%</u>	<u>05.68%</u>	55.70%
CE & Dirichlet	05.44%	07.77%	51.12%
CE & Softmax	04.95%	06.48%	43.72%
MSE	05.07%	08.13%	59.70%
MSE & Prop	04.17%	05.45%	<u>56.10%</u>
SUnSAL	19.70%	28.89%	17.50%
UnDIP	36.30%	54.79%	06.73%
HyperAE	28.10%	52.00%	00.68%

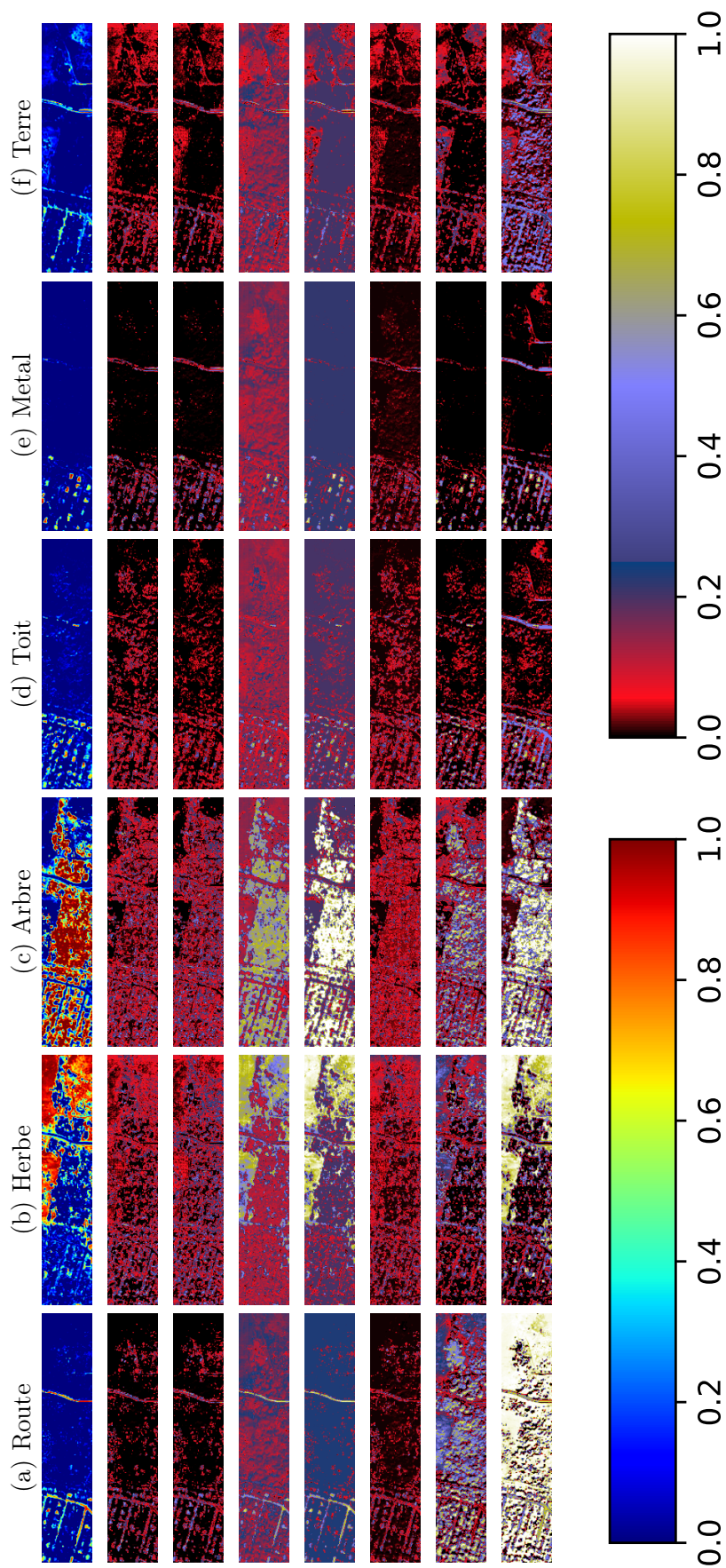
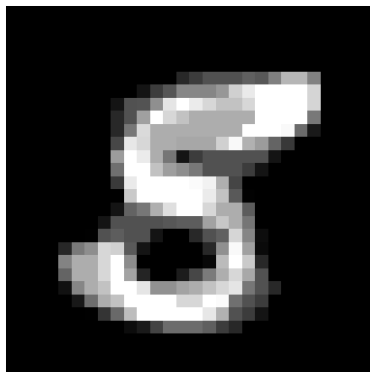


FIGURE 4.14 – Pour chaque classe de l’image Urban : Zone test (1^e ligne) et différence absolue entre la vérité et la prédiction respectivement pour les fonctions : “MSE & DIR” (2^e ligne), “CE & Softmax” (3^e ligne), “CE & Softmax” (4^e ligne), “MSE” (5^e ligne), “MSE & Prop” (6^e ligne) et méthodes alternatives : SUnSAL (7^e ligne), UnDIP (8^e ligne) et HyperAE (9^e ligne).

4.4 Images MNIST

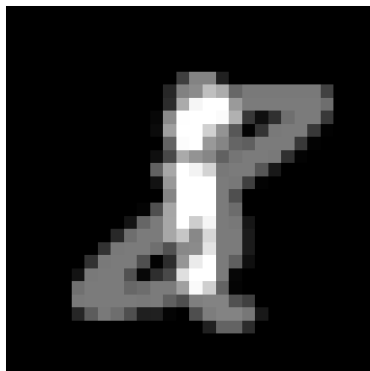
La dernière application que nous proposons concerne des images construites à partir de la base de données MNIST (LECUN 1998). Il s’agit à travers cet exemple synthétique de valider la méthodologie proposée pour l’inférence de proportions. Cette base de données contient des images (28×28 pixels) de chiffres manuscrits. C’est une base de données répandue pour la classification. La base est composée d’images de chiffres manuscrits de 0 à 9.

Nous avons construit la base de données “MNIST mélange” dont chaque élément est une combinaison linéaire de deux images (i.e deux chiffres) issues de la base initiale MNIST. Il s’agit ici de considérer un exemple simple pour l’inférence de proportions sur des images. Quelques exemples d’images artificielles créées sont données à la figure 4.15. La base de données est limitée à des mélanges à 2 classes au plus. L’objectif est de retrouver le vecteur de proportions des chiffres mélangés.



	0	1	2	3	4	5	6	7	8	9
GT	0%	0%	0%	0%	0%	68%	0%	0%	32%	0%
Pred	1%	6%	1%	1%	1%	61%	1%	1%	30%	1%

(a)



	0	1	2	3	4	5	6	7	8	9
GT	0%	52%	0%	0%	0%	0%	0%	0%	48%	0%
Pred	0%	42%	19%	0%	0%	0%	0%	0%	37%	0%

(b)

FIGURE 4.15 – Exemples de deux images traitées par le réseau (à gauche) et comparaison des vecteurs de proportions entre la vérité (GT) et la prédiction du réseau (Pred) (à droite). Les valeurs en gras indiquent les chiffres mélangés. Le réseau est entraîné avec la fonction “MSE & Dirichlet”.

La base d’images MNIST initiale a été découpée en trois sous-ensembles qui ont permis de générer des bases de données différentes. Les images utilisées pour la création de chacune des bases ne sont pas les mêmes afin de réduire le biais entre les sous-ensembles. Au total 70 000 images ont été créées, réparties en trois sous-ensembles : 60 000 images pour la base d’entraînement, 5 000 pour la validation et 5 000 pour le test.

On a utilisé un réseau de neurones convolutif. L'architecture est reprise du tutoriel PyTorch⁴ pour la classification. Le réseau possède deux couches de convolution, puis deux couches linéaires avec Dropout. Pour l'entraînement, on a considéré 200 époques et les images ont été groupées en *batch* de taille 64. L'optimizer est Adam, avec un *learning rate* constant de 0.001. Comme pour les expériences précédentes, le réseau est entraîné 5 fois par fonction avec la même initialisation pour chaque numéro d'entraînement (de 1 à 5).

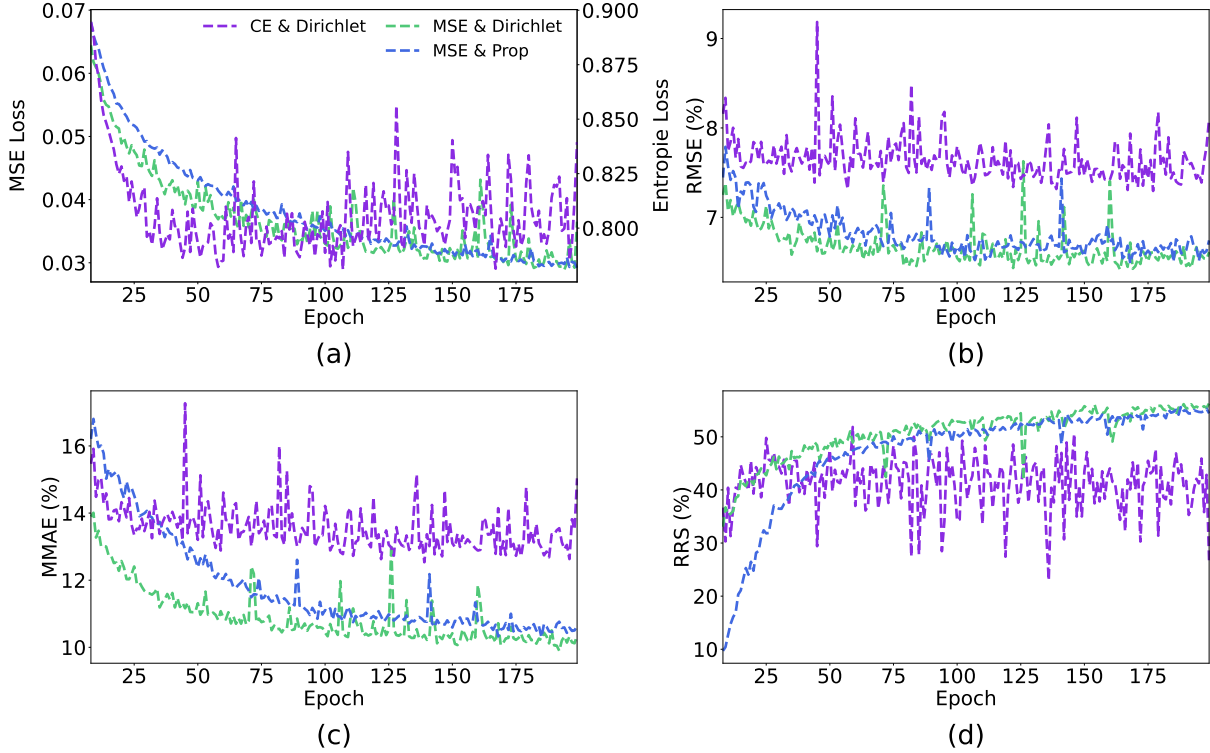


FIGURE 4.16 – Images “MNIST mélange” : (a) Evolution des valeurs des fonctions de perte, l’échelle de gauche correspond aux valeurs des fonctions utilisant l’erreur des moindres carrés (*MSE*, *MSE & Dirichlet* et *MSE & Prop*). Les figures (b), (c) et (d) décrivent les valeurs des mesures (RMSE, MMAE et RRS respectivement) sur l’ensemble de validation au cours des époques d’entraînement.

Résultats La figure 4.16 montre l’évolution de la fonction de perte, ainsi que des mesures d’erreurs sur l’ensemble de validation. Seul le meilleur entraînement (dont la fonction de perte converge) a été retenu pour chacune des fonctions concernées. Le tableau 4.4 montre les mesures pour chacune des fonctions de perte respectivement pour les 5 entraînements, ceux ayant fonctionné, et le meilleur. On a testé deux valeurs de ε pour la RRS, le seuil à 10% permettant une tolérance un peu plus grande. Ce seuil correspond à peu près à la valeur de la MMAE. La figure 4.15 donne deux exemples d’images comparant la vérité et la prédiction du modèle.

A partir de ces résultats, plusieurs remarques peuvent être faites. Premièrement, on remarque que les fonctions de perte sans une modélisation adaptée (“MSE” et “CE & Softmax”) ne présentent aucun entraînement où la fonction de perte converge. Ainsi pour une meilleure lisibilité, elles ne sont pas affichées sur les graphiques de la figure 4.16.

4. <https://github.com/pytorch/examples/blob/main/mnist/main.py>

Ensuite, on voit que la fonction “MSE & Dirichlet” donne les meilleures performances, aussi bien en terme de stabilité (avec un très faible écart-type sur les cinq entraînements), qu’en terme de performances (une erreur des moindres carrés de 6% en moyenne, et une erreur maximale inférieure à 10% en moyenne). La RRS au seuil de 1% est légèrement supérieure à 50%. Mais comme on peut le voir sur les deux exemples de la figure 4.15, le réseau a tendance à mettre des petites valeurs dans chaque coordonnée du vecteur de proportion prédit. Ainsi en augmentant le seuil ε à la valeur de la MMAE (10%), il y a un gain de 5% sur la valeur de la RRS. La modélisation de Dirichlet proposée permet au réseau d’être performant pour retrouver les proportions des composants.

Enfin, on peut noter que sur les entraînements qui ont fonctionné, la “MSE & Prop” est également très performante avec des chiffres proches de la fonction “MSE & Dirichlet”. Cependant elle offre moins de stabilité sur les résultats. Enfin la “CE & Dirichlet” présente une bonne stabilité et de bonnes performances, bien qu’inférieures à celles des fonctions évoquées précédemment.

TABLE 4.4 – MNIST : Résultats des différentes mesures (RMSE, MMAE et RRS) sur l’ensemble de test pour chacune des fonctions. Le tableau est séparé en 3 parties, avec d’abord les moyennes et écarts-types sur 5 entraînements. La deuxième partie ne retient que les entraînements performants (dont la fonction de perte a convergé), leur nombre étant indiqué entre parenthèses. Enfin, seul le meilleur entraînement est considéré, i.e celui avec la valeur de MMAE la plus faible sur l’ensemble de validation.

Fonction de perte	RMSE ↓	MMAE ↓	RRS _{1%} ↑	RRS ₁₀ ↑
5 entraînements				
MSE & Dirichlet	06.43% ± 00.01	10.15% ± 00.26	51.40% ± 01.54	55.91% ± 01.04
<u>CE & Dirichlet</u>	<u>06.95% ± 00.09</u>	<u>11.59% ± 00.27</u>	<u>50.52% ± 02.08</u>	61.98% ± 01.34
CE & Softmax	23.83% ± 00.01	64.85% ± 00.00	00.00% ± 00.00	00.00% ± 00.00
MSE	24.22% ± 00.53	64.87 ± 00.15%	00.00% ± 00.00	00.00% ± 00.00
MSE & Prop	16.90% ± 08.42	43.21% ± 26.54	20.74% ± 25.44	22.43% ± 23.87
Entraînements performants				
MSE & Dirichlet (5)	06.43% ± 00.01	10.15% ± 00.26	<u>51.40% ± 01.54</u>	<u>56.37% ± 01.22</u>
CE & Dirichlet	06.95% ± 00.09	11.59% ± 00.27	50.52% ± 02.08	58.92% ± 01.18
CE & Softmax (0)	X	X	X	X
MSE (0)	X	X	X	X
<u>MSE & Prop (2)</u>	<u>06.59% ± 00.11</u>	<u>10.70% ± 00.40</u>	51.86% ± 02.18	<u>55.49% ± 00.89</u>
Meilleur entraînement				
MSE & Dirichlet	06.30%	09.94%	52.72%	56.82%
CE & Dirichlet	06.90%	11.33%	52.18%	60.70%
CE & Softmax	X	X	X	X
MSE	X	X	X	X
<u>MSE & Prop</u>	<u>06.48%</u>	<u>10.30%</u>	54.04%	<u>56.86%</u>

4.5 Conclusion

Nous avons proposé dans ce chapitre une méthode générique d’inférence de proportions. Cette dernière a été appliquée avec succès à trois problèmes distincts.

Nous avons proposé deux fonctions de perte spécifiques basées sur la modélisation de Dirichlet, “MSE & DIR” décrit par l’équation (4.13) et “CE & DIR” décrit par l’équa-

tion (4.14). Elles ont été comparées avec d’autres fonctions plus courantes : “MSE” (équation (4.15)), “CE & Softmax” (équation (4.16)) et “MSE & Prop” (équation (4.17)), et lorsque cela a été possible, avec des méthodes concurrentes de *machine learning* ou de *deep learning*.

D’après les trois expériences menées sur les signaux de diffraction, les spectres des images hyperspectrales et les images MNIST, la modélisation à l’aide de la distribution de Dirichlet semble avoir un impact à la fois sur les performances et sur la stabilité. Tout d’abord comme le montre l’exemple des images hyperspectrales et la comparaison avec des méthodes concurrentes, l’utilisation d’un réseau dont la tâche spécifique est uniquement l’inférence de proportions surpasse d’autres approches dont les tâches sont multiples.

Ensuite, la “MSE & DIR” surpasse de manière significative les autres fonctions de perte, et notamment celles utilisant simplement les termes de minimisation (MSE et CE). On observe que la fonction “MSE & Prop” basée sur la même modélisation mais dont le terme de variance est retiré (équation (4.13)), présente également d’excellentes performances lorsque les entraînements fonctionnent. Au vu des résultats des expériences, il apparaît que le terme de variance apporte une stabilité au modèle. Lorsqu’on la retire dans la fonction “MSE & Prop”, un nombre important d’entraînements ne fonctionne pas.

En s’appuyant sur des modèles dont les architectures sont déjà performantes pour la classification et en choisissant une fonction de perte adaptée, notre approche est capable de quantifier les composants efficacement. Nous constatons aussi à travers les différentes expériences que les résultats varient en fonction de la complexité de la donnée, avec des erreurs inférieures à 5% pour des signaux et inférieures à 10% sur des images.

Finalement, pour la quantification de phases minérales à partir des diffractogrammes, l’application qui a motivé ces travaux, les résultats sont excellents. La modélisation proposée pour les $K = 4$ phases minérales montre notamment des erreurs de l’ordre de 2% sur la base de test simulée. Ces résultats montrent l’efficacité de la méthode, et laissent envisager le traitement de données expérimentales.

Chapitre 5

DRX : Analyse d'un jeu de données de laboratoire

5.1 Les données

Les bons résultats pour la quantification de phases minérales sur des diffractogrammes simulés présentés dans le chapitre précédent 4, laissent penser qu'un problème plus complexe peut être traité avec une approche similaire.

Une base de données expérimentale a été créée dans cette optique. Ces données sont acquises dans un laboratoire du BRGM, dont le montage expérimental est décrit dans la section "Acquisition de données en laboratoire" du chapitre 2. Elles sont acquises sur une plage angulaire entre 4.0001° à $90.020\,055^\circ$ (selon 2θ), par pas de $0.029\,611^\circ$, pour un total de 2 905 points d'acquisition.

Les diffractogrammes sont des mélanges des quatre phases minérales de la section précédente : Calcite (MARKGRAF & REEDER 1985) (CaCO_3), Dolomite (STEINFINK & SANS 1959) (CaMgC_2O_6), Gibbsite (BALAN *et al.* 2006) (AlO_3H_3) et Hématite (BLAKE *et al.* 1966) (Fe_2O_3). La composition précise des 32 échantillons est donnée par le tableau 5.1. La figure 5.1 montre deux exemples de ces diffractogrammes.

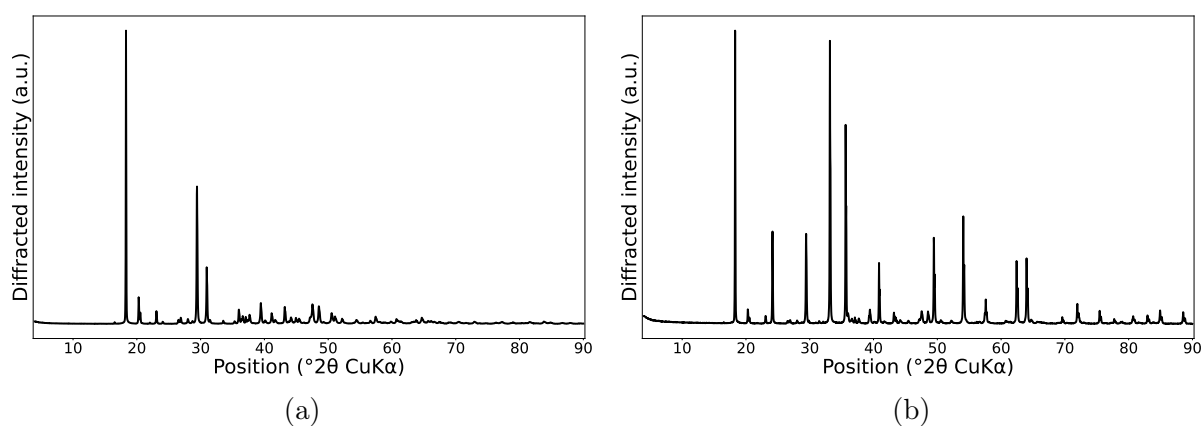


FIGURE 5.1 – Exemples de diffractogrammes expérimentaux dont la composition est proche de : (a) 40% Calcite, 40% Gibbsite 20% Dolomite et (b) 20% Calcite, 20% Gibbsite 60% Hématite. La composition précise est donnée par le tableau 5.1 : (a) correspond à l'échantillon 11, et (b) à l'échantillon 13.

Pour rappel, l'objectif est d'identifier et de quantifier les phases minérales dans cha-

TABLE 5.1 – Composition en termes de phases minérales des 32 données expérimentales acquises en laboratoire. La proportion massique de chaque phase dans le mélange est déterminée via pesée avant mélange.

Données	Calcite	Gibbsite	Dolomite	Hematite
Echantillon 1	0	1	0	0
Echantillon 2	0.203	0.797	0	0
Echantillon 3	0.403	0.597	0	0
Echantillon 4	0.600	0.400	0	0
Echantillon 5	0.800	0.200	0	0
Echantillon 6	1	0	0	0
Echantillon 7	0.204	0.2001	0.595	0
Echantillon 8	0.402	0.200	0.398	0
Echantillon 9	0.602	0.201	0.196	0
Echantillon 10	0.201	0.402	0.397	0
Echantillon 11	0.401	0.399	0.200	0
Echantillon 12	0.202	0.599	0.199	0
Echantillon 13	0.202	0.200	0	0.598
Echantillon 14	0.402	0.198	0	0.399
Echantillon 15	0.599	0.201	0	0.199
Echantillon 16	0.199	0.401	0	0.400
Echantillon 17	0.401	0.400	0	0.199
Echantillon 18	0.202	0.597	0	0.200
Echantillon 19	0.204	0.199	0.197	0.401
Echantillon 20	0.400	0.202	0.199	0.199
Echantillon 21	0.204	0.398	0.199	0.199
Echantillon 22	0.201	0.202	0.398	0.199
Echantillon 23	0	0	1	0
Echantillon 24	0	0	0	1
Echantillon 25	0.824	0	0	0.176
Echantillon 26	0.838	0	0.162	0
Echantillon 27	0.234	0.766	0	0
Echantillon 28	0	0.736	0	0.264
Echantillon 29	0	0.790	0.210	0
Echantillon 30	0.202	0	0	0.798
Echantillon 31	0	0	0.220	0.78
Echantillon 32	0	0.174	0	0.826

cune des données expérimentales. Encore une fois, cela revient à chercher le vecteur de proportion $\mathbf{p}$ de l'équation (4.1). Le plus fréquemment, dans la littérature scientifique, cette tâche est effectuée grâce à l'affinement Rietveld (RIETVELD 1969).

5.1.1 Rietveld

Cette méthode minimise point par point la différence entre le diffractogramme réel, et celui qu'elle calcule. Pour cela, Rietveld prend des paramètres initiaux afin de proposer une première reconstruction du signal, puis à chaque itération, l'algorithme minimise la

différence entre la reconstruction et le signal original en optimisant les paramètres choisis initialement. Afin de faciliter la convergence du modèle, il est important que l’initialisation soit proche de la solution. Une connaissance assez précise de l’échantillon est donc nécessaire. En particulier, toutes les phases composant le mélange doivent être connues. Cependant, étant donné le faible nombre de phases minérales dans les échantillons de cette expérience, il n’est pas obligatoire de connaître avec précision chaque donnée. L’algorithme convergera de manière précise avec seulement quatre phases minérales en entrée.

Grâce au logiciel Profex (DOEBELIN & KLEEBERG 2015), nous avons analysé les 32 diffractogrammes expérimentaux avec la méthode d’affinement Rietveld, comme le montre la dernière partie du tableau 5.3. Pour ces diffractogrammes, l’affinement Rietveld est capable de quantifier au pourcentage près les phases minérales. Et pour les 32 diffractogrammes, la méthode retrouve le bon support ($\text{RRS}_{1\%} = 100\%$), i.e les composants sont bien identifiés.

La comparaison avec le modèle de réseaux de neurones que nous proposons ne sera pas faite par la suite, car les méthodes ont des objectifs différents. La méthode Rietveld est très performante pour quantifier et ajuster les paramètres de maille, si le nombre de phases minérales est faible. Elle l’est également si l’échantillon est déjà connu de manière précise, par exemple si les phases minérales et une bonne approximation de leurs paramètres de maille sont connues.

Cependant, lorsque le nombre de phases minérales dans l’échantillon est grand, alors le nombre de paramètres à affiner pour proposer une reconstruction pertinente du diffractogramme devient important et la convergence de l’algorithme devient dépendante de l’initialisation.

Il est nécessaire de développer une méthode alternative, lorsque le nombre de phases dans le jeu de données est plus important, et que les données ne sont que peu connues. Nous verrons un exemple de ce type de jeu de données dans le chapitre 6 qui concerne l’analyse d’un jeu de données obtenu par tomographie de rayons X. Ainsi, nous proposons d’utiliser les réseaux de neurones pour effectuer cette tâche, et éventuellement venir en complément de Rietveld.

5.1.2 Une base de test pour l’approche des réseaux de neurones

Nous revenons donc à notre approche de *deep learning* et à la manière dont les données expérimentales seront utilisées. Elles serviront uniquement à tester les capacités du réseau une fois entraîné. Pour la base d’entraînement, il est nécessaire d’avoir un nombre de données important, qui couvre l’ensemble des variations naturelles. Cependant, expérimentalement, il est impossible d’obtenir un nombre de données suffisant et surtout, avec une telle variabilité. Même si cela était possible, c’est-à-dire qu’une bonne quantité d’échantillons serait disponible, alors le temps d’acquisition serait une nouvelle limite.

Une approche utilisant la *data augmentation* est également compliquée, même si cela permettrait d’augmenter la taille de la base à partir des 32 données expérimentales. Mais dans le cas de maille non cubique (c’est le cas avec les quatre phases minérales considérées), la dépendance entre les pics n’est pas linéaire. La variation d’un ou plusieurs paramètres de maille ne correspond pas à une translation du signal. C’est principalement ce facteur qui limite l’approche de *data augmentation*.

C’est pourquoi le réseau ne peut pas être entraîné sur des données réelles. L’objectif est donc de développer une méthode se basant sur des simulations pour quantifier les phases minérales sur des données réelles.

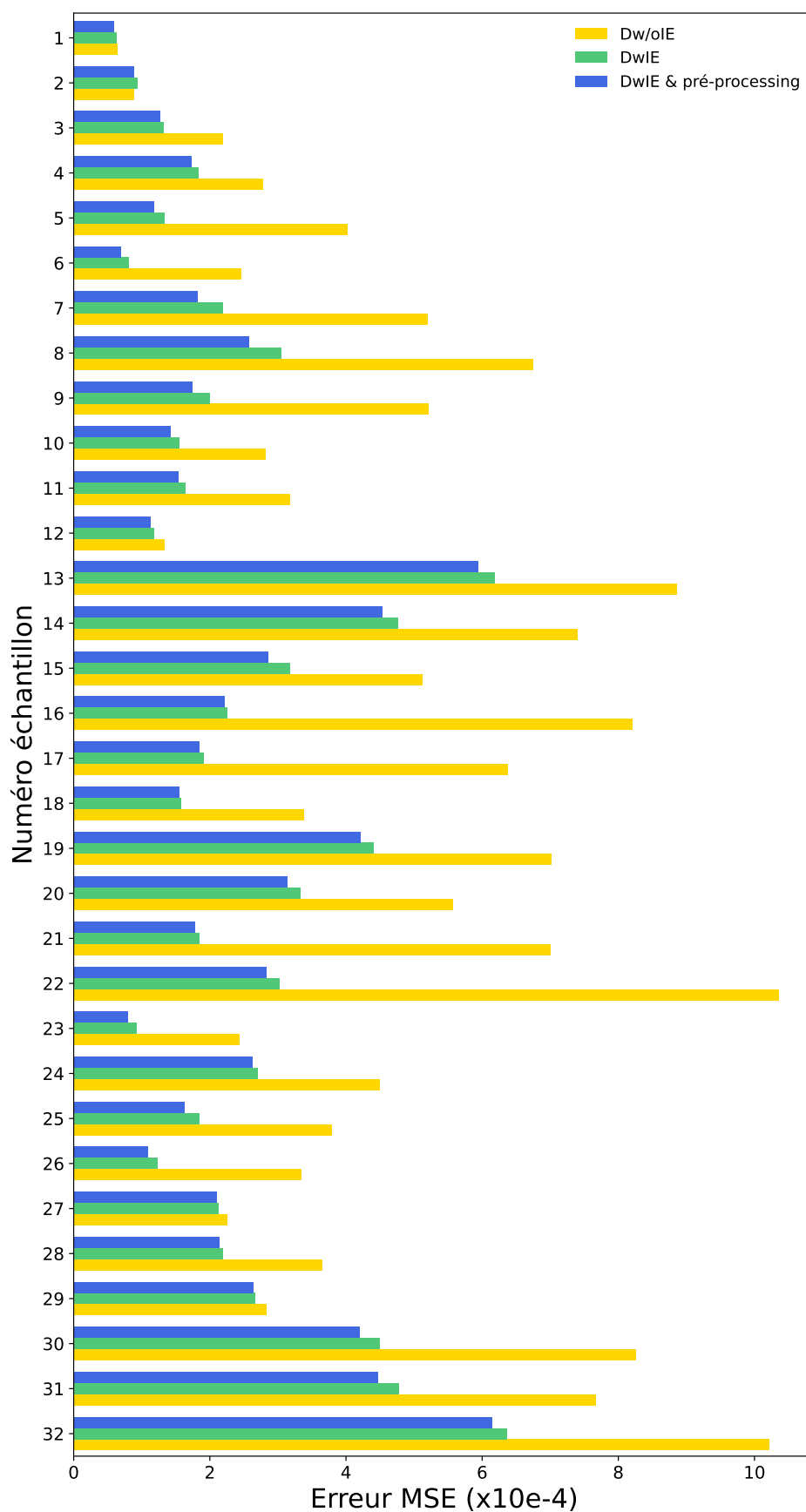


FIGURE 5.2 – Comparaison de l’erreur des moindres carrés entre les simulations provenant des méthodes “Dw/oIE” et “DwIE” avec les données expérimentales. Une comparaison en ajoutant un pré-traitement du signal réel est ajoutée.

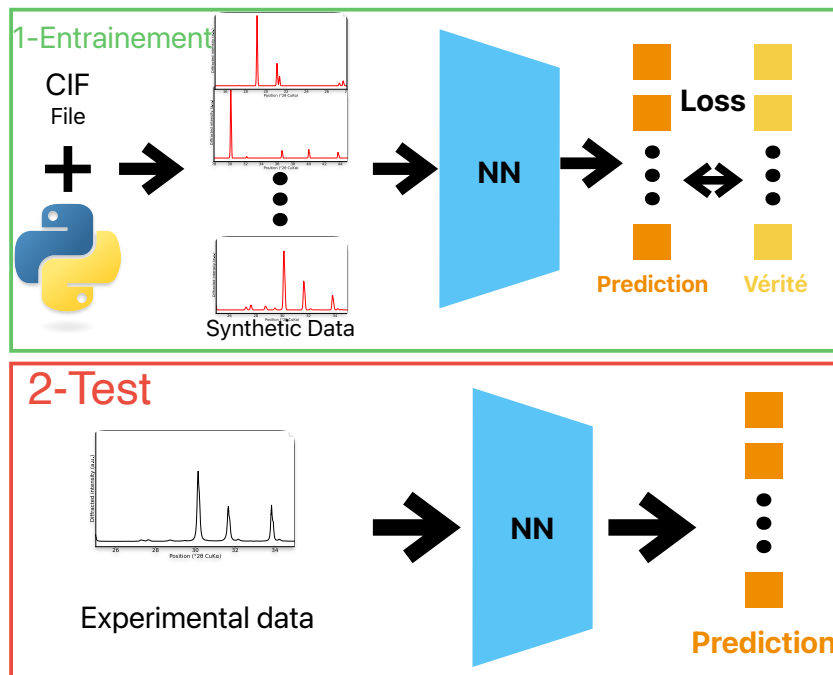


FIGURE 5.3 – Représentation schématique de la méthode proposée pour l’analyse de données expérimentales.

5.2 Méthode proposée

L’approche développée est schématisée par la figure 5.3. L’objectif est d’entraîner le réseau sur des simulations, puis de le confronter aux données réelles dans un second temps.

La première étape est donc de simuler les données pour l’entraînement. Nous pouvons signaler que les phases contenues dans les données sont connues au préalable, ce qui est généralement le cas, quel que soit le jeu de données. Ainsi, connaissant les phases minérales, nous pouvons simuler des milliers de données grâce aux algorithmes présentés dans le chapitre 3. Avec les quatre phases (Calcite, Gibbsite, Dolomite et Hématite), nous pouvons utiliser le code simplifié (algorithme 1) qui offre une vitesse de calcul plus importante, même si les simulations sont moins précises.

D’après les résultats du chapitre précédent, avec la modélisation adaptée, un réseau de neurones est capable de quantifier les phases minérales à partir du diffractogramme avec une erreur inférieure à 2%. Mais les données réelles sont plus complexes que les simulations. L’enjeu est donc de simuler des données le plus proche possible des diffractogrammes expérimentaux afin d’avoir les meilleures performances possibles.

5.2.1 Comparaison de deux simulations

Afin de voir comment la qualité des données d’entraînement influe sur le modèle, deux types de simulations sont comparées. La première est une base constituée uniquement à l’aide du code rapide et simplifié (algorithme 1), sans effets instrumentaux. Elle est nommée “Dw/oIE” par la suite (pour *Database WithOut Instrumental Effects*).

La seconde est simulée avec le même code mais inclut les effets instrumentaux : la fonction d’onde et le facteur d’absorption présentés dans la section 2.4.3 du premier chapitre.

Cette base est nommée “DwIE”, qui correspond à *Database With Instrumental Effects*.

Une comparaison entre les deux bases est d’abord effectuée par rapport aux données expérimentales. Les graphiques de la figure 2.13 du chapitre 2 montrent déjà l’apport des effets instrumentaux sur des diffractogrammes de simple phases.

Pour pousser plus loin la comparaison, nous utilisons cette fois les 32 données à disposition. Elles correspondent à des échantillons de simple phase ou à des mélanges de plusieurs phases. Le but est de synthétiser des signaux de diffraction le plus proche possible des données réelles. Pour chaque diffractogramme expérimental, nous simulons un signal selon les deux méthodes (“Dw/oIE” et “DwIE”). Pour ces simulations, les paramètres de maille ont été ajustés grâce au logiciel Profex et au refinement Rietveld. Le facteur d’agitation thermique est identique pour l’ensemble des simulations. La valeur B de l’équation (2.3) est fixée à 0.2 et la largeur à mi-hauteur est également gardée constante avec une valeur de 0.07. Enfin, les coefficients de la combinaison linéaire pour les mélanges de phases sont connus grâce au tableau 5.1.

Pour faciliter la comparaison, les données (Expérimentales, “Dw/oIE” et “DwIE”) sont normalisées de telle sorte que la valeur maximum soit 1. Nous rappelons que comme pour la classification, seule la forme du signal est intéressante (voir section 3.2).

La figure 5.2 compare l’erreur des moindres carrés pour chaque type de simulation avec le signal mesuré en laboratoire. Il apparaît clairement que les effets instrumentaux permettent aux simulations d’être plus réalistes. Les erreurs sont systématiquement (sauf sur l’échantillon 2) plus faibles avec la base “DwIE” (en vert). Nous remarquons que la différence avec les effets instrumentaux est plus importante sur les données de mélange, cela s’explique avec le facteur d’absorption. Il permet d’avoir les bons rapports d’intensité entre les signaux provenant des différentes phases.

Sur la figure 5.2 a été ajoutée, en bleu, la comparaison avec “DwIE” en réalisant un traitement de la donnée expérimentale. La donnée après traitement, notée $\mathbf{x}^n$, est obtenue par :

$$\mathbf{x}_i^n = \mathbf{x}_i - \min(\mathbf{x}_i). \quad (5.1)$$

Ce traitement correspond au retrait d’un fond continu sur le signal. En effet, divers effets instrumentaux contribuent au signal en y ajoutant un bruit de fond, qui n’est pas présent dans les simulations. Etant donné que ce bruit de fond a une origine complexe, nous avons choisi d’essayer de le retirer aux données réelles plutôt que de l’ajouter aux simulations. Avec la normalisation des signaux, ce sont les seuls traitements appliqués aux données réelles.

Toujours d’après la figure 5.2, ce traitement de la donnée réelle permet de réduire encore un peu plus la différence avec la simulation. Ainsi, avant de tester les données expérimentales, ce fond continu leur sera retiré, ce qui permettra d’améliorer les performances du modèle.

Des bases de données d’entraînement ont donc été générées suivant les deux simulations : “DwIE” et “Dw/oIE”. Chacune servira à entraîner le réseau de manière indépendante.

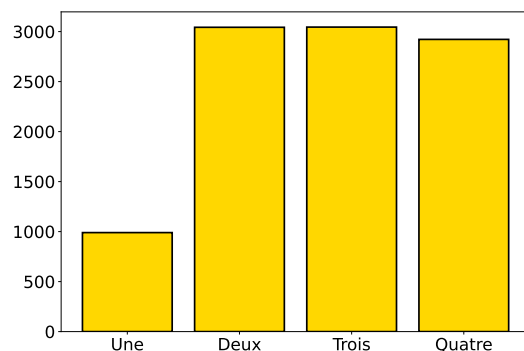


FIGURE 5.4 – Barplot de la répartition des diffractogrammes simulés de la base d’entraînement en fonction du nombre de phases présentes dans le mélange.

5.2.2 Génération des bases de données

La procédure pour générer les bases de données est quasiment identique à celle décrite dans le chapitre précédent par la figure 4.3. A partir de diffractogrammes simulés avec une unique phase minérale, des combinaisons linéaires sont effectuées pour produire les mélanges. Toujours dans le but de minimiser la corrélation entre les sous-ensembles, les données à simple phase qui génèrent chacune des sous-ensembles sont indépendantes. Au total 15 000 données ont été générées, toujours découpées en trois sous-ensembles : entraînement, validation et test. Cette procédure a été effectuée pour chacune des bases : “DwIE” et “Dw/oIE”.

La première différence par rapport à la section 4.2 du chapitre précédent concerne la simulation des phases simples. Les paramètres de maille (a, b, c) sont tirés aléatoirement et de manière uniforme avec 2% (et non plus 5% comme précédemment) de variation autour des valeurs fournies par les fichiers CIF, afin de se rapprocher d’une variation des signaux réaliste.

La seconde différence concerne le nombre k de classes dans les mélanges qui n’est plus tiré de manière uniforme. Un biais est ajouté afin de limiter le nombre de diffractogrammes avec une seule phase minérale. La répartition est donnée par le diagramme en barre de la figure 5.4. Nous pouvons voir que les diffractogrammes simple phase sont trois fois moins présents dans la base d’entraînement que ceux qui mélangent deux, trois ou quatre phases, qui sont représentés de manière équilibrée.

Finalement, les réseaux sont entraînés cinq fois pour chaque base d’entraînement, avec la même initialisation. L’architecture est toujours celle de la figure 3.4. La fonction de perte retenue est la “MSE & Dirichlet”, choisie à la fois pour les performances et la stabilité. Comme pour le chapitre précédent, les données sont regroupées par batch de taille 4, et l’entraînement se déroule en 100 époques. L’optimiseur est l’algorithme Adam, avec un taux d’apprentissage constant et égal à 0.001.

Nous gardons toujours une base de validation et une base de test simulées pour vérifier que l’entraînement fonctionne correctement. L’entraînement du réseau, et les résultats sur la base de validation peuvent être suivis sur la figure 5.5. Sur ces graphiques, seul l’entraînement avec la meilleure MMAE sur la base de validation est retenu.

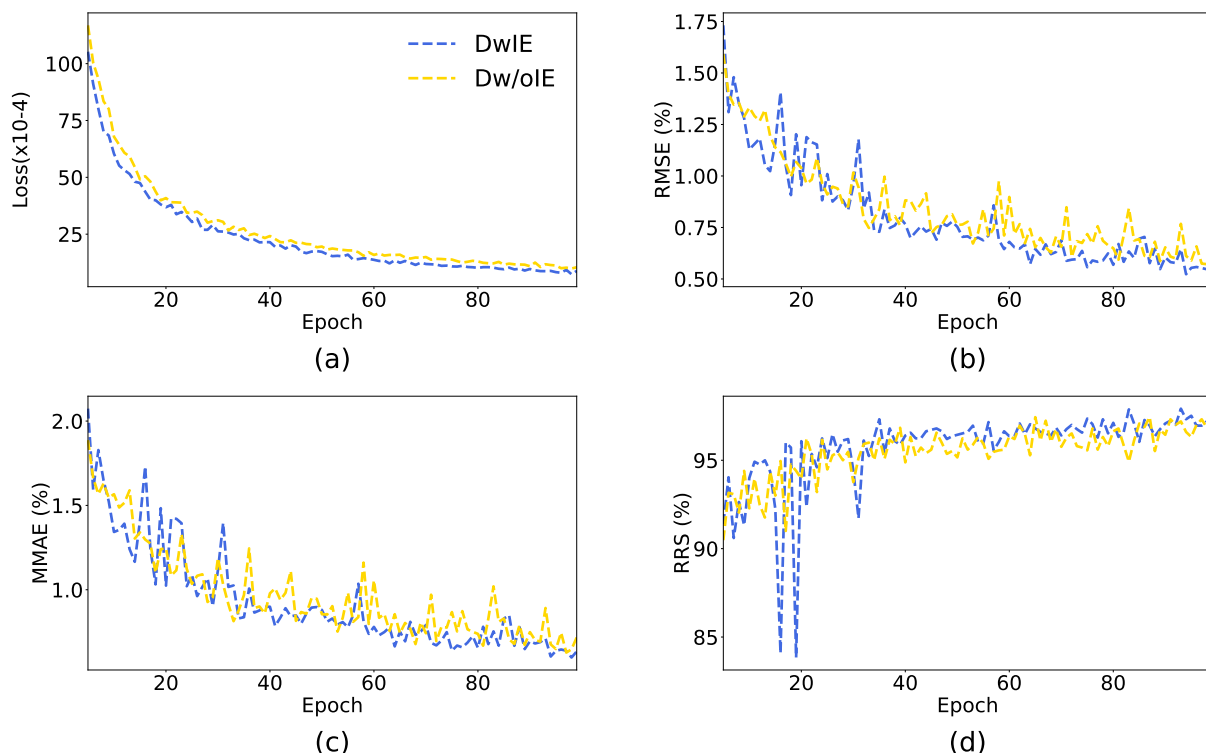


FIGURE 5.5 – Analyse de diffractogrammes simulés : (a) Evolution des valeurs de la *loss* en fonction de la base d’entraînement, avec (b), (c) et (d) qui décrivent les valeurs des mesures (RMSE, MMAE et RRS respectivement) sur l’ensemble de validation au cours des époques d’entraînement.

Les résultats montrent une bonne convergence des fonctions de perte, avec une décroissance qui reste marquée même dans les dernières époques. On remarque aussi que les trois mesures d’erreur (RMSE, MMAE et RRS) continuent de progresser. Ceci laisse penser que le réseau pourrait être entraîné sur plus d’époques, et également avec un taux d’apprentissage programmé. C’est-à-dire en ajustant sa valeur en fonction de l’époque. Des travaux sur ce sujet sont présentés à la fin du chapitre.

Pour les résultats obtenus, la RMSE et la MMAE se rapprochent des 0.5%, et la RRS dépasse les 95%. Les gains par rapport au chapitre précédent s’expliquent car la variation des paramètres de maille est limitée à 2%. Les chiffres sont confirmés par le tableau 5.2 des résultats sur l’ensemble de test des diffractogrammes simulés. Nous pouvons également voir que la fonction de perte est toujours stable, bien que certains entraînements n’aient pas convergé : 2 pour la “Dw/oIE” et 1 pour la “DwIE”. La précision est également légèrement supérieure avec la base incluant les effets instrumentaux. Par exemple, sur le meilleur entraînement, les RMSE sont de 0.58% pour la “Dw/oIE” contre 0.49% pour la “DwIE”. De même pour la MMAE 0.62% pour 0.55% et la RRS 97.20% pour 97.40%. Ces chiffres peuvent s’expliquer par un écart plus important entre les intensités des simples phases, avec le facteur d’absorption, donc un meilleur discernement pour les quantifier de manière plus précise.

Nous allons maintenant voir comment les réseaux entraînés traitent les données expérimentales.

TABLE 5.2 – Résultats des différentes mesures (RMSE, MMAE et RRS) sur les ensembles de test simulés des bases de données “DwIE” et “Dw/oIE”. Le tableau est séparé en 3 parties, avec d’abord les moyennes et écarts types sur 5 entraînements. La deuxième partie ne retient que les entraînements performants (où les fonctions de perte ont convergé), leur nombre est indiqué entre parenthèses. Enfin, seul le meilleur entraînement est considéré, i.e celui avec la valeur de MMAE la plus faible sur l’ensemble de validation.

	RMSE ↓	MMAE ↓	RRS _{1%} ↑
5 entraînements			
Dw/oIE	10.44% ± 12.07	14.46% ± 16.91	70.16% ± 32.89
DwIE	05.49% ± 09.98	07.58% ± 14.00	83.49% ± 27.49
Entraînements performants			
Dw/oIE (3)	00.58% ± 00.00	00.65% ± 00.20	97.01% ± 00.24
DwIE (4)	00.50% ± 00.02	00.58% ± 00.03	97.22% ± 00.47
Meilleur entraînement			
Dw/oIE	0.58%	0.62%	97.20%
DwIE	0.49 %	0.55%	97.40%

5.3 Résultats sur les données expérimentales

La base de données expérimentales se compose de 32 diffractogrammes. Ils sont analysés avec des réseaux de neurones déjà entraînés sur des données simulées, comme cela a été décrit au début du chapitre. Cet ensemble sert de base de test pour le réseau. Les résultats peuvent être décrits de différentes manières.

Premièrement, le tableau 5.3 compare les mesures sur les différents types de simulation : “DwIE” et “Dw/oIE”. Les mesures sont également données pour l’affinement Rietveld afin d’illustrer les capacités de cette méthode. Comme nous pouvions l’attendre, sur les diffractogrammes expérimentaux, les résultats sont moins bons que sur les bases de test synthétiques. Et les entraînements qui n’avaient pas fonctionné sur les simulations ne fonctionnent pas non plus sur les données réelles.

Ensuite, en comparant les valeurs sur les entraînements performants nous constatons une nette différence entre les deux bases d’entraînements. La base “DwIE” montre des valeurs inférieures pour la MMAE, en moyenne 6.47% contre 17.08% pour la “Dw/oIE”. De même pour la RMSE où les valeurs sont respectivement 4.82% et 12.96%. Ces chiffres sont attendus puisqu’en améliorant la qualité de nos simulations, le réseau doit être plus performant. Cependant, les valeurs de la RRS avec un seuil ε de 1% montrent un net avantage pour la base “Dw/oIE”. Cet écart de plus de 20% peut s’expliquer par le facteur d’atténuation qui facilite la distinction des phases minérales, mais qui ne permet pas de faire une bonne identification. Sans ce facteur les rapports d’intensité entre phases ne sont pas bons mais permettent une meilleure identification.

Une bonne manière de voir cet effet est de considérer la RRS à une valeur proche de la MMAE, en fixant le seuil $\varepsilon = 7\%$. Cette fois, c’est la base “DwIE” qui présente la RRS la plus haute. Nous pouvons aussi regarder les histogrammes des fréquences cumulatives de la figure 5.7 qui comparent les deux bases de données sur les erreurs absolues. La base présentant les meilleurs résultats est celle dont l’histogramme atteint le plus vite sa valeur maximale. Cela veut dire que le maximum des erreurs est atteint à une valeur plus faible.

Cela permet de voir que les entraînements avec la “DwIE” restent plus performants malgré des valeurs de $RRS_{1\%}$ plus faibles.

TABLE 5.3 – Résultats des différentes mesures (RMSE, MMAE et RRS) sur les 32 données expérimentales en fonction de la base d’entraînement. Le tableau est séparé en 4 parties, avec d’abord les moyennes et écarts types sur 5 entraînements. La deuxième partie ne retient que les entraînements performants (où les fonctions de perte ont convergé), leur nombre est indiqué entre parenthèses. Ensuite, seul le meilleur entraînement est considéré, i.e celui avec la valeur de MMAE la plus faible sur l’ensemble de validation. Enfin sont présentés les résultats du refinement Rietveld.

	RMSE ↓	MMAE ↓	$RRS_{1\%}$ ↑	$RRS_{7\%}$ ↑
5 entraînements				
Dw/oIE	18.79% ± 07.15	26.92% ± 12.06	45.62% ± 27.36	49.09 ± 30.33
DwIE	09.37% ± 09.12	13.51% ± 14.12	38.75% ± 13.35	65.62% ± 27.67
Entraînements performants				
Dw/oIE (3)	12.96% ± 00.24	17.08% ± 00.36	67.71% ± 05.31	75% ± 02.55
DwIE (4)	04.82% ± 00.74	06.47% ± 01.15	45.31% ± 02.71	78.91% ± 08.71
Meilleur entraînement				
Dw/oIE	12.71%	16.68%	75.00%	75.00%
DwIE	05.16%	06.96%	43.75%	71.87%
Rietveld	01.27%	01.78%	100%	100%

Ces résultats peuvent être complétés avec la figure 5.6 qui permet de visualiser l’écart entre l’erreur absolue et la vérité. Plus les points sont proches de la droite identité, meilleure est la prédiction. Cela confirme l’efficacité de la base “DwIE” qui est capable de quantifier les phases minérales avec environ 5% d’erreur. Cependant au vu des expériences, nous pourrions imaginer quelques améliorations de la méthode.

5.4 Conclusion et discussion

L’objectif de cette première approche sur un jeu de données expérimentales est d’évaluer le potentiel de la méthode. Les premiers résultats obtenus sont très satisfaisants comme cela a été décrit dans la section précédente 5.3. Cependant, nous constatons que certaines pistes d’améliorations peuvent être considérées.

Premièrement, comme l’entraînement se fait avec des simulations, nous pourrions imaginer augmenter le nombre de données dans la base d’entraînement, voir si cela influe ensuite sur les performances en évaluant les données expérimentales. Le produit nombre d’époques × taille de la base d’apprentissage est gardé constant afin de ne pas augmenter les temps de calculs. Cela permet tout de même d’entraîner le réseau sur une diversité de données plus importante.

Les résultats du tableau 5.4 sur les données réelles montrent une légère augmentation des performances du réseau pour la MMAE, une nette amélioration sur la RRS et une légère baisse pour la RMSE. Finalement, nous pouvons considérer que cela augmente les performances du modèle. D’autant plus en regardant les deux premières lignes de la

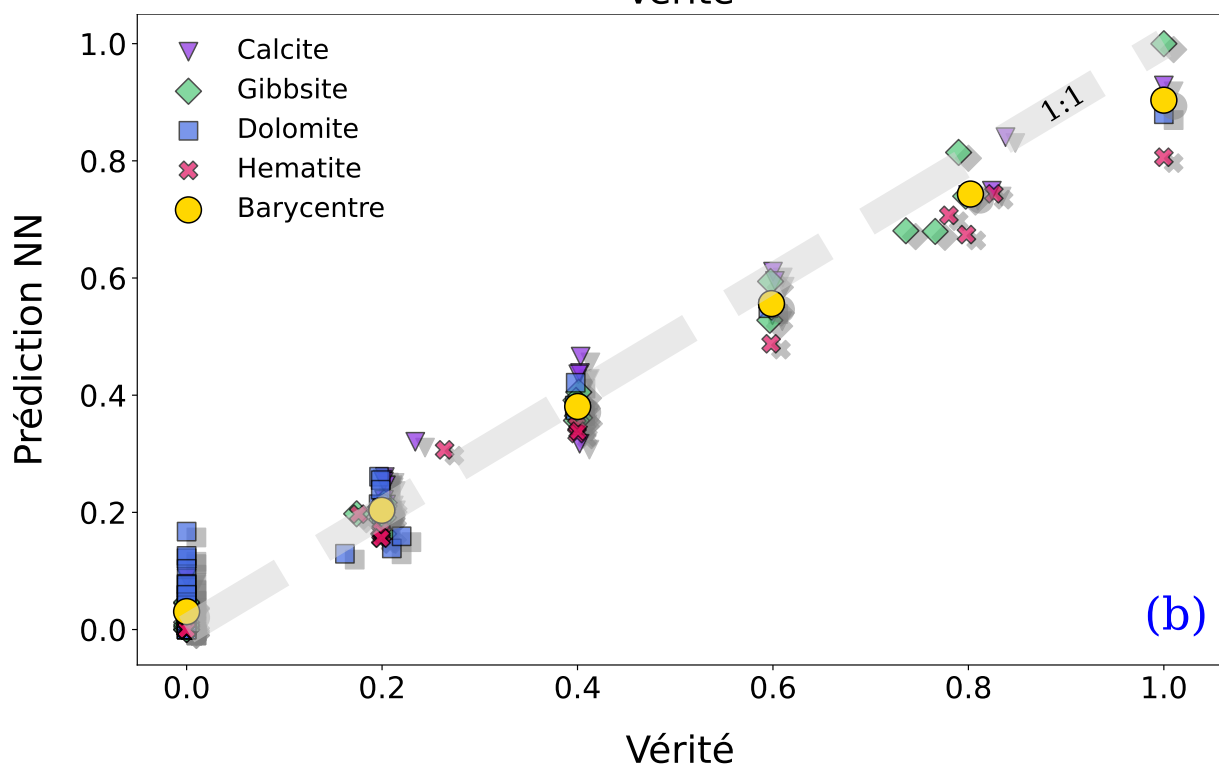
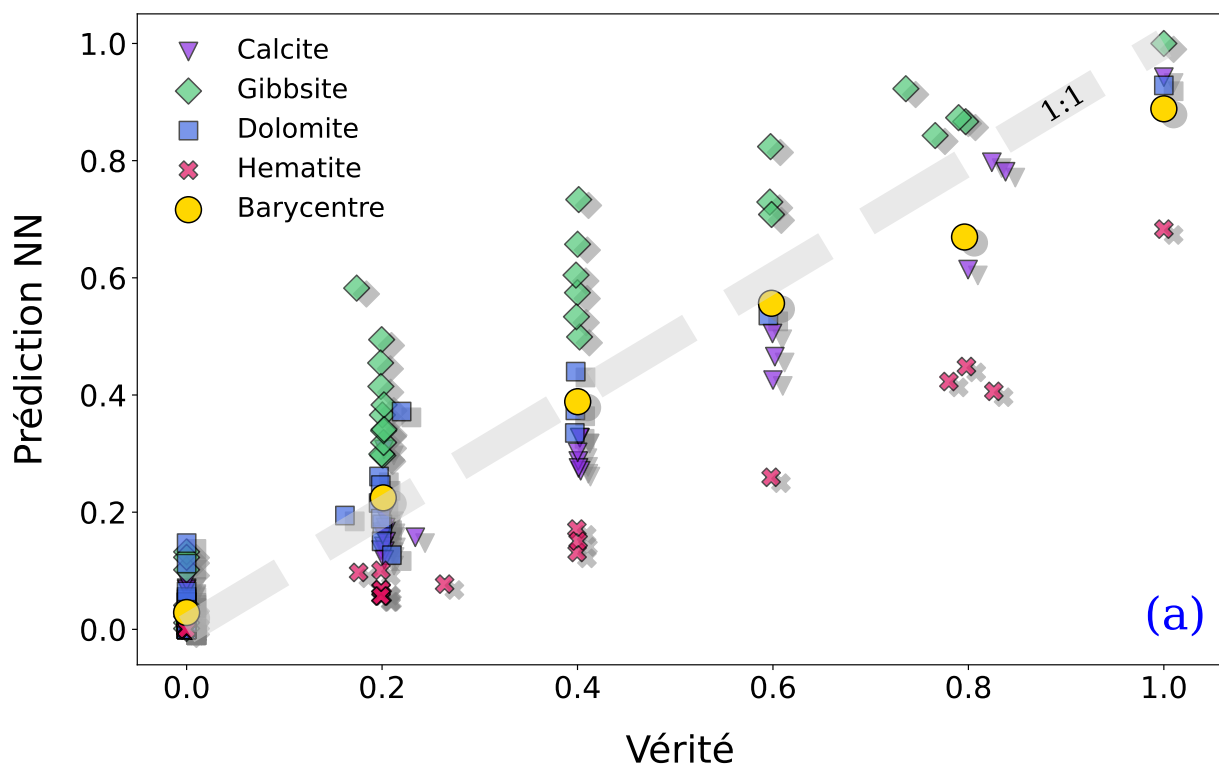


FIGURE 5.6 – Nuage de points des prédictions en fonction de la vérité. (a) est entraîné avec la base “Dw/oIE” et (b) avec “DwIE” .

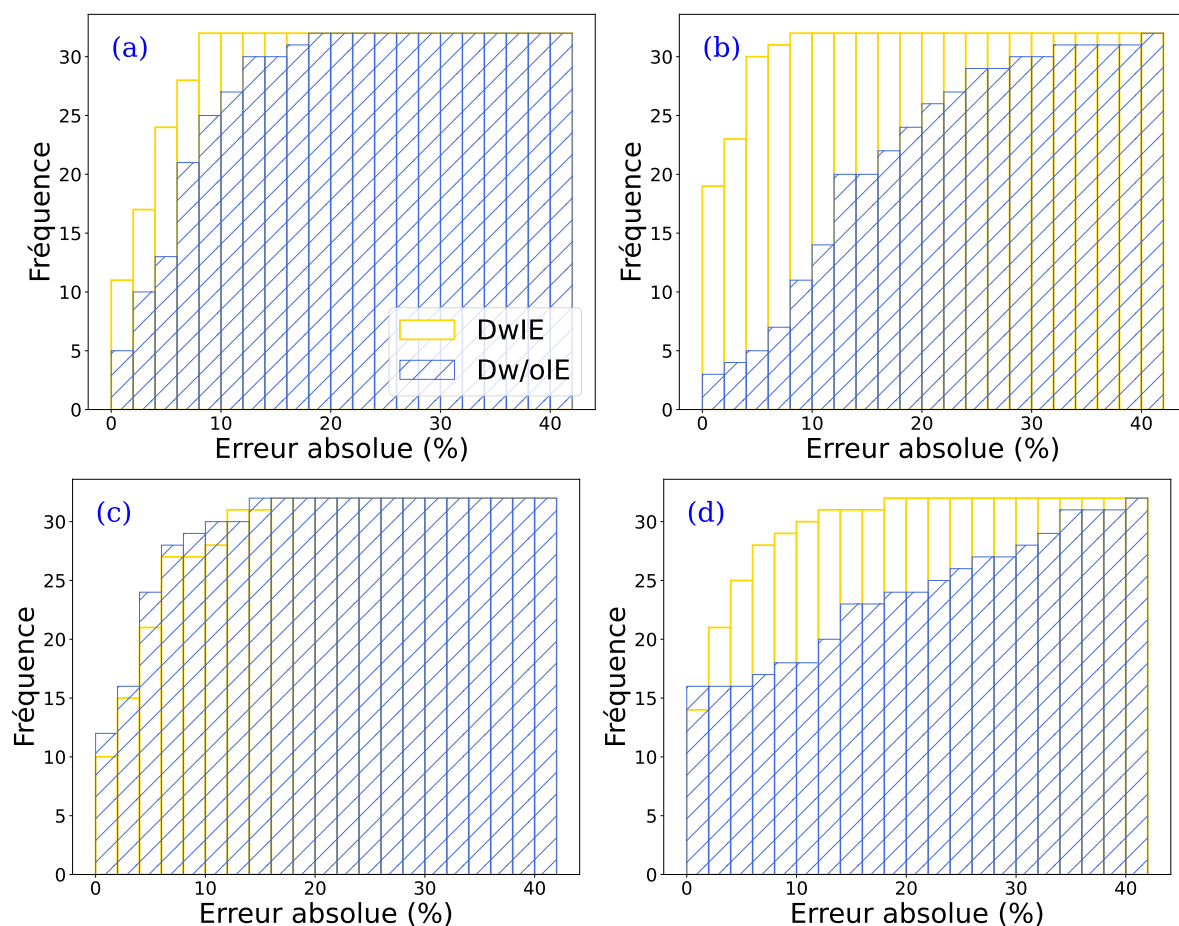


FIGURE 5.7 – Histogrammes cumulatifs de l'erreur absolue par classes pour les 32 données expérimentales : (a) Calcite, (b) Gibbsite, (c) Dolomite et (d) Hematite. Chaque barre de l'histogramme correspond à un intervalle de 2%, (de 0% à 40%).

TABLE 5.4 – Comparaison des performances sur l’ensemble de test des données expérimentales, en fonction de la taille de base d’entraînement et du nombre d’époques.

Trainset size	Epochs	RMSE	MMAE	RRS _{1%}
100	10 000	34.14	46.69	12.5
1 000	1 000	45.22	68.42	16.67
<u>10 000</u>	100	5.16	<u>6.96</u>	<u>43.75</u>
100 000	10	<u>5.84</u>	6.68	51.52

table, où la taille de la base d’entraînement est réduite ce qui mène à de très mauvaises performances pour quantifier les phases minérales.

Dans un second temps, nous avons augmenté le nombre d’époques d’entraînement, en utilisant un *learning rate* dont la valeur s’ajuste en fonction des époques. La valeur de ce dernier est fixé à 0.001 sur les premières époques, puis sa valeur est divisée par 10 toutes les 100 époques. Le réseau est entraîné sur 300 époques au total comme nous le montre la figure 5.8. Cette figure permet de voir les apports, avec une fonction de perte qui continue de décroître tout au long des 300 époques, tout comme les mesures sur la base de validation. A l’exception de la RRS avec un seuil $\varepsilon = 1\%$, qui semble légèrement décroître à partir de la 100^e époque. Pour ce qui est de la base de test expérimentale, cela améliore les résultats de quelques pourcents. La valeur de la RMSE est réduite à 4.75 %, la MMAE à 6.21 % et la $RRS_{7\%}$ atteint 87.5%.

D’autres pistes d’optimisation peuvent encore être apportées pour perfectionner le modèle. Nous pourrions notamment considérer d’autres types de modèle de *deep learning* tels que les **Transformers**.

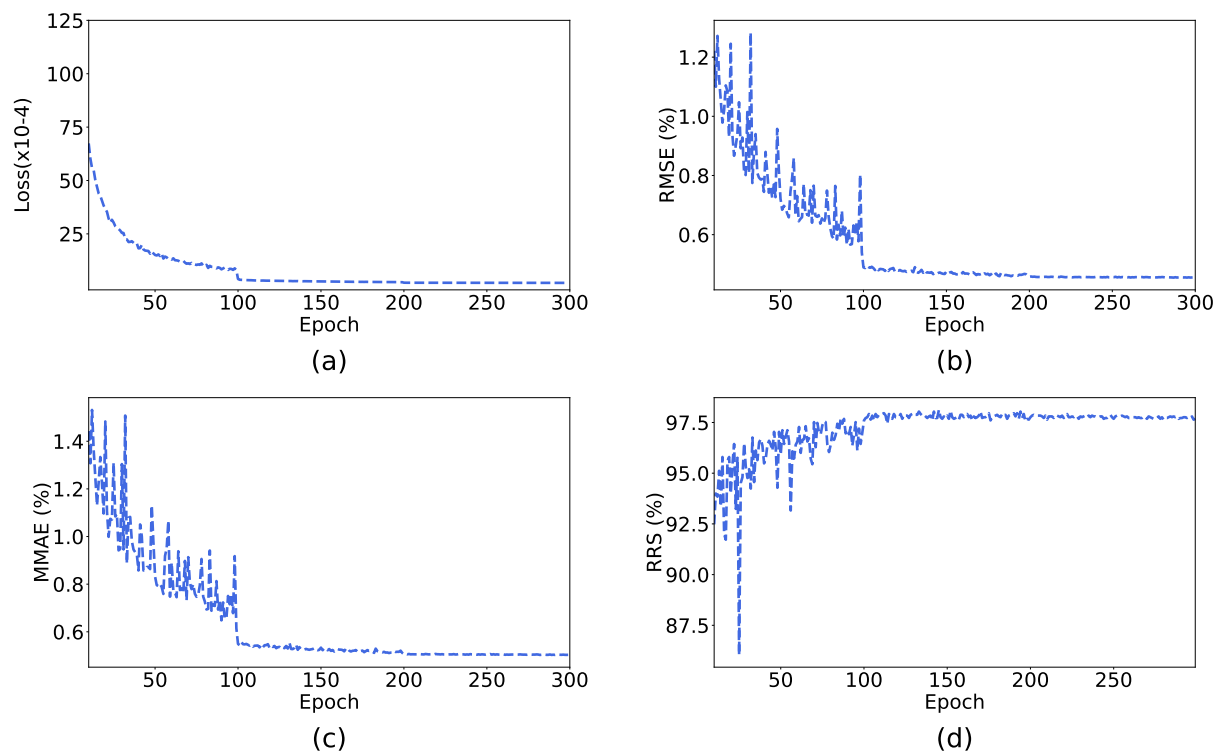


FIGURE 5.8 – Analyse de diffractogrammes simulés : (a) Evolution des valeurs de la *loss* en fonction de la base d’entrainement, avec (b), (c) et (d) qui décrivent les valeurs des mesures (RMSE, MAE et RRS respectivement) sur l’ensemble de validation au cours des époques d’entrainement.

Chapitre 6

DRX : Analyse d'un jeu de données tomographie

6.1 Description des données

La finalité de cette thèse est l'analyse de données provenant d'une analyse par tomographie par diffraction de rayons X. La méthode d'analyse a été développée plus en détail à la section 2.2.2.

L'échantillon analysé a été décrit par CLARET *et al.* (2018) qui proposent une analyse minéralogique qualitative. Il s'agit d'une carotte de ciment d'un diamètre de 1.6 mm sur une hauteur d'environ 1 cm. Les données dont nous disposons sont trois images de 250×250 pixels, mais comme montré sur la figure 6.2, la zone d'intérêt est un disque d'un diamètre de 85 pixels. Ces trois images correspondent à différentes tranches de la carotte analysée, d'une hauteur égale à celle du faisceau de rayons X (10 μm). Pour la suite des travaux, nous concentrons notre analyse sur une seule image.

La procédure pour obtenir une image est schématisée par la figure 6.1. Les résultats après analyse sont obtenus sous forme d'images, qui sont des diffractogrammes. C'est à partir de l'ensemble de ces diffractogrammes qu'un sinogramme est obtenu. Il contient toute l'information spatiale de l'échantillon, avec des couleurs plus chaudes lorsque l'intensité de diffraction est plus importante. Une projection spatiale du sinogramme pour obtenir une donnée spatialisée est ensuite effectuée grâce à une transformée de Fourier.

Les diffractogrammes sont, eux, calculés avec une moyenne azimutale à partir des images obtenues par la caméra. Chaque cercle de l'image correspond à un angle de diffraction et permet d'obtenir l'intensité à cet angle. Les rayons les plus petits correspondent aux petits angles de diffraction.

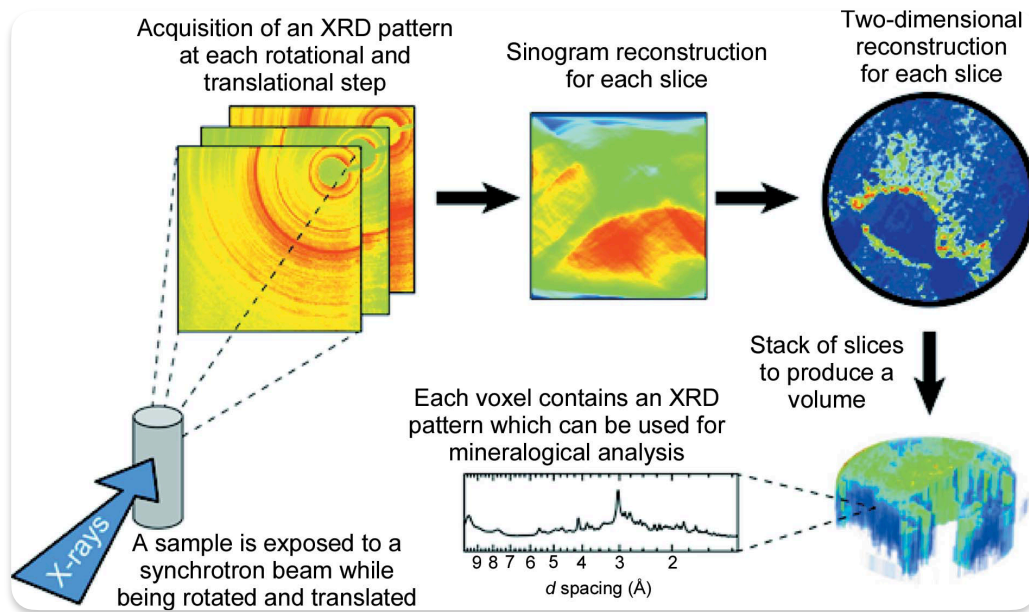


FIGURE 6.1 – Illustration de la procédure pour obtenir les diffractogrammes et les images du jeu de données de tomographie. Extrait de *Deciphering mineralogical changes and carbonation development during hydration and ageing of a consolidated ternary blended cement paste* de Claret *et al.* (CLARET *et al.* 2018).

Chacun des pixels du volume étudié contient un signal de taille 2048, qui couvre une plage angulaire (2θ) approximativement de 0.5° à 10° . Ces angles sont faibles mais avec une longueur d'onde des rayons X plus faible $\lambda = 0.1897\text{\AA}$ (pour $\lambda = 1.5418\text{\AA}$ sur le diffractomètre D8 Bruker). D'après la loi de Bragg (équation (2.5)), cela permet de couvrir une plage de distances inter-réticulaire (d), qui sera comprise entre $18.5\text{\AA}$ et $1.07\text{\AA}$, comparable aux données de laboratoire sur poudre du chapitre 5. Un total de 10 phases minérales ont été identifiées grâce à des précédents travaux (CLARET *et al.* 2018) :

- Alite (ANGELES *et al.* 2008b) (Ca_3SiO_5),
- Calcite (MARKGRAF & REEDER 1985) (CaCO_3),
- C-S-H (GRANGEON *et al.* 2013) ($\text{Ca}_4\text{Si}_6\text{O}_{18}\text{W}_2$),
- Ettringite (MOORE & TAYLOR 1970b) ($\text{Al}_2\text{Ca}_6\text{O}_{49.68}\text{S}_3\text{H}_{63.36}$),
- Hydrotalcite (ALLMANN & JEPSEN 1969) ($\text{Mg}_{0.667}\text{Al}_{0.333}\text{O}_3\text{H}_3\text{C}_{0.167}$),
- Monosulfoaluminate (ALLMANN 1977) ($\text{Ca}_2\text{AlO}_{11}\text{H}_{12}\text{S}_{0.5}$),
- Mullite (ANGEL *et al.* 1991) ($\text{Al}_{2.4}\text{Si}_{0.6}\text{O}_{4.8}$),
- Portlandite (HENDERSON & GUTOWSKY 1962) (CaO_2H_2),
- Quartz (LEVIEN *et al.* 1980) (SiO_2),
- Amorph (GRANGEON *et al.* 2013) ($\text{Ca}_4\text{Si}_6\text{O}_{18}\text{W}_2$),

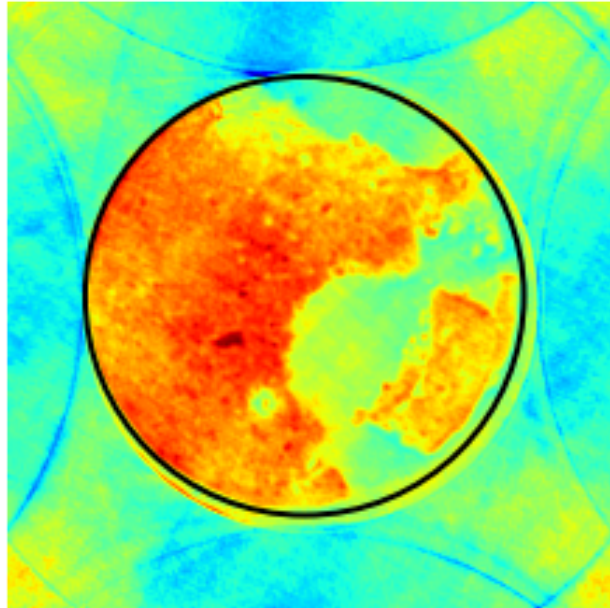


FIGURE 6.2 – Visualisation de l’une des images issues du jeu de données tomographie avec la colormap *jet*.

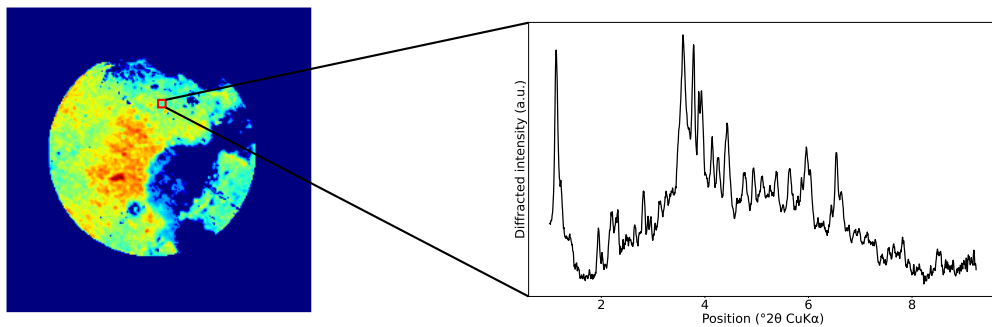


FIGURE 6.3 – Exemple d’un diffractogramme contenu dans un pixel d’une image du jeu de données tomographie. La taille du rectangle de sélection (en rouge) ne correspond pas à la dimension réelle d’un pixel de l’image.

6.1.1 Traitement des diffractogrammes

Le traitement préliminaire des données que nous décrivons ci-dessous comporte plusieurs étapes.

Extraction du disque de données Cette première étape, à partir de l’image originale, vise à extraire au mieux les données d’intérêt. Nous analysons une carotte de béton quasi-circulaire, incluse dans un porte échantillon en plastique (capillaire en polymide, qui apparaît à l’image). Il s’agit donc de supprimer les signaux qui correspondent au plastique,

ainsi que les artefacts de projections inhérents à l’usage de la transformée de Fourier d’un signal fini, qui se trouvent dans ce type de données rejetés à l’extérieur de l’échantillon. Nous avons choisi de retenir les données contenues dans un cercle centré en (119,119) et de rayon 85 pixels. Avec une taille de pixel de $10\mu m$, cela correspond à un diamètre d’échantillon de 1.7 mm , proche de ce qui était recherché lors du carottage (1.6 mm).

Retrait des signaux parasites Des signaux de vide sont au sein de l’échantillon et correspondent à de la porosité, en bleu foncé dans le disque de la figure 6.3. Un exemple de ces signaux est donné par la figure 6.4. Pour retirer ces signaux, des filtres sont ajoutés. Le premier enlève les signaux dont la valeur minimale est inférieure à zéro, et le second retire ceux dont la valeur moyenne est inférieure à une certaine intensité. La valeur seuil de l’intensité est fixée à 2 000, ce qui correspond approximativement à 10% de l’intensité des diffractogrammes sans porosité pour cet échantillon.

Une autre partie des signaux a été retirée à cause d’un monograin d’Alite dont la diffraction ne peut pas être gérée lors du passage du sinogramme à une image spatialisée.

Au final, ce sont 15 561 diffractogrammes qui sont retenus pour une image.

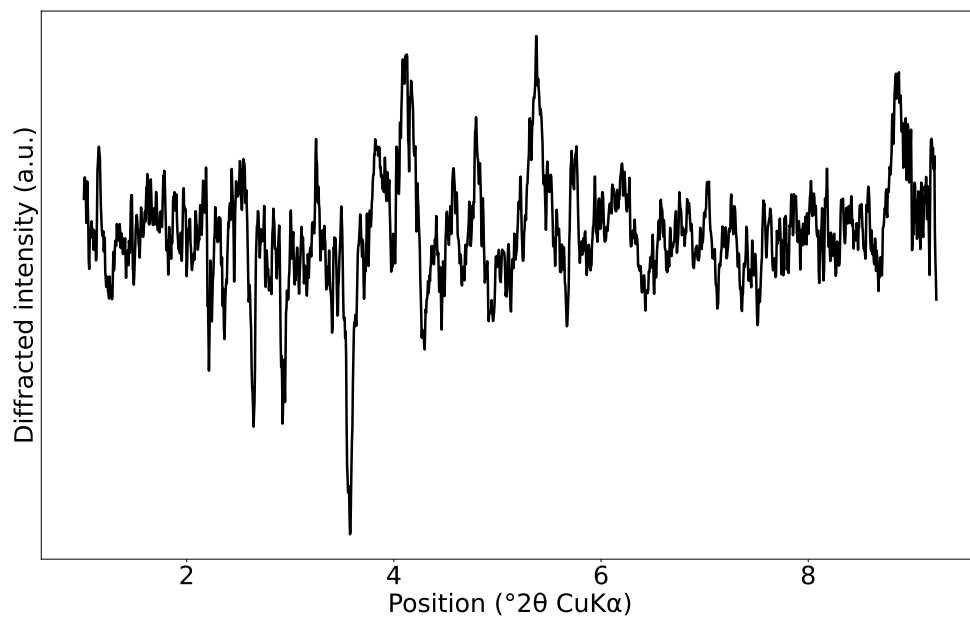


FIGURE 6.4 – Exemple de signal issu de la base de données tomographie correspondant à de la porosité.

Suppression des artefacts à petits et grands angles Il ne reste, désormais, que les signaux correspondant à de la matière. Il faut maintenant traiter ces diffractogrammes pour retirer certains artefacts présents à grands et petits angles en 2θ . Ainsi les premiers et derniers canaux sont retirés, pour ne retenir finalement que 1748 points sur chaque signal sur les 2048 originaux. Les artefacts à bas angles sont liés à l’empreinte du faisceau direct sur le récepteur.

Pour les grands angles, cela s’explique par la méthode d’acquisition des données à partir des diffractogrammes 2D. Les diffractogrammes 1D sont obtenus par moyenne azimutale.

Mais comme cela est affiché sur la figure 6.1, l'information est plus rare à grands angles à cause d'un effet de troncature (seules les données des angles du détecteur sont disponibles). Ainsi pour limiter ces artefacts, les 200 derniers canaux sont retirés.

Normalisation et retrait d'un fond continu Comme dans le cas des données de laboratoire, un fond continu est retiré, toujours selon l'équation (5.1) où la valeur minimale est retirée du signal. Puis, l'intensité est normalisée pour que la valeur maximale soit 1.

6.2 Analyse des données grâce aux réseaux de neurones

La méthode pour analyser les diffractogrammes sera la même que celle décrite dans le chapitre 5 sur le jeu de données de laboratoire par la figure 5.3. Dans un premier temps, un réseau de neurones est entraîné sur des diffractogrammes simulés. Puis dans un second temps, il est confronté aux données réelles.

6.2.1 Construction de la base de données

Pour entraîner le réseau, il faut d'abord constituer la base d'entraînement.

Simulation des diffractogrammes à une phase Pour la base d'entraînement, nous commençons par simuler 1 250 diffractogrammes pour les dix phases minérales identifiées. Comme les dix phases minérales ne sont pas centro-symétriques, que le C-S-H est turbostratique, et qu'une phase amorphe est présente, il est nécessaire d'utiliser le code de simulation complet (algorithme 2), même si cela implique des temps de calculs plus longs. Avec cet algorithme, les paramètres qui varient sont légèrement différents. Il y a, cependant, toujours les paramètres de maille avec 2% de variation autour des valeurs fournies par les fichiers CIF, le tirage de la valeur se faisant de manière uniforme sur l'intervalle. De même pour le paramètre d'agitation thermique dont la valeur B est tirée de manière uniforme entre 0.1 et 2.

Par contre, ce n'est plus directement la largeur à mi hauteur des pics que l'on fait varier. Ce sont les nombres de répétitions de maille N_1 , N_2 et N_3 dans les directions de l'espace (ce qui permet la simulation de phases turbostratiques et/ou anisotropes). Les variations des valeurs N_1 , N_2 et N_3 sont spécifiques à chaque phase minérale, car elles appartiennent à des groupes d'espace différents. L'ensemble des détails concernant ces variations sont donnés en annexe par le tableau A.2. La phase Amorphe est traitée de manière différente. C'est une phase dont le signal de diffraction ne varie pas. Une seule simulation est donc réalisée à partir du code du C-S-H, en limitant le nombre de répétitions de mailles dans les directions de l'espace. La phase Amorphe est en effet relativement proche du C-S-H comme le montre la figure 6.5.

Ces 1 250 diffractogrammes par phase minérale sont cette fois séparés en deux sous-ensembles indépendants pour générer des bases d'entraînement et de validation le moins corrélées possibles.

Constitution des diffractogrammes contenant le signal de plusieurs phases Le processus de construction des diffractogrammes multi-phases est identique à celui décrit dans les chapitres précédents. Celui-ci repose toujours sur la figure 4.3, avec une adaptation afin d'avoir une meilleure répartition des données comme le montre les histogrammes des

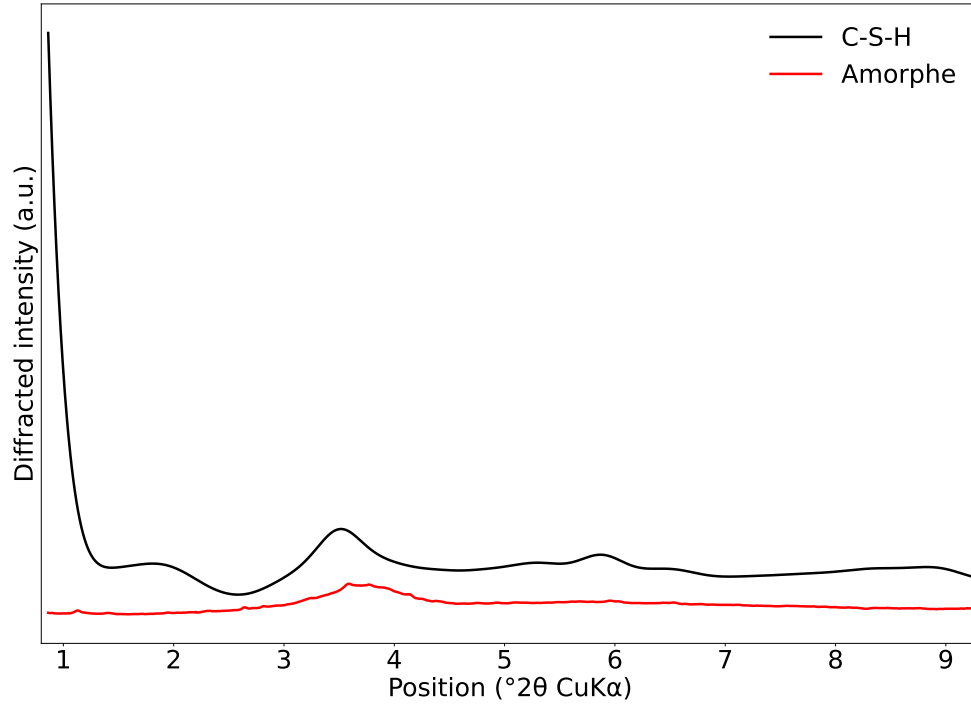


FIGURE 6.5 – Comparaison d’un diffractogramme Amorphe obtenu de manière expérimentale avec un C-S-H simulé avec un nombre de répétitions de mailles tel que $N_1 = N_2 = 2$ et $N_3 = 1$.

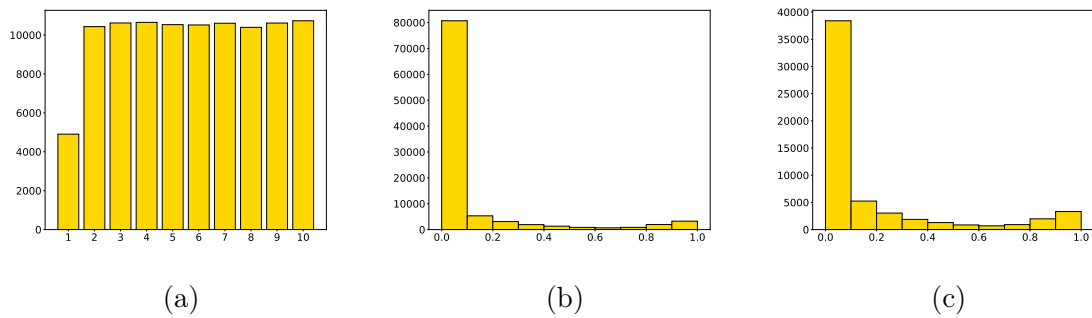


FIGURE 6.6 – (a) Barplot de la répartition des données dans la base d’entraînement, en fonction du nombre de phases présentes dans le mélange. (b) Dans la base d’entraînement, histogramme de la répartition des proportions d’une phase. (c) décrit le même histogramme sans les données qui ne contiennent pas la phase.

figures 6.6(a) et (b). Cette adaptation dans le tirage du vecteur de proportion permet de tirer plus facilement des valeurs importantes dans les mélanges contenant un nombre de phases important.

Aussi, comme le nombre de phases minérales est plus important, la base d’entraînement contient dix fois plus de données, soit un total de 100 000 diffractogrammes, dont la répartition selon le nombre de phases dans le mélange est donnée par le diagramme en barres 6.6(c).

6.2.2 Entraînement du réseau de neurones

L’architecture du réseau est similaire à celle décrite dans le chapitre 3 par la figure 3.4 en s’adaptant à la taille du signal et au nombre de classes K . Le modèle est entraîné sur 500 époques, avec les données regroupées par *batch* de 100. L’*optimizer* est toujours Adam, avec un learning rate constant à 0.001. Comme dans le chapitre précédent, la fonction de perte est la “MSE & Dirichlet”.

La base de validation contenant 10 000 diffractogrammes permet de suivre la phase d’entraînement, et d’évaluer la capacité de l’approche sur ce problème plus complexe. La figure 6.7 permet de suivre l’évolution de la fonction de perte (a) qui décroît nettement sur les 500 époques. Tout comme la RMSE (b) et la MMAE (c) qui diminuent au fur et à mesure de l’entraînement et qui atteignent respectivement des valeurs proches de 0.6% et 1%. Ces valeurs sont proches de celles obtenues lors des expériences précédentes.

Cependant pour la RRS, les valeurs sont moins bonnes. Comme précédemment sur les données synthétiques, le seuil ε est fixé à 1%. Nous pouvons voir que les valeurs de la RRS augmentent lors de l’entraînement, mais la valeur atteinte au final est inférieure à 53%, contre plus de 90% précédemment. Cette différence s’explique par l’augmentation du nombre de classes passant de $K = 4$ à $K = 10$.

Les résultats sur l’ensemble de validation montrent que l’entraînement a fonctionné correctement. Le réseau peut donc être confronté aux données réelles.

6.3 Résultats

Les résultats que nous présentons ne concernent qu’une seule tranche d’échantillon parmi les trois issues du jeu de données. Comme les données ne présentent aucune vérité terrain, les résultats seront présentés de manière différente.

Il est possible de comparer nos résultats à la méthode développée par CLARET *et al.* (2018). Elle s’appuie sur les pics de différentes phases minérales. Il s’agit de trouver une plage angulaire spécifique à une phase, où il n’y a qu’une seule phase minérale qui puisse produire de l’intensité diffractée dans les canaux sélectionnés. Les abondances de chaque phase sont ensuite déterminées à partir des valeurs de l’intensité dans ces plages angulaires, en normalisant à 1 l’abondance la plus forte pour chaque phase. Ainsi, plus la valeur est importante plus la phase est présente. Pour le maximum, l’estimation sera de 100%. Cette méthode est cependant très sensible au bruit. De plus, comme l’échantillon (pour une image) ne subit qu’une rotation uni-axiale (selon l’axe x), des effets d’orientation préférentielle peuvent exister. Cela peut amener à masquer certains pics de diffraction, ce qui constitue une limite supplémentaire. La figure 6.8 compare les cartes d’abondances de notre approche avec les résultats du réseau de neurones, et celle de CLARET *et al.* (2018).

Pour rappel, chacune donne des résultats différents, l’une obtient un rapport d’abondance en fonction du pixel où la phase est la plus présente (méthode du maximum d’in-

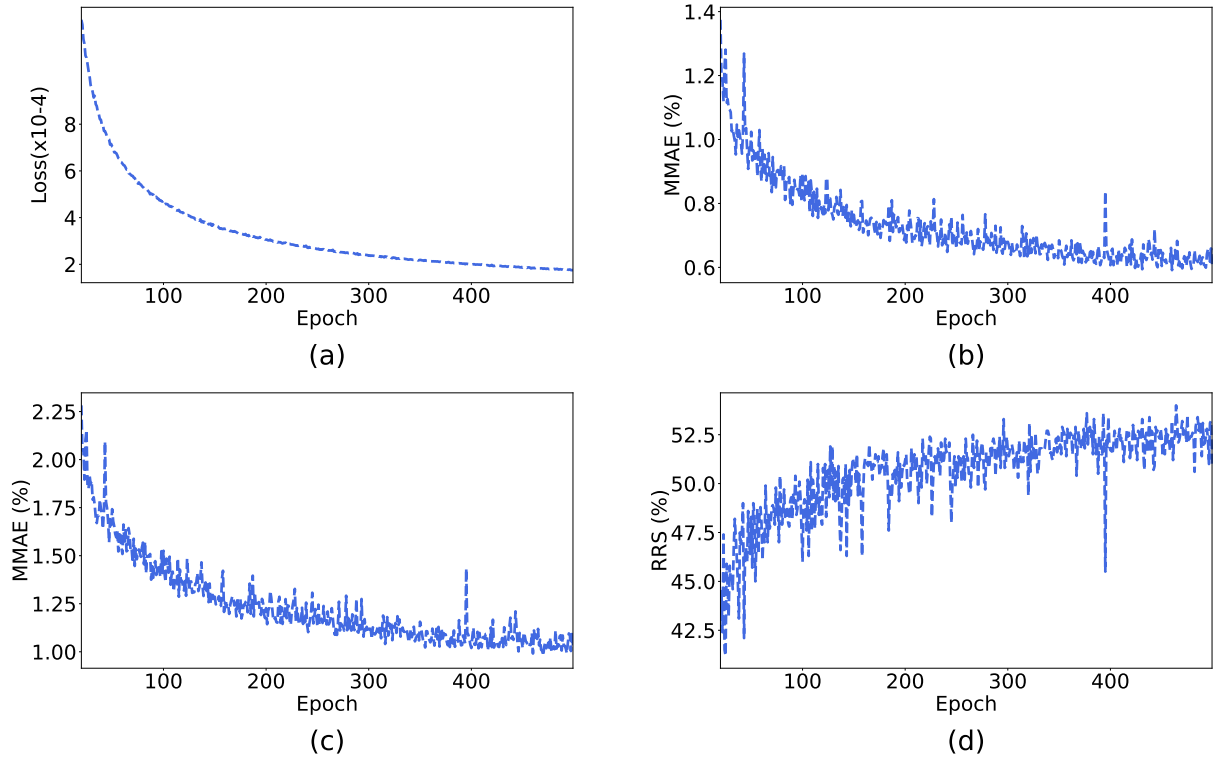


FIGURE 6.7 – Analyse de diffractogrammes simulés : (a) Evolution des valeurs de la *loss* en fonction de la base d’entraînement avec (b), (c) et (d) qui décrivent les valeurs des mesures (RMSE, MMAE et RRS respectivement) sur l’ensemble de validation au cours des époques d’entraînement.

tensité), l’autre permet d’obtenir la proportion (méthode neuronale). De nombreuses similitudes sont observées. Comme par exemple, pour la Calcite, où les zones majoritaires sont similaires, ou encore le Quartz qui est identifié principalement sur le même pixel. C’est également le cas pour la Mullite et la Portlandite, dont une majorité des zones d’abondances importantes sont identiques. Le C-S-H est également trouvé en abondance partout dans l’image pour les deux méthodes.

Il faut cependant noter que la phase Amorphe semble mal identifiée par le réseau de neurones, notamment car les zones où sa présence est détectée sont proches de celles de la Calcite. Une explication pourrait venir de la proximité entre le C-S-H et l’Amorphe qui pourraient être confondus par le réseau. Cette affirmation sera confirmée par la suite.

Une autre comparaison des méthodes est établie avec les régressions linéaires de la figure 6.9. Pour chacune des phases, un nuage de point est dessiné, il correspond aux abondances du maximum d’intensité en fonction des proportions trouvées par le réseau de neurones. Idéalement, la droite de régression doit correspondre à la fonction identité. Les phases pour lesquelles les deux méthodes ont des résultats relativement proches sont : la Calcite, le C-S-H, l’Ettringite et la Mullite. L’écart est important pour la phase Amorphe, le Quartz, l’Alite et le Monosulfoaluminate, où les droites de régression linéaire sont presque horizontales. Pour ces phases minérales, présentes à l’état de trace, il est normal et attendu de faire de telles observations.

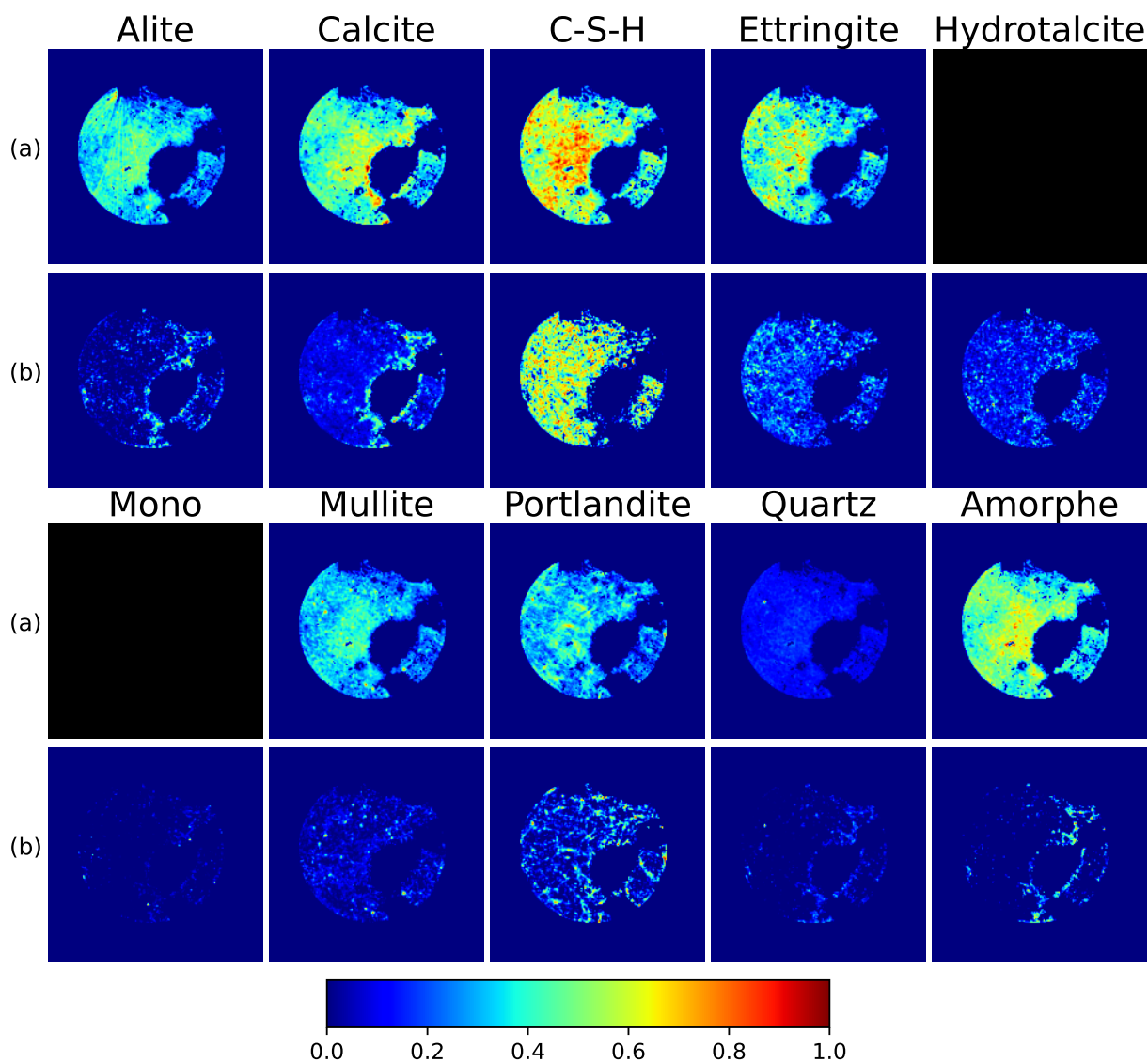


FIGURE 6.8 – Cartes d’abondances pour les dix phases minérales : (a) décrit les résultats de la méthode alternative présentée dans les travaux de CLARET *et al.* (2018) et (b) les résultats de l’approche avec les réseaux de neurones. Les cartes d’abondances pour l’Hydrotalcite et le Monosulfoaluminate avec la méthode de CLARET *et al.* (2018) ne sont pas présentées car ces phases sont simplement à l’état de trace.

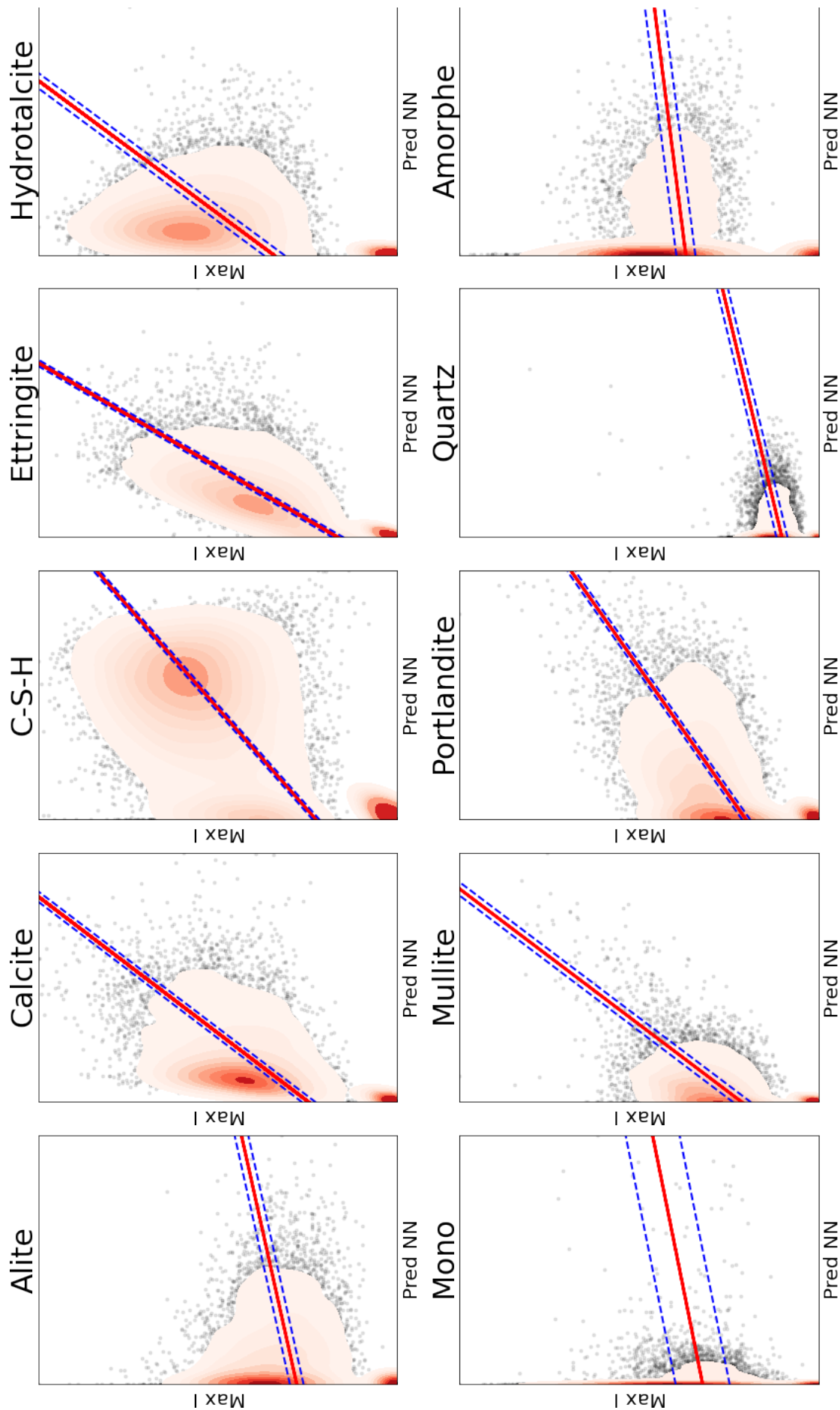


FIGURE 6.9 – Nuages de points des proportions prédites par le réseau de neurones en fonction de la prédiction de la méthode du maximum d'intensité (CLARET *et al.* 2018). Une droite de régression du nuage de points est proposée. Enfin, la densité des nuages est représentée afin de faciliter la lecture du graphique.

Même si la comparaison des méthodes est intéressante et que les similitudes sont des bons indicateurs, il est compliqué de tirer des conclusions sur la capacité du réseau de neurones pour l'analyse du jeu de données en se basant uniquement sur la comparaison. D'autres figures sont utilisées pour décrire et interpréter les résultats.

Nous proposons la figure 6.10 qui montre pour chaque phase minérale, les dix diffractogrammes identifiés par le réseau où les proportions sont les plus importantes. Ces résultats sont intéressants, et plusieurs phases semblent très bien identifiées par le modèle, comme la Calcite, l'Ettringite, l'Hydrotalcite, la Mullite et la Portlandite. Mais pour d'autres phases, les résultats sont plus contrastés. Pour la Alite, par exemple, même si ses proportions sont de l'ordre 30%, on reconnaît les pics de Calcite assez nettement. De même pour la phase Amorphe, ce qui vient confirmer les remarques faites précédemment à partir des cartes d'abondances (figure 6.8).

Une autre remarque concernant la figure 6.10 est à faire sur la phase de Monosulfoaluminate. Sur les 10 signaux identifiés par le modèle pour cette phase, une grande partie ne correspond pas du tout à un diffractogramme de la phase. De plus, ces données sont très majoritairement bruitées. Ainsi, comme le Monosulfoaluminate est présent à l'état de trace, il semblerait que le modèle ait une tendance à ranger dans cette phase les signaux très bruités pour lesquels il ne distingue pas clairement une autre phase.

Pour compléter l'analyse des résultats, de nouveaux graphiques ont été ajoutés sur la figure 6.11. Ils comparent une simulation avec une moyenne de signaux issus de la base de données tomographie, pour chacune des phases minérales. Pour chaque phase, l'ensemble des signaux à moyenner est obtenu en rassemblant les diffractogrammes dont la proportion prédite par le réseau est supérieure à un seuil ε . Ce seuil est variable selon la phase. Pour en avoir un nombre suffisant, il varie entre 25% et 50%. Ce graphique est un bon indicateur des performances du réseau, et permet également de voir quels sont les pics que le réseau utilise pour identifier et quantifier les phases minérales. Une majorité des phases est bien identifiée, puisqu'on reconnaît les pics caractéristiques. C'est le cas notamment de la Calcite, du C-S-H, de l'Ettringite, de la Mullite et du Quartz.

Ce graphique permet également de voir certaines proximités entre phases. Par exemple, sur le signal moyen de C-S-H, nous pouvons voir les pics de l'Ettringite. L'inverse est également vrai, ce qui montre une proximité spatiale et un lien entre les deux phases.

Cependant, et cela confirme les résultats précédents, la phase Amorphe et le Monosulfoaluminate ne sont pas très bien reconnus par le modèle. Le Monosulfoaluminate semble même servir de « classe fourre-tout », où les signaux bruités peu reconnaissables et les signaux de porosité encore présents sont rangés.

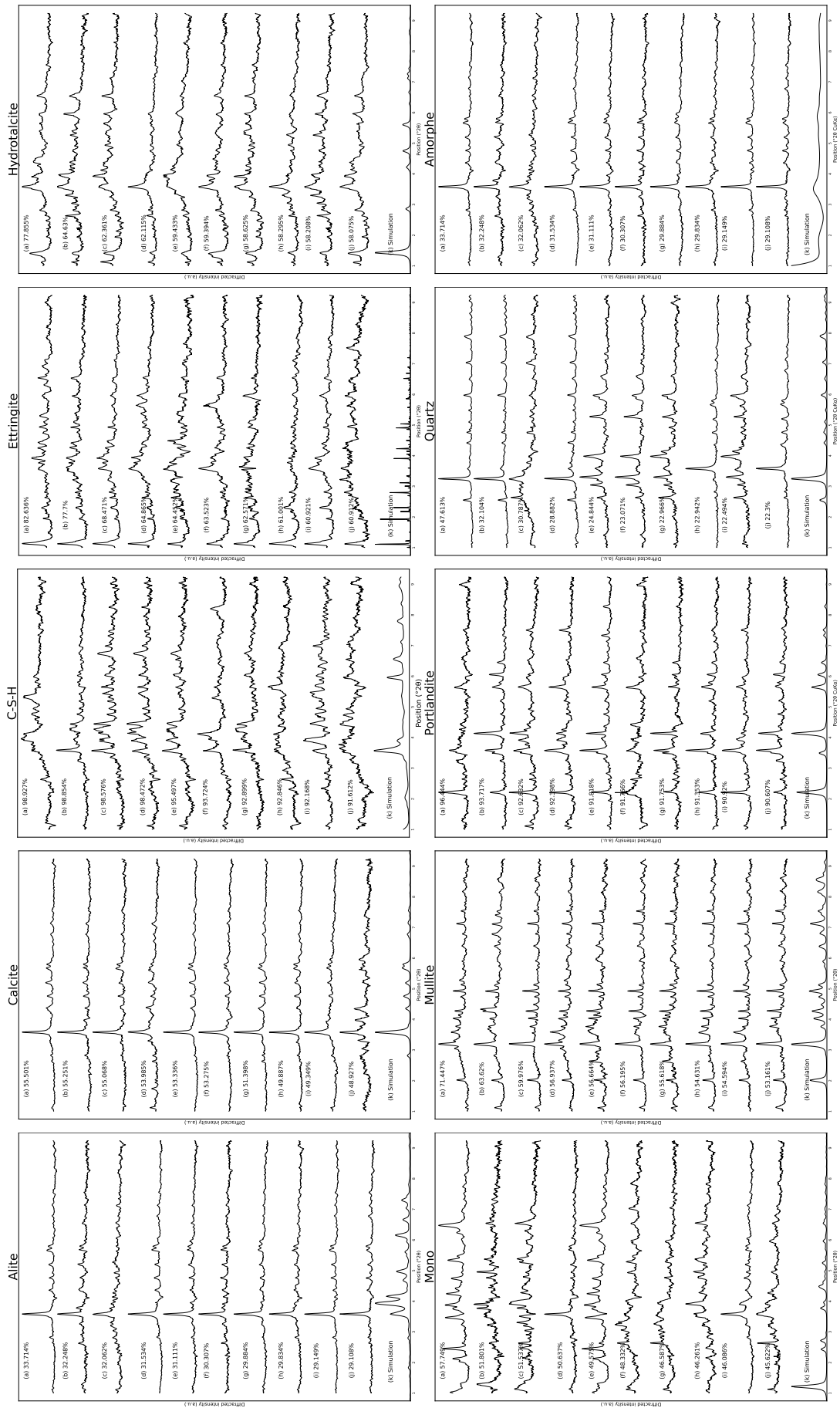


FIGURE 6.10 – Pour chaque phase minérale, les 10 diffractogrammes avec les proportions les plus importantes identifiés par le modèle.

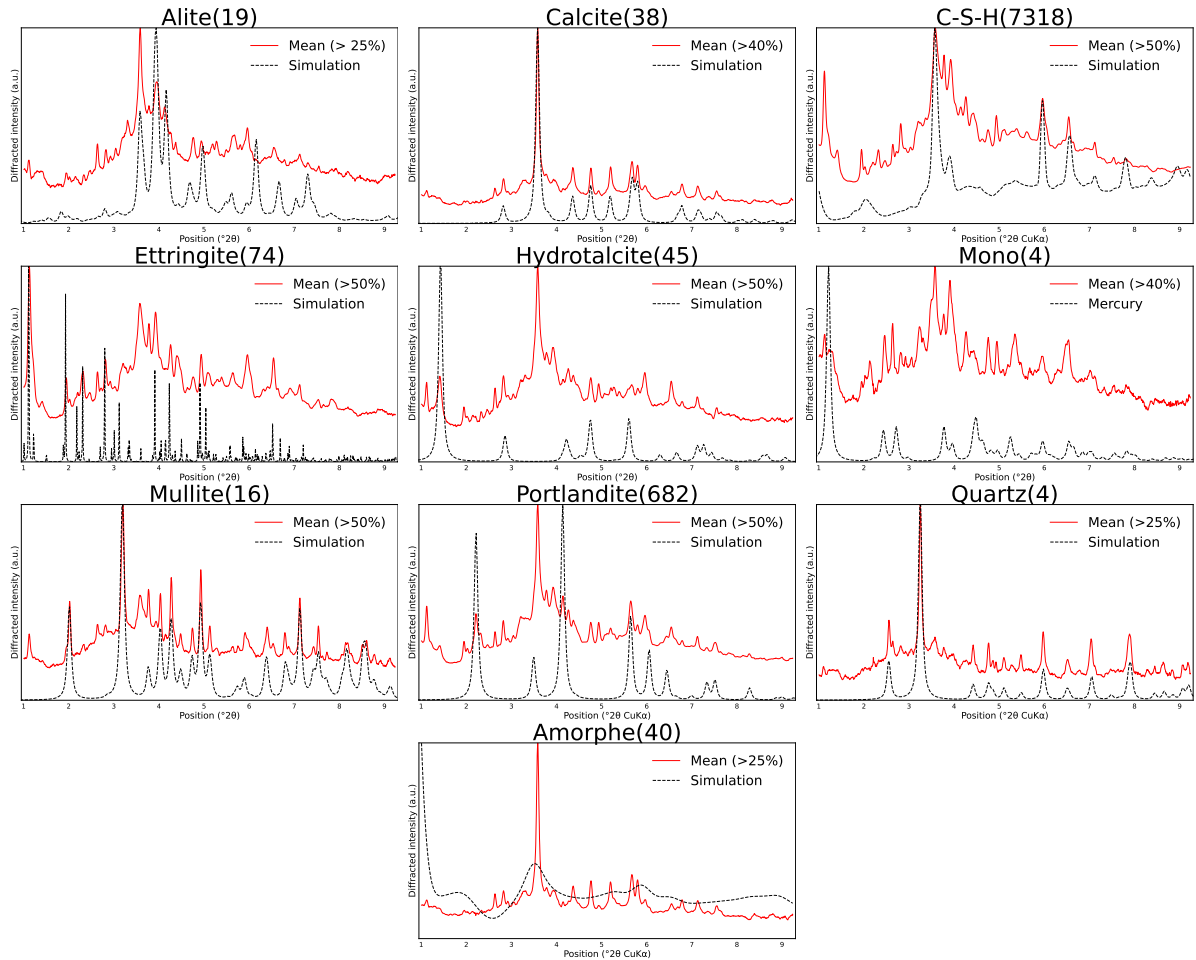


FIGURE 6.11 – Pour chaque phase minérale, moyenne des signaux avec une proportion prédite par le modèle supérieure à un certain seuil ε . Ce seuil est indiqué dans la légende. Il est choisi pour que la moyenne contienne un nombre minimum de données, indiqué entre parenthèses à coté du nom de la phase minérale.

6.3.1 Un modèle à 11 classes

Face au phénomène évoqué au-dessus, nous avons décidé d’entraîner le CNN en ajoutant une onzième classe qui, dans la base d’apprentissage, est constituée des signaux qui avaient été retirés lors du pré-traitement des données. Des signaux similaires à celui de la figure 6.4 qui peuvent être assimilés à de la porosité. Cette phase permet de ranger tous les signaux les plus bruités et ceux associés à la porosité, afin d’éviter qu’une phase minérale serve de “classe poubelle”.

Cependant, les résultats ne sont pas ceux attendus. En effet, comme les cartes d’abondances de la figure 6.12 le montrent, la quasi-totalité des pixels est identifiée avec une proportion importante de la classe supplémentaire ajoutée. Ces résultats peuvent s’expliquer par le fait que les données de la base d’apprentissage de la classe supplémentaire soient extraites du jeu de données. Donc même si ces diffractogrammes correspondent à la porosité de l’échantillon, le bruit et la structure sont similaires à ceux observables dans les données à analyser. Le modèle considère donc cette classe comme majoritaire dans les pixels de l’image. Pour conforter cette hypothèse, nous observons que les diffractogrammes associés à la Calcite et au Quartz, qui sont très peu bruités, sont bien reconnus par le modèle comme le montre la figure 6.13.

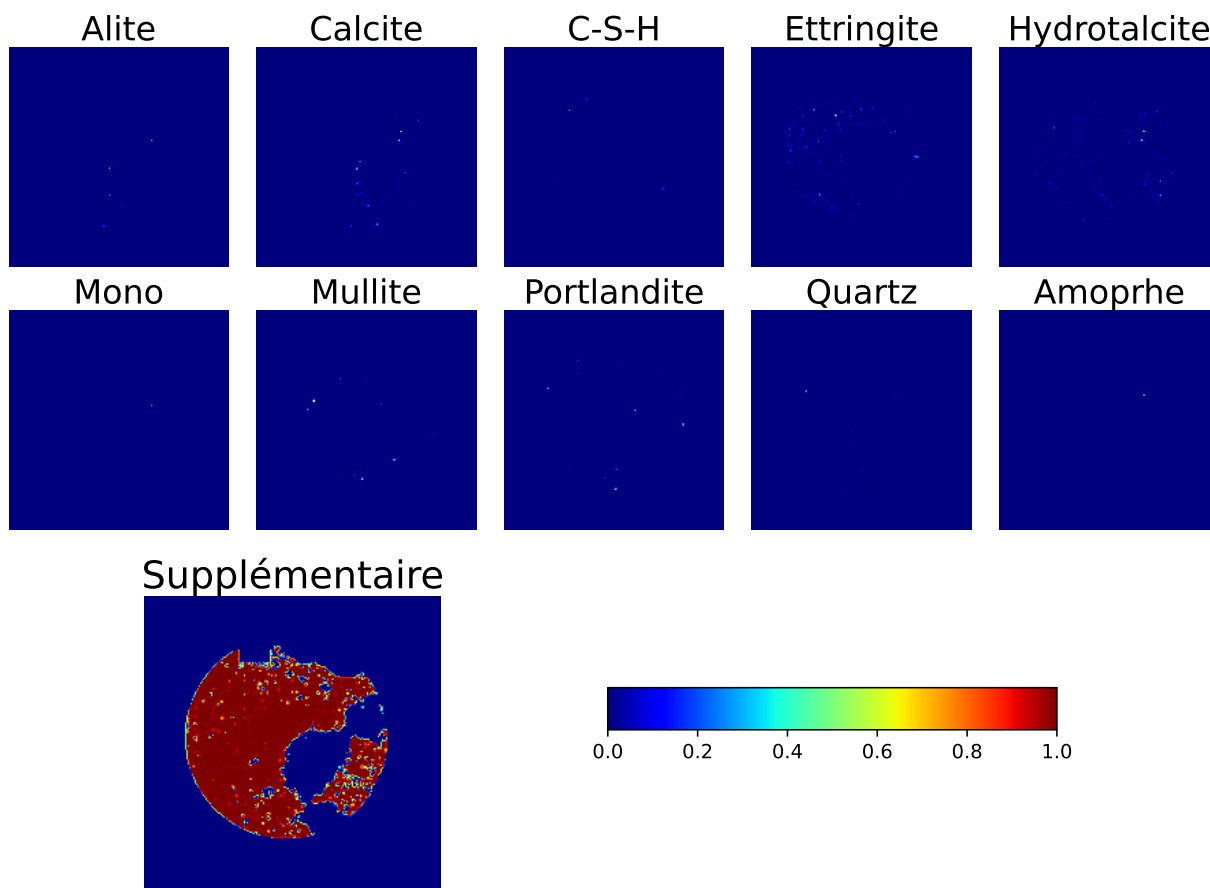


FIGURE 6.12 – Cartes d’abondances obtenus avec le système à 11 classes : les dix phases minérales et la classe supplémentaire.

6.3.2 Reconstitution

À partir du modèle entraîné avec dix phases minérales, i.e sans la phase supplémentaire, deux reconstitutions de signaux sont proposées sur la figure 6.14. Elles sont faites à partir de combinaisons linéaires de simulations, dont les coefficients sont déterminés par la prédiction du réseau. Il est important de préciser que le réseau ne prédit qu’un vecteur de proportions, qui correspond aux abondances des phases minérales. Mais il ne prédit ni les paramètres de maille, ni le coefficient d’agitation thermique, ni le nombre de répétitions de maille. Les simulations ont donc été choisies avec les paramètres de maille indiqués par le fichier CIF, l’agitation thermique avec $B = 0.75$ et le nombre de répétitions de maille dépendant des phases. Sur les données réelles, le fond continu est retiré, puis le signal normalisé.

Sur la figure 6.14a la reconstitution est bonne avec une très grande majorité des pics bien repris et des ratios d’intensité entre eux respectés. La différence réside notamment sur le fond continu qui est plus important sur le diffractogramme de la donnée réelle. Sur la seconde figure 6.14b, la reconstruction est moins bonne. Les pics sont pour la plupart bien identifiés, mais les ratios ne sont pas bons. Cependant, les largeurs à mi-hauteur des pics n’ont pas été ajustées, ce qui peut expliquer en partie ces différences.

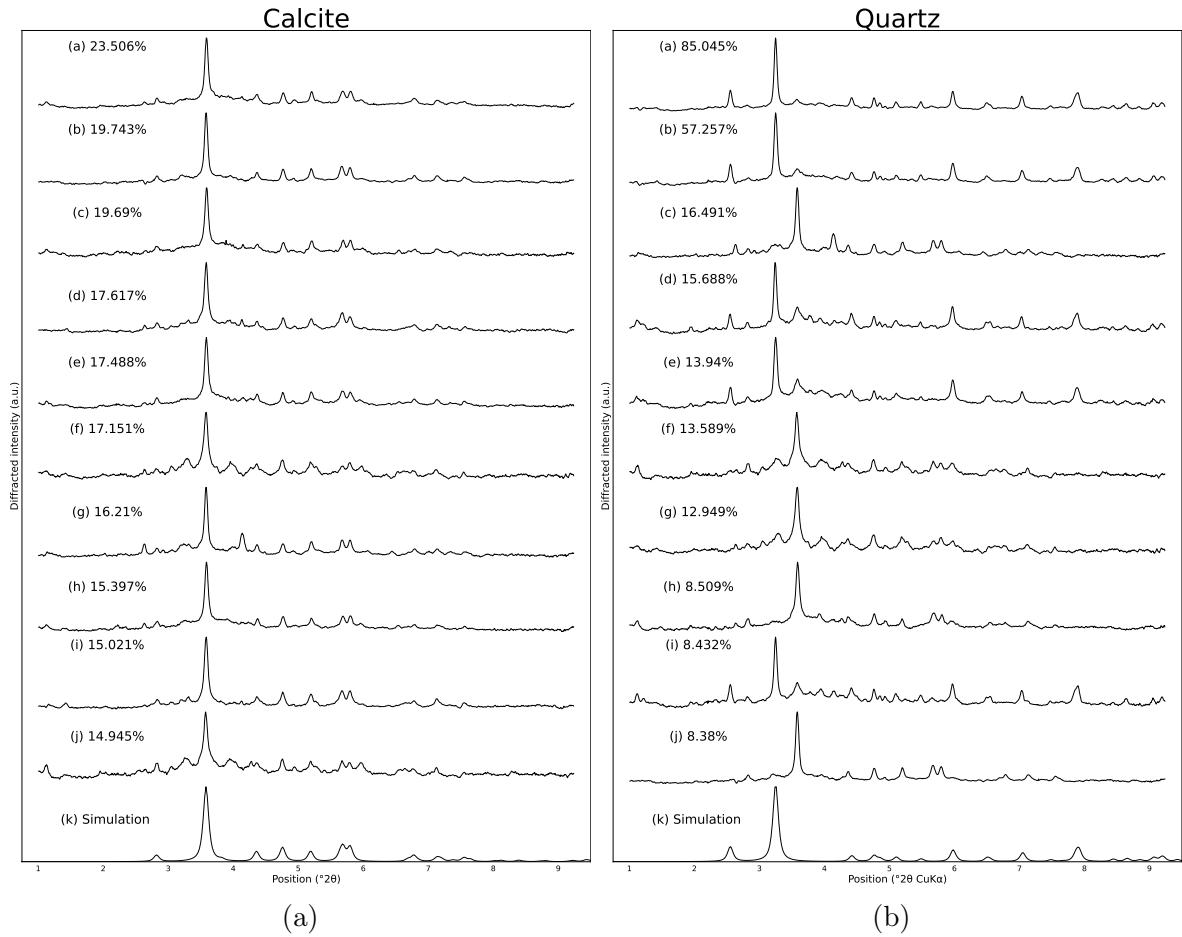


FIGURE 6.13 – Pour le modèle entraîné à 11 classes, les 10 diffractogrammes avec les proportions les plus importantes identifiées pour la Calcite (a) et le Quartz (b).

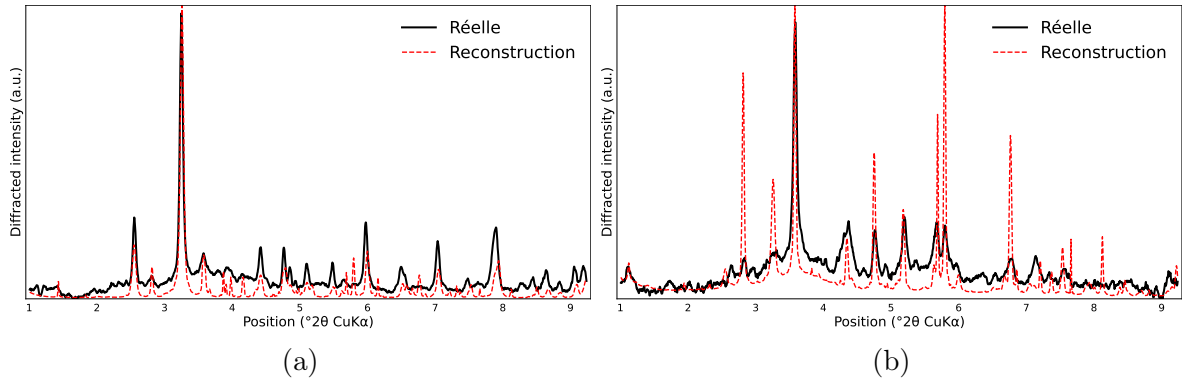


FIGURE 6.14 – Reconstitution des signaux à partir de simulations et des proportions prédites par le réseau. Les principales phases prédites par le réseau sont pour la figure (a) 48% Quartz, 14% Hydrotalcite, 7% Calcite, 7% C-S-H, 7% Monosulaluminate, pour les phases les plus présentes. Et (b) 55% Calcite, 18% Quartz, 16% Amorphe, 9% Ettringite.

6.4 Conclusion et discussion

L'analyse du jeu de données tomographie consiste à quantifier les $K = 10$ phases minérales identifiées au préalable sur plusieurs dizaines de milliers de diffractogrammes.

La méthode reprend l'approche du chapitre précédent, avec simplement un entraîne-

ment du réseau et non plus cinq. La base d'entraînement est uniquement constituée de données simulées, ce qui permet d'avoir un important nombre de données avec toute la variabilité souhaitée.

Comme précédemment, l'entraînement est suivi grâce à la base de validation sur laquelle les mesures RMSE et MMAE sont bonnes, avec des valeurs proches de 1%. La RRS présente de moins bons chiffres que sur les expériences précédentes, avec des valeurs autour de 50%. Cette différence s'explique par le nombre de phases plus important. Finalement ces chiffres montrent que l'entraînement s'est bien déroulé et que le modèle peut être confronté aux données réelles.

Le jeu de données réelles ne possédant pas de vérité terrain, il ne permet pas de quantifier l'erreur comme précédemment. Différents graphiques permettent donc de voir que le modèle s'en sort également très bien pour analyser les données. Par exemple, le Quartz et la Calcite sont bien reconnus. Cependant, certaines phases (Amorphe et Monosulfoaluminate) restent mal identifiées par le réseau.

Comme évoqué dans le chapitre précédent 5, des pistes d'amélioration peuvent être considérées. La constitution de la base de données pourrait être améliorée. Nous pouvons également optimiser certains paramètres du modèle, par exemple en utilisant un taux d'apprentissage programmé. Une architecture différente pourrait être utilisée, en construisant un réseau plus profond, ou avec un autre type d'architecture évoqué en introduction : les Transformers.

Chapitre 7

Une nouvelle approche : les Transformers

7.1 Introduction

Les Transformers correspondent à un nouveau type d'architecture de réseaux de neurones, développé ces dernières années. Le principe est similaire à celui des réseaux de neurones classiques décrit dans le chapitre 3. Le principal changement concerne le modèle utilisé, qui est basé sur le principe d'**attention** (VASWANI *et al.* 2017), décrit dans la suite du chapitre (section 7.2.2).

A l'origine les Transformers sont utilisés par Vaswani *et al.* (VASWANI *et al.* 2017) pour le traitement automatique du langage (CHOWDHARY & CHOWDHARY 2020) (TAL ou *Natural Language Processing (NLP)* en anglais). Dans ce domaine, les méthodes les plus avancées utilisent les Transformers, comme BERT (*Bidirectional Encoder Representation from Transformers*) (DEVLIN 2018) qui propose d'utiliser des modèles pré-entraînés, qui nécessitent un simple affinage, pour répondre à diverses tâches de TAL. Egalement le modèle GPT (*Generative Pretrained Transformer*) (BROWN *et al.* 2020) qui est directement capable d'effectuer de nombreuses tâches. Mais les performances des Transformers ne se limitent pas au TAL, et sont aussi état de l'art dans d'autres domaines.

C'est le cas pour le traitement d'images, avec notamment les Visions Transformers (ViT) introduits par DOSOVITSKIY *et al.* (2021) qui atteignent d'excellentes performances pour la classification d'images. Il y a de nombreux autres exemples de l'utilisation des Transformers pour le traitement d'image, par exemple la détection d'objet (CARION *et al.* 2020), ou la segmentation (LIU *et al.* 2021 ; W. WANG *et al.* 2021). Les ViT offrent en plus des outils de visualisation comme le *positional encoding* (DOSOVITSKIY *et al.* 2021), ou des méthodes pour suivre la propagation de l'information dans le modèle (ABNAR & ZUIDEMA 2020 ; LEEM & SEO 2024).

Le traitement du signal est aussi possible avec les Transformers, pour le diagnostic de maladies cardiaques (CHE *et al.* 2021), dans le domaine hyperspectral (GHOSH *et al.* 2022 ; SCHEIBENREIF *et al.* 2023) ou la séparation de source (ROUARD *et al.* 2023). L'analyse des signaux de XRD avec les Transformers est pour le moment peu développée, seuls quelques travaux récents ont été publiés en 2024. CHEN *et al.* (2024) proposent un ViT pour l'identification de réseaux métallo-organiques. CAO *et al.* (2024) utilisent les transformers pour prédire la résistance compressive du ciment de Portland.

L'objectif de ce chapitre est de tester les Transformers pour l'identification et la quantification de phases à partir des diffractogrammes. Cette étude est appuyée par les outils

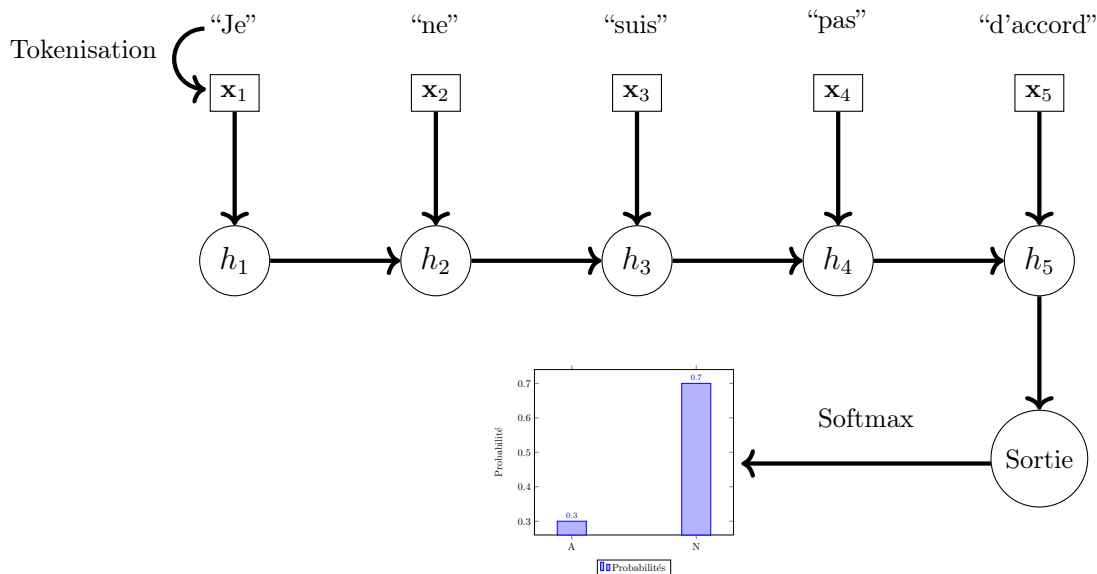


FIGURE 7.1 – Schéma simplifié d'un réseau de neurones récurrent pour une tâche de classification. Le processus de *tokenization* est décrit ci-dessous (section 7.2.1)

de visualisation que les Transformers peuvent offrir.

7.2 L'architecture des Transformers

7.2.1 Les réseaux de neurones récurrents

Afin d'introduire les Transformers, il est intéressant de commencer par une description des **réseaux de neurones récurrents** (MEDSKER, JAIN *et al.* 2001) (RNN), qui présentent des similitudes.

L'objectif de cette section est uniquement d'introduire les modèles des RNN. Afin d'illustrer le fonctionnement global de ce type de réseaux, une représentation schématique est donnée par la figure 7.1. Elle présente une tâche simple de TAL, la classification de phrases, en phrases affirmatives ou négatives.

La première étape avec ces modèles consiste à segmenter la donnée en plusieurs morceaux x_i , c'est la *tokenization*. Dans l'exemple de classification de la figure 7.1, il s'agit de découper la phrase avec les mots qui la composent. Puis d'encoder chaque mot à travers un processus de *word embedding* (VASWANI *et al.* 2017) qui vise à convertir un mot en un vecteur. Pour un vecteur, la *tokenization* consiste donc à découper le signal en plusieurs segments x_i .

Ensuite, au niveau de chaque segment x_i , un état caché h_i est associé. Le premier état caché h_1 dépend uniquement de x_1 , le deuxième h_2 dépend de x_2 et de h_1 , le troisième de x_3 et de h_2 , et ainsi de suite. Le dernier état caché est utilisé pour réaliser la tâche de classification, par exemple en utilisant la fonction Softmax (équation (3.1)). L'apprentissage se fait, une nouvelle fois, en optimisant une fonction de perte lors d'une phase d'entraînement, toujours à travers une base de données labellisées.

Cependant, bien que les RNN aient été très développés et performants pour certaines tâches comme la traduction automatique de texte (SUTSKEVER *et al.* 2014) ou la reconnaissance vocale (GRAVES *et al.* 2013), ils souffrent de certaines limitations. La principale étant que l'information ne se propage que dans un sens. Si bien que, dans notre exemple,

le mot de début de phrase aura une influence bien moindre que le dernier mot pour la classification finale. L'opération d'**attention**, centrale pour les Transformers, permet de remédier à cela.

7.2.2 Attention

L'opération d'**attention** permet de calculer d'une autre manière que les RNN les liens entre les différents segments $\mathbf{x}_i$ de la donnée d'entrée. A la place des états cachés h_i , ce sont des scores d'**attention** qui sont calculés. Pour ces calculs, nous commençons par associer différents vecteurs à chaque segment de l'entrée du réseau. Ce sont les vecteurs $\mathbf{q}$, $\mathbf{k}$, et $\mathbf{v}$. Le premier $\mathbf{q}$ (pour *queries* en anglais) correspond aux requêtes, $\mathbf{k}$ (pour *keys*) correspond aux clés et $\mathbf{v}$ (pour *values*) aux valeurs. Ces vecteurs sont simplement obtenus après une opération linéaire (dont les poids sont appris par le modèle) à partir des vecteurs $\mathbf{x}_i$.

Le score d'**attention** est calculé à partir des vecteurs $\mathbf{q}$, $\mathbf{k}$, et $\mathbf{v}$, avec la formule ci-dessous :

$$\text{Attention}(\mathbf{q}, \mathbf{k}, \mathbf{v}) = \text{Softmax} \left(\frac{\mathbf{q}\mathbf{k}^T}{\sqrt{d_{\mathbf{k}}}} \right) \mathbf{v}, \quad \text{où } d_{\mathbf{k}} \text{ est la dimension de } \mathbf{k}. \quad (7.1)$$

Ce score compare d'abord les requêtes avec les clés en calculant $\mathbf{q}\mathbf{k}^T$. Ce produit est normalisé par la dimension de $\mathbf{k}$. Puis une fonction Softmax est appliquée pour revenir à un vecteur de proportion. Après ces deux opérations, nous obtenons les **poids d'attention** qui servent à pondérer les valeurs $\mathbf{v}$. Le score d'**attention** résulte finalement du produit des poids d'**attention** et de $\mathbf{v}$. La figure 7.2(b) montre le processus de calcul des scores d'**attention**, avec des schémas simplifiés.

Ces calculs permettent de voir les corrélations entre chaque segment de la donnée d'entrée (voir figure 7.2(a)). Ainsi, l'information n'est pas propagée dans un sens unique, ce qui permet une meilleure diffusion de celle-ci dans le modèle. Les Transformers reposent donc sur ce principe d'**attention** pour offrir une approche plus performante.

En général ce principe est appliqué plusieurs fois en parallèle, chaque application étant désignée par une tête (*head*) d'**attention**. Nous parlerons de *Multi-Head Attention*. Cela permet de capturer les informations de la donnée de différentes manières.

Tous ces éléments permettent d'introduire les principaux éléments du Transformers. L'application spécifique à notre problématique (DRX) peut maintenant être décrite.

7.2.3 Encodage du signal et vecteurs de positions

L'approche que nous proposons utilise un ViT (DOSOVITSKIY *et al.* 2021). Les ViT se distinguent des Transformers pour le NLP présenté précédemment, par les données qu'ils traitent, qui ne sont plus des textes, mais des images, ou dans notre cas des signaux. Le ViT permet de capturer les relations entre les différentes parties du signal grâce à l'opération d'**attention**. C'est donc pour la première étape de *tokenization* que les différences entre les modèles sont importantes, où il s'agit de convertir la donnée en entrée (i.e *tokens*) pour le Transformer.

Pour les signaux, il existe plusieurs manières de réaliser la *tokenization*. L'une des méthodes répandues consiste à segmenter le signal en N segments de longueur C . C'est l'approche choisie par CHEN *et al.* (2024) pour les diffractogrammes, mais aussi par SCHEIBENREIF *et al.* (2023) pour le traitement d'images hyperspectrales. Pour cette méthode, le choix de la longueur C des tokens est un hyperparamètre du modèle à déterminer.

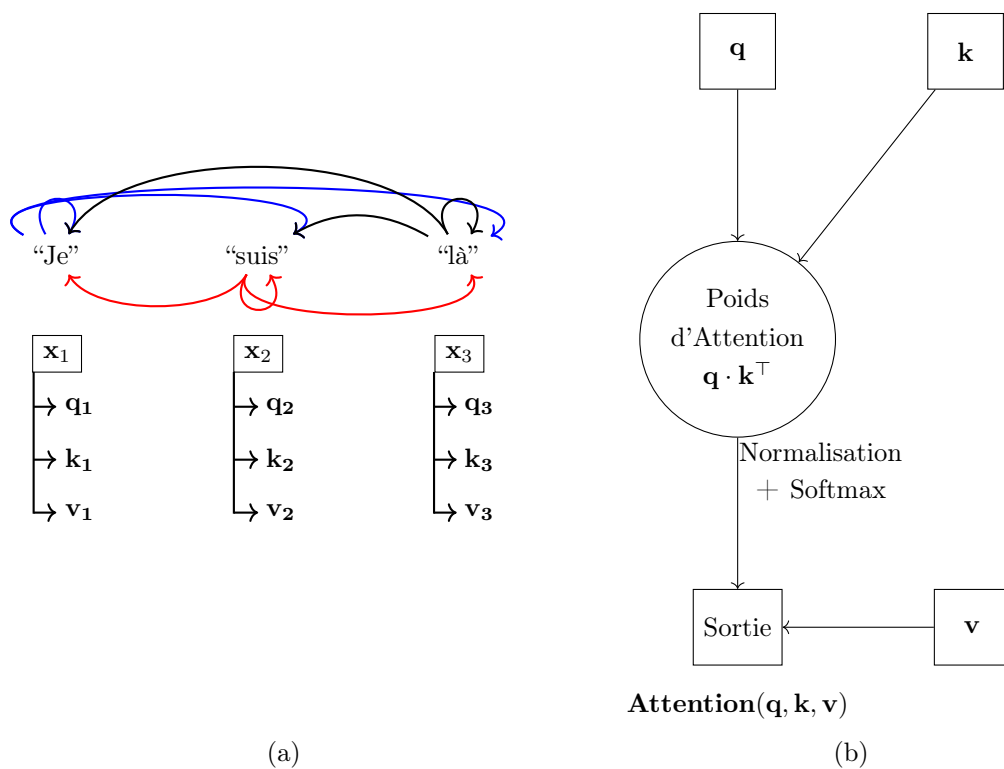


FIGURE 7.2 – (a) Lien entre chaque segment de l’entrée grâce aux vecteurs requêtes ($\mathbf{q}$), clés ($\mathbf{k}$) et valeurs ($\mathbf{v}$) grâce au principe d’**attention**. Les vecteurs $\mathbf{q}, \mathbf{k}, \mathbf{v}$ sont obtenus avec partir d’une opération linéaire à partir des vecteurs $\mathbf{x}_i$. (b) Schéma simplifié du calcul des scores d’**attention**.

Une alternative consiste à utiliser un réseau de neurones pour la *tokenization*. C'est notamment ce qui est proposé pour le traitement d'électrocardiogrammes, où les tokens sont obtenus à partir du signal grâce à un CNN (CHE *et al.* 2021). C'est aussi l'approche choisie par GHOSH *et al.* (2022) pour traiter des images hyperspectrales.

Pour ces deux méthodes de *tokenization*, un *token* supplémentaire est ajouté au début de la séquence. Il est nommé CLS *token* (pour *class token*) et permet de concentrer les informations du signal pour réaliser la tâche finale (identification ou quantification).

Pour compléter ces *tokens* et fournir plus d'informations, des vecteurs sont associés à chacun des tokens afin d'encoder leur position. Le terme anglais de *positional embedding* (VASWANI *et al.* 2017) désigne l'assignation d'un vecteur de position à chaque token. Cela permet d'apporter encore plus de robustesse au modèle. En effet, un signal avec une permutation de pics sera traité de la même manière que le signal original. L'encodage de la position de tokens permet donc de différencier des signaux qui auraient des pics identiques, mais à des positions différentes. Il existe plusieurs manières pour encoder les vecteurs de position. A l'origine, pour les problèmes de génération ou de traitement de texte, ces vecteurs sont encodés grâce aux fonctions sinus et cosinus (VASWANI *et al.* 2017). Mais avec les ViT et le traitement de signal, une majorité des études laisse ces vecteurs comme des paramètres du réseau à apprendre (CHEN *et al.* 2024 ; GHOSH *et al.* 2022 ; GUPTA *et al.* 2023 ; SCHEIBENREIF *et al.* 2023).

7.2.4 Les Transformers pour l'analyse des signaux de diffraction

Le schéma 7.3 résume notre approche avec un ViT pour traiter les signaux de diffraction. Notre approche utilise des éléments introduits par CHEN *et al.* (2024). Ces travaux se concentrent sur l'identification (i.e classification) de réseaux métallo-organiques, donc avec une variabilité intra-classe assez faible.

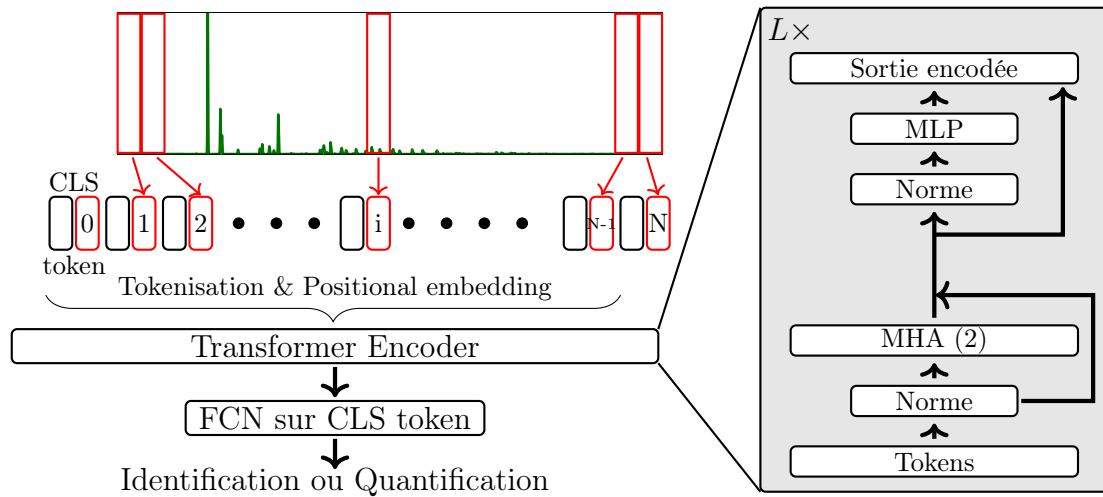


FIGURE 7.3 – Schéma représentant le modèle d'un Transformers.

La première étape pour les diffractogrammes est donc de réaliser la *tokenization* du signal. Comme les résultats des expériences décrites par les tableaux 7.4(b) et 7.1 le montrent, plusieurs manières de réaliser la *tokenization* ont été testées. Les meilleurs résultats sont obtenus par segmentation du signal, en fixant le paramètre $C = 80$. Pour les vecteurs de positions, ils sont initialisés avec des valeurs à 0, puis appris au fur et à mesure de l'entraînement du modèle.

A la suite de la *tokenization*, les *tokens* et les vecteurs de positions sont utilisés comme les entrées pour passer à travers un encodeur. Cette partie du modèle est décrite par la partie droite de la figure 7.3. Ce bloc se compose d’abord d’une normalisation, puis de la partie **attention**, nommée MHA pour *Multi-Head Attention*. Ici, 2 têtes d’**attention** sont utilisées. A la sortie du MHA, une nouvelle normalisation est appliquée avant d’utiliser un MLP qui termine l’encodage des tokens. Ce bloc encodeur est répété $L = 12$ fois, pour apporter une meilleure compréhension du signal. De plus, des connexions résiduelles sont ajoutées pour apporter de la stabilité au modèle.

Finalelement, la tâche finale est effectuée grâce au CLS token de la sortie encodée avec un MLP. Nous considérons ici à la fois le problème classification introduit dans la section 3.2, et le problème de quantification pour des données simulées et expérimentales présenté dans le chapitre 5.

Dans le cadre de nos travaux, nous avons choisi d’utiliser les Transformers pour plusieurs raisons. En se basant sur la physique des diffractogrammes, l’utilisation d’un ViT est pertinente pour capturer les informations entre les différents pics du signal, telles que la co-occurrence de pics caractéristiques. Ensuite, grâce aux performances des Transformers qui font partie de l’état de l’art pour les tâches de TAL, et pour la classification d’images, qui amènent à penser que l’analyse des diffractogrammes avec cette approche est pertinente. Et ce, même si leur utilisation n’est pour le moment pas très répandue en traitement du signal, malgré des premières études montrant des résultats prometteurs. Enfin, grâce à différents outils, il est possible de visualiser la propagation de l’information dans le modèle (ABNAR & ZUIDEMA 2020 ; DOSOVITSKIY *et al.* 2021).

7.2.5 Des outils de visualisation

Les outils de visualisation introduits ci-dessous ont pour but d’offrir une meilleure compréhension des diffractogrammes. Ils permettent également de voir comment le modèle interprète et analyse la donnée. Nous utilisons, pour cela deux visualisations.

Positional encoding : Il repose sur les vecteurs de positions associés aux tokens. Chaque token correspond à une plage angulaire sur l’ensemble des angles de diffraction considérés par le diffractogramme. En étudiant les vecteurs de position, des similarités ou des dépendances peuvent être observées entre les différentes plages angulaires. C’est la similarité cosinus que nous utilisons comme mesure de proximité entre chacun des vecteurs. Entre deux vecteurs **a** et **b**, elle est définie par l’équation suivante :

$$\text{cossimilarity}(\mathbf{a}, \mathbf{b}) = \frac{\mathbf{a} \cdot \mathbf{b}}{\|\mathbf{a}\| \|\mathbf{b}\|}. \quad (7.2)$$

Plus concrètement, le *positional encoding* appliqué aux diffractogrammes permet de voir les dépendances de certains pics. Il s’agit d’observer, si la présence d’un pic sur une plage angulaire donnée influe sur la présence ou l’absence d’un pic sur une autre plage angulaire.

Attention Rollout : Le second outil que nous utilisons est l’*Attention Rollout* (ABNAR & ZUIDEMA 2020). Il permet de voir comment le modèle interprète la donnée grâce aux matrices d’**attention**. L’*Attention Rollout* utilise les matrices d’**attention** afin de suivre la propagation de l’information à travers les couches d’**attention**. Pour cela il faut d’abord récupérer les matrices d’**attention** des deux têtes, que l’on peut traiter séparément ou

fusionner comme le suggère ABNAR & ZUIDEMA (2020) La matrice fusionnée à la couche $\ell \in \{1, \dots, 12\}$ sera notée A_ℓ , et la fusion peut être réalisée de différentes manières : moyenne, minimum, maximum. Dans notre cas, nous considérons le minimum qui semble offrir de meilleurs résultats comme le montre les travaux de GILDENBLAT (2020). Les matrices fusionnées sont alors normalisées selon le principe suivant :

$$FA_\ell = (A_\ell + I) / \max(A_\ell + I), \quad (7.3)$$

où I est la matrice identité. Les matrices d'*Attention Rollout* sont alors définies de manière récursive avec l'équation suivante :

$$AR_1 = FA_1 \text{ et } AR_\ell = (FA_\ell + I) \times AR_{\ell-1}, \ell = 2, \dots, L. \quad (7.4)$$

L'initialisation se fait simplement avec la matrice d'**attention** de la première couche. De plus, les matrices identités sont ajoutées afin de considérer les connexions résiduelles entre chacune des couches.

Finalement, seule la dernière matrice à la couche $L = 12$ est considérée car elle synthétise la propagation de l'information. La première ligne contient les informations du CLS token qui permet d'effectuer la tâche du modèle. Elle permet de voir les tokens (i.e les plages angulaires du diffractogramme) qui ont le plus contribué à la construction du CLS token.

Les deux outils présentés dans cette section sont utilisés par la suite pour visualiser l'interprétation du signal par le ViT pour l'identification et la quantification de phases minérales.

7.3 Classification de diffractogrammes simple phase

Dans ce contexte, nous choisissons d'abord de tester les Transformers sur les diffractogrammes avec un problème d'identification de phases minérales.

7.3.1 Entraînement et résultats de la classification

Le problème est identique à celui décrit dans la section 3.2 du chapitre 3, avec $K = 6$ phases minérales : Calcite (MARKGRAF & REEDER 1985) (CaCO_3), Dolomite (STEINFINK & SANS 1959) (CaMgC_2O_6), Gibbsite (BALAN *et al.* 2006) (AlO_3H_3), Hematite (BLAKE *et al.* 1966) (Fe_2O_3), Halite (WALKER *et al.* 2004) (NaCl) et Quartz (LEVIEN *et al.* 1980) (SiO_2). Les bases d'entraînement et de test sont identiques.

Le modèle est celui introduit dans la section précédente 7.2. La fonction de perte est l'entropie croisée (voir équation (3.9)). Elle est optimisée avec l'algorithme Adam (KINGMA & BA 2014), avec un *learning rate* constant égal à 0.001. Le modèle est entraîné sur 100 époques, comme les réseaux de neurones du chapitre 3.

Plusieurs manières pour *tokeniser* le signal sont comparées : la segmentation du signal avec différentes valeurs de C (20, 40, 80 et 100), et la *tokenisation* avec un CNN. Les résultats et comparaisons sur l'ensemble de test sont reportés dans le tableau 7.4(b). Cela montre que la *tokenisation* du signal est un point critique pour le traitement du signal avec les modèles de Transformers. Les chiffres du CNN du chapitre 3 sont également reportés. D'après les résultats du tableau, la méthode la plus efficace est celle qui consiste à segmenter le signal en $N = 36$ segments de longueur $C = 80$, juste devant le CNN. En effet, le ViT-80 atteint des valeurs 99.8% sur les trois mesures, soit presque 3% d'efficacité

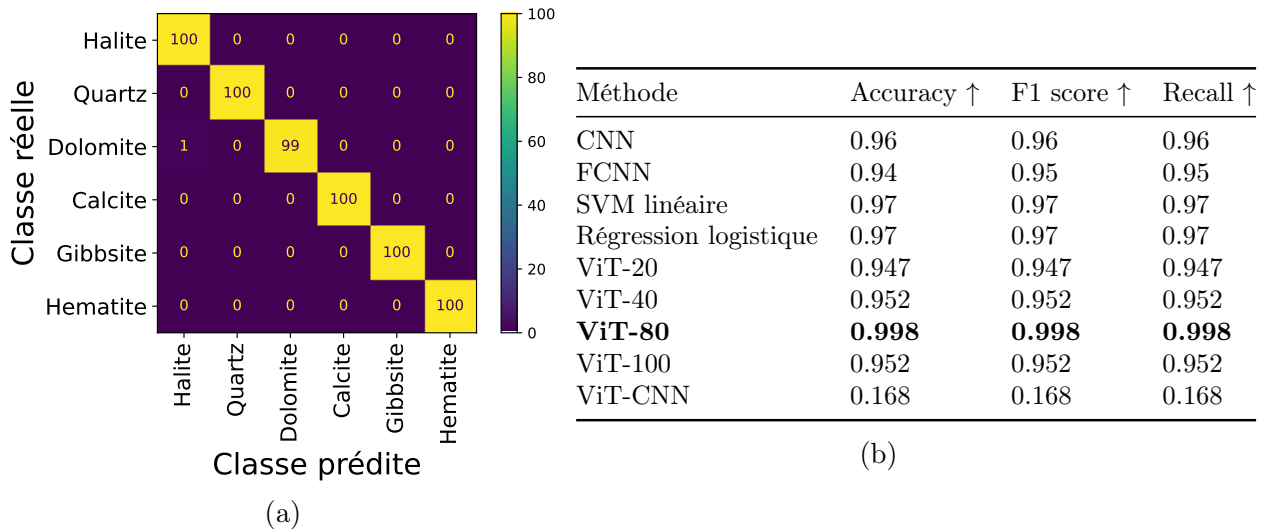


FIGURE 7.4 – (a) Matrice de confusion avec le ViT-80 et (b) Tableau de comparaison des performances de classification pour les différentes approches. Trois mesures d’erreurs sont utilisées pour la comparaison : Accuracy, F1 score et Recall. Les ViT-20, ViT-40, etc... indiquent une *tokenisation* par segmentation, la valeur de C est celle indiquée après le trait d’union, et ViT-CNN correspond à une *tokenisation* avec un CNN.

gagnée sur les autres méthodes : CNN, FCNN, SVM ou régression logistique. Les résultats sur la base de test sont décrits par la figure 7.4(a), où la matrice de confusion permet de voir l’unique donnée mal classée. Il s’agit d’un diffractogramme de Dolomite classé en Halite.

Une autre remarque concerne les approches utilisant la segmentation pour *tokeniser* le signal. Cette méthode semble efficace quel que soit le choix de C , puisque l’ensemble des modèles réussissent la classification pour plus de 95% des diffractogrammes. Cependant, au vu des différences de performances observées, l’impact sur le choix du paramètre C semble important. L’utilisation d’un CNN pour *tokeniser* le diffractogramme ne semble pas pertinente dans ce cas, les différentes expériences menées n’ayant pas abouti à une convergence du modèle.

Une fois le modèle entraîné, les visualisations (*Attention Rollout* et *Positional encoding*) sont présentées pour la meilleure approche : le ViT utilisant la segmentation avec $C = 80$.

7.3.2 Visualisation des résultats

Positional encoding Nous considérons d’abord le *positional encoding*. La figure 7.5 (a) montre la similarité cosinus des vecteurs de positions entraînés pour l’identification de phases minérales. Comme décrit précédemment, cela permet de voir les liens entre les plages angulaires.

Les liens les plus forts sur cette figure (en valeur absolue) sont entre les tokens 11-15 (~ 0.45) et 7-10 (~ 0.41). La figure 7.6 permet de voir que ces tokens liés correspondent à des pics de diffraction de certaines phases minérales : Les tokens 11-15 correspondent à des pics de Dolomite et d’Hématite, et les tokens 7-10 à ceux de la Dolomite et de la Calcite.

Cet outil permet de vérifier que le modèle apprend correctement les co-occurrences des pics de diffraction spécifiques à une phase minérale, ce qui ne garantit pas un CNN. Cela

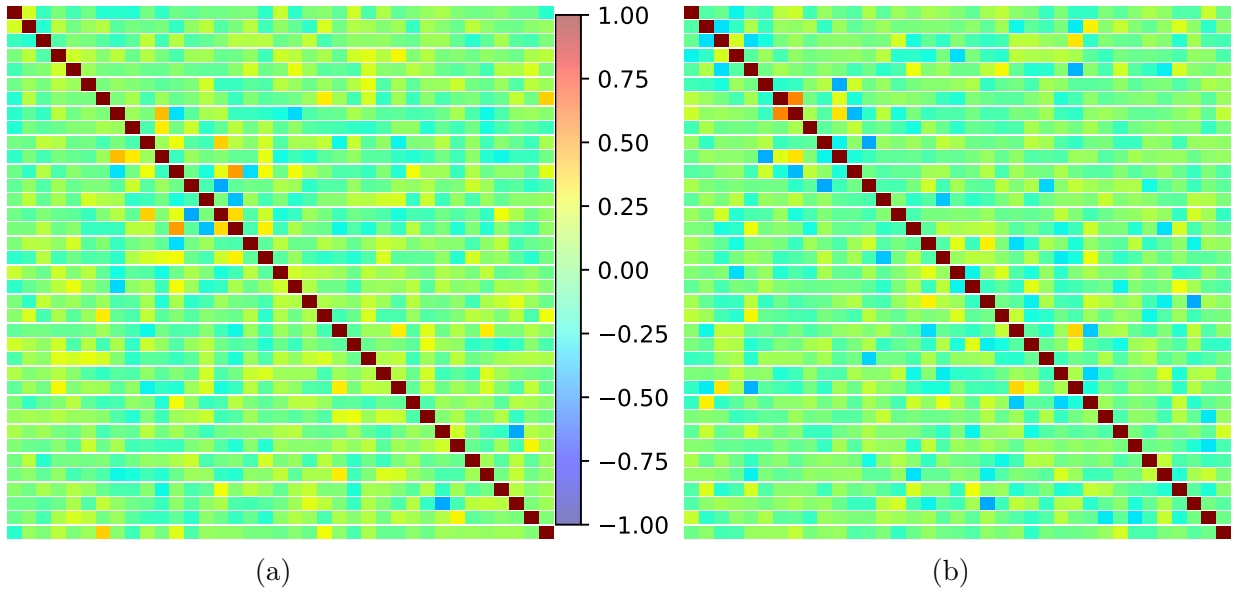


FIGURE 7.5 – Similarité cosinus des 36 vecteurs position appris par le modèle ViT. (a) est entraîné pour identifier les phases à partir du diffractogrammes, et (b) pour quantifier les proportions des phases minérales.

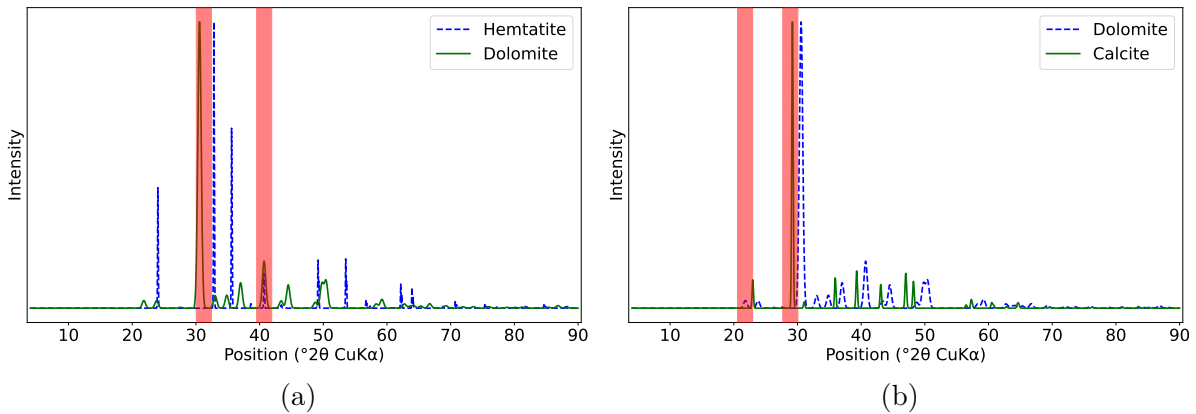


FIGURE 7.6 – Classification : visualisation des relations les plus importantes entre les vecteurs positions. (a) correspond aux tokens 11-15 et (b) aux tokens 7-10.

revient à mettre en évidence les dépendances de pics sur des plages angulaires différentes.

Attention Rollout Nous considérons ensuite l'*Attention Rollout* pour ce problème d'identification. La figure 7.7 montre des graphiques obtenus par ce moyen. Chaque graphique concerne l'une des six classes et représente les contributions des différentes plages angulaires sur le CLS token, donc sur la classification.

Il est nécessaire de préciser qu'il est compliqué de bien comprendre le fonctionnement du modèle. En effet, les ViT possèdent plusieurs centaines de milliers (voir millions) de paramètres qui rendent les modèles complexes. L'interprétation des graphiques de la figure 7.7 est plus difficile.

Pour trois des six classes, le modèle semble utiliser une majorité du signal pour construire le CLS token (qui effectue la tâche de classification). Ces trois phases sont la Hématite, la Halite et le Quartz. Un grand nombre de tokens (et donc de plages angulaires) sont en rouge sur ces graphiques. Il est donc compliqué de comprendre comment

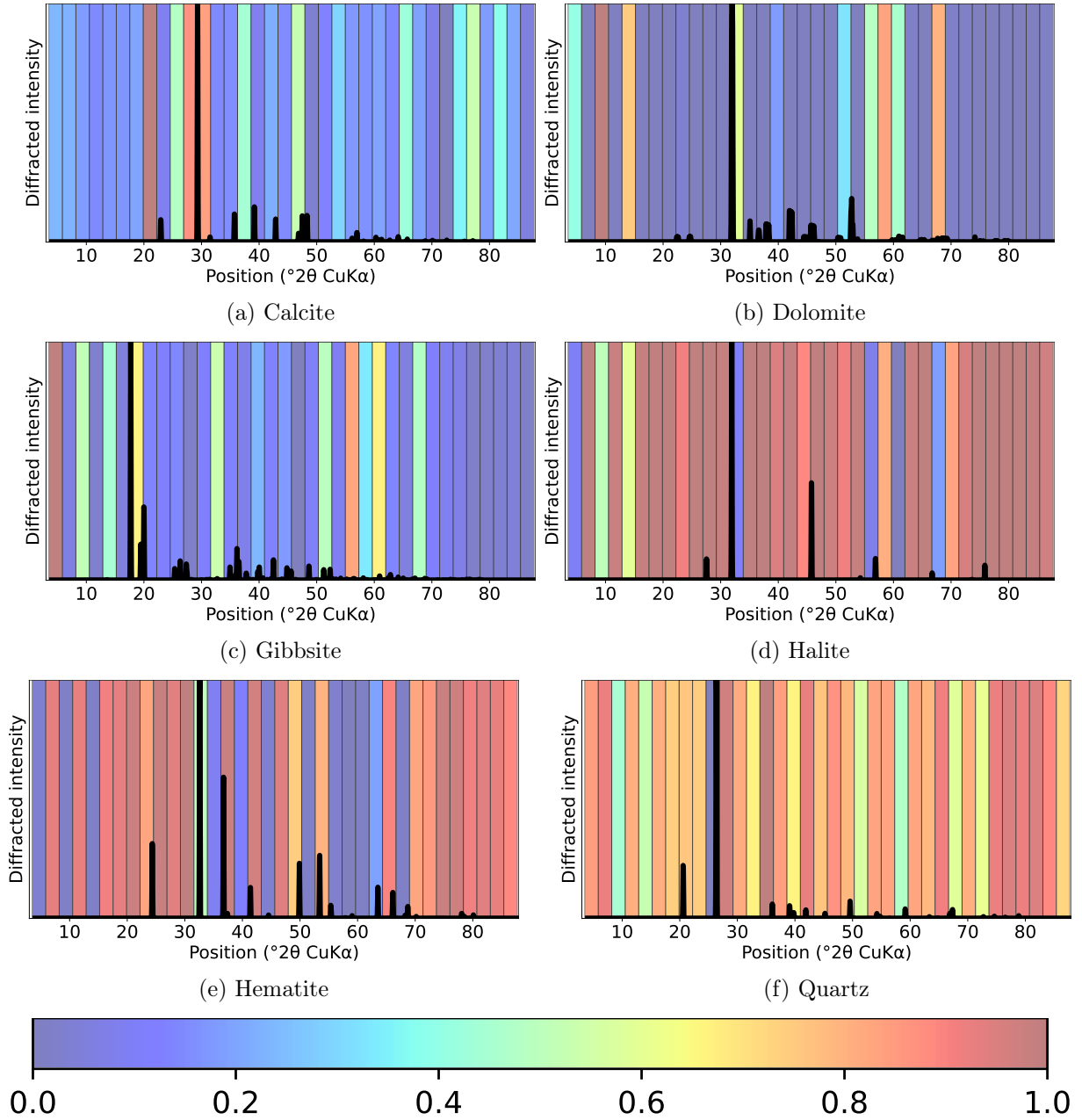


FIGURE 7.7 – Visualisation de l'*Attention Rollout* pour chacune des classes. L'*Attention Rollout* est moyenné sur l'ensemble des diffractogrammes de la phase minérale. Un exemple de signal est ajouté au dessus pour une meilleure lecture du graphique.

le réseau fonctionne sur ces données.

Pour les autres phases minérales (Calcite, Dolomite et Gibbsite), il y a là uniquement quelques segments du signal qui se distinguent avec des valeurs d'*Attention Rollout* importantes. Dans les trois cas, nous remarquons que ces canaux correspondent à des pics de diffraction. Même, les plages angulaires contenant les pics les plus importants participent beaucoup à la construction du CLS token. Ils influent donc directement sur l'identification de la phase minérale par le modèle.

Nous pouvons aussi remarquer la proximité dans la façon de traiter la Dolomite et la Halite par le réseau. Les plages angulaires, qui se distinguent avec des valeurs hautes pour la Dolomite, et des valeurs faibles pour la Halite, sont très proches. C'est probablement

cette proximité qui a mené à la mauvaise classification d’un diffractogramme de Dolomite en Halite. Il faut aussi mentionner le décalage de la position du pic dû à la variabilité des paramètres de maille. La figure 7.8 permet de visualiser l’*Attention Rollout* pour cette donnée mal classée.

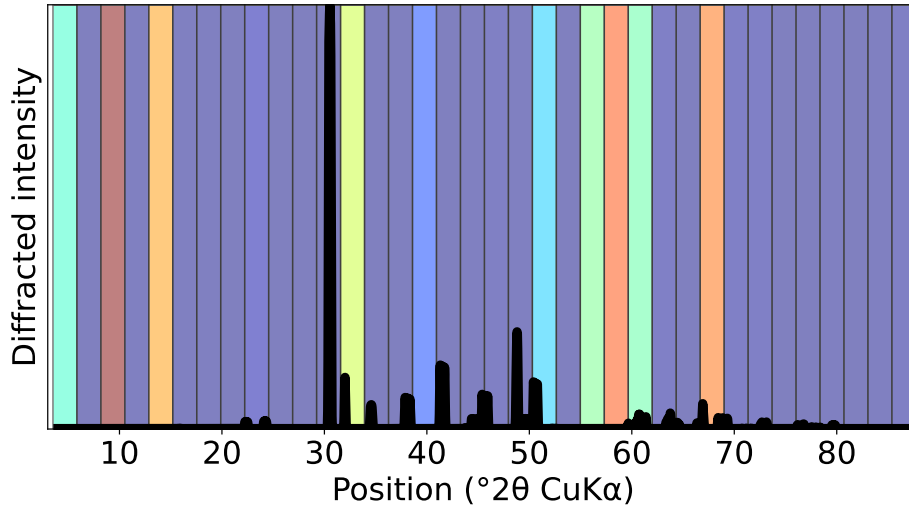


FIGURE 7.8 – Visualisation de l’*Attention Rollout* pour la donnée mal classée. Il s’agit d’une Dolomite classée en Halite. La colorbar est identique à celle de la figure 7.7

Finalement, pour la tâche de classification, le réseau présente de meilleurs résultats que toute autre méthode testée auparavant, incluant les CNN traditionnels. Nous allons maintenant confronter le ViT au problème plus complexe de quantification des phases minérales.

7.4 Inférence de proportions : Quantification des phases minérales des données de laboratoire

Le second problème est celui de la quantification de phases minérales, notamment avec l’analyse du jeu de données de laboratoire décrit dans le chapitre 5.

7.4.1 Entraînement et résultats pour la quantification des phases minérales

L’approche est similaire à celle des CNN. Le principal changement concerne le modèle, qui n’est plus un CNN, mais un Transformers. Le modèle est identique au ViT décrit précédemment dans la partie classification de la section 7.3.

Comme dans l’approche des CNN, le ViT est d’abord entraîné et testé sur des données synthétiques, avant d’être confronté aux données de laboratoire. Pour rappel, le modèle est entraîné sur 10 000 diffractogrammes synthétiques et la fonction de perte utilisée est celle combinant l’erreur des moindres carrés avec la modélisation de Dirichlet : “MSE & DIR”. Au vu des résultats sur la classification, les résultats et les graphiques présentés dans cette section sont obtenus avec un Transformer dont la *tokenisation* est réalisée par segmentation avec $C = 80$, même si le tableau 7.1 compare les résultats avec différentes *tokenisation* du signal.

Pour suivre l'entraînement du ViT-80, une base de validation est toujours utilisée. Cela permet de se rendre compte que 100 époques ne sont pas suffisantes pour un bon entraînement du modèle. Le Transformer est donc entraîné sur 1 000 époques.

La figure 7.9 permet de suivre l'entraînement à travers quatre graphiques. Le premier (a) permet de voir l'évolution de la fonction de perte au fur et à mesure des époques. Celle-ci décroît nettement durant les 200 premières époques. La décroissance se poursuit de manière moins marquée tout au long des 1 000 époques de l'entraînement. Les mesures de la RMSE (b) et de la MMAE (c) suivent ce même schéma, ce qui ne suggère pas un effet de sur-apprentissage. Au final, la RMSE atteint des valeurs légèrement inférieures à 2%, et la MMAE des valeurs inférieures à 2.5%.

Le dernier graphique (d) concerne la RRS. Ses valeurs sont très basses durant les 150 premières époques, avant de croître de manière significative sur les 100 époques suivantes jusqu'à être supérieures à 90%. Sur les époques restantes, la croissance de la RRS est plus faible.

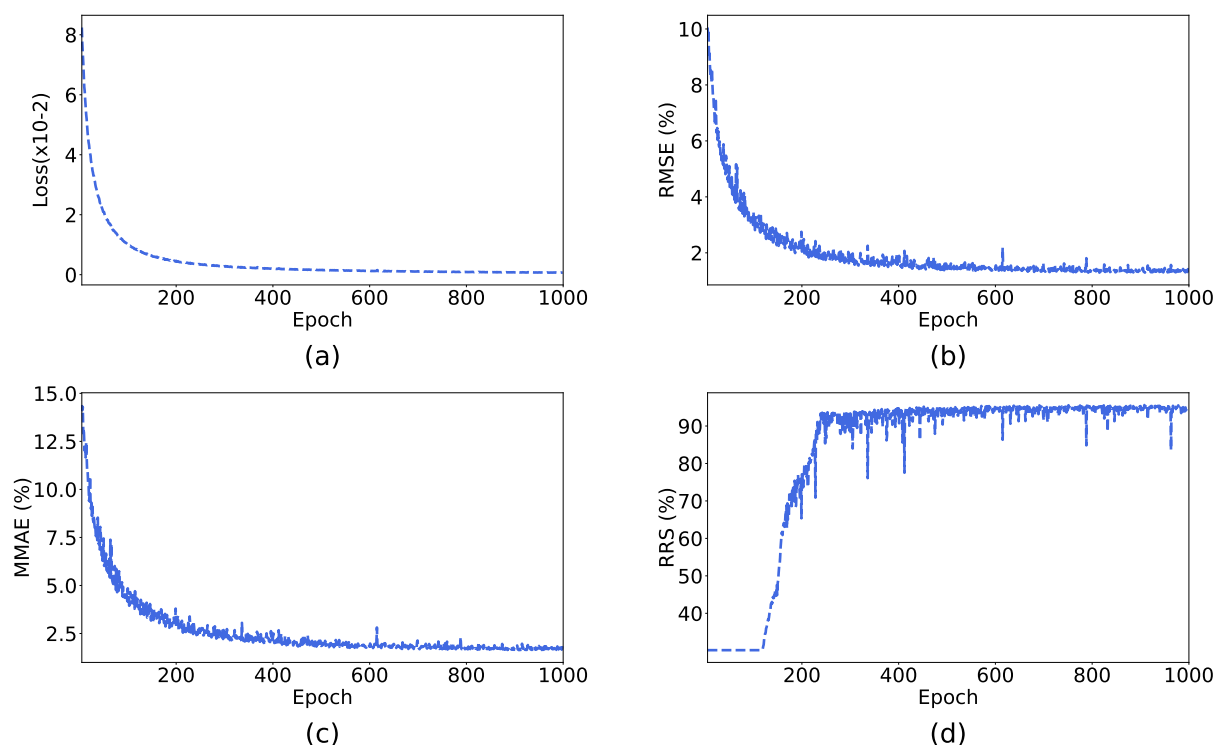


FIGURE 7.9 – Analyse de l'entraînement avec un modèle ViT : (a) Evolution des valeurs de la fonction de perte. Les figures (b), (c) et (d) décrivent les valeurs des mesures (RMSE, MMAE et RRS respectivement) sur l'ensemble de validation au cours des époques d'entraînement.

Les résultats de la base de validation sont confirmés par ceux obtenus sur l'ensemble de test contenant les données simulées. Les valeurs des différentes mesures sont décrites par le tableau 7.1. Comme pour la classification, ces résultats sont comparés avec ceux du CNN décrit dans le chapitre 5, mais aussi avec différentes approches pour la *tokenisation*. Sur les données simulées, ce tableau permet de voir que le modèle de ViT-CNN n'a pas convergé, et que l'influence du choix de C pour la segmentation est importante. Les grandes valeurs de C ($C = 80$ ou $C = 100$) produisent de meilleurs résultats que des valeurs plus faibles ($C = 20$ ou $C = 40$). Même si les valeurs sont proches c'est encore la segmentation avec le choix de $C = 80$ qui obtient les meilleurs résultats avec le ViT,

Méthode	Paramètres	Données simulées			Données réelles		
		RMSE ↓	MMAE ↓	RRS _{1%} ↑	RMSE↓	MMAE↓	RRS _{7%} ↑
CNN	832,868	00.49%	00.55%	97.40%	05.16%	6.96%	71.87%
ViT-20	60 844	03.02%	03.72%	86.28%	08.32%	10.95%	93.75%
ViT-40	236 884	02.55%	03.12%	86.60%	10.70%	12.00%	100%
ViT-80	934,564	<u>01.23%</u>	<u>01.56%</u>	87.20%	04.56%	05.81%	<u>96.90%</u>
ViT-100	1 456 204	01.29%	01.57%	<u>87.52%</u>	<u>04.69%</u>	<u>05.97%</u>	90.62%
ViT-CNN	215 466	Pas de convergence du modèle					

TABLE 7.1 – Résultats sur les ensembles de test simulés et réels pour les ViT. La comparaison est effectuée avec le CNN présenté dans le chapitre 5. Le nombre de paramètres de chaque modèle est indiqué.

même s'ils sont assez loin de ceux obtenus avec le CNN. Sur les simulations, le ViT-80 obtient une valeur de la RMSE égale à 1.23%, une valeur plus importante que les 0.49% du CNN. De même pour la MMAE (1.56%) et la RRS (87.2%), où les valeurs obtenues avec un CNN sont meilleures.

Les résultats sur les données de laboratoire sont également reportés sur le tableau 7.1. Ce sont les valeurs obtenues pour l'analyse de ce jeu de données qui montrent le potentiel des ViT pour la quantification de phases minérales.

Ainsi, même si les résultats sont moins bons sur les simulations, les modèles ViT-80 et ViT-100 ont réalisé un meilleur apprentissage des diffractogrammes puisque les valeurs des mesures d'erreurs sont meilleures sur l'ensemble des 32 données expérimentales. Les meilleures valeurs sont obtenues avec le ViT-80. La RMSE est meilleure de 0.5% pour atteindre 4.56%, la MMAE est améliorée de plus de 1% avec une valeur de 5.81%. Enfin la RRS (avec $\varepsilon = 7\%$) est améliorée de 20%, et vaut 96.9%. Il est également important de souligner que le ViT-40 obtient la meilleure valeur de RRS_{7%} avec 100%.

Les Transformers sont donc plus efficaces notamment pour retrouver les bonnes phases minérales dans chaque mélange. Ceci est confirmé par la figure 7.10 qui permet de visualiser le nuage de points des prédictions en fonction de la vérité, pour le ViT-80. La majorité des points et les centres de gravité sont situés sur la droite identité.

Surtout si l'on compare aux précédents résultats des CNN (figure 5.6), c'est aux extrémités que le ViT se distingue le plus. C'est-à-dire quand la donnée contient 0% ou 100% d'une phase minérale. Ces points sont plus concentrés autour des coordonnées (0, 0) et (1, 1), tout comme les centres de gravité.

Nous constatons également que les points correspondant à la Gibbsite sont les plus éloignés. Le modèle a donc plus de difficultés à quantifier cette phase. Cela peut s'expliquer par la plus grande complexité du diffractogramme de la Gibbsite, qui possède un nombre important de pics de diffraction (pour un exemple voir figure 7.11(a)).

Ce tableau compare également le nombre de paramètres et permet d'observer que la *tokenization* avec un CNN limite le nombre de paramètres. On constate également que le nombre de paramètres d'un ViT augmente lorsque la valeur de C pour la segmentation augmente, et que pour les ViT les plus performants, le nombre de paramètres est supérieur à celui des CNN. En effet, le ViT-100 possède plus de 1,4 millions de paramètres à optimiser, et le ViT-80, 934 564 paramètres, contre 832 868 pour le CNN. Même si ces chiffres sont relativement proches et du même ordre de grandeur pour le ViT-80 et le CNN, les temps d'entraînement sont différents. Cela s'explique par un nombre d'époques 10 fois plus important, qui mène approximativement à un entraînement dix fois plus long,

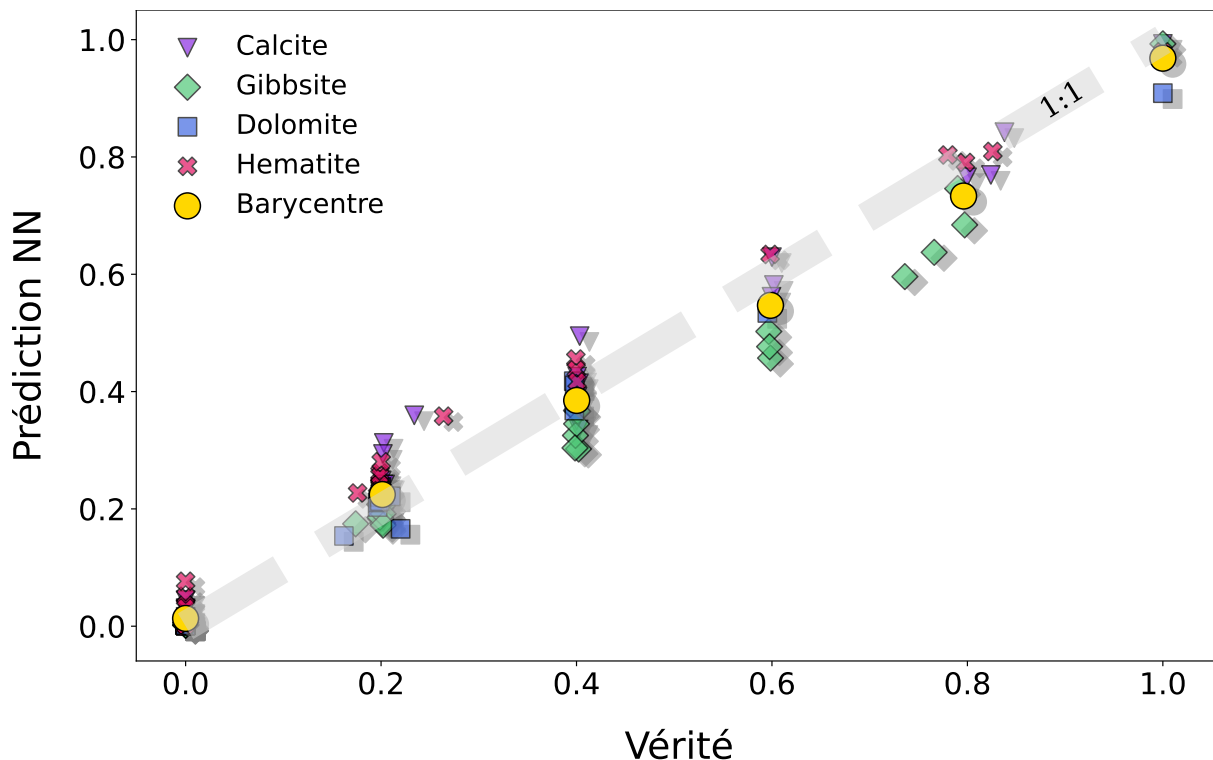


FIGURE 7.10 – Nuage de points des prédictions en fonction de la vérité.

Des expériences menées sur GPU (NVIDIA Tesla P4), ont conclu à un temps d'entraînement de 13 minutes pour le CNN sur 100 époques, et 170 minutes pour le ViT-80 sur 1 000 époques.

Finalement, même si le ViT est moins performant sur les simulations, il semble mieux généraliser pour traiter les données réelles. Puisque c'est bien les Transformers qui répondent le mieux à la problématique en fournissant une meilleure analyse sur les données réelles. De plus, les outils de visualisation évoqués précédemment sont disponibles et permettent de voir comment le ViT traite la donnée.

7.4.2 Visualisation des résultats

Les graphiques proposés dans cette section sont obtenus à partir de l'entraînement du ViT-80.

Positional encoding Le premier outil est le *positional encoding*. La matrice des similarités cosinus est visible sur la figure 7.5(b). Plusieurs plages angulaires sont particulièrement liées. Les tokens 6-7, et 9-12 sont ceux dont les valeurs absolues des similarités cosinus sont les plus grandes avec des valeurs proches de 0.4.

Les graphiques de la figure 7.11 montrent que ces tokens sont liés à des concordances de pics. Les tokens 6 et 7 correspondent à une concordance des deux pics les plus intenses de la Gibbsite. Pour les tokens 9 et 12 qui sont corrélés négativement, on observe que la présence de pics sur une plage angulaire entraîne l'absence de pics sur l'autre plage angulaire. Ce phénomène de non co-occurrence est observable sur les diffractogrammes de Gibbsite et de Dolomite.

Encore une fois, le positional encoding met en évidence des dépendances de pics,

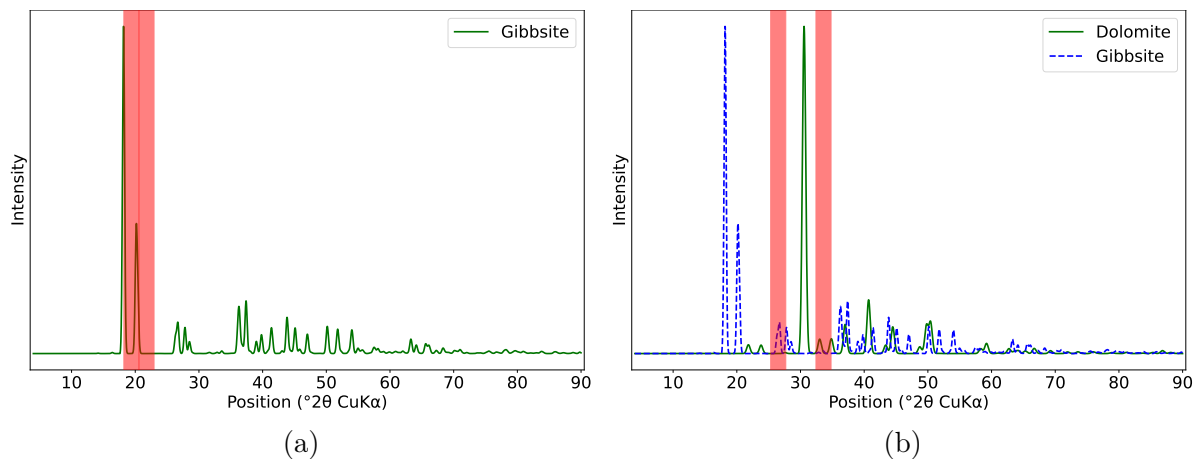


FIGURE 7.11 – Quantification de phases : visualisation des relations les plus importantes entre les vecteurs positions. (a) correspond aux tokens 6-7 (cossim = 0.42) et (b) aux tokens 9-12 (cossim = -0.41).

facilitant la compréhension du signal.

Attention Rollout Le second outil est l'*Attention Rollout*. Il permet de voir la propagation de l'information au sein du ViT. Les graphiques de la figure 7.12 permettent de visualiser cela.

Nous pouvons voir sur les deux graphiques que ce sont des plages angulaires possédant des pics de diffraction qui participent le plus à la construction du CLS token, donc à la quantification des phases. Le réseau semble se baser sur ces pics là pour effectuer la tâche qui lui est demandée.

Cependant, comme cela a été dit, la lecture de ce type de graphique n'est pas évidente étant donnée la complexité du modèle. C'est d'autant plus valable que nous considérons le problème d'inférence de proportion.

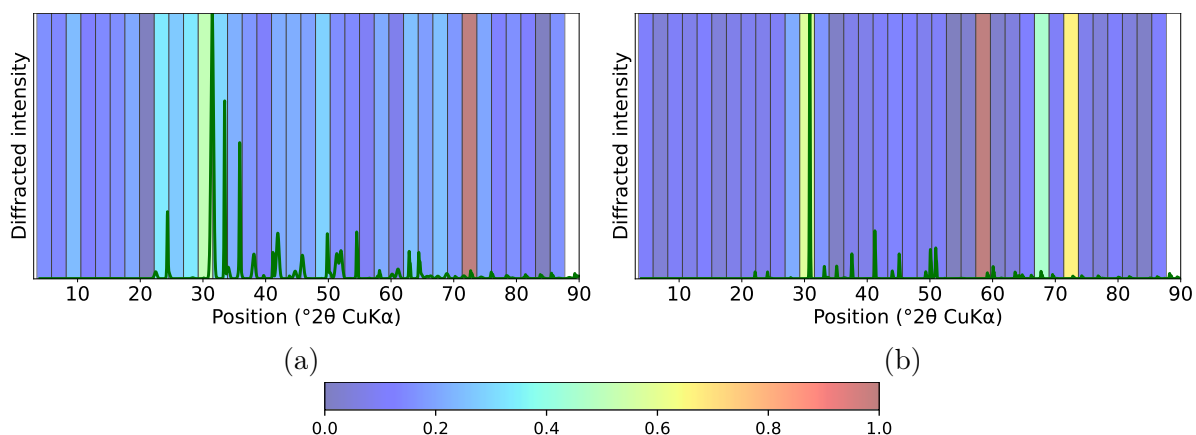


FIGURE 7.12 – Visualisation de l'*Attention Rollout* à travers l'exemple d'un diffractogramme pour chacune des classes. Les données sont issues de la base de validation, dont nous donnerons la vérité et la prédiction en termes de phases minérales. (a) est composé de 73% de Dolomite et 27% de Gibbsite, la prédiction du modèle est 71% de Dolomite et 29% Hématite. (b) est composé de 97% de Dolomite et 3% de Gibbsite, la prédiction du modèle est 97% de Dolomite et 3% Hématite.

7.5 Conclusion

Pour conclure ce chapitre, plusieurs éléments sont à souligner. Le plus important est de mettre en valeur les performances des Transformers (plus spécifiquement des ViT) pour l'identification et la quantification de phases minérales. La classification des diffractogrammes selon la phase minérale a obtenu de meilleurs résultats que les réseaux de neurones convolutifs. La phase est bien identifiée pour plus de 99% des signaux, contre 96% pour un CNN.

Il faut aussi considérer la quantification de phases minérales, même si concernant les données simulées, les valeurs des mesures sont moins bonnes que le CNN. Le ViT semble mieux apprendre la donnée car en passant aux données réelles, toutes les mesures sont meilleures. Cela s'observe notamment pour retrouver le support de la donnée, avec plus de 96% de réussite (au seuil $\varepsilon = 7\%$).

De plus, passer de l'utilisation des CNN à celle d'un Transformers n'est pas très coûteux, car l'unique changement important est celui du modèle. D'un point de vue méthodologique l'approche est similaire et les temps de calcul sont comparables. Même si le nombre d'époques nécessaire est plus important, les 170 minutes d'entraînement du modèle sur GPU restent acceptables.

Il est aussi important de considérer les outils supplémentaires liés au principe d'attention. Premièrement avec le *positional encoding* qui offre une meilleure compréhension des signaux et permet de visualiser les concordances de pics. Surtout, il permet de voir si le Transformer est capable de les mettre en évidence.

Ensuite, grâce à l'*Attention Rollout*, il est possible d'observer la propagation de l'information au sein des couches du modèle. Il permet de visualiser l'interprétation de la donnée par le modèle, par exemple pour se rendre compte de l'importance de certains pics pour l'identification de la phase minérale.

Dans le cadre de nos travaux, l'utilisation des Transformers a un aspect novateur pour l'analyse des diffractogrammes, car les travaux sur ce sujet sont peu développés. Afin de voir leur potentiel, l'objectif était de tester et d'introduire leur utilisation pour ce type de tâches en traitement du signal. Les bons résultats obtenus permettent de valider cette première approche où le ViT surpasse déjà les méthodes alternatives.

Même si une importante partie du travail concernant la génération de bases d'apprentissage synthétiques avait été réalisée auparavant, les résultats laissent penser que les ViT présentent plus de garanties en termes de résultats, mais aussi pour l'interprétation et la compréhension du modèle.

De plus, de nombreux paramètres restent encore à optimiser pour obtenir de meilleurs résultats avec des Transformers, sur le choix du paramètre C , de profondeur du réseau L ou encore sur le choix du nombre de têtes d'Attention.

Chapitre 8

Conclusion et perspectives

8.1 Conclusion

L’objectif de la thèse était de proposer des méthodes d’apprentissage de type *machine learning* pour quantifier et identifier des phases minérales dans des grands jeux de données, à l’image de ceux obtenus sur synchrotron grâce à l’analyse par tomographie par diffraction de rayons X qui permet de faire une étude “par tranche” d’un solide.

La première étape s’est bâtie autour de la nécessité de construire une base de données de diffractogrammes pour entraîner un réseau de neurones. Afin que le réseau apprenne correctement, il était important que les données d’entraînement représentent le plus possible les données auxquelles le réseau sera confronté plus tard. Pour des raisons présentées auparavant, il n’était pas possible d’utiliser des données expérimentales. Le choix s’est donc porté sur la synthétisation de données. Des algorithmes de simulations ont été développés ou améliorés pour synthétiser des diffractogrammes.

Les résultats de nos simulations montrent la robustesse des algorithmes. La position et la forme sont respectées, les ratios d’intensité entre les pics également. La différence majeure concerne plutôt les artefacts liés aux appareils de mesures qui sont compliqués à mesurer. Au final, cela a permis de générer plusieurs bases de données avec une importante quantité de diffractogrammes, tout en incluant une grande variabilité. Cette approche permet donc de contourner le manque de données expérimentales, ainsi que les limitations que cela engendre.

Dans un deuxième temps, il a été question du problème d’inférence de proportions. La quantification des phases minérales revient à chercher un vecteur de proportion dont chaque élément correspond à l’abondance d’une phase minérale. Or, cette tâche est pour le moment assez peu traitée avec les réseaux de neurones. En s’appuyant sur la distribution de Dirichlet, des fonctions de perte adaptées ont été créées. L’une d’entre elles, (“MSE & Dirichlet”) montrait d’excellentes performances, elle a aussi révélé une stabilité et une robustesse importante quel que soit le type de données que le modèle doit traiter.

La dernière étape était de confronter la méthode à des jeux de données expérimentales. Le choix a été fait d’utiliser une architecture de réseau déjà existante, puis d’entraîner le modèle en utilisant une base de données synthétiques avec une fonction de perte adaptée. Un perfectionnement des simulations et des bases de données d’entraînement a été réalisé au fur et à mesure des travaux.

Finalement, les résultats suivants ont été obtenus. Pour les 32 diffractogrammes mesurés avec un diffractomètre classique, le réseau est capable de quantifier les phases à environ 5-6% près, et d’identifier les phases pour presque 80% des données. Nous avons

aussi montré qu'un autre type de modèle, les Transformers, montrait de meilleures performances. Les erreurs de quantification sont diminuées d'un pourcent, et l'identification est réussie pour plus de 96% des diffractogrammes.

Pour les données issues de la tomographie par DRX, les erreurs ne peuvent pas être quantifiées de la même manière car la vérité n'est pas connue. Cependant, la comparaison effectuée avec un traitement utilisé préalablement, basé sur l'intensité des pics des différentes phases permet de valider l'efficacité de l'approche.

La méthode développée dans ces travaux permet donc de faire face à des jeux de données complexes obtenus par DRX. La quantification des phases grâce aux réseaux de neurones est réalisable, avec une marge d'erreur de l'ordre de quelques pourcents. Des gains sur ces performances sont tout de même envisageables.

8.2 Perspectives

8.2.1 Amélioration des résultats pour la quantification de phases minérales

Un certain nombre de pistes pour améliorer les résultats obtenus ont déjà été évoquées dans les chapitres précédents.

Par exemple, l'optimisation du choix de l'architecture du modèle. En effet, pour le moment, notre approche se base sur des architectures de réseaux déjà existantes pour des problèmes similaires à celui de la quantification de phases minérales à partir de diffractogrammes. Des travaux concernant l'amélioration de l'architecture devraient donc mener à de meilleurs résultats. C'est notamment ce qui peut être observé par les gains de performances que les Transformers amènent (voir chapitre 7).

L'amélioration de la qualité de la base d'apprentissage devrait également mener à de meilleurs résultats. Cela peut passer par une amélioration des simulations, par exemple en incluant du bruit lors de la synthétisation des diffractogrammes. Cela peut aussi passer par une meilleure construction de la base de données, simplement en augmentant la taille de la base, ou en cherchant une meilleure répartition des mélanges. Toujours dans l'optique d'améliorer les données, un travail plus important peut être réalisé sur le traitement des données réelles, pour réduire les effets de l'instrument de mesure, ou du bruit sur le signal.

8.2.2 Autres pistes exploratoires

En parallèle des améliorations envisagées précédemment, d'autres pistes peuvent être explorées.

Reconstruction du signal avec des auto-encodeurs La reconstruction d'un signal en s'appuyant sur le vecteur de proportion prédit, ou avec un auto-encodeur, serait un bon moyen d'évaluer les performances de la méthode. Cela serait aussi utile pour la détection de nouvelles phases minérales. Cependant, pour reconstruire à partir du vecteur de proportion de l'équation (4.1), les vecteurs $\mathbf{C}_i$ doivent également être connus. Ce n'est pas le cas, puisqu'il existe une infinité de diffractogrammes (variation intra-classe) correspondant à une phase minérale, avec la variation des paramètres de maille, des répétitions de maille dans les trois directions de l'espace, l'agitation thermique, etc.

L'approche des Auto-encodeurs (AE) (GOODFELLOW *et al.* 2016) pour proposer une reconstruction peut être envisagée, à la manière des études dans le domaine des images

hyperspectrales, qui proposaient de reconstruire des signaux. Un auto-encodeur, notamment qui permettait à la fois une reconstruction du signal, l'extraction des signaux des *endmembers* (i.e composants) et enfin d'obtenir les abondances (PALSSON *et al.* 2018). Ce type d'approche serait pertinent pour les diffractogrammes, car tous les éléments importants seraient fournis. Une étude de UTIMULA *et al.* (2020) propose même une méthode d'identification des pics caractéristiques des diffractogrammes en construisant un espace latent grâce à un AE. Mais, ces travaux se concentrent sur la projection des signaux de diffraction dans l'espace latent, et ne cherchent pas à produire une reconstruction.

Pour le moment nos travaux réalisés pour la reconstruction de diffractogrammes n'ont pas abouti. Ils ne permettent pas une bonne reconstruction du signal, bien que de nombreux tests aient été effectués, notamment sur l'architecture où des AE avec des couches linéaires, et des architectures U-Net (RONNEBERGER *et al.* 2015) ont été testées. L'extraction des signaux correspondant aux vecteurs $\mathbf{C}_i$ a aussi été considérée, mais les travaux n'ont pas abouti non plus. Pour des données simulées, représentant un mélange de $K = 4$ classes, les figures 8.1 et 8.2 ont, par exemple, pu être obtenues.

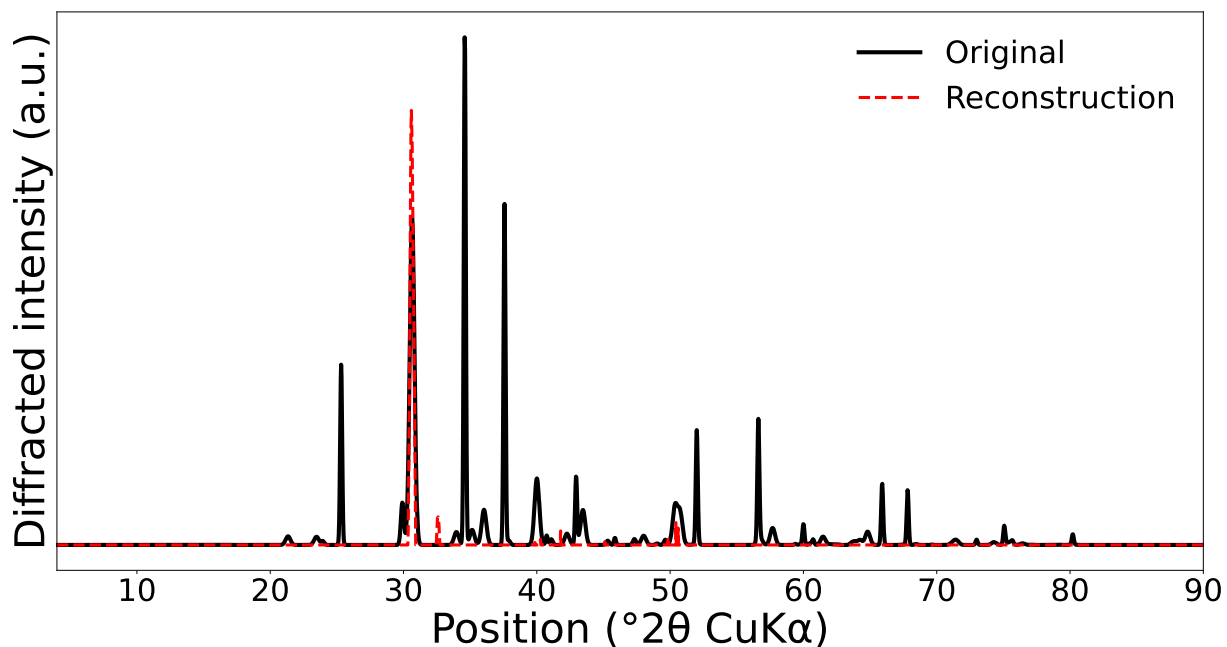


FIGURE 8.1 – Comparaison d'un diffractogramme et de sa reconstruction avec un auto-encodeur. La donnée est synthétique et composée de 4% de Calcite, 68% de Dolomite et 27% d'Hématite.

Estimation des paramètres de maille Une autre piste pour arriver à identifier les vecteurs $\mathbf{C}_i$ serait de trouver les paramètres nécessaires à la génération des données, soit les paramètres de maille, le coefficient de Debye-Waller et la largeur à mi hauteur des pics. Des travaux similaires ont été réalisés à l'aide de réseaux de neurones convolutifs (H. DONG *et al.* 2021). Ils permettent d'identifier différents paramètres à l'aide du diffractogramme : les paramètres de maille, le facteur d'échelle et la taille de cristallite. Cependant cette étude se limite à des phases dont la maille est cubique. Le diffractogramme est donc plus simple à comprendre pour le modèle, car il y a une dépendance linéaire entre la variation des paramètres de maille et la position des pics de diffraction.

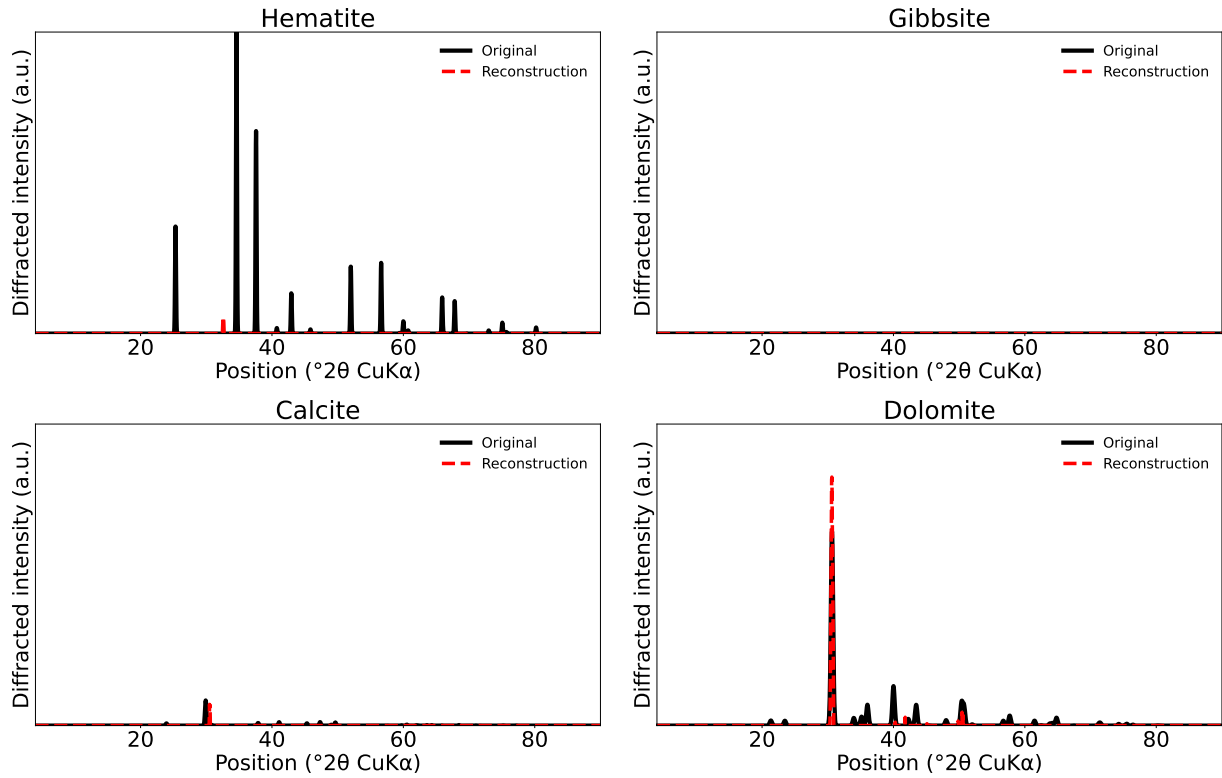


FIGURE 8.2 – Comparaison des signaux C_i reconstruits par l’auto-encodeur avec ceux originaux. Comme pour la figure 8.1, la donnée est synthétique et composée de 4% de Calpajene, 68% de Dolomite et 27% d’Hématite.

Analyse en dimension spatiale Une perspective concernant l’étude des jeux de données obtenus par tomographie de diffraction de rayons X peut être considérée. Comme l’analyse permet de trouver des images, on pourrait imaginer s’appuyer sur deux dimensions, la dimension spectrale (liée au signal) d’une part et la dimension spatiale d’autre part, qui tiendrait compte des liens entre les pixels proches. Le traitement des images se ferait alors par cube de pixels de dimensions choisies, comme les travaux que ZHAO *et al.* (2022) proposent pour les images hyperspectrales. Dans l’application à la diffraction, la difficulté réside dans la construction d’une base de données adaptée pour cette tâche, puisque des cartes d’abondances physiquement cohérentes devraient être créées.

Complémentarité avec l’affinement Rietveld Finalement, on pourrait imaginer combiner nos travaux avec l’affinement Rietveld. Considérant un jeu de données, la première étape serait d’utiliser notre méthode pour effectuer l’identification et proposer une première quantification des phases minérales. À partir de ces résultats, il s’agirait de proposer une écriture de fichier automatique pour l’entrée de l’affinement Rietveld. Chaque donnée aurait son propre fichier pour simplifier la convergence de l’algorithme. Alors la technique développée par Rietveld pourrait fournir une quantification précise, et s’affranchir de la tâche d’identification. De plus, une reconstruction et des incertitudes seraient fournies.

Les bons résultats obtenus et présentés dans ce manuscrit ouvrent la porte aux nombreuses perspectives évoquées ci-dessus. Elles ont pour but d’améliorer ou de compléter les résultats fournis par le réseau de neurones, en donnant toujours plus d’informations

sur la minéralogie des cristaux, que ce soit à travers des reconstructions, ou avec une estimation plus précise.

Bibliographie

Les numéros des pages où les références sont citées sont indiqués à la suite de la référence.

ABNAR, S. & ZUIDEMA, W. (juill. 2020). “Quantifying Attention Flow in Transformers”. In : *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*. Sous la dir. de D. JURAFSKY, J. CHAI, N. SCHLUTER & J. TETREAULT. Online : Association for Computational Linguistics, p. 4190-4197. DOI : 10.18653/v1/2020.acl-main.385. URL : <https://aclanthology.org/2020.acl-main.385> (cf. p. 8, 91, 96, 97).

ALLMANN, R (1977). “Refinement of the hybrid layer structure $[\text{Ca}_2\text{Al}(\text{OH})_6] + [1/2\text{SO}_4 \cdot 3\text{H}_2\text{O}]$ ”, *Neues Jahrb. Für Mineral*. *Monatshefte* 1977, p. 136-143 (cf. p. 76).

ALLMANN, R & JEPSEN, HP (1969). “Die struktur des hydrotalkits”. *Neues Jahrbuch für Mineralogie Monatshefte* 1969, 544-551 (cf. p. 76).

ANGEL, RJ, McMULLAN, RK & PREWITT, Charles T (1991). “Substructure and superstructure of mullite by neutron diffraction”. *American Mineralogist* 76.3-4, p. 332-342 (cf. p. 76).

ÁNGELES, G, DE VERA, Ruth N, CUBEROS, Antonio JM & ARANDA, Miguel AG (2008a). “Crystal structure of low magnesium-content alite : Application to Rietveld quantitative phase analysis”. *Cement and Concrete Research* 38.11, p. 1261-1269 (cf. p. 20).

ÁNGELES, G, DE VERA, Ruth N, CUBEROS, Antonio JM & ARANDA, Miguel AG (2008b). “Crystal structure of low magnesium-content alite : Application to Rietveld quantitative phase analysis”. *Cement and Concrete Research* 38.11, p. 1261-1269 (cf. p. 76).

BALAN, E., LAZZERI, M., MORIN, G. & MAURI, F. (2006). “First-principles study of the OH-stretching modes of gibbsite”. *American Mineralogist* 91.1, p. 115-119 (cf. p. 32, 42, 61, 97).

BIOUCAS-DIAS, J. M, PLAZA, A., DOBIGEON, N., PARENTE, M., DU, Q., GADER, P. & CHANUSSOT, J. (2012). “Hyperspectral unmixing overview : Geometrical, statistical, and sparse regression-based approaches”. *IEEE journal of selected topics in applied earth observations and remote sensing* 5.2, p. 354-379 (cf. p. 4).

BIOUCAS-DIAS, J. M. & FIGUEIREDO, M. A. T. (2010). “Alternating direction algorithms for constrained sparse regression : Application to hyperspectral unmixing”. In : *2nd Workshop on Hyperspectral Image and Signal Processing : Evolution in Remote Sensing*, p. 1-4. DOI : 10.1109/WHISPERS.2010.5594963 (cf. p. 4, 47).

BIRD, Dennis K & SPIELER, Abigail R (2004). “Epidote in geothermal systems”. *Reviews in Mineralogy and Geochemistry* 56.1, p. 235-300 (cf. p. 1).

BISH, D. L. & POST, J. E. (1990). *Modern powder diffraction*. Walter de Gruyter GmbH & Co KG (cf. p. 2, 16, 20, 42).

BISHOP, Christopher M & NASRABADI, Nasser M (2006). *Pattern recognition and machine learning*. 4. 4. Springer (cf. p. 27).

BLAKE, R. L., HESSEVICK, R. E., ZOLTAI, T. & FINGER, L. W. (1966). “Refinement of the hematite structure”. *American Mineralogist : Journal of Earth and Planetary Materials* 51.1-2, p. 123-129 (cf. p. 32, 42, 61, 97).

BRAHIM, Abdelbasset, JENNANE, Rachid, RIAD, Rabia, JANVIER, Thomas, KHEDHER, Laila, TOUMI, Hechmi & LESPESSAILLES, Eric (2019). “A decision support tool for early detection of knee OsteoArthritis using X-ray imaging and machine learning : Data from the OsteoArthritis Initiative”. *Computerized Medical Imaging and Graphics* 73, p. 11-18. ISSN : 0895-6111. DOI : <https://doi.org/10.1016/j.compmedimag.2019.01.007>. URL : <https://www.sciencedirect.com/science/article/pii/S0895611119300035> (cf. p. 27).

BROWN, Tom, MANN, Benjamin, RYDER, Nick, SUBBIAH, Melanie, KAPLAN, Jared D, DHARIWAL, Prafulla, NEELAKANTAN, Arvind, SHYAM, Pranav, SASTRY, Girish, ASKELL, Amanda, AGARWAL, Sandhini, HERBERT-VOSS, Ariel, KRUEGER, Gretchen, HENIGHAN, Tom, CHILD, Rewon, RAMESH, Aditya, ZIEGLER, Daniel, WU, Jeffrey, WINTER, Clemens, HESSE, Chris, CHEN, Mark, SIGLER, Eric, LITWIN, Mateusz, GRAY, Scott, CHESSE, Benjamin, CLARK, Jack, BERNER, Christopher, MCCANDLISH, Sam, RADFORD, Alec, SUTSKEVER, Ilya & AMODEI, Dario (2020). “Language Models are Few-Shot Learners”. In : *Advances in Neural Information Processing Systems*. Sous la dir. de H. LAROCHELLE, M. RANZATO, R. HADSELL, M.F. BALCAN & H. LIN. **33**. Curran Associates, Inc., p. 1877-1901. URL : https://proceedings.neurips.cc/paper_files/paper/2020/file/1457c0d6bfc4967418bfb8ac142f64a-Paper.pdf (cf. p. 8, 91).

CAO, Fuli, YE, Jiayuan, SUN, Zitang, REN, Xuehong, YAO, Xinyue, ZHAI, Munan & ZHANG, Wensheng (2024). “Compressive strength prediction of Portland cement clinker using machine learning by XRD data”. *Journal of Sustainable Cement-Based Materials*, p. 1-12 (cf. p. 91).

CARION, Nicolas, MASSA, Francisco, SYNNAEVE, Gabriel, USUNIER, Nicolas, KIRILLOV, Alexander & ZAGORUYKO, Sergey (2020). “End-to-end object detection with transformers”. In : *European conference on computer vision*. Springer, p. 213-229 (cf. p. 91).

CHE, C., ZHANG, P., ZHU, M., QU, Y. & JIN, B. (2021). “Constrained transformer network for ECG signal processing and arrhythmia classification”. *BMC Medical Informatics and Decision Making* 21.1, p. 184 (cf. p. 91, 95).

CHEN, Z., XIE, Y., WU, Y., LIN, Y., TOMIYA, S. & LIN, J. (2024). “An interpretable and transferrable vision transformer model for rapid materials spectra classification”. *Digital Discovery* 3.2, p. 369-380 (cf. p. 3, 91, 93, 95).

CHITTURI, Sathya R., RATNER, Daniel, WALROTH, Richard C., THAMPY, Vivek, REED, Evan J., DUNNE, Mike, TASSONE, Christopher J. & STONE, Kevin H. (déc. 2021). “Automated prediction of lattice parameters from X-ray powder diffraction patterns”. *Journal of Applied Crystallography* 54.6, p. 1799-1810. DOI : [10.1107/S1600576721010840](https://doi.org/10.1107/S1600576721010840). URL : <https://doi.org/10.1107/S1600576721010840> (cf. p. 3).

CHOI, Keunwoo, FAZEKAS, György, SANDLER, Mark & CHO, Kyunghyun (2017). “Convolutional recurrent neural networks for music classification”. In : *2017 IEEE International conference on acoustics, speech and signal processing (ICASSP)*. IEEE, p. 2392-2396 (cf. p. 27).

CHOWDHARY, KR1442 & CHOWDHARY, KR (2020). “Natural language processing”. *Fundamentals of artificial intelligence*, p. 603-649 (cf. p. 91).

CLARET, F., GRANGEON, S., LOSCHETTER, A., TOURNASSAT, C., DE NOLF, W., HARKER, N., BOULAHYA, F., GABOREAU, S., LINARD, Y., BOURBON, X. *et al.* (2018). “Deciphe-

ring mineralogical changes and carbonation development during hydration and ageing of a consolidated ternary blended cement paste”. *IUCrJ* 5.2, p. 150-157 (cf. p. 2, 14, 75, 76, 81, 83, 84).

CROMER, Don T & MANN, Joseph B (1968). “X-ray scattering factors computed from numerical Hartree–Fock wave functions”. *Acta Crystallographica Section A : Crystal Physics, Diffraction, Theoretical and General Crystallography* 24.2, p. 321-324 (cf. p. 15).

DEBYE, P (1915). “Dispersion of Roentgen rays”. *Ann. Phys* 46.809, p. 21 (cf. p. 16).

DEVLIN, Jacob (2018). “Bert : Pre-training of deep bidirectional transformers for language understanding”. *arXiv preprint arXiv :1810.04805* (cf. p. 8, 91).

DOBIGEON, Nicolas, MOUSSAOUI, Saïd, COULON, Martial, TOURNERET, Jean-Yves & HERO, Alfred O. (2009). “Joint Bayesian Endmember Extraction and Linear Unmixing for Hyperspectral Imagery”. *IEEE Transactions on Signal Processing* 57.11, p. 4355-4368. DOI : 10.1109/TSP.2009.2025797 (cf. p. 4).

DOEBELIN, N. & KLEEGERG, R. (2015). “Profex : a graphical user interface for the Rietveld refinement program BGMN”. *Journal of applied crystallography* 48.5, p. 1573-1580 (cf. p. 1, 22, 63).

DONG, Chao, LOY, Chen Change, HE, Kaiming & TANG, Xiaoou (2015). “Image super-resolution using deep convolutional networks”. *IEEE transactions on pattern analysis and machine intelligence* 38.2, p. 295-307 (cf. p. 27).

DONG, Hongyang, BUTLER, Keith T, MATRAS, Dorota, PRICE, Stephen WT, ODARCHENKO, Yaroslav, KHATRY, Rahul, THOMPSON, Andrew, MIDDELKOOP, Vesna, JACQUES, Simon DM, BEALE, Andrew M *et al.* (2021). “A deep convolutional neural network for real-time full profile analysis of big powder diffraction data”. *npj Computational Materials* 7.1, p. 74 (cf. p. 109).

DOSOVITSKIY, A., BEYER, L., KOLESNIKOV, A.I, WEISSENBORN, D., ZHAI, X., UNTERTHINER, T., DEGHANI, M., MINDERER, M., HEIGOLD, G., GELLY, S. *et al.* (2021). “An Image is Worth 16x16 Words : Transformers for Image Recognition at Scale”. In : *International Conference on Learning Representations*. URL : <https://openreview.net/forum?id=YicbFdNTTy> (cf. p. 8, 91, 93, 96).

FALL, Mame Diarra (2024). “QUANTIFYING UNCERTAINTY IN KNEE OSTEOARTHRITIS DIAGNOSIS”. In : *2024 IEEE 20th International Symposium on Biomedical Imaging (ISBI)* (cf. p. 38).

FERRAGE, Eric, LANSON, Bruno, MALIKOVA, Natalie, PLANÇON, Alain, SAKHAROV, Boris A. & DRITS, Victor A. (2005a). “New Insights on the Distribution of Interlayer Water in Bi-Hydrated Smectite from X-ray Diffraction Profile Modeling of 00l Reflections”. *Chemistry of Materials* 17.13, p. 3499-3512. DOI : 10.1021/cm047995v. eprint : <https://doi.org/10.1021/cm047995v>. URL : <https://doi.org/10.1021/cm047995v> (cf. p. 5).

FERRAGE, Eric, LANSON, Bruno, MICHOT, Laurent J. & ROBERT, Jean-Louis (2010). “Hydration Properties and Interlayer Organization of Water and Ions in Synthetic Na-Smectite with Tetrahedral Layer Charge. Part 1. Results from X-ray Diffraction Profile Modeling”. *The Journal of Physical Chemistry C* 114.10, p. 4515-4526. DOI : 10.1021/

jp909860p. eprint : <https://doi.org/10.1021/jp909860p>. URL : <https://doi.org/10.1021/jp909860p> (cf. p. 5).

FERRAGE, Eric, LANSON, Bruno, SAKHAROV, Boris A. & DRITS, Victor A. (août 2005b). “Investigation of smectite hydration properties by modeling experimental X-ray diffraction patterns : Part I. Montmorillonite hydration properties”. *American Mineralogist* 90.8-9, p. 1358-1374. ISSN : 0003-004X. DOI : 10.2138/am.2005.1776. eprint : https://pubs.geoscienceworld.org/msa/ammin/article-pdf/90/8-9/1358/3621208/15/_1776Ferrage.indd.pdf. URL : <https://doi.org/10.2138/am.2005.1776> (cf. p. 5).

GAMBOA, John Cristian Borges (2017). “Deep learning for time-series analysis”. *arXiv preprint arXiv :1701.01887* (cf. p. 2).

GATYS, Leon A, ECKER, Alexander S & BETHGE, Matthias (2015). “A neural algorithm of artistic style”. *arXiv preprint arXiv :1508.06576* (cf. p. 27).

GHOSH, P., ROY, S. K., KOIRALA, B., RASTI, B. & SCHEUNDERS, P. (2022). “Hyperspectral Unmixing Using Transformer Network”. *IEEE Transactions on Geoscience and Remote Sensing* 60, p. 1-16. DOI : 10.1109/TGRS.2022.3196057 (cf. p. 91, 95).

GILDENBLAT, Jacob (2020). *Exploring Explainability for Vision Transformers*. <https://jacobgil.github.io/deeplearning/vision-transformer-explainability>. Consulté le : 05 09 2024 (cf. p. 97).

GOODFELLOW, I., BENGIO, Y. & COURVILLE, A. (2016). *Deep learning*. MIT press (cf. p. 2, 27, 108).

GOODFELLOW, Ian, POUGET-ABADIE, Jean, MIRZA, Mehdi, XU, Bing, WARDE-FARLEY, David, OZAI, Sherjil, COURVILLE, Aaron & BENGIO, Yoshua (2020). “Generative adversarial networks”. *Communications of the ACM* 63.11, p. 139-144 (cf. p. 27, 30).

GRANGEON, Sylvain, BATAILLARD, Philippe & COUSSY, Samuel (2020). “The nature of manganese oxides in soils and their role as scavengers of trace elements : Implication for soil remediation”. *Environmental Soil Remediation and Rehabilitation : Existing and Innovative Solutions*, p. 399-429 (cf. p. 1).

GRANGEON, Sylvain, CLARET, Francis, LINARD, Yannick & CHIABERGE, Christophe (oct. 2013). “X-ray diffraction : a powerful tool to probe and understand the structure of nanocrystalline calcium silicate hydrates”. *Acta Crystallographica Section B* 69.5, p. 465-473. DOI : 10.1107/S2052519213021155. URL : <https://doi.org/10.1107/S2052519213021155> (cf. p. 76).

GRAVES, Alex, MOHAMED, Abdel-rahman & HINTON, Geoffrey (mars 2013). “Speech Recognition with Deep Recurrent Neural Networks”. *ICASSP, IEEE International Conference on Acoustics, Speech and Signal Processing - Proceedings* 38. DOI : 10.1109/ICASSP.2013.6638947 (cf. p. 92).

GRIFFEN, D. T (1999). “Quantitative phase analysis of clay minerals by X-ray powder diffraction using artificial neural networks. I. Feasibility study with calculated powder patterns”. *Clay minerals* 34.1, p. 117-126 (cf. p. 2).

GUPTA, Anup Kumar, KUMAR, Rupesh, BIRLA, Lokendra & GUPTA, Puneet (2023). “Radiant : Better rppg estimation using signal embeddings and transformer”. In : *Proceedings*

of the *IEEE/CVF winter conference on applications of computer vision*, p. 4976-4986 (cf. p. 95).

HALL, Sydney R, ALLEN, Frank H & BROWN, I David (1991). "The crystallographic information file (CIF) : a new standard archive file for crystallography". *Acta Crystallographica Section A : Foundations of Crystallography* 47.6, p. 655-685 (cf. p. 9).

HAYKIN, Simon (2009). *Neural networks and learning machines, 3/E*. Pearson Education India (cf. p. 2).

HE, Kaiming, ZHANG, Xiangyu, REN, Shaoqing & SUN, Jian (2016). "Deep residual learning for image recognition". In : *Proceedings of the IEEE conference on computer vision and pattern recognition*, p. 770-778 (cf. p. 2, 6).

HENDERSON, DM & GUTOWSKY, HS (1962). "A nuclear magnetic resonance determination of the hydrogen positions in Ca (OH) ²⁺". *American Mineralogist : Journal of Earth and Planetary Materials* 47.11-12, p. 1231-1251 (cf. p. 76).

HINTON, Geoffrey E & SALAKHUTDINOV, Ruslan R (2006). "Reducing the dimensionality of data with neural networks". *science* 313.5786, p. 504-507 (cf. p. 30).

INCE, Taner & DOBIGEON, Nicolas (2022). "Weighted Residual NMF With Spatial Regularization for Hyperspectral Unmixing". *IEEE Geoscience and Remote Sensing Letters* 19, p. 1-5. DOI : 10.1109/LGRS.2022.3182042 (cf. p. 51).

JACQUES, S. DM., DI MICHIEL, M., KIMBER, S. AJ., YANG, X., CERNIK, R. J., BEALE, A. M. & BILLINGE, S. JL. (2013). "Pair distribution function computed tomography". *Nature Communications* 4.1, p. 2536 (cf. p. 2).

JAHANKHANI, Pari, KODOGIANNIS, Vassilis & REVETT, Kenneth (2006). "EEG signal classification using wavelet feature extraction and neural networks". In : *IEEE John Vincent Atanasoff 2006 International Symposium on Modern Computing (JVA'06)*. IEEE, p. 120-124 (cf. p. 6).

JENSEN, K. MØ., YANG, X., LAVEDA, J. V., ZEIER, W. G., SEE, K. A., DI MICHIEL, M., MELOT, B. C., CORR, S. A. & BILLINGE, S. JL. (2015). "X-ray diffraction computed tomography for structural analysis of electrode materials in batteries". *Journal of The Electrochemical Society* 162.7, A1310 (cf. p. 2).

JIANG, Wei & KONG, Seong G (2007). "Block-based neural networks for personalized ECG signal classification". *IEEE Transactions on Neural Networks* 18.6, p. 1750-1761 (cf. p. 6).

KINGMA, Diederik P & BA, Jimmy (2014). "Adam : A method for stochastic optimization". *arXiv preprint arXiv :1412.6980* (cf. p. 31, 39, 97).

KRIVOVICHEV, Sergey V, KRIVOVICHEV, Vladimir G, HAZEN, Robert M, AKSENOV, Sergey M, AVDONTCEVA, Margarita S, BANARU, Alexander M, GORELOVA, Liudmila A, ISMAGILOVA, Rezeda M, KORNYAKOV, Ilya V, KUPOREV, Ivan V *et al.* (2022). "Structural and chemical complexity of minerals : an update". *Mineralogical Magazine* 86.2, p. 183-204 (cf. p. 1).

LECUN, Y., BOTTOU, L., BENGIO, Y. & HAFFNER, P. (1998). "Gradient-based learning applied to document recognition". *Proceedings of the IEEE* 86.11, p. 2278-2324. DOI : 10.1109/5.726791 (cf. p. 29).

- LECUN, Yann (1998). "The MNIST database of handwritten digits". *http://yann.lecun.com/exdb/mnist/* (cf. p. 4, 57).
- LEE, J. W., PARK, W. B., KIM, M., SINGH, S. P., PYO, M. & SOHN, K. S. (2021). "A data-driven XRD analysis protocol for phase identification and phase-fraction prediction of multiphase inorganic compounds". *Inorganic Chemistry Frontiers* 8.10, p. 2492-2504 (cf. p. 3, 4, 42).
- LEE, J. W., PARK, W. B., LEE, J. H., SINGH, S. P. & SOHN, K. S. (2020). "A deep-learning technique for phase identification in multiphase inorganic compounds using synthetic XRD powder patterns". *Nature communications* 11.1, p. 1-11 (cf. p. 3, 4, 42).
- LEEM, Saebom & SEO, Hyunseok (2024). "Attention Guided CAM : Visual Explanations of Vision Transformer Guided by Self-Attention". In : *Proceedings of the AAAI Conference on Artificial Intelligence*. **38**. 4, p. 2956-2964 (cf. p. 91).
- LEVIEN, Louise, PREWITT, Charles T & WEIDNER, Donald J (1980). "Structure and elastic properties of quartz at pressure". *American Mineralogist* 65.9-10, p. 920-930 (cf. p. 20, 32, 76, 97).
- LIU, Ze, LIN, Yutong, CAO, Yue, HU, Han, WEI, Yixuan, ZHANG, Zheng, LIN, Stephen & GUO, Baining (2021). "Swin transformer : Hierarchical vision transformer using shifted windows". In : *Proceedings of the IEEE/CVF international conference on computer vision*, p. 10012-10022 (cf. p. 91).
- MARKGRAF, S. A. & REEDER, R. J. (1985). "High-temperature structure refinements of calcite and magnesite". *American Mineralogist* 70.5-6, p. 590-600 (cf. p. 32, 42, 61, 76, 97).
- MEDSKER, Larry R, JAIN, Lakhmi *et al.* (2001). "Recurrent neural networks". *Design and Applications* 5.64-67, p. 2 (cf. p. 92).
- MITCHELL, Tom M & MITCHELL, Tom M (1997). *Machine learning*. **1**. 9. McGraw-hill New York (cf. p. 2).
- MOORE, AE & TAYLOR, HFW (1970a). "Crystal structure of ettringite". *Acta Crystallographica Section B : Structural Crystallography and Crystal Chemistry* 26.4, p. 386-393 (cf. p. 20).
- MOORE, AE & TAYLOR, HFW (1970b). "Crystal structure of ettringite". *Acta Crystallographica Section B : Structural Crystallography and Crystal Chemistry* 26.4, p. 386-393 (cf. p. 76).
- NDLOVU, Bulelwa, BECKER, Megan, FORBES, Elizaveta, DEGLON, David & FRANZIDIS, Jean-Paul (2011). "The influence of phyllosilicate mineralogy on the rheology of mineral slurries". *Minerals engineering* 24.12, p. 1314-1322 (cf. p. 1).
- OVIEDO, F., REN, Z., SUN, S., SETTENS, C., LIU, Z., HARTONO, N. TP., RAMASAMY, S., DECOST, B. L., TIAN, S. IP., ROMANO, G. *et al.* (2019). "Fast and interpretable classification of small X-ray diffraction datasets using data augmentation and deep neural networks". *npj Computational Materials* 5.1, p. 1-9 (cf. p. 3, 27, 32, 42).
- PALSSON, B., SIGURDSSON, J., SVEINSSON, J. R. & ULFARSSON, M. O. (2018). "Hyper-spectral unmixing using a neural network autoencoder". *IEEE Access* 6, p. 25646-25656 (cf. p. 47, 109).

- PARK, W. B., CHUNG, J., JUNG, J., SOHN, K., SINGH, S. P., PYO, M., SHIN, N. & SOHN, K. S. (2017). “Classification of crystal structure using a convolutional neural network”. *IUCrJ* 4.4, p. 486-494 (cf. p. 3, 27, 42).
- POLINE, Victor, PUROHIT, Ravi Raj Purohit Purushottam Raj, BORDET, Pierre, BLANC, Nils & MARTINETTO, Pauline (2024). “Neural networks for rapid phase quantification of cultural heritage X-ray powder diffraction data”. *Journal of Applied Crystallography* 57.Pt 3, p. 831 (cf. p. 3, 42).
- PRASIANAKIS, Nikolaos I (2024). “AI-enhanced X-ray diffraction analysis : towards real-time mineral phase identification and quantification”. *IUCrJ* 11.5, p. 647-648 (cf. p. 8).
- RASTI, B., KOIRALA, B., SCHEUNDERS, P. & GHAMISI, P. (2021). “UnDIP : Hyperspectral unmixing using deep image prior”. *IEEE Transactions on Geoscience and Remote Sensing* 60, p. 1-15 (cf. p. 47).
- RENÉ DE COTRET, Laurent P, OTTO, Martin R, STERN, Mark J & SIWICK, Bradley J (2018). “An open-source software ecosystem for the interactive exploration of ultrafast electron scattering data”. *Advanced Structural and Chemical Imaging* 4.1, p. 1-11 (cf. p. 20).
- RIETVELD, H. M. (1969). “A profile refinement method for nuclear and magnetic structures”. *Journal of applied Crystallography* 2.2, p. 65-71 (cf. p. 1, 62).
- RONNEBERGER, Olaf, FISCHER, Philipp & BROX, Thomas (2015). “U-net : Convolutional networks for biomedical image segmentation”. In : *Medical image computing and computer-assisted intervention–MICCAI 2015 : 18th international conference, Munich, Germany, October 5-9, 2015, proceedings, part III* 18. Springer, p. 234-241 (cf. p. 27, 109).
- ROUARD, Simon, MASSA, Francisco & DÉFOSSEZ, Alexandre (2023). “Hybrid transformers for music source separation”. In : *ICASSP 2023-2023 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*. IEEE, p. 1-5 (cf. p. 91).
- RUMELHART, David E, HINTON, Geoffrey E & WILLIAMS, Ronald J (1986). “Learning representations by back-propagating errors”. *nature* 323.6088, p. 533-536 (cf. p. 28).
- SALGADO, Jerardo E, LERMAN, Samuel, DU, Zhaotong, XU, Chenliang & ABDOLRAHIM, Niaz (2023). “Automated classification of big X-ray diffraction data using deep learning models”. *npj Computational Materials* 9.1, p. 214 (cf. p. 3).
- SCHEIBENREIF, Linus, MOMMERT, Michael & BORTH, Damian (2023). “Masked Vision Transformers for Hyperspectral Image Classification”. In : *2023 IEEE/CVF Conference on Computer Vision and Pattern Recognition Workshops (CVPRW)*, p. 2166-2176. DOI : 10.1109/CVPRW59228.2023.00210 (cf. p. 4, 91, 93, 95).
- SENSOY, M., KAPLAN, L. & KANDEMIR, M. (2018). “Evidential deep learning to quantify classification uncertainty”. *Advances in neural information processing systems* 31 (cf. p. 38).
- SIMONNET, T., FALL, M. D., GALERNE, B., CLARET, F. & GRANGEON, S. (2023). “Proportion Inference Using Deep Neural Networks. Applications to X-Ray Diffraction and Hyperspectral Imaging”. In : *2023 31st European Signal Processing Conference (EUSIPCO)*, p. 1310-1314. DOI : 10.23919/EUSIPC058844.2023.10289954 (cf. p. 8).

SIMONNET, Titouan, FALL, Mame Diarra, GRANGEON, Sylvain & GALERNE, Bruno (2024a). “VISION TRANSFORMERS FOR X-RAY DIFFRACTION PATTERNS ANALYSIS” (cf. p. 8).

SIMONNET, Titouan, GRANGEON, Sylvain, CLARET, Francis, MAUBEC, Nicolas, FALL, Mame Diarra, HARBA, Rachid & GALERNE, Bruno (2024b). “Phase quantification using deep neural network processing of XRD patterns”. *IUCrJ* 11.5 (cf. p. 8).

STEINFINK, H. & SANS, F. J. (1959). “Refinement of the crystal structure of dolomite”. *American Mineralogist* 44, p. 679-682 (cf. p. 32, 42, 61, 97).

SURDU, V. A. & GYÖRGY, R. (2023). “X-ray Diffraction Data Analysis by Machine Learning Methods—A Review”. *Applied Sciences* 13.17, p. 9992 (cf. p. 2).

SUTSKEVER, Ilya, VINYALS, Oriol & LE, Quoc V. (2014). “Sequence to sequence learning with neural networks”. In : *Proceedings of the 27th International Conference on Neural Information Processing Systems - Volume 2*. NIPS’14. Montreal, Canada : MIT Press, p. 3104-3112 (cf. p. 92).

TATLIER, Melkon (2011). “Artificial neural network methods for the prediction of framework crystal structures of zeolites from XRD data”. *Neural Computing and Applications* 20.3, p. 365-371 (cf. p. 3).

UTIMULA, Keishu, YANO, Masao, KIMOTO, Hiroyuki, HONGO, Kenta, NAKANO, Kosuke & MAEZONO, Ryo (mai 2020). “Feature space of XRD patterns constructed by auto-encoder” (cf. p. 109).

VASWANI, A., SHAZEER, N., PARMAR, N., USZKOREIT, J., JONES, L., GOMEZ, A.N., KAISER, Ł. & POLOSUKHIN, I. (2017). “Attention is all you need”. *Advances in neural information processing systems* 30 (cf. p. 3, 91, 92, 95).

VECSEI, P. M., CHOO, K., CHANG, J. & NEUPERT, T. (2019). “Neural network based classification of crystal symmetries from x-ray diffraction patterns”. *Physical Review B* 99.24, p. 245120 (cf. p. 3, 27).

VIDAL, Olivier, GOFFÉ, Bruno & ARNDT, Nicholas (2013). “Metals for a low-carbon society”. *Nature Geoscience* 6.11, p. 894-896 (cf. p. 1).

WALKER, David, VERMA, Pramod K, CRANSWICK, Lachlan MD, JONES, Raymond L, CLARK, Simon M & BUHRE, Stephan (2004). “Halite-sylvite thermoelasticity”. *American Mineralogist* 89.1, p. 204-210 (cf. p. 32, 97).

WAN, L., CHEN, T., PLAZA, A. & CAI, H. (2021). “Hyperspectral Unmixing Based on Spectral and Sparse Deep Convolutional Neural Networks”. *IEEE Journal of Selected Topics in Applied Earth Observations and Remote Sensing* 14, p. 11669-11682. DOI : 10.1109/JSTARS.2021.3126755 (cf. p. 4, 7, 51).

WANG, Hong, XIE, Yunchao, LI, Dawei, DENG, Heng, ZHAO, Yunxin, XIN, Ming & LIN, Jian (2020). “Rapid identification of X-ray diffraction patterns based on very limited data by interpretable convolutional neural networks”. *Journal of chemical information and modeling* 60.4, p. 2004-2011 (cf. p. 3, 42).

WANG, Wenhai, XIE, Enze, LI, Xiang, FAN, Deng-Ping, SONG, Kaitao, LIANG, Ding, LU, Tong, LUO, Ping & SHAO, Ling (2021). “Pyramid vision transformer : A versatile

- backbone for dense prediction without convolutions”. In : *Proceedings of the IEEE/CVF international conference on computer vision*, p. 568-578 (cf. p. 91).
- WERBOS, P.J. (1990). “Backpropagation through time : what it does and how to do it”. *Proceedings of the IEEE* 78.10, p. 1550-1560. DOI : 10.1109/5.58337 (cf. p. 31).
- WOOD, BJ & STRENS, RGJ (1979). “Diffuse reflectance spectra and optical properties of some sulphides and related minerals”. *Mineralogical Magazine* 43.328, p. 509-518 (cf. p. 1).
- ZALOGA, A., STANOVOV, V., BEZRUKOVA, O., DUBININ, P. & YAKIMOV, I. (2020). “Crystal symmetry classification from powder X-ray diffraction patterns using a convolutional neural network”. *Materials Today Communications* 25, p. 101662 (cf. p. 3, 27).
- ZHANG, X., SUN, Y., ZHANG, J., WU, P. & JIAO, L. (2018). “Hyperspectral unmixing via deep convolutional neural networks”. *IEEE Geoscience and Remote Sensing Letters* 15.11, p. 1755-1759 (cf. p. 4, 7, 48).
- ZHAO, Min, SHI, Shuaikai, CHEN, Jie & DOBIGEON, Nicolas (2022). “A 3-D-CNN framework for hyperspectral unmixing with spectral variability”. *IEEE Transactions on Geoscience and Remote Sensing* 60, p. 1-14 (cf. p. 51, 110).
- ZHU, Feiyun (2017). “Hyperspectral unmixing : ground truth labeling, datasets, benchmark performances and survey”. *arXiv preprint arXiv :1708.05125* (cf. p. 47, 48, 51).

Annexe A

Tableaux des amplitudes de variations pour les simulations

TABLE A.1 – Tableau des variations des paramètres pour les simulations avec l’algorithme 1. Ce tableau concerne les travaux concernant l’identification de phases minérales (section 3.2), la quantification sur données synthétiques (section 4.2 et sur données de laboratoire (chapitre 5), en limitant les variations de a , b et c à 2%. Les valeurs des simulations sont tirées de manière uniforme dans les intervalles donnés.

Phase minérale	a	b	c	LMH	Debye-Waller
Calcite	$4.988 \pm 5\%$	$4.988 \pm 5\%$	$17.061 \pm 5\%$	0.05-0.5	0.1-2.0
Dolomite	$4.815 \pm 5\%$	$4.815 \pm 5\%$	$16.119 \pm 5\%$	0.05-0.5	0.1-2.0
Gibbsite	$8.742 \pm 5\%$	$5.112 \pm 5\%$	$9.801 \pm 5\%$	0.05-0.5	0.1-2.0
Halite	$5.6401 \pm 5\%$	$5.6401 \pm 5\%$	$5.6401 \pm 5\%$	0.05-0.5	0.1-2.0
Hématite	$5.038 \pm 5\%$	$5.038 \pm 5\%$	$13.772 \pm 5\%$	0.05-0.5	0.1-2.0
Quartz	$4.916 \pm 5\%$	$4.916 \pm 5\%$	$5.4054 \pm 5\%$	0.05-0.5	0.1-2.0

TABLE A.2 – Tableau des variations des paramètres pour les simulations avec l’algorithme 2. Ce tableau concerne l’analyse du jeu de données obtenu par tomographie de diffraction de rayons X (chapitre 6). Les valeurs des simulations sont tirées de manière uniforme dans les intervalles donnés.

Phase minérale	a	b	c	N_1	N_2	N_3	Debye-Waller
Alite	$11.6389 \pm 2\%$	$14.1716 \pm 2\%$	$13.6434 \pm 2\%$	25-51	25-51	25-51	0.65-0.75
Calcite	$4.988 \pm 2\%$	$4.988 \pm 2\%$	$17.061 \pm 2\%$	25-51	25-51	25-51	0.65-0.75
C-S-H	$6.7320 \pm 2\%$	$7.3690 \pm 2\%$	$11.3400 \pm 2\%$	2-18	2-18	1-5	0.65-0.75
Ettringite	$11.26 \pm 2\%$	$11.26 \pm 2\%$	$21.48 \pm 2\%$	25-51	25-51	25-51	0.65-0.75
Hydrotalcite	$3.054 \pm 2\%$	$3.054 \pm 2\%$	$22.81 \pm 2\%$	25-51	25-51	25-51	0.65-0.75
Monosulfoaluminate	$5.7586 \pm 2\%$	$5.7586 \pm 2\%$	$26.7946 \pm 2\%$	25-51	25-51	25-51	0.65-0.75
Mullite	$7.5785 \pm 2\%$	$7.6817 \pm 2\%$	$2.8864 \pm 2\%$	25-51	25-51	25-51	0.65-0.75
Portlandite	$3.5925 \pm 2\%$	$3.5925 \pm 2\%$	$4.905 \pm 2\%$	25-51	25-51	25-51	0.65-0.75
Quartz	$4.916 \pm 2\%$	$4.916 \pm 2\%$	$5.4054 \pm 2\%$	25-51	25-51	25-51	0.65-0.75
Amorphe (C-S-H)	$6.7320 \pm 2\%$	$7.3690 \pm 2\%$	$11.3400 \pm 2\%$	3	3	1	0.65-0.75

Apprentissage et réseaux de neurones en tomographie par diffraction de rayons X. Application à l'identification minéralogique

La compréhension du comportement chimique et mécanique des matériaux compactés (par exemple sol, sous-sol, matériaux ouvragés) nécessite de se baser sur une description quantitative de structuration du matériau, et en particulier de la nature des différentes phases minéralogiques et de leur relation spatiale. Or, les matériaux naturels sont composés de nombreux minéraux de petite taille, fréquemment mixés à petite échelle. Les avancées récentes en tomographie de diffraction des rayons X sur source synchrotron (à différencier de la tomographie en contraste de phase) permettent maintenant d'obtenir des volumes tomographiques avec des voxels de taille nanométrique, avec un diffractogramme pour chacun de ces voxels (là où le contraste de phase ne donne qu'un niveau de gris). En contrepartie, le volume de données (typiquement de l'ordre de 100 000 diffractogrammes par tranche d'échantillon), associé au grand nombre de phases présentes, rend le traitement quantitatif virtuellement impossible sans codes numériques appropriés.

Cette thèse vise à combler ce manque, en utilisant des approches de type réseaux de neurones pour identifier et quantifier des minéraux dans un matériau. L'entraînement de tels modèles nécessite la construction de bases d'apprentissage de grande taille, qui ne peuvent pas être constituées uniquement de données expérimentales. Des algorithmes capables de synthétiser des diffractogrammes pour générer ces bases ont donc été développés. L'originalité de ce travail a également porté sur l'inférence de proportions avec des réseaux de neurones. Pour répondre à cette tâche, nouvelle et complexe, des fonctions de perte adaptées ont été conçues. Le potentiel des réseaux de neurones a été testé sur des données de complexités croissantes : (i) à partir de diffractogrammes calculés à partir des informations cristallographiques, (ii) en utilisant des diffractogrammes expérimentaux de poudre mesurés au laboratoire, (iii) sur les données obtenues par tomographie de rayons X.

Différentes architectures de réseaux de neurones ont aussi été testées. Si un réseau de neurones convolutifs semble apporter des résultats intéressants, la structure particulière du signal de diffraction (qui n'est pas invariant par translation) a conduit à l'utilisation de modèles comme les Transformers.

L'approche adoptée dans cette thèse a démontré sa capacité à quantifier les phases minérales dans un solide. Pour les données les plus complexes, tomographie notamment, des pistes d'amélioration ont été proposées.

Mots clés : Réseaux de neurones, Diffraction par rayons X, Inférence de proportions, Quantification de phases minérales, Modélisation de Dirichlet, Transformers

Neural networks for X-ray diffraction computed tomography. Application to Mineralogical Identification

Understanding the chemical and mechanical behavior of compacted materials (e.g. soil, subsoil, engineered materials) requires a quantitative description of the material's structure, and in particular the nature of the various mineralogical phases and their spatial relationships. Natural materials, however, are composed of numerous small-sized minerals, frequently mixed on a small scale. Recent advances in synchrotron-based X-ray diffraction tomography (to be distinguished from phase contrast tomography) now make it possible to obtain tomographic volumes with nanometer-sized voxels, with a XRD pattern for each of these voxels (where phase contrast only gives a gray level). On the other hand, the sheer volume of data (typically on the order of 100 000 XRD patterns per sample slice), combined with the large number of phases present, makes quantitative processing virtually impossible without appropriate numerical codes.

This thesis aims to fill this gap, using neural network approaches to identify and quantify minerals in a material. Training such models requires the construction of large-scale learning bases, which cannot be made up of experimental data alone. Algorithms capable of synthesizing XRD patterns to generate these bases have therefore been developed. The originality of this work also concerned the inference of proportions using neural networks. To meet this new and complex task, adapted loss functions were designed. The potential of neural networks was tested on data of increasing complexity : (i) from XRD patterns calculated from crystallographic information, (ii) using experimental powder XRD patterns measured in the laboratory, (iii) on data obtained by X-ray tomography.

Different neural network architectures were also tested. While a convolutional neural network seemed to provide interesting results, the particular structure of the diffraction signal (which is not translation invariant) led to the use of models such as Transformers.

The approach adopted in this thesis has demonstrated its ability to quantify mineral phases in a solid. For more complex data, such as tomography, improvements have been proposed.

Keywords : Neural Network, X-ray diffraction, Proportion inference, Phases quantification, Dirichlet Modeling, Transformers