



Résolution des équations de Maxwell avec des éléments finis de Galerkin continus

Erell Jamelot

► To cite this version:

Erell Jamelot. Résolution des équations de Maxwell avec des éléments finis de Galerkin continus. Mathématiques [math]. Ecole Polytechnique X, 2005. Français. NNT: . tel-00440043

HAL Id: tel-00440043

<https://pastel.hal.science/tel-00440043>

Submitted on 9 Dec 2009

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

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



Thèse de Doctorat de l'École Polytechnique

Domaine :

Mathématiques et Informatique

Spécialité :

Mathématiques de la Modélisation, Simulation et Applications de la Physique

Présentée par :

Erell JAMELOT

Pour obtenir le titre de :

DOCTEUR DE L'ÉCOLE POLYTECHNIQUE

**RÉSOLUTION DES ÉQUATIONS DE MAXWELL AVEC DES
ÉLÉMENTS FINIS DE GALERKIN CONTINUS**

Thèse déposée le Jeudi 8 Septembre 2005 - Soutenue le Jeudi 17 Novembre 2005
Devant le jury composé de :

M.	Franck Assous	Examineur, Professeur, Université de Bar Ilan, Israël
Mme	Christine Bernardi	Président, Directeur de Recherche au CNRS, Université Paris VI
Mme	Annalisa Buffa	Rapporteur, Directeur de Recherche au CNR, IMATI, Pavia, Italie
M.	Patrick Ciarlet	Directeur de thèse, Enseignant-Chercheur, ENSTA, Paris
M.	Martin Costabel	Examineur, Professeur, Université de Rennes I
M.	Éric Sonnendrücker	Rapporteur, Professeur, Université Louis Pasteur, Strasbourg

Thèse réalisée au Laboratoire de Mathématiques Appliquées de l'ENSTA.



Table des matières

Remerciements	11
Introduction	13
I Modélisation	17
1 Modélisation et méthodes de résolution	19
1.1 Introduction	19
1.2 L'électrodynamique classique	19
1.3 Les équations de Maxwell	20
1.3.1 Champs et sources	20
1.3.2 Lois de comportement	21
1.3.3 Énergie électromagnétique	22
1.3.4 Conditions de transmission entre deux milieux matériels	22
1.3.5 Condition aux limites en domaine borné	23
1.3.6 Condition de radiation en domaine non-borné	23
1.3.7 Effet de pointe	24
1.4 Hypothèses mathématiques sur les données	24
1.5 Différents modèles	26
1.5.1 Le modèle de Darwin	26
1.5.2 Les modèles quasi-électrostatique et quasi-magnétostatique	27
1.5.3 Les équations harmoniques	28
1.5.4 Contrôlabilité	29
1.5.5 Équations instationnaires	29
1.6 Méthodes de résolution	30
1.6.1 Différences finies	30
1.6.2 Éléments finis de Nédélec	31
1.6.3 Éléments finis discontinus	32
1.6.4 Éléments finis hp	32
1.6.5 Méthodes spectrales	32
1.6.6 Obtention d'un domaine borné	33
2 Intérêt de l'étude du problème bidimensionnel	35
2.1 Le guide d'onde	35
2.2 La ligne microruban	36
2.3 Le filtre à stubs	37

II	Le problème bidimensionnel	39
1	Notations et résultats préliminaires 2D	41
1.1	Notations relatives au domaine d'étude	41
1.2	Opérateurs bidimensionnels	42
1.2.1	Notations	42
1.2.2	Relations entre opérateurs bidimensionnels	43
1.3	Espaces de Hilbert usuels et leurs normes associées	43
1.3.1	Espaces des champs scalaires	44
1.3.2	Espaces de champs vectoriels	44
1.3.3	Espaces des traces	45
1.4	Espaces de Sobolev à poids	46
1.5	Espaces de Sobolev classiques	46
1.6	Espaces duaux	47
1.7	Formules d'intégration par parties classiques dans un ouvert	47
1.7.1	Formules de Green	47
1.7.2	Généralisation des formules de Green	47
2	Le problème statique 2D direct continu	49
2.1	Introduction	49
2.2	Champ électrique 2D : décomposition de Helmholtz	51
2.2.1	CL non-homogènes	51
2.2.2	CL homogènes	52
2.2.3	Le problème aux potentiels, à la <i>Grisvard</i>	53
2.2.4	Le problème aux potentiels, à la <i>Nazarov-Plamenevsky</i>	58
2.3	Champ électrique 2D : CL naturelles	59
2.4	Champ électrique 2D : le complément singulier	60
2.4.1	Introduction	60
2.4.2	La λ -approche	61
2.4.3	Simplification du calcul de λ	63
2.4.4	Le complément singulier orthogonal	67
2.4.5	Autres décompositions conformes	72
2.4.6	Régularité dans les espaces singuliers conformes	75
2.4.7	Conclusions sur la méthode du complément singulier	77
2.5	Champ électrique 2D : la régularisation à poids	77
2.5.1	Condition pour obtenir la coercivité	77
2.5.2	Condition pour obtenir la densité des éléments finis	79
2.5.3	Choix du poids	80
2.5.4	Cas où il existe plusieurs coins rentrants	81
2.6	Milieux inhomogènes	81
2.7	Conclusion	82
3	Le problème statique 2D mixte continu	85
3.1	Rappel sur les multiplicateurs de Lagrange	85
3.2	Formulation mixte 2D : CL naturelles	86
3.3	Formulation mixte 2D : MCSO	87
3.4	Formulation mixte 2D : régularisation à poids	87

4	Le problème statique 2D direct discret	89
4.1	Discrétisation par les éléments finis de Galerkin continus	89
4.1.1	L'approximation de Galerkin	89
4.1.2	Les éléments finis de Lagrange continus d'ordre k	90
4.2	Discrétisation du domaine d'étude	91
4.3	Décomposition de Helmholtz : discrétisation	92
4.3.1	Le Laplacien avec CL de Dirichlet	92
4.3.2	Le Laplacien avec CL de Neumann	94
4.3.3	Dérivation du champ électrique	96
4.4	Calcul des singularités du Laplacien	96
4.4.1	Singularités duales : CL de Dirichlet	97
4.4.2	Singularités duales : CL de Neumann	98
4.4.3	Coefficients des singularités duales	99
4.4.4	Coefficients des singularités primales	99
4.5	Calcul des singularités électromagnétiques	100
4.5.1	Singularités primales : CL de Dirichlet	100
4.5.2	Singularités primales : CL de Neumann	100
4.6	Méthode avec CL naturelles : discrétisation	101
4.6.1	Matrice de raideur interne	103
4.6.2	Matrice de masse du bord	104
4.6.3	Second membre	104
4.7	Élimination des CL essentielles	105
4.7.1	Discrétisation de l'espace avec CL essentielles	105
4.7.2	Matrice de raideur interne	106
4.7.3	Second membre	108
4.8	λ -approche : discrétisation	109
4.8.1	Champ électrique : partie régulière	109
4.8.2	Relèvement de la CL	113
4.9	Complément singulier orthogonal : discrétisation	114
4.9.1	Champ électrique : partie régulière	114
4.9.2	Champ électrique : partie singulière	115
4.9.3	Conclusion	116
4.10	Régularisation à poids : discrétisation	116
4.11	Analyse d'erreur	117
4.11.1	Méthode avec CL naturelles : convergence	117
4.11.2	Méthode du complément singulier : convergence	117
4.11.3	Régularisation à poids : convergence	121
4.11.4	Conclusion	122
5	Le problème statique 2D mixte discret	123
5.1	L'élément fini mixte de Taylor-Hood P_2 - P_1	123
5.2	Condition inf-sup discrète	124
5.2.1	Méthode avec CL naturelles	124
5.2.2	Méthode avec CL essentielles	124
5.2.3	Régularisation à poids	124
5.3	Méthode avec CL naturelles mixte : discrétisation	125
5.3.1	Matrice mixte	126
5.3.2	Second membre	126

5.4	Complément singulier orthogonal mixte : discrétisation	126
5.4.1	Matrices mixtes	128
5.4.2	Second membre	128
5.5	Régularisation à poids mixte : discrétisation	129
5.5.1	Matrices mixtes	130
5.5.2	Second membre	130
5.6	Analyse d'erreur	131
5.7	Optimalité de l'algorithme d'Uzawa	133
6	Résultats numériques du problème statique 2D	135
6.1	Introduction	135
6.2	Calcul direct	136
6.2.1	Cas régulier	136
6.2.2	Premier cas singulier	138
6.2.3	Second cas singulier	141
6.3	Utilisation du multiplicateur de Lagrange	142
6.4	Allure du champ quasi-électrostatique	145
6.5	Conclusions	149
III	Le problème tridimensionnel	151
7	Notations et résultats préliminaires 3D	153
7.1	Notations relatives au domaine d'étude	153
7.2	Opérateurs tridimensionnels	154
7.3	Espaces de Hilbert usuels et leurs normes associées	155
7.3.1	Espaces de champs scalaires	156
7.3.2	Espaces de champs vectoriels	157
7.3.3	Espaces des traces	158
7.4	Espaces à poids	159
7.5	Formules d'intégration par parties classiques dans un ouvert	160
7.5.1	Formules de Green	160
7.5.2	Généralisation des formules de Green	160
7.6	Espaces fonctionnels du problème en temps	161
8	Le problème statique 3D direct continu	163
8.1	Introduction	163
8.2	Champ électrostatique 3D	166
8.3	Champ électrique 3D : CL naturelles	166
8.4	Champ électrique 3D : CL essentielles	168
8.5	Champ électrique 3D : régularisation à poids	168
8.6	Champ électrique $2D\frac{1}{2}$: le complément singulier	170
8.6.1	Introduction	170
8.6.2	Cas prismatique : CL essentielles	174
8.6.3	Cas prismatique : CL presque essentielles	175

9	Le problème statique 3D mixte continu	177
9.1	Formulation mixte 3D : CL naturelles	177
9.2	Formulation mixte 3D : CL essentielles	178
9.3	Formulation mixte 3D : régularisation à poids	178
9.4	Formulation mixte $2D\frac{1}{2}$: le complément singulier	179
10	Le problème statique 3D direct discret	181
10.1	Discrétisation du domaine d'étude	181
10.2	Électrostatique : discrétisation	182
10.2.1	Le Laplacien avec CL de Dirichlet	182
10.2.2	Dérivation du champ électrique	184
10.3	Méthode avec CL naturelles : discrétisation	184
10.3.1	Matrice de raideur interne	187
10.3.2	Matrice de masse du bord	187
10.3.3	Second membre	188
10.4	Méthode avec CL essentielles : discrétisation	189
10.4.1	Élimination des conditions aux limites essentielles	189
10.4.2	Champ électrique	193
10.4.3	Relèvement de la CL	195
10.5	Régularisation à poids : discrétisation	195
10.6	Cas prismatique : discrétisation	196
10.6.1	Champ électrique transverse : CL essentielles	196
10.6.2	Champ électrique transverse : CL presque essentielles	197
10.6.3	Champ électrique longitudinal	200
11	Le problème statique 3D mixte discret	201
11.1	L'élément fini mixte de Taylor-Hood P_2 - P_1	201
11.2	Méthode avec CL naturelles mixte : discrétisation	202
11.3	Méthode avec CL essentielles mixte : discrétisation	203
11.4	Régularisation à poids mixte : discrétisation	204
11.5	Cas prismatique mixte : discrétisation	206
11.5.1	CL essentielles	206
11.5.2	CL presque essentielles	207
12	Le problème temporel 3D	209
12.1	Introduction	209
12.2	Champ électrique : formulations variationnelles	211
12.2.1	Introduction	211
12.2.2	Formulation variationnelle classique	212
12.2.3	Formulation variationnelle augmentée	214
12.2.4	Formulation variationnelle augmentée mixte	215
12.2.5	Existence d'une frontière artificielle	218
12.3	Champ magnétique : formulation variationnelle	219
12.4	Semi-discrétisation en temps	220
12.5	Discrétisation complète	221
12.5.1	Élimination des CL presque essentielles	221
12.5.2	Formation des second-membres	223
12.5.3	Champ électromagnétique	223

12.5.4 Étude de la stabilité	225
12.6 Présentation du code de calcul	227
13 Problème temporel 3D : résultats numériques	229
13.1 Méthode avec CL essentielles	229
13.2 Méthode de régularisation à poids	231
13.2.1 Introduction	231
13.2.2 Évolution spatiale	232
13.2.3 Évolution temporelle	237
13.3 Conclusions sur les méthodes utilisés	237
Conclusions et perspectives	239
IV Annexe	241
14 Calculs complémentaires pour la MCS	243
14.1 Calcul de β_D et β_N	243
14.1.1 Cas d'un unique coin rentrant	243
14.1.2 Cas de plusieurs coins rentrants	244
14.2 Calcul de λ_D et λ_N	246
14.3 Simplification du calcul de λ	246
14.3.1 Preuve du lemme 2.37	246
14.3.2 Preuve du lemme 2.38	247
14.4 Preuve du lemme 4.16	247
15 Calculs du problème discrétisé	251
15.1 Éléments finis P_k 2D	251
15.1.1 Éléments finis P_0	251
15.1.2 Éléments finis P_1	251
15.1.3 Éléments finis P_2	252
15.2 Intégration numérique 2D	252
15.2.1 Schémas d'intégration numérique intérieure	252
15.2.2 Schémas d'intégration numérique sur la frontière	253
15.3 Éléments finis P_k 3D	254
15.3.1 Éléments finis P_0	254
15.3.2 Éléments finis P_1	254
15.3.3 Éléments finis P_2	255
15.3.4 Éléments finis $\tilde{P}_2$	255
15.3.5 Schémas d'intégration numérique 3D	256
15.4 Algorithme du gradient conjugué	257
15.4.1 Gradient conjugué non-préconditionné	258
15.4.2 Gradient conjugué préconditionné	258
15.5 Réduction de la matrice de masse	259
15.5.1 Cas général	260
15.5.2 Triangulation ou tétraèdrisation régulière et quasi-uniforme	261
15.6 Réduction de la matrice de masse pondérée en 2D	262
15.6.1 Triangles intérieurs	262
15.6.2 Triangles touchant le coin rentrant	262

15.6.3	Matrice équivalente	263
15.7	Équivalence entre matrice de masse et matrice mixte	263
16	Le champ magnétique	265
16.1	Le problème quasi-magnétostatique $2D$	265
16.1.1	Introduction	265
16.1.2	Le problème direct	265
16.2	Le problème quasi-magnétostatique $3D$	268
16.2.1	Introduction	268
16.2.2	Le problème direct	268
16.3	Champ quasi-magnétostatique : discrétisation $3D$	270
16.3.1	Méthode avec CL naturelles	270
16.3.2	Méthode avec CL essentielles	270
17	Rappels de notations	279
17.1	Rappel des espaces fonctionnels et des formes bilinéaires $2D$	279
17.2	Rappel des espaces fonctionnels et des formes bilinéaires $3D$	280
	Bibliographie	280

Remerciements

Tout d'abord, je tiens à remercier la Délégation Générale pour l'Armement pour le financement de cette thèse pendant trois ans; ainsi qu'Éric Lunéville, directeur du Laboratoire de Mathématiques Appliquées de l'ENSTA (le LMA) de m'avoir accueillie dans son Laboratoire, au sein duquel j'ai bénéficié de très bonnes conditions de travail.

Merci à Patrick Ciarlet de m'avoir proposé cette thèse, qui m'a ouverte à l'Analyse Numérique. Sa disponibilité, son dynamisme et ses compétences scientifiques ont grandement contribué à la réalisation de ce travail.

J'exprime toute ma reconnaissance envers Annalisa Buffa et Éric Sonnendrücker, qui ont accepté la lourde tâche d'être les rapporteurs de ce long manuscrit. Je remercie vivement Franck Assous, Christine Bernardi et Martin Costabel de m'avoir fait l'honneur et le plaisir de participer au jury.

Je salue chaleureusement les membres du LMA pour l'ambiance de conviviale de travail, entretenue par les "PSAUMES". Je tiens à remercier en particulier Jean-Luc Commeau, Maurice Diamantini et Fabrice Roy qui m'ont régulièrement aidée à régler mes soucis informatiques; ainsi qu'Annie Marchal, notre bienveillante et organisée secrétaire. Je n'oublie pas mes camarades thésards et thésardes pour leur amitié et leur soutien, avec un clin d'oeil spécial à ma "co-bureau" Ève-Marie Duclairoir.

Enfin, je sais gré aux membres de l'atelier de reprographie de l'ENSTA d'avoir imprimé ce document avec soin.

Introduction

Les équations de Maxwell sont résolues aisément lorsque le domaine de calcul est convexe, ou à bord régulier, mais s'il présente des singularités géométriques (coins et/ou arêtes rentrants), le champ électromagnétique est localement intense et très difficile à calculer. On dit qu'il est singulier. Or, dans de nombreuses applications de l'électromagnétisme, comme les guides d'ondes, ou les filtres à stubs utilisés en télécommunication, il existe des singularités géométriques pouvant induire un champ électromagnétique intense. La modélisation et le calcul numérique du champ permettent alors de calculer l'intensité du champ au voisinage des singularités et de détecter les effets délétères. L'enjeu est donc de construire des méthodes numériques qui réussissent à capturer les singularités du champ, et qui soient efficaces en terme de précision et de coût calcul.

Les éléments finis d'arêtes permettent d'approcher les singularités, mais doivent être maniés avec précaution lorsqu'on a besoin d'une approximation continue, quand on résout le système couplé Maxwell-Vlasov par exemple. Par contre, à l'aide d'éléments finis nodaux, il est possible de calculer des approximations continues du champ électromagnétique. Pour les équations quasi-statiques bidimensionnelles, nous présentons principalement l'étude de trois différentes méthodes d'éléments finis nodaux, codées en Matlab. Nous étudions la généralisation de ces méthodes en $3D$, et nous présentons deux méthodes de résolution pour les équations de Maxwell tridimensionnelles instationnaires, codées en Fortran 77.

La première méthode $2D$ est une nouvelle version de la **méthode du complément singulier**, développée par F. Assous et al. dans [9] et dans la thèse d'E. Garcia [63]. Les conditions limites sont traitées de façon essentielle : la condition aux limites de conducteur parfait est prise en compte explicitement. En dimension deux, les singularités du champ électromagnétique sont connues exactement (à un facteur près). On peut séparer le champ en une partie régulière et une partie analytique. Cette méthode, qui montre d'excellents résultats en $2D$ peut s'étendre aux cas de dimension $2D\frac{1}{2}$ [38, 76], ou alors lorsque les seules singularités géométriques tridimensionnelles sont des pointes coniques [63]. Dans les domaines tridimensionnels généraux, les singularités électromagnétiques sont difficilement discrétisables, les singularités de coins et d'arêtes étant liées entre elles.

Pour les deux autres méthodes $2D$, le découplage n'est pas nécessaire car l'espace fonctionnel usuel des solutions est modifié, de façon à retrouver la densité des éléments finis de Lagrange. Ainsi, elles fonctionnent dans les domaines tridimensionnels généraux.

La seconde méthode est la **méthode à poids**, développée par M. Costabel et M. Dauge dans [53] pour le champ électrique. Les conditions limites sont essentielles. L'équation de Maxwell-Gauss est multipliée par un poids qui dépend de la distance aux singularités géométriques.

Enfin, la dernière méthode est la **méthode avec conditions aux limites naturelles**, développée par P. Ciarlet, Jr. dans [36, 40]. On relaxe la condition aux limites de conducteur parfait.

Composition du document

Après une introduction sur les méthodes de modélisation des équations de Maxwell, développée dans la partie I, le document est constitué de deux parties et d'une annexe, constituée par la partie IV. La partie II traite le problème bidimensionnel, et la partie III traite le problème tridimensionnel. Chacune de ces parties contient tout d'abord une étude du problème statique direct, puis du problème statique mixte (ajout d'un multiplicateur de Lagrange). Pour le problème tridimensionnel, on fait de plus une étude du problème instationnaire. Nous proposons et analysons différentes méthodes, en partant du problème continu pour aboutir à la discrétisation et la mise oeuvre numérique. Enfin, nous présentons des résultats numériques que nous confrontons aux tentatives théoriques.

Modélisation

Dans cette partie, on présente dans un premier chapitre différents problèmes liés à l'électromagnétisme et quelques méthodes de résolution. Dans un second chapitre, on détaille les intérêts physique et mathématique de l'étude du problème bidimensionnel.

La partie 2D

La partie 2D s'étend des chapitres un à six. Le problème quasi-statique consiste à manipuler les équations de Maxwell sans se soucier de la dépendance en temps. Dans le cas bidimensionnel, les problèmes quasi-électrostatique et quasi-magnétostatique sont similaires, aussi nous ne détaillons que le problème quasi-électrostatique.

Dans le premier chapitre, nous définissons les notations et les espaces fonctionnels dont nous avons besoin.

Dans le second chapitre, nous décrivons en détail le problème quasi-électrostatique et nous présentons quatre méthodes de résolution par éléments finis nodaux :

- La méthode aux potentiels : le calcul du champ électrique est indirect, dérivé des potentiels électrostatiques. Cette méthode sert de référence dans certains cas tests.
- La méthode aux conditions aux limites naturelles, pour laquelle la condition aux limites est incluse dans la formulation variationnelle.
- La méthode du complément singulier : nous présentons une nouvelle décomposition non conforme pour le cas statique, la λ -approche, ainsi que la décomposition orthogonale usuelle. Nous montrons comment utiliser la λ -approche pour le calcul des fonctions de base singulières de la décomposition orthogonale.
- La méthode de régularisation à poids, développée par M. Costabel et M. Dauge dans [53].

Nous montrons que les formulations variationnelles sont bien posées pour les espaces fonctionnels choisis.

Dans le troisième chapitre, nous introduisons un multiplicateur de Lagrange sur la divergence (étude pour les méthodes directes), afin de préparer la résolution de problèmes transitoires.

Le quatrième chapitre est consacré à la résolution numérique du problème quasi-électrostatique par les éléments finis de Lagrange P_k . Dans le cinquième chapitre, nous détaillons la résolution numérique du problème mixte par les éléments finis de Taylor-Hood P_{k+1} - P_k .

Enfin, dans le sixième chapitre, nous présentons et analysons les résultats numériques obtenus avec le code Matlab.

La partie 3D

La partie 3D est décomposée en huit chapitres, numérotés de sept à quatorze.

Dans le chapitre sept, nous définissons les notations, les espaces fonctionnels utilisés.

Pour le problème statique, détaillé dans le chapitre huit, nous étudions cinq méthodes de résolution par éléments finis nodaux :

- La méthode aux potentiels, valable lorsque le champ électrique est à rotationnel nul. Comme en 2D, cette méthode consiste à dériver le champ électrique du potentiel électrostatique.
- La méthode avec conditions aux limites naturelles.
- La méthode avec conditions aux limites essentielles, valable lorsque le domaine de calcul est convexe [69]. Ceci prépare la méthode suivante.
- La méthode de régularisation à poids [53].
- La méthode du complément singulier en domaine prismatique, présentée dans [38].

Dans le chapitre neuf, on donne les formulations mixtes des quatre dernières méthodes. Les chapitres dix et onze sont consacrés à la résolution numérique des problèmes directs et mixtes, ce qui prépare la discrétisation en espace du problème instationnaire.

Le problème dépendant du temps est traité dans le douzième chapitre : à l'aide de bonnes hypothèses sur les données, on montre que la formulation mixte augmentée, pour le champ électrique, est équivalente au système des équations de Maxwell et admet une unique solution. Ensuite, on discrétise le problème avec un schéma aux différences finies pour la discrétisation en temps, et les éléments finis de Taylor-Hood P_{k+1} - P_k pour la discrétisation spatiale. Une étude de stabilité met en exergue une condition de type CFL à respecter entre le pas de temps et les caractéristiques du maillage. Dans le treizième chapitre, nous présentons les résultats obtenus en 3D.

Première partie

Modélisation

Chapitre 1

Modélisation et méthodes de résolution

1.1 Introduction

L'objectif de ce chapitre est d'une part de faire quelques rappels de Physique sur les équations de Maxwell (sections 1.2 et 1.3), et de justifier ainsi l'étude mathématique (qui semble parfois lointaine de la Physique) qui va suivre.

D'autre part, on donne les hypothèses mathématiques requises (section 1.4), on présente quelques modèles de simplification de ces équations (section 1.5) chers aux mathématiciens et numériciens, puis on indique différentes méthodes (non exhaustives) de résolution (section 1.6).

Pour la partie physique, on se réfère essentiellement aux ouvrages classiques suivants : les cours de R. P. Feynman [62], pédagogiques et ludiques, le livre de J. D. Jackson [71], précis et complet. Le livre d'A. Bossavit [25] permet (avec quelques connaissances mathématiques) de passer de la Physique à la modélisation.

1.2 L'électrodynamique classique

Bien que les phénomènes électromagnétiques soient connus depuis l'Antiquité, les premières expériences sur l'électricité et le magnétisme remontent seulement au XVII^{ème} siècle. L'analyse scientifique de ces phénomènes commencèrent avec les travaux de Coulomb sur l'électrisation, qui furent publiés en 1785, et qui conduisirent à la théorie dynamique du champ électromagnétique de Maxwell, publiée en 1864. Cette théorie fut validée en 1888 par Hertz, qui avait découvert des ondes électromagnétiques se propageant à la vitesse de la lumière.

Depuis les années soixante, notre compréhension sur les constituants fondamentaux de la matière et les forces qui interagissent entre eux a révolutionné la Physique. Cela a donné lieu à la formulation du *modèle standard* des particules physiques, qui décrit les particules et leurs interactions. L'électrodynamique classique héritée de Maxwell est une forme limite de l'électrodynamique quantique contenue dans le modèle standard, c'est-à-dire valide lorsque le nombre de photons impliqués est suffisamment grand.

L'électrodynamique classique permet de décrire les phénomènes électromagnétiques qui se manifestent dans de nombreuses technologies modernes : télécommunication, micro-onde, radar, antenne. Afin de simuler les effets produits, qui peuvent être destructeurs, il est nécessaire de faire une analyse mathématique des équations de Maxwell, pour les résoudre numériquement, la résolution analytique étant dans la majorité des cas actuellement hors de notre portée.

1.3 Les équations de Maxwell

1.3.1 Champs et sources

Dans les milieux continus, les phénomènes électromagnétiques sont décrits par quatre fonctions qui dépendent du temps t et des coordonnées d'espace $\mathbf{x}$, à valeurs dans $\mathbb{R}^3$:

- le champ électrique $\mathcal{E}$, qui est de la dimension d'une force par unité de charge ou $\text{V} \cdot \text{m}^{-1}$ (Volts par mètre),
- l'induction magnétique $\mathcal{B}$, qui est de la dimension d'une force par unité de courant ou T (Tesla),
- le champ magnétique $\mathcal{H}$, en $\text{A} \cdot \text{m}^{-1}$ (Ampères par mètre),
- le déplacement électrique $\mathcal{D}$, en $\text{C} \cdot \text{m}^{-2}$ (Coulombs par mètre carré).

La force agissant sur une charge ponctuelle q en présence d'un champ électromagnétique est décrite par l'équation de la force de Lorentz :

$$\mathcal{F} = q(\mathcal{E} + \mathbf{v} \times \mathcal{B}). \quad (1.1)$$

Le champ électrique $\mathcal{E}$ et l'induction magnétique $\mathcal{B}$ furent initialement introduits à partir de l'équation de force de Lorentz. Cette équation permet de décrire le mouvement d'une particule chargée.

Les fonctions électromagnétiques sont régies par les équations de Maxwell (système d'unité SI) :

$$\text{div } \mathcal{D} = \rho, \text{ équation de Maxwell-Gauss,} \quad (1.2)$$

$$\text{rot } \mathcal{H} - \partial_t \mathcal{D} = \mathcal{J}, \text{ équation de Maxwell-Ampère,} \quad (1.3)$$

$$\text{rot } \mathcal{E} + \partial_t \mathcal{B} = 0, \text{ équation de Maxwell-Faraday,} \quad (1.4)$$

$$\text{div } \mathcal{B} = 0, \text{ absence de monopôle magnétique.} \quad (1.5)$$

Les équations (1.3) et (1.4) sont des équations d'évolution, alors que les équations (1.2) et (1.5) sont des équations de contrainte. Bien que nous présentons ces équations d'un seul bloc, elles ont été élaborées pas à pas, par plusieurs physiciens.

ρ (en $\text{C} \cdot \text{m}^{-3}$, à valeurs dans $\mathbb{R}$) est la densité volumique de charges électriques dans le milieu, $\mathcal{J}$ (en $\text{A} \cdot \text{m}^{-2}$ à valeurs dans $\mathbb{R}^3$) est la densité de courant, qui est non nulle dès qu'il y a un courant électrique.

Le champ électromagnétique peut exister dans des régions dépourvues de sources, lorsque $\mathcal{J} = 0$ et $\rho = 0$. Son existence est indépendante des charges et du courant, puisqu'il transporte de l'énergie, de la quantité de mouvement et du moment cinétique.

Les équations (1.2)-(1.5) contiennent implicitement l'équation de continuité de la charge, reliant les densités de charges ρ et de courant $\mathcal{J}$; obtenue en combinant la dérivée temporelle de (1.2) et la divergence de (1.3) :

$$\partial_t \rho + \text{div } \mathcal{J} = 0. \quad (1.6)$$

On peut remarquer une certaine redondance dans ce système d'équations. En effet, si on applique l'opérateur divergence à (1.4), on obtient que $\partial_t(\text{div } \mathcal{B}) = 0$. Ainsi, il suffit que (1.5) soit vérifiée à un seul instant (par exemple à $t = 0$) pour qu'elle le soit à tout temps t . De même, si on applique l'opérateur divergence à (1.3), on obtient que $\partial_t(\text{div } \mathcal{E}) = \partial_t \rho$, à condition toutefois que (1.6) soit vérifiée.

Lorsque $\mathcal{J}$ et ρ sont connus, et satisfont (1.6), le système (1.2)-(1.5) comporte six équations scalaires indépendantes, et les équations (1.2) et (1.5) jouent le rôle de conditions initiales. Comme on a douze inconnues scalaires, on sera amené à ajouter à ces équations des relations, appelées *lois de comportement* ou relations constitutives, qui permettent de décrire la nature du milieu dans lequel ont lieu les phénomènes électromagnétiques.

1.3.2 Lois de comportement

Lorsque le milieu est conducteur, la loi d'Ohm est vérifiée :

$$\mathcal{J}(\mathbf{x}, t) = \sigma(\mathbf{x}) \mathcal{E}(\mathbf{x}, t), \text{ en milieu isotrope,} \quad (1.7)$$

$$\mathcal{J}_i(\mathbf{x}, t) = \sum_{j=1}^3 \sigma_{ij}(\mathbf{x}) \mathcal{E}_j(\mathbf{x}, t), \quad i = 1, 3, \text{ en milieu anisotrope.} \quad (1.8)$$

σ est la conductivité électrique. En milieu anisotrope, c'est un tenseur.

Lorsque le milieu est isolant, σ est nulle, donc $\mathcal{J} = 0$: il n'y a pas de courant circulant dans le milieu. Dans le cas d'un conducteur parfait, σ est infinie : les champs $\mathcal{E}$ et $\mathcal{H}$ sont nuls.

Dans les milieux parfaits, c'est-à-dire les milieux pour lesquels les lois de comportement sont linéaires, les relations suivantes sont vérifiées :

$$\mathcal{D}(\mathbf{x}, t) = \varepsilon(\mathbf{x}) \mathcal{E}(\mathbf{x}, t), \text{ en milieu isotrope,} \quad (1.9)$$

$$\mathcal{D}_i(\mathbf{x}, t) = \sum_{j=1}^3 \varepsilon_{ij}(\mathbf{x}) \mathcal{E}_j(\mathbf{x}, t), \quad i = 1, 3, \text{ en milieu anisotrope.} \quad (1.10)$$

$$\mathcal{B}(\mathbf{x}, t) = \mu(\mathbf{x}) \mathcal{H}(\mathbf{x}, t), \text{ en milieu isotrope,} \quad (1.11)$$

$$\mathcal{B}_i(\mathbf{x}, t) = \sum_{j=1}^3 \mu_{ij}(\mathbf{x}) \mathcal{H}_j(\mathbf{x}, t), \quad i = 1, 3, \text{ en milieu anisotrope.} \quad (1.12)$$

ε est la permittivité diélectrique et μ la perméabilité magnétique. En milieu anisotrope, ce sont des tenseurs.

Les équations de Maxwell en milieu isotrope se réécrivent alors en fonctions de $\mathcal{E}$ et $\mathcal{H}$ seulement :

$$\operatorname{div}(\varepsilon \mathcal{E}) = \rho, \quad (1.13)$$

$$\mathbf{rot} \mathcal{H} - \varepsilon \partial_t \mathcal{E} = \mathcal{J}, \quad (1.14)$$

$$\mathbf{rot} \mathcal{E} + \mu \partial_t \mathcal{H} = 0, \quad (1.15)$$

$$\operatorname{div}(\mu \mathcal{H}) = 0. \quad (1.16)$$

Lorsque le milieu est de plus homogène, ε et μ sont constants.

Le vide est un cas particulier de milieu *parfait, isotrope, homogène et isolant*, pour lequel la permittivité diélectrique, notée $\varepsilon_0 \simeq (36 \pi \cdot 10^9)^{-1} \text{C}^2 \cdot \text{N}^{-1} \cdot \text{m}^{-2}$ et la perméabilité magnétique notée $\mu_0 = 4 \pi \cdot 10^{-7} \text{F} \cdot \text{m}^{-1}$ sont telles que : $c^2 \varepsilon_0 \mu_0 = 1$, où $c \simeq 3 \cdot 10^8 \text{m} \cdot \text{s}^{-1}$ est la vitesse de propagation des ondes électromagnétiques dans le vide.

Dans le cas d'un milieu isotrope, homogène, isolant et non chargé ($\mathcal{J} = 0$, $\rho = 0$), $\mathcal{E}$ et $\mathcal{H}$ satisfont l'équation des ondes :

$$\Delta \mathcal{E} - \varepsilon \mu \partial_t^2 \mathcal{E} = 0, \text{ et } \Delta \mathcal{H} - \varepsilon \mu \partial_t^2 \mathcal{H} = 0. \quad (1.17)$$

Cette équation s'obtient en injectant (1.15) dans $\partial_t(1.14)$ pour le champ électrique et (1.14) dans $\partial_t(1.15)$ pour le champ magnétique, et en utilisant le fait que $\mathbf{rot} \mathbf{rot} - \mathbf{grad} \operatorname{div} = -\Delta$. L'onde électromagnétique se propage à la vitesse $(\varepsilon \mu)^{-1/2}$. Ainsi, les équations de Maxwell sont de nature *hyperbolique*.

1.3.3 Énergie électromagnétique

Le flux d'énergie du champ électromagnétique est représenté par le *vecteur de Poynting*, de dimension $\text{J.m}^{-2}.\text{s}^{-1}$:

$$\mathcal{S} = \mathcal{E} \times \mathcal{H}. \quad (1.18)$$

La densité d'énergie électromagnétique totale, de dimension J.m^{-3} est donnée par :

$$w = \frac{1}{2} (\mathcal{E} \cdot \mathcal{D} + \mathcal{B} \cdot \mathcal{H}). \quad (1.19)$$

Dans le vide, l'équation de conservation de l'énergie, appelée aussi *théorème de Poynting*, s'écrit ainsi :

$$\partial_t w + \text{div } \mathcal{S} = -\mathcal{J} \cdot \mathcal{E}. \quad (1.20)$$

La variation instantannée de l'énergie électromagnétique à l'intérieur d'un volume donné plus l'énergie s'écoulant, par unité de temps, à travers la surface délimitant le volume correspond au travail accompli par le champ électromagnétique sur les sources dans le volume. En milieu linéaire dispersif, il faut tenir compte des pertes ohmiques, et de l'absorption du milieu.

L'énergie électromagnétique totale dans un volume $\Omega \subset \mathbb{R}^3$ quelconque est donc :

$$W_\Omega(t) = \int_\Omega w \, d\Omega = \frac{1}{2} \int_\Omega (\varepsilon(\mathbf{x}) |\mathcal{E}(\mathbf{x}, t)|^2 + \mu(\mathbf{x}) |\mathcal{H}(\mathbf{x}, t)|^2) \, d\Omega \quad (1.21)$$

En intégrant (1.20) sur Ω , l'équation de conservation de l'énergie totale dans le vide s'écrit alors :

$$\frac{dW_\Omega}{dt} + \int_{\partial\Omega} \mathcal{S} \cdot \boldsymbol{\nu} \, d\Sigma + \int_\Omega \mathcal{E} \cdot \mathcal{J} \, d\Omega = 0, \quad (1.22)$$

où $\boldsymbol{\nu}$ est le vecteur normal sortant de $\partial\Omega$.

1.3.4 Conditions de transmission entre deux milieux matériels

Le théorème de Stokes et le théorème de la divergence permettent d'écrire les équations de Maxwell sous forme intégrale et de déduire les relations entre les composantes normales et tangentielles des champs d'une part et d'autre d'une surface Γ séparant deux milieux différents, et éventuellement porteuse d'une densité de charge surfacique σ et d'une densité de courant surfacique $\mathcal{K}$:

$$(\mathcal{D}_2 - \mathcal{D}_1) \cdot \boldsymbol{\nu} = \sigma \text{ sur } \Gamma, \quad (1.23)$$

$$(\mathcal{B}_2 - \mathcal{B}_1) \cdot \boldsymbol{\nu} = 0 \text{ sur } \Gamma, \quad (1.24)$$

$$\boldsymbol{\nu} \times (\mathcal{E}_2 - \mathcal{E}_1) = 0 \text{ sur } \Gamma, \quad (1.25)$$

$$\boldsymbol{\nu} \times (\mathcal{H}_2 - \mathcal{H}_1) = \mathcal{K} \text{ sur } \Gamma, \quad (1.26)$$

$\boldsymbol{\nu}$ étant le vecteur normal à Γ , dirigé du milieu 1 vers le milieu 2. Ainsi, à la traversée de l'interface :

- la composante normale de $\mathcal{B}$ est continue,
- la discontinuité de la composante normale de $\mathcal{D}$ en un point est égale à la densité de charge surfacique en ce point,
- la composante tangentielle de $\mathcal{E}$ est continue,
- la composante tangentielle de $\mathcal{H}$ subit une discontinuité égale à la densité de courant surfacique, et de direction $\mathcal{K} \times \boldsymbol{\nu}$.

Dans le cas particulier où le milieu 2 est un conducteur parfait (on note alors $\Gamma = \Gamma_C$), le champ électromagnétique est nul dans ce milieu : $\mathcal{E}_2 = \mathcal{B}_2 = 0$. On pose alors $\mathcal{E} = \mathcal{E}_1$, $\mathcal{B} = \mathcal{B}_1$ et on a dans le milieu 1 des conditions aux limites de conducteur parfait :

$$\mathcal{B} \cdot \boldsymbol{\nu} = 0 \text{ sur } \Gamma_C, \quad (1.27)$$

$$\boldsymbol{\nu} \times \mathcal{E} = 0 \text{ sur } \Gamma_C. \quad (1.28)$$

Notons que l'expression (1.27) apparaît comme redondante. En effet, si $\boldsymbol{\nu} \times \mathcal{E} = 0$ sur Γ_C et à $t = 0$, $\mathcal{B}(0) \cdot \boldsymbol{\nu} = 0$ sur Γ_C , alors d'après (1.4), $\mathcal{B}(t) \cdot \boldsymbol{\nu} = 0$ sur Γ_C pour tout t .

1.3.5 Condition aux limites en domaine borné

Lorsqu'on borne le domaine d'étude par une frontière artificielle Γ_A , il faut imposer des conditions aux limites au bord du domaine borné Ω obtenu. Au premier ordre, les conditions aux limites sur Γ_A s'écrivent ainsi :

$$(\mathcal{E} - c\mathcal{B} \times \boldsymbol{\nu}) \times \boldsymbol{\nu} = \mathbf{e}^* \times \boldsymbol{\nu}, \text{ sur } \Gamma_A, \mathbf{e}^* \text{ donné}, \quad (1.29)$$

$$\text{ou bien, de façon analogue : } (c\mathcal{B} - \mathcal{E} \times \boldsymbol{\nu}) \times \boldsymbol{\nu} = \mathbf{b}^* \times \boldsymbol{\nu}, \text{ sur } \Gamma_A, \mathbf{b}^* \text{ donné}. \quad (1.30)$$

Cette condition s'obtient en approchant localement la frontière Γ_A par son plan tangent, et en écrivant qu'une onde plane à incidence normale sort du domaine sans être réfléchie, pour $\mathbf{e}^* = 0$ ou $\mathbf{b}^* = 0$. Elle est donc exacte dans ce cas de figure, sinon, c'est une approximation. Lorsque $\mathbf{e}^* \neq 0$ ou $\mathbf{b}^* \neq 0$, cette condition traduit la pénétration d'une onde plane à incidence normale dans le domaine. Cette condition aux limites est souvent appelée *condition de Silver-Müller*. Elle est en générale suffisante pour les problèmes intérieurs. Pour les problèmes extérieurs, il est préférable d'utiliser une approximation d'ordre deux, afin de ne pas polluer la solution par des réflexions parasites.

1.3.6 Condition de radiation en domaine non-borné

Pour les problèmes stationnaires en domaine non-borné, les équations de Maxwell doivent être complétées par des conditions de radiation qui éliminent les ondes venant de l'infini. Considérons le cas où on envoie une onde électromagnétique incidente $\mathcal{E}^i$ (par exemple provenant d'un radar) sur un objet inhomogène borné (tel qu'un avion). L'onde est réfléchie et diffusée en une onde $\mathcal{E}^s$ et le champ électromagnétique total est : $\mathcal{E}^i + \mathcal{E}^s$. L'onde incidente est une onde plane : $\mathcal{E}^i = \mathbf{p} e^{i\kappa \mathbf{x} \cdot \mathbf{d}}$, où $\mathbf{p} \in \mathbb{R}^3$ est le vecteur de polarisation et $\mathbf{d} \in \mathbb{R}^3$ est le vecteur unitaire de direction de propagation de l'onde. $\mathbf{p}$ et $\mathbf{d}$ sont orthogonaux. $\mathcal{E}^i$ satisfait les équations de Maxwell en l'absence de diffusion :

$$\mathbf{rot} \mathbf{rot} \mathcal{E}^i - \kappa^2 \mathcal{E}^i = 0, \text{ dans } \mathbb{R}^3.$$

Le champ électromagnétique diffusé doit alors satisfaire la *condition de radiation de Silver-Müller* suivante [85] :

$$\lim_{R \rightarrow 0} R (\mathbf{rot} \mathcal{E}^s \times \hat{\mathbf{x}} - i\kappa \mathcal{E}^s) = 0, \quad (1.31)$$

où $R = |\mathbf{x}|$ et $\hat{\mathbf{x}} = \mathbf{x}/R$.

La difficulté pour discrétiser ce problème est qu'il est posé en domaine infini. On peut borner le domaine par une frontière artificielle loin de l'objet réfléchissant, et imposer la condition de radiation de Silver-Müller sur cette frontière.

1.3.7 Effet de pointe

Considérons deux sphères conductrices dans le vide, soumises au même potentiel. On observe que le champ est plus grand à la surface de la plus petite sphère.

En effet, si la plus grande sphère, de rayon a porte une charge Q , son potentiel électrique vaut environ $\phi_1 = \frac{1}{4\pi\epsilon_0} \frac{Q}{a}$. Le potentiel électrique de la petite sphère, de rayon b et portant une charge q est de l'ordre de $\phi_2 = \frac{1}{4\pi\epsilon_0} \frac{q}{b}$. Ainsi, si les sphères sont au même potentiel, on a : $\frac{Q}{a} = \frac{q}{b}$.

Sur la surface d'une des sphères, le champ électrique est proportionnel à la densité surfacique de charges σ , qui vaut à peu près la charge totale divisée par la surface de la sphère. D'où : $\frac{\|\mathcal{E}_1\|}{\|\mathcal{E}_2\|} \simeq \frac{Q/a^2}{q/b^2} = \frac{b}{a}$. Ainsi, les champs sont inversement proportionnels aux rayons, et le champ électrique de la petite sphère est plus important que celui de la grosse sphère.

Le même effet se produit si on charge un conducteur qui a une pointe ou une extrémité très aiguë : le champ électromagnétique au voisinage de la pointe est beaucoup plus grand que le champ dans les autres régions. Les charges s'étendent le plus possible sur la surface d'un conducteur, certaines charges sur le conducteur sont poussées vers l'extrémité, qui a une surface petite devant la surface de tout le conducteur.

On en déduit que la densité surfacique de charges est plus importante localement, sur l'extrémité, que sur le reste du conducteur, ce qui implique un champ intense au voisinage de la pointe.

On appelle *singularités géométriques* les arêtes et/ou coins que forme un conducteur. Le champ électromagnétique est alors localement intense, il se produit un effet de pointe.

La présence de singularités géométriques dans le milieu matériel peut être un choix délibéré du constructeur pour générer des champs intenses (singularité active), ou une contrainte de conception (singularité passive). Dans les deux cas, la solution des équations de Maxwell est singulière, et la valeur précise du champ est cruciale pour une analyse correcte des phénomènes physiques observés.

1.4 Hypothèses mathématiques sur les données

Pour être en mesure de résoudre numériquement les équations de Maxwell, il faut faire des hypothèses mathématiques sur le champ électromagnétique et les données. Dans cette section, nous établissons de façon formelle ces hypothèses, qui reposent sur des considérations physiques, en particulier sur le fait l'énergie électromagnétique soit finie. Nous noterons $L^2(\Omega)$ l'espace des fonctions scalaires de carré intégrable sur $\Omega \subset \mathbb{R}^3$, et $\mathcal{L}^2(\Omega) = L^2(\Omega)^3$ l'espace des fonctions vectorielles dont les composantes sont de carré intégrable.

Dans le cas statique, le champ électromagnétique et les données ne dépendent pas du temps. Les équations statiques en $(\mathcal{E}, \mathcal{H})$ s'écrivent :

$$\mathbf{rot} \mathcal{E} = 0, \quad \text{et } \mathbf{div} \mathcal{E} = \rho/\epsilon, \quad (1.32)$$

$$\mathbf{rot} \mathcal{H} = \mathcal{J}, \quad \text{et } \mathbf{div} (\mu \mathcal{H}) = 0. \quad (1.33)$$

Pour simplifier, nous supposons ϵ et μ constants (cas d'un milieu homogène).

- Le champ électrostatique $\mathcal{E}$ est à rotationnel nul. Il existe donc un potentiel scalaire ϕ tel que : $\mathcal{E} = -\mathbf{grad} \phi$. Ce potentiel satisfait l'équation de Poisson : $-\Delta\phi = \rho/\epsilon$. Par intégration par parties

sur un volume Ω , on peut réécrire l'énergie électrostatique totale en fonction de ρ et ϕ dans Ω :

$$\begin{aligned} W_{ES} &= \frac{1}{2} \int_{\Omega} \varepsilon |\mathcal{E}|^2 d\Omega = \frac{1}{2} \int_{\Omega} \varepsilon \mathbf{grad} \phi \cdot \mathbf{grad} \phi d\Omega \\ &= \frac{1}{2} \int_{\Omega} \varepsilon (-\Delta \phi) \phi d\Omega + \frac{1}{2} \int_{\partial\Omega} \varepsilon \phi \partial_{\nu} \phi d\Sigma, \\ &= \frac{1}{2} \int_{\Omega} \rho \phi d\Omega + \frac{1}{2} \int_{\partial\Omega} \varepsilon \phi \partial_{\nu} \phi d\Sigma, \end{aligned}$$

où $\partial_{\nu} \phi$ est la dérivée partielle normale à ϕ sur $\partial\Omega$. Supposons que le champ soit généré par une source, et que Ω soit une sphère de centre la source. Comme le potentiel ϕ est en $1/r$, où r est la distance à la source, $\partial_{\nu} \phi$ est en $1/r^2$. On en déduit que $\lim_{r \rightarrow \infty} \int_{\partial\Omega} \varepsilon \phi \partial_{\nu} \phi d\Sigma \approx \lim_{r \rightarrow \infty} 1/r^3 = 0$. Cela se généralise à d'autres formes de volume : l'intégrale du bord s'annule, et on a :

$$W_{ES} = \frac{1}{2} \int_{\mathbb{R}^3} \rho \phi d\mathbf{x}.$$

Comme $W_{ES} < \infty$, on a $\mathcal{E} \in \mathcal{L}^2(\mathbb{R}^3)$. Ainsi, les dérivées partielles de ϕ sont de carré intégrable, d'où : $\phi \in H^1(\mathbb{R}^3)$. Comme $\rho \phi$ est intégrable, on en déduit que ρ est dans le dual de $H^1(\Omega)$, noté $H^{-1}(\mathbb{R}^3)$. De plus, comme $\mathcal{E}$ est à rotationnel nul, $\mathcal{E}$ est un vecteur de $\mathcal{H}(\mathbf{rot}, \mathbb{R}^3)$, l'espace des fonctions vectorielles de rotationnel de carré intégrable.

• L'induction magnétostatique $\mathcal{B} = \mu \mathcal{H}$ est à divergence nulle. Il existe alors un potentiel vecteur $\mathcal{A}$ tel que : $\mathcal{B} = \mathbf{rot} \mathcal{A}$, c'est-à-dire $\mathcal{H} = \frac{1}{\mu} \mathbf{rot} \mathcal{A}$. Si on choisit $\mathcal{A}$ à divergence nulle (ce choix est appelé jauge de Coulomb), alors $\mathcal{A}$ satisfait l'équation : $-\Delta \mathcal{A} = \mathcal{J}$. On peut réécrire l'énergie magnétostatique totale en fonction de $\mathcal{A}$ et $\mathcal{J}$:

$$\begin{aligned} W_{MS} &= \frac{1}{2} \int_{\Omega} \mu |\mathcal{H}|^2 d\Omega = \frac{1}{2} \int_{\Omega} \mathcal{H} \cdot \mathbf{rot} \mathcal{A} d\Omega \\ &= \frac{1}{2} \int_{\Omega} \mathbf{rot} \mathcal{H} \cdot \mathcal{A} d\Omega + \frac{1}{2} \int_{\partial\Omega} \mathcal{A} \cdot (\mathbf{rot} \mathcal{H} \times \boldsymbol{\nu}) d\Sigma, \\ &= \frac{1}{2} \int_{\Omega} \mathcal{J} \cdot \mathcal{A} d\Omega + \frac{1}{2} \int_{\partial\Omega} \mu \mathcal{A} \cdot (\mathbf{rot} \mathcal{B} \times \boldsymbol{\nu}) d\Sigma, \end{aligned}$$

De nouveau, on suppose que le champ soit généré par une source, et que Ω soit une sphère de centre la source. Le potentiel vecteur $\mathcal{A}$ est en $1/r$ et $\mathbf{rot} \mathcal{B} \times \boldsymbol{\nu}$ est en $1/r^2$. Ainsi, l'intégrale au bord s'annule lorsque $r \rightarrow \infty$, et on a :

$$W_{MS} = \frac{1}{2} \int_{\mathbb{R}^3} \mathcal{J} \cdot \mathcal{A} d\mathbf{x}.$$

Comme $W_{MS} < \infty$, on a $\mathcal{B} \in \mathcal{L}^2(\mathbb{R}^3)$. Donc le rotationnel de $\mathcal{A}$ est de carré intégrable, c'est-à-dire : $\mathcal{A} \in \mathcal{H}(\mathbf{rot}, \mathbb{R}^3)$. On en déduit alors que $\mathcal{J} \cdot \mathcal{A}$ est intégrable, et que $\mathcal{J}$ est dans le dual de $\mathcal{H}(\mathbf{rot}, \mathbb{R}^3)$, noté $\mathcal{H}(\mathbf{rot}, \mathbb{R}^3)'$. Notons que $\mathcal{J}$ étant un rotationnel, $\mathcal{J}$ est à divergence nulle. De plus, comme $\mu \mathcal{H}$ est à divergence nulle, $\mathcal{H}$ est un vecteur de $\mathcal{H}(\mathbf{div}, \mathbb{R}^3)$, l'espace des fonctions vectorielles de divergence de carré intégrable.

Les hypothèses de régularités minimales du problème statique sont donc :

$$\rho \in H^{-1}(\mathbb{R}^3), \mathcal{J} \in \mathcal{H}(\mathbf{rot}, \mathbb{R}^3)' \text{ avec } \mathbf{div} \mathcal{J} = 0 \text{ et } \mathcal{E} \in \mathcal{H}(\mathbf{rot}, \mathbb{R}^3), \mathcal{H} \in \mathcal{H}(\mathbf{div}, \mathbb{R}^3).$$

Notons que l'hypothèse mathématique faite sur ρ ne peut être vérifiée que si on exclut certaines singularités de charge telles que les singularités de charges ponctuelles ou linéiques. Il faut alors changer de modélisation et passer au modèle microscopique. Par la suite, on fait souvent l'hypothèse supplémentaire que $\rho \in L^2(\mathbb{R}^3)$, c'est-à-dire que $\mathcal{E} \in \mathcal{H}(\text{div}, \mathbb{R}^3)$. De même, par commodité, on suppose en général que : $\mathcal{J} \in \mathcal{L}^2(\mathbb{R}^3)$, d'où : $\mathcal{H} \in \mathcal{H}(\mathbf{rot}, \mathbb{R}^3)$. Les hypothèses de régularités usuelles du problème statique sont alors :

$$\rho \in L^2(\mathbb{R}^3), \mathcal{J} \in \mathcal{L}^2(\mathbb{R}^3) \text{ avec } \text{div } \mathcal{J} = 0 \text{ et } \mathcal{E}, \mathcal{H} \in \mathcal{H}(\mathbf{rot}, \mathbb{R}^3) \cap \mathcal{H}(\text{div}, \mathbb{R}^3).$$

Dans le cas général, où ε et μ ne sont pas constants, avec les hypothèses suivantes : ε et μ sont mesurables uniformément bornées et strictement positives, on obtient que :

$$\mathcal{E} \in \mathcal{H}(\mathbf{rot}, \mathbb{R}^3) \cap \mathcal{H}(\text{div } \varepsilon, \mathbb{R}^3) \text{ et } \mathcal{H} \in \mathcal{H}(\mathbf{rot}, \mathbb{R}^3) \cap \mathcal{H}(\text{div } \mu, \mathbb{R}^3),$$

soit : $\text{div } (\varepsilon \mathcal{E}) \in L^2(\mathbb{R}^3)$ et $\text{div } (\mu \mathcal{H}) \in L^2(\mathbb{R}^3)$.

Dans un ouvert borné Ω , on arrive à des conclusions semblables : $\mathcal{E} \in \mathcal{H}(\mathbf{rot}, \Omega) \cap \mathcal{H}(\text{div } \varepsilon, \Omega)$ et $\mathcal{H} \in \mathcal{H}(\mathbf{rot}, \Omega) \cap \mathcal{H}(\text{div } \mu, \Omega)$, la différence portant sur les *conditions aux limites*. Dans le cas où Ω est l'intérieur d'un conducteur parfait, on a alors : $\mathcal{E} \times \boldsymbol{\nu}_{|\partial\Omega} = 0$, et $\mathcal{B} \cdot \boldsymbol{\nu}_{|\partial\Omega} = 0$. On a alors : $\mathcal{E} \in \mathcal{H}_0(\mathbf{rot}, \Omega) \cap \mathcal{H}(\text{div } \varepsilon, \Omega)$ et $\mathcal{H} \in \mathcal{H}(\mathbf{rot}, \Omega) \cap \mathcal{H}_0(\text{div } \mu, \Omega)$.

Nous ne détaillons pas ici le cas instationnaire. On a des hypothèses similaires sur la dépendance en espace, plus des hypothèses de continuité sur la dépendance en temps. Pour plus de détails, on peut lire notamment le document de F. Assous et P. Ciarlet, Jr. [6] (chap. 2).

Notons que la régularité du champ dépend entre autre de la présence de singularités.

1.5 Différents modèles

Selon le type de problème que l'on souhaite étudier, il existe plusieurs façon de réécrire les équations de Maxwell. Dans cette section, nous indiquons différents modèles permettant de les simplifier.

1.5.1 Le modèle de Darwin

Lorsqu'on simule des faisceaux de particules chargées, pour lesquels il n'y a pas de phénomène haute fréquence ou de changement rapide de courant, on peut négliger la composante transverse (orthogonale à la direction du faisceau) du courant de déplacement. C'est le modèle Darwin, qui correspond à une approximation d'ordre un en terme de développement asymptotique des équations de Maxwell [97].

Le champ électrique se décompose de la façon suivante (décomposition de Helmholtz) :

$$\mathcal{E} = \mathcal{E}^T + \mathcal{E}^L.$$

$\mathcal{E}^L$, la partie longitudinale est telle que : $\mathbf{rot } \mathcal{E}^L = 0$.

$\mathcal{E}^T$, la partie transverse se caractérise par $\text{div } \mathcal{E}^T = 0$.

Lorsque la vitesse caractéristique du phénomène étudié est petite devant la vitesse de la lumière, on néglige $\partial_t \mathcal{E}^T$, la composante transverse du courant de déplacement. On obtient alors un système

de trois problèmes elliptiques en $\mathcal{E}^L$, $\mathcal{B}$ et $\mathcal{E}^T$:

$$-\Delta\phi = \rho/\varepsilon, \mathcal{E}^L = -\mathbf{grad}\phi,$$

$$\begin{aligned} \mathbf{rot}\mathbf{rot}\mathcal{B} &= \mathbf{rot}\mathcal{J}, \\ \operatorname{div}\mathcal{B} &= 0, \end{aligned}$$

$$\begin{aligned} \mathbf{rot}\mathbf{rot}\mathcal{E}^T &= -\mu\mathbf{rot}\mathcal{B}, \\ \operatorname{div}\mathcal{E}^T &= 0. \end{aligned}$$

$\mathbf{rot}\mathcal{B}$ est considéré comme une donnée pour calculer $\mathcal{E}^T$. La difficulté de ce problème est la prise en compte des conditions aux limites, qu'il faut partager entre les deux composantes du champ électrique [59]. Dans [96], P.-A. Raviart et E. Sonnendrücker présentent une méthode asymptotique pour déterminer les conditions aux limites sur $\mathcal{B}$ qui permettent d'obtenir que le problème en $\mathcal{B}$ soit bien posé. Dans [103], E. Sonnendrücker et al. utilisent le système couplé Vlasov-Darwin comme approximation du système couplé Vlasov-Maxwell. Enfin, une analyse de la convergence des éléments finis pour ce modèle a été faite par P. Ciarlet, Jr. et J. Zou dans [43].

1.5.2 Les modèles quasi-électrostatique et quasi-magnétostatique

Lorsqu'on peut négliger le terme $\partial_t\mathcal{B}$ dans l'équation de Faraday (1.4), on obtient le modèle quasi-électrostatique. Ce modèle est valide lorsque la vitesse caractéristique des phénomènes magnétiques est petite devant la vitesse de propagation de l'onde. On obtient le problème suivant :

$$\operatorname{div}(\varepsilon\mathcal{E}) = \rho, \mathbf{rot}\mathcal{E} = 0, \quad (1.34)$$

avec des conditions aux limites adéquates, ρ dépendant du temps.

Avec l'hypothèse supplémentaire $\partial_t\mathcal{E} = 0$, ρ est indépendant du temps, c'est le problème électrostatique. L'équation (1.34) devient elliptique. La solution $\mathcal{E}$ est telle que : $\mathcal{E} = -\mathbf{grad}\phi$, avec $-\Delta\phi = \rho/\varepsilon$. C'est le *modèle de Poisson*, bien moins coûteux à résoudre que le système complet des équations de Maxwell. L'une des applications de ce modèle est l'analyse de distribution de charges (en anglais C.D.A. : charge distribution analysis). Le but est d'analyser la charge électrique présente dans un isolant, afin d'en mesurer la qualité d'isolation. À cette fin, on plonge l'isolant dans un champ électrique non uniforme, et on mesure la force à laquelle il est soumis.

L'étude de ce modèle est nécessaire lorsqu'on s'intéresse au problème quasi-électrostatique. En effet, la dépendance en temps du second membre requiert la résolution d'une suite de problèmes électrostatiques. Cependant, on n'a pas toujours la même condition aux limites.

Si on peut négliger le courant de déplacement $\partial_t\mathcal{D}$ par rapport au courant induit dans l'équation d'Ampère (1.3), on obtient le modèle quasi-magnétostatique, dont l'une des applications est le calcul des courants induits dans un matériau (ou courants de Foucault). Dans ce cas, on dispose de la relation complémentaire $\mathcal{J} = \sigma\mathcal{E}$, de sorte que $\mathcal{B}$ peut être écrit comme la solution d'une équation parabolique :

$$\begin{aligned} \partial_t\mathcal{B} + \mathbf{rot}\left(\frac{1}{\sigma}\mathbf{rot}\frac{1}{\mu}\mathcal{B}\right) &= 0, \\ \operatorname{div}\mathcal{B} &= 0. \end{aligned}$$

Le modèle magnétostatique, valable lorsqu'on peut négliger $\partial_t\mathcal{B}$, c'est-à-dire lorsque $\mathcal{J}$ est indépendant du temps s'écrit alors :

$$\operatorname{div}\mathcal{B} = 0, \mathbf{rot}\frac{1}{\mu}\mathcal{B} = \mathcal{J}. \quad (1.35)$$

Dans le cas où μ est constant, $\mathcal{B}$ est alors solution de l'équation elliptique suivante : $-\Delta \mathcal{B} = \mu \mathbf{rot} \mathcal{J}$.

D'un point de vue numérique, les modèles statiques permettent de faire l'étude spatiale des méthodes de résolution envisagées, avant de les appliquer au cas instationnaire ou au cas harmonique. Ces modèles fournissent de précieuses informations pour comprendre l'influence des singularités géométriques sur le comportement en espace et sur le comportement en temps du champ électromagnétique.

1.5.3 Les équations harmoniques

On suppose que les fonctions électromagnétiques ont une dépendance harmonique en temps, de pulsation ω , c'est-à-dire qu'elles sont de la forme :

$$\begin{aligned}\mathcal{E}(\mathbf{x}, t) &= \Re(\mathbf{e}(\mathbf{x})e^{i\omega t}), \\ \mathcal{H}(\mathbf{x}, t) &= \Re(\mathbf{h}(\mathbf{x})e^{i\omega t}), \\ \rho(\mathbf{x}, t) &= \Re(r(\mathbf{x})e^{i\omega t}), \\ \mathcal{J}(\mathbf{x}, t) &= \Re(\mathbf{j}(\mathbf{x})e^{i\omega t}),\end{aligned}$$

où $\mathbf{e}$, $\mathbf{h}$, $\mathbf{j}$ sont à valeur dans $\mathbb{C}^3$ et r est à valeur dans $\mathbb{C}$. Les équations de Maxwell se réécrivent alors :

$$\begin{aligned}i\omega \varepsilon \mathbf{e} - \mathbf{rot} \mathbf{h} &= -\mathbf{j}, \\ i\omega \mu \mathbf{h} + \mathbf{rot} \mathbf{e} &= 0, \\ \operatorname{div}(\varepsilon \mathbf{e}) &= r, \\ \operatorname{div}(\mu \mathbf{h}) &= 0.\end{aligned}$$

On peut découpler ces équations ainsi, en injectant $\mathbf{h} = (i\omega \mu)^{-1} \mathbf{rot} \mathbf{e}$ dans la première équation, ou $\mathbf{e} = (i\omega \varepsilon)^{-1}(\mathbf{rot} \mathbf{h} - \mathbf{j})$ dans la seconde équation :

$$\begin{aligned}-\omega^2 \varepsilon \mathbf{e} + \mathbf{rot}(\mu^{-1} \mathbf{rot} \mathbf{e}) &= -i\omega \mathbf{j}, \\ \omega^2 \mu \mathbf{h} - \mathbf{rot}(\mu^{-1}(\mathbf{rot} \mathbf{h} - \mathbf{j})) &= 0.\end{aligned}$$

Lorsque $\mathbf{j}$ est nul, on doit résoudre un problème aux valeurs propres. Les équations harmoniques permettent de modéliser le comportement d'une onde électromagnétique dans une cavité. Différents modèles ont été développés, selon que le milieu soit conducteur ou non. La première difficulté est de trouver une formulation variationnelle du problème qui satisfasse l'alternative de Fredholm ([27], chap. 6), qui vérifie les propriétés de coercivité et de compacité, c'est-à-dire pour la seconde, l'injection compacte de l'espace fonctionnel choisi dans $\mathcal{L}^2(\Omega)$.

Comme l'injection de $\mathcal{H}_0(\mathbf{rot}, \Omega)$ dans $\mathcal{L}^2(\Omega)$ n'est pas compacte, une méthode consiste à ajouter un terme d'intégration en div-div dans la formulation variationnelle (qui contient un terme d'intégration en $\mathbf{rot} \cdot \mathbf{rot}$ et un terme d'intégration $\mathcal{L}^2$), afin d'éliminer les fonctions propres non physiques. On parle alors de *régularisation* ou de *pénalisation*. La coercivité de la forme bilinéaire pénalisée a été montrée par M. Costabel dans [50]. On peut discrétiser la formulation variationnelle obtenue par des éléments finis nodaux [68, 53]. Une autre méthode consiste à ajouter un multiplicateur de Lagrange portant sur la condition de divergence. Ces méthodes ont été introduites par F. Kikuchi dans [74], qui prouve dans [75] des propriétés de compacité discrète pour les éléments finis de degré le plus bas de Nédélec. Une étude comparative des éléments finis d'arêtes avec des éléments finis nodaux est menée par D. Boffi et al. dans [22] (pour le 2D) et [23]. Les éléments finis nodaux proposés dans ces articles sont biquadratiques et permettent de capter les singularités du champ par une méthode de projection. L'étude de la compacité discrète pour une méthode *hp* d'éléments finis d'arêtes 2D est faite par D. Boffi et al. dans [21] et [20]. Dans [32], A. Buffa et al. prouvent la compacité discrète pour les éléments finis d'arêtes sur un maillage anisotrope.

A. Alonso s'est intéressée au cas où le milieu est hétérogène, et se comporte comme un conducteur dans une partie du domaine et un isolant parfait dans une autre partie. Il existe deux modèles :
- Le modèle basse fréquence, pour lequel on élimine le terme en ω^2 et on prend en compte la loi d'Ohm. Le champ électrique harmonique $\mathbf{e}$ satisfait alors :

$$\mathbf{rot}(\mu^{-1}\mathbf{rot}\mathbf{e}) + i\omega\sigma\mathbf{e} = 0,$$

avec les conditions aux limites adéquates. Dans [1], A. Alonso établit une preuve mathématique de ce modèle, pour lequel l'alternative de Fredholm a été prouvée pour la forme bilinéaire :

$$\int_{\Omega} \mu^{-1}\mathbf{rot}\mathbf{u} \cdot \mathbf{rot}\mathbf{v} \, d\Omega + i\omega \int_{\Omega} \sigma\mathbf{u} \cdot \mathbf{v} \, d\Omega.$$

- Le modèle haute fréquence : on conserve le terme en ω^2 et on prend en compte la loi d'Ohm. Le champ électrique harmonique $\mathbf{e}$ satisfait :

$$\mathbf{rot}(\mu^{-1}\mathbf{rot}\mathbf{e}) - \alpha^2(\varepsilon - i\omega^{-1}\sigma)\mathbf{e} = 0,$$

avec les conditions aux limites adéquates. Dans [3], A. Alonso et A. Valli prouvent l'alternative de Fredholm, en étendant le résultat de compacité de C. Weber [107] (injection compacte de $\mathcal{X}_{\varepsilon}^0$ dans $\mathcal{L}^2(\Omega)$) pour l'espace $\mathcal{H}(\text{div}\eta, \Omega) \cap \mathcal{H}_0(\mathbf{rot}, \Omega)$, avec $\eta = \varepsilon - i\omega^{-1}\sigma$.

1.5.4 Contrôlabilité

Dans certaines situations physiques (en particulier : les antennes) on ne dispose pas de données volumiques, mais surfaciques. Il faut alors réécrire les équations de Maxwell, en considérant que l'évolution temporelle du champ électromagnétique est gouvernée par un courant surfacique tangentiel de densité $\mathcal{J}$. On veut résoudre :

$$\begin{aligned} \varepsilon\partial_t\mathcal{E} - \mathbf{rot}\mathcal{H} &= 0 \text{ et } \mu\partial_t\mathcal{H} + \mathbf{rot}\mathcal{E} = 0, & \text{dans } \Omega, \\ \text{div}(\varepsilon\mathcal{E}) &= 0 \text{ et } \text{div}(\mu\mathcal{H}) = 0, & \text{dans } \Omega, \\ \mathcal{H} \times \boldsymbol{\nu} &= \mathcal{J}, & \text{sur } \Gamma_J, \\ \mathcal{H} \times \boldsymbol{\nu} &= 0 \text{ et } \mathcal{E} \cdot \boldsymbol{\nu} = 0, & \text{sur } \partial\Omega \setminus \Gamma_J, \end{aligned}$$

où $\Gamma_J \subset \partial\Omega$ est le support (ouvert) du courant. Le champ électromagnétique étant dans $\mathcal{H}(\mathbf{rot}, \Omega)$, il est intéressant d'étudier l'espace des traces de $\mathcal{H}(\mathbf{rot}, \Omega)$. A. Alonso et A. Valli décrivent dans [2] des opérateurs d'extensions continus de $\mathcal{H}(\mathbf{rot}, \Omega)$ dans l'espace de ses traces tangentielles, ce qui en pratique permet l'approximation numérique du problème surfacique. Le cas d'un milieu hétérogène avec une frontière non régulière est étudié par S. Nicaise dans [89], qui établit des estimations d'énergie. Ces résultats s'appliquent au problème inverse qu'est la reconstitution d'antenne.

1.5.5 Équations instationnaires

Pour la modélisation numérique des plasmas, l'étude de dispositifs sélectifs en fréquence, il est nécessaire de disposer de codes résolvant le système couplé Maxwell-Vlasov. Pour cela, on doit être en mesure de résoudre les équations de Maxwell dépendant du temps (1.2)-(1.5). La discrétisation de ces équations, qui sont de premier ordre en temps mène à des algorithmes instables. Des oscillations apparaissent et on ne peut pas les contrôler.

Considérons le cas d'un milieu isotrope. En injectant (1.15) dans $\partial_t(1.14)$ pour le champ électrique et (1.14) dans $\partial_t(1.15)$ pour le champ magnétique, on obtient les équations suivantes, de

second ordre en temps et en espace, similaires à l'équation des ondes :

$$\varepsilon \partial_t^2 \mathcal{E} + \mathbf{rot} \left(\frac{1}{\mu} \mathbf{rot} \mathcal{E} \right) = -\partial_t \mathcal{J}, \quad (1.36)$$

$$\mu \partial_t^2 \mathcal{H} + \mathbf{rot} \left(\frac{1}{\varepsilon} (\mathbf{rot} \mathcal{H} - \mathcal{J}) \right) = 0. \quad (1.37)$$

La discrétisation de ces équations de second ordre permet de contrôler les dérivées des solutions.

P. Monk [81, 80] a contribué à l'étude de la discrétisation de (1.36)-(1.37) par des éléments finis discontinus :

- Dans [81], des estimations d'erreur sont prouvées pour la discrétisation en espace de (1.36) avec des éléments finis de Nédélec. Les estimations en norme $\mathcal{H}(\mathbf{rot}, \Omega)$ sont en h^k pour un champ suffisamment régulier (h est le pas du maillage et $k > 0$ est lié à la régularité des solutions). On a des estimations en norme $\mathcal{L}^2$ d'ordre h^{k+1} lorsque le domaine est convexe.

- Dans [80], trois méthodes d'éléments finis mixtes sont détaillées pour le calcul du champ électrique (et de la composante longitudinale du champ magnétique) en $2D$. La première méthode, consiste à prendre un champ électrique constant par élément, mais la condition inf-sup n'est pas vérifiée. Pour la seconde méthode, les degrés de liberté du champ électrique sont les composantes normales aux arêtes. Pour permettre à la permittivité diélectrique d'être discontinue, on a deux valeurs de la composante normale du champ électrique de part et d'autre de l'arête. Enfin, pour la dernière méthode, les degrés de liberté du champ électrique sont les composantes tangentielles aux arêtes : on est conforme dans $\mathcal{H}(\mathbf{rot}, \Omega)$, et ε peut être discontinue.

Enfin, dans [44], P. Ciarlet, Jr. et J. Zou prouvent des estimations d'erreurs pour la discrétisation en temps et en espace de (1.36) avec des éléments finis de Nédélec. De plus, les résultats de [81] sont étendus à des solutions moins régulières. Dans [109], J. Zhao obtient des résultats similaires, avec ε et μ discontinus.

La discrétisation de (1.36)-(1.37) par des éléments finis continus a fait l'objet de la thèse d'É. Heintzé, en 1992 [69, 10] (voir aussi [26]) et a abouti à un code de calcul Maxwell-Vlasov $3D$, valable lorsque le domaine de calcul est convexe (ou la solution suffisamment régulière). Dix ans plus tard (grâce à la méthode du complément singulier, introduite par Assous et al. [9] en 1998), à l'issue de la thèse d'E. Garcia [63], on dispose d'un code de résolution avec des éléments finis continus du problème Maxwell-Vlasov couplé en domaine non-convexe $2D$.

1.6 Méthodes de résolution

Différentes méthodes de résolution de ces équations ont été étudiées : différences finies, éléments finis, volumes finis pour les domaines bornés ; méthodes intégrales, éléments infinis pour les problèmes de diffraction. Nous n'en présentons que quelques unes.

1.6.1 Différences finies

Considérons $\Omega =]0, 1[^3$. Soit $N \in \mathbb{N}$. La méthode des différences finies en espace consiste à calculer le champ électromagnétique aux points : $(x_i, y_j, z_k) = (i/(N+1), j/(N+1), k/(N+1))$, où $(i, j, k) \in \{0, \dots, N+1\}$. Posons $h = 1/(N+1)$, le pas du maillage, et $\mathcal{E}_{i,j,k}(t) \approx \mathcal{E}(x_i, y_j, z_k; t)$ (de même pour $\mathcal{H}$). Les dérivées partielles sont approchées par différences finies par exemple de la façon suivante (de même pour $\mathcal{H}$) :

$$\partial_x \mathcal{E}(x_i, y_j, z_k; t) \approx \frac{\mathcal{E}_{i+1,j,k}(t) - \mathcal{E}_{i,j,k}(t)}{h}, \text{ de même pour } \partial_y \mathcal{E} \text{ et } \partial_z \mathcal{E}.$$

Les valeurs au bord (en $i = 0$, ou $i = N + 1$, etc) sont données par les conditions aux limites. Plus N est grand, plus la méthode est précise. On peut choisir des pas différents selon les directions x , y , z . Il existe bien sûr d'autres schémas de discrétisation. L'inconvénient de la méthode des différences finies en espace est qu'elle s'utilise difficilement dans le cas de géométries complexes.

En général, on applique la méthode des différences finies à la discrétisation en temps. Soit T_f le temps final, et N_T le nombre de pas de temps. On pose $\Delta t = T_f/N_T$ et $\mathcal{E}^n \approx \mathcal{E}(\mathbf{x}; t_n)$, où $t_n = n \Delta t$, $n \in \{0, \dots, N_T\}$. On approche alors $\partial_t \mathcal{E}$ et $\partial_t^2 \mathcal{E}$ (de même pour $\mathcal{H}$) par les schémas suivants :

$$\partial_t \mathcal{E}(\mathbf{x}; t_n) \approx \frac{\mathcal{E}(\mathbf{x})^{n+1} - \mathcal{E}(\mathbf{x})^n}{\Delta t}; \quad \partial_t^2 \mathcal{E}(\mathbf{x}; t_n) \approx \frac{\mathcal{E}(\mathbf{x})^{n+1} - 2\mathcal{E}(\mathbf{x})^n + \mathcal{E}(\mathbf{x})^{n-1}}{\Delta t^2} \quad (1.38)$$

Le second schéma est appelé schéma saute-mouton. Pour la discrétisation des équations de Maxwell par les différences finies en temps et en espace, le schéma le plus utilisé est le schéma de Yee [108]. Dans ce cas, on utilise un schéma d'approximation centré en espace ($\partial_x \mathcal{E} \simeq \mathcal{E}_{i+1/2,j,k} - \mathcal{E}_{i-1/2,j,k}$), et le schéma ci-dessus en temps. D'autres schémas sont disponibles dans [104].

Lorsqu'on discrétise par les différences finies en espace l'équation des ondes (1.17), on obtient un système linéaire de la forme :

$$\partial_t^2 \mathbb{M} \underline{\mathbf{E}}(t) + c^2 \mathbb{K} \underline{\mathbf{E}}(t) = 0,$$

où $\underline{\mathbf{E}}(t)$ est le vecteur des valeurs $\mathcal{E}_{i,j,k}$, $\mathbb{M}$ est appelée la matrice de masse et $\mathbb{K}$ la matrice de raideur interne. $\mathbb{M}$ est plus creuse que $\mathbb{K}$. Pour la discrétisation en temps de ce système linéaire, nous avons deux schémas possibles, selon qu'on prenne $\mathbb{K} \underline{\mathbf{E}}(t)$ au pas de temps n (schéma explicite) ou au pas de temps $n + 1$ (schéma implicite) :

$$\begin{aligned} \text{Schéma explicite :} \quad & \mathbb{M}(\underline{\mathbf{E}}^{n+1} - 2\underline{\mathbf{E}}^n + \underline{\mathbf{E}}^{n-1}) + c^2 \Delta t^2 \mathbb{K} \underline{\mathbf{E}}^n = 0 \rightsquigarrow \mathbb{M} \underline{\mathbf{E}}^{n+1} = \underline{\mathbf{L}}_e^n, \\ \text{Schéma implicite :} \quad & \mathbb{M}(\underline{\mathbf{E}}^{n+1} - 2\underline{\mathbf{E}}^n + \underline{\mathbf{E}}^{n-1}) + c^2 \Delta t^2 \mathbb{K} \underline{\mathbf{E}}^{n+1} = 0 \rightsquigarrow (\mathbb{M} + c^2 \Delta t^2 \mathbb{K}) \underline{\mathbf{E}}^{n+1} = \underline{\mathbf{L}}_i^n. \end{aligned}$$

Dans le premier cas, le pas d'espace et le pas de temps sont liés par la condition de Courant-Friedrichs-Lewy (appelée condition CFL), qui assure la stabilité du schéma et qui s'écrit de la façon suivante :

$$\Delta t \leq \frac{K}{c},$$

où K dépend de la finesse du maillage et de l'ordre d'approximation choisi. Ainsi, on doit diminuer simultanément le pas spatial et le pas temporel. Dans certains cas, la matrice de masse $\mathbb{M}$ est équivalente à une matrice diagonale, ce qui réduit bien évidemment le coût de calcul. Dans [45], G. Cohen développe différents types d'éléments finis permettant la réduction de la matrice de masse interne, en $2D$ et $3D$ (voir aussi l'article de G. Cohen et al. [46] pour les triangles).

Dans le second cas, on peut choisir indépendamment le pas spatial et le pas temporel, mais il faut inverser une matrice moins creuse.

1.6.2 Éléments finis de Nédélec

Les degrés de liberté des éléments finis de Nédélec, introduits en 1980 dans [87] (et améliorés dans [88]), sont portés par les arêtes, ce qui permet de conserver les propriétés des matériaux discontinus de façon transparente (conservation de la composante tangentielle du champ électrique). Cependant, par leur nature même, les éléments finis de Nédélec sont conformes dans $\mathcal{H}(\mathbf{rot}, \Omega)$ seulement. Ce sont les espaces fonctionnels qui interviennent naturellement dans la formulation variationnelle, mais les champs obtenus ne sont pas continus. Ainsi, il faut manipuler ces éléments finis avec prudence et adresse lorsqu'il s'agit de résoudre le système d'équations couplées Maxwell-Vlasov.

De plus, pour obtenir une bonne approximation de la solution, il faut soit raffiner localement le maillage près des arêtes et des coins, soit adapter le maillage et le degré des éléments finis par des techniques *hp* (voir le paragraphe 1.6.4). Dans [82], P. Monk détaille la résolution de différents problèmes électromagnétiques harmoniques par les éléments finis d'arêtes.

Pour les problèmes instationnaires, la difficulté des éléments finis de Nédélec consiste à réduire la matrice de masse. Dans [47], G. Cohen et P. Monk proposent de coupler les éléments finis d'arêtes avec une méthode d'éléments finis discontinus pour condenser la matrice de masse. Cette méthode est modifiée dans [48] pour les milieux anisotropes. Dans [77], P. Lacoste propose une autre nouvelle technique, pour laquelle on doit ajouter des éléments au second membre.

Notons que dans [87], J.-C. Nédélec décrit aussi des éléments finis de faces, introduits initialement en 2D par P.-A. Raviart et J.-M. Thomas [99, 98]. Les degrés de liberté sont des flux moyens à travers les faces, ce qui permet de conserver la composante normale du champ : ces éléments finis sont conformes dans $\mathcal{H}(\text{div}, \Omega)$. Ils sont parfois utilisés pour discrétiser le champ magnétique.

1.6.3 Éléments finis discontinus

La méthode de Galerkin discontinue, introduite dans les années soixante-dix pour simuler le transport de neutrons, consiste à approcher le champ électromagnétique avec des fonctions polynômiales discontinues. Des termes de sauts surfaciques sur les faces des éléments apparaissent alors dans la formulation variationnelle du problème. En pratique, cette méthode est bien adaptée aux méthodes *hp*, et aux cas où le milieu matériel est inhomogène. Dans [91] I. Perugia et al. étudient la discrétisation des équations de Maxwell harmoniques par les éléments finis de Galerkin discontinus. La contrainte de divergence est dualisée. Une analyse d'erreur a priori du problème non-mixte en milieu homogène est faite dans [70]. D'autre part, I. Perugia et D. Schötzau étudient le problème régularisé (ajout d'un terme en div-div) dans [90].

1.6.4 Éléments finis *hp*

La méthode des éléments finis *hp* consiste à faire varier localement à la fois le pas du maillage h et l'ordre d'approximation p . Par exemple, loin des singularités géométriques, la solution est régulière, on utilise un maillage grossier et un ordre d'approximation élevé pour assurer la précision. En particulier, on approche très correctement les fonctions analytiques. En revanche, près des coins ou des arêtes rentrants, on utilise un maillage fin et des éléments finis d'ordre bas. La programmation de cette méthode est assez complexe. Dans [92, 93], W. Rachowicz et L. Demkowicz décrivent comment programmer cette méthode pour les éléments finis d'arêtes (voir aussi [60]). L'étude de la convergence des éléments finis pour la méthode de régularisation à poids en est faite par M. Costabel et al. dans [56].

1.6.5 Méthodes spectrales

La solution des équations de Maxwell est approchée par des polynômes de haut degré N . Le domaine est partagé en K sous-domaines Ω_k . On obtient les points de discrétisation, et les poids associés, à l'aide des zéros de la première dérivée du polynôme de Legendre de degré N : l'élément de base est un cube, éventuellement déformé, et chaque noeud est projeté par translation et homothétie dans Ω_k dans chaque direction. Il y a donc N points de discrétisations dans chaque direction par sous-domaine. Dans [15], F. Ben Belgacem et C. Bernardi établissent des estimations d'erreur pour le cas instationnaire (la discrétisation en temps se fait à l'aide d'un schéma saute-mouton), avec des conditions aux limites absorbantes sur une partie de la frontière. L'avantage de cette méthode est que l'ordre de convergence n'est limité que par la régularité de la solution exacte. Cette méthode

est moins bien adaptée au calcul de champ électromagnétique intense, intrinsèquement peu régulier au voisinage des singularités géométriques.

1.6.6 Obtention d'un domaine borné

Lorsque le domaine Ω , dans lequel on étudie l'évolution du champ électromagnétique, change de forme ou se déplace au cours du temps, il est intéressant d'utiliser la méthode des domaines fictifs, qui consiste à considérer le problème dans un domaine $\square$ fixe contenant Ω . Les conditions aux limites sur Ω sont alors prises en compte avec un multiplicateur de Lagrange. Cette méthode est aussi utilisée pour les problèmes de contrôles. Dans [57], W. Dahmen et al. montrent que les formulations variationnelles mixtes issues des problèmes harmoniques et temporelles sont bien posées dans $\mathcal{H}(\mathbf{rot}, \square)$.

Pour les problèmes de diffraction (voir le paragraphe 1.3.6), la difficulté est réduire l'étude à un domaine borné. Dans [105], T. Van et A. Wood étudient le problème de diffraction par une cavité d'ouverture Γ , prise dans un plan parfaitement conducteur. L'espace est borné en couplant la solution du demi-espace supérieur à celle de la cavité au niveau de la frontière Γ . Les auteurs utilisent les éléments finis d'arêtes pour la discrétisation en espace et les différences finies pour la discrétisation en temps. Ils montrent que leur problème mixte est bien posé, et donnent des estimations d'erreurs.

Chapitre 2

Intérêt de l'étude du problème bidimensionnel

Pourquoi s'intéresser au problème bidimensionnel avant d'étudier le problème tridimensionnel ? D'une part, cela permet de baliser les difficultés numériques liées aux conditions limites, à la prise en compte des discrétisations temporelle et spatiales simultanées (cas instationnaire). D'autre part, on détermine ainsi les méthodes efficaces de résolution. En particulier, on peut étudier l'influence des singularités géométriques sur la qualité de l'approximation. Enfin, les équations bidimensionnelles permettent de modéliser des objets invariants selon une direction tels que les guides d'onde, ou symétrique selon un plan, comme les filtres à stubs. Ceci permet de préparer la résolution des cas $2D^{\frac{1}{2}}$, moins coûteuse que la résolution $3D$.

2.1 Le guide d'onde

Un guide d'onde est une cavité ouverte en ses extrémités, dont le bord est un conducteur parfait. Dans le cas où cette cavité est prismatique, le guide est modélisé par le domaine $\Omega = \omega \times \mathbb{R}$, où ω est un polygone (voir la figure 2.1) et représente une section du guide d'onde. La frontière $\partial\Omega$ est un conducteur parfait, sur la frontière $\partial\Omega$, $\mathcal{E}$ satisfait l'équation (1.28). Considérons une onde harmonique plane, se déplaçant le long de l'axe z . Décomposons le champ en une partie transverse et une partie longitudinale :

$$\begin{cases} \mathcal{E}(x, y, z; t) &= (\mathbf{E}(x, y) + E_z(x, y) \mathbf{e}_3) e^{i(wt - kz)}, \\ \mathcal{H}(x, y, z; t) &= (\mathbf{H}(x, y) + H_z(x, y) \mathbf{e}_3) e^{i(wt - kz)}, \end{cases}$$

où $\mathbf{E}$ et $\mathbf{H}$ sont dans le plan $(\mathbf{e}_1, \mathbf{e}_2)$.

Supposons que l'intérieur du guide d'onde soit le vide : il n'y a pas de charges, donc ρ et $\mathcal{J}$ sont nuls, et $\varepsilon = \varepsilon_0$, $\mu = \mu_0$. Dans le domaine bidimensionnel ω , $\mathbf{E}$ satisfait les équations suivantes :

$$\text{rot } \mathbf{E} = -i w \mu_0 \mathbf{H}_z, \quad (2.1)$$

$$\text{div } \mathbf{E} = i k E_z, \quad (2.2)$$

$$\mathbf{E} \cdot \boldsymbol{\tau}|_{\partial\omega} = 0. \quad (2.3)$$

Pour une onde transverse électrique (mode T. E.), $E_z = 0$, et la divergence de $\mathbf{E}$ est nulle. Pour une onde transverse magnétique (mode T. M.), $H_z = 0$, et le rotationnel de $\mathbf{E}$ est nul.

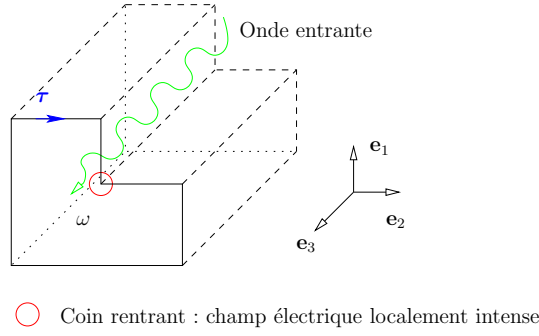


FIG. 2.1 – Modélisation d'un guide d'onde.

2.2 La ligne microruban

La ligne microruban est un guide d'onde particulier, utilisée en microélectronique pour confectionner des circuits planaires (miniaturisation) réalisant des fonctions données (filtrage, amplification, etc). Il est constitué d'un plan de masse parfaitement conducteur sur lequel est déposé un substrat diélectrique, sur la surface duquel est posé une ligne conductrice. L'ensemble est enfermé dans un boîtier (blindage). Le champ électromagnétique est guidé dans le substrat, entre le plan de masse et la ligne. La permittivité diélectrique et la perméabilité magnétique ne dépendent alors que des directions transverses au guide : $\varepsilon = \varepsilon(x, y)$ et $\mu = \mu(x, y)$. Le champ électromagnétique se propage selon la direction z . On peut donc se ramener à l'étude d'une section transverse de la ligne. Dans sa thèse [94], K. Ramdani fait l'étude de trois modèles pour une ligne supraconductrice (figure 2.2). Il propose des formulations variationnelles régularisées.

Une première approche consiste à assimiler la ligne conductrice à un conducteur parfait. La difficulté consiste à borner le domaine de calcul.

Le modèle d'impédance permet de prendre en compte plus précisément de la ligne conductrice. La difficulté est de gérer la condition de saut discontinue de la composante normale du déplacement électrique à la traversée de la ligne.

Dans le cas le plus réaliste, le modèle de London, la permittivité diélectrique ε n'est pas partout de signe positif, et dépend de la pulsation w , ce qui cause des problèmes de coercivité et de linéarité.

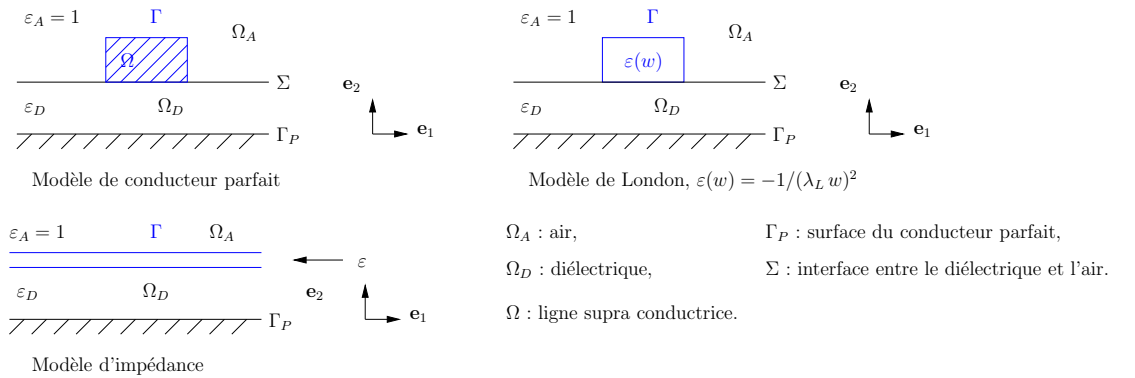


FIG. 2.2 – Modélisation d'une ligne supraconductrice.

2.3 Le filtre à stubs

Un filtre à stubs (figure 2.3) est un guide à l'apparence d'un peigne, où ne passent que certaines fréquences, qui dépendent de la taille des dents. Pour ces fréquences particulières, tout se passe comme si le filtre était fermé transversalement et devenait un simple guide d'onde de section rectangulaire : rien ne passe dans les dents. Ce guide est éclairé par une onde dont le signal temporel, dirigé selon x , se situe dans la plage des fréquences passantes.

Soit L la largeur du filtre. Le filtre est représenté par le domaine $\Omega = \omega \times]0, L[$, où ω est un polygone. La frontière $\partial\Omega$ est composée de Γ_C , le conducteur parfait, et Γ_A , la frontière artificielle qui borne le domaine. $\Gamma_A = \Gamma_1 \cup \Gamma_2$, où $\Gamma_1 := \partial\Omega \cap \{x = 0\}$ (condition aux limites d'onde entrante), $\Gamma_2 := \partial\Omega \cap \{x = A\}$ (condition aux limites absorbante).

Soit $u \in L^2(\Omega)$. De par la nature tensorielle du domaine, on peut décomposer u en une série de Fourier telle que : $u = \sum_{k \geq 0} u^k(x, y; t) \sin\left(k\pi \frac{z}{L}\right)$ ou bien $u = \sum_{k \geq 0} u^k(x, y; t) \cos\left(k\pi \frac{z}{L}\right)$

avec $\|u\|_0^2 = \frac{L}{2} \sum_{k \geq 0} \|u^k\|_{0,\omega}^2$. Dans notre cas, d'après les conditions aux limites portant sur les composantes tangentielles pour le champ électrique et normale pour le champ magnétique, on a en particulier : $E_x = E_y = 0$ et $H_z = 0$ sur $\Gamma_z := (\Gamma_C \cap \{z = 0\}) \cup (\Gamma_C \cap \{z = L\})$. On choisit donc de décomposer les champs électrique et magnétique en une partie transverse et une partie longitudinale selon :

$$\begin{cases} \mathcal{E}(x, y, z; t) = \sum_{k \geq 0} \mathbf{E}^k(x, y; t) \sin\left(k\pi \frac{z}{L}\right) + E_z^k(x, y; t) \cos\left(2k\pi \frac{z}{L}\right) \mathbf{z}, \\ \mathcal{H}(x, y, z; t) = \sum_{k \geq 0} \mathbf{H}^k(x, y; t) \cos\left(k\pi \frac{z}{L}\right) + H_z^k(x, y; t) \sin\left(k\pi \frac{z}{L}\right) \mathbf{z}. \end{cases}$$

$\mathbf{E}$ et $\mathbf{H}$ sont dans le plan $(\mathbf{e}_1, \mathbf{e}_2)$. De même, la densité électrique est décomposée en une série de Fourier, par exemple :

$$\rho = \sum_{k \geq 0} \rho^k(x, y; t) \sin\left(k\pi \frac{z}{L}\right).$$

Remarque 2.1 ρ est seulement L^2 . Le choix de la décomposition de ρ en série de sin plutôt qu'en série de cos permet l'identification mode à mode.

Pour le problème quasi-électrostatique, on obtient alors une série de problèmes variationnels posés dans ω et détaillés dans la section 8.6 de la partie III, de la forme :

Trouver $\mathbf{E}^k \in \mathbf{X}$ et $E_z^k \in X$ tels que,

$$\begin{aligned} (\mathbf{E}^k, \mathbf{F})_{\mathbf{X}} + \frac{k^2 \pi^2}{L^2} (\mathbf{E}^k, \mathcal{F})_{0,\omega} &= \mathcal{L}_k(\mathbf{F}), \quad \forall \mathbf{F} \in \mathbf{X}, \\ (E_z^k, v)_X + \frac{k^2 \pi^2}{L^2} (E_z^k, v)_{0,\omega} &= l_k(v), \quad \forall v \in X, \end{aligned}$$

où $\mathbf{X}$ est l'espace fonctionnel vectoriel 2D choisi pour la discrétisation des $\mathbf{E}^k$, X est l'espace fonctionnel vectoriel 1D choisi pour la discrétisation des E^k ; $\mathcal{L}_k$ est une forme linéaire de $\mathbf{X}$ dans $\mathbb{R}$, et l_k est une forme linéaire de X dans $\mathbb{R}$. Ces formes linéaires dépendent des ρ^k .

On ne peut résoudre qu'un nombre fini K de problèmes. On est alors en mesure de reconstituer une partie du champ électrique :

$$\mathcal{E}^K = \sum_{k=0}^K \mathbf{E}^k(x, y; t) \sin\left(k\pi \frac{z}{L}\right) + E_z^k(x, y; t) \cos\left(k\pi \frac{z}{L}\right) \mathbf{z}$$

Notons que plus K est élevé, meilleure est l'approximation.

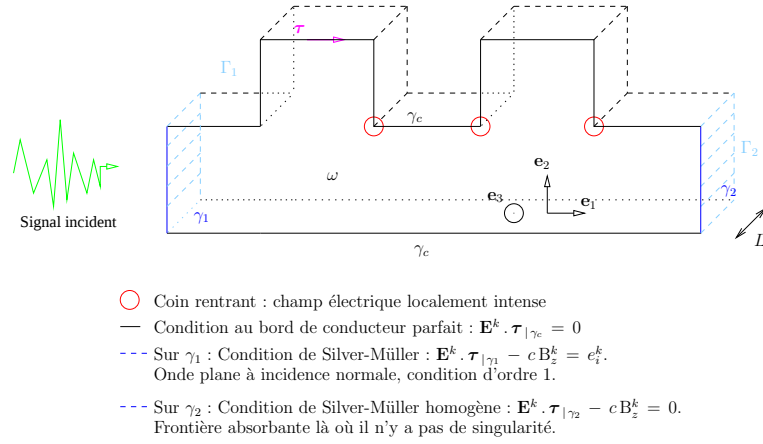


FIG. 2.3 – Modélisation d'un filtre à stubs.

Deuxième partie

Le problème bidimensionnel

Chapitre 1

Notations et résultats préliminaires 2D

1.1 Notations relatives au domaine d'étude

Le domaine d'étude $\omega \subset \mathbb{R}^2$ est un polygone de frontière lipschitzienne $\partial\omega$. On suppose en particulier que ω est simplement connexe et que $\partial\omega$ est connexe. Dans le cas général, on renvoie le lecteur à [5]. On utilisera les notations suivantes :

$(\mathbf{e}_1, \mathbf{e}_2) = (\mathbf{x}, \mathbf{y})$: base orthonormale canonique de $\mathbb{R}^2$.

$(x_1, x_2) = (x, y)$: coordonnées d'un point de $\mathbb{R}^2$.

$\mathbf{u} = (u_1, u_2) = (u_x, u_y)$: composantes d'un champ de vecteurs 2D.

$\mathbf{u} \cdot \mathbf{v} = u_1 v_1 + u_2 v_2 \in \mathbb{R}$: produit scalaire 2D entre $\mathbf{u}$ et $\mathbf{v}$.

$\boldsymbol{\nu} = (\nu_1, \nu_2)$: vecteur unitaire sortant normal à $\partial\omega$.

$\boldsymbol{\tau} = (\nu_2, -\nu_1)$: vecteur tangentiel associé, tel que $(\boldsymbol{\tau}, \boldsymbol{\nu})$ forme une base orthonormée directe.

$u_\nu = \mathbf{u} \cdot \boldsymbol{\nu}|_{\partial\omega}$: composante normale à $\partial\omega$ du vecteur $\mathbf{u}$.

$u_\tau = \mathbf{u} \cdot \boldsymbol{\tau}|_{\partial\omega}$: composante tangentielle à $\partial\omega$ du vecteur $\mathbf{u}$.

La frontière $\partial\omega$ est séparée en K arêtes ouvertes A_k , $k \in \{1, \dots, K\}$: $\partial\omega = \cup_k \overline{A_k}$. ω étant un polygone, ω est dit singulier lorsqu'il n'est pas convexe, c'est-à-dire lorsqu'il existe un ou plusieurs coin(s) rentrant(s). Considérons le cas où le domaine ω présente N_{cr} coins rentrants $(O_i)_{i=1, \dots, N_{cr}}$, d'angles $\Theta_i = \frac{\pi}{\alpha_i}$, $\alpha_i \in]1/2, 1[$, $i \in \{1, \dots, N_{cr}\}$. On pose :

$$\boxed{\alpha = \min_i \alpha_i .} \quad (1.1)$$

Soient A_i^0 et A_i^Θ les arêtes du coin rentrant O_i , et (r_i, θ_i) , les coordonnées polaires associées, avec $\theta_i \in [0, \Theta_i]$ tel que : $\theta_i = 0$ sur A_i^0 , et $\theta = \Theta_i$ sur A_i^Θ (voir la figure 1.1). Soient $\mathbf{e}_{r_i}$ et $\mathbf{e}_{\theta_i}$ les

vecteurs de base associés aux coordonnées polaires. On remarque que :

$$\begin{cases} \tau|_{A_i^0} = -\mathbf{e}_{r_i}, \\ \nu|_{A_i^0} = -\mathbf{e}_{\theta_i=0}, \end{cases} \quad \text{et} \quad \begin{cases} \tau|_{A_i^\Theta} = \mathbf{e}_{r_i}, \\ \nu|_{A_i^\Theta} = \mathbf{e}_{\theta_i=\Theta_i}. \end{cases} \quad (1.2)$$

Définitions 1.1 Pour tout $i \in \{1, \dots, N_{cr}\}$, nous appellerons η_i , une fonction de troncature $C^\infty(\omega)$, ne dépendant que de r_i telle que $\eta_i(r_i) = 1$ pour $r_i \leq \epsilon_i/2$, et $\eta_i(r_i) = 0$ pour $r_i \geq \epsilon_i$, où ϵ_i est un petit nombre strictement positif donné.

Nous appellerons $\gamma_i = A_i^0 \cup A_i^\Theta \cup O_i$: les arêtes du coin rentrant O_i et le coin rentrant O_i lui-même.

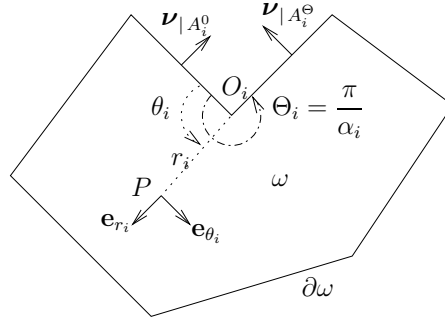


FIG. 1.1 – Coordonnées polaires par rapport à un coin rentrant, et données associées.

1.2 Opérateurs bidimensionnels

1.2.1 Notations

On utilisera les notations suivantes :

t : la variable temporelle.

$\partial_t(.) = \frac{\partial(.)}{\partial t}$: dérivée partielle par rapport au temps.

$\partial_t^m(.) = \frac{\partial^m(.)}{\partial t^m}$: dérivée partielle $m^{\text{ième}}$ par rapport au temps.

$\partial_i(.) = \frac{\partial(.)}{\partial x_i}$: dérivée partielle suivant x_i .

$\partial_i^m(.) = \frac{\partial^m(.)}{\partial x_i^m}$: dérivée partielle $m^{\text{ième}}$ suivant x_i .

$\partial_{\mathbf{x}}^m v = \sum_{m_1, m_2 \geq 0 \mid m_1 + m_2 = m} \frac{\partial^m v}{\partial x_1^{m_1} \partial x_2^{m_2}}.$

$\mathbf{grad} v = (\partial_1 v, \partial_2 v)$: gradient de v .

$\partial_\nu v|_{\partial\omega} = \mathbf{grad} v \cdot \nu|_{\partial\omega}$: dérivée de v sur $\partial\omega$ dans la direction normale à $\partial\omega$.

$\partial_\tau v|_{\partial\omega} = \mathbf{grad} v \cdot \tau|_{\partial\omega}$: dérivée de v sur $\partial\omega$ dans la direction tangentielle à $\partial\omega$, aussi appelé gradient tangentiel.

$\mathbf{rot} v = (\partial_2 v, -\partial_1 v)$: le rotationnel vectoriel de v .

$\Delta v = \partial_1^2 v + \partial_2^2 v$: le laplacien de v .

$\operatorname{div} \mathbf{u} = \partial_1 u_1 + \partial_2 u_2$: la divergence de $\mathbf{u}$.

$\operatorname{rot} \mathbf{u} = \partial_1 u_2 - \partial_2 u_1$: le rotationnel scalaire de $\mathbf{u}$.

$\Delta \mathbf{u} = (\Delta u_1, \Delta u_2)$: le laplacien vectoriel de $\mathbf{u}$.

$\mathbf{grad} : \mathbf{u} = \begin{pmatrix} \partial_1 u_1 & \partial_2 u_1 \\ \partial_1 u_2 & \partial_2 u_2 \end{pmatrix}$: le gradient matriciel de $\mathbf{u}$.

1.2.2 Relations entre opérateurs bidimensionnels

$$\operatorname{div} \mathbf{rot} (.) = 0, \quad (1.3)$$

$$\operatorname{rot} \mathbf{grad} (.) = 0, \quad (1.4)$$

$$\operatorname{div} \mathbf{grad} (.) = \Delta (.), \quad (1.5)$$

$$-\operatorname{rot} \mathbf{rot} (.) = \Delta (.), \quad (1.6)$$

$$-\mathbf{rot} \operatorname{rot} (.) + \mathbf{grad} \operatorname{div} (.) = \Delta (.). \quad (1.7)$$

Propriété 1.2 Soit u une fonction régulière. On a la relation suivante :

$$\mathbf{rot} u \cdot \boldsymbol{\tau}|_{\partial\omega} = \partial_\nu u|_{\partial\omega}.$$

DÉMONSTRATION. D'après les définitions des opérateurs $\mathbf{rot}$, $\mathbf{grad}$, et des vecteurs $\boldsymbol{\nu}$ et $\boldsymbol{\tau}$, on a :
 $\mathbf{rot} u \cdot \boldsymbol{\tau}|_{\partial\omega} = \partial_2 u \tau_1|_{\partial\omega} - \partial_1 u \tau_2|_{\partial\omega} = \partial_2 u \nu_2|_{\partial\omega} + \partial_1 u \nu_1|_{\partial\omega} = \mathbf{grad} u \cdot \boldsymbol{\nu}|_{\partial\omega}.$

□

1.3 Espaces de Hilbert usuels et leurs normes associées

De façon générale, les espaces désignés par des lettres majuscules en caractère gras sont des espaces de champs vectoriels, et les espaces désignés par des lettres majuscules italiques en caractère normal sont des espaces de champs scalaires. Soit $H(.)$ un espace de champs scalaires. On notera $\mathbf{H}(.) = H(.)^2$, l'espace de champs vectoriels correspondant.

$d\omega$ désigne la mesure de l'ouvert ω et $d\sigma$ désigne la mesure de l'ouvert $\partial\omega$.

$\mathcal{D}(\omega)$ est l'espace des fonctions $C^\infty(\omega)$, à support compact dans ω . On appelle $\mathcal{D}'(\omega)$ son dual, l'espace des distributions (la définition du dual d'un espace est donnée en 1.6).

1.3.1 Espaces des champs scalaires

$$L^2(\omega) = \left\{ u \text{ mesurable sur } \omega : \int_{\omega} u^2 d\omega < \infty \right\}, \quad \|u\|_0 = \left(\int_{\omega} u^2 d\omega \right)^{1/2}.$$

$$L^2_{loc}(\omega) = \left\{ u \text{ mesurable sur } \omega : \forall \omega_c | \bar{\omega}_c \subset \omega, \int_{\omega_c} u^2 d\omega < \infty \right\}.$$

$$L^2_0(\omega) = \left\{ u \in L^2(\omega) : \int_{\omega} u d\omega = 0 \right\}.$$

$$L^2_{\Delta}(\omega) = \{ u \in L^2(\omega) : \Delta u \in L^2(\omega) \}, \quad \|u\|_{0,\Delta} = (\|u\|_0^2 + \|\Delta u\|_0^2)^{1/2}.$$

$$H^1(\omega) = \{ u \in L^2(\omega) : \mathbf{grad} u \in \mathbf{L}^2(\omega) \}, \quad \|u\|_{H^1} = (\|u\|_0^2 + \|\mathbf{grad} u\|_0^2)^{1/2}.$$

$$H^1_0(\omega) = \{ u \in H^1(\omega) : u|_{\partial\omega} = 0 \}, \quad \|u\|_{H^1_0} = \|\mathbf{grad} u\|_0.$$

$$H^1_{\Delta}(\omega) = \{ u \in H^1(\omega) : \Delta u \in L^2(\omega) \}, \quad \|u\|_{H^1_{\Delta}} = (\|u\|_{H^1}^2 + \|\Delta u\|_0^2)^{1/2}.$$

$$H^2(\omega) = \{ u \in H^1(\omega) : \mathbf{grad} u \in \mathbf{H}^1(\omega) \}, \quad \|u\|_{H^2} = (\|u\|_0^2 + \|\mathbf{grad} u\|_{\mathbf{H}^1}^2)^{1/2}.$$

L'équivalence entre la norme du graphe et la semi-norme dans $H^1_0(\omega)$ due à l'inégalité de Poincaré. $H^{-1}(\omega)$ est le dual de $H^1_0(\omega)$.

Φ_N et Φ_D sont les espaces des solutions du Laplacien :

$$\Phi_N = \{ u \in H^1(\omega) \cap L^2_0(\omega) : \Delta u \in L^2(\omega), \partial_{\nu} u|_{\partial\omega} = 0 \},$$

$$\Phi_D = \{ u \in H^1_0(\omega) : \Delta u \in L^2(\omega) \}.$$

Dans Φ_N et Φ_D , la semi-norme est équivalente à la norme du graphe : $\|u\|_{\Phi_{D,N}} = \|u\|_{\Phi} := \|\Delta u\|_0$ (inégalité de Poincaré-Friedrichs pour Φ_D ; inégalité de Poincaré-Wirtinger et théorème de Lax-Milgram pour Φ_N).

1.3.2 Espaces de champs vectoriels

$$\mathbf{L}^2(\omega) = \{ \mathbf{u} \in L^2(\omega)^2 \}, \quad \|\mathbf{u}\|_0 = (\|u_1\|_0^2 + \|u_2\|_0^2)^{1/2}.$$

$$\mathbf{H}^1(\omega) = H^1(\omega)^2, \quad \|\mathbf{u}\|_{\mathbf{H}^1} = (\|\mathbf{u}\|_0^2 + \|\mathbf{grad} : \mathbf{u}\|_0^2)^{1/2}.$$

$$\mathbf{H}(\text{rot}, \omega) = \{ \mathbf{u} \in \mathbf{L}^2(\omega) : \text{rot } \mathbf{u} \in L^2(\omega) \}, \quad \|\mathbf{u}\|_{\mathbf{H}(\text{rot})} = (\|\mathbf{u}\|_0^2 + \|\text{rot } \mathbf{u}\|_0^2)^{1/2}.$$

$$\mathbf{H}_0(\text{rot}, \omega) = \{ \mathbf{u} \in \mathbf{H}(\text{rot}, \omega) : \mathbf{u}_{\tau} = 0 \}.$$

$$\mathbf{H}(\text{rot}^0, \omega) = \{ \mathbf{u} \in \mathbf{L}^2(\omega) : \text{rot } \mathbf{u} = 0 \}.$$

$$\mathbf{H}_0(\text{rot}^0, \omega) = \{ \mathbf{u} \in \mathbf{L}^2(\omega) : \text{rot } \mathbf{u} = 0 \text{ et } \mathbf{u}_{\tau} = 0 \}.$$

Proposition 1.3 On a : $\text{grad}(H_0^1(\omega)) = \mathbf{H}_0(\text{rot}^0, \omega)$.

DÉMONSTRATION. Voir [63], (prop. 10.1.2, p. 174). □

$$\mathbf{H}(\text{div}, \omega) = \{\mathbf{u} \in \mathbf{L}^2(\omega) : \text{div } \mathbf{u} \in L^2(\omega)\}, \quad \|\mathbf{u}\|_{\mathbf{H}(\text{div})} = (\|\mathbf{u}\|_0^2 + \|\text{div } \mathbf{u}\|_0^2)^{1/2}.$$

$$\mathbf{H}_0(\text{div}, \omega) = \{\mathbf{u} \in \mathbf{H}(\text{div}, \omega) : \mathbf{u}_\nu = 0\}.$$

$$\mathbf{H}(\text{div}^0, \omega) = \{\mathbf{u} \in \mathbf{L}^2(\omega) : \text{div } \mathbf{u} = 0\}.$$

$$\mathbf{H}_0(\text{div}^0, \omega) = \{\mathbf{u} \in \mathbf{L}^2(\omega) : \text{div } \mathbf{u} = 0 \text{ et } \mathbf{u}_\nu = 0\}.$$

$\mathbf{X}_\mathbf{E}$ et $\mathbf{X}_\mathbf{H}$ sont les espaces des solutions des équations de Maxwell quasi-statiques :

$$\mathbf{X}_\mathbf{E} = \{\mathbf{u} \in \mathbf{H}(\text{rot}, \omega) \cap \mathbf{H}(\text{div}, \omega) : \mathbf{u}_\tau \in L^2(\partial\omega)\}.$$

$$\mathbf{X}_\mathbf{H} = \{\mathbf{u} \in \mathbf{H}(\text{rot}, \omega) \cap \mathbf{H}(\text{div}, \omega) : \mathbf{u}_\nu \in L^2(\partial\omega)\}.$$

Proposition 1.4 Dans $\mathbf{X}_\mathbf{E}$ et $\mathbf{X}_\mathbf{H}$, la semi-norme est équivalente à la norme du graphe, d'où :

$$\|\mathbf{u}\|_{\mathbf{X}_\mathbf{E}} = (\|\text{rot } \mathbf{u}\|_0^2 + \|\text{div } \mathbf{u}\|_0^2 + \|\mathbf{u}_\tau\|_{0, \partial\omega}^2)^{1/2}, \quad \|\mathbf{u}\|_{\mathbf{X}_\mathbf{H}} = (\|\text{rot } \mathbf{u}\|_0^2 + \|\text{div } \mathbf{u}\|_0^2 + \|\mathbf{u}_\nu\|_{0, \partial\omega}^2)^{1/2}.$$

DÉMONSTRATION. On applique la preuve du théorème 8.5 de la partie III au cas $2D$. □

On considère les sous-espaces de $\mathbf{X}_\mathbf{E}$ suivants :

$$\mathbf{X}_\mathbf{E}^0 = \{\mathbf{u} \in \mathbf{X}_\mathbf{E} : \mathbf{u}_\tau = 0\}, \quad \|\mathbf{u}\|_{\mathbf{X}_\mathbf{E}^0} = (\|\text{rot } \mathbf{u}\|_0^2 + \|\text{div } \mathbf{u}\|_0^2)^{1/2}.$$

$$\mathbf{V}_\mathbf{E} = \{\mathbf{u} \in \mathbf{X}_\mathbf{E}^0 : \text{div } \mathbf{u} = 0\}, \quad \|\mathbf{u}\|_{\mathbf{V}_\mathbf{E}} = \|\text{rot } \mathbf{u}\|_0.$$

$$\mathbf{W}_\mathbf{E} = \{\mathbf{u} \in \mathbf{X}_\mathbf{E}^0 : \text{rot } \mathbf{u} = 0\}, \quad \|\mathbf{u}\|_{\mathbf{W}_\mathbf{E}} = \|\text{div } \mathbf{u}\|_0.$$

On considère les sous-espaces de $\mathbf{X}_\mathbf{H}$ suivants :

$$\mathbf{X}_\mathbf{H}^0 = \{\mathbf{u} \in \mathbf{X}_\mathbf{H} : \mathbf{u}_\nu = 0\}, \quad \|\mathbf{u}\|_{\mathbf{X}_\mathbf{H}^0} = (\|\text{rot } \mathbf{u}\|_0^2 + \|\text{div } \mathbf{u}\|_0^2)^{1/2}.$$

$$\mathbf{V}_\mathbf{H} = \{\mathbf{u} \in \mathbf{X}_\mathbf{H}^0 : \text{div } \mathbf{u} = 0\}, \quad \|\mathbf{u}\|_{\mathbf{V}_\mathbf{H}} = \|\text{rot } \mathbf{u}\|_0.$$

$$\mathbf{W}_\mathbf{H} = \{\mathbf{u} \in \mathbf{X}_\mathbf{H}^0 : \text{rot } \mathbf{u} = 0\}, \quad \|\mathbf{u}\|_{\mathbf{W}_\mathbf{H}} = \|\text{div } \mathbf{u}\|_0.$$

Notons que dans [63] (pp. 48-49), on a une preuve directe de l'équivalence entre la norme du graphe et la semi-norme dans $\mathbf{X}_\mathbf{E}^0$ et dans $\mathbf{X}_\mathbf{H}^0$.

1.3.3 Espaces des traces

$$L^2(\partial\omega) = \left\{ u \text{ mesurable sur } \partial\omega : \int_{\partial\omega} u^2 d\sigma < \infty \right\}, \quad \|u\|_{0, \partial\omega} = \left(\int_{\partial\omega} u^2 d\sigma \right)^{1/2}.$$

$\forall u \in H^1(\omega), u|_{\partial\omega} \in H^{1/2}(\partial\omega)$, où :

$$H^{1/2}(\partial\omega) = \left\{ u \in L^2(\partial\omega) : \int_{\partial\omega} \int_{\partial\omega} \frac{(u(\mathbf{x}) - u(\mathbf{y}))^2}{\|\mathbf{x} - \mathbf{y}\|^2} d\sigma(\mathbf{x}) d\sigma(\mathbf{y}) < \infty \right\}.$$

$\forall u \in H^2(\omega)$, $u|_{\partial\omega} \in H^{3/2}(\partial\omega)$, où [29] :

$$H^{3/2}(\partial\omega) = \left\{ u \in H^1(\partial\omega) : \int_{\partial\omega} \int_{\partial\omega} \frac{|\mathbf{grad} u(\mathbf{x}) \cdot \boldsymbol{\tau} - \mathbf{grad} u(\mathbf{y}) \cdot \boldsymbol{\tau}|^2}{\|\mathbf{x} - \mathbf{y}\|^2} d\sigma(\mathbf{x}) d\sigma(\mathbf{y}) < \infty \right\}.$$

Pour $s = 1$ ou 3 , on désigne par $H^{-s/2}(\partial\omega)$ le dual de $H^{s/2}(\partial\omega)$.

Pour $s = 1$ ou 3 , pour chaque arête A_k du bord $\partial\omega$, on définit les espaces $H_{00}^{s/2}(A_k)$ tels que :

$$H_{00}^{s/2}(A_k) := \{ u \in H^{s/2}(A_k) : \tilde{u} \in H^{s/2}(\partial\omega) \}, \tilde{u} \text{ étant un prolongement de } u \text{ par } 0 \text{ sur } \partial\omega.$$

$$H_{00}^{-s/2}(A_k) \text{ est le dual de } H_{00}^{s/2}(A_k).$$

1.4 Espaces de Sobolev à poids

Supposons que le domaine ω contienne un coin rentrant. Soit r la distance à ce coin rentrant. On définit alors les espaces à poids suivants, pour $\gamma \in \mathbb{R}$ et $l \in \mathbb{N}$:

$$V_\gamma^l(\omega) = \left\{ u \in L_{loc}^2(\omega) : \sum_{|j| \leq l} \int_\omega r^{2(\gamma-l+|j|)} |\partial_{\mathbf{x}}^j u|^2 d\omega < \infty \right\}. \quad (1.8)$$

$V_\gamma^{l-1/2}(\partial\omega)$ est l'espace des traces sur $\partial\omega$ des fonctions de V_γ^l .

Pour $l = 0$, on utilisera aussi la notation : $L_\gamma^2(\omega) := V_\gamma^0(\omega)$. Le produit scalaire de $L_\gamma^2(\omega)$ et la norme associée seront notés :

$$(u, v)_{0, \gamma} := \int_\omega r^{2\gamma} u v d\omega, \quad \|u\|_{0, \gamma} := \left(\int_\omega r^{2\gamma} u^2 d\omega \right)^{1/2}.$$

On considère l'espace vectoriel suivant, avec $0 < \gamma < 1$:

$$\mathbf{X}_{\mathbf{E}, \gamma}^0 = \{ \mathbf{u} \in \mathbf{H}_0(\text{rot}, \omega) : \text{div } \mathbf{u} \in L_\gamma^2(\omega) \}, \quad \|\mathbf{u}\|_{\mathbf{X}_{\mathbf{E}, \gamma}^0} = (\|\text{rot } \mathbf{u}\|_0^2 + \|\text{div } \mathbf{u}\|_{0, \gamma}^2)^{1/2}.$$

1.5 Espaces de Sobolev classiques

On définit ici les espaces $W^{m,p}$, (m entier ou non, p entier) dont on aura besoin pour certaines estimations.

- Pour $m \in \mathbb{N}$ et $p \in \mathbb{N}$, on définit :

$$W^{m,p}(\omega) = \left\{ u \in \mathcal{D}'(\omega) : \sum_{|j| \leq m} \int_\omega |\partial_{\mathbf{x}}^j u|^p d\omega < \infty \right\}, \quad \|u\|_{m,p} = \left(\sum_{|j| \leq m} \int_\omega |\partial_{\mathbf{x}}^j u|^p d\omega \right)^{1/p}.$$

- Pour $s = m + \sigma$, $m \in \mathbb{N}$, $0 < 1 < \sigma$, et $p \in \mathbb{N}$ on définit :

$$W^{s,p}(\omega) = \left\{ u \in W^{m,p}(\omega) : |u|_{s,p} := \left(\sum_{|j|=m} \int_\omega \int_\omega \frac{|\partial_{\mathbf{x}}^j u(\mathbf{x}) - \partial_{\mathbf{x}}^j u(\mathbf{y})|^p}{\|\mathbf{x} - \mathbf{y}\|^{2+\sigma p}} d\mathbf{x} d\mathbf{y} \right)^{1/p} < \infty \right\}.$$

La norme associée est : $\|u\|_{s,p} = (\|u\|_{m,p}^p + |u|_{s,p}^p)^{1/p}$.

On remarque que $H^m(\omega) = W^{m,2}(\omega)$, et on posera de même : $H^s(\omega) = W^{s,2}(\omega)$.

On a le théorème suivant [67] (thm. 1.4.2, p. 12) :

Théorème 1.5 Soit $s \in \mathbb{R}$. Soit $k \in \mathbb{N}$ tel que : $k < s - 1/2$. Alors, pour tout $j \in \{1, \dots, K\}$, l'application :

$$\begin{aligned} H^s(\omega) &\rightarrow \prod_{p=0}^k H^{s-p-1/2}(A_j) \\ u &\mapsto (u|_{A_j}, \partial_{\nu_j} u|_{A_j}, \dots, \partial_{\nu_j}^k u|_{A_j}) \end{aligned}$$

est continue.

Par abus de notation, on écrira lorsque $1 < s < 2$ que pour $u \in H^s(\omega)$, $u|_{\partial\omega} \in H^{s-1/2}(\partial\omega)$.

1.6 Espaces duaux

Définition 1.6 On désignera par H' le dual de H , c'est-à-dire l'ensemble des formes linéaires continues sur H . Le produit scalaire de dualité s'écrit ainsi : $\langle u', u \rangle_{H', H}$. La norme duale sur H' est définie par :

$$\|u'\|_{H'} = \sup_{u \in H, \|u\|_H \leq 1} \langle u', u \rangle_{H', H}. \quad (1.9)$$

On peut montrer le théorème suivant, appelé théorème de représentation de Riesz-Fréchet [27] :

Théorème 1.7 Soit H un Hilbert et H' son dual. Alors pour tout $u' \in H'$, il existe un unique $f \in H$ tel que pour tout $u \in H$: $\langle u', u \rangle_{H', H} = (f, u)_H$.

1.7 Formules d'intégration par parties classiques dans un ouvert

1.7.1 Formules de Green

$$\forall u \in \mathcal{D}(\bar{\omega}), \forall v \in \mathcal{D}(\bar{\omega}), \quad \int_{\omega} \mathbf{grad} u \cdot \mathbf{grad} v \, d\omega + \int_{\omega} \Delta u v \, d\omega = \int_{\partial\omega} \partial_{\nu} u v \, d\sigma. \quad (1.10)$$

$$\forall \mathbf{u} \in \mathcal{D}(\bar{\omega})^2, \forall v \in \mathcal{D}(\bar{\omega}), \quad \int_{\omega} \mathbf{u} \cdot \mathbf{grad} v \, d\omega + \int_{\omega} \operatorname{div} \mathbf{u} v \, d\omega = \int_{\partial\omega} u_{\nu} v \, d\sigma. \quad (1.11)$$

$$\forall \mathbf{u} \in \mathcal{D}(\bar{\omega})^2, \forall v \in \mathcal{D}(\bar{\omega}), \quad \int_{\omega} \mathbf{u} \cdot \mathbf{rot} v \, d\omega - \int_{\omega} \operatorname{rot} \mathbf{u} v \, d\omega = \int_{\partial\omega} u_{\tau} v \, d\sigma. \quad (1.12)$$

1.7.2 Généralisation des formules de Green

Les preuves de ces formules se trouvent dans [65]. L'argument principal est la densité des fonctions régulières dans les espaces de Hilbert considérés. Ici, $H = H^{1/2}(\partial\omega)$.

$$\forall u \in H_{\Delta}^1(\omega), \forall v \in H^1(\omega), \quad \int_{\omega} \mathbf{grad} u \cdot \mathbf{grad} v \, d\omega + \int_{\omega} \Delta u v \, d\omega = \langle \partial_{\nu} u, v \rangle_{H', H}. \quad (1.13)$$

$$\forall \mathbf{u} \in \mathbf{H}(\operatorname{div}, \omega), \forall v \in H^1(\omega), \quad \int_{\omega} \mathbf{u} \cdot \mathbf{grad} v \, d\omega + \int_{\omega} \operatorname{div} \mathbf{u} v \, d\omega = \langle u_{\nu}, v \rangle_{H', H}. \quad (1.14)$$

$$\forall \mathbf{u} \in \mathbf{H}(\operatorname{rot}, \omega), \forall v \in H^1(\omega), \quad \int_{\omega} \mathbf{u} \cdot \mathbf{rot} v \, d\omega - \int_{\omega} \operatorname{rot} \mathbf{u} v \, d\omega = \langle u_{\tau}, v \rangle_{H', H}. \quad (1.15)$$

Chapitre 2

Le problème statique 2D direct continu

2.1 Introduction

Le domaine ω représente l'intérieur *vide* d'un conducteur parfait, qu'on borne si besoin est par une frontière artificielle γ_A ne coupant pas de singularités géométriques. La frontière $\partial\omega$ est dans ce cas composée de deux parties : $\partial\omega = \overline{\gamma_C} \cup \overline{\gamma_A}$, où γ_C est le conducteur parfait. Dans cette configuration, le champ électrique bidimensionnel satisfait $E_\tau = 0$ sur γ_C et $E_\tau = \pm cB_\nu + e^*$ sur γ_A , avec :

- Pour une condition aux limites d'onde entrante : $e^* = e_i$, où e_i est une donnée,
- Pour une condition aux limites absorbante : $e^* = 0$.

Posons $E_\tau = e$. On suppose que e est connu (c'est-à-dire que B_ν et e^* sont connus). Considérons $\mathcal{V}_{\gamma_A}$, un voisinage de γ_A ne contenant pas de singularité.

Localement, $\mathbf{E}|_{\mathcal{V}_{\gamma_A}} \in \mathbf{H}^1(\mathcal{V}_{\gamma_A})$ [36], (rem. 1, p. 562), d'après [5], (thm. 2.9 p. 829 et 2.12 p. 830). De plus, $e|_{\gamma_C} = 0$, donc $e|_{\gamma_C} \in H^{1/2}(\gamma_C)$. D'où, en prolongeant $e|_{\gamma_A}$ par 0 sur γ_C , on a : $e \in H^{1/2}(\partial\omega)$. Ainsi, dans la suite de cette partie, on considère toujours que :

E_τ sur $\partial\omega$ est donné par $e \in H^{1/2}(\partial\omega)$, s'annulant au voisinage des coins rentrants de ω .

Cette hypothèse n'est en aucun cas restrictive, et correspond à une réalité de la modélisation. En effet, comme on l'a mentionné plus haut, on peut toujours placer la frontière artificielle de façon à ne pas couper les singularités géométriques.

Nous étudions le problème quasi-électrostatique bidimensionnel, qui s'écrit mathématiquement de la façon suivante :

Trouver $\mathbf{E} \in \mathbf{L}^2(\omega)$ tel que :

$$\text{rot } \mathbf{E} = f \text{ dans } \omega, f \in L^2(\omega), \quad (2.1)$$

$$\text{div } \mathbf{E} = g \text{ dans } \omega, g \in L^2(\omega), \quad (2.2)$$

$$E_\tau = e \text{ sur } \partial\omega, e \in L^2(\partial\omega), \quad (2.3)$$

Dans le cas électrostatique, $f = 0$. Pour une onde transverse électrique, $g = \rho/\varepsilon_0$.

Pour que le problème soit bien posé, il faut que :

$$\int_\omega f \, d\omega = - \int_{\partial\omega} e \, d\sigma. \quad (2.4)$$

En effet, par la formule d'intégration par parties (1.15), en prenant $\mathbf{u} = \mathbf{E}$ et $v = 1$, on a la relation (2.4).

Ce problème peut être résolu d'une part en utilisant la décomposition de Helmholtz : $\mathbf{E} = -\mathbf{grad} \phi_D + \mathbf{rot} \phi_N$, où ϕ_D satisfait (2.9) et ϕ_N satisfait (2.10). Les potentiels ϕ_D et ϕ_N sont calculés par les éléments finis continus P_k de Lagrange, et $\mathbf{E}$ est approché par les éléments finis discontinus P_{k-1} composante par composante. Le développement de cette méthode fera l'objet de la section 2.2 pour le problème continu et de la section 4.3 pour le problème discrétisé.

D'autre part, le problème (2.1)-(2.3) peut être résolu par les éléments finis continus de Lagrange P_k composante par composante, en calculant $\mathbf{E}$ directement dans $\mathbf{X}_{\mathbf{E}}$. Cette méthode sera détaillée dans la section 2.3 pour le problème continu et dans la section 4.6 pour le problème discrétisé.

Définition 2.1 *Le vecteur $\mathbf{e}$ est un relèvement de la donnée tangentielle e si $\mathbf{e}_\tau = e$. Comme $e \in H^{1/2}(\partial\omega)$, il existe un relèvement régulier de e : $\mathbf{e} \in \mathbf{H}^1(\omega)$ [29] (prop. 2.7, p. 15).*

Proposition 2.2 *Pour $e \in L^2(\partial\omega)$, on peut construire un relèvement $\mathbf{e} \in \mathbf{X}_{\mathbf{E}}$ de e .*

DÉMONSTRATION. Soit $\phi_N^e \in H^1(\omega)$ tel que :

$$\begin{cases} \Delta \phi_N^e = \frac{\int_{\partial\omega} e \, d\sigma}{|\omega|} \text{ dans } \omega, \\ \partial_\nu \phi_N^e|_{\partial\omega} = e \text{ sur } \partial\omega. \end{cases} \quad (2.5)$$

Posons $\mathbf{e} = \mathbf{rot} \phi_N^e$. Alors $\mathbf{e} \in \mathbf{L}^2(\omega)$, $\text{div} \mathbf{e} = 0 \in L^2(\omega)$, et $\text{rot} \mathbf{e} = \frac{-\int_{\partial\omega} e \, d\sigma}{|\omega|} \in L^2(\omega)$, d'après (1.6). La propriété 1.2 permet de déduire que $\mathbf{e} \cdot \boldsymbol{\tau}|_{\partial\omega} = e$. □

Posons $\mathbf{E}^0 = \mathbf{E} - \mathbf{e}$. $\mathbf{E}^0$ satisfait le problème suivant : Trouver $\mathbf{E}^0 \in \mathbf{X}_{\mathbf{E}}^0$ tel que :

$$\text{rot} \mathbf{E}^0 = f^0 \in L^2(\omega), \quad (2.6)$$

$$\text{div} \mathbf{E}^0 = g^0 \in L^2(\omega), \quad (2.7)$$

avec $f^0 = f - \text{rot} \mathbf{e}$, et $g^0 = g - \text{div} \mathbf{e}$. Si on prend $\mathbf{e}$ comme dans la proposition 2.2, alors $g^0 = g$. La condition de compatibilité (2.4) devient :

$$\int_{\omega} f^0 \, d\omega = 0, \quad (2.8)$$

c'est-à-dire $f^0 \in L_0^2(\omega)$. $\mathbf{E}^0$ peut être calculé par la méthode des potentiels, ou par les éléments finis continus de Lagrange P_k et le complément singulier, dans l'espace $\mathbf{X}_{\mathbf{E}}^0$. Nous présentons deux approches possibles de la méthode du complément singulier, exposées dans la section 2.4 pour le problème continu ; et dans les sections 4.8 et 4.9 pour le problème discret. On peut encore calculer $\mathbf{E}^0$ dans l'espace à poids $\mathbf{X}_{\mathbf{E},\gamma}^0$. Cela fait l'objet de la section 2.5 pour le problème continu et dans la section 4.10 pour le problème discret. On reconstruit ensuite $\mathbf{E} = \mathbf{E}^0 + \mathbf{e}$.

2.2 Champ électrique 2D : décomposition de Helmholtz

2.2.1 CL non-homogènes

Soient $\phi_D \in H_0^1(\omega)$ et $\phi_N \in H^1(\omega) \cap L_0^2(\omega)$ les potentiels quasi-électrostatiques tels que :
 - ϕ_D est solution du problème de Dirichlet suivant :
 Trouver $\phi_D \in H_0^1(\omega)$ tel que :

$$-\Delta \phi_D = g \text{ dans } \omega. \quad (2.9)$$

- ϕ_N est solution du problème de Neumann suivant :
 Trouver $\phi_N \in H^1(\omega) \cap L_0^2(\omega)$ tel que :

$$\begin{aligned} -\Delta \phi_N &= f \text{ dans } \omega, \\ \partial_\nu \phi_N|_{\partial\omega} &= e \text{ sur } \partial\omega, \end{aligned} \quad (2.10)$$

où f et e satisfont (2.4).

Proposition 2.3 *Le champ électrique $\mathbf{E}$ satisfaisant (2.1)-(2.3) se décompose de la façon suivante :*

$$\mathbf{E} = -\mathbf{grad} \phi_D + \mathbf{rot} \phi_N. \quad (2.11)$$

DÉMONSTRATION. Posons $\mathbf{E}_D = -\mathbf{grad} \phi_D \in \mathbf{L}^2(\omega)$. D'après (1.4) $\mathbf{rot} \mathbf{E}_D = 0 \in L^2(\omega)$; d'après (1.5) $\mathbf{div} \mathbf{E}_D = g \in L^2(\omega)$; de plus $\mathbf{E}_D \cdot \boldsymbol{\tau}|_{\partial\omega} = 0$. On en déduit que : $\mathbf{E}_D \in \mathbf{X}_{\mathbf{E}}^0$. Posons $\mathbf{E}_N = \mathbf{rot} \phi_N \in \mathbf{L}^2(\omega)$. D'après (1.6) $\mathbf{rot} \mathbf{E}_N = f \in L^2(\omega)$; d'après (1.3), $\mathbf{div} \mathbf{E}_N = 0 \in L^2(\omega)$; et d'après la propriété 1.2, $\mathbf{E}_N \cdot \boldsymbol{\tau}|_{\partial\omega} = \partial_\nu \phi_N|_{\partial\omega} = e$. On en déduit que : $\mathbf{E}_N \in \mathbf{X}_{\mathbf{E}}$, et que : $\mathbf{E}_D + \mathbf{E}_N$ est solution de (2.1)-(2.3). □

On peut réécrire (2.11) ainsi :

$$\begin{aligned} \mathbf{E} &= \mathbf{E}_D + \mathbf{E}_N \text{ où :} \\ \mathbf{E}_D &\in \mathbf{X}_{\mathbf{E}}^0 \text{ est tel que } \mathbf{rot} \mathbf{E}_D = 0, \mathbf{div} \mathbf{E}_D = g, \text{ et :} \\ \mathbf{E}_N &\in \mathbf{X}_{\mathbf{E}} \text{ est tel que } \mathbf{rot} \mathbf{E}_N = f, \mathbf{div} \mathbf{E}_N = 0, \mathbf{E}_{N,\tau}|_{\partial\omega} = e. \end{aligned} \quad (2.12)$$

Comme ϕ_D et ϕ_N sont dans $H^1(\omega)$, on les approchera par les éléments finis continus P_k de Lagrange. Le champ électrique $\mathbf{E}$ sera donc approché par les éléments finis discontinus P_{k-1} , composante par composante. De façon plus précise :

Proposition 2.4 *Le problème (2.9) est équivalent au problème variationnel suivant (FVD) :
 Trouver $\phi_D \in H_0^1(\omega)$ tel que :*

$$\forall u \in H_0^1(\omega), a_D(\phi_D, u) = l_D(u), \quad (2.13)$$

où a_D est la forme bilinéaire symétrique suivante :

$$\begin{aligned} a_D : H_0^1(\omega) \times H_0^1(\omega) &\rightarrow \mathbb{R} \\ (u, v) &\mapsto \int_{\omega} \mathbf{grad} u \cdot \mathbf{grad} v \, d\omega, \end{aligned}$$

et l_D est la forme linéaire dépendant de g suivante :

$$\begin{aligned} l_D : H_0^1(\omega) &\rightarrow \mathbb{R} \\ u &\mapsto \int_{\omega} g u \, d\omega. \end{aligned}$$

a_D est la norme de $H_0^1(\omega)$. D'après le théorème de Lax-Milgram [65], (thm. 1.7, p. 10), le problème (FVD) admet une unique solution $\phi_D \in H_0^1(\omega)$.

Proposition 2.5 *Le problème (2.10) est équivalent au problème variationnel suivant (FVN) : Trouver $\phi_N \in H^1(\omega) \cap L_0^2(\omega)$ tel que :*

$$\forall u \in H^1(\omega) \cap L_0^2(\omega), a_N(\phi_N, u) = l_N(u), \quad (2.14)$$

où a_N est la forme bilinéaire symétrique suivante :

$$\begin{aligned} a_N : H^1(\omega) \times H^1(\omega) &\rightarrow \mathbb{R} \\ (u, v) &\mapsto \int_{\omega} \mathbf{grad} u \cdot \mathbf{grad} v \, d\omega, \end{aligned}$$

et l_N est la forme linéaire dépendant des données f et e suivante :

$$\begin{aligned} l_N : H^1(\omega) &\rightarrow \mathbb{R} \\ u &\mapsto \int_{\omega} f u \, d\omega + \int_{\partial\omega} e u \, d\sigma. \end{aligned}$$

D'après le théorème de Lax-Milgram, sous réserve que (2.4) soit vérifiée, le problème (FVN) admet une unique solution $\phi_N \in H^1(\omega) \cap L_0^2(\omega)$.

2.2.2 CL homogènes

Dans le cas où l'on calcule $\mathbf{E}_0$, il s'agit de déterminer $\phi_D^0 \in H_0^1(\omega)$ et $\phi_N^0 \in H^1(\omega) \cap L_0^2(\omega)$ tels que :

$$-\Delta \phi_D^0 = g^0 \text{ dans } \omega, \quad (2.15)$$

et

$$\begin{cases} -\Delta \phi_N^0 = f^0 \text{ dans } \omega, \\ \partial_{\nu} \phi_N^0|_{\partial\omega} = 0 \text{ sur } \partial\omega, \end{cases} \quad (2.16)$$

avec $f^0 \in L_0^2(\omega)$. On remarque que $\phi_D^0 \in \Phi_D$, et $\phi_N^0 \in \Phi_N$.

Si on choisit $\mathbf{e}$ comme dans la proposition 2.2, alors $g^0 = g$, et $\phi_D^0 = \phi_D$; et $\phi_N^0 = \phi_N - \phi_N^e$. On reconstruit alors : $\mathbf{E}^0 = \mathbf{E}_D^0 + \mathbf{E}_N^0$, où : $\mathbf{E}_D^0 = -\mathbf{grad} \phi_D^0$ et $\mathbf{E}_N^0 = \mathbf{rot} \phi_N^0$.

Proposition 2.6 *On a la décomposition orthogonale suivante :*

$$\mathbf{X}_{\mathbf{E}}^0 = \mathbf{V}_{\mathbf{E}} \oplus \mathbf{W}_{\mathbf{E}}.$$

(2.17)

DÉMONSTRATION. On remarque que $\mathbf{E}_D^0 \in \mathbf{V}_E$, $\mathbf{E}_N^0 \in \mathbf{W}_E$, et que par ailleurs, $(\mathbf{E}_D^0, \mathbf{E}_N^0)_{\mathbf{X}_E^0} = 0$.

□

Le problème du Laplacien avec conditions limites de Dirichlet ou Neumann dans un domaine non convexe a fait l'objet de nombreuses recherches. Nous nous contentons de reprendre les résultats de P. Grisvard dans [67] (conditions aux limites homogènes) et de S. Nazarov et B. Plamenevsky [86] (conditions aux limites non-homogènes). Nous allons voir que pour obtenir plus de précision, on peut décomposer les potentiels en une partie qui est dans $H^2(\omega)$, et une partie explicite qui est seulement $H^1(\omega)$. Cette décomposition repose sur la décomposition de $L^2(\omega)$ en une partie régulière et une partie singulière, orthogonales entre elles [18, 66].

2.2.3 Le problème aux potentiels, à la Grisvard

Lorsque ω est convexe, $\Phi_{D,N}$ est inclus dans $H^2(\omega)$. Mais cette inclusion n'est plus vérifiée lorsqu'il existe un coin rentrant. On introduit alors le sous-espace régularisé de $\Phi_{D,N}$, noté $\Phi_{D,N}^R$ tel que : $\Phi_{D,N}^R = \Phi_{D,N} \cap H^2(\omega)$.

Théorème 2.7 *Lorsque ω n'est pas convexe, $\Phi_{D,N}^R$ est fermé et strictement inclus dans $\Phi_{D,N}$.*

DÉMONSTRATION. Soit $\phi \in \Phi_{D,N}^R$. Comme dans $\Phi_{D,N}$, la norme du graphe est équivalente à la semi-norme, on a :

$$\|\phi\|_{H^1} \leq C \|\Delta\phi\|_0.$$

Or, on a :

$$|\phi|_{H^2} = |\mathbf{grad} \phi|_{H^1} = \|\mathbf{grad} : \mathbf{grad} \phi\|_0 \equiv \|\Delta\phi\|_0.$$

Ainsi, sur $\Phi_{D,N}^R$, $\|\Delta \cdot\|_0 \equiv |\cdot|_2$. Comme $H^2(\omega)$ est complet pour sa norme canonique, $\Phi_{D,N}^R$ est fermé dans $\Phi_{D,N}$.

□

Proposition 2.8 *L'opérateur Δ définit un isomorphisme de Φ_N dans $L_0^2(\omega)$ et de Φ_D dans $L^2(\omega)$. Φ_D et Φ_N étant munis de la norme L^2 du Laplacien, et $L_0^2(\omega)$ étant muni de la norme L^2 , ces isomorphismes préservent l'orthogonalité.*

DÉMONSTRATION. Pour le problème de Neuman, voir [9] (lem. 2.1, p. 361).

□

Ainsi $\Delta\Phi_D^R$ (resp. $\Delta\Phi_N^R$) est un sous-espace fermé de $L^2(\omega)$ (resp. $L_0^2(\omega)$).

Soit $S_{D,N}$, l'espace orthogonal à $\Delta\Phi_{D,N}^R$ dans $L_{(0)}^2(\omega)$. On en déduit la décomposition directe orthogonale suivante : $L^2(\omega) = \Delta\Phi_D^R \oplus S_D$ (resp. $L_0^2(\omega) = \Delta\Phi_N^R \oplus S_N$). Ce résultat nous permet de caractériser les fonctions singulières de $\Phi_{D,N}$: on obtient la décomposition directe orthogonale suivante entre parties régulière et singulière : $\Phi_{D,N} = \Phi_{D,N}^R \oplus \Phi_{D,N}^S$ où $\Phi_{D,N}^S = \Delta^{-1}S_{D,N}$.

Les espaces $\Phi_{D,N}^S$ sont appelés *espace des singularités primales du Laplacien* et les espace $S_{D,N}$ sont appelés *espace des singularités duales du Laplacien*.

Proposition 2.9 *Les espaces des singularités duales S_D et S_N sont engendrés par les fonctions qui satisfont les problèmes de Dirichlet et de Neumann non standards suivants :*

$$s_D \in L^2(\omega), s_D \in S_D \iff \begin{cases} \Delta s_D &= 0 \text{ dans } \omega, \\ s_D|_{A_k} &= 0 \forall k \in \{1, \dots, K\} \text{ dans } H_{00}^{-1/2}(A_k), \end{cases} \quad (2.18)$$

$$s_N \in L^2(\omega), s_N \in S_N \iff \begin{cases} \Delta s_N &= 0 \text{ dans } \omega, \\ \partial_\nu s_N|_{A_k} &= 0 \forall k \in \{1, \dots, K\} \text{ dans } H_{00}^{-3/2}(A_k). \end{cases} \quad (2.19)$$

DÉMONSTRATION. Soit $s_N \in S_N$. Alors $\forall \phi \in \Phi_N^R$, $\int_\omega \Delta \phi s_N d\omega = 0$. On en déduit que $\Delta s_N = 0$ au sens des distributions. On obtient la condition limite au bord en utilisant une formule de Green [67] (thm. 1.5.3, p. 26) et le fait que $\forall k$, $\partial_\nu \phi|_{A_k} = 0$, et $\forall k$, $\phi|_{A_k} \in H_{00}^{3/2}(A_k)$. La preuve est similaire pour s_D . $\square$

Les fonctions de $S_{D,N}$ sont H^1 -régulières en dehors d'un voisinage de chaque coin rentrant. $S_{D,N}$ est de dimension égale au nombre de singularités géométriques, c'est-à-dire le nombre de coins rentrants : $\dim(S_{D,N}) = N_{cr}$, et $S_D = \text{vect}(s_{D,i})$ (resp. $S_N = \text{vect}(s_{N,i})$) où $s_{D,i}$ (resp. $s_{N,i}$) satisfait (2.18) (resp. (2.19)). Qui plus est, $s_{D,i}$ et $s_{N,i}$ se décomposent en une partie régulière $\tilde{s}_{D,i}$, $\tilde{s}_{N,i} \in H^1(\omega)$ et une partie principale $s_{D,i}^P = r_i^{-\alpha_i} \sin(\alpha_i \theta_i)$, $s_{N,i}^P = r_i^{-\alpha_i} \cos(\alpha_i \theta_i)$ qui n'est pas dans $H^1(\omega)$ [66] (2.3.6, p. 51.) :

$$s_{D,i} = \tilde{s}_{D,i} + s_{D,i}^P \text{ et } s_{N,i} = \tilde{s}_{N,i} + s_{N,i}^P.$$

$s_{D,i}^P$ et $s_{N,i}^P$ sont harmoniques et régulières en dehors d'un voisinage du coin rentrant O_i .

Proposition 2.10 *$\tilde{s}_{D,i}$ satisfait le problème de Dirichlet suivant :*

Trouver $\tilde{s}_{D,i} \in H^1(\omega)$ tel que :

$$\begin{cases} \Delta \tilde{s}_{D,i} &= 0 \text{ dans } \omega, \\ \tilde{s}_{D,i}|_{A_k} &= -s_{D,i}^P|_{A_k} \forall k \in \{1, \dots, K\}, \text{ dans } H^{1/2}(\partial\omega). \end{cases} \quad (2.20)$$

DÉMONSTRATION. Montrons qu'on a : $\tilde{s}_{D,i}|_{\gamma_i} = 0$ dans $H^{1/2}(\gamma_i)$. On fait la preuve par contradiction. Supposons que $\tilde{s}_{D,i}|_{\gamma_i} \neq 0$ dans $H^{1/2}(\gamma_i)$. Comme $H^{1/2}(\gamma_i) \subset L^2(\gamma_i)$, cela implique que $\tilde{s}_{D,i}|_{\gamma_i} \neq 0$ dans $L^2(\gamma_i)$, c'est-à-dire que $\tilde{s}_{D,i}|_{A_i^0} \neq 0$ dans $L^2(A_i^0)$, ou $\tilde{s}_{D,i}|_{A_i^\Theta} \neq 0$ dans $L^2(A_i^\Theta)$. Retenons la première possibilité : $\tilde{s}_{D,i}|_{A_i^0} \neq 0$ dans $L^2(A_i^0)$. Comme $L^2(A_i^0) \subset H_{00}^{-1/2}(A_i^0)$, on a donc : $\tilde{s}_{D,i}|_{A_i^0} \neq 0$ dans $H_{00}^{-1/2}(A_i^0)$. Or, par définition, on a : $\tilde{s}_{D,i}|_{A_i^0} = (s_{D,i} - r_i^{-\alpha_i} \sin(\alpha_i \theta_i))|_{A_i^0}$, dans $H_{00}^{-1/2}(A_i^0)$. Comme $s_{D,i} \in S_D$, et que $(r_i^{-\alpha_i} \sin(\alpha_i \theta_i))|_{A_i^0}$ s'annule (car $\theta_i = 0$), on en conclut que $\tilde{s}_{D,i}|_{A_i^0} = 0$ dans $H_{00}^{-1/2}(A_i^0)$, ce qui contredit le précédent résultat. Par ailleurs, sur les autres arêtes, on a : $s_{D,i}^{(P)} \in C^\infty(\partial\omega \setminus \gamma_i)$. Les traces de $s_{D,i}^{(P)}$ et $(1 - \eta_i(r_i)) s_{D,i}^{(P)}$ se raccordent aux extrémités des arêtes. On en déduit que $\tilde{s}_{D,i}$ et $(1 - \eta_i(r_i)) \tilde{s}_{D,i}$ ont même trace sur $\partial\omega$. $\square$

Proposition 2.11 *$\tilde{s}_{N,i}$ satisfait le problème de Neumann suivant :*

Trouver $\tilde{s}_{N,i} \in H^1(\omega)$ tel que :

$$\begin{cases} \Delta \tilde{s}_{N,i} &= 0 \text{ dans } \omega, \\ \partial_\nu \tilde{s}_{N,i}|_{A_k} &= -\partial_\nu s_{N,i}^P|_{A_k} \forall k \in \{1, \dots, K\}, \text{ dans } L^2(\partial\omega). \end{cases} \quad (2.21)$$

DÉMONSTRATION. Montrons, par contradiction, qu'on a : $\partial_\nu \tilde{s}_{N,i}|_{\gamma_i} = 0$ dans $L^2(\gamma_i)$. Supposons que $\partial_\nu \tilde{s}_{N,i}|_{\gamma_i} \neq 0$ dans $L^2(\gamma_i)$, c'est-à-dire $\partial_\nu \tilde{s}_{N,i}|_{A_i^0} \neq 0$ dans $L^2(A_i^0)$, ou $\partial_\nu \tilde{s}_{N,i}|_{A_i^\Theta} \neq 0$ dans $L^2(A_i^\Theta)$. Retenons la première possibilité. Comme $L^2(A_i^0) \subset H_{00}^{-3/2}(A_i^0)$, on a donc : $\partial_\nu \tilde{s}_{N,i}|_{A_i^0} \neq 0$ dans $H_{00}^{-3/2}(A_i^0)$.

Or, on remarque que : $\mathbf{grad} s_{N,i}^P(r_i, \theta_i) = -\alpha_i r_i^{-\alpha_i-1} \begin{pmatrix} \cos(\alpha_i \theta_i) \\ \sin(\alpha_i \theta_i) \end{pmatrix}$, exprimé en coordonnées polaires par rapport au coin rentrant. Ainsi, sur les arêtes du coin rentrant :

$$\partial_\nu s_{N,i}^P|_{A_i^0} = (\partial_\nu s_{N,i} - \mathbf{grad} s_{N,i}^P) \cdot \boldsymbol{\nu}|_{A_i^0} = (\partial_\nu s_{N,i} - \alpha_i r_i^{-\alpha_i-1} \sin(\alpha_i \theta_i)), \text{ dans } H_{00}^{-3/2}(A_i^0).$$

Comme $\partial_\nu s_{N,i} \in S_N$, et que $(r_i^{-\alpha_i-1} \sin(\alpha_i \theta_i))|_{A_i^0}$ s'annule (car $\theta = 0$), on en conclut que $\partial_\nu \tilde{s}_{N,i}|_{A_i^0} = 0$ dans $H_{00}^{-3/2}(A_i^0)$, ce qui contredit le précédent résultat.

Sur les autres arêtes, $\partial_\nu s_{N,i}^{(P)}|_{\partial\omega \setminus \gamma_i} \in C^\infty(\partial\omega \setminus \gamma_i)$. Les traces de $\partial_\nu s_{N,i}^{(P)}|_{\partial\omega \setminus \gamma_i}$ et $(1 - \eta_i(r_i)) \partial_\nu s_{N,i}^{(P)}$ se raccordent aux extrémités des arêtes. On en déduit que $\partial_\nu \tilde{s}_{N,i}$ et $(1 - \eta_i(r_i)) \partial_\nu \tilde{s}_{N,i}$ ont même trace sur $\partial\omega$. □

Théorème 2.12 *On peut établir les estimations suivantes [67] (thm. 1.2.18, p. 7) :*

$\forall \epsilon > 0$ $s \in S_{D,N}$ appartient à $H^{1-\alpha-\epsilon}(\omega)$ et n'appartient pas à $H^{1-\alpha}(\omega)$.

D'après la proposition 2.8, $\forall i \in \{1, \dots, N_{cr}\}$ il correspond à $s_{D,i} \in S_D$ (resp. $s_{N,i} \in S_N$) un unique $\varphi_{D,i}$ (resp. $\varphi_{N,i}$) de $H^1(\omega)$ tels que :

$$\begin{cases} -\Delta \varphi_{D,i} = s_{D,i} \text{ dans } \omega, \\ \varphi_{D,i}|_{\partial\omega} = 0 \text{ sur } \partial\omega, \end{cases} \quad \text{et} \quad \begin{cases} -\Delta \varphi_{N,i} = s_{N,i} \text{ dans } \omega, \\ \partial_\nu \varphi_{N,i}|_{\partial\omega} = 0 \text{ sur } \partial\omega. \end{cases} \quad (2.22)$$

Ainsi, les espaces des singularités primales Φ_D^S et Φ_N^S sont de dimension finie égale à N_{cr} , et on a :

$$\Phi_D^S = \text{vect}(\varphi_{D,i}) \text{ et } \Phi_N^S = \text{vect}(\varphi_{N,i}). \quad (2.23)$$

$\forall j \in \{1, \dots, N_{cr}\}$, on pose : $\varphi_{D,j}^P = r_j^{\alpha_j} \sin(\alpha_j \theta_j)$, et $\varphi_{N,j}^P = r_j^{\alpha_j} \cos(\alpha_j \theta_j)$.

Les $\varphi_{D,j}^P$ et $\varphi_{N,j}^P$, $j \in \{1, \dots, N_{cr}\}$ formeront les parties principales des $\varphi_{D,i}$ et $\varphi_{N,i}$, $i \in \{1, \dots, N_{cr}\}$.

Proposition 2.13 $\varphi_{D,i}$ et $\varphi_{N,i}$ se décomposent en une partie H^2 -régulière et une partie singulière de la façon suivante :

$$\varphi_{D,i} = \tilde{\varphi}_{D,i} + \sum_{j=1}^{N_{cr}} \beta_D^{i,j} \varphi_{D,j}^P \text{ avec } \tilde{\varphi}_{D,i} \in H^2(\omega) \text{ et } \varphi_{D,j}^P \notin H^2(\omega) \quad (2.24)$$

$$\varphi_{N,i} = \tilde{\varphi}_{N,i} + \sum_{j=1}^{N_{cr}} \beta_N^{i,j} \varphi_{N,j}^P \text{ avec } \tilde{\varphi}_{N,i} \in H^2(\omega) \text{ et } \varphi_{N,j}^P \notin H^2(\omega), \quad (2.25)$$

où $\beta_D^{i,j} = (s_{D,i}, s_{D,j})_0 / \pi$ et $\beta_N^{i,j} = (s_{N,i}, s_{N,j})_0 / \pi$.

Les $\varphi_{D,i}^P$ et $\varphi_{N,i}^P$ sont harmoniques et régulières en dehors du voisinage des coins rentrants. Le calcul des $\beta_{D,N}^{i,j}$ est détaillé en dans la section 14.1 de la partie IV.

Théorème 2.14 *On peut établir les estimations suivantes [67] :*

$\forall \epsilon > 0$, $u \in \Phi_{D,N}^S$ appartient à $H^{1+\alpha-\epsilon}(\omega)$ et n'appartient pas à $H^{1+\alpha}(\omega)$; $\Phi_{D,N}^R \subset H^2(\omega)$.

Proposition 2.15 $\tilde{\varphi}_{D,i}$ satisfait le problème de Dirichlet suivant :
 Trouver $\tilde{\varphi}_{D,i} \in H^1(\omega)$ tel que :

$$\begin{cases} -\Delta \tilde{\varphi}_{D,i} = s_{D,i} \text{ dans } \omega, \\ \tilde{\varphi}_{D,i}|_{A_k} = -\sum_{j=1}^{N_{cr}} \beta_D^{i,j} \varphi_{D,j}|_{A_k}^P \quad \forall k \in \{1, \dots, K\} \text{ au sens } H^{1/2}(\partial\omega). \end{cases} \quad (2.26)$$

DÉMONSTRATION. En procédant comme pour la démonstration de 2.10, on montre que sur les arêtes du coin rentrant O_i :

$$\begin{aligned} - \varphi_{D,i}^P|_{A_i^0} &= \alpha_i r_i^{\alpha_i} \sin(\alpha_i \theta_i)_{\theta_i=0} = 0, \text{ dans } H^{1/2}(A_i^0), \\ - \varphi_{D,i}^P|_{A_i^\Theta} &= \alpha_i r_i^{\alpha_i} \sin(\alpha_i \theta_i)_{\theta_i=\pi/\alpha_i} = 0, \text{ dans } H^{1/2}(A_i^\Theta). \end{aligned}$$

Sur les autres arêtes, on a : $\varphi_{D,i}^P|_{\partial\omega \setminus \gamma_i} \in C^\infty(\partial\omega \setminus \gamma_i)$. Par ailleurs, les traces de $\varphi_{D,i}^{(P)}$ se raccordent aux extrémités des arêtes. On en déduit que $\tilde{\varphi}_{D,i}$ et $(1 - \eta_i(r_i)) \tilde{\varphi}_{D,i}$ ont même trace sur $\partial\omega$. □

Proposition 2.16 $\tilde{\varphi}_{N,i}$ satisfait le problème de Neumann suivant :
 Trouver $\tilde{\varphi}_{N,i} \in H^1(\omega)$ tel que :

$$\begin{cases} -\Delta \tilde{\varphi}_{N,i} = s_{N,i} \text{ dans } \omega, \\ \partial_\nu \tilde{\varphi}_{N,i}|_{A_k} = -\sum_{j=1}^{N_{cr}} \beta_N^{i,j} \partial_\nu \varphi_{N,j}|_{A_k}^P \quad \forall k \in \{1, \dots, K\}, \text{ au sens } L^2(\partial\omega). \end{cases} \quad (2.27)$$

DÉMONSTRATION. Notons d'abord que : $\mathbf{grad} \varphi_{N,i}^P(r_i, \theta_i) = \alpha_i r_i^{\alpha_i-1} \begin{pmatrix} \cos(\alpha_i \theta_i) \\ -\sin(\alpha_i \theta_i) \end{pmatrix}$, exprimé en coordonnées polaires par rapport au coin rentrant O_i . En procédant comme pour la démonstration de 2.11, on montre que sur les arêtes du coin rentrant O_i :

$$\begin{aligned} - \partial_\nu \varphi_{N,i}^P|_{A_i^0} &= \alpha_i r_i^{\alpha_i-1} \sin(\alpha_i \theta_i)_{\theta_i=0} = 0, \text{ dans } L^2(A_i^0), \\ - \partial_\nu \varphi_{N,i}^P|_{A_i^\Theta} &= -\alpha_i r_i^{\alpha_i-1} \sin(\alpha_i \theta_i)_{\theta_i=\pi/\alpha_i} = 0, \text{ dans } L^2(A_i^\Theta). \end{aligned}$$

Sur les autres arêtes, on a : $\partial_\nu \varphi_{N,i}^P|_{\partial\omega \setminus \gamma_i} \in C^\infty(\partial\omega \setminus \gamma_i)$. Par ailleurs, les traces de $\partial_\nu \varphi_{N,i}^{(P)}$ se raccordent aux extrémités des arêtes. On en déduit que $\partial_\nu \tilde{\varphi}_{N,i}$ et $(1 - \eta_i(r_i)) \partial_\nu \tilde{\varphi}_{N,i}$ ont même trace sur $\partial\omega$. □

Nous allons maintenant étudier $\phi_D^0 \in \Phi_D$, et $\phi_N^0 \in \Phi_N$, les solutions des problèmes (2.15) et (2.16). Les coefficients définis ci-dessous (définition 2.17) vont apparaître naturellement dans l'expression de ϕ_D^0 et ϕ_N^0 sous forme de somme entre une partie régulière et une partie singulière, décomposée sur les parties principales.

Définition 2.17 Pour tout $j \in \{1, \dots, N_{cr}\}$, on pose :

$$\lambda_D^j := (g^0, s_{D,j})_0/\pi \text{ et } \lambda_N^j := (f^0, s_{N,j})_0/\pi.$$

Proposition 2.18 $\phi_D^0 \in \Phi_D$, la solution du problème (2.15) s'écrit ainsi :

$$\boxed{\phi_D^0 = \tilde{\phi}_D + \sum_{j=1}^{N_{cr}} \lambda_D^j \varphi_{D,j}^P, \text{ où } \tilde{\phi}_D \in H^2(\omega).} \quad (2.28)$$

Cette décomposition a été à l'origine introduite pour le cas d'un unique coin rentrant par M. Moussaoui dans [83].

DÉMONSTRATION. ϕ_D^0 se décompose dans Φ_D en une partie régulière $\hat{\phi}_D \in \Phi_D^R$ et une partie singulière, décomposée sur les vecteurs de base de Φ_D^S ainsi :

$$\phi_D^0 = \hat{\phi}_D + \sum_{i=1}^{N_{cr}} c_D^i \varphi_{D,i}, \text{ où : } \forall i, c_D^i \in \mathbb{R},$$

ce qui se réécrit sous la forme :

$$\phi_D^0 = \tilde{\phi}_D + \sum_{j=1}^{N_{cr}} \left(\sum_{i=1}^{N_{cr}} c_D^i \beta_D^{i,j} \right) \varphi_{D,j}^P, \text{ où : } \tilde{\phi}_D = \hat{\phi}_D + \sum_{i=1}^{N_{cr}} c_D^i \tilde{\varphi}_{D,i} \in H^2(\omega).$$

D'après la section 14.2 de la partie IV, on a :

$$\sum_{i=1}^{N_{cr}} c_D^i \beta_D^{i,j} = (g^0, s_{D,j})_0/\pi := \lambda_D^j. \quad (2.29)$$

□

Proposition 2.19 $\phi_N^0 \in \Phi_N$, la solution du problème (2.16) s'écrit ainsi :

$$\boxed{\phi_N^0 = \tilde{\phi}_N + \sum_{j=1}^{N_{cr}} \lambda_N^j \varphi_{N,j}^P, \text{ où } \tilde{\phi}_N \in H^2(\omega).} \quad (2.30)$$

DÉMONSTRATION. Changer les indices D par N dans la démonstration de la proposition 2.18.

□

Théorème 2.20 Supposons que pour $\epsilon > 0$ tel que $2\alpha - 1 - \epsilon > 0$, $f, g \in H^{2\alpha-1-\epsilon}(\omega)$ et $e \in H^{2\alpha-1/2-\epsilon}(\partial\omega)$. Alors $\tilde{\phi}_{D,N} \in H^{2\alpha+1-\epsilon}(\omega)$.

DÉMONSTRATION. D'après les hypothèses sur les données, e est la trace d'un vecteur $\mathbf{e}$ de $\mathbf{H}^{2\alpha-\epsilon}(\omega)$, et donc $\operatorname{div} \mathbf{e} \in H^{2\alpha-1-\epsilon}(\omega)$, d'où : $g^0 = g - \operatorname{div} \mathbf{e} \in H^{2\alpha-1-\epsilon}(\omega)$. D'après [67] (p. 82), lorsque $g^0 \in H^s(\omega)$, $s \geq -1$, on peut décomposer ϕ_D^0 sous la forme :

$$\phi_D^0 = \phi_D^R + \sum_{j=1}^{N_{cr}} \left(\sum_{0 < m\alpha_j < s+1} c_{j,m} r_j^{m\alpha_j} \sin(m\alpha_j \theta_j) \right), \text{ avec } \phi_D^R \in H^{s+2}(\omega), c_{j,m} \in \mathbb{R}.$$

Le calcul des $c_{j,m}$ est donné dans [86]. On a $c_{j,1} = \lambda_D^j$. Afin d'obtenir la décomposition de la proposition 2.18, c'est-à-dire $\phi_D^R = \tilde{\phi}_D$, il faut limiter la somme sur m à $m = 1$ pour tout j , ce que est possible si et seulement si :

$$\forall j \in \{1, \dots, N_{cr}\}, s+1 < 2\alpha_j \iff s < s_{\max} \text{ avec } s_{\max} = 2\alpha - 1. \quad (2.31)$$

Ainsi, $s = 2\alpha - 1 - \epsilon$ convient, et dans ce cas, $\phi_D^R = \tilde{\phi}_D$, comme annoncé. Le même résultat est obtenu pour le problème de Neumann. Pour conclure, notons qu'on peut obtenir une partie régulière de plus en plus régulière sous réserve de la régularité des données. $\square$

2.2.4 Le problème aux potentiels, à la Nazarov-Plamenevsky

La décomposition partie régulière-partie singulière est possible également pour les potentiels ϕ_D et ϕ_N , c'est-à-dire avec une condition aux limites non-homogènes pour le problème de Neumann. Cette approche a été développée par S. Nazarov et B. Plamenevsky dans [86]. On considère le cas où il n'existe qu'un seul coin rentrant. Comme ϕ_D satisfait un problème de Dirichlet homogène, l'approche de [86] est identique à celle de [67] :

Théorème 2.21 *Supposons que $g \in L^2(\omega)$. Alors ϕ_D , la solution de (2.9) s'écrit de la façon suivante [86] (thm. 3.4, p. 33) :*

$$\phi_D = w_D + \frac{1}{\pi} (g, s_D)_0 \varphi_D^P, \quad (2.32)$$

où $w_D \in H^2(\omega)$ satisfait le problème de Dirichlet non homogène suivant :

$$\begin{cases} -\Delta w_D &= g \text{ dans } \omega \\ w_D|_{\partial\omega} &= -\frac{1}{\pi} (g, s_D)_0 \varphi_D^P|_{\partial\omega} \text{ sur } \partial\omega. \end{cases}$$

On a un résultat intéressant pour ϕ_N qui satisfait un problème de Neumann non-homogène :

Théorème 2.22 *Supposons que $f \in L^2(\omega)$ et $e \in V_0^{1/2}(\partial\omega)$. Alors ϕ_N , la solution de (2.10) s'écrit de la façon suivante [86] (thm. 4.4, p. 41) :*

$$\phi_N = w_N + \frac{1}{\pi} ((f, s_N)_0 + \langle e, s_N \rangle_{H,H'}) \varphi_N^P, \quad (2.33)$$

où H désigne $V_0^{1/2}(\partial\omega)$ et $\langle \cdot, \cdot \rangle_{H,H'}$ est le crochet de dualité entre H et son dual H' ; et où $w_N \in H^2(\omega)$ satisfait le problème de Neumann suivant :

$$\begin{cases} -\Delta w_N &= f \text{ dans } \omega \\ w_N|_{\partial\omega} &= e - \frac{1}{\pi} \left((f, s_N)_0 + \langle e, s_N \rangle_{H,H'} \right) \varphi_N^P \text{ sur } \partial\omega, \end{cases}$$

Rappelons que $V_0^{1/2}(\partial\omega)$ est l'espace des traces des fonctions de $V_0^1(\omega)$, avec :

$$V_0^1(\omega) = \left\{ u \in H^1(\omega) \mid \int_{\omega} \left| \frac{u}{r} \right|^2 d\omega < \infty \right\}. \quad (2.34)$$

On en déduit que e est ici la trace d'une fonction de $H^1(\omega)$ qui s'annule par exemple en $r^{1/2}$ au voisinage du coin rentrant, c'est-à-dire que $e \in H^{1/2}(\omega)$, s'annulant au voisinage du coin rentrant.

2.3 Champ électrique 2D : CL naturelles

Le cas électrostatique (g quelconque, $f = 0$, $e = 0$, paragraphe 2.2.2) se généralise sans difficulté au problème tridimensionnel. Cependant, la qualité de l'approximation du champ électrostatique est limitée car elle est calculée par dérivation d'éléments continus P_k (conformité uniquement dans $\mathcal{H}(\text{rot}, \Omega)$ pour le problème tridimensionnel, $\mathbf{H}(\text{rot}, \omega)$ pour le problème bidimensionnel). On s'intéresse donc au calcul direct du champ électrique dans l'espace $\mathbf{X}_{\mathbf{E}}$.

Proposition 2.23 *Le problème (2.1)-(2.3) est équivalent au problème variationnel suivant : Trouver $\mathbf{E} \in \mathbf{X}_{\mathbf{E}}$ tel que :*

$$\forall \mathbf{F} \in \mathbf{X}_{\mathbf{E}}, \mathcal{A}_{\mathbf{E}}(\mathbf{E}, \mathbf{F}) = \mathcal{L}_{\mathbf{E}}(\mathbf{F}), \quad (2.35)$$

où $\mathcal{A}_{\mathbf{E}}$ est la forme bilinéaire symétrique coercitive suivante :

$$\begin{aligned} \mathcal{A}_{\mathbf{E}} : \mathbf{X}_{\mathbf{E}} \times \mathbf{X}_{\mathbf{E}} &\rightarrow \mathbb{R} \\ (\mathbf{E}, \mathbf{F}) &\mapsto (\text{div } \mathbf{E}, \text{div } \mathbf{F})_0 + (\text{rot } \mathbf{E}, \text{rot } \mathbf{F})_0 + \int_{\partial\omega} E_{\tau} F_{\tau} d\sigma, \end{aligned}$$

et $\mathcal{L}_{\mathbf{E}}$ est la forme linéaire suivante :

$$\begin{aligned} \mathcal{L}_{\mathbf{E}} : \mathbf{X}_{\mathbf{E}} &\rightarrow \mathbb{R} \\ \mathbf{F} &\mapsto (f, \text{rot } \mathbf{F})_0 + (g, \text{div } \mathbf{F})_0 + \int_{\partial\omega} e F_{\tau} d\sigma. \end{aligned}$$

DÉMONSTRATION. Il est clair que si $\mathbf{E}$ satisfait (2.1)-(2.3) alors $\mathbf{E}$ satisfait (2.35). Réciproquement, montrons que si $\mathbf{E}$ est solution de (2.35), alors $\mathbf{E}$ satisfait (2.1)-(2.3). Soit $u \in L^2(\omega)$. D'après la décomposition (2.12), il existe $\mathbf{F}_D \in \mathbf{X}_{\mathbf{E}}^0$ tel que : $\text{rot } \mathbf{F}_D = 0$ et $\text{div } \mathbf{F}_D = u$, où u est la donnée de (2.9). En injectant $\mathbf{F}_D$ dans (2.35), on a : $(\text{div } \mathbf{E}, u)_0 = (g, u)_0$, où $\text{div } \mathbf{E}$ et g sont dans $L^2(\omega)$. On en déduit que $\text{div } \mathbf{E} = g$. L'équation (2.35) se réduit à :

Trouver $\mathbf{E} \in \mathbf{X}_{\mathbf{E}}$ tel que :

$$\forall \mathbf{F} \in \mathbf{X}_{\mathbf{E}}, (\text{rot } \mathbf{E}, \text{rot } \mathbf{F})_0 + \int_{\partial\omega} E_{\tau} F_{\tau} d\sigma = (f, \text{rot } \mathbf{F})_0 + \int_{\partial\omega} e F_{\tau} d\sigma.$$

Considérons le vecteur $\mathbf{E}_N$ de la décomposition (2.12), on a $\text{rot } \mathbf{E}_N = f$, $\text{div } \mathbf{E}_N = 0$, $E_{N,\tau} = e$. Ainsi :

$$\forall \mathbf{F} \in \mathbf{X}_{\mathbf{E}}, (\text{rot } \mathbf{E}_N, \text{rot } \mathbf{F})_0 + \int_{\partial\omega} E_{N,\tau} F_{\tau} d\sigma = (f, \text{rot } \mathbf{F})_0 + \int_{\partial\omega} e F_{\tau} d\sigma.$$

On en déduit que : $\forall \mathbf{F} \in \mathbf{X}_{\mathbf{E}}, (\text{rot}(\mathbf{E} - \mathbf{E}_N), \text{rot} \mathbf{F})_0 + \int_{\partial\omega} (\mathbf{E}_\tau - \mathbf{E}_{N,\tau}) \mathbf{F}_\tau d\sigma = 0$.
 Prenons comme fonction test $\mathbf{F} = \mathbf{E} - \mathbf{E}_N$. On a alors : $\mathbf{F}_\tau = \mathbf{E}_\tau - \mathbf{E}_{N,\tau}$, d'où :

$$\|\text{rot}(\mathbf{E} - \mathbf{E}_N)\|_0^2 + \|\mathbf{E}_\tau - \mathbf{E}_{N,\tau}\|_{0,\partial\omega}^2 = 0.$$

D'où $\text{rot}(\mathbf{E} - \mathbf{E}_N) = 0$ dans ω et $\mathbf{E}_\tau - \mathbf{E}_{N,\tau} = 0$ sur $\partial\omega$.

On obtient finalement : $\text{rot} \mathbf{E} = \text{rot} \mathbf{E}_N = f$, et $\mathbf{E}_\tau = \mathbf{E}_{N,\tau} = e$.

Pour conclure, d'après le théorème 1.7, p. 47, (2.35) admet une unique solution $\mathbf{E} \in \mathbf{X}_{\mathbf{E}}$, qui dépend continûment des données f, g et e . □

Notons pour finir que, comme $\mathbf{X}_{\mathbf{E}} \cap \mathbf{H}^1(\omega)$ est dense dans $\mathbf{X}_{\mathbf{E}}$ [40, 51], on peut approcher la solution de (2.35) par les éléments finis de Lagrange continus P_k .

Remarques 2.24 *Au contraire de la méthode des potentiels, cette méthode est applicable au problème tridimensionnel général (1.2)-(1.5), et l'on peut aussi la développer en 2D pour le cas dépendant du temps. Notons que dans [54], M. Costabel et al. étudient le problème harmonique dans $\mathbf{X}_{\mathbf{H}}$ pour le champ magnétique.*

2.4 Champ électrique 2D : le complément singulier

Dans la formulation (2.35), la condition aux limites est naturelle. Que se passe-t-il si on veut résoudre (2.1)-(2.3) dans l'espace fonctionnel $\mathbf{X}_{\mathbf{E}}^0$ pour lequel la condition aux limites est prise en compte de façon essentielle ?

2.4.1 Introduction

L'espace $\mathbf{X}_{\mathbf{E}}^0$ permet de prendre en compte de façon exacte la condition aux limites tangentielle nulle. Lorsque le domaine ω n'est pas convexe, la solution du problème (2.6)-(2.7) peut être singulière : la valeur des composantes du champ électrostatique est localement non bornée au voisinage des points de non-convexité. Mathématiquement, les composantes du champ ne sont pas dans $H^1(\omega)$.

Soit $\mathbf{X}_{\mathbf{E}}^{0,R} := \mathbf{X}_{\mathbf{E}}^0 \cap \mathbf{H}^1(\omega)$, l'espace *régularisé* de $\mathbf{X}_{\mathbf{E}}^0$. On a le théorème suivant, montré par M. Costabel dans [50] :

Théorème 2.25 *Dans l'espace $\mathbf{X}_{\mathbf{E}}^{0,R}$, la norme de $\mathbf{X}_{\mathbf{E}}^0$ est équivalente à la norme de $\mathbf{H}^1(\omega)$:*

$$\forall \mathbf{F} \in \mathbf{X}_{\mathbf{E}}^{0,R}, \|\mathbf{F}\|_{\mathbf{X}_{\mathbf{E}}^0} \equiv \|\mathbf{F}\|_{\mathbf{H}^1}.$$

De plus, d'après M. Costabel et al. [55], les fonctions régulières sont denses dans $\mathbf{X}_{\mathbf{E}}^{0,R}$.

Lorsque ω est convexe, $\mathbf{X}_{\mathbf{E}}^{0,R}$ est égal à $\mathbf{X}_{\mathbf{E}}^0$ et l'on peut résoudre (2.6)-(2.7) par les éléments finis de Lagrange de la même façon que dans $\mathbf{X}_{\mathbf{E}}$. En revanche, lorsque ω n'est pas convexe, on a le théorème suivant :

Théorème 2.26 *Lorsque ω n'est pas convexe, $\mathbf{X}_{\mathbf{E}}^{0,R}$ est fermé et strictement inclus dans $\mathbf{X}_{\mathbf{E}}^0$.*

DÉMONSTRATION. Soit $\mathbf{u} \in \mathbf{X}_{\mathbf{E}}^{0,R}$. Comme $\mathbf{u} \in \mathbf{X}_{\mathbf{E}}^0$, et que dans $\mathbf{X}_{\mathbf{E}}^0$, la semi-norme est équivalente à la norme du graphe, il existe une constante $C > 0$ telle que :

$$\|\mathbf{u}\|_0^2 \leq C (\|\text{div} \mathbf{u}\|_0^2 + \|\text{rot} \mathbf{u}\|_0^2).$$

D'après [50, 84], on a pour $\mathbf{u} \in \mathbf{X}_{\mathbf{E}}^{0,R}$: $\|\mathbf{grad} : \mathbf{u}\|_0^2 = \|\operatorname{div} \mathbf{u}\|_0^2 + \|\mathbf{rot} \mathbf{u}\|_0^2$. Comme $\mathbf{H}^1(\omega)$ est complet pour sa norme canonique, on en déduit que $\mathbf{X}_{\mathbf{E}}^{0,R}$ est fermé dans $\mathbf{X}_{\mathbf{E}}^0$. $\square$

Tout sous-espace de $\mathbf{X}_{\mathbf{E}}^0$ généré par les éléments finis P_k de Lagrange est un sous-espace de $\mathbf{X}_{\mathbf{E}}^{0,R}$. Ainsi, les éléments finis P_k de Lagrange ne convergent pas vers la solution de (2.6)-(2.7).

Posons $\mathbf{X}_{\mathbf{E}}^0 = \mathbf{X}_{\mathbf{E}}^{0,R} \oplus \mathbf{X}_{\mathbf{E}}^{0,S}$, où $\mathbf{X}_{\mathbf{E}}^{0,S}$ est appelé *partie singulière* de $\mathbf{X}_{\mathbf{E}}^0$. Soit $\mathbf{E}^0 \in \mathbf{X}_{\mathbf{E}}^0$, que l'on décompose en une partie régulière $\mathbf{E}^{0,R} \in \mathbf{X}_{\mathbf{E}}^{0,R}$ et une partie singulière $\mathbf{E}^{0,S} \in \mathbf{X}_{\mathbf{E}}^{0,S}$. Soit $\mathbf{E}_h^0$ l'approximation de $\mathbf{E}^0$ par les éléments finis P_k de Lagrange. On a alors :

$$\|\mathbf{E}^0 - \mathbf{E}_h^0\|_{\mathbf{X}_{\mathbf{E}}^0}^2 = \|\mathbf{E}^{0,S} - \mathbf{E}_h^0\|_{\mathbf{X}_{\mathbf{E}}^0}^2 + \|\mathbf{E}^{0,S}\|_{\mathbf{X}_{\mathbf{E}}^0}^2 \geq \|\mathbf{E}^{0,S}\|_{\mathbf{X}_{\mathbf{E}}^0}^2.$$

L'erreur que l'on commet localement en n'approchant pas la singularité se répercute sur toute la solution. Ceci souligne le caractère non local des équations de Maxwell.

La résolution de (2.6)-(2.7) est alors basée sur la décomposition de la solution en une partie régulière $\mathbf{H}^1(\omega)$, et une partie singulière qui n'est pas $\mathbf{H}^1(\omega)$ connue entièrement ou en partie explicitement, appelée *le complément singulier*. Cette technique a été développée pour le cas instationnaire par F. Assous, P. Ciarlet et E. Sonnendrücker dans [9], et a fait l'objet de la thèse d'E. Garcia [63]. Dans [68], C. Hazard et S. Lohrengel étudient la méthode du complément singulier appliquée au cas magnétostatique.

Dans le paragraphe 2.4.2, nous présentons une nouvelle méthode de décomposition *non-orthogonale et non-conforme* dans $\mathbf{X}_{\mathbf{E}}^0$ du champ quasi-électrostatique. Le champ électrique est décomposé en une partie $\mathbf{H}^1$ -régulière, et une partie singulière, connue explicitement à un coefficient près, λ qui ne dépend que du domaine et des données. Nous montrons comment simplifier le calcul de ce paramètre dans le paragraphe 2.4.3.

La λ -approche ne peut être généralisée dans le cas instationnaire, il faut choisir une décomposition conforme de $\mathbf{X}_{\mathbf{E}}^0$. L'objet du paragraphe 2.4.4 est la méthode du complément singulier orthogonal, qui s'applique au cas dépendant du temps. Nous comparons notre approche à d'autres méthodes de décompositions conformes dans le paragraphe 2.4.5. Le paragraphe 2.4.6 est consacré à l'étude de la régularité dans les espaces singuliers conformes.

Finalement, nous apportons quelques conclusions dans le paragraphe 2.4.7.

2.4.2 La λ -approche

Pour tout $j \in \{1, \dots, N_{cr}\}$, on notera $\mathbf{x}_j^P := -\alpha_j r_j^{\alpha_j-1} \begin{pmatrix} \sin(\alpha_j \theta_j) \\ \cos(\alpha_j \theta_j) \end{pmatrix}$, exprimé en coordonnées polaires par rapport au coin rentrant O_j .

Lemme 2.27 *Les parties principales $\varphi_{D,j}^P$ et $\varphi_{N,j}^P$ sont telles que $-\mathbf{grad} \varphi_{D,j}^P = \mathbf{rot} \varphi_{N,j}^P = \mathbf{x}_j^P$. On remarque que par construction, on a : $\operatorname{div} \mathbf{x}_j^P = \operatorname{rot} \mathbf{x}_j^P = 0$.*

Théorème 2.28 $\mathbf{E}^0$ se décompose en une partie régulière $\mathbf{E}^R \in \mathbf{H}^1(\omega)$ et une partie singulière de la façon suivante :

$$\mathbf{E}^0 = \mathbf{E}^R + \sum_{j=1}^{N_{cr}} \lambda^j \mathbf{x}_j^P \text{ avec } \forall j, \lambda^j = \lambda_D^j + \lambda_N^j. \quad (2.36)$$

DÉMONSTRATION. D'après l'étude des potentiels ϕ_D^0 et ϕ_N^0 (section 2.2), on a :

$$\mathbf{E}^0 = -\mathbf{grad} \tilde{\phi}_D + \mathbf{rot} \tilde{\phi}_N + \sum_{j=1}^{N_{cr}} \left(\lambda_D^j (-\mathbf{grad} \varphi_{D,j}^P) + \lambda_N^j \mathbf{rot} \varphi_{N,j}^P \right).$$

Posons : $\mathbf{E}^R = -\mathbf{grad} \tilde{\phi}_D + \mathbf{rot} \tilde{\phi}_N \in \mathbf{H}^1(\omega)$. On en déduit (2.36) du lemme précédent. $\square$

Le champ électrique $\mathbf{E} = \mathbf{E}^0 + \mathbf{e}$ se décompose ainsi :

$$\boxed{\mathbf{E} = \tilde{\mathbf{E}} + \sum_{j=1}^{N_{cr}} \lambda^j \mathbf{x}_j^P,} \quad (2.37)$$

où $\tilde{\mathbf{E}} = \mathbf{E}^R + \mathbf{e} \in \mathbf{H}^1(\omega)$ et les $\mathbf{x}_j^P \notin \mathbf{H}^1(\omega)$ sont connus explicitement.

On en déduit que $\mathbf{X}_{\mathbf{E}}$ se décompose en un espace régulier et un espace singulier de dimension N_{cr} . $\tilde{\mathbf{E}}$ satisfait :

$$\operatorname{div} \tilde{\mathbf{E}} = g \text{ dans } \omega, \quad (2.38)$$

$$\operatorname{rot} \tilde{\mathbf{E}} = f \text{ dans } \omega, \quad (2.39)$$

$$\tilde{\mathbf{E}} \cdot \boldsymbol{\tau}_{|\partial\omega} = e - \sum_{j=1}^{N_{cr}} \lambda^j \mathbf{x}_j^P \cdot \boldsymbol{\tau}_{|\partial\omega}. \quad (2.40)$$

Proposition 2.29 *Il existe un relèvement $\mathbf{e}_\lambda$ régulier de $e - \sum_{j=1}^{N_{cr}} \lambda^j \mathbf{x}_j^P \cdot \boldsymbol{\tau}_{|\partial\omega}$.*

DÉMONSTRATION. On a déjà vu qu'il existe un relèvement régulier de e . Fixons j . On a sur les arêtes du coin rentrant O_j :

$$- \mathbf{x}_j^P \cdot \boldsymbol{\tau}_{|A_j^0} = \alpha_j r_j^{\alpha_j-1} \sin(\alpha_j \theta_j)|_{\theta_j=0} = 0, \text{ dans } L^2(A_j^0),$$

$$- \mathbf{x}_j^P \cdot \boldsymbol{\tau}_{|A_j^\Theta} = -\alpha_j r_j^{\alpha_j-1} \sin(\alpha_j \theta_j)|_{\theta_j=\pi/\alpha_j} = 0 \text{ dans } L^2(A_j^\Theta).$$

Sur les autres arêtes, on a : $\mathbf{x}_j^P \cdot \boldsymbol{\tau}_{|\partial\omega \setminus \gamma_i} \in C^\infty(\partial\omega \setminus \gamma_i)$. Par ailleurs, la trace tangentielle de $\mathbf{x}_j^P$ se raccorde aux extrémités des arêtes. On en déduit que $\mathbf{x}_j^P$ et $(1 - \eta_j(r_j)) \mathbf{x}_j^P$ ont même trace tangentielle sur $\partial\omega$. Posons $e_j = (1 - \eta_j(r_j)) \mathbf{x}_j^P \cdot \boldsymbol{\tau}_{|\partial\omega}$. Il existe un relèvement régulier $\mathbf{e}_j$ de e_j .

D'où $\mathbf{e}_\lambda = \mathbf{e} - \sum_{j=1}^{N_{cr}} \lambda^j \mathbf{e}_j$ est un relèvement régulier de $\tilde{\mathbf{E}} \cdot \boldsymbol{\tau}_{|\partial\omega}$. $\square$

Soit $\tilde{\mathbf{E}}^0 \in \mathbf{X}_{\mathbf{E}}^{0,R}$ tel que : $\tilde{\mathbf{E}} = \tilde{\mathbf{E}}^0 + \mathbf{e}_\lambda$. $\tilde{\mathbf{E}}^0 \in \mathbf{X}_{\mathbf{E}}^{0,R}$ satisfait :

$$\operatorname{div} \tilde{\mathbf{E}}^0 = g - \operatorname{div} \mathbf{e}_\lambda \text{ dans } \omega, \quad (2.41)$$

$$\operatorname{rot} \tilde{\mathbf{E}}^0 = f - \operatorname{rot} \mathbf{e}_\lambda \text{ dans } \omega. \quad (2.42)$$

Pour approximer $\mathbf{E}$, on peut procéder comme suit (voir aussi la section 4.8) :

- Calcul des fonctions de bases $s_{D,j}$ et $s_{N,j}$. On approchera $\tilde{s}_{N,j}$ (resp. $\tilde{s}_{D,j}$) par les éléments finis de Lagrange continus P_k .
- Calcul des λ^j pour chaque coin rentrant. Le calcul des intégrales se fera par un schéma d'intégration numérique. La partie singulière $\sum_j \lambda^j \mathbf{x}_j^P$ est maintenant connue.
- Calcul d'un relèvement régulier $\mathbf{e}_\lambda$. On a donc $\mathbf{e}_\lambda + \sum_j \lambda^j \mathbf{x}_j^P$.
- Calcul de $\tilde{\mathbf{E}}^0$ par les éléments finis de Lagrange continus P_k composante par composante. On obtient alors $\mathbf{E} = \tilde{\mathbf{E}}^0 + \mathbf{e}_\lambda + \sum_j \lambda^j \mathbf{x}_j^P$.

Proposition 2.30 *Le problème (2.41)-(2.42) est équivalent à la formulation variationnelle suivante : Trouver $\tilde{\mathbf{E}}^0 \in \mathbf{X}_{\mathbf{E}}^{0,R}$ tel que :*

$$\forall \mathbf{F} \in \mathbf{X}_{\mathbf{E}}^{0,R}, \mathcal{A}_0(\tilde{\mathbf{E}}^0, \mathbf{F}) = \mathcal{L}_0(\mathbf{F}) - \mathcal{A}_0(\mathbf{e}_\lambda, \mathbf{F}), \quad (2.43)$$

où $\mathcal{A}_0$ est la forme bilinéaire symétrique suivante :

$$\begin{aligned} \mathcal{A}_0 : \mathbf{X}_{\mathbf{E}} \times \mathbf{X}_{\mathbf{E}} &\rightarrow \mathbb{R} \\ (\mathbf{E}, \mathbf{F}) &\mapsto (\operatorname{div} \mathbf{E}, \operatorname{div} \mathbf{F})_0 + (\operatorname{rot} \mathbf{E}, \operatorname{rot} \mathbf{F})_0, \end{aligned}$$

et $\mathcal{L}_0$ est la forme linéaire suivante :

$$\begin{aligned} \mathcal{L}_0 : \mathbf{X}_{\mathbf{E}}^0 &\rightarrow \mathbb{R} \\ \mathbf{F} &\mapsto (g, \operatorname{div} \mathbf{F})_0 + (f, \operatorname{rot} \mathbf{F})_0. \end{aligned}$$

Notons que $\mathcal{A}_0$ restreinte à $\mathbf{X}_{\mathbf{E}}^0 \times \mathbf{X}_{\mathbf{E}}^0$ est coercitive (c'est le produit scalaire dans $\mathbf{X}_{\mathbf{E}}^0$).

On peut alors approcher la solution de (2.41)-(2.42) par les éléments finis continus P_k de Lagrange composante par composante dans l'espace $\mathbf{X}_{\mathbf{E}}^{0,R}$ discrétisé. En pratique, on n'a pas besoin de déterminer un relèvement explicite $\mathbf{e}_\lambda$. Nous allons montrer dans la section 4.8 qu'on peut se contenter de l'interpolation sur $\partial\omega$ de la condition aux limites tangentielle au bord, et qu'on peut ainsi se ramener au calcul direct de $\tilde{\mathbf{E}}$.

Étudions la régularité de $\tilde{\mathbf{E}}$. Nous avons vu que lorsque f et $g \in L^2(\omega)$, $\tilde{\mathbf{E}} \in \mathbf{H}^1(\omega)$. On voudrait retrouver le résultat de [52] pour la résolution du problème magnétostatique, c'est-à-dire $\tilde{\mathbf{E}} \in \mathbf{H}^{2\alpha-\epsilon}(\omega)$, $\forall \epsilon > 0$. Pour cela, il faut que les données f et g soient plus régulières. On a le résultat suivant :

Théorème 2.31 *Supposons que pour $\epsilon > 0$ tel que $2\alpha - 1 - \epsilon > 0$, $f, g \in H^{2\alpha-1-\epsilon}(\omega)$ et $e \in H^{2\alpha-1/2-\epsilon}(\omega)$. Alors $\tilde{\mathbf{E}} \in \mathbf{H}^{2\alpha-\epsilon}(\omega)$.*

DÉMONSTRATION. On a $\tilde{\mathbf{E}} = -\mathbf{grad} \tilde{\phi}_D + \mathbf{rot} \tilde{\phi}_N + \mathbf{e}$. Ce théorème est donc un corollaire du théorème 2.20. □

Remarque 2.32 *Cette hypothèse est plus générale que celle de M. Costabel et M. Dauge dans [52] et de C. Hazard et S. Lohrengel dans [68] qui supposent que $\operatorname{div} \mathbf{E} = 0$ et $\operatorname{rot} \mathbf{E} \in H^1(\omega)$. De plus, $e = 0$ dans les travaux sus-mentionnés.*

2.4.3 Simplification du calcul de λ

Dans ce paragraphe, on considère le cas où ω ne présente qu'un seul coin rentrant. Nous allons montrer qu'on peut simplifier le calcul de λ grâce à l'hypothèse que e s'annule au voisinage du coin rentrant. Montrons d'abord le théorème suivant :

Théorème 2.33 *Supposons que $f, g \in L^2(\omega)$ et $e \in V_0^{1/2}(\partial\omega)$. Le champ électrique solution de (2.1)-(2.3) se décompose alors de la façon suivante :*

$$\mathbf{E} = \mathbf{E}_w^R + \frac{1}{\pi} \left((g, s_D)_0 + (f, s_N)_0 + \langle e, s_N \rangle_{H,H'} \right) \mathbf{x}^P \text{ avec } \mathbf{E}_w^R \in \mathbf{H}^1(\omega), \quad (2.44)$$

où H désigne $V_0^{1/2}(\partial\omega)$ et $\langle \cdot, \cdot \rangle_{H,H'}$ est le crochet de dualité entre H et son dual H' .

DÉMONSTRATION. On injecte (2.32) et (2.33) dans la décomposition de Helmholtz (2.11) en utilisant le lemme (2.27). On a alors $\mathbf{E}_w^R = -\mathbf{grad} w_D + \mathbf{rot} w_N \in \mathbf{H}^1(\omega)$. $\square$

Par identification des parties régulières et singulières de (2.44) et (2.37), on en déduit que lorsque $e \in V_0^{1/2}(\partial\omega)$, on a :

$$\lambda = \frac{1}{\pi} ((g^0, s_D)_0 + (f^0, s_N)_0) = \frac{1}{\pi} ((g, s_D)_0 + (f, s_N)_0 + \langle e, s_N \rangle_{H, H'}) .$$

On cherche à retrouver ce résultat dans le cadre de nos hypothèses de modélisation, c'est-à-dire $e \in H^{1/2}(\partial\omega)$ s'annulant au voisinage du coin rentrant. Cela revient donc à montrer le théorème suivant :

Théorème 2.34 *Soit $e \in H^{1/2}(\partial\omega)$ s'annulant dans un voisinage du coin rentrant. Alors :*

$$\lambda = \frac{1}{\pi} \left((f, s_N)_0 + (g, s_D)_0 + \int_{\gamma_0} e s_N d\sigma \right) \quad (2.45)$$

où γ_0 désigne la partie de $\partial\omega$ sur laquelle e ne s'annule pas.

Ce théorème est une généralisation des résultats de [86], importante en pratique car la condition imposée sur e correspond à notre modélisation.

Il s'agit donc de démontrer la proposition suivante :

Proposition 2.35 *Soit $e \in H^{1/2}(\partial\omega)$ s'annulant dans un voisinage du coin rentrant. Il existe un relèvement régulier de e $\mathbf{e} \in \mathbf{H}^1(\omega)$, et tel qu'on ait la formule d'intégration suivante :*

$$-(\mathbf{rot} \mathbf{e}, s_N)_0 - (\mathbf{div} \mathbf{e}, s_D)_0 = \int_{\gamma_0} e s_N d\sigma. \quad (2.46)$$

Pour cela, nous avons besoin des lemmes 2.37 et 2.38 ci-dessous, dont les preuves sont détaillées dans la section 14.3. Nous définissons d'abord la portion de disque $\mathcal{B}_\varepsilon$ de la façon suivante :

Définition 2.36 *On appelle $\mathcal{B}_\varepsilon := \omega \cap B(O, \varepsilon)$: l'intersection de ω avec la boule ouverte de centre le coin rentrant et de rayon $\varepsilon > 0$. Notons que :*

$$\mathcal{B}_\varepsilon = \left\{ (r, \theta) \in]0, \varepsilon[\times \left] 0, \frac{\pi}{\alpha} \right[\right\}.$$

Lemme 2.37 *Soit $\varepsilon > 0$. Soit $u \in H^1(\mathcal{B}_\varepsilon)$, dont la trace sur $\partial\mathcal{B}_\varepsilon$ s'annule au voisinage du coin rentrant. Posons $\phi = r^{-\alpha/2} u$. Alors si on note $\epsilon_\alpha = (1 - \alpha)/2 \in]0, 1/4[$, on a $\phi \in H^{1/2+\epsilon_\alpha}(\mathcal{B}_\varepsilon)$.*

La preuve de ce lemme repose sur deux théorèmes d'injections continues, énoncés dans la thèse de G. Raugel [95] et dans le livre de P. Grisvard [66].

Lemme 2.38 *Soient $\epsilon_0 > 0$ et $\varepsilon > 0$. Posons $H_\varepsilon = H^{1/2+\epsilon_0}(\mathcal{B}_\varepsilon)$. Alors on a :*

$$\forall f \in H'_\varepsilon, g \in H_\varepsilon \lim_{\varepsilon \rightarrow 0} | \langle f, g \rangle_{H'_\varepsilon, H_\varepsilon} | = 0.$$

Nous avons besoin de plus de la propriété suivante, démontrée par E. Garcia dans [63] (p. 87) :

Propriété 2.39 *On a les relations suivantes :*

$$\begin{aligned} \mathbf{rot} s_N &= \mathbf{grad} s_D, \\ \mathbf{rot} s_D &= -\mathbf{grad} s_N. \end{aligned}$$

Nous allons maintenant démontrer la proposition 2.35.

DÉMONSTRATION DE LA PROPOSITION 2.35. Soit $\varepsilon > 0$ destiné à tendre vers 0. Soit $\omega_\varepsilon = \omega \setminus \mathcal{B}_\varepsilon$ (voir la figure 2.1). Posons $b_\varepsilon = \partial\omega_\varepsilon \setminus \partial\omega$. Notons que :

$$b_\varepsilon = \left\{ (r, \theta) : r = \varepsilon, \theta \in \left] 0, \frac{\pi}{\alpha} \right[\right\}.$$

Comme e s'annule au voisinage du coin rentrant, il existe $\varepsilon_0 > 0$ tel que presque pour tout $r \leq \varepsilon_0$, $e(r, 0) = e(r, \pi/\alpha) = 0$. On appelle γ_0 l'ensemble formé de $\partial\omega$, privé des points de $\partial\omega \cap \partial\mathcal{B}_{\varepsilon_0}$:

$$\gamma_0 = \{ (r, \theta) \in \partial\omega : r > \varepsilon_0 \}.$$

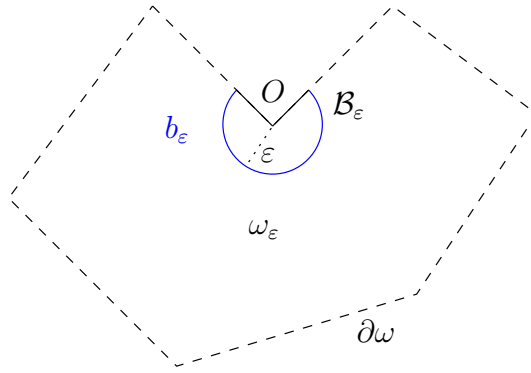


FIG. 2.1 – Décomposition du domaine ω .

• Montrons que :

$$-(\operatorname{rot} \mathbf{e}, s_N)_0 - (\operatorname{div} \mathbf{e}, s_D)_0 = \int_{\gamma_0} e s_N d\sigma + \lim_{\varepsilon \rightarrow 0} \int_{b_\varepsilon} (e_\tau r^{-\alpha} \cos(\alpha\theta) - e_\nu r^{-\alpha} \sin(\alpha\theta)) d\sigma. \quad (2.47)$$

Écrivons l'intégrale sur ω ainsi : $-(\operatorname{rot} \mathbf{e}, s_N)_0 - (\operatorname{div} \mathbf{e}, s_D)_0 = -\lim_{\varepsilon \rightarrow 0} \int_{\omega_\varepsilon} (\operatorname{rot} \mathbf{e} s_N d\omega + \operatorname{div} \mathbf{e} s_D) d\omega$.

Par intégration par parties sur ω_ε , on obtient :

$$-\int_{\omega_\varepsilon} \operatorname{rot} \mathbf{e} s_N d\omega = -\int_{\omega_\varepsilon} \mathbf{e} \cdot \operatorname{rot} s_N d\omega + \int_{\partial\omega_\varepsilon} e_\tau s_N d\sigma,$$

et :

$$-\int_{\omega_\varepsilon} \operatorname{div} \mathbf{e} s_D d\omega = \int_{\omega_\varepsilon} \mathbf{e} \cdot \operatorname{grad} s_D d\omega - \int_{\partial\omega_\varepsilon} e_\nu s_D d\sigma.$$

D'après la propriété 2.39, on en déduit :

$$-\int_{\omega_\varepsilon} \operatorname{rot} \mathbf{e} s_N d\omega - \int_{\omega_\varepsilon} \operatorname{div} \mathbf{e} s_D d\omega = \int_{\partial\omega_\varepsilon} e_\tau s_N d\sigma - \int_{\partial\omega_\varepsilon} e_\nu s_D d\sigma.$$

De plus, comme $s_D|_{\partial\omega} = 0$, on a : $\int_{\partial\omega_\varepsilon} e_\nu s_D d\sigma = \int_{b_\varepsilon} e_\nu s_D d\sigma$. D'où :

$$-\int_{\omega_\varepsilon} (\operatorname{rot} \mathbf{e} s_N d\omega + \operatorname{div} \mathbf{e} s_D) d\omega = \int_{\partial\omega_\varepsilon} e_\tau s_N d\sigma - \int_{b_\varepsilon} e_\nu s_D d\sigma.$$

Lorsque ε tend vers 0, le terme de gauche admet pour limite $-(\operatorname{rot} \mathbf{e}, s_N)_0 - (\operatorname{div} \mathbf{e}, s_D)_0$. Il en est donc de même pour le terme de droite :

$$-(\operatorname{rot} \mathbf{e}, s_N)_0 - (\operatorname{div} \mathbf{e}, s_D)_0 = \lim_{\varepsilon \rightarrow 0} \left(\int_{\partial\omega_\varepsilon} e_\tau s_N d\sigma - \int_{b_\varepsilon} e_\nu s_D d\sigma \right),$$

ce que l'on peut écrire sous la forme :

$$\lim_{\varepsilon \rightarrow 0} \left(\int_{\partial\omega_\varepsilon \setminus b_\varepsilon} e s_N d\sigma + \int_{b_\varepsilon} (e_\tau s_N - e_\nu s_D) d\sigma \right) = -(\operatorname{rot} \mathbf{e}, s_N)_0 - (\operatorname{div} \mathbf{e}, s_D)_0. \quad (2.48)$$

Comme $s_N|_{\partial\omega_\varepsilon}$ et $e|_{\partial\omega_\varepsilon} \in L^2(\partial\omega_\varepsilon)$, et que e s'annule au voisinage du coin rentrant, le premier terme de gauche de (2.48) admet une limite égale à l'intégrale sur γ_0 :

$$\forall \varepsilon < \varepsilon_0, \int_{\partial\omega_\varepsilon \setminus b_\varepsilon} e s_N d\sigma = \int_{\gamma_0} e s_N d\sigma.$$

Par conséquent, le second terme de gauche de (2.48) : $\lim_{\varepsilon \rightarrow 0} \int_{b_\varepsilon} (e_\tau s_N - e_\nu s_D) d\sigma$ existe.

Séparons cette intégrale en une partie singulière et une partie régulière :

$$\int_{b_\varepsilon} (e_\tau s_N - e_\nu s_D) d\sigma = \int_{b_\varepsilon} (e_\tau r^{-\alpha} \cos(\alpha \theta) - e_\nu r^{-\alpha} \sin(\alpha \theta)) d\sigma + \int_{b_\varepsilon} (e_\tau \tilde{s}_N - e_\nu \tilde{s}_D) d\sigma.$$

Montrons que la partie régulière s'annule lorsque $\varepsilon \rightarrow 0$. Pour cela, nous allons nous ramener à une intégrale volumique. Tout d'abord, comme $\tilde{s}_D|_{A^0} = \tilde{s}_D|_{A^\Theta} = 0$, on a :

$$- \int_{b_\varepsilon} e_\nu \tilde{s}_D d\sigma = - \int_{\partial\mathcal{B}_\varepsilon} e_\nu \tilde{s}_D d\sigma. \quad (2.49)$$

D'autre part, pour $r < \varepsilon_0$, $e_\tau|_{A^0} = e_\tau|_{A^\Theta} = 0$, d'où : $\forall \varepsilon < \varepsilon_0$,

$$\int_{b_\varepsilon} e_\tau \tilde{s}_N d\sigma = \int_{\partial\mathcal{B}_\varepsilon} e_\tau \tilde{s}_N d\sigma. \quad (2.50)$$

En sommant (2.49) et (2.50), et en intégrant par parties, on obtient alors :

$$\int_{b_\varepsilon} (e_\tau \tilde{s}_N - e_\nu \tilde{s}_D) d\sigma = \int_{\mathcal{B}_\varepsilon} ((\mathbf{e} \cdot \operatorname{rot} \tilde{s}_N - \operatorname{rot} \mathbf{e} \tilde{s}_N) - (\mathbf{e} \cdot \operatorname{grad} \tilde{s}_D + \operatorname{div} \mathbf{e} \tilde{s}_D)) d\omega.$$

Toutes les fonctions sous intégrale sont dans $L^2(\mathcal{B}_\varepsilon)$, d'où $\int_{b_\varepsilon} (e_\tau \tilde{s}_N - e_\nu \tilde{s}_D) d\sigma = 0$.

• Nous avons donc obtenu (2.47). Nous allons maintenant montrer que la limite dans (2.47) s'annule. Pour cela, nous allons nous ramener à une intégrale volumique sur $\mathcal{B}_\varepsilon$. Comme $\sin(\alpha \theta)$ est nul sur les arêtes du coin rentrant A^0 et A^Θ , alors :

$$- \int_{b_\varepsilon} e_\nu r^{-\alpha} \sin(\alpha \theta) d\sigma = - \int_{\partial\mathcal{B}_\varepsilon} e_\nu r^{-\alpha} \sin(\alpha \theta) d\sigma.$$

De plus, pour la même raison que (2.50), on a : $\forall \varepsilon < \varepsilon_0$,

$$\int_{b_\varepsilon} e_\tau r^{-\alpha} \cos(\alpha \theta) d\sigma = \int_{\partial\mathcal{B}_\varepsilon} e_\tau r^{-\alpha} \cos(\alpha \theta) d\sigma.$$

Soit $t_N = r^{-\alpha/2} \cos(\alpha \theta)$, et $t_D = r^{-\alpha/2} \sin(\alpha \theta)$. Posons $\phi = r^{-\alpha/2} \mathbf{e}$. Il s'agit d'évaluer :

$$\int_{\partial \mathcal{B}_\varepsilon} (\phi \cdot \boldsymbol{\tau} t_N - \phi \cdot \boldsymbol{\nu} t_D) d\sigma. \quad (2.51)$$

D'après [67] (p. 7, thm. 1.2.18), t_N et $t_D \in H^s(\omega)$, pour tout s tel que $1 - \alpha/2 > s$.

Soit $\epsilon' \in]0, (1 - \alpha)/2[$, de sorte que $s = 1/2 + \epsilon'$ convienne.

On a donc t_N et $t_D \in H^{1/2+\epsilon'}(\omega)$, c'est-à-dire $t_D|_{\partial \mathcal{B}_\varepsilon}$ et $t_N|_{\partial \mathcal{B}_\varepsilon} \in H^{\epsilon'}(\partial \mathcal{B}_\varepsilon)$.

D'après le lemme 2.37, pour $\epsilon_\alpha = (1 - \alpha)/2 > 0$, $\phi \in \mathbf{H}^{1/2+\epsilon_\alpha}(\mathcal{B}_\varepsilon)$, d'où $\phi|_{\partial \mathcal{B}_\varepsilon} \in \mathbf{H}^{\epsilon_\alpha}(\partial \mathcal{B}_\varepsilon)$. Ainsi, $\phi \cdot \boldsymbol{\tau}|_{\partial \mathcal{B}_\varepsilon}$ et $\phi \cdot \boldsymbol{\nu}|_{\partial \mathcal{B}_\varepsilon} \in L^2(\partial \mathcal{B}_\varepsilon)$. L'intégrale (2.51) est bien posée.

Par densité de $C^\infty(\bar{\mathcal{B}}_\varepsilon)$ dans $H^{1/2+\epsilon_\alpha}(\mathcal{B}_\varepsilon)$, on généralise les formules de Green (1.11) et (1.12) aux fonctions $\mathbf{u} \in \mathbf{H}^{1/2+\epsilon_\alpha}(\mathcal{B}_\varepsilon)$ et $v \in H^{1/2+\epsilon_\alpha}(\mathcal{B}_\varepsilon)$. D'où :

$$\begin{aligned} \int_{\partial \mathcal{B}_\varepsilon} (\phi \cdot \boldsymbol{\tau} t_N - \phi \cdot \boldsymbol{\nu} t_D) d\sigma &= \langle \mathbf{rot} t_N, \phi \rangle_{\mathbf{H}'_\varepsilon, \mathbf{H}_\varepsilon} - \langle \mathbf{rot} \phi, t_N \rangle_{\mathbf{H}'_\varepsilon, \mathbf{H}_\varepsilon} \\ &\quad - \langle \mathbf{div} \phi, t_D \rangle_{\mathbf{H}'_\varepsilon, \mathbf{H}_\varepsilon} - \langle \mathbf{grad} t_D, \phi \rangle_{\mathbf{H}'_\varepsilon, \mathbf{H}_\varepsilon}, \end{aligned}$$

où $\mathbf{H}_\varepsilon = H^{1/2+\epsilon_\alpha}(\mathcal{B}_\varepsilon)$, et $\langle \cdot, \cdot \rangle_{\mathbf{H}'_\varepsilon, \mathbf{H}_\varepsilon}$ est le crochet de dualité entre $\mathbf{H}_\varepsilon$ et son dual $\mathbf{H}'_\varepsilon$.

D'après le lemme 2.38, cette intégrale s'annule. On en déduit la proposition 2.35. □

2.4.4 Le complément singulier orthogonal

Nous avons exhibé une décomposition non orthogonale du champ électrique $\mathbf{E}^0$ en une partie $\mathbf{H}^1(\omega)$ -régulière et une partie singulière. Cette décomposition n'est pas conforme dans $\mathbf{X}_{\mathbf{E}}^0$ car aucune de ces parties ne satisfait la condition aux limites de conducteur parfait. Quand on travaille sur les équations de Maxwell instationnaires ou harmoniques, il est plus facile d'utiliser une décomposition conforme $\mathbf{X}_{\mathbf{E}}^0$. Dans ce paragraphe, nous détaillons une méthode de décomposition conforme et orthogonale, appelée la méthode du complément singulier orthogonal (notée MCSO).

La MCSO a été largement développée par E. Garcia dans [63], à partir des travaux de F. Assous et al. [9, 8]. Afin de déterminer l'espace $\mathbf{X}_{\mathbf{E}}^{0,S}$, le supplémentaire orthogonal de $\mathbf{X}_{\mathbf{E}}^{0,R}$ dans $\mathbf{X}_{\mathbf{E}}^0$, étudions $\mathbf{W}_{\mathbf{E}}$ et $\mathbf{V}_{\mathbf{E}}$.

Proposition 2.40 *On a les isomorphismes suivants (voir la figure (2.2)) :*

- L'opérateur **grad** définit un isomorphisme de Φ_D dans $\mathbf{W}_{\mathbf{E}}$,
- L'opérateur **rot** définit un isomorphisme de Φ_N dans $\mathbf{V}_{\mathbf{E}}$,
- L'opérateur **div** définit un isomorphisme de $\mathbf{W}_{\mathbf{E}}$ dans $L^2(\omega)$,
- L'opérateur **rot** définit un isomorphisme de $\mathbf{V}_{\mathbf{E}}$ dans $L_0^2(\omega)$.

$\mathbf{W}_{\mathbf{E}}$ étant muni de la norme L^2 de la divergence, $\mathbf{V}_{\mathbf{E}}$ étant muni de la norme L^2 du rotationnel ; Φ_D et Φ_N étant munis de la norme L^2 du Laplacien, et $L_0^2(\omega)$ étant muni de la norme L^2 , ces isomorphismes préservent l'orthogonalité.

DÉMONSTRATION. Voir le document [6] ou l'article [8]. □

Lorsque ω est convexe, $\mathbf{W}_{\mathbf{E}}$ et $\mathbf{V}_{\mathbf{E}}$ sont inclus dans $\mathbf{H}^1(\omega)$, mais ce n'est plus vrai lorsqu'il existe un ou plusieurs coin(s) rentrant(s). Soit $\mathbf{W}_{\mathbf{E}}^R = \mathbf{W}_{\mathbf{E}} \cap \mathbf{H}^1(\omega)$ (resp. $\mathbf{V}_{\mathbf{E}}^R = \mathbf{V}_{\mathbf{E}} \cap \mathbf{H}^1(\omega)$) le sous-espace régularisé de $\mathbf{W}_{\mathbf{E}}$ (resp. $\mathbf{V}_{\mathbf{E}}$). $\mathbf{W}_{\mathbf{E}}^R$ (resp. $\mathbf{V}_{\mathbf{E}}^R$) est fermé et strictement inclus dans $\mathbf{W}_{\mathbf{E}}$ (resp. $\mathbf{V}_{\mathbf{E}}$). D'après les propositions 2.8 et 2.40, on a $\Delta \Phi_D^R = \mathbf{div} \mathbf{W}_{\mathbf{E}}^R$ et $\Delta \Phi_N^R = \mathbf{rot} \mathbf{V}_{\mathbf{E}}^R$, d'où :



FIG. 2.2 – Isomorphismes entre espaces potentiels et espaces vectoriels.

$S_D = (\operatorname{div} \mathbf{W}_E^R)^\perp$, et $S_N = (\operatorname{rot} \mathbf{V}_E^R)^\perp$. Ce résultat nous permet de caractériser les fonctions singulières de $\mathbf{W}_E$ et $\mathbf{V}_E$:

$$\mathbf{W}_E = \mathbf{W}_E^R \overset{\perp_{\operatorname{div}}}{\oplus} \mathbf{W}_E^S \text{ où } \mathbf{W}_E^S = \operatorname{div}^{-1} S_D \text{ et } \mathbf{V}_E = \mathbf{V}_E^R \overset{\perp_{\operatorname{rot}}}{\oplus} \mathbf{V}_E^S \text{ où } \mathbf{V}_E^S = \operatorname{rot}^{-1} S_N.$$

On peut donc appliquer les isomorphismes de la proposition 2.40 aux espaces réguliers et singuliers (voir la figure 2.3). On en déduit : $\mathbf{W}_E^{R,S} = \operatorname{grad} \Phi_D^{R,S}$ et $\mathbf{V}_E^{R,S} = \operatorname{rot} \Phi_N^{R,S}$.

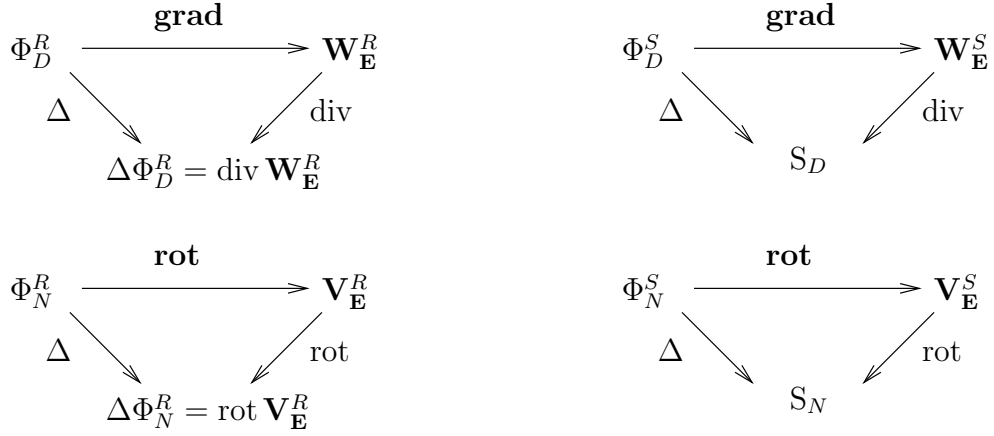


FIG. 2.3 – Isomorphismes entre espaces réguliers/singuliers.

Théorème 2.41 *D'après le théorème 2.14, on établit les estimations suivantes :*

- $\forall \varepsilon > 0$, $\mathbf{u} \in \mathbf{W}_E^S$ appartient à $\mathbf{H}^{\alpha-\varepsilon}(\omega)$, et n'appartient pas à $\mathbf{H}^\alpha(\omega)$; $\mathbf{W}_E^R \subset \mathbf{H}^1(\omega)$.

- $\forall \varepsilon > 0$, $\mathbf{u} \in \mathbf{V}_E^S$ appartient à $\mathbf{H}^{\alpha-\varepsilon}(\omega)$, et n'appartient pas à $\mathbf{H}^\alpha(\omega)$; $\mathbf{V}_E^R \subset \mathbf{H}^1(\omega)$.

Ce théorème est un corollaire trivial du théorème 2.14.

D'après la décomposition (2.17), $\mathbf{X}_E^0$ se décompose quant à lui de la façon suivante :

$$\mathbf{X}_E^0 = \operatorname{grad} \Phi_D^R \overset{\perp_{\mathbf{X}_E^0}}{\oplus} \operatorname{rot} \Phi_N^R \overset{\perp_{\mathbf{X}_E^0}}{\oplus} \operatorname{grad} \Phi_D^S \overset{\perp_{\mathbf{X}_E^0}}{\oplus} \operatorname{rot} \Phi_N^S.$$

Comme $\operatorname{grad} \Phi_D^R \overset{\perp}{\oplus} \operatorname{rot} \Phi_N^R \subset \mathbf{X}_E^{0,R}$, alors $\mathbf{X}_E^{0,S} \subset \operatorname{grad} \Phi_D^S \overset{\perp}{\oplus} \operatorname{rot} \Phi_N^S$. On remarque que $\operatorname{grad} \Phi_D^S \overset{\perp}{\oplus} \operatorname{rot} \Phi_N^S$ est de dimension $\geq 2 N_{cr}$. D'après la λ -approche, $\mathbf{X}_E^{0,S}$ est de dimension N_{cr} . Il existe donc un sous-espace de $\operatorname{grad} \Phi_D^S \overset{\perp}{\oplus} \operatorname{rot} \Phi_N^S$ de dimension N_{cr} , composé de champs réguliers.

Ceci étant dit, on peut écrire $\mathbf{x} \in \mathbf{X}_E^{0,S}$ sous la forme : $\mathbf{x} = \sum_{i=1}^{N_{cr}} (a_i \operatorname{grad} \varphi_{D,i} + b_i \operatorname{rot} \varphi_{N,i})$,

les a_i et b_i étant des réels. Les a_i et b_i ne sont pas quelconques, puisqu'il existe des combinaisons linéaires qui sont régulières. Ceci est précisé dans le théorème 2.43, pour lequel on a besoin du lemme suivant :

Lemme 2.42 *La matrice $\mathbb{X} \in \mathbb{R}^{N_{cr} \times N_{cr}}$ telle que : $\forall i, j \ \mathbb{X}^{i,j} = \beta^{i,j}$ est symétrique définie positive.*

DÉMONSTRATION. Posons $\mathbb{X} = \mathbb{X}_D + \mathbb{X}_N$, avec $\mathbb{X}_{D,N}^{i,j} = \beta_{D,N}^{i,j}$. $\mathbb{X}_{D,N}$ est la matrice des produits scalaires des fonctions de base de $S_{D,N}$. Il est clair qu'elle est symétrique. Montrons que $\mathbb{X}_D$ est définie positive. Soit $\underline{x} \in \mathbb{R}^{N_{cr}}$. On a :

$$\begin{aligned} (\mathbb{X}_D \underline{x} | \underline{x}) &= \sum_{i=1}^{N_{cr}} \sum_{j=1}^{N_{cr}} \mathbb{X}_D^{i,j} \underline{x}_i \underline{x}_j = \sum_{i=1}^{N_{cr}} \sum_{j=1}^{N_{cr}} (s_{D,i}, s_{D,j})_0 / \pi \underline{x}_i \underline{x}_j, \\ &= \sum_{i=1}^{N_{cr}} \sum_{j=1}^{N_{cr}} (\underline{x}_i s_{D,i}, \underline{x}_j s_{D,j})_0 / \pi = \left(\sum_{i=1}^{N_{cr}} \underline{x}_i s_{D,i}, \sum_{j=1}^{N_{cr}} \underline{x}_j s_{D,j} \right)_0 / \pi, \\ &= (x, x)_0 \geq 0, \text{ avec } x = \sum_{i=1}^{N_{cr}} \underline{x}_i s_{D,i}. \end{aligned}$$

Or, comme $(s_{D,i})_{i=1,\dots,N_{cr}}$ est une base de S_D , $x \in S_D$ est nul si et seulement si $\forall i, \underline{x}_i = 0$. Donc $\forall \underline{x} \in \mathbb{R}^{N_{cr}}$, tel que $\underline{x} \neq 0$, $(\mathbb{X}_D \underline{x} | \underline{x}) > 0$. On montre exactement de la même façon que $\mathbb{X}_N$ est définie positive. D'où, $\mathbb{X}_D$, $\mathbb{X}_N$ et $\mathbb{X}$ sont symétriques définies positives. $\square$

Théorème 2.43 *On considère les vecteurs $\mathbf{x}_i^S = -\mathbf{grad} \varphi_{D,i} + \mathbf{rot} \varphi_{N,i}$, $i \in \{1, \dots, N_{cr}\}$. Alors $\mathbf{X}_{\mathbf{E}}^{0,S}$ est généré par les $\mathbf{x}_i^S : \mathbf{X}_{\mathbf{E}}^{0,S} = \text{vect}(\mathbf{x}_i^S)$.*

DÉMONSTRATION. On remarque que : $\text{div} \mathbf{x}_i^S = -\Delta \varphi_{D,i} = s_{D,i}$, et que $\text{rot} \mathbf{x}_i^S = -\Delta \varphi_{N,i} = s_{N,i}$. On en déduit que $\forall i \in \{1, \dots, N_{cr}\}$, $\mathbf{x}_i^S \in \mathbf{X}_{\mathbf{E}}^0$ satisfait :

$$\begin{aligned} \text{div} \mathbf{x}_i^S &= s_{D,i} \text{ dans } \omega, \\ \text{rot} \mathbf{x}_i^S &= s_{N,i} \text{ dans } \omega. \end{aligned} \quad (2.52)$$

- Montrons que pour tout i , $\mathbf{x}_i^S \notin \mathbf{X}_{\mathbf{E}}^{0,R}$.

D'après la λ -approche, on a : $\mathbf{x}_i^S = \tilde{\mathbf{x}}_i + \sum_{j=1}^{N_{cr}} \beta^{i,j} \mathbf{x}_j^P$, où $\tilde{\mathbf{x}}_i = -\mathbf{grad} \tilde{\varphi}_{D,i} + \mathbf{rot} \tilde{\varphi}_{N,i} \in \mathbf{H}^1(\omega)$ et

$\beta^{i,j} = (s_{D,i}, s_{D,j})_0 / \pi + (s_{N,i}, s_{N,j})_0 / \pi$. D'où, au voisinage du coin rentrant O_i , $\mathbf{x}_i^S \simeq \beta^{i,i} \mathbf{x}_i^P$ avec $\beta^{i,i} = (\|s_{D,i}\|_0^2 + \|s_{N,i}\|_0^2) / \pi \neq 0$, et donc $\mathbf{x}_i^S \notin \mathbf{X}_{\mathbf{E}}^{0,R}$.

- Montrons que pour tout i , $\mathbf{x}_i^S \in (\mathbf{X}_{\mathbf{E}}^{0,R})^\perp$.

Soit $\mathbf{x}^R \in \mathbf{X}_{\mathbf{E}}^{0,R}$. D'après la λ -approche, on peut décomposer $\mathbf{x}^R$ sous la forme : $\mathbf{x}^R = \tilde{\mathbf{x}} + \sum_j \lambda_R^j \mathbf{x}_j^P$, où $\tilde{\mathbf{x}} \in \mathbf{H}^1(\omega)$ et : $\lambda_R^j = ((s_{D,j}, \text{div} \mathbf{x}^R)_0 + (s_{N,j}, \text{rot} \mathbf{x}^R)_0) / \pi$. Comme $\mathbf{x}^R$ est régulier au voisinage de chaque coin O_j , on en déduit que tous les λ_R^j sont nuls. Or, d'après (2.52) on a que pour tout j , $\lambda_R^j = (\mathbf{x}_j^S, \mathbf{x}^R)_{\mathbf{X}_{\mathbf{E}}^0} / \pi$ est nul. Ainsi, par construction, les $\mathbf{x}_j^S$ sont dans $\mathbf{X}_{\mathbf{E}}^{0,S}$.

- Montrons que la familles des $\mathbf{x}_i^S$ est libre.

Soit $(a^i)_{i=1,\dots,N_{cr}} \in \mathbb{R}^{N_{cr}}$ tel que : $\sum_i a^i \mathbf{x}_i^S = 0$. D'où : $\sum_i a^i \tilde{\mathbf{x}}_i + \sum_i a^i \sum_j \beta^{i,j} \mathbf{x}_j^P = 0$. En se plaçant une nouvelle fois au voisinage de chaque coin O_j , on en déduit que $\forall j \ \sum_i a^i \beta^{i,j} = 0$.

Soit $\mathbf{a} \in \mathbb{R}^{N_{cr}}$ tel que : $\forall i \mathbf{a}^i = a^i$. Il s'agit de résoudre le système linéaire : $\mathbb{X}\mathbf{a} = 0$. Or $\mathbb{X}$ est inversible, et la seule solution est $\mathbf{a} = 0$. La famille $(\mathbf{x}_i^S)_{i=1, \dots, N_{cr}}$ est libre dans $\mathbf{X}_{\mathbf{E}}^{0,S}$. Comme son cardinal est la dimension de $\mathbf{X}_{\mathbf{E}}^{0,S}$, elle génère $\mathbf{X}_{\mathbf{E}}^{0,S}$.

□

On déduit de la décomposition de $\mathbf{X}_{\mathbf{E}}^0$ le résultat suivant :

Théorème 2.44 *Le champ électrique $\mathbf{E}^0$ se décompose en une partie régulière $\hat{\mathbf{E}}^0 \in \mathbf{H}^1(\omega)$ et une partie singulière dans $\mathbf{X}_{\mathbf{E}}^{0,S}$ de la façon suivante :*

$$\mathbf{E}^0 = \hat{\mathbf{E}}^0 + \sum_{i=1}^{N_{cr}} c_i \mathbf{x}_i^S \quad \text{où } \forall i \in \{1, \dots, N_{cr}\}, c_i \in \mathbb{R}. \quad (2.53)$$

Notons que $\hat{\mathbf{E}}^0$ et les c_i satisfont les équations :

$$\operatorname{div} \hat{\mathbf{E}}^0 + \sum_{i=1}^{N_{cr}} c_i \operatorname{div} \mathbf{x}_i^S = g - \operatorname{div} \mathbf{e} \text{ dans } \omega, \quad (2.54)$$

$$\operatorname{rot} \hat{\mathbf{E}}^0 + \sum_{i=1}^{N_{cr}} c_i \operatorname{rot} \mathbf{x}_i^S = f - \operatorname{rot} \mathbf{e} \text{ dans } \omega. \quad (2.55)$$

Ainsi, le champ électrique $\mathbf{E} = \mathbf{E}^0 + \mathbf{e}$ se décompose de la façon suivante :

$$\boxed{\mathbf{E} = \hat{\mathbf{E}} + \sum_{i=1}^{N_{cr}} c_i \mathbf{x}_i^S,} \quad (2.56)$$

avec $\hat{\mathbf{E}} = \hat{\mathbf{E}}^0 + \mathbf{e}$.

Soient $\mathbf{c}$ et $\boldsymbol{\lambda} \in \mathbb{R}^{N_{cr}}$ tels que : $\forall i \in \{1, \dots, N_{cr}\}, (\mathbf{c})_i = c_i$ et $(\boldsymbol{\lambda})_i = \lambda^i$.

Proposition 2.45 *Le problème (2.54)-(2.55) est équivalent à la formulation variationnelle suivante : Trouver $\hat{\mathbf{E}}^0 \in \mathbf{X}_{\mathbf{E}}^{0,R}$ et $\mathbf{c} \in \mathbb{R}^{N_{cr}}$ tels que :*

$$\mathcal{A}_0(\hat{\mathbf{E}}^0, \mathbf{F}) = \mathcal{L}_0(\mathbf{F}) - \mathcal{A}_0(\mathbf{e}, \mathbf{F}), \quad \forall \mathbf{F} \in \mathbf{X}_{\mathbf{E}}^{0,R}, \quad (2.57)$$

$$\mathbb{X}\mathbf{c} = \boldsymbol{\lambda}. \quad (2.58)$$

DÉMONSTRATION. Le problème (2.54)-(2.55), posé dans $\mathbf{X}_{\mathbf{E}}^{0,R} \overset{\perp \mathbf{x}_{\mathbf{E}}^0}{\oplus} \mathbf{X}_{\mathbf{E}}^{0,S}$ est équivalent à la formulation variationnelle suivante :

Trouver $\hat{\mathbf{E}}^0 \in \mathbf{X}_{\mathbf{E}}^{0,R}$ et $\mathbf{c} \in \mathbb{R}^{N_{cr}}$ tels que :

$$\mathcal{A}_0(\hat{\mathbf{E}}^0, \mathbf{F}) + \sum_{i=1}^{N_{cr}} c_i \mathcal{A}_0(\mathbf{x}_i^S, \mathbf{F}) = \mathcal{L}_0(\mathbf{F}) - \mathcal{A}_0(\mathbf{e}, \mathbf{F}), \quad \forall \mathbf{F} \in \mathbf{X}_{\mathbf{E}}^{0,R},$$

$$\mathcal{A}_0(\hat{\mathbf{E}}^0, \mathbf{x}_j^S) + \sum_{i=1}^{N_{cr}} c_i \mathcal{A}_0(\mathbf{x}_i^S, \mathbf{x}_j^S) = \mathcal{L}_0(\mathbf{x}_j^S) - \mathcal{A}_0(\mathbf{e}, \mathbf{x}_j^S), \quad \forall j \in \{1, \dots, N_{cr}\}.$$

- Par orthogonalité dans $\mathbf{X}_{\mathbf{E}}^0$, on a : $\forall \mathbf{F} \in \mathbf{X}_{\mathbf{E}}^{0,R}$ et $\forall i$, $\mathcal{A}_0(\mathbf{x}_i^S, \mathbf{F}) = 0$; et $\forall j$, $\mathcal{A}_0(\hat{\mathbf{E}}^0, \mathbf{x}_j^S) = 0$.
- D'après (2.52) on a : $\forall i$, $\mathcal{L}_0(\mathbf{x}_i^S) - \mathcal{A}_0(\mathbf{e}, \mathbf{x}_i^S) = (g - \operatorname{div} \mathbf{e}, \operatorname{div} \mathbf{x}_i^S)_0 + (f - \operatorname{rot} \mathbf{e}, \operatorname{rot} \mathbf{x}_i^S)_0 = \pi \lambda^i$; et $\forall i, j$, $\mathcal{A}_0(\mathbf{x}_i^S, \mathbf{x}_j^S) = \pi (\beta_D^{i,j} + \beta_N^{i,j})$. En simplifiant par π , on obtient alors le système (2.57)-(2.58).

□

Une fois que l'on a calculé $\hat{\mathbf{E}}^0$ et les c_i , il faut construire la base de $\mathbf{X}_{\mathbf{E}}^{0,S}$, constituée des $\mathbf{x}_i^S$ qui ne sont connus explicitement qu'en partie. Notons que pour tout i :

$$\mathbf{x}_i^S = \tilde{\mathbf{x}}_i + \sum_{j=1}^{N_{cr}} \beta^{i,j} \mathbf{x}_j^P \text{ avec } \tilde{\mathbf{x}}_i = -\mathbf{grad} \tilde{\varphi}_D + \mathbf{rot} \tilde{\varphi}_N \in \mathbf{H}^1(\omega).$$

Pour connaître complètement les $\mathbf{x}_i^S$, il faut approcher leurs parties régulières, les $\tilde{\mathbf{x}}_i$. Cela peut se faire par dérivation des potentiels $\varphi_{D,i}$ et $\varphi_{N,i}$, ou par la λ -approche. Comme les $\mathbf{x}_j^P$ sont à rotationnel et divergence nuls, on remarque que : $\forall i \in \{1, \dots, N_{cr}\}$,

$$\operatorname{div} \tilde{\mathbf{x}}_i = s_{D,i} \text{ dans } \omega, \quad (2.59)$$

$$\operatorname{rot} \tilde{\mathbf{x}}_i = s_{N,i} \text{ dans } \omega, \quad (2.60)$$

$$\tilde{\mathbf{x}}_i \cdot \boldsymbol{\tau}_{|\partial\omega} = - \sum_{j=1}^{N_{cr}} \beta^{i,j} \mathbf{x}_j^P \cdot \boldsymbol{\tau}_{|\partial\omega} \text{ sur } \partial\omega. \quad (2.61)$$

Posons : $\tilde{\mathbf{x}}_i = \tilde{\mathbf{x}}_i^0 + \mathbf{e}_i$, où $\tilde{\mathbf{x}}_i^0 \in \mathbf{X}_{\mathbf{E}}^{0,R}$ et $\mathbf{e}_i$ est un relèvement régulier de la condition aux limites tangentielle $-\sum_{j=1}^{N_{cr}} \beta^{i,j} \mathbf{x}_j^P \cdot \boldsymbol{\tau}_{|\partial\omega}$. $\tilde{\mathbf{x}}_i^0$ satisfait :

$$\operatorname{div} \tilde{\mathbf{x}}_i^0 = s_{D,i} - \operatorname{div} \mathbf{e}_i \text{ dans } \omega, \quad (2.62)$$

$$\operatorname{rot} \tilde{\mathbf{x}}_i^0 = s_{N,i} - \operatorname{rot} \mathbf{e}_i \text{ dans } \omega. \quad (2.63)$$

Proposition 2.46 *Pour tout i , le problème (2.59)-(2.61) est équivalent à la formulation variationnelle : Trouver $\tilde{\mathbf{x}}_i^0 \in \mathbf{X}_{\mathbf{E}}^{0,R}$ tel que :*

$$\forall \mathbf{F} \in \mathbf{X}_{\mathbf{E}}^{0,R}, \mathcal{A}_0(\tilde{\mathbf{x}}_i^0, \mathbf{F}) = -\mathcal{A}_0(\mathbf{e}_i, \mathbf{F}). \quad (2.64)$$

DÉMONSTRATION. La preuve est similaire à la preuve de la proposition 2.45.

□

Proposition 2.47 *La λ -approche et la MCSO donnent tout calcul fait le même résultat :*

$$\boxed{\hat{\mathbf{E}}^0 + \sum_{i=1}^{N_{cr}} c_i \tilde{\mathbf{x}}_i^0 = \tilde{\mathbf{E}}^0.}$$

DÉMONSTRATION. L'addition des formulations variationnelles (2.57) et (2.64) donne :

$$\forall \mathbf{F} \in \mathbf{X}_{\mathbf{E}}^{0,R}, \mathcal{A}_0 \left(\hat{\mathbf{E}}^0 + \sum_{i=1}^{N_{cr}} c_i \tilde{\mathbf{x}}_i^0, \mathbf{F} \right) = \mathcal{L}_0(\mathbf{F}) - \mathcal{A}_0 \left(\mathbf{e} + \sum_{i=1}^{N_{cr}} c_i \mathbf{e}_i, \mathbf{F} \right).$$

Or : $(\mathbf{e} + \sum_i c_i \mathbf{e}_i) \cdot \boldsymbol{\tau}|_{\partial\omega} = e - \sum_i c_i \sum_j \beta^{i,j} \mathbf{x}_j^P \cdot \boldsymbol{\tau}|_{\partial\omega} = e - \sum_j \sum_i c_i \beta^{i,j} \mathbf{x}_j^P \cdot \boldsymbol{\tau}|_{\partial\omega}$. D'où en utilisant la relation (2.58), on obtient : $(\mathbf{e} + \sum_i c_i \mathbf{e}_i) \cdot \boldsymbol{\tau}|_{\partial\omega} = e - \sum_j \lambda^j \mathbf{x}_j^P \cdot \boldsymbol{\tau}|_{\partial\omega} = \mathbf{e}_\lambda \cdot \boldsymbol{\tau}|_{\partial\omega}$. Ainsi $\hat{\mathbf{E}}^0 + \sum_i c_i \tilde{\mathbf{x}}_i^0$, qui appartient à $\mathbf{X}_{\mathbf{E}}^{0,R}$ est solution de la formulation variationnelle (2.43). $\square$

Notons que cette décomposition permet un calcul optimal de $\mathbb{X}$, la matrice d'interaction entre les fonctions de base singulières. L'intérêt de cette approche est qu'elle permet de conserver l'orthogonalité lorsqu'on discrétise la formulation variationnelle.

2.4.5 Autres décompositions conformes

Dans [68], les auteurs proposent deux décompositions conformes de $\mathbf{X}_{\mathbf{E}}^0$:

- La première décomposition proposée est non-orthogonale, et consiste à écrire :

$$\mathbf{X}_{\mathbf{E}}^0 = \mathbf{X}_{\mathbf{E}}^{0,R} \oplus \mathbf{grad} S, \text{ où : } S = \text{vect}\{\eta_j \varphi_{D,j}^P\},$$

η_j étant une fonction de troncature (voir la définition 1.1) telle que $\eta_j \varphi_{D,j}^P \in H_0^1(\omega) \setminus H^2(\omega)$ et $\Delta(\eta_j \varphi_{D,j}^P) = 0$ dans un voisinage du coin rentrant O_j .

Le champ électrique $\mathbf{E}^0$ se décompose alors de la façon suivante :

$$\mathbf{E}^0 = \mathbf{E}^{0,R} + \sum_{i=1}^{N_{cr}} a_i \mathbf{grad}(\eta_j \varphi_{D,j}^P) \text{ avec } \mathbf{E}^{0,R} \in \mathbf{X}_{\mathbf{E}}^{0,R}.$$

On a $\mathbf{grad}(\eta_j \varphi_{D,j}^P) = \mathbf{x}_j^P$ au voisinage du coin rentrant O_j . Soit $\mathbf{a} \in \mathbb{R}^{N_{cr}}$ tel que : $\forall i, (\mathbf{a})_i = a_i$. Posons $\forall i, v_i = \Delta(\eta_i \varphi_{D,i})$.

Proposition 2.48 *Le problème (2.6)-(2.7) est équivalent à la formulation variationnelle suivante : Trouver $\mathbf{E}^{0,R} \in \mathbf{X}_{\mathbf{E}}^{0,R}$ et $\mathbf{a} \in \mathbb{R}^N$ tels que :*

$$\begin{aligned} \mathcal{A}_0(\mathbf{E}^{0,R}, \mathbf{F}) + \sum_{i=1}^{N_{cr}} a_i (v_i, \text{div } \mathbf{F})_0 &= \mathcal{L}_0(\mathbf{F}) - \mathcal{A}_0(\mathbf{e}, \mathbf{F}), \quad \forall \mathbf{F} \in \mathbf{X}_{\mathbf{E}}^{0,R}, \\ (\text{div } \mathbf{E}^{0,R}, v_j)_0 + \sum_{i=1}^{N_{cr}} a_i (v_i, v_j)_0 &= (g, v_j)_0 - (\text{div } \mathbf{e}, v_j)_0, \quad \forall j \in \{1, \dots, N_{cr}\}. \end{aligned}$$

DÉMONSTRATION. Reprendre la démonstration de la proposition 2.45, en utilisant (1.4). $\square$

Comme la décomposition est non orthogonale, notons qu'il existe un couplage entre parties régulière et singulière. Dans cette méthode, on a besoin de connaître seulement les $v_j = \Delta(\eta_j \varphi_{D,j})$ pour calculer la partie régulière du champ électrique $\mathbf{E}^{0,R}$ et les coefficients de singularité a_i . L'inconvénient de cette méthode est que lorsqu'on discrétise la formulation variationnelle, il faut approcher la fonction de troncature, qui varie entre zéro et un, et dont les dérivées secondes peuvent prendre des valeurs très élevées.

- La seconde décomposition proposée est orthogonale, et consiste à écrire : $\mathbf{X}_{\mathbf{E}}^{0,S}$ sous la forme : $\mathbf{X}_{\mathbf{E}}^{0,S} = \text{vect}\{\mathbf{x}_i^{HL}\}$, où $\forall i \in \{1, \dots, N_{cr}\}$, $\mathbf{x}_i^{HL}$ est tel que : $\mathbf{x}_i^{HL} = \tilde{\mathbf{x}}_i^{HL} + \mathbf{x}_i^P$, avec $\tilde{\mathbf{x}}_i^{HL} \in \mathbf{H}^1(\omega)$.

La partie régulière du champ électrique est bien sûr la même que pour la MCSO et le champ électrique $\mathbf{E}^0$ se décompose de la façon suivante :

$$\mathbf{E}^0 = \hat{\mathbf{E}}^0 + \sum_{i=1}^{N_{cr}} b_i \mathbf{x}_i^{HL}.$$

Soit $\mathbf{b} \in \mathbb{R}^{N_{cr}}$ tel que : $\forall i, \mathbf{b}^i = b_i$.

Proposition 2.49 *Le problème (2.6)-(2.7) est équivalent à la formulation variationnelle suivante : Trouver $\hat{\mathbf{E}}^0 \in \mathbf{X}_{\mathbf{E}}^{0,R}$ et $\mathbf{b} \in \mathbb{R}^N$ tels que :*

$$\begin{aligned} \mathcal{A}_0(\hat{\mathbf{E}}^0, \mathbf{F}) &= \mathcal{L}_0(\mathbf{F}) - \mathcal{A}_0(\mathbf{e}, \mathbf{F}), \forall \mathbf{F} \in \mathbf{X}_{\mathbf{E}}^{0,R}, \\ \mathbb{X}_{HL} \mathbf{b} &= \boldsymbol{\lambda}_{HL}, \end{aligned}$$

avec $\mathbb{X}_{HL} \in \mathbb{R}^{N_{cr} \times N_{cr}}$ telle que :

$$\forall i, j \in \{1, \dots, N_{cr}\}, \mathbb{X}_{HL}^{i,j} = (\mathbf{x}_i^{HL}, \mathbf{x}_j^{HL})_{\mathbf{X}_{\mathbf{E}}^0} / \pi,$$

et :

$$\forall i \in \{1, \dots, N_{cr}\}, \boldsymbol{\lambda}_{HL}^i = \lambda_{HL}^i := ((g^0, \operatorname{div} \mathbf{x}_i^{HL})_0 + (f^0, \operatorname{rot} \mathbf{x}_i^{HL})_0) / \pi.$$

Pour approcher les b_i , il faut auparavant calculer les $\tilde{\mathbf{x}}_i^{HL}$. Les auteurs résolvent le problème suivant : $\forall i$, trouver $\tilde{\mathbf{x}}_i^{HL} \in \mathbf{H}^1(\omega)$ tel que :

$$\begin{aligned} \mathcal{A}_0(\tilde{\mathbf{x}}_i^{HL}, \mathbf{F}) &= 0, \quad \forall \mathbf{F} \in \mathbf{X}_{\mathbf{E}}^{0,R}, \\ \tilde{\mathbf{x}}_i^{HL} \cdot \boldsymbol{\tau}|_{\partial\omega} &= -\mathbf{x}_i^P \cdot \boldsymbol{\tau}|_{\partial\omega}. \end{aligned}$$

Ainsi, numériquement, l'orthogonalité discrète est conservée. Notons que le calcul des coefficients de singularité c_i que nous proposons pour la MCSO est plus précis. En effet, nous utilisons le fait que $\operatorname{div} \mathbf{x}_i^S = s_{D,i}$ et $\operatorname{rot} \mathbf{x}_i^S = s_{N,j}$ pour calculer $\mathbb{X}$ et $\boldsymbol{\lambda}$: il n'est pas nécessaire de connaître les $\tilde{\mathbf{x}}_i^S$. Au contraire, dans [68], $\mathbb{X}_{HL}$ et $\boldsymbol{\lambda}_{HL}$ ne peuvent être obtenus qu'après avoir calculer les $\tilde{\mathbf{x}}_i^{HL}$, et les avoir dérivés. Pour améliorer le calcul de $\mathbb{X}_{HL}$ et $\boldsymbol{\lambda}_{HL}$, il faut donc connaître les $\operatorname{rot} \mathbf{x}_i^{HL}$ et $\operatorname{div} \mathbf{x}_i^{HL}$ directement. Soient $(\zeta^{i,j})_{i,j \in \{1, \dots, N_{cr}\}}$ les coefficients de la matrice inverse de $\mathbb{X}$: $\forall i, j$, $\zeta^{i,j} = (\mathbb{X}^{-1})^{i,j}$. On a la proposition suivante :

Proposition 2.50 $\forall i \in \{1, \dots, N_{cr}\}$, $\mathbf{x}_i^{HL} \in \mathbf{X}_{\mathbf{E}}^{0,S}$ satisfait le problème suivant :

$$\begin{aligned} \operatorname{div} \mathbf{x}_i^{HL} &= \sum_{j=1}^{N_{cr}} \zeta^{i,j} s_{D,j}, \text{ dans } \omega, \\ \operatorname{rot} \mathbf{x}_i^{HL} &= \sum_{j=1}^{N_{cr}} \zeta^{i,j} s_{N,j}, \text{ dans } \omega. \end{aligned} \tag{2.65}$$

DÉMONSTRATION. Soit $i \in \{1, \dots, N_{cr}\}$. Comme $\mathbf{x}_i^{HL} \in \mathbf{X}_{\mathbf{E}}^{0,S}$, on peut le décomposer dans la base des $\mathbf{x}_j^S$ de la façon suivante :

$$\mathbf{x}_i^{HL} := \tilde{\mathbf{x}}_i^{HL} + \mathbf{x}_i^P = \sum_{j=1}^{N_{cr}} \zeta^{i,j} \mathbf{x}_j^S, \text{ où } \forall j, \zeta^{i,j} \in \mathbb{R},$$

de sorte que $\mathbf{x}_i^{HL}$ satisfasse les équations (2.65). Déterminons les $\zeta^{i,j}$. En identifiant partie régulière et partie singulière, on obtient que :

$$\tilde{\mathbf{x}}_i^{HL} = \sum_{j=1}^{N_{cr}} \zeta^{i,j} \tilde{\mathbf{x}}_j^S, \text{ et } : \mathbf{x}_i^P = \sum_{j=1}^{N_{cr}} \zeta^{i,j} \sum_{k=1}^{N_{cr}} \beta^{j,k} \mathbf{x}_k^P = \sum_{k=1}^{N_{cr}} \left(\sum_{j=1}^{N_{cr}} \zeta^{i,j} \beta^{j,k} \right) \mathbf{x}_k^P.$$

Afin de satisfaire l'égalité sur les parties singulières, on doit avoir : $\forall k, \sum_{j=1}^{N_{cr}} \zeta^{i,j} \beta^{j,k} = \delta_{i,k}$.

Soit $\mathbb{X}_{HL} \in \mathbb{R}^{N_{cr} \times N_{cr}}$ la matrice des coefficients $\zeta^{i,j} : \forall i, j, \mathbb{X}_{HL}^{i,j} = \zeta^{i,j}$. Il s'agit alors de résoudre le système linéaire : $\mathbb{X}_{HL} \mathbb{X} = \mathbb{I}_{N_{cr}}$, où $\mathbb{I}_{N_{cr}}$ est la matrice identité de $\mathbb{R}^{N_{cr} \times N_{cr}}$. Comme $\mathbb{X}$ est inversible, il existe une unique solution telle que : $\mathbb{X}_{HL} = \mathbb{X}^{-1}$. □

Dans [63], E. Garcia propose les deux décompositions conformes non-orthogonales.

Posons $\forall i, \mathbf{x}_i^N = \mathbf{rot} \varphi_{N,i}$ et $\mathbf{x}_i^D = -\mathbf{grad} \varphi_{D,i}$.

$\forall i, \mathbf{x}_i^D \in \mathbf{X}_{\mathbf{E}}^0$ est tel que : $\text{div} \mathbf{x}_i^D = s_{D,i}$, et $\text{rot} \mathbf{x}_i^D = 0$; et $\mathbf{x}_i^N \in \mathbf{X}_{\mathbf{E}}^0$ est tel que : $\text{div} \mathbf{x}_i^N = 0$, et $\text{rot} \mathbf{x}_i^N = s_{N,i}$.

• La première décomposition est la suivante : $\mathbf{X}_{\mathbf{E}}^0 = \mathbf{X}_{\mathbf{E}}^{0,R} \oplus \text{vect}\{\mathbf{x}_i^D\}$. Le champ électrique $\mathbf{E}^0$ s'écrit alors ainsi :

$$\mathbf{E}^0 = \mathbf{E}^{0,D} + \sum_{i=1}^{N_{cr}} a_i^D \mathbf{x}_i^D \text{ avec } \mathbf{E}^{0,D} \in \mathbf{X}_{\mathbf{E}}^{0,R}.$$

On a alors la proposition suivante :

Proposition 2.51 *Le problème (2.6)-(2.7) est équivalent à la formulation variationnelle suivante : Trouver $\mathbf{E}^{0,D} \in \mathbf{X}_{\mathbf{E}}^{0,R}$ et $(a_i^D)_{i=1,\dots,N_{cr}} \in \mathbb{R}^{N_{cr}}$ tels que :*

$$\begin{aligned} \mathcal{A}_0(\mathbf{E}^{0,D}, \mathbf{F}) + \sum_{i=1}^{N_{cr}} a_i^D (s_{D,i}, \text{div} \mathbf{F})_0 &= \mathcal{L}_0(\mathbf{F}) - \mathcal{A}_0(\mathbf{e}, \mathbf{F}), \quad \forall \mathbf{F} \in \mathbf{X}_{\mathbf{E}}^{0,R}, \\ (\text{div} \mathbf{E}^{0,D}, s_{D,j})_0 + \sum_{i=1}^{N_{cr}} a_i^D \beta_D^{i,j} &= \lambda_D^j, \quad \forall j \in \{1, \dots, N_{cr}\}. \end{aligned}$$

La démonstration est similaire à celle de la proposition 2.4.5.

• La seconde décomposition est la suivante : $\mathbf{X}_{\mathbf{E}}^0 = \mathbf{X}_{\mathbf{E}}^{0,R} \oplus \text{vect}\{\mathbf{x}_i^N\}$. Le champ électrique $\mathbf{E}^0$ s'écrit alors ainsi

$$\mathbf{E}^0 = \mathbf{E}^{0,N} + \sum_{i=1}^{N_{cr}} a_i^N \mathbf{x}_i^N \text{ avec } \mathbf{E}^{0,N} \in \mathbf{X}_{\mathbf{E}}^{0,R}.$$

On a alors la proposition suivante :

Proposition 2.52 *Le problème (2.6)-(2.7) est équivalent à la formulation variationnelle suivante : Trouver $\mathbf{E}^{0,N} \in \mathbf{X}_{\mathbf{E}}^{0,R}$ et $(a_i^N)_{i=1,\dots,N_{cr}} \in \mathbb{R}^{N_{cr}}$ tels que :*

$$\begin{aligned} \mathcal{A}_0(\mathbf{E}^{0,N}, \mathbf{F}) + \sum_{i=1}^{N_{cr}} a_i^N (s_{N,i}, \text{rot} \mathbf{F})_0 &= \mathcal{L}_0(\mathbf{F}) - \mathcal{A}_0(\mathbf{e}, \mathbf{F}), \quad \forall \mathbf{F} \in \mathbf{X}_{\mathbf{E}}^{0,R}, \\ (\text{rot} \mathbf{E}^{0,N}, s_{N,j})_0 + \sum_{i=1}^{N_{cr}} a_i^N \beta_N^{i,j} &= \lambda_N^j, \quad \forall j \in \{1, \dots, N_{cr}\}. \end{aligned}$$

La démonstration est similaire à celle de la proposition 2.4.5.

Après avoir obtenu $\mathbf{E}^{0,D}$ (resp. $\mathbf{E}^{0,N}$), la calcul des parties régulières des $\mathbf{x}_i^D$ (resp. $\mathbf{x}_i^N$) se fait numériquement par dérivation des potentiels associés $\tilde{\varphi}_{D,i}$ (resp. $\tilde{\varphi}_{N,i}$). Les fonctions de base de $\Phi_{D,N}$ sont obtenues par la méthode *Dirichlet-Neumann*. On améliore l'approximation obtenue en utilisant une projection L^2 , ce qui est moins précis que l'approche que nous proposons dans la proposition 2.46.

2.4.6 Régularité dans les espaces singuliers conformes

Étudions la régularité de $\hat{\mathbf{E}}^0 + \sum_{i=1}^{N_{cr}} c_i \tilde{\mathbf{x}}_i^0 = \tilde{\mathbf{E}}^0$. Le théorème suivant se déduit directement du théorème 2.31 :

Théorème 2.53 *Supposons que pour $\epsilon > 0$, tel que $2\alpha - 1 - \epsilon > 0$, $f, g \in H^{2\alpha-1-\epsilon}(\omega)$, et $e \in H^{2\alpha-1/2-\epsilon}(\omega)$. Alors $\hat{\mathbf{E}}^0 + \sum_{i=1}^{N_{cr}} c_i \tilde{\mathbf{x}}_i^0 \in \mathbf{H}^{2\alpha-\epsilon}(\omega)$.*

Cependant $\hat{\mathbf{E}}^0$ et les $\tilde{\mathbf{x}}_i^0$ ne sont pas toujours individuellement dans $\mathbf{H}^{2\alpha-\epsilon}(\omega)$. En effet :

Théorème 2.54 *Soit $i \in \{1, \dots, N_{cr}\}$. Lorsque $\alpha \leq 2/3$, $\tilde{\mathbf{x}}_i^0 \in \mathbf{H}^{2\alpha-\epsilon}(\omega)$, mais lorsque $\alpha > 2/3$, $\tilde{\mathbf{x}}_i^0 \notin \mathbf{H}^{2\alpha-\epsilon}(\omega)$.*

DÉMONSTRATION. D'après [67] (p. 7, thm. 1.2.18), les $s_{D,i}$ et $s_{N,i}$ appartiennent à $H^{1-\alpha-\epsilon}(\omega)$, pour tout $\epsilon > 0$ tel que $1 - \alpha - \epsilon > 0$. Or, on a l'inclusion $H^{1-\alpha-\epsilon}(\omega) \subset H^{2\alpha-1-\epsilon}(\omega)$ si et seulement si :

$$1 - \alpha - \epsilon \geq 2\alpha - 1 - \epsilon \iff \alpha \leq 2/3.$$

□

Ainsi, lorsque les coins rentrants sont d'angles plus petits que $3\pi/2$, les $\tilde{\mathbf{x}}_i^0$ et $\hat{\mathbf{E}}^0$ ne sont pas individuellement dans $\mathbf{H}^{2\alpha-\epsilon}(\omega)$. Cependant, il y a compensation des parties moins régulières des $\tilde{\mathbf{x}}_i^0$ et de $\hat{\mathbf{E}}^0$. Nous verrons que numériquement, dans le cas statique, la méthode du complément singulier orthogonale donne les mêmes résultats que la λ -approche, quelques soient les angles des coins rentrants.

Est-il possible de construire une base singulière conforme dans $\mathbf{X}_{\mathbf{E}}^0$ qui soit dans $\mathbf{H}^{2\alpha-\epsilon}(\omega)$? Considérons le cas d'un unique coin rentrant. Soient $\mathbf{w}^S$ et $\mathbf{v}^S$ dans $\mathbf{X}_{\mathbf{E}}^0$ tels que :

$$\begin{cases} \operatorname{div} \mathbf{w}^S = s_D \text{ dans } \omega, \\ \operatorname{rot} \mathbf{w}^S = 0 \text{ dans } \omega, \end{cases} \quad \text{et} \quad \begin{cases} \operatorname{div} \mathbf{v}^S = 0 \text{ dans } \omega, \\ \operatorname{rot} \mathbf{v}^S = s_N \text{ dans } \omega, \end{cases} \quad (2.66)$$

$\mathbf{w}^S$ engendre $\mathbf{W}_{\mathbf{E}}^S$ et $\mathbf{v}^S$ engendre $\mathbf{V}_{\mathbf{E}}^S$. D'après F. Assous et al. [9], on a localement :

$$\mathbf{v}^S = \sum_{n \geq 1} B_n r^{n\alpha-1} \begin{pmatrix} \sin(n\alpha\theta) \\ \cos(n\alpha\theta) \end{pmatrix} + \sum_{n \geq -1} A_n r^{n\alpha+1} \begin{pmatrix} \frac{n\alpha}{4n\alpha+4} \sin(n\alpha\theta) \\ \frac{n\alpha+2}{4n\alpha+4} \cos(n\alpha\theta) \end{pmatrix}.$$

La somme sur les $(B_n)_{n \geq 1}$ est une solution homogène de (2.66), et par construction, les $(A_n)_{n \geq -1}$ sont localement tels que : $s_N = \sum_{n \geq -1} A_n r^{n\alpha} \cos(n\alpha\theta)$, avec $A_{-1} = 1$.

Comme $\mathbf{v}^S = \tilde{\mathbf{v}} + \beta_N \mathbf{x}^P$, avec $\tilde{\mathbf{v}} \in \mathbf{H}^1(\omega)$, on a : $B_1 = -\alpha\beta_N$. Décomposons $\tilde{\mathbf{v}}$ ainsi :

$$\tilde{\mathbf{v}} = \tilde{\mathbf{v}}_+ + \tilde{\mathbf{v}}_- \text{ avec : } \tilde{\mathbf{v}}_- = r^{1-\alpha} \begin{pmatrix} \frac{\alpha}{-4\alpha+4} \sin(\alpha\theta) \\ \frac{-\alpha+2}{-4\alpha+4} \cos(\alpha\theta) \end{pmatrix}.$$

Le terme le plus singulier de $\tilde{\mathbf{v}}_+$ est en $r^{2\alpha-1}$, ainsi on a toujours $\tilde{\mathbf{v}}_+ \in \mathbf{H}^{2\alpha-\epsilon}(\omega)$. En revanche, lorsque $\alpha > 2/3$, $\tilde{\mathbf{v}}_- \notin \mathbf{H}^{2\alpha-\epsilon}(\omega)$.

Pour déterminer $\mathbf{w}^S$, on part de la formule de représentation $s_D = \sum_{n \geq -1} \bar{A}_n r^{n\alpha} \sin(n\alpha\theta)$, avec $\bar{A}_{-1} = -1$. On sait que : $\mathbf{w}^S = \tilde{\mathbf{w}} + \beta_D \mathbf{x}^P$, avec $\tilde{\mathbf{w}} \in \mathbf{H}^1(\omega)$. Une solution particulière du second

système d'équations de (2.66) est : $\sum_{n \geq -1} \bar{A}_n r^{n\alpha+1} \begin{pmatrix} \frac{n\alpha+2}{4n\alpha+4} \sin(n\alpha\theta) \\ \frac{n\alpha}{4n\alpha+4} \cos(n\alpha\theta) \end{pmatrix}.$

Ainsi, localement, on a :

$$\mathbf{w}^S = \sum_{n \geq 1} \bar{B}_n r^{n\alpha-1} \begin{pmatrix} \sin(n\alpha\theta) \\ \cos(n\alpha\theta) \end{pmatrix} + \sum_{n \geq -1} \bar{A}_n r^{n\alpha+1} \begin{pmatrix} \frac{n\alpha+2}{4n\alpha+4} \sin(n\alpha\theta) \\ \frac{n\alpha}{4n\alpha+4} \cos(n\alpha\theta) \end{pmatrix},$$

avec $\bar{B}_1 = -\alpha\beta_D$. D'où :

$$\tilde{\mathbf{w}} = \tilde{\mathbf{w}}_+ + \tilde{\mathbf{w}}_- \text{ avec : } \tilde{\mathbf{w}}_- = -r^{1-\alpha} \begin{pmatrix} \frac{-\alpha+2}{-4\alpha+4} \sin(\alpha\theta) \\ \frac{\alpha}{-4\alpha+4} \cos(\alpha\theta) \end{pmatrix}.$$

Le terme le plus singulier de $\tilde{\mathbf{w}}_+$ est en $r^{2\alpha-1}$, ainsi $\tilde{\mathbf{w}}_+ \in \mathbf{H}^{2\alpha-\epsilon}(\omega) \forall \epsilon > 0$, et pour $\alpha > 2/3$, $\tilde{\mathbf{w}}_- \notin \mathbf{H}^{2\alpha-\epsilon}(\omega)$. D'après [67] (thm. 1.2.18), on a : $\tilde{\mathbf{w}} \in \mathbf{H}^{2-\alpha-\epsilon}(\omega)$, d'où :

$$\begin{aligned} \forall \alpha \leq 2/3, \quad \forall \epsilon > 0 \mid 2\alpha - \epsilon - 1 > 0, \quad \tilde{\mathbf{v}}, \tilde{\mathbf{w}} &\in \mathbf{H}^{2\alpha-\epsilon}(\omega), \\ \forall \alpha > 2/3, \quad \forall \epsilon > 0 \mid 1 - \alpha - \epsilon > 0, \quad \tilde{\mathbf{v}}, \tilde{\mathbf{w}} &\in \mathbf{H}^{2-\alpha-\epsilon}(\omega). \end{aligned}$$

Soit $\mathbf{x} \in \mathbf{X}_{\mathbf{E}}^0 \setminus \mathbf{X}_{\mathbf{E}}^{0,R}$, tel que : $\mathbf{x} = \mathbf{w}^S + a_\alpha \mathbf{v}^S$, $a_\alpha \in \mathbb{R}$. On souhaite déterminer a_α de sorte que $\mathbf{x} \in \mathbf{H}^{2\alpha-\epsilon}(\omega)$. Pour cela, il faut éliminer le terme en $r^{1-\alpha}$, c'est-à-dire choisir a_α tel que :

$$\begin{cases} \frac{\alpha}{-4\alpha+4} - a_\alpha \frac{-\alpha+2}{-4\alpha+4} = 0, \\ \frac{-\alpha+2}{-4\alpha+4} - a_\alpha \frac{\alpha}{-4\alpha+4} = 0. \end{cases}$$

Le système est compatible si et seulement si $\alpha = 1$, ce qui contredit $\alpha < 1$ qui est l'hypothèse de départ pour le cas d'un coin rentrant. Ainsi, lorsque $\alpha > 2/3$, il est impossible de créer une base de $\mathbf{X}_{\mathbf{E}}^0 \setminus \mathbf{X}_{\mathbf{E}}^{0,R}$ dont la partie régulière soit dans $\mathbf{H}^{2\alpha-\epsilon}(\omega)$. On a le corollaire immédiat suivant :

Corollaire 2.55 *Soit $i \in \{1, \dots, N_{cr}\}$. Alors $\tilde{\mathbf{x}}_i \in \mathbf{H}^{2-\alpha-\epsilon}(\Omega)$.*

2.4.7 Conclusions sur la méthode du complément singulier

La λ -approche est une nouvelle méthode de décomposition de $\mathbf{X}_{\mathbf{E}}^0$ non-conforme, et dont l'originalité par rapport au travail de C. Hazard et S. Lohrengel [68] repose sur le fait qu'on ne requiert pas de fonction de troncature. Celle-ci est remplacée par une condition aux limites non-homogène. L'intérêt de la λ -approche pour la MCSO, qui est incontournable si on veut traiter le cas instationnaire, est qu'elle permet de calculer une meilleure approximation des $\mathbf{x}_i^S$ que l'approximation par dérivation des potentiels.

Pour les problèmes quasi-statiques, la λ -approche et la MCSO sont identiques. Elles donnent des résultats numériques compétitifs en terme de précision (voir le chapitre 6), et se généralisent au cas de pointes coniques à base régulière [63]. Dans les domaines tridimensionnels généraux, les singularités électromagnétiques sont difficilement discrétisables, les singularités de coin et d'arête étant liées entre elles [52]. Néanmoins, dans le cas d'un ouvert prismatique [42], ainsi que dans le cas d'un domaine axisymétrique [76], on peut adapter ces méthodes en utilisant une décomposition en série de Fourier (voir la section 8.6).

Dans le cas dépendant du temps, on n'a plus les mêmes résultats de régularité pour la MCSO, car on ne sait pas si les parties de régularité $\mathbf{H}^{1-\alpha-\epsilon}(\omega)$ de $\tilde{\mathbf{E}}^0$ et de $\sum_i c_i \tilde{\mathbf{x}}_i^S$ se compensent. Cependant, pour améliorer la précision, on peut faire les calculs en exprimant explicitement les termes en $r^{1-\alpha}$.

2.5 Champ électrique 2D : la régularisation à poids

Dans cette section, nous présentons essentiellement les résultats obtenus par M. Costabel et M. Dauge dans [53] dans le cas harmonique. Nous avons vu que lorsque ω est non convexe, l'espace $\mathbf{X}_{\mathbf{E}}^0 \cap \mathbf{H}^1(\omega)$ est fermé et strictement inclus dans $\mathbf{X}_{\mathbf{E}}^0$, ce qui implique que la solution du problème discrétisé par les éléments finis de Lagrange P_k dans l'espace $\mathbf{X}_{\mathbf{E}}^0$ ne converge pas vers la bonne solution. D'autre part, dans l'espace $\mathbf{X}_{\mathbf{E}}$, la condition aux limites tangentielle est naturelle, mais pas essentielle, ce qui limite en pratique la vitesse de convergence de la solution discrétisée. L'idée est de déterminer un espace intermédiaire $\mathbf{X}$, dans lequel la condition aux limites tangentielle est essentielle, et tel que $\mathbf{X} \cap \mathbf{H}^1(\omega)$ soit dense dans $\mathbf{X}$ (afin d'obtenir la convergence des éléments finis de Galerkin dans $\mathbf{X}$). On s'intéresse aux espaces $\mathbf{X}$ de la forme : $\mathbf{X} = \{ \mathbf{u} \in \mathbf{H}_0(\text{rot}, \omega) : \text{div } \mathbf{u} \in V \}$. La forme bilinéaire associée à cet espace pour la résolution de (2.6)-(2.7) est la suivante :

$$\begin{aligned} \mathcal{A}_{\mathbf{X}} : \mathbf{X} \times \mathbf{X} &\rightarrow \mathbb{R} \\ (\mathbf{E}, \mathbf{F}) &\mapsto (\text{div } \mathbf{E}, \text{div } \mathbf{F})_V + (\text{rot } \mathbf{E}, \text{rot } \mathbf{F})_0, \end{aligned}$$

où $(\cdot, \cdot)_V$ est le produit scalaire dans V .

La formulation variationnelle de (2.1)-(2.3) dans $\mathbf{X}$ s'écrit (FVX) : Trouver $\mathbf{E}^0 \in \mathbf{X}$ tel que :

$$\forall \mathbf{F} \in \mathbf{X}, \mathcal{A}_{\mathbf{X}}(\mathbf{E}^0, \mathbf{F}) = (g^0, \text{div } \mathbf{F})_V + (f^0, \text{rot } \mathbf{F})_0.$$

Nous voulons déterminer les espaces V tels que :

- il existe une solution unique à (FVX), qui soit équivalente à (2.1)-(2.3) ;
- $\mathbf{X} \cap \mathbf{H}^1(\omega)$ soit dense dans $\mathbf{X}$.

2.5.1 Condition pour obtenir la coercivité

Pour appliquer le théorème de Lax-Milgram à (FVX), il faut rendre $\mathcal{A}_{\mathbf{X}}$ coercitive sur $\mathbf{X}$. D'une part, V doit être un espace de type L^2 . D'autre part, comme les éléments de $\mathbf{H}_0(\text{rot}, \omega)$ sont à

divergence dans $H^{-1}(\omega)$, on cherche $V \subset H^{-1}(\omega)$. Pour des raisons pratiques, on limite le choix aux espaces à poids de la forme : $V = \{u \in L^2_{\text{loc}}(\omega) : wu \in L^2(\omega)\}$, où $w \in C^\infty(\omega, \mathbb{R}^+)$ est borné dans $\bar{\omega}$. A priori V contient $L^2(\omega) : L^2(\omega) \subset V \subset H^{-1}(\omega)$. Le produit scalaire dans V s'écrit :

$$(u, v)_0 = \int_{\omega} w^2 uv \, d\omega, \text{ et la norme correspondante est : } \|u\|_V = \left(\int_{\omega} w^2 u^2 \, d\omega \right)^{1/2}.$$

Proposition 2.56 *Lorsque l'injection de $\mathbf{X}$ dans $\mathbf{L}^2(\omega)$ est compacte, la norme du graphe dans $\mathbf{X}$ est équivalente à la semi-norme. La norme dans $\mathbf{X}$ est alors définie ainsi :*

$$\|\mathbf{u}\|_{\mathbf{X}} = (\|\text{rot } \mathbf{u}\|_0^2 + \|\text{div } \mathbf{u}\|_V^2)^{1/2}. \quad (2.67)$$

DÉMONSTRATION. On procède par l'absurde. Soit $(\mathbf{u}_n)_n$ une suite de $\mathbf{X}$ telle que : $\|\mathbf{u}_n\|_{\text{rot, div}} \rightarrow 0$ et $\|\mathbf{u}_n\|_0 = 1$, c'est-à-dire :

$$\begin{aligned} \text{div } \mathbf{u}_n &\rightarrow 0 \text{ dans } V, \\ \text{rot } \mathbf{u}_n &\rightarrow 0 \text{ dans } L^2. \end{aligned}$$

La suite $(\mathbf{u}_n)_n$ étant bornée dans $\mathbf{X}$, et l'injection de $\mathbf{X}$ dans $\mathbf{L}^2(\omega)$ étant compacte, il existe une sous-suite, encore notée $(\mathbf{u}_n)_n$ de $\mathbf{X}$ qui converge dans $\mathbf{L}^2(\omega)$. Soit $\mathbf{u}$ sa limite. On a la convergence forte de $(\mathbf{u}_n)_n$ vers $\mathbf{u}$ dans $\mathbf{L}^2(\omega)$ à cause de l'injection compacte. On en déduit que $\|\mathbf{u}\|_0 = 1$. En utilisant les formules d'intégration par parties classiques, les conditions aux limites homogènes et par passage à la limite et unicité de celle-ci, on a :

$$\begin{aligned} \langle \text{div } \mathbf{u}, \phi \rangle_{\mathcal{D}', \mathcal{D}} &= - \int_{\omega} \mathbf{u} \cdot \mathbf{grad } \phi \, d\omega = - \lim_{n \rightarrow +\infty} \int_{\omega} \mathbf{u}_n \cdot \mathbf{grad } \phi \, d\omega \\ &= - \lim_{n \rightarrow +\infty} \int_{\omega} \text{div } \mathbf{u}_n \phi \, d\omega = 0, \forall \phi \in \mathcal{D}(\omega), \\ \langle \text{rot } \mathbf{u}, \phi \rangle_{\mathcal{D}', \mathcal{D}} &= \int_{\omega} \mathbf{u} \cdot \mathbf{rot } \phi \, d\omega = \lim_{n \rightarrow +\infty} \int_{\omega} \mathbf{u}_n \cdot \mathbf{rot } \phi \, d\omega \\ &= - \lim_{n \rightarrow +\infty} \int_{\omega} \mathbf{rot } \mathbf{u}_n \phi \, d\omega = 0, \forall \phi \in \mathcal{D}(\omega). \end{aligned}$$

Ainsi, $\text{rot } \mathbf{u} = 0$ et $\text{div } \mathbf{u} = 0$ dans $\mathbf{X}$. D'après [65], la trace tangentielle est continue dans $\mathbf{H}(\text{rot}, \omega)$. Comme $\mathbf{u}_n \cdot \boldsymbol{\tau}|_{\partial\Omega} = 0$, on en déduit que : $\mathbf{u} \cdot \boldsymbol{\tau}|_{\partial\Omega} = 0$. Ainsi, $\mathbf{u} \in \mathbf{H}_0(\text{rot}^0, \omega)$, et donc d'après la proposition 1.3, il existe un unique $\phi \in H_0^1(\omega)$ tel que $\mathbf{grad } \phi = \mathbf{u}$. Comme $\Delta\phi = \text{div } \mathbf{u} = 0$, alors $\phi = 0$, et donc $\|\mathbf{u}\|_0 = 0$, ce qui contredit l'hypothèse de départ. $\square$

En pratique, cela permet d'obtenir la coercivité de la forme bilinéaire $\mathcal{A}_{\mathbf{X}}$, qui est alors le produit scalaire dans $\mathbf{X}$. Déterminons V tel que l'injection de $\mathbf{X}$ dans $\mathbf{L}^2(\omega)$ soit compacte. Pour cela, considérons la décomposition de Helmholtz suivante :

Lemme 2.57 *On peut décomposer $\mathbf{X}$ sous la forme : $\mathbf{X} = \mathbf{V}_{\mathbf{E}} \oplus^{\perp} \mathbf{W}$, où :*

$$\mathbf{W} = \{ \mathbf{u} := \mathbf{grad } u, u \in \Phi_V \} \text{ avec } \Phi_V = \{ u \in H_0^1(\omega) : \Delta u \in V \},$$

ce qui s'écrit aussi : $\mathbf{W} = \{ \mathbf{u} \in \mathbf{H}_0(\text{rot}^0, \omega) : \text{div } \mathbf{u} \in V \}$.

DÉMONSTRATION. Cette décomposition correspond à la décomposition de Helmholtz (2.17) de $\mathbf{X}_{\mathbf{E}}^0$: considérons $\mathbf{F} \in \mathbf{X}$. Alors $\mathbf{F} = \mathbf{F}_N + \mathbf{F}_D$, où :

- $\mathbf{F}_N = \mathbf{rot } \phi_N$, ϕ_N satisfaisant (2.16), avec la donnée $\text{rot } \mathbf{F} \in L_0^2(\omega)$;
- $\mathbf{F}_D = -\mathbf{grad } \phi_D$, ϕ_D satisfaisant : Trouver $\phi_D \in H_0^1(\omega)$ tel que : $-\Delta\phi_D = \text{div } \mathbf{F}$ dans V .

$\square$

L'injection de $\mathbf{V}_{\mathbf{E}}$ dans $\mathbf{L}^2(\omega)$ est compacte [8]. Comment obtenir l'injection compacte de $\mathbf{W}$ dans $\mathbf{L}^2(\omega)$? Montrons les propriétés suivantes :

Proposition 2.58 *Si l'injection de V dans $H^{-1}(\omega)$ est compacte, alors l'injection de $\mathbf{W}$ dans $\mathbf{L}^2(\omega)$ l'est aussi.*

DÉMONSTRATION. Comme $V \subset\subset H^{-1}(\omega)$, alors $\Phi_V \subset\subset H^1(\omega)$, d'où $\mathbf{W} \subset\subset \mathbf{L}^2(\omega)$. □

Des propositions 2.56 et 2.58, on déduit le théorème suivant :

Théorème 2.59 *Si l'injection de V dans $H^{-1}(\omega)$ est compacte, alors (2.6)-(2.7) est équivalent à (FVX).*

DÉMONSTRATION. D'après la proposition 2.58, l'injection compacte de V dans $H^{-1}(\omega)$ permet d'obtenir l'injection compacte de $\mathbf{W}$ dans $\mathbf{L}^2(\omega)$, et donc d'après la proposition 2.56, $\mathcal{A}_{\mathbf{X}}$ est coercitive. L'équivalence entre le système d'équations (2.6)-(2.7) et la formulation variationnelle (FVX) se montre de la même façon que la proposition 2.23 et nécessite la décomposition du lemme 2.57. □

En conclusion, afin d'obtenir l'équivalence entre la norme du graphe et la semi-norme (2.67) dans $\mathbf{X}$, et donc la coercivité de $\mathcal{A}_{\mathbf{X}}$, on cherche V tel que l'injection de V dans $H^{-1}(\omega)$ soit compacte. Nous allons maintenant voir sous quelles conditions sur V l'espace $\mathbf{X} \cap \mathbf{H}^1(\omega)$ est dense dans $\mathbf{X}$.

2.5.2 Condition pour obtenir la densité des éléments finis

Nous avons vu qu'il existe une décomposition de Helmholtz de $\mathbf{X}$. On peut aussi décomposer les éléments de $\mathbf{X}$ en une partie $\mathbf{H}^1$ -régulière et une partie qui n'est pas $\mathbf{H}^1$.

Théorème 2.60 *On a la décomposition suivante :*

$$\mathbf{X} = \mathbf{X}_{\mathbf{E}}^{0,R} \oplus \mathbf{W} := \mathbf{X}_{\mathbf{E}}^{0,R} \oplus \mathbf{grad} \Phi_V.$$

DÉMONSTRATION. Soit $\mathbf{u} \in \mathbf{X}$. D'après le lemme 2.57, il existe $\phi_{\mathbf{u}} \in \Phi_V$ et $\mathbf{v} \in \mathbf{V}_{\mathbf{E}}$ tels que : $\mathbf{u} = \mathbf{v} + \mathbf{grad} \phi_{\mathbf{u}}$. D'après [18], comme $\mathbf{v} \in \mathbf{V}_{\mathbf{E}}$, il existe $\mathbf{u}_0 \in \mathbf{X}_{\mathbf{E}}^{0,R}$ et $\phi_{\mathbf{v}} \in \Phi_D$ tels que : $\mathbf{v} = \mathbf{u}_0 + \mathbf{grad} \phi_{\mathbf{v}}$. Or $\Phi_D \subset \Phi_V$. En posant $\phi = \phi_{\mathbf{u}} + \phi_{\mathbf{v}}$, on a donc :

$$\mathbf{u} = \mathbf{u}_0 + \mathbf{grad} \phi, \text{ avec } \mathbf{u}_0 \in \mathbf{X}_{\mathbf{E}}^{0,R} \text{ et } \phi \in \Phi_V \text{ soit : } \mathbf{grad} \phi \in \mathbf{W}.$$

□

Pour avoir la densité des fonctions régulières dans $\mathbf{X}$, on cherche à avoir $\mathbf{X}_{\mathbf{E}}^{0,R}$ dense dans $\mathbf{X}$. Pour cela, d'après la décomposition du théorème 2.60, il nous faut la densité de $\mathbf{X}_{\mathbf{E}}^{0,R}$ dans $\mathbf{W}$. Ceci est possible si $\Phi_V \cap H^2(\omega)$ est dense dans Φ_V . Notons que : $\Phi_V \cap H^2(\omega) = \Phi_D^R := \{\phi \in H^2(\omega) \mid \phi|_{\partial\omega} = 0\}$. Plus précisément, on a le théorème fondamental suivant qui est un corollaire de la décomposition du théorème 2.60 :

Théorème 2.61

- 1) Si Φ_D^R est dense dans Φ_V pour la norme du graphe, alors $\mathbf{X}_{\mathbf{E}}^{0,R}$ est dense dans $\mathbf{X}$.
- 2) Φ_D^R est fermé dans Φ_V si et seulement si $\mathbf{X}_{\mathbf{E}}^{0,R}$ est fermé dans $\mathbf{X}$.
- 3) $\Phi_V \subset \Phi_D^R$ si et seulement si $\mathbf{X} = \mathbf{X}_{\mathbf{E}}^{0,R}$.

Dans le second cas, bien sûr si Φ_D^R est différent de Φ_V , on ne peut pas approcher la solution de (FVX) par les éléments finis, car on ne capte pas la partie singulière du champ. Le dernier cas correspond à un domaine ω convexe, car $\Phi_V \cap H^2(\omega) = \Phi_V$, il n'y a pas de singularité. Nous sommes intéressés par le premier cas. On se limite aux poids de la forme $w = r^\gamma$, où $\gamma > 0$ et l'on pose : $V = V_\gamma^0$, défini par (1.8). V_γ^0 peut se définir aussi de la façon suivante :

$$V_\gamma^0 = \{ u \in \mathcal{D}'(\omega) : u \in L^2(\omega \setminus \omega_c) \text{ et } r^\gamma u \in L^2(\omega_c) \},$$

où ω_c est un voisinage donné du coin rentrant inclus dans ω . On écrira le produit scalaire dans V_γ^0 ainsi :

$$\forall u, v \in V_\gamma^0, (v, u)_{0,\gamma} = \int_\omega r^{2\gamma} u v \, d\omega \quad (2.68)$$

Lorsque $\gamma > 1$, la condition $V_\gamma^0 \subset H^{-1}(\omega)$ n'est plus respectée. La valeur limite de γ est donc 1.

Soit $\mathbf{X}_{\mathbf{E},\gamma}^0$ l'espace vectoriel suivant : $\mathbf{X}_{\mathbf{E},\gamma}^0 = \{ u \in \mathbf{H}_0(\text{rot}, \omega) : \text{div } u \in V_\gamma^0 \}$.

Considérons $\Phi_\gamma = \{ \phi \in H_0^1(\omega) : \Delta \phi \in V_\gamma^0 \}$. D'après le théorème 2.60, $\mathbf{X}_{\mathbf{E},\gamma}^0 = \mathbf{X}_{\mathbf{E}}^{0,R} \oplus \mathbf{W}_\gamma$, où $\mathbf{W}_\gamma = \mathbf{grad } \Phi_\gamma$.

Posons $\mathbf{X}_{\mathbf{E},\gamma} = \{ \mathbf{u} \in \mathbf{H}(\text{rot}, \omega) : \text{div } \mathbf{u} \in V_\gamma^0, u_\tau \in L^2(\partial\omega) \}$. On appellera $\mathcal{A}_\gamma$ la forme bilinéaire suivante :

$$\mathcal{A}_\gamma : \mathbf{X}_{\mathbf{E},\gamma} \times \mathbf{X}_{\mathbf{E},\gamma} \rightarrow \mathbb{R}$$

$$(\mathbf{E}, \mathbf{F}) \mapsto (\text{div } \mathbf{E}, \text{div } \mathbf{F})_{0,\gamma} + (\text{rot } \mathbf{E}, \text{rot } \mathbf{F})_0.$$

$\mathcal{A}_\gamma$ restreinte à $\mathbf{X}_{\mathbf{E},\gamma}^0 \times \mathbf{X}_{\mathbf{E},\gamma}^0$ est coercitive, à condition qu'on obtienne l'équivalence entre la norme du graphe et la semi-norme dans $\mathbf{X}_{\mathbf{E},\gamma}^0$. $\mathcal{A}_\gamma$ définit alors le produit scalaire dans $\mathbf{X}_{\mathbf{E},\gamma}^0$.

2.5.3 Choix du poids

Proposition 2.62 *L'injection de V_γ^0 dans $H^{-1}(\omega)$ est compacte si et seulement si $\gamma < 1$.*

D'après le théorème 2.59, cette première condition sur γ entraîne la coercivité de $\mathcal{A}_\gamma$. Notons que lorsque $\gamma = 1$, $\mathbf{X}_{\mathbf{E},\gamma}^0 \simeq \mathbf{H}_0(\text{rot}, \omega)$. Or l'injection de $\mathbf{H}_0(\text{rot}, \omega)$ dans $\mathbf{L}^2(\omega)$ n'est pas compacte. Concernant la condition sur γ pour avoir la densité de Φ_D^R dans Φ_γ , intéressons-nous au résultat suivant :

Proposition 2.63 *Quel que soit γ , les fonctions $C^\infty(\bar{\omega})$ à trace nulle sur $\partial\omega$ sont denses dans $V_\gamma^2 \cap H_0^1(\omega)$.*

Le théorème fondamental suivant nous donne les valeurs de γ telles que $\Phi_\gamma \subset V_\gamma^2$:

Théorème 2.64 *Pour tout γ tel que : $1 - \alpha < \gamma \leq 1$, l'opérateur Δ est un isomorphisme de $V_\gamma^2 \cap H_0^1(\omega)$ dans V_γ^0 . De plus, $H^2(\omega) \cap H_0^1(\omega)$ est dense dans Φ_γ .*

Lorsque γ vérifie l'encadrement du théorème, pour tout $u \in V_\gamma^0$, il existe donc une unique solution dans $V_\gamma^2 \cap H_0^1(\omega)$ au problème :

$$\text{Trouver } \phi \in H_0^1(\omega) \text{ tel que } -\Delta \phi = u \text{ dans } \omega.$$

D'après la proposition 2.63, on obtient ainsi la densité des fonctions régulières dans Φ_γ , et par conséquent, la densité de $\mathbf{X}_{\mathbf{E},\gamma}^0 \cap \mathbf{H}^1(\omega)$ dans $\mathbf{X}_{\mathbf{E},\gamma}^0$.

Le poids γ doit donc vérifier :

$$1 - \alpha < \gamma < 1.$$

Remarque 2.65 Les expériences numériques confirment le fait que lorsque $\gamma \leq 1 - \alpha$, on ne capte plus les singularités du champ électrique.

Proposition 2.66 Le problème (2.6)-(2.7) est équivalent à la formulation variationnelle suivante : Trouver $\mathbf{E}_\gamma^0 \in \mathbf{X}_{\mathbf{E},\gamma}^0$ tel que :

$$\forall \mathbf{F} \in \mathbf{X}_{\mathbf{E},\gamma}^0, \mathcal{A}_\gamma(\mathbf{E}_\gamma^0, \mathbf{F}) = \mathcal{L}_\gamma(\mathbf{F}) - \mathcal{A}_\gamma(\mathbf{e}, \mathbf{F}), \quad (2.69)$$

où $\mathcal{L}_\gamma$ est la forme linéaire suivante :

$$\begin{aligned} \mathcal{L}_\gamma : \mathbf{X}_{\mathbf{E},\gamma}^0 &\rightarrow \mathbb{R} \\ \mathbf{F} &\mapsto (g, \operatorname{div} \mathbf{F})_{0,\gamma} + (f, \operatorname{rot} \mathbf{F})_0. \end{aligned}$$

On remarque que cette formulation variationnelle est similaire à (2.43), au produit scalaire de la divergence près. Ainsi, la discrétisation de (2.69) et celle de (2.43) seront analogues.

En pratique, la convergence des éléments finis est d'autant meilleure que γ est proche de 1, mais la norme sur la divergence est moins bonne. On choisira donc $\gamma \simeq 0.95$ dans les tests numériques.

Remarques 2.67 Cette méthode consiste à relaxer la condition sur la divergence, ce qui permet de contrer les singularités du champ électrique, de type $\mathbf{grad} \varphi_D^S$, qui dérivent des singularités primales du Laplacien. Elle se transpose aisément au problème tridimensionnel. D'après M. Dauge (communication privée), la méthode à poids peut s'appliquer au problème quasi-magnétostatique.

2.5.4 Cas où il existe plusieurs coins rentrants

Lorsqu'il existe plusieurs coins rentrants, on considère l'espace V_γ^0 , où γ est le multi-indice $\{\gamma_1, \dots, \gamma_{N_{cr}}\}$, tel que :

$$V_\gamma^0 = \{u \in L_{loc}^2(\omega) : \forall i \in \{1, \dots, N_{cr}\}, \prod_{i=1}^{N_{cr}} r_i^{\gamma_i} u \in L^2(\omega)\}. \quad (2.70)$$

Le produit scalaire associé s'écrit ainsi :

$$\forall u, v \in V_\gamma^0, (u, v)_{0,\gamma} = \int_\omega \prod_{i=1}^{N_{cr}} r_i^{2\gamma_i} u v \, d\omega.$$

En pratique, chaque γ_i doit être tel que : $1 - \alpha_i < \gamma_i < 1$.

2.6 Milieux inhomogènes

Les méthodes de résolution que nous avons présenté concernent les milieux homogènes. Que se passe-t-il si ε varie ? Posons :

$$\begin{aligned} \mathbf{H}(\operatorname{div} \varepsilon, \omega) &:= \{\mathbf{u} \in \mathbf{L}^2(\omega) : \operatorname{div}(\varepsilon \mathbf{E}) \in L^2(\omega)\}, \text{ et} \\ \mathbf{X}_{\mathbf{E},\varepsilon} &:= \{\mathbf{u} \in \mathbf{H}(\operatorname{rot}, \omega) \cap \mathbf{H}(\operatorname{div} \varepsilon, \omega) : u_\tau \in L^2(\partial\omega)\}. \end{aligned}$$

Supposons que ε soit constante par morceaux : ω est composé de J sous-domaines $\omega_1, \dots, \omega_J$ tels que $\varepsilon(\mathbf{x}) = \varepsilon_j > 0$ dans ω_j . Si les interfaces présentent des coins rentrants, le champ électromagnétique est singulier au voisinage de ces points singuliers. Dans [55], M. Costabel et al. font l'étude de ces singularités. Notons en particulier que la partie régulière du champ peut être beaucoup moins régulière que dans le cas homogène, même avec des données très régulières. D'après S. Lohrengel et S. Nicaise [79], si $J < 3$, les champs réguliers par morceaux sont denses dans $\mathbf{X}_{\mathbf{E},\varepsilon}$. Ce résultat permet de généraliser la méthode avec conditions aux limites naturelles et la méthode du complément singulier à certains milieux composites. Pour la discrétisation, détaillée par F. Assous et al. dans [11], il faut ajouter des conditions de saut à l'interface du domaine. Les valeurs du champ électrique sont alors doublées aux noeuds de l'interface.

2.7 Conclusion

Nous avons présenté cinq méthodes de résolution de (2.1)-(2.3) :

Méthode des potentiels : Calcul indirect du champ électrique, par dérivation des potentiels.

- C. L. naturelles $\mathbf{E}_\phi = -\mathbf{grad} \phi_D + \mathbf{rot} \phi_N$.
- C. L. essentielles $\mathbf{E}_{\phi^0} = -\mathbf{grad} \phi_D^0 + \mathbf{rot} \phi_N^0 + \mathbf{e}$.

Méthode avec C. L. naturelles : Calcul du champ électrique $\mathbf{E}_{nat}$, dans $\mathbf{X}_{\mathbf{E}}$.

λ -approche : Calcul du champ électrique, dans $\mathbf{X}_{\mathbf{E}}^{0,R} \oplus \text{vect}(\mathbf{x}_i^P)$.

- C. L. essentielles, $\mathbf{E}_\lambda = \tilde{\mathbf{E}}^0 + \mathbf{e}_\lambda + \sum_i \lambda_i \mathbf{x}_i^P$.

MCSO : Calcul du champ électrique, dans $\mathbf{X}_{\mathbf{E}}^{0,R} \oplus^\perp \text{vect}(\mathbf{x}_i^S)$.

- C. L. essentielles, $\mathbf{E}_\perp = \hat{\mathbf{E}}^0 + \mathbf{e} + \sum_i c_i \mathbf{x}_i^S$.

Méthode à poids : Calcul du champ électrique, dans $\mathbf{X}_{\mathbf{E},\gamma}^0$.

- C. L. essentielles, $\mathbf{E}_\gamma = \mathbf{E}_\gamma^0 + \mathbf{e}$ dépend du poids γ choisi.

• La **méthode des potentiel** consiste à dériver le champ électrique à partir des potentiels de la décomposition de Helmholtz, calculés par les éléments finis de Lagrange continus P_k . Cette méthode est la plus ancienne. Le champ ainsi approché étant P_{k-1} discontinu composante par composante, cette méthode est moins précise que les trois autres méthodes, pour lesquels on calcule directement le champ électrique. Cette méthode sert de référence pour vérifier les méthodes directes. La discrétisation de cette méthode fait l'objet de la section 4.3.

• La **méthode avec conditions aux limites naturelles** est la plus simple à mettre en oeuvre et se transpose aisément au problème tridimensionnel, mais elle manque de précision quant à l'approximation de la condition aux limites de conducteur parfait. La discrétisation de cette méthode fait l'objet de la section 4.6.

• La **λ -approche** donne d'excellents résultats numériques, mais elle ne se généralise pas à tous les problèmes trimensionnels, et ne s'applique pas au cas instationnaire. La discrétisation de cette méthode fait l'objet de la section 4.8 .

• La **MCSO** est équivalente à la λ -approche dans le cas quasi-électrostatique, et s'applique au cas instationnaire. Nous avons montré comment utiliser la λ -approche pour obtenir une approximation plus précise du complément singulier orthogonal que les approximations proposées dans [9, 63, 68].

Cette méthode promet des résultats compétitifs en terme de précision dans le cas prismatique. La discrétisation de cette méthode fait l'objet de la section 4.9.

Pour programmer ces deux méthodes, il faut approcher les singularités du Laplacien, ce qui complexifie la programmation. La discrétisation de ces calculs est détaillée dans la section 4.4

- La **méthode à poids** est la méthode qui se rapproche le plus des éléments finis d'arêtes (rappelons que pour $\gamma = 1$, $\mathbf{X}_{\mathbf{E},\gamma}^0 = \mathbf{H}_0(\text{rot}, \omega)$). Cette méthode donne d'excellents résultats numériques, mais dans une norme plus faible que la norme de $\mathbf{X}_{\mathbf{E}}^0$. Notons qu'elle a l'avantage d'être facilement généralisable au problème tridimensionnel, et qu'elle se programme aisément. La discrétisation de cette méthode fait l'objet de la section 4.10.

Chapitre 3

Le problème statique 2D mixte continu

Lorsqu'on veut résoudre le problème dépendant du temps, on doit conserver la relation $\partial_t \rho + \operatorname{div} \mathcal{J} = 0$. Après discrétisation, cette propriété n'est plus exactement vérifiée, notamment lorsque la densité de charge ρ et le vecteur densité de courant $\mathcal{J}$ sont créés par des particules chargées ($\partial_t \rho_h + \operatorname{div} \mathcal{J}_h \neq 0$). Afin d'avoir un meilleur contrôle de la divergence au sens $L^2_{(\gamma)}(\Omega)$ (conservation de la charge), on ajoute donc comme inconnue au problème un multiplicateur de Lagrange sur la divergence.

Nous rappelons à la section 17.1 de la partie IV les espaces fonctionnels des différentes méthodes et les formes bilinéaires associées.

3.1 Rappel sur les multiplicateurs de Lagrange

Ce rappel s'applique aussi bien au problème tridimensionnel qu'au problème bidimensionnel. Dans la perspective de traiter le problème dépendant du temps, nous allons introduire un multiplicateur de Lagrange sur la divergence de $\mathbf{E}$ dans la formulation variationnelle du problème (2.1)-(2.3).

La condition $\operatorname{div} \mathbf{E} = g$ est alors additionnellement prise en compte comme une contrainte. $\mathbf{X}$ désignera l'un des espaces étudiés $\mathbf{X}_{\mathbf{E}}$, $\mathbf{X}_{\mathbf{E}}^{0,R}$ ou $\mathbf{X}_{\mathbf{E},\gamma}^0$. Q désignera l'espace du multiplicateur de Lagrange, c'est-à-dire $L^2(\omega)$ pour les espaces $\mathbf{X}_{\mathbf{E}}$ et $\mathbf{X}_{\mathbf{E}}^{0,R}$, et $L^2_{\gamma}(\omega)$ pour l'espace $\mathbf{X}_{\mathbf{E},\gamma}^0$. Ces espaces sont définis dans la section 1.3.2 et dans le paragraphe 1.4. On considère la formulation variationnelle mixte suivante :

Trouver $(\mathbf{E}, p) \in \mathbf{X} \times Q$ solution de :

$$\begin{aligned} \mathcal{A}(\mathbf{E}, \mathbf{F}) + \mathcal{B}(\mathbf{F}, p) &= \mathcal{L}(\mathbf{F}) \quad \forall \mathbf{F} \in \mathbf{X}, \\ \mathcal{B}(\mathbf{E}, q) &= \mathcal{G}(q) \quad \forall q \in Q, \end{aligned} \tag{3.1}$$

où $\mathcal{A}$ est la forme bilinéaire coercitive associée à $\mathbf{X}$, $\mathcal{L}$ la forme linéaire associée à la formulation variationnelle de (2.1)-(2.3) dans $\mathbf{X}$.

$\mathcal{B}$ est la forme bilinéaire continue suivante :

$$\begin{aligned} \mathcal{B} : \mathbf{X} \times Q &\rightarrow \mathbb{R} \\ (\mathbf{F}, q) &\mapsto (\operatorname{div} \mathbf{F}, q)_Q, \end{aligned} \tag{3.2}$$

et $\mathcal{G}$ est la forme linéaire telle que :

$$\begin{aligned} \mathcal{G} : Q &\rightarrow \mathbb{R} \\ q &\mapsto (g, q)_Q. \end{aligned}$$

D'après [12] et [28], le problème est bien posé si la condition inf-sup suivante est satisfaite :

Propriété 3.1 *Il existe une constante $\kappa > 0$ telle que :*

$$\inf_{q \in Q} \sup_{\mathbf{F} \in \mathbf{X}} \frac{\mathcal{B}(\mathbf{F}, q)}{\|\mathbf{F}\|_{\mathbf{X}} \|q\|_Q} \geq \kappa.$$

Nous allons montrer que la propriété 3.1 est vérifiée pour les espaces considérés, puis que la formulation mixte est équivalente à la formulation variationnelle, c'est à dire que $p = 0$ dans l'espace Q . Notons que ce qui suit se transpose automatiquement au problème tridimensionnel.

3.2 Formulation mixte 2D : CL naturelles

Considérons le problème (3.1) avec $\mathbf{X} = \mathbf{X}_{\mathbf{E}}$. On a alors : $\mathcal{A} = \mathcal{A}_{\mathbf{E}}$ et $\mathcal{L} = \mathcal{L}_{\mathbf{E}}$. Comme $\operatorname{div} \mathbf{E} \in L^2(\omega)$, on prend $Q = L^2(\omega)$, et l'on notera $\mathcal{B}_{\mathbf{E}}$ la forme bilinéaire, et $\mathcal{G}_{\mathbf{E}}$ la forme linéaire associées :

$$\begin{aligned} \mathcal{B}_{\mathbf{E}} : \mathbf{X}_{\mathbf{E}} \times L^2(\omega) &\rightarrow \mathbb{R} \\ (\mathbf{F}, q) &\mapsto (\operatorname{div} \mathbf{F}, q)_0; \end{aligned} \quad (3.3)$$

$$\begin{aligned} \mathcal{G}_{\mathbf{E}} : L^2(\omega) &\rightarrow \mathbb{R} \\ q &\mapsto (g, q)_0. \end{aligned} \quad (3.4)$$

Proposition 3.2 *Il existe une constante $\kappa_{\mathbf{E}} \geq 1$ telle que :*

$$\inf_{q \in L^2(\omega)} \sup_{\mathbf{F} \in \mathbf{X}_{\mathbf{E}}} \frac{\mathcal{B}_{\mathbf{E}}(\mathbf{F}, q)}{\|\mathbf{F}\|_{\mathbf{X}_{\mathbf{E}}} \|q\|_0} \geq \kappa_{\mathbf{E}}.$$

DÉMONSTRATION. Soit $q \in L^2(\omega)$. Considérons $\phi \in H_0^1(\omega)$ tel que $-\Delta\phi = q$, et $\mathbf{F} = -\mathbf{grad} \phi$ (voir la section 2.2). Alors $\mathbf{F} \in L^2(\omega)$ est à rotationnel nul, à divergence dans $L^2(\omega)$, et vérifie $\mathbf{F}_{\tau} = 0$. On en déduit que $\mathbf{F} \in \mathbf{X}_{\mathbf{E}}$ et que : $\mathcal{B}_{\mathbf{E}}(\mathbf{F}, q) = \|\operatorname{div} \mathbf{F}\|_0^2 = \|\mathbf{F}\|_{\mathbf{X}_{\mathbf{E}}}^2$, et aussi : $\mathcal{B}_{\mathbf{E}}(\mathbf{F}, q) = \|q\|_0^2$, d'où : $\mathcal{B}_{\mathbf{E}}(\mathbf{F}, q) = \|q\|_0 \|\mathbf{F}\|_{\mathbf{X}_{\mathbf{E}}}$. □

Théorème 3.3 *Le problème (2.1)-(2.3) est équivalent à la formulation variationnelle mixte suivante :*

Trouver $(\mathbf{E}, p) \in \mathbf{X}_{\mathbf{E}} \times L^2(\omega)$ tels que :

$$\begin{aligned} \mathcal{A}_{\mathbf{E}}(\mathbf{E}, \mathbf{F}) + \mathcal{B}_{\mathbf{E}}(\mathbf{F}, p) &= \mathcal{L}_{\mathbf{E}}(\mathbf{F}) \quad \forall \mathbf{F} \in \mathbf{X}_{\mathbf{E}}, \\ \mathcal{B}_{\mathbf{E}}(\mathbf{E}, q) &= \mathcal{G}_{\mathbf{E}}(q) \quad \forall q \in L^2(\omega). \end{aligned} \quad (3.5)$$

DÉMONSTRATION. D'après 3.2, la condition inf-sup est satisfaite. D'après la section 3.1, le problème (3.5) est donc bien posé et admet un unique couple $(\mathbf{E}, p) \in \mathbf{X}_{\mathbf{E}} \times L^2(\omega)$ solution.

Pour montrer que (3.5) est équivalente à (2.35), il faut alors prouver que $p = 0$. Soit $\mathbf{F} = -\mathbf{grad} \phi$, où : $\phi \in H_0^1(\omega)$ est solution de $-\Delta\phi = p$. La première équation de (3.5) devient alors :

$$(\operatorname{div} \mathbf{E}, p)_0 + (p, p)_0 = (g, p)_0.$$

D'après la seconde équation de (3.5), $(\operatorname{div} \mathbf{E}, p)_0 = (g, p)_0$, puisque $p \in L^2(\omega)$.

D'où : $\|p\|_0^2 = 0$, et $p = 0$ au sens $L^2(\omega)$. □

3.3 Formulation mixte 2D : MCSO

Considérons le problème (3.1) avec $\mathbf{X} = \mathbf{X}_{\mathbf{E}}^0$. On a alors : $\mathcal{A} = \mathcal{A}_0$ et $\mathcal{L} = \mathcal{L}_0$. Comme $\operatorname{div} \mathbf{E} \in L^2(\omega)$, on prend $Q = L^2(\omega)$, et l'on notera $\mathcal{B}_0$ la restriction à $\mathbf{X}_{\mathbf{E}}^0 \times L^2(\omega)$ de $\mathcal{B}_{\mathbf{E}}$:

$$\begin{aligned} \mathcal{B}_0 : \mathbf{X}_{\mathbf{E}}^0 \times L^2(\omega) &\rightarrow \mathbb{R} \\ (\mathbf{F}, q) &\mapsto (\operatorname{div} \mathbf{F}, q)_0 ; \end{aligned} \quad (3.6)$$

$\mathcal{G}_{\mathbf{E}}$ définira la contrainte. Décomposons $\mathbf{E} = \hat{\mathbf{E}}^0 + \sum_{i=1}^{N_{cr}} c_i \mathbf{x}_i^S$. Nous arrivons alors à :

Théorème 3.4 *Le problème (2.6)-(2.7) est équivalent à la formulation variationnelle mixte suivante :*

Trouver $(\hat{\mathbf{E}}^0, p) \in \mathbf{X}_{\mathbf{E}}^{0,R} \times L^2(\omega)$ et $(c_i)_{i=1,\dots,N_{cr}} \in \mathbb{R}^{N_{cr}}$ tels que :

$$\begin{aligned} \mathcal{A}_0(\hat{\mathbf{E}}^0, \mathbf{F}) + \mathcal{B}_0(\mathbf{F}, p) &= \mathcal{L}_0(\mathbf{F}) - \mathcal{A}_0(\mathbf{e}, \mathbf{F}) \quad \forall \mathbf{F} \in \mathbf{X}_{\mathbf{E}}^{0,R}, \\ \pi \sum_{i=1}^{N_{cr}} c_i \beta^{i,j} + \mathcal{B}_0(\mathbf{x}_j^S, p) &= \pi \lambda^j, \quad \forall j \in \{1, \dots, N_{cr}\}. \end{aligned} \quad (3.7)$$

$$\mathcal{B}_0(\hat{\mathbf{E}}^0, q) + \sum_{i=1}^{N_{cr}} c_i \mathcal{B}_0(\mathbf{x}_i^S, q) = \mathcal{G}_{\mathbf{E}}(q) - (\operatorname{div} \mathbf{e}, q)_0 \quad \forall q \in L^2(\omega).$$

DÉMONSTRATION. $\mathbf{X}_{\mathbf{E}}^0$ est un sous espace de $\mathbf{X}_{\mathbf{E}}$, donc la condition inf-sup est vérifiée automatiquement, puisque le champ $\mathbf{F}$ construit appartient toujours à $\mathbf{X}_{\mathbf{E}}^0$. L'équivalence avec le problème (2.6)-(2.7) se montre de la même façon que pour le théorème (3.3). $\square$

Remarque 3.5 *Dans [39], P. Ciarlet, Jr. et V. Girault donnent une autre preuve de la condition inf-sup continue dans $\mathbf{X}_{\mathbf{E}}^0$ en construisant un relèvement de la condition aux limites normales, utilisé pour la preuve de la condition inf-sup discrète.*

3.4 Formulation mixte 2D : régularisation à poids

Considérons le problème (3.1) avec $\mathbf{X} = \mathbf{X}_{\mathbf{E},\gamma}^0$. On a alors : $\mathcal{A} = \mathcal{A}_{\gamma}$ et $\mathcal{L} = \mathcal{L}_{\gamma}$. Comme $\operatorname{div} \mathbf{E} \in L_{\gamma}^2(\omega)$, on prend $Q = L_{\gamma}^2(\omega)$. On note $\mathcal{B}_{\gamma}$ la forme bilinéaire, et $\mathcal{G}_{\gamma}$ la forme linéaire associées :

$$\begin{aligned} \mathcal{B}_{\gamma} : \mathbf{X}_{\mathbf{E},\gamma}^0 \times L_{\gamma}^2(\omega) &\rightarrow \mathbb{R} \\ (\mathbf{F}, q) &\mapsto (\operatorname{div} \mathbf{F}, q)_{0,\gamma} ; \end{aligned} \quad (3.8)$$

$$\begin{aligned} \mathcal{G}_{\gamma} : L_{\gamma}^2(\omega) &\rightarrow \mathbb{R} \\ q &\mapsto (g, q)_{0,\gamma}. \end{aligned} \quad (3.9)$$

Proposition 3.6 *Il existe une constante $\kappa_{\gamma} \geq 1$ telle que :*

$$\inf_{q \in L_{\gamma}^2(\omega)} \sup_{\mathbf{F} \in \mathbf{X}_{\mathbf{E},\gamma}^0} \frac{\mathcal{B}_{\gamma}(\mathbf{F}, q)}{\|\mathbf{F}\|_{\mathbf{X}_{\mathbf{E},\gamma}^0} \|q\|_0} \geq \kappa_{\gamma}.$$

DÉMONSTRATION. Soit $q \in L^2_\gamma(\omega)$. Considérons $\phi \in H^1_0(\omega)$ tel que $-\Delta\phi = q$. Comme $L^2_\gamma(\omega) \subset H^{-1}(\omega)$ (voir la section 2.5), il existe une unique solution ϕ . Posons $\mathbf{F} = -\mathbf{grad} \phi$. Alors $\mathbf{F} \in \mathbf{L}^2(\omega)$ est à rotationnel nul, à divergence dans $L^2_\gamma(\omega)$, et vérifie $F_\tau = 0$. On en déduit que $\mathbf{F} \in \mathbf{X}^0_{\mathbf{E},\gamma}$ et que : $\mathcal{B}_\gamma(\mathbf{F}, q) = \|\operatorname{div} \mathbf{F}\|_{0,\gamma} = \|\mathbf{F}\|_{\mathbf{X}^0_{\mathbf{E},\gamma}}$.

□

Théorème 3.7 *Le problème (2.6)-(2.7) est équivalent à la formulation variationnelle mixte suivante :*

Trouver $(\mathbf{E}^0_\gamma, p) \in \mathbf{X}^0_{\mathbf{E},\gamma} \times L^2_\gamma(\omega)$ tels que :

$$\begin{aligned} \mathcal{A}_\gamma(\mathbf{E}^0_\gamma, \mathbf{F}) + \mathcal{B}_\gamma(\mathbf{F}, p) &= \mathcal{L}_\gamma(\mathbf{F}) - \mathcal{A}_\gamma(\mathbf{e}, \mathbf{F}) \quad \forall \mathbf{F} \in \mathbf{X}^0_{\mathbf{E},\gamma}, \\ \mathcal{B}_\gamma(\mathbf{E}^0_\gamma, q) &= \mathcal{G}_\gamma(q) - (\operatorname{div} \mathbf{e}, q)_{0,\gamma} \quad \forall q \in L^2_\gamma(\omega). \end{aligned} \tag{3.10}$$

DÉMONSTRATION. D'après 3.6, la condition inf-sup est satisfaite. D'après la section 3.1, le problème (3.10) est donc bien posé et admet un unique couple $(\mathbf{E}, p) \in \mathbf{X}_{\mathbf{E},\gamma} \times L^2_\gamma(\omega)$ solution.

Montrons que $p = 0$, et donc que (3.10) est équivalente à (2.69). Prenons $\mathbf{F} = -\mathbf{grad} \phi$, où : $\phi \in H^1_0(\omega)$ est la solution de $-\Delta\phi = p$. La première équation de (3.10) devient alors :

$$(\operatorname{div} \mathbf{E}, p)_{0,\gamma} + (p, p)_{0,\gamma} = (g, p)_{0,\gamma}.$$

D'après la seconde équation de (3.10), $(\operatorname{div} \mathbf{E}, p)_{0,\gamma} = (g, p)_{0,\gamma}$, puisque $p \in L^2_\gamma(\omega)$. D'où : $\|p\|_{0,\gamma}^2 = 0$, et on a bien $p = 0$ au sens $L^2_\gamma(\omega)$.

□

Chapitre 4

Le problème statique 2D direct discret

4.1 Discrétisation par les éléments finis de Galerkin continus

Cette section provient du cours de A.-S. Bonnet-Ben Dhia et É. Luneville [24] : “*Résolution numérique des équations aux dérivées partielles*”.

4.1.1 L’approximation de Galerkin

Considérons le problème variationnel général suivant :
Trouver $w \in V$ tel que :

$$\forall u \in V, a(w, u) = l(u), \quad (4.1)$$

où V est un espace de Hilbert sur ω . a est une forme bilinéaire continue et coercitive sur $V \times V$ et l une forme linéaire continue. D’après le théorème de Lax-Milgram, ce problème admet une unique solution $w \in V$.

Considérons V_h , un sous-espace de V de dimension finie $N(h)$, où h est un paramètre positif tel que $\lim_{h \rightarrow 0} N(h) = +\infty$. On introduit le problème variationnel approché suivant :
Trouver $w_h \in V_h$ tel que :

$$\forall u_h \in V_h, a(w_h, u_h) = l(u_h). \quad (4.2)$$

V_h est un espace de Hilbert, comme sous-espace fermé de V , on peut donc appliquer le théorème de Lax-Milgram pour affirmer que le problème approché (4.2) admet une unique solution $v_h \in V_h$. On dit que (4.2) est une *approximation interne* ou une *approximation de Galerkin* de (4.1). Le lemme de Céa 4.1 ci-dessous permet de démontrer la convergence de l’approximation interne (4.2).

Lemme 4.1 *Il existe une constante $C > 0$ indépendante de V_h telle que :*

$$\|w - w_h\|_V \leq C \inf_{u_h \in V_h} \|w - u_h\|_V. \quad (4.3)$$

Théorème 4.2 *On suppose qu’il existe un sous-espace W de V dense dans V , et pour tout h une application $r_h : W \rightarrow V_h$ tels que : $\forall u \in W$,*

$$\lim_{h \rightarrow 0} \|u - r_h(u)\|_V = 0. \quad (4.4)$$

Alors, on a : $\lim_{h \rightarrow 0} \|w - w_h\|_V = 0$.

En général, V est un espace de Sobolev, dans lequel les fonctions régulières sont denses. On choisit alors pour W un espace de type $C^m(\omega)$, $m \geq 0$.

On appelle base hilbertienne de V toute suite $(v_i)_{i \geq 1}$ d'éléments de V telle que :

- $\forall i, j \geq 1$ tels que $i \neq j$, $(v_i, v_j)_V = 0$: les v_i sont orthogonaux entre eux,
- les espaces vectoriels engendrés par des combinaisons linéaires finies des $(v_i)_{i \geq 1}$ sont denses dans V .

L'approximation de Ritz-Galerkin consiste à prendre $V_h = \text{vect}(v_1, v_2, \dots, v_N)$, où $h = \frac{1}{N}$, et w_h l'unique solution de (4.2).

Proposition 4.3 *Le problème approché (4.2) est équivalent au système linéaire suivant :*

$$\mathbb{A} \underline{w} = \underline{l}, \quad (4.5)$$

où la matrice $\mathbb{A} \in \mathbb{R}^{N \times N}$ est telle que : $\mathbb{A}_{I,J} = a(v_J, v_I)$; $\underline{l} \in \mathbb{R}^N$ est le vecteur de composantes :

$$\underline{l}_I = l(v_I) ; \text{ et le vecteur } \underline{w} \in \mathbb{R}^N \text{ est tel que : } w_h = \sum_{J=1}^N \underline{w}_J v_J.$$

Autrement dit, les $\underline{w}_J$ sont les composantes de w_h dans la base $(v_1, v_2, \dots, v_N)$.

DÉMONSTRATION. Au lieu de considérer $u_h \in V_h$ quelconque, il est équivalent de choisir $u_h = v_i$ dans la formulation variationnelle (4.2), ce qui conduit aux N équations : $\forall i \in \{1, \dots, N\}$,

$$\sum_{j=1}^N \underline{w}_j a(v_j, v_i) = l(v_i).$$

□

4.1.2 Les éléments finis de Lagrange continus d'ordre k

Définition 4.4 *On appelle élément fini de Lagrange continu d'ordre k tout triplet (K, Σ, P) où K est un ensemble fermé borné non vide de $\mathbb{R}^2$, Σ un ensemble de N_K points $(M_I^K)_{I=1, N_K}$ de K , et P un espace vectoriel de polynômes contenant $\mathbb{P}^k(K)$ Σ -unisolvant :*

$$\forall (c_1, c_2, \dots, c_{N_K}) \in \mathbb{R}^{N_K}, \exists ! p \in P \mid \forall I \in \{1, \dots, N_K\}, p(M_I^K) = c_i.$$

Il existe donc N_K fonctions de base *locales* telles que :

$$\forall I, J \in \{1, \dots, N_K\}, v_I^K(M_J^K) = \delta_{IJ}, \quad (4.6)$$

où δ_{IJ} est le symbole de Kronecker. Les $(v_I^K)_{I=1, N_K}$ forment une famille libre et engendrent le sous-espace vectoriel P .

Remarque 4.5 *L'élément K peut prendre n'importe quelle forme (triangle, quadrangle), et les points $(M_I^K)_{I=1, N_K}$ ne sont pas nécessairement des sommets.*

L'erreur de calcul est liée à l'ordre de l'élément fini k .

On crée alors une partition du domaine d'étude ω en L éléments $(K_l)_{l=1, L}$, on définit $(\Sigma_l, P_l)_{l=1, L}$, et l'on numérote l'ensemble total $\left((M_I^{K_l})_{I=1, N_{K_l}} \right)_{l=1, L}$ des points de façon globale : $(M_i)_{i=1, N}$.

Deux points numérotés localement M_I^K et $M_{I'}^{K'}$ lorsque $\partial K \cap \partial K' \neq \emptyset$, peuvent correspondre au

même point de numéro global i .

À chaque M_i , on associe une fonction *continue* v_i , dont la restriction à chaque élément K est un polynôme de degré k , telle que : $\forall j \in \{1, \dots, N\}, v_i(M_j) = \delta_{ij}$. Supposons que le point M_i , appartenant à l'élément K_l ait la numérotation locale $M_1^{K_l}$. Alors la fonction globale v_i associée à M_i est telle que : $v_i|_{K_l} = v_1^{K_l}$.

On construit ainsi l'espace d'approximation :

$$V_h = \{ u_h \in V \cap C^0(\bar{\omega}) \mid \forall l \in \{1, \dots, L\}, u_h|_{K_l} \in P_l \}. \quad (4.7)$$

Cet espace est généré par les fonctions $(v_i)_{i=1, N}$: v_i a pour support l'union des éléments K_l contenant M_i .

Nous allons ainsi discrétiser les problèmes variationnels posés dans le chapitre 2. Les éléments K seront des triangles vérifiant certaines propriétés.

4.2 Discrétisation du domaine d'étude

On construit un maillage du domaine ω constitué de L triangles $\{T_l, l = 1, \dots, L\}$, d'intérieurs non vides, tels que $\cup_l T_l = \bar{\omega}$, et vérifiant :

$$\forall l, l' \in \{1, \dots, L\} T_l \cap T_{l'} = \begin{cases} \text{soit } \emptyset, \\ \text{soit un sommet commun,} \\ \text{soit une arête commune.} \end{cases}$$

Ces propriétés définissent un maillage conforme. On note $h = \max_l h_l$, où h_l est le rayon du cercle



FIG. 4.1 – Allure du maillage.

circonscrit au triangle T_l .

N_ω désignera le nombre de points de discrétisation intérieurs à ω , $N_{\partial\omega}$, le nombre de points de discrétisation sur le bord $\partial\omega$, et $N = N_\omega + N_{\partial\omega}$ le nombre total de points.

Ainsi, on ordonne les points de discrétisation $(M_i)_{i=1, N}$ de la façon suivante :

- $\forall i \in I_\omega, M_i \in \omega$, et
- $\forall i \in I_{\partial\omega}, M_i \in \partial\omega$,

où on a défini les ensembles d'indices suivants :

$$I_\omega = \{1, \dots, N_\omega\}, I_{\partial\omega} = \{N_\omega + 1, \dots, N_\omega + N_{\partial\omega}\} \text{ et } I = I_\omega \cup I_{\partial\omega}.$$

$C^0(\bar{\omega})$ et les éléments finis de Lagrange sont denses [33] dans $H^1(\omega)$. On considère :

$$V_k = \{ v_h \in C^0(\bar{\omega}) \mid \forall l \in \{1, \dots, L\} u_h|_{T_l} \in P_k \}, \quad (4.8)$$

où P_k est l'ensemble des fonctions continues, polynômiale, de degré au plus k .

- Si $k = 1$, $u_h \in V_1$ est une fonction continue, affine par triangle. Les points de discrétisation

du maillage sont les sommets des triangles. Il y a donc trois degrés de liberté par triangle. u_h est complètement définie par ses valeurs aux sommets des triangles T_l .

- Si $k = 2$, $u_h \in V_2$ est une fonction continue, quadratique par triangle. Les points de discrétisation du maillage sont les sommets des triangles et les milieux des arêtes des triangles. Il y a donc six degrés de liberté par triangle. u_h est complètement définie par ses valeurs aux sommets de la triangulation, et aux milieux des arêtes.

V_k est un sous-espace de dimension finie de $H^1(\omega)$. Il est engendré par les fonctions $(v_i)_{i=1,N}$ décrites dans la section 4.1 : $V_k = \text{vect}(v_1, \dots, v_N)$. Pour tout i , v_i a pour support les triangles T_l contenant M_i .

Enfin, les S_q , $q \in \{1, \dots, N_{\partial\omega}\}$ désigneront les arêtes du bord (segments constitués par la jonction de deux points consécutifs de la discrétisation de $\partial\omega$). On introduit l'opérateur d'interpolation Π_k tel que :

$$\begin{aligned} \Pi_k : H^1(\omega) \cap C^0(\bar{\omega}) &\rightarrow V_k \\ u &\mapsto \sum_{i=1}^N u(M_i) v_i. \end{aligned}$$

On notera par la suite : $\underline{u} = (u(M_1), u(M_2), \dots, u(M_N))^T$.

Soit π_k l'opérateur d'interpolation au bord :

$$\begin{aligned} \pi_k : H^{1/2}(\partial\omega) \cap C^0(\partial\omega) &\rightarrow V_k \\ u &\mapsto \sum_{i \in I_{\partial\omega}} u(M_i) v_i. \end{aligned}$$

Dans la suite de cette partie, nous ne traitons que la discrétisation du problème quasi-électrostatique. Pour le problème quasi-magnétostatique, tout se passe de façon similaire, en remplaçant τ par ν pour les conditions aux limites.

4.3 Décomposition de Helmholtz : discrétisation

4.3.1 Le Laplacien avec CL de Dirichlet

Nous allons discrétiser (2.13), qui donne une solution faible de (2.9). Il s'agit d'un problème de Dirichlet homogène, nous allons donc introduire V_k^0 le sous-espace de $H_0^1(\omega)$, de dimension N_ω :

$$V_k^0 := V_k \cap H_0^1(\omega) = \{u_h \in V_k : u_h|_{\partial\omega} = 0\}. \quad (4.9)$$

Les $(v_i)_{i \in I_\omega}$ forment une base de V_k^0 . Soit Π_k^0 , l'opérateur d'interpolation associé à cet espace :

$$\begin{aligned} \Pi_k^0 : H_0^1(\omega) \cap C^0(\bar{\omega}) &\rightarrow V_k^0 \\ u &\mapsto \sum_{i \in I_\omega} u(M_i) v_i. \end{aligned}$$

Soit $\phi_{D,h} \in V_k^0$ l'approximation de ϕ_D , telle que : $\phi_{D,h} = \sum_{j \in I_\omega} \phi_{D,h}(M_j) v_j$.

La formulation variationnelle (2.13) discrétisée dans V_k^0 s'écrit :
Trouver $\phi_{D,h} \in V_k^0$ tel que :

$$\forall i \in I_\omega, a_D(\phi_{D,h}, v_i) = l_D(v_i), \quad (4.10)$$

où : $\phi_{D,h} = \sum_{j \in I_\omega} \phi_{D,h}(M_j) v_j$.

Ce problème linéaire est de dimension $\mathbb{R}^{N_\omega}$. Il est résolu *en pratique* dans $\mathbb{R}^N$: on procède par pseudo-élimination des degrés de liberté liés au bord $\partial\omega$, en modifiant le système linéaire (4.10) de la façon suivante :

Trouver $\phi_{D,h} \in V_k$ tel que :

$$\begin{aligned} \forall i \in I_\omega \quad , \quad a_D(\phi_{D,h}, v_i) &= l_D(v_i), \\ \forall i \in I_{\partial\omega} \quad , \quad \phi_{D,h}(M_i) &= 0. \end{aligned}$$

Nous allons exprimer ce problème linéaire sous forme matricielle.

Le second membre peut être calculé de deux façons :

- Si la fonction d'interpolation $g_h := \Pi_k^0(g)$ existe, on calcule le second membre par un produit matrice-vecteur. En pratique, on utilise une interpolation de ce type (notamment lors de couplages avec des particules [63]).
- Sinon, on construit directement le vecteur $\underline{G} \in \mathbb{R}^N$, de composantes :

$$\begin{aligned} \underline{G}_i &= l_D(v_i) = (g, v_i)_0 \text{ si } i \in I_\omega, \\ &= 0 \text{ si } i \in I_{\partial\omega} \end{aligned}$$

Pour un exemple d'approche générale, voir les exemples numériques du paragraphe 6.2.3.

Supposons que g_h existe. On écrit alors l'équation (4.10) sous la forme :

Trouver $\phi_{D,h} \in V_k^0$ tel que :

$$\forall i \in I_\omega, \quad \sum_{j \in I_\omega} \phi_{D,h}(M_j) a_D(v_j, v_i) = (g_h, v_i)_0 = \sum_{j \in I_{\partial\omega}} g(M_j) (v_j, v_i)_0, \quad (4.11)$$

avec : $a_D(v_j, v_i) = (\mathbf{grad} v_j, \mathbf{grad} v_i)_0 = (\partial_1 v_j, \partial_1 v_i)_0 + (\partial_2 v_j, \partial_2 v_i)_0$.

Considérons $\mathbb{K}_D \in \mathbb{R}^{N \times N}$, appelée *matrice de raideur interne* (ou matrice de rigidité interne), la matrice d'éléments :

$$\begin{aligned} \mathbb{K}_D^{i,j} &= a_D(v_j, v_i) \text{ si } i \text{ et } j \in I_\omega, \\ &= \delta_{ij} \text{ si } i \text{ ou } j \in I_{\partial\omega}. \end{aligned} \quad (4.12)$$

$\forall i, j, \mathbb{K}_D^{i,j} = \mathbb{K}_D^{j,i}$: $\mathbb{K}_D$ est symétrique. $\mathbb{K}_D$ est formée de 2 sous-blocs :

$$\mathbb{K}_D = \begin{pmatrix} \mathbb{K}_\omega & 0 \\ 0 & \mathbb{I}_{\partial\omega} \end{pmatrix},$$

où $\mathbb{K}_\omega \in \mathbb{R}^{N_\omega \times N_\omega}$ est telle que : $\forall i, j \in I_\omega, \mathbb{K}_\omega^{i,j} = \mathbb{K}_D^{i,j}$; et $\mathbb{I}_{\partial\omega} \in \mathbb{R}^{N_{\partial\omega} \times N_{\partial\omega}}$ est la matrice identité.

On remarque de plus que pour $i, j \in I_\omega$, si M_i et M_j n'appartiennent pas à un même triangle, l'élément $\mathbb{K}_D^{i,j}$ nul. La matrice $\mathbb{K}_D$ est donc *creuse*. Lorsque M_i et M_j appartiennent à un même triangle, on a :

$$\mathbb{K}_D^{i,j} = \int_\omega (\partial_1 v_j \partial_1 v_i + \partial_2 v_j \partial_2 v_i) d\omega = \sum_{T_l | M_i, M_j \in T_l} \int_{T_l} (\partial_1 v_j \partial_1 v_i + \partial_2 v_j \partial_2 v_i) d\omega.$$

Considérons la matrice $\mathbb{M}_D \in \mathbb{R}^{N \times N}$ appelée *matrice de masse interne*, la matrice d'éléments :

$$\begin{aligned} \mathbb{M}_D^{i,j} &= (v_j, v_i)_0 \text{ si } i \in I_\omega, \forall j, \\ &= 0 \text{ si } i \in I_{\partial\omega}, \forall j. \end{aligned} \quad (4.13)$$

$\mathbb{M}_D$ est formée de deux sous-blocs :

$$\mathbb{M}_D = \begin{pmatrix} \mathbb{M}_\omega & \mathbb{M}_{\omega, \partial\omega} \\ 0 & 0 \end{pmatrix},$$

où $\mathbb{M}_\omega \in \mathbb{R}^{N_\omega \times N_\omega}$, et $\mathbb{M}_{\omega, \partial\omega} \in \mathbb{R}^{N_\omega \times N_{\partial\omega}}$. On remarque que le sous-bloc $\mathbb{M}_\omega$ est symétrique. De même que pour $\mathbb{K}_D$, lorsque M_i et M_j n'appartiennent pas à un même triangle, l'élément $\mathbb{M}_D^{i,j}$ est nul : $\mathbb{M}_D$ est une matrice creuse. Comme précédemment, $\forall i \in I_\omega, \forall j \in I$,

$$\mathbb{M}_D^{i,j} = \sum_{T_l \mid M_i, M_j \in T_l} \int_{T_l} v_j v_i d\omega.$$

Supposons que $\Pi_k g$ existe. Soit $\underline{g} \in \mathbb{R}^N$ tel que : $\underline{g} = (g(M_1), \dots, g(M_N))^T$.

Posons $\underline{\phi}_D \in \mathbb{R}^N$ tel que $\underline{\phi}_D = (\phi_{D,h}(M_1), \dots, \phi_{D,h}(M_{N_\omega}), 0, \dots, 0)^T$.

On peut alors écrire le problème (4.11) sous la forme matricielle suivante :

$$\boxed{\mathbb{K}_D \underline{\phi}_D = \mathbb{M}_D \underline{g}.} \quad (4.14)$$

4.3.2 Le Laplacien avec CL de Neumann

ϕ_N est défini à une constante près, ou à moyenne nulle (voir la section 2.2). Pour résoudre le problème numérique, nous devons imposer la valeur de cette constante. On choisit alors d'imposer $\phi_N(M_N) = 0$. On construit alors le sous-espace de $H^1(\omega)$, de dimension $N - 1$:

$$V_k^N = \{u_h \in V_k : u_h(M_N) = 0\}.$$

Les $(v_i)_{i=1, N-1}$ forment une base de V_k^N . Soit Π_k^N , l'opérateur d'interpolation associé à cet espace :

$$\begin{aligned} \Pi_k^N : H_0^1(\omega) \cap C^0(\overline{\omega}) &\rightarrow V_k^N \\ u &\mapsto \sum_{i \in I_\omega} u(M_i) v_i. \end{aligned}$$

Soit $\phi_{N,h} \in V_k^N$ l'approximation de ϕ_N , telle que : $\phi_{N,h} = \sum_{j=1}^{N-1} \phi_{N,h}(M_j) v_j$.

On note $\underline{\phi}_N = (\phi_{N,h}(M_1), \dots, \phi_{N,h}(M_{N-1}), 0)^T$.

La formulation variationnelle (2.14) discrétisée s'écrit cette fois :

Trouver $\phi_{N,h} \in V_k^N$ tel que :

$$\begin{aligned} \forall i \in I \setminus \{N\}, a_N(\phi_{N,h}, v_i) &= (f, v_i)_0 + (e, v_i)_{0, \partial\omega} \\ \phi_{N,h}(M_N) &= 0. \end{aligned} \quad (4.15)$$

Ce problème linéaire est de dimension $\mathbb{R}^{N-1}$. Il est résolu *en pratique* dans $\mathbb{R}^N$: on procède par pseudo-élimination du degré de liberté liés en M_N . Nous allons exprimer ce problème linéaire sous forme matricielle.

De même que pour la discrétisation du problème de Dirichlet, le second membre peut être calculé de deux façons :

- Si les fonctions d'interpolation $f_h := \Pi_k^N(f)$ et $e_h = \pi_k^N(e)$ existent, on calcule le second membre par un produit matrice-vecteur.
- Sinon, on construit directement le vecteur $\underline{E} \in \mathbb{R}^N$, de composantes :

$$\begin{aligned}\underline{E}_i &= l_N(v_i) = (f, v_i)_0 \text{ si } i \in I \setminus \{N\}, \\ &= 0 \text{ si } i = N,\end{aligned}$$

et le vecteur $\underline{E} \in \mathbb{R}^N$, de composantes :

$$\begin{aligned}\underline{E}_i &= (e, v_i)_0 \text{ si } i \in I_{\partial\omega} \setminus \{N\}, \\ &= 0 \text{ si } i \in I_\omega \cup \{N\},\end{aligned}$$

Supposons que f_h et e_h existent. On écrit alors l'équation (4.15) sous la forme :

Trouver $\phi_{N,h} \in V_k^N$ tel que : $\forall i \in I \setminus \{N\}$,

$$\begin{aligned}\sum_{j \in I \setminus \{N\}} \phi_{N,h}(M_j) a_N(v_j, v_i) &= (f_h, v_j)_0 + (e_h, v_j)_{0,\partial\omega}, \\ &= \sum_{j \in I \setminus \{N\}} f_h(M_j) (v_j, v_i)_0 \\ &\quad + \sum_{j \in I_{\partial\omega} \setminus \{N\}} e_h(M_j) (v_j, v_i)_{0,\partial\omega},\end{aligned} \tag{4.16}$$

avec : $a_N(v_j, v_i) = (\mathbf{grad} v_j, \mathbf{grad} v_i)_0 = (\partial_1 v_j, \partial_1 v_i)_0 + (\partial_2 v_j, \partial_2 v_i)_0$.

Considérons $\mathbb{K}_N \in \mathbb{R}^{N \times N}$ appelée *matrice de raideur interne* la matrice d'éléments :

$$\begin{aligned}\mathbb{K}_N^{i,j} &= a_N(v_j, v_i) \text{ si } i \text{ et } j \in \{1, \dots, N-1\}, \\ &= \delta_{ij} \text{ si } i \text{ ou } j = N.\end{aligned} \tag{4.17}$$

Considérons $\mathbb{M}_N \in \mathbb{R}^{N \times N}$, appelée *matrice de masse interne* la matrice d'éléments :

$$\begin{aligned}\mathbb{M}_N^{i,j} &= (v_j, v_i)_0 \text{ si } i \in \{1, \dots, N-1\} \text{ et } j \in \{1, \dots, N\}, \\ &= 0 \text{ si } j = N.\end{aligned} \tag{4.18}$$

Enfin, considérons $\mathbb{B}_N \in \mathbb{R}^{N \times N}$, appelée *matrice de masse frontière* la matrice d'éléments :

$$\begin{aligned}\mathbb{B}_N^{i,j} &= (v_j, v_i)_{0,\partial\omega} \text{ si } i \text{ et } j \in I_{\partial\omega} \setminus \{N\}, \\ &= 0 \text{ si } i \text{ ou } j \in I_\omega \cup \{N\}.\end{aligned} \tag{4.19}$$

Soit $\underline{f} \in \mathbb{R}^N$ tel que : $\underline{f} = (f(M_1), \dots, f(M_N))^T$.

Soit $\underline{e} \in \mathbb{R}^N$ tel que : $\underline{e} = (0, \dots, 0, e(M_{N_\omega+1}), \dots, e(M_N))^T$.

Posons $\underline{\phi}_N \in \mathbb{R}^N$ tel que $\underline{\phi}_N = (\phi_{N,h}(M_1), \dots, \phi_{N,h}(M_{N-1}), 0)^T$.

On peut alors écrire le problème (4.11) sous la forme matricielle suivante :

$$\mathbb{K}_N \underline{\phi}_N = \mathbb{M}_N \underline{f} + \mathbb{B}_N \underline{e}.$$

(4.20)

4.3.3 Dérivation du champ électrique

Une fois que l'on a évalué $\phi_{D,h}$ et $\phi_{N,h}$, il faut calculer $\mathbf{E}_h$, l'approximation de $\mathbf{E}$, composante par composante. Pour cela, il faut calculer $\mathbf{grad} \phi_{D,h}$ et $\mathbf{rot} \phi_{N,h}$. $\phi_{D,h}$ et $\phi_{N,h}$ sont P_k -continus, donc par dérivation, $\mathbf{E}_h \in (P_{k-1})^2$ -discontinu. Par exemple, si on a utilisé les éléments finis P_1 pour le calcul des potentiels, ceux-ci sont continus, affines par triangles. Ainsi, le champ $\mathbf{E}_h$ dérivé des potentiels sera discontinu, constant par triangle. On a :

$$\begin{aligned} \mathbf{E}_h &= -\mathbf{grad} \phi_{D,h} + \mathbf{rot} \phi_{N,h}, \text{ avec :} \\ &\bullet \phi_{D,h}, \text{ approximation de } \phi_D, \text{ solution de (4.14),} \\ &\bullet \phi_{N,h}, \text{ approximation de } \phi_N, \text{ solution de (4.20).} \end{aligned} \quad (4.21)$$

Pour avoir une représentation continue de $\mathbf{E}_h$ par les éléments finis de Lagrange P_k , on procède à une projection P_{k-1} - P_k [6, 63], c'est-à-dire qu'au lieu de calculer directement les dérivées des potentiels, on résout :

$$(\mathbf{E}_h, v_i \mathbf{e}_\alpha)_0 = -(\mathbf{grad} \phi_{D,h}, v_i \mathbf{e}_\alpha)_0 + (\mathbf{rot} \phi_{N,h}, v_i \mathbf{e}_\alpha)_0, \forall i \in \mathbf{I}, \forall \alpha \in \{1, 2\}. \quad (4.22)$$

Soient $\mathbb{G}_D$ et $\mathbb{R}_N \in (\mathbb{R}^{2 \times 1})^{N \times N}$ les matrices telle que : $\forall i, j \in \mathbf{I}$,

$$\mathbb{G}_D^{i,j} = \begin{pmatrix} (v_i, \partial_1 v_j)_0 \\ (v_i, \partial_2 v_j)_0 \end{pmatrix}, \quad \text{et} \quad \mathbb{R}_N^{i,j} = \begin{pmatrix} (v_i, \partial_2 v_j)_0 \\ -(v_i, \partial_1 v_j)_0 \end{pmatrix}.$$

Soit $\underline{\mathbf{E}}$ la représentation de $\mathbf{E}_h$ suivante : $\underline{\mathbf{E}} = (\mathbf{E}_{h,1}(M_1), \mathbf{E}_{h,2}(M_1), \dots, \mathbf{E}_{h,1}(M_N), \mathbf{E}_{h,2}(M_N))^T$ (voir la section 4.6). On a alors :

$$\underline{\mathbf{E}} = -\mathbb{G}_D \underline{\phi}_D + \mathbb{R}_N \underline{\phi}_N. \quad (4.23)$$

4.4 Calcul des singularités du Laplacien

$\mathbf{X}_{\mathbf{E}}^0 \cap \mathbf{H}^1(\omega)$ n'est pas dense dans $\mathbf{X}_{\mathbf{E}}^0$. Pour approcher le champ électrique $\mathbf{E} = \mathbf{E}^0 + \mathbf{e}$, où $\mathbf{E}^0$ est solution de (2.6)-(2.7), dans $\mathbf{X}_{\mathbf{E}}^0$ par les éléments finis de Lagrange continus, on doit donc décomposer $\mathbf{E}^0$ en une partie $\mathbf{H}^1$ -régulière, et une partie singulière, appelée complément singulier. Deux méthodes sont possibles :

- La λ -approche, pour laquelle l'approximation du champ électrique s'écrit :

$$\mathbf{E}_h = \tilde{\mathbf{E}}_h + \sum_{l=1}^{N_{cr}} \lambda_l^l \mathbf{x}_l^P.$$

La discrétisation du champ électrique se déroule en deux temps :

- Calculs des λ_h^l , qui sont les approximations des λ^l .
- Calcul de $\tilde{\mathbf{E}}_h$, l'approximation de $\tilde{\mathbf{E}}$ dans $\mathbf{X}_{\mathbf{E},k}^{0,R}$ avec relèvement de la condition tangentielle sur $\partial\omega$, décrit dans la section 4.8.

En sommant les contributions $\lambda_h^l \mathbf{x}_l^P$, on obtient la partie singulière du champ électrique.

- La MCSO, pour laquelle l'approximation du champ électrique s'écrit :

$$\mathbf{E}_h = \hat{\mathbf{E}}_h + \sum_{l=1}^{N_{cr}} c_h^l \mathbf{x}_l^S \text{ où : } \mathbf{x}_l^S = \tilde{\mathbf{x}}_l + \sum_{m=1}^{N_{cr}} \beta^{l,m} \mathbf{x}_l^P.$$

La discrétisation du champ électrique se déroule en trois temps :

- Calcul des $\beta^{k,l}$ et les λ^l , afin de déterminer les c_h^l .
- Calcul de $\widehat{\mathbf{E}}_h$, l'approximation de $\widehat{\mathbf{E}}$ dans $\mathbf{X}_{\mathbf{E},k}^{0,R}$ avec relèvement de la condition aux limites tangentielle sur $\partial\omega$, et calcul des c_h^l décrit dans la section 4.9.
- Calcul des $\widetilde{\mathbf{x}}_{l,h}$, les parties régulières des $\mathbf{x}_l^S$. Deux méthodes sont possibles : la méthode des potentiels, ou la λ -approche.

En sommant les contributions $c_h^l \widetilde{\mathbf{x}}_{l,h}$ et $\widehat{\mathbf{E}}_h$, on obtient la partie régulière du champ électrique. La partie singulière est obtenue en ajoutant les contributions $\lambda_h^l \mathbf{x}_l^P$, $l \in \{1, \dots, N_{cr}\}$.

Pour les deux méthodes, il faut d'abord approcher les fonctions $\widetilde{s}_{N,l}$ et $\widetilde{s}_{D,l}$ en chaque point du maillage.

4.4.1 Singularités duales : CL de Dirichlet

On rappelle que $s_{D,l} = \widetilde{s}_{D,l} + s_{D,l}^P$, avec $\widetilde{s}_{D,l} \in H^1(\omega)$ et $s_{D,l}^P = r_l^{-\alpha_l} \sin(\alpha_l \theta_l)$. $\widetilde{s}_{D,l}$ satisfait un problème de Dirichlet non-homogène (2.20). On a : $\widetilde{s}_{D,l}|_{\partial\omega} = -s_{D,l}^P|_{\partial\omega}$. On rappelle que $s_{D,l}^P|_{\partial\omega} = 0$ sur les arêtes du coin rentrant A_l^0 et A_l^Θ , et que $s_{D,l}^P|_{\partial\omega}(M_{N_{cl}}) = 0$, où $M_{N_{cl}}$ est le $l^{\text{ième}}$ coin rentrant.

On peut donc se contenter de calculer $\widetilde{s}_{D,l}$ aux points intérieurs à ω : les $(M_i)_{i \in I_\omega}$, et procéder à une pseudo-élimination des degrés de liberté du bord.

Décomposons $\widetilde{s}_{D,l} \in H^1(\omega)$ de la façon suivante : $\widetilde{s}_{D,l} = s_{D,l}^0 + s_{D,l}^*$, où $s_{D,l}^*$ est un relèvement suffisamment régulier de $(1 - \eta_l(r_l)) \widetilde{s}_{D,l}|_{\partial\omega} = -(1 - \eta_l(r_l)) s_{D,l}^P|_{\partial\omega}$ [42], et $s_{D,l}^0 \in H_0^1(\omega)$ satisfait :

$$\Delta s_{D,l}^0 = -\Delta s_{D,l}^* \quad \text{dans } \omega, \quad (4.24)$$

Ce problème est similaire à (2.9), et on peut l'écrire sous la formulation variationnelle suivante : Trouver $s_{D,l}^0 \in H_0^1(\omega)$ tel que $\forall u \in H_0^1(\omega)$,

$$a_D(s_{D,l}^0, u) = -a_D(s_{D,l}^*, u) \quad \text{dans } \omega, \quad (4.25)$$

On en déduit que $\widetilde{s}_{D,l}$ est solution du problème :

Trouver $\widetilde{s}_{D,l} \in H^1(\omega)$ tel que $\forall u \in H_0^1(\omega)$,

$$a_D(\widetilde{s}_{D,l}, u) = 0 \quad \text{dans } \omega, \quad (4.26)$$

avec $\widetilde{s}_{D,l}|_{\partial\omega} = -s_{D,l}^P|_{\partial\omega}$. La formulation discrète de ce problème est la suivante :

Trouver $\widetilde{s}_{D,l,h} \in V_k$ tel que :

$$\begin{aligned} \forall i \in I_\omega, \quad \sum_{j \in I_\omega} \widetilde{s}_{D,l,h}(M_j) a_D(v_j, v_i) &= \sum_{j \in I_{\partial\omega}} s_{D,l}^P(M_j) a_D(v_j, v_i), \\ \forall j \in I_{\partial\omega} \setminus N_{cl}, \quad \widetilde{s}_{D,l,h}(M_j) &= -s_{D,l}^P(M_j) \end{aligned} \quad (4.27)$$

$$\widetilde{s}_{D,l,h}(M_{N_{cl}}) = 0.$$

Soit $\widetilde{s}_{D,l,h} = \sum_{j \in I_\omega} \widetilde{s}_{D,l,h}(M_j) v_j - \sum_{j \in I_{\partial\omega} \setminus N_{cl}} s_{D,l}^P(M_j) v_j$, l'approximation de $\widetilde{s}_{D,l}$ dans

V_k ainsi construite. On considère $\underline{s}_{D,l} \in \mathbb{R}^N$ le vecteur associé :

$$\begin{aligned} \underline{s}_{D,l} &= (\widetilde{s}_{D,l,h}(M_1), \dots, \widetilde{s}_{D,l,h}(M_N))^T, \\ &= (\widetilde{s}_{D,l,h}(M_1), \dots, \widetilde{s}_{D,l,h}(M_{N_\omega}), -s_{D,l}^P(M_{N_\omega+1}), \dots, -s_{D,l}^P(M_N))^T. \end{aligned}$$

Soit $\mathbb{K}_D^P$ la matrice telle que :

$$\mathbb{K}_D^P = \begin{pmatrix} 0 & \mathbb{K}_{\omega, \partial\omega} \\ 0 & -\mathbb{I}_{\partial\omega} \end{pmatrix},$$

où $\mathbb{I}_{\partial\omega}$ est la matrice identité de $\mathbb{R}^{N_{\partial\omega} \times N_{\partial\omega}}$; et $\mathbb{K}_{\omega, \partial\omega} \in \mathbb{R}^{N_{\omega} \times N_{\partial\omega}}$ est telle que :

$$\mathbb{K}_{\omega, \partial\omega}^{i,j} = a_D(v_j, v_i) \text{ si } i \in I_{\omega}, \text{ et } j \in I_{\partial\omega}.$$

Soit $\underline{s}_{D,l}^P \in \mathbb{R}^N$ le vecteur tel que :

$$\underline{s}_{D,l}^P = \begin{cases} 0 & \text{si } i \in I_{\omega} \text{ ou } i = N_{cl}, \\ r_l(M_i)^{-\alpha_l} \sin(\alpha_l \theta_l(M_i)) & \text{si } i \in I_{\partial\omega} \setminus N_{cl}. \end{cases}$$

Les équations (4.27) s'écrivent sous la forme matricielle suivante :

$$\boxed{\mathbb{K}_D \underline{s}_{D,l} = \mathbb{K}_D^P \underline{s}_{D,l}^P.} \quad (4.28)$$

4.4.2 Singularités duales : CL de Neumann

On rappelle que $s_{N,l} = \tilde{s}_{N,l} + s_{N,l}^P$, avec $\tilde{s}_{N,l} \in H^1(\omega)$ et $s_{N,l}^P = r_l^{-\alpha_l} \cos(\alpha_l \theta_l)$. $\tilde{s}_N$ satisfait le problème de Neumann non-homogène (2.21). Son calcul est similaire au calcul de ϕ_N , décrit dans le paragraphe 4.3.2.

Soit $\tilde{s}_{N,l,h} = \sum_{j \in I} \tilde{s}_{N,l,h}(M_j) v_j$, l'approximation de $\tilde{s}_{N,l}$ dans V_k .

Posons $\underline{s}_{N,l} = (\tilde{s}_{N,l,h}(M_1), \dots, \tilde{s}_{N,l,h}(M_N))^T \in \mathbb{R}^N$.

$s_{N,l}$ est donc connu à une constante près. Comme $s_{N,l} \in S_N \subset L_0^2(\omega)$, $\tilde{s}_{N,l,h}$ est choisi tel que :

$$\int_{\omega} (\tilde{s}_{N,l,h} + s_N^P) d\omega = 0.$$

Soit $\underline{b}_{N,l} \in \mathbb{R}^N$ le vecteur tel que :

$$\underline{b}_{N,l} = \begin{cases} 0 & \text{si } i \in I_{\omega}, \\ - \sum_{S_q | M_i \in S_q} \int_{S_q} \partial_{\nu} s_{N,l}^P v_i d\sigma & \text{si } i \in I_{\partial\omega}, \end{cases}$$

où les intégrales seront calculées par un schéma d'intégration numérique (voir le paragraphe 15.2.1).

Rappelons que si $S_q \in A_l^{0,\Theta}$, alors $\partial_{\nu} s_{N,l}^P|_{S_q} = 0$.

Il s'agit alors de résoudre le problème matriciel :

$$\boxed{\mathbb{K}_N \underline{s}_{N,l} = \underline{b}_{N,l}.} \quad (4.29)$$

Posons $s_h = \sum_{i \in I} (\underline{s}_{N,l})_i v_i$, avec $(\underline{s}_{N,l})_N = (\underline{b}_{N,l})_N$, car $\mathbb{K}_N$ est telle que $\mathbb{K}_N^{i,N} = \delta_{i,N}$ et $\mathbb{K}_N^{N,j} = \delta_{N,j}$.

On a alors : $\tilde{s}_{N,l,h} = s_h - \frac{1}{|\omega|} \left(\int_{\omega} (s_h + s_{N,l}^P) d\omega \right)$.

Pour l'intégration numérique de $s_{N,l}^P$, on décompose ω en $\omega_{R_l} = B(M_{N_{cl}}, R_l)$, où $M_{N_{cl}}$ désigne le

coin rentrant d'indice l , et son complémentaire dans ω . Soit $R_l > 0$ le rayon maximal que ω_{R_l} puisse avoir. On écrit $\int_{\omega} = \int_{\omega_{R_l}} + \int_{\omega \setminus \overline{\omega_{R_l}}}$, ce qui donne ici :

$$\int_{\omega} s_{N,l}^P d\omega = \frac{2R_l^{2-\alpha_l}}{\alpha_l(2-\alpha_l)} + \int_{\omega \setminus \overline{\omega_{R_l}}} s_{N,l}^P d\omega.$$

L'intégrale sur $\omega \setminus \overline{\omega_{R_l}}$ est approchée par un schéma d'intégration numérique (voir le paragraphe 15.2.1).

4.4.3 Coefficients des singularités duales

Pour évaluer les $\beta_{D,N}^{l,l}$, on utilise la méthode de P. Ciarlet, Jr. et J. He, décrite dans [41] pour optimiser la précision de calcul. Cette méthode repose sur l'intégration explicite sur une portion de disque de la partie analytique de la singularité duale. Il s'agit de procéder à la décomposition suivante :

$$\beta_{D,h}^{l,l} = (\|\tilde{s}_{D,l,h}\|_0^2 + 2(\tilde{s}_{D,l,h}, s_{D,l}^P)_0 + \|s_{D,l}^P\|_0^2)/\pi.$$

On décompose de nouveau ω en $\omega_{R_l} = B(M_{N,l}, R_l)$, (voir le paragraphe 4.4.2). On a alors :

$$\|s_{D,l}^P\|_0^2 = \int_{\omega_{R_l}} (s_{D,l}^P)^2 d\omega + \int_{\omega \setminus \overline{\omega_{R_l}}} (s_{D,l}^P)^2 d\omega.$$

On sait calculer analytiquement la première partie de cette intégrale :

$$\int_{\omega_{R_l}} (s_{D,l}^P)^2 d\omega = \int_0^R r_l^{-2\alpha_l} r_l dr \int_0^{\pi/\alpha_l} \sin^2(\alpha_l \theta_l) d\theta = \frac{\pi}{2\alpha_l} \frac{R_l^{2-2\alpha_l}}{(2-2\alpha_l)}.$$

$s_{D,l}^P$ est plus régulier dans $\omega \setminus \overline{\omega_{R_l}}$, on utilise alors la formule d'intégration numérique (voir le paragraphe 15.2.1) pour la seconde partie de l'intégrale, ainsi que pour les termes dépendant de $\tilde{s}_{D,l,h}$. On procède de la même façon pour le calcul de $\beta_N^{l,l}$. La partie analytique est similaire (remplacer sin par cos). Quant au calcul des $\beta_{D,N}^{l,m}$, $l \neq m$, on procède à un schéma d'intégration numérique à sept points (paragraphe 15.2.1).

4.4.4 Coefficients des singularités primales

D'après les calculs sur la simplification du calcul de λ , il y a deux façon d'approcher λ^l . Soit λ_h^l , l'approximation de λ^l . Les deux calculs suivants donnent les mêmes résultats numériques :

$$\begin{aligned} \lambda_h^l = & \frac{1}{\pi} \left(\int_{\omega} (\tilde{s}_{D,l,h} + \Pi_k(s_{D,l}^P)) \Pi_k(g) d\omega + \int_{\omega} (\tilde{s}_{N,l,h} + \Pi_k(s_{N,l}^P)) \Pi_k(f) d\omega \right. \\ & \left. + \int_{\partial\omega} (\tilde{s}_{N,l,h} + \pi_k(s_{N,l}^P)) \pi_k(e) d\sigma \right), \end{aligned}$$

ou bien :

$$\begin{aligned} \lambda_h^l = & \frac{1}{\pi} \left(\int_{\omega} (\tilde{s}_{D,l,h} + \Pi_k(s_{D,l}^P)) (\Pi_k(g) - \Pi_k(\operatorname{div} \mathbf{e})) d\omega \right. \\ & \left. + \int_{\omega} (\tilde{s}_{N,l,h} + \Pi_k(s_{N,l}^P)) (\Pi_k(f) - \Pi_k(\operatorname{rot} \mathbf{e})) d\omega \right). \end{aligned}$$

Cette approximation sera calculée par un schéma d'intégration numérique à sept points (paragraphe 15.2.1).

4.5 Calcul des singularités électromagnétiques

Pour calculer les $\mathbf{x}_l^S$, on peut dériver les fonctions singulières primales $\varphi_{D,l}$ et $\varphi_{N,l}$. Nous indiquons dans cette section comment approcher les parties régulières de ces fonctions. Le calcul est similaire à celui des singularités primales.

4.5.1 Singularités primales : CL de Dirichlet

La partie régulière de la $l^{\text{ième}}$ fonction singulière primale $\tilde{\varphi}_{D,l}$ satisfait l'équation (2.26), et son calcul s'apparente cette fois à celui de $\tilde{s}_{D,l}$. La formulation variationnelle de (2.26), discrétisée dans V_k^0 s'écrit : Trouver $\tilde{\varphi}_{D,h}^l \in V_k^0$ tel que :

$$\forall i \in I_\omega, \quad \int_\omega \mathbf{grad} \tilde{\varphi}_{D,h}^l \cdot \mathbf{grad} v_i d\omega = \int_\omega s_{D,h}^l v_i d\omega, \quad (4.30)$$

$$\forall i \in I_{\partial\omega}, \quad \tilde{\varphi}_{D,h}^l(M_i) = -\beta_{D,h}^{l,m} \varphi_{D,m}^P(M_i).$$

Soit $\underline{\tilde{\varphi}}_{D,l} \in \mathbb{R}^N$, tel que :

$$\underline{\tilde{\varphi}}_{D,l} = (\tilde{\varphi}_{D,l,h}(M_1), \dots, \tilde{\varphi}_{D,l,h}(M_{N_\omega}), -\sum_{l=1}^{N_{cr}} \beta_{D,h}^{l,m} \varphi_{D,m}^P(M_{N_\omega+1}), \dots, -\sum_{l=1}^{N_{cr}} \beta_{D,h}^{l,m} \varphi_{D,m}^P(M_N))^T.$$

On considère $\underline{d}_{D,l} \in \mathbb{R}^N$ le vecteur tel que :

$$\underline{d}_{D,l} = \begin{cases} \sum_{T_i | M_i \in T_i} \int_{T_i} s_{D,l,h} v_i d\omega & \text{si } i \in I_\omega, \\ 0 & \text{si } i \in I_{\partial\omega}. \end{cases}$$

Posons $\underline{\varphi}_{D,l}^P$ le vecteur tel que :

$$\underline{\varphi}_{D,l}^P = \begin{cases} 0 & \text{si } i \in I_\omega, \\ -\sum_{l=1}^{N_{cr}} \beta_{D,h}^{l,m} r_m^{\alpha_m}(M_i) \sin(\alpha_m \theta_m(M_i)) & \text{si } i \in I_{\partial\omega}. \end{cases}$$

Il s'agit alors de résoudre le système matriciel :

$$\boxed{\mathbb{K}_D \underline{\tilde{\varphi}}_{D,l} = \underline{d}_{D,l} + \mathbb{K}_D^P \underline{\varphi}_{D,l}^P.} \quad (4.31)$$

4.5.2 Singularités primales : CL de Neumann

La partie régulière de la $l^{\text{ième}}$ fonction singulière primale $\tilde{\varphi}_{N,l}$ satisfait l'équation (2.27), et son calcul s'apparente à celui de $\tilde{s}_{N,l}$. La formulation variationnelle de (2.27), discrétisée dans V_k s'écrit : Trouver $\tilde{\varphi}_{N,l,h} \in V_k$ tel que : $\forall i \in \{1, \dots, N-1\}$,

$$\int_\omega \mathbf{grad} \tilde{\varphi}_{N,l,h} \cdot \mathbf{grad} v_i d\omega = \int_\omega s_{N,l,h} v_i d\omega + \sum_{l=1}^{N_{cr}} \beta_{N,h}^{l,m} \int_{\partial\omega} \partial_\nu \varphi_{N,m}^P v_i d\sigma. \quad (4.32)$$

Posons $\tilde{\varphi}_{N,l} = (\tilde{\varphi}_{N,l,h}(M_1), \dots, \tilde{\varphi}_{N,l,h}(M_{N-1}), 0)^T$. Le premier membre de (4.32) correspond alors aux lignes du produit $\mathbb{K}_N \tilde{\varphi}_{N,l}$. Soit $\underline{d}_l \in \mathbb{R}^N$ le vecteur tel que :

$$\underline{d}_{N,l} = \begin{cases} \sum_{T_l | M_i \in T_l} \int_{T_l} s_{N,l,h} v_i d\omega & \text{si } i \in I_\omega, \\ \sum_{T_l | M_i \in T_l} \int_{T_l} s_{N,l,h} v_i d\omega + \sum_{S_q | M_i \in S_q} \sum_{l=1}^{N_{cr}} \beta_{N,h}^{l,m} \int_{S_q} \partial_\nu \varphi_{N,m}^P v_i d\sigma & \text{si } i \in I_{\partial\omega}. \end{cases}$$

Les intégrales sont calculées par des schémas d'intégration numérique. Il s'agit alors de résoudre le système matriciel :

$$\boxed{\mathbb{K}_N \tilde{\varphi}_{N,l} = \underline{d}_{N,l}.} \quad (4.33)$$

4.6 Méthode avec CL naturelles : discrétisation

$\mathbf{X}_E \cap \mathbf{H}^1(\omega)$ est dense dans $\mathbf{X}_E$ [40, 51]. On définit alors une discrétisation de $\mathbf{X}_E \cap \mathbf{H}^1(\omega)$ afin de résoudre les problèmes posés dans la section 2.3 à l'aide des éléments finis continus de Lagrange P_k . Une bonne définition de cet espace discrétisé est, a priori, l'espace $\mathbf{X}_k$:

$$\boxed{\mathbf{X}_k = \{ \mathbf{v}_h \in C^0(\bar{\omega})^2 \mid \mathbf{v}_h|_{T_l} \in P_k(T_l)^2, \forall l \in \{1, \dots, L\} \}.} \quad (4.34)$$

Par construction, $\mathbf{X}_k \subset \mathbf{H}^1(\Omega)$. Lorsque $k = 1$, les fonctions $\mathbf{v}_h$ de $\mathbf{X}_k$ sont telles que : $\mathbf{v}_h|_{T_l} = \mathbf{a}|_{T_l} + \mathbf{b}|_{T_l}(\mathbf{x})$, où $\mathbf{a}|_{T_l} \in \mathbb{R}^2$; et $\mathbf{b}|_{T_l}(\mathbf{x})$ est une application linéaire de $\mathbb{R}^2$ dans $\mathbb{R}^2$. On choisit alors comme espace d'approximation de $\mathbf{X}_E \cap \mathbf{H}^1(\omega)$ l'espace : $\mathbf{X}_k = (V_k)^2$, pour lequel deux approches sont possibles :

- D'une part, on peut considérer $\mathbf{X}_k$ comme un espace de dimension finie $2N$ sur $\mathbb{R}$. Soit $(\mathbf{e}_1, \mathbf{e}_2)^T$ une base orthonormale directe de $\mathbb{R}^2$. Les $(v_i \mathbf{e}_\alpha)_{i \in \{1, \dots, N\}, \alpha \in \{1, 2\}}$ forment une base de $\mathbf{X}_k$, ce qui va nous permettre de discrétiser la formulation variationnelle (2.35). On cherche $\mathbf{E}_h$, l'approximation de $\mathbf{E}$ sous la forme :

$$\mathbf{E}_h = \sum_{j \in I} \sum_{\beta=1}^2 E_{h,\beta}(M_j) v_j \mathbf{e}_\beta.$$

- D'autre part, on peut considérer $\mathbf{X}_k$ comme un espace de dimension finie N sur $\mathbb{R}^2$, ce qui consiste à chercher les N valeurs *physiques* de $\mathbf{E}_h$ aux points de la discrétisation, c'est-à-dire :

$$\mathbf{E}_h = \sum_{j \in I} \mathbf{E}_h(M_j) v_j,$$

avec $\mathbf{E}_h(M_j) \in \mathbb{R}^2$. Bien sûr, on a : $\mathbf{E}_h(M_j) = E_{h,1}(M_j) \mathbf{e}_1 + E_{h,2}(M_j) \mathbf{e}_2$ pour tout j , et les deux choix d'approximation coïncident.

En pratique, on utilise la première représentation, qui consiste à considérer des inconnues scalaires. Nous allons la détailler dans ce qui suit.

La formulation variationnelle (2.35) discrétisée dans $\mathbf{X}_k$ s'écrit :
 Trouver $\mathbf{E}_h \in \mathbf{X}_k$ tel que :

$$\forall i \in I, \forall \alpha \in \{1, 2\}, \mathcal{A}_{\mathbf{E}}(\mathbf{E}_h, v_i \mathbf{e}_\alpha) = \mathcal{L}_{\mathbf{E}}(v_i \mathbf{e}_\alpha). \quad (4.35)$$

En décomposant $\mathbf{E}_h$ selon la base canonique de $\mathbf{X}_k$ dans (4.35), on obtient :

$$\forall i \in I, \forall \alpha \in \{1, 2\}, \sum_{j \in I} \sum_{\beta=1}^2 E_{h,\beta}(M_j) \mathcal{A}_{\mathbf{E}}(v_j \mathbf{e}_\beta, v_i \mathbf{e}_\alpha) = \mathcal{L}_{\mathbf{E}}(v_i \mathbf{e}_\alpha).$$

Soit $\mathbb{A}_{\mathbf{E}} \in (\mathbb{R}^{2 \times 2})^{N \times N}$ la matrice ainsi construite, définie par les *sous-blocs* 2×2 $\mathbb{A}_{\mathbf{E}}^{i,j} \in \mathbb{R}^{2 \times 2}$ tels que :

$$\mathbb{A}_{\mathbf{E}}^{i,j} = \begin{pmatrix} \mathcal{A}_{\mathbf{E}}(v_i \mathbf{e}_1, v_j \mathbf{e}_1) & \mathcal{A}_{\mathbf{E}}(v_i \mathbf{e}_2, v_j \mathbf{e}_1) \\ \mathcal{A}_{\mathbf{E}}(v_i \mathbf{e}_1, v_j \mathbf{e}_2) & \mathcal{A}_{\mathbf{E}}(v_i \mathbf{e}_2, v_j \mathbf{e}_2) \end{pmatrix}.$$

Ces sous-blocs agissent sur les vecteurs de $\mathbb{R}^2$ suivants : $\mathbf{E}_h(M_j) = \begin{pmatrix} E_{h,1}(M_j) \\ E_{h,2}(M_j) \end{pmatrix}$.

Par soucis de consistance, on appelle $\underline{\mathbf{L}} \in (\mathbb{R}^2)^N$ le vecteur composé des valeurs $\mathcal{L}_{\mathbf{E}}(v_i \mathbf{e}_\alpha)$, $i \in I$, $\alpha \in \{1, 2\}$, formé de N vecteurs de $\mathbb{R}^2$:

$$\underline{\mathbf{L}} = (\mathcal{L}_{\mathbf{E}}(v_1 \mathbf{e}_1), \mathcal{L}_{\mathbf{E}}(v_1 \mathbf{e}_2), \dots, \mathcal{L}_{\mathbf{E}}(v_N \mathbf{e}_1), \mathcal{L}_{\mathbf{E}}(v_N \mathbf{e}_2))^T.$$

Selon la même idée, posons $\underline{\mathbf{E}} = (E_{h,1}(M_1), E_{h,2}(M_1), \dots, E_{h,1}(M_N), E_{h,2}(M_N))^T \in (\mathbb{R}^2)^N$. Si

on regroupe explicitement les composantes deux à deux, on a aussi : $\underline{\mathbf{E}} = \begin{pmatrix} \mathbf{E}_h(M_1) \\ \vdots \\ \mathbf{E}_h(M_N) \end{pmatrix}$.

Le problème (4.35) s'écrit alors : $\mathbb{A}_{\mathbf{E}} \underline{\mathbf{E}} = \underline{\mathbf{L}}$.

Remarque 4.6 Dans les sections 4.8, 4.9 et 4.10, on utilisera une autre base de décomposition de $\mathbf{X}_k$ car il faudra alors éliminer explicitement les valeurs tangentielles.

Nous allons étudier le calcul des composantes de $\mathbb{A}_{\mathbf{E}}$ et $\underline{\mathbf{L}}$. Décomposons $\mathbb{A}_{\mathbf{E}}$ en 3 matrices : $\mathbb{A}_{\mathbf{E}} = \mathbb{D}iv + \mathbb{R}ot + \mathbb{B}$ telles que : $\forall i, j \in I$,

$$\mathbb{D}iv^{i,j} = \begin{pmatrix} (\operatorname{div} v_j \mathbf{e}_1, \operatorname{div} v_i \mathbf{e}_1)_0 & (\operatorname{div} v_j \mathbf{e}_2, \operatorname{div} v_i \mathbf{e}_1)_0 \\ (\operatorname{div} v_j \mathbf{e}_1, \operatorname{div} v_i \mathbf{e}_2)_0 & (\operatorname{div} v_j \mathbf{e}_2, \operatorname{div} v_i \mathbf{e}_2)_0 \end{pmatrix} = \begin{pmatrix} (\partial_1 v_j, \partial_1 v_i)_0 & (\partial_2 v_j, \partial_1 v_i)_0 \\ (\partial_1 v_j, \partial_2 v_i)_0 & (\partial_2 v_j, \partial_2 v_i)_0 \end{pmatrix},$$

$$\mathbb{R}ot^{i,j} = \begin{pmatrix} (\operatorname{rot} v_j \mathbf{e}_1, \operatorname{rot} v_i \mathbf{e}_1)_0 & (\operatorname{rot} v_j \mathbf{e}_2, \operatorname{rot} v_i \mathbf{e}_1)_0 \\ (\operatorname{rot} v_j \mathbf{e}_1, \operatorname{rot} v_i \mathbf{e}_2)_0 & (\operatorname{rot} v_j \mathbf{e}_2, \operatorname{rot} v_i \mathbf{e}_2)_0 \end{pmatrix} = \begin{pmatrix} (\partial_2 v_j, \partial_2 v_i)_0 & -(\partial_1 v_j, \partial_2 v_i)_0 \\ -(\partial_2 v_j, \partial_1 v_i)_0 & (\partial_1 v_j, \partial_1 v_i)_0 \end{pmatrix},$$

et enfin, avec $\boldsymbol{\tau} = \begin{pmatrix} \tau_1 \\ \tau_2 \end{pmatrix}$, le vecteur tangentiel à $\partial\omega$:

$$\begin{aligned} \mathbb{B}^{i,j} &= \begin{pmatrix} \int_{\partial\omega} v_j v_i \tau_1^2 d\sigma & \int_{\partial\omega} v_j v_i \tau_1 \tau_2 d\sigma \\ \int_{\partial\omega} v_j v_i \tau_1 \tau_2 d\sigma & \int_{\partial\omega} v_j v_i \tau_2^2 d\sigma \end{pmatrix}, \forall i, j \in I_{\partial\omega}, \\ &= 0 \text{ sinon.} \end{aligned}$$

Notons que : $\mathbb{B}^{i,j} = \mathbb{B}^{j,i}$. $\mathbb{B}$ est donc symétrique par blocs 2×2 .

4.6.1 Matrice de raideur interne

On constate que :

$$(\mathbb{R}ot + \mathbb{D}iv)^{i,j} = \begin{pmatrix} c_{i,j} & c'_{i,j} \\ -c'_{i,j} & c_{i,j} \end{pmatrix},$$

où : $c_{i,j} = (\partial_1 v_j, \partial_1 v_i)_0 + (\partial_2 v_j, \partial_2 v_i)_0$, et $c'_{i,j} = (\partial_2 v_j, \partial_1 v_i)_0 - (\partial_1 v_j, \partial_2 v_i)_0$.
On en déduit que : $\forall i \in I, c'_{i,i} = 0$.

Proposition 4.7 $c'_{i,j}$ s'écrit comme une intégrale sur $\partial\omega$:

$$c'_{i,j} = \langle \mathbf{grad} v_i \cdot \boldsymbol{\tau}, v_j \rangle_{H^{-1/2}(\partial\omega), H^{1/2}(\partial\omega)}. \quad (4.36)$$

D'où : $c'_{i,j} = 0$ dès lors que M_i ou M_j n'est pas sur $\partial\omega$.

DÉMONSTRATION. On a en effet :

$$\int_{\omega} (\partial_1 v_i \partial_2 v_j - \partial_2 v_i \partial_1 v_j) d\omega = \int_{\omega} \mathbf{grad} v_i \cdot \mathbf{rot} v_j d\omega.$$

Or, d'après la formule d'intégration par parties (1.15) :

$$\int_{\omega} \mathbf{u} \cdot \mathbf{rot} v d\omega = \int_{\omega} \mathbf{rot} \mathbf{u} \cdot v d\omega + \langle \mathbf{u} \cdot \boldsymbol{\tau}, v \rangle_{H^{-1/2}(\partial\omega), H^{1/2}(\partial\omega)}.$$

Supposons que $\mathbf{u} = \mathbf{grad} w$ avec $w \in H^1(\omega)$, alors $\mathbf{rot} \mathbf{u} = 0$, et on a :

$$\int_{\omega} \mathbf{grad} w \cdot \mathbf{rot} v d\omega = \langle \mathbf{grad} w \cdot \boldsymbol{\tau}, v \rangle_{H^{-1/2}(\partial\omega), H^{1/2}(\partial\omega)}.$$

D'où le résultat. □

On en déduit immédiatement $c'_{i,j} = 0$ si M_i ou M_j n'est pas sur le bord. Ainsi, en regroupant les matrices $\mathbb{R}ot$ et $\mathbb{D}iv$, on peut limiter le calcul des $c'_{i,j}$ aux $i, j \in I_{\partial\omega}$. D'autre part, lorsque $i \neq j$, $c_{i,j}$ et $c'_{i,j}$ sont non nuls seulement si les supports de v_i et v_j se recouvrent, c'est-à-dire si M_i et M_j appartiennent à un même triangle. D'où : $\forall i, j \in I$,

$$c_{i,j} = \sum_{T_l \mid M_i, M_j \in T_l} \int_{T_l} (\partial_1 v_i \partial_1 v_j + \partial_2 v_i \partial_2 v_j) d\omega,$$

et pour $i, j \in I_{\partial\omega}$:

$$c'_{i,j} = \sum_{S_q \mid M_i, M_j \in S_q} \int_{S_q} (\partial_1 v_i \partial_2 v_j - \partial_2 v_i \partial_1 v_j) d\sigma.$$

4.6.2 Matrice de masse du bord

Le bloc $\mathbb{B}^{i,j}$ est non nul seulement si M_i et M_j sont les extrémités d'une même arête du bord $\partial\omega$. On a donc : $\forall i$ ou $j \notin \mathcal{I}_{\partial\omega}$, $\mathbb{B}^{i,j} = 0$ et : $\forall i, j \in \mathcal{I}_{\partial\omega}$,

$$\begin{aligned}\mathbb{B}^{i,j} &= \sum_{S_q \mid M_i, M_j \in S_q} \int_{S_q} v_j v_i d\sigma \begin{pmatrix} \tau_{q,1}^2 & \tau_{q,1} \tau_{q,2} \\ \tau_{q,1} \tau_{q,2} & \tau_{q,2}^2 \end{pmatrix}, \\ &= \sum_{S_q \mid M_i, M_j \in S_q} \int_{S_q} v_j v_i d\sigma (\boldsymbol{\tau}_q \cdot \boldsymbol{\tau}_q^T); \end{aligned}$$

où $\boldsymbol{\tau}_q$ est le vecteur tangentiel à l'arête S_q .

4.6.3 Second membre

De même que pour le calcul des potentiels, le second membre peut-être calculé directement, à l'aide de schémas d'intégration numérique, ou avec les fonctions d'interpolation g_h , f_h et e_h , à condition que l'on puisse définir leur valeur en chaque point de discrétisation (voir la section 4.3). Nous allons développer cette approche. Ceci permet d'écrire le second membre sous forme d'une somme de produits matrice-vecteur, en utilisant $\underline{g}$, $\underline{f}$ et $\underline{e}$.

Soit $\mathbb{L}_g \in (\mathbb{R}^2)^{N \times N}$ la matrice formée des $N \times N$ sous-blocs $\mathbb{L}_g^{i,j} \in \mathbb{R}^2$ tels que : $\forall i, j \in \mathcal{I}$,

$$\mathbb{L}_g^{i,j} = \begin{pmatrix} (v_j, \partial_1 v_i)_0 \\ (v_j, \partial_2 v_i)_0 \end{pmatrix} = \sum_{T_l \mid M_i, M_j \in T_l} \begin{pmatrix} \int_{T_l} v_j \partial_1 v_i d\omega \\ \int_{T_l} v_j \partial_2 v_i d\omega \end{pmatrix}.$$

De même, on considère $\mathbb{L}_f \in (\mathbb{R}^2)^{N \times N}$ la matrice formée des $N \times N$ sous-blocs $\mathbb{L}_f^{i,j} \in \mathbb{R}^2$ tels que : $\forall i, j \in \mathcal{I}$,

$$\mathbb{L}_f^{i,j} = \begin{pmatrix} (v_j, \partial_2 v_i)_0 \\ -(v_j, \partial_1 v_i)_0 \end{pmatrix} = \sum_{T_l \mid M_i, M_j \in T_l} \begin{pmatrix} \int_{T_l} v_j \partial_2 v_i d\omega \\ - \int_{T_l} v_j \partial_1 v_i d\omega \end{pmatrix}.$$

Posons enfin $\mathbb{L}_e \in (\mathbb{R}^2)^{N \times N}$ la matrice définie par les $N \times N$ sous-blocs : $\forall i, j \in \mathcal{I}$,

$$\begin{aligned}\mathbb{L}_e^{i,j} &= \begin{pmatrix} \int_{\partial\omega} v_j v_i \tau_1 d\sigma \\ \int_{\partial\omega} v_j v_i \tau_2 d\sigma \end{pmatrix} = \sum_{S_q \mid M_i, M_j \in S_q} \int_{S_q} v_j v_i d\sigma \boldsymbol{\tau}_q, \forall i, j \in \mathcal{I}_{\partial\omega}, \\ &= 0 \text{ sinon.} \end{aligned}$$

On peut donc calculer $\underline{\mathbf{L}}$ sous la forme matricielle suivante : $\underline{\mathbf{L}} = \mathbb{L}_g \underline{\mathbf{g}} + \mathbb{L}_f \underline{\mathbf{f}} + \mathbb{L}_e \underline{\mathbf{e}}$. Pour approcher le champ électrique, on résout donc le système $2N \times 2N$ suivant :

$$\begin{aligned} \mathbb{A}_{\mathbf{E}} \underline{\mathbf{E}} &= \mathbb{L}_g \underline{\mathbf{g}} + \mathbb{L}_f \underline{\mathbf{f}} + \mathbb{L}_e \underline{\mathbf{e}}, \\ \text{avec : } \mathbb{A}_{\mathbf{E}} &= \mathbb{D}\text{iv} + \mathbb{R}\text{ot} + \mathbb{B}. \end{aligned} \quad (4.37)$$

4.7 Élimination des CL essentielles

Dans cette section, on considère le cas général, pour lequel il existe plusieurs coins rentrants. Soit N_c le nombre de sommets du domaine polygonal ω . Chaque sommet est un point de discrétisation du bord, on appelle $\mathcal{C}$ l'ensemble de ces sommets. On pose $N_a = N_{\partial\omega} - N_c$, le nombre de points de discrétisation du bord qui ne sont pas des sommets. On ordonne les points du bord $\partial\omega$ de la façon suivante :

- $\forall i \in I_a, M_i \in \partial\omega \setminus \mathcal{C}$,
- $\forall i \in I_c, M_i \in \mathcal{C}$,

où on a décomposé $I_{\partial\omega}$ en : $I_a = \{N_\omega + 1, \dots, N_\omega + N_a\}$, et $I_c = \{N_\omega + N_a + 1, \dots, N\}$.

On remarque que : $N_\omega + N_a + N_c = N$.

On considère comme espace de discrétisation de $\mathbf{X}_{\mathbf{E}}^{0,R}$ le sous-espace de $\mathbf{X}_{\mathbf{E},k}$, conforme dans $\mathbf{X}_{\mathbf{E}}^0$:

$$\mathbf{X}_{\mathbf{E},k}^{0,R} = \{ \mathbf{v} \in \mathbf{X}_k \mid \mathbf{v} \cdot \boldsymbol{\tau}|_{\partial\omega} = 0 \text{ sur } \partial\omega \}. \quad (4.38)$$

La discrétisation de (2.43) et (2.69) se fait de la même façon que celle de (2.35), mais il faut en plus prendre en compte la condition aux limites tangentielle dans les matrices $\mathbb{D}\text{iv}$ et $\mathbb{R}\text{ot}$, et dans les termes du second membre. Pour cela, on procède alors à une pseudo-élimination de tous les degrés de liberté connus.

4.7.1 Discrétisation de l'espace avec CL essentielles

Nous allons montrer que l'élimination des composantes tangentielles réduit l'espace d'approximation du champ électrique.

Propriété 4.8 *L'espace $\mathbf{X}_{\mathbf{E},k}^{0,R}$ est de dimension finie $2N - N_a - 2N_c$.*

DÉMONSTRATION. Soit $\mathbf{u} \in \mathbf{X}_{\mathbf{E},k}^{0,R}$. Comme $\mathbf{X}_{\mathbf{E},k}^{0,R}$ est un sous-espace de $\mathbf{X}_k$, il est de dimension

inférieure ou égale à $2N$ et on peut écrire $\mathbf{u} = \sum_{i \in I} \sum_{\alpha=1}^2 u_\alpha(M_i) v_i \mathbf{e}_\alpha$.

Considérons un coin M_i , $i \in I_c$. Alors M_i est l'intersection de 2 arêtes distinctes de la frontière $\partial\omega$: il existe k et $k' \in \{1, \dots, K\}$ tels que $k \neq k'$ et $M_i = A_k \cap A_{k'}$. $\mathbf{u}$ étant dans $\mathbf{X}_{\mathbf{E},k}^{0,R}$, on a :

$$\begin{cases} \mathbf{u}(M_i) \cdot \boldsymbol{\tau}_k &= 0, \\ \mathbf{u}(M_i) \cdot \boldsymbol{\tau}_{k'} &= 0. \end{cases}$$

Comme les vecteurs $\boldsymbol{\tau}_k$ et $\boldsymbol{\tau}_{k'}$ ne sont pas colinéaires, ceci correspond exactement à $\mathbf{u}(M_i) = 0$.

Ceci étant valable pour tous les coins, on peut écrire : $\mathbf{u} = \sum_{i=1}^{N-N_c} \sum_{\alpha=1}^2 u_\alpha(M_i) v_i \mathbf{e}_\alpha$. Dit autrement,

pour chaque vecteur de $\mathbf{X}_{\mathbf{E},k}^{0,R}$, nous éliminons alors les 2 degrés de liberté liés à chaque coin. Il y en a N_c , d'où : $\dim(\mathbf{X}_{\mathbf{E},k}^{0,R}) \leq 2N - 2N_c$.

On peut encore éliminer d'autres degrés de liberté. En effet, considérons un point du bord M_i qui n'est pas un coin : $i \in I_a$. M_i appartient à une arête du bord A_k .

Ecrivons $\mathbf{u}(M_i)$ dans la base locale $(\boldsymbol{\tau}_k, \boldsymbol{\nu}_k)$:

$$\mathbf{u}(M_i) = u_\tau(M_i) \boldsymbol{\tau}_k + u_\nu(M_i) \boldsymbol{\nu}_k.$$

Comme $\mathbf{u} \in \mathbf{X}_{\mathbf{E},k}^{0,R}$, on a : $\mathbf{u}(M_i) \cdot \boldsymbol{\tau}_k = 0$, d'où : $u_\tau(M_i) = 0$.

Posons $(\boldsymbol{\tau}_i, \boldsymbol{\nu}_i) = (\boldsymbol{\tau}_k, \boldsymbol{\nu}_k)$ (il n'y a pas d'ambiguïté car l'arête ne touche pas de coin). On a alors : $\mathbf{u}(M_i) = u_\nu(M_i) \boldsymbol{\nu}_i$. On généralise à tous les points du bord qui ne sont pas des coins, d'où :

$$\begin{aligned} \mathbf{u} &= \sum_{i \in I_\omega} \sum_{\alpha=1}^2 u_\alpha(M_i) v_i \mathbf{e}_\alpha + \sum_{i \in I_a} u_\nu(M_i) v_i \boldsymbol{\nu}_i, \\ &= \sum_{i \in I_\omega} \mathbf{u}(M_i) v_i + \sum_{i \in I_a} u_\nu(M_i) v_i \boldsymbol{\nu}_i. \end{aligned} \quad (4.39)$$

Ainsi, pour chaque vecteur de $\mathbf{X}_{\mathbf{E},k}^{0,R}$, nous pouvons éliminer un degré de liberté par point du bord qui n'est pas un coin : $\dim(\mathbf{X}_{\mathbf{E},k}^{0,R}) \leq 2N - N_a - 2N_c$.

Pour conclure, la famille $B_0 := (v_i \mathbf{e}_\alpha)_{i \in I_\omega, \alpha \in \{1,2\}} \cup (v_i \boldsymbol{\nu}_i)_{i \in I_a}$ est contenue dans $\mathbf{X}_{\mathbf{E},k}^{0,R}$, d'où : $\dim(\mathbf{X}_{\mathbf{E},k}^{0,R}) \geq 2N - N_a - 2N_c$. □

Notons qu'alors $B_0 = (v_i \mathbf{e}_\alpha)_{i \in I_\omega, \alpha \in \{1,2\}} \cup (v_i \boldsymbol{\nu}_i)_{i \in I_a}$ engendre $\mathbf{X}_{\mathbf{E},k}^{0,R}$. On utilisera donc deux types de base de $\mathbb{R}^2$ selon l'emplacement de M_i :

- $\forall i \in I_\omega$, $M_i \in \omega$, on travaille dans la base canonique $(\mathbf{e}_1, \mathbf{e}_2)$,
- $\forall i \in I_a$, $M_i \in \partial\omega \setminus \mathcal{C}$, on travaille dans la base locale $(\boldsymbol{\tau}_i, \boldsymbol{\nu}_i)$,
- Afin de lever l'ambiguïté sur $\boldsymbol{\tau}$ et $\boldsymbol{\nu}$ aux coins du domaine, on travaille dans la base canonique $(\mathbf{e}_1, \mathbf{e}_2)$ pour les $M_i \in \mathcal{C}$, $i \in I_c$.

On appelle B la base de $\mathbf{X}_{\mathbf{E},k}$ ainsi formée :

$$B = (\mathbf{v}_i, \alpha)_{i \in I_\omega \cup I_c, \alpha \in \{1,2\}} \cup (v_i \boldsymbol{\tau}_i)_{i \in I_a} \cup (v_i \boldsymbol{\nu}_i)_{i \in I_a}. \quad (4.40)$$

On décomposera $\mathbf{E}_h^0 \in \mathbf{X}_{\mathbf{E},k}^{0,R}$ dans la base B_0 et $\mathbf{E}_h \in \mathbf{X}_{\mathbf{E},k}$ dans la base B .

4.7.2 Matrice de raideur interne

Soit $\mathbb{A} \in (\mathbb{R}^{2 \times 2})^{N \times N}$ la décomposition de la matrice $\mathbb{D}iv + \mathbb{R}ot$ dans B , que l'on écrit de la façon suivante, les sommets du maillage étant ordonnés comme d'habitude (intérieurs, arêtes, coins) :

$$\mathbb{A} = \begin{pmatrix} \mathbb{A}_{\omega, \omega} & \mathbb{A}_{\omega, a} & \mathbb{A}_{\omega, c} \\ \mathbb{A}_{a, \omega} & \mathbb{A}_{a, a} & \mathbb{A}_{a, c} \\ \mathbb{A}_{c, \omega} & \mathbb{A}_{c, a} & \mathbb{A}_{c, c} \end{pmatrix}.$$

Décrivons les sous-matrices de $\mathbb{A}$.

Les sous-blocs de $\mathbb{A}^{i,j} \in \mathbb{R}^{2 \times 2}$ tels que i et $j \in \mathbb{I}_\omega \cup \mathbb{I}_c$ sont définis dans la base $(\mathbf{e}_1, \mathbf{e}_2)$ et ont été calculés dans le paragraphe 4.6.1. Ils sont contenus dans les sous-matrices suivantes :

- $\mathbb{A}_{\omega, \omega} = (\text{Div} + \text{Rot})^{i \in \mathbb{I}_\omega, j \in \mathbb{I}_\omega} \in (\mathbb{R}^{2 \times 2})^{N_\omega \times N_\omega}$;
- $\mathbb{A}_{c, c} = (\text{Div} + \text{Rot})^{i \in \mathbb{I}_c, j \in \mathbb{I}_c} \in (\mathbb{R}^{2 \times 2})^{N_c \times N_c}$;
- $\mathbb{A}_{\omega, c} = (\text{Div} + \text{Rot})^{i \in \mathbb{I}_\omega, j \in \mathbb{I}_c} \in (\mathbb{R}^{2 \times 2})^{N_\omega \times N_c}$;
- $\mathbb{A}_{c, \omega} = (\text{Div} + \text{Rot})^{i \in \mathbb{I}_c, j \in \mathbb{I}_\omega} \in (\mathbb{R}^{2 \times 2})^{N_c \times N_\omega}$.

La sous-matrice $\mathbb{A}_{a, a} \in (\mathbb{R}^{2 \times 2})^{N_a \times N_a}$ est composée des $N_a \times N_a$ sous-blocs de $\mathbb{R}^{2 \times 2}$, tels que : $\forall i \in \mathbb{I}_a, \forall j \in \mathbb{I}_a$:

$$\mathbb{A}_{a, a}^{i, j} = \begin{pmatrix} \mathcal{A}_0(v_j \boldsymbol{\tau}_j, v_i \boldsymbol{\tau}_i) & \mathcal{A}_0(v_j \boldsymbol{\nu}_j, v_i \boldsymbol{\tau}_i) \\ \mathcal{A}_0(v_j \boldsymbol{\tau}_j, v_i \boldsymbol{\nu}_i) & \mathcal{A}_0(v_j \boldsymbol{\nu}_j, v_i \boldsymbol{\nu}_i) \end{pmatrix}.$$

Ces sous-blocs agissent sur les vecteurs de $\mathbb{R}^2$ décomposés dans la base locale $(\boldsymbol{\tau}_j, \boldsymbol{\nu}_j)$:

$$\mathbf{E}_h(M_j) = \begin{pmatrix} \mathbf{E}_\tau(M_j) \\ \mathbf{E}_\nu(M_j) \end{pmatrix}, \text{ où } \mathbf{E}_\tau(M_j) = \mathbf{E}_h(M_j) \cdot \boldsymbol{\tau}_j \text{ et } \mathbf{E}_\nu(M_j) = \mathbf{E}_h(M_j) \cdot \boldsymbol{\nu}_j.$$

La sous-matrice $\mathbb{A}_{\omega, a} \in (\mathbb{R}^{2 \times 2})^{N_\omega \times N_a}$ est composée des $N_\omega \times N_a$ sous-blocs de $\mathbb{R}^{2 \times 2}$ suivants : $\forall i \in \mathbb{I}_\omega, \forall j \in \mathbb{I}_a$,

$$\mathbb{A}_{\omega, a}^{i, j} = \begin{pmatrix} \mathcal{A}_0(v_j \boldsymbol{\tau}_j, v_i \mathbf{e}_1) & \mathcal{A}_0(v_j \boldsymbol{\nu}_j, v_i \mathbf{e}_1) \\ \mathcal{A}_0(v_j \boldsymbol{\tau}_j, v_i \mathbf{e}_2) & \mathcal{A}_0(v_j \boldsymbol{\nu}_j, v_i \mathbf{e}_2) \end{pmatrix}.$$

Symétriquement, la sous-matrice $\mathbb{A}_{a, \omega} \in (\mathbb{R}^{2 \times 2})^{N_a \times N_\omega}$ est composée des $N_a \times N_\omega$ sous-blocs de $\mathbb{R}^{2 \times 2}$ suivants : $\forall i \in \mathbb{I}_a, \forall j \in \mathbb{I}_\omega$,

$$\mathbb{A}_{a, \omega}^{i, j} = \mathbb{A}_{\omega, a}^{j, i} = \begin{pmatrix} \mathcal{A}_0(v_j \mathbf{e}_1, v_i \boldsymbol{\tau}_i) & \mathcal{A}_0(v_j \mathbf{e}_2, v_i \boldsymbol{\tau}_i) \\ \mathcal{A}_0(v_j \mathbf{e}_1, v_i \boldsymbol{\nu}_i) & \mathcal{A}_0(v_j \mathbf{e}_2, v_i \boldsymbol{\nu}_i) \end{pmatrix}.$$

La sous-matrice $\mathbb{A}_{a, c} \in (\mathbb{R}^{2 \times 2})^{N_a \times N_c}$ est composée des $N_a \times N_c$ sous-blocs de $\mathbb{R}^{2 \times 2}$ suivants : $\forall i \in \mathbb{I}_a, \forall j \in \mathbb{I}_c$,

$$\mathbb{A}_{a, c}^{i, j} = \begin{pmatrix} \mathcal{A}_0(v_j \mathbf{e}_1, v_i \boldsymbol{\tau}_i) & \mathcal{A}_0(v_j \mathbf{e}_2, v_i \boldsymbol{\tau}_i) \\ \mathcal{A}_0(v_j \mathbf{e}_1, v_i \boldsymbol{\nu}_i) & \mathcal{A}_0(v_j \mathbf{e}_2, v_i \boldsymbol{\nu}_i) \end{pmatrix}.$$

Symétriquement, la sous-matrice $\mathbb{A}_{c, a} \in (\mathbb{R}^{2 \times 2})^{N_c \times N_a}$ est composée des $N_c \times N_a$ sous-blocs de $\mathbb{R}^{2 \times 2}$ suivants : $\forall i \in \mathbb{I}_c, \forall j \in \mathbb{I}_a$,

$$\mathbb{A}_{c, a}^{i, j} = \mathbb{A}_{a, c}^{j, i} = \begin{pmatrix} \mathcal{A}_0(v_j \boldsymbol{\tau}_j, v_i \mathbf{e}_1) & \mathcal{A}_0(v_j \boldsymbol{\tau}_j, v_i \mathbf{e}_2) \\ \mathcal{A}_0(v_j \boldsymbol{\nu}_j, v_i \mathbf{e}_1) & \mathcal{A}_0(v_j \boldsymbol{\nu}_j, v_i \mathbf{e}_2) \end{pmatrix}.$$

Nous ne détaillerons pas les calculs des éléments de $\mathbb{A}_{a, a}$ etc, car ils sont similaires aux calculs des éléments de Rot et Div , expliqués dans le paragraphe 4.6.1, où l'on a remplacé $v_i \mathbf{e}_1$ par $v_i \boldsymbol{\tau}_i$ et $v_i \mathbf{e}_2$ par $v_i \boldsymbol{\nu}_i$.

4.7.3 Second membre

Soit $\mathbf{F} \in \mathbf{X}_k$. Afin de définir le produit matrice-vecteur $\mathbb{A} \underline{\mathbf{F}}$, où $\underline{\mathbf{F}}$ est le vecteur de $(\mathbb{R}^2)^N$ associé à $\mathbf{F}$ dans la base B , on décompose $\mathbf{F}$ dans cette même base selon :

$$\mathbf{F} = \sum_{j \in I_\omega \cup I_c} \sum_{\beta=1}^2 F_\beta(M_j) v_j \mathbf{e}_\beta + \sum_{j \in I_a} (F_\tau(M_j) v_j \boldsymbol{\tau}_j + F_\nu(M_j) v_j \boldsymbol{\nu}_j),$$

où $\forall i \in I_a$, $F_\tau(M_j) = \mathbf{F}(M_i) \cdot \boldsymbol{\tau}_i$ et $F_\nu(M_j) = \mathbf{F}(M_i) \cdot \boldsymbol{\nu}_i$.

On considère alors $\underline{\mathbf{F}} = (\underline{\mathbf{F}}_\omega, \underline{\mathbf{F}}_a, \underline{\mathbf{F}}_c)^T$ le vecteur de $\mathbb{R}^{2N}$ associé, dont les sous-composantes sont :

- $\underline{\mathbf{F}}_\omega = (F_1(M_1), F_2(M_1), \dots, F_1(M_{N_\omega}), F_2(M_{N_\omega}))^T \in (\mathbb{R}^2)^{N_\omega}$,
- $\underline{\mathbf{F}}_a = (F_\tau(M_{N_\omega+1}), F_\nu(M_{N_\omega+1}), \dots, F_\tau(M_{N-N_c}), F_\nu(M_{N-N_c}))^T \in (\mathbb{R}^2)^{N_a}$,
- $\underline{\mathbf{F}}_c = (F_1(M_{N-N_c+1}), F_2(M_{N-N_c+1}), \dots, F_1(M_N), F_2(M_N))^T \in (\mathbb{R}^2)^{N_c}$.

Le produit matrice-vecteur $\mathbb{A} \underline{\mathbf{F}}$ est ainsi bien défini. Pour construire la matrice de la formulation variationnelle (2.43) discrétisées dans $\mathbf{X}_{\mathbf{E},k}^{0,R}$, on peut alors procéder à l'élimination des conditions aux limites essentielles dans $\mathbb{A}$. Cela correspond à éliminer les lignes et les colonnes de $\mathbb{A}$ dont les éléments dépendent de $\boldsymbol{\tau}$ pour les degrés de liberté d'arêtes, et les couples de ligne et de colonnes de $\mathbb{A}$ pour les degrés de liberté de coin. On appelle $\mathbb{A}_0$ la matrice ainsi obtenue.

Soit $\mathbb{A}_{\omega,\nu} \in (\mathbb{R}^{2 \times 2})^{N_\omega \times N_a}$ la matrice $\mathbb{A}_{\omega,a}$ dont on a éliminé les colonnes dépendant de $\boldsymbol{\tau}$: $\forall i \in I_\omega, \forall j \in I_a$,

$$\mathbb{A}_{\omega,\nu}^{i,j} = \begin{pmatrix} 0 & \mathcal{A}_0(v_j \boldsymbol{\nu}_j, v_i \mathbf{e}_1) \\ 0 & \mathcal{A}_0(v_j \boldsymbol{\nu}_j, v_i \mathbf{e}_2) \end{pmatrix}.$$

Symétriquement, $\mathbb{A}_{\nu,\omega} \in (\mathbb{R}^{2 \times 2})^{N_a \times N_\omega}$ est la matrice $\mathbb{A}_{a,\omega}$ dont on a éliminé les lignes qui dépendent de $\boldsymbol{\tau}$:

$$\mathbb{A}_{\nu,\omega}^{i,j} = \begin{pmatrix} 0 & 0 \\ \mathcal{A}_0(v_j \mathbf{e}_1, v_i \boldsymbol{\nu}_i) & \mathcal{A}_0(v_j \mathbf{e}_2, v_i \boldsymbol{\nu}_i) \end{pmatrix}.$$

Ce faisant, on n'a pas modifié de terme extra diagonal, puisque $\mathbb{A}_{\omega,a}$ et $\mathbb{A}_{a,\omega}$ sont des sous-blocs de $\mathbb{A}$ extra-diagonaux.

On considère $\mathbb{A}_{\nu,\nu}$ la matrice $\mathbb{A}_{a,a}$ dont on a éliminé les termes dépendant de $\boldsymbol{\tau}$, sauf les termes diagonaux, pour lesquels on a imposé la valeur 1 :

$$\mathbb{A}_{\nu,\nu}^{i,j} = \begin{pmatrix} \delta_{ij} & 0 \\ 0 & \mathcal{A}_0(v_j \boldsymbol{\nu}_j, v_i \boldsymbol{\nu}_i) \end{pmatrix}.$$

En effet, $\mathbb{A}_{a,a}$ correspondant à un bloc diagonal de $\mathbb{A}$, il faut faire attention à préserver l'inversibilité de la matrice ainsi modifiée. Ensuite, on remplace tous les sous-blocs liés aux coins, sauf le bloc diagonal, par zéro ; et enfin, on remplace $\mathbb{A}_{c,c}$ par $\mathbb{I}_{c,c}$, la matrice identité de $(\mathbb{R}^{2 \times 2})^{N_c \times N_c}$.

On en déduit que la structure par blocs de la matrice $\mathbb{A}_0 \in (\mathbb{R}^{2 \times 2})^{N \times N}$ est la suivante :

$$\mathbb{A}_0 = \begin{pmatrix} \mathbb{A}_{\omega, \omega} & \mathbb{A}_{\omega, \nu} & 0 \\ \mathbb{A}_{\nu, \omega} & \mathbb{A}_{\nu, \nu} & 0 \\ 0 & 0 & \mathbb{I}_{c, c} \end{pmatrix}.$$

Proposition 4.9 *La matrice $\mathbb{A}_0$ est inversible.*

DÉMONSTRATION. Comme $\mathcal{A}_0$ est une forme bilinéaire symétrique, les sous-blocs diagonaux de $\mathbb{A}_{\omega, \omega}$ et $\mathbb{A}_{\nu, \nu}$ sont strictement positifs. On a :

$$\mathbb{A}_{\omega, \omega}^{i, i} = \begin{pmatrix} \|\mathbf{grad} v_i\|_0^2 & 0 \\ 0 & \|\mathbf{grad} v_i\|_0^2 \end{pmatrix}, \quad \text{et } \mathbb{A}_{\nu, \nu}^{i, i} = \begin{pmatrix} 1 & 0 \\ 0 & \|\mathbf{grad} v_i\|_0^2 \end{pmatrix}.$$

Le sous-bloc diagonal $\mathbb{I}_{c, c}$ garantit donc l'inversibilité de $\mathbb{A}_0$. □

4.8 λ -approche : discrétisation

4.8.1 Champ électrique : partie régulière

Pour la λ -approche, l'approximation du champ électrique s'écrit : $\mathbf{E}_h = \tilde{\mathbf{E}}_h + \sum_{l=1}^{N_{cr}} \lambda_h^l \mathbf{x}_l^P$. Nous avons montré dans le paragraphe 4.4.4 comment calculer les λ_h^l . Nous décrivons ici la seconde étape de la discrétisation de la λ -approche : le calcul de $\tilde{\mathbf{E}}_h$. Pour cela, il faut d'abord discrétiser la condition aux limites tangentielles.

Soit $\mathbf{e}_{\lambda_h}$ un relèvement régulier de $e - \sum_{l=1}^{N_{cr}} \lambda_h^l \mathbf{x}_l^P$. Soit $\mathbf{e}_{\lambda, h} = \Pi_k(\mathbf{e}_{\lambda_h})$, l'interpolation de $\mathbf{e}_{\lambda_h}$ que l'on décompose dans B comme suit :

$$\mathbf{e}_{\lambda, h} = \sum_{j \in I_\omega \cup I_c} \sum_{\beta=1}^2 \varepsilon_{j, \beta} v_j \mathbf{e}_\beta + \sum_{j \in I_a} (\varepsilon_{j, \tau} v_j \boldsymbol{\tau}_j + \varepsilon_{j, \nu} v_j \boldsymbol{\nu}_j), \quad \text{avec :}$$

- $\forall j \in I_\omega \cup I_c, \forall \beta \in \{1, 2\}, \varepsilon_{j, \beta} = \mathbf{e}_{\lambda, h}(M_j) \cdot \mathbf{e}_\beta$;

- $\forall j \in I_a, \varepsilon_{j, \tau} = \mathbf{e}_{\lambda, h}(M_j) \cdot \boldsymbol{\tau}_j, \varepsilon_{j, \nu} = \mathbf{e}_{\lambda, h}(M_j) \cdot \boldsymbol{\nu}_j$.

Pour $j \in I_c, \beta = 1$ ou 2 , les $\varepsilon_{j, \beta}$ sont explicitement connus. De même, pour $j \in I_a$, les $\varepsilon_{j, \tau}$ sont explicitement connus.

Soit $\tilde{\mathbf{E}}_h^0$ l'approximation de $\tilde{\mathbf{E}}^0$, décomposée dans la base de $\mathbf{X}_{\mathbf{E}, k}^{0, R}$:

$$\tilde{\mathbf{E}}_h^0 = \sum_{j \in I_\omega} \sum_{\beta=1}^2 \tilde{\mathbf{E}}_{h, \beta}^0(M_j) v_j \mathbf{e}_\beta + \sum_{j \in I_a} \tilde{\mathbf{E}}_{h, \nu}^0(M_j) v_j \boldsymbol{\nu}_j.$$

Considérons $\tilde{\mathbf{E}}_h$ l'approximation de $\tilde{\mathbf{E}}$ dans $\mathbf{X}_{\mathbf{E},k}$.

On a : $\tilde{\mathbf{E}} = \tilde{\mathbf{E}}^0 + \mathbf{e}_\lambda$, et $\tilde{\mathbf{E}}_h = \tilde{\mathbf{E}}_h^0 + \mathbf{e}_{\lambda,h}$. On en déduit :

$$\begin{aligned}\tilde{\mathbf{E}}_h &= \sum_{j \in \mathbf{I}_\omega} \sum_{\beta=1}^2 \left(\tilde{\mathbf{E}}_{h,\beta}^0(M_j) + \varepsilon_{j,\beta} \right) v_j \mathbf{e}_\beta \\ &\quad + \sum_{j \in \mathbf{I}_a} \left(\varepsilon_{j,\tau} v_j \boldsymbol{\tau}_j + (\tilde{\mathbf{E}}_{h,\nu}^0(M_j) + \varepsilon_{j,\nu}) v_j \boldsymbol{\nu}_j \right) \\ &\quad + \sum_{j \in \mathbf{I}_c} \sum_{\beta=1}^2 \varepsilon_{j,\beta} v_j \mathbf{e}_\beta.\end{aligned}$$

Par identification, on trouve :

- $\forall j \in \mathbf{I}_\omega, \forall \beta \in \{1, 2\}, \tilde{\mathbf{E}}_{h,\beta}(M_j) = \tilde{\mathbf{E}}_{h,\beta}^0(M_j) + \varepsilon_{j,\beta},$
- $\forall j \in \mathbf{I}_a, \tilde{\mathbf{E}}_{h,\tau}(M_j) := \tilde{\mathbf{E}}_h(M_j) \cdot \boldsymbol{\tau} = \varepsilon_{j,\tau},$
- $\forall j \in \mathbf{I}_a, \tilde{\mathbf{E}}_{h,\nu}(M_j) := \tilde{\mathbf{E}}_h(M_j) \cdot \boldsymbol{\nu} = \tilde{\mathbf{E}}_{h,\nu}^0(M_j) + \varepsilon_{j,\nu},$
- $\forall j \in \mathbf{I}_c, \forall \beta \in \{1, 2\}, \tilde{\mathbf{E}}_{h,\beta}(M_j) = \varepsilon_{j,\beta}.$

En décomposant $\tilde{\mathbf{E}}_h$ dans la base B , on obtient donc :

$$\begin{aligned}\tilde{\mathbf{E}}_h &= \sum_{j \in \mathbf{I}_\omega} \sum_{\beta=1}^2 \tilde{\mathbf{E}}_{h,\beta}(M_j) v_j \mathbf{e}_\beta + \sum_{j \in \mathbf{I}_a} \tilde{\mathbf{E}}_{h,\nu}(M_j) v_j \boldsymbol{\nu}_j \\ &\quad + \sum_{j \in \mathbf{I}_c} \sum_{\beta=1}^2 \varepsilon_{j,\beta} + \sum_{j \in \mathbf{I}_a} \varepsilon_{j,\tau} v_j \boldsymbol{\tau}_j.\end{aligned}$$

Soit $\tilde{\underline{\mathbf{E}}} = (\tilde{\underline{\mathbf{E}}}_\omega, \tilde{\underline{\mathbf{E}}}_a, \tilde{\underline{\mathbf{E}}}_c)^T$ le vecteur de $\mathbb{R}^{2N}$ associé dont les composantes sont égales aux coordonnées de $\tilde{\mathbf{E}}_h$ dans B .

Comment construit-on le relèvement discret ? On peut par exemple considérer le relèvement discret suivant :

$$\underline{\varepsilon} = (\underline{\mathbf{Z}}_\omega, \underline{\varepsilon}_\tau, \underline{\varepsilon}_c)^T, \text{ où :}$$

- $\underline{\mathbf{Z}}_\omega$ est le vecteur nul de $(\mathbb{R}^2)^{N_\omega}$;
- $\underline{\varepsilon}_\tau = (\varepsilon_{N_\omega+1,\tau}, 0, \dots, \varepsilon_{N-N_c,\tau}, 0)^T \in (\mathbb{R}^2)^{N_a}$, est le vecteur contenant les valeurs des composantes tangentielle de $\mathbf{e}_{\lambda,h}$ sur $\partial\omega \setminus \mathcal{C}$;
- $\underline{\varepsilon}_c = (\varepsilon_{N-N_c+1,1}, \varepsilon_{N-N_c+1,2}, \dots, \varepsilon_{N,1}, \varepsilon_{N,2})^T \in (\mathbb{R}^2)^{N_c}$ est le vecteur contenant les valeurs des composantes de $\mathbf{e}_\lambda$ aux coins du maillage.

En d'autres termes, on conserve exactement les valeurs tangentielles aux points de la discrétisation du bord, et on impose zéro sur tous les autres degrés de liberté. Le calcul de ce vecteur sera détaillé dans le paragraphe 4.8.2.

On a : $\tilde{\underline{\mathbf{E}}}_c = \underline{\varepsilon}_c$, et $\tilde{\underline{\mathbf{E}}}_a = \tilde{\underline{\mathbf{E}}}_\nu + \tilde{\underline{\mathbf{E}}}_\tau$, en posant :

$$\tilde{\underline{\mathbf{E}}}_\nu = (0, \tilde{\underline{\mathbf{E}}}_{N_\omega+1, \nu}, \dots, 0, \tilde{\underline{\mathbf{E}}}_{N-N_c, \nu})^T \in (\mathbb{R}^2)^{N_a} \text{ et } \tilde{\underline{\mathbf{E}}}_\tau = \underline{\varepsilon}_\tau.$$

$\tilde{\underline{\mathbf{E}}}_\nu$ est le vecteur contenant les valeurs des composantes normales de $\tilde{\underline{\mathbf{E}}}_h$ sur $\partial\omega \setminus \mathcal{C}$.

Si on revient à la formulation variationnelle sur $\tilde{\underline{\mathbf{E}}}^0$, on peut déterminer $\tilde{\underline{\mathbf{E}}}_\omega$ et $\tilde{\underline{\mathbf{E}}}_\nu$, en réécrivant le problème variationnel (2.43) discrétisé dans $\mathbf{X}_{\mathbf{E},k}^{0,R}$ comme suit :

Trouver $\tilde{\underline{\mathbf{E}}}_h^0 \in \mathbf{X}_{\mathbf{E},k}^{0,R}$ tel que :

$$\begin{aligned} \forall i \in I_\omega, \forall \alpha \in \{1, 2\}, \quad \mathcal{A}_0(\tilde{\underline{\mathbf{E}}}_h^0, v_i \mathbf{e}_\alpha) &= \mathcal{L}_0(v_i \mathbf{e}_\alpha) - \mathcal{A}_0(\mathbf{e}_{\lambda,h}, v_i \mathbf{e}_\alpha), \\ \forall i \in I_a, \quad \mathcal{A}_0(\tilde{\underline{\mathbf{E}}}_h^0, v_i \boldsymbol{\nu}_i) &= \mathcal{L}_0(v_i \boldsymbol{\nu}_i) - \mathcal{A}_0(\mathbf{e}_{\lambda,h}, v_i \boldsymbol{\nu}_i). \end{aligned}$$

En regroupant $\tilde{\underline{\mathbf{E}}}_h^0$ et $\mathbf{e}_{\lambda,h}$, on peut réécrire l'ensemble de ces équations sous la forme :

Trouver $\tilde{\underline{\mathbf{E}}}_h \in \mathbf{X}_k$ tel que :

$$\begin{aligned} \forall i \in I_\omega, \forall \alpha \in \{1, 2\}, \quad \mathcal{A}_0(\tilde{\underline{\mathbf{E}}}_h, v_i \mathbf{e}_\alpha) &= \mathcal{L}_0(v_i \mathbf{e}_\alpha), \\ \forall i \in I_a, \quad \mathcal{A}_0(\tilde{\underline{\mathbf{E}}}_h, v_i \boldsymbol{\nu}_i) &= \mathcal{L}_0(v_i \boldsymbol{\nu}_i), \end{aligned} \quad (4.41)$$

sachant que :

$$\begin{aligned} \forall i \in I_a, \quad \tilde{\underline{\mathbf{E}}}_h(M_i) \cdot \boldsymbol{\tau}_i &= \mathbf{e}_{\lambda_h}(M_i) \cdot \boldsymbol{\tau}_i, \\ \forall i \in I_c, \quad \tilde{\underline{\mathbf{E}}}_h(M_i) &= \mathbf{e}_{\lambda_h}(M_i). \end{aligned} \quad (4.42)$$

Soit $\mathbb{A}_{\nu,a} \in (\mathbb{R}^{2 \times 2})^{N_a \times N_c}$ la matrice $\mathbb{A}_{a,a}$ dont on a éliminé les lignes dépendant de $\boldsymbol{\tau}$: $\forall i, j \in I_a$,

$$\mathbb{A}_{\nu,a}^{i,j} = \begin{pmatrix} 0 & 0 \\ \mathcal{A}_0(v_j \boldsymbol{\tau}_j, v_i \boldsymbol{\nu}_i) & \mathcal{A}_0(v_j \boldsymbol{\nu}_j, v_i \boldsymbol{\nu}_i) \end{pmatrix}.$$

Soit $\mathbb{A}_{\nu,c} \in (\mathbb{R}^{2 \times 2})^{N_a \times N_c}$, la matrice $\mathbb{A}_{a,c}$ dont on a éliminé les lignes dépendant de $\boldsymbol{\tau}$: $\forall i \in I_a, \forall j \in I_c$,

$$\mathbb{A}_{\nu,c}^{i,j} = \begin{pmatrix} 0 & 0 \\ \mathcal{A}_0(v_j \mathbf{e}_1, v_i \boldsymbol{\nu}_i) & \mathcal{A}_0(v_j \mathbf{e}_2, v_i \boldsymbol{\nu}_i) \end{pmatrix}.$$

Les lignes des équations (4.41) correspondent aux produit matriciels suivants :

$$\mathbb{A}_{\omega,\omega} \tilde{\underline{\mathbf{E}}}_\omega + \mathbb{A}_{\omega,a} \tilde{\underline{\mathbf{E}}}_a + \mathbb{A}_{\omega,c} \tilde{\underline{\mathbf{E}}}_c = \underline{\mathbf{L}}_\omega,$$

$$\mathbb{A}_{\nu,\omega} \tilde{\underline{\mathbf{E}}}_\omega + \mathbb{A}_{\nu,a} \tilde{\underline{\mathbf{E}}}_a + \mathbb{A}_{\nu,c} \tilde{\underline{\mathbf{E}}}_c = \underline{\mathbf{L}}_\nu,$$

où $\underline{\mathbf{L}}_\omega \in (\mathbb{R}^2)^{N_\omega}$ la restriction de $\underline{\mathbf{L}}$ aux $i \in I_\omega$; et $\underline{\mathbf{L}}_\nu \in (\mathbb{R}^2)^{N_a}$ est le vecteur tel que :

$$\underline{\mathbf{L}}_\nu = (0, \mathcal{L}_0(v_{N_\omega+1} \boldsymbol{\nu}_{N_\omega+1}), \dots, 0, \mathcal{L}_0(v_{N-N_c} \boldsymbol{\nu}_{N-N_c})).$$

On obtient que $\tilde{\underline{\mathbf{E}}}$ est solution du système :

$$\mathbb{A}_{\omega,\omega} \tilde{\underline{\mathbf{E}}}_\omega + \mathbb{A}_{\omega,a} \tilde{\underline{\mathbf{E}}}_\nu = \underline{\mathbf{L}}_\omega - \mathbb{A}_{\omega,a} \underline{\varepsilon}_\tau - \mathbb{A}_{\omega,c} \underline{\varepsilon}_c,$$

$$\mathbb{A}_{\nu,\omega} \tilde{\underline{\mathbf{E}}}_\omega + \mathbb{A}_{\nu,a} \tilde{\underline{\mathbf{E}}}_\nu = \underline{\mathbf{L}}_\nu - \mathbb{A}_{\nu,a} \underline{\varepsilon}_\tau - \mathbb{A}_{\nu,c} \underline{\varepsilon}_c,$$

$$\tilde{\underline{\mathbf{E}}}_\tau = \underline{\varepsilon}_\tau,$$

$$\tilde{\underline{\mathbf{E}}}_c = \underline{\varepsilon}_c.$$

On remarque que :

$$\mathbb{A}_{\omega,a} \tilde{\underline{\mathbf{E}}}_{\nu} = \mathbb{A}_{\omega,\nu} \tilde{\underline{\mathbf{E}}}_{\nu} = \mathbb{A}_{\omega,\nu} \tilde{\underline{\mathbf{E}}}_a, \text{ et } \mathbb{A}_{\nu,a} \tilde{\underline{\mathbf{E}}}_{\nu} = \mathbb{A}_{\nu,\nu} \tilde{\underline{\mathbf{E}}}_{\nu} = \mathbb{A}_{\nu,\nu} \tilde{\underline{\mathbf{E}}}_a - \tilde{\underline{\mathbf{E}}}_{\tau}.$$

D'où, en passant $-\mathbb{A}_{\nu,\nu} \tilde{\underline{\mathbf{E}}}_{\tau} = -\underline{\varepsilon}_{\tau}$ au second membre dans le seconde équation, $\tilde{\underline{\mathbf{E}}}$ satisfait :

$$\begin{aligned} \mathbb{A}_{\omega,\omega} \tilde{\underline{\mathbf{E}}}_{\omega} + \mathbb{A}_{\omega,\nu} \tilde{\underline{\mathbf{E}}}_a &= \underline{\mathbf{L}}_{\omega} - \mathbb{A}_{\omega,a} \underline{\varepsilon}_{\tau} - \mathbb{A}_{\omega,c} \underline{\varepsilon}_c, \\ \mathbb{A}_{\nu,\omega} \tilde{\underline{\mathbf{E}}}_{\omega} + \mathbb{A}_{\nu,\nu} \tilde{\underline{\mathbf{E}}}_a &= \underline{\mathbf{L}}_{\nu} - \mathbb{A}_{\nu,a} \underline{\varepsilon}_{\tau} + \underline{\varepsilon}_{\tau} - \mathbb{A}_{\nu,c} \underline{\varepsilon}_c, \\ \tilde{\underline{\mathbf{E}}}_c &= \underline{\varepsilon}_c. \end{aligned}$$

Soit $\mathbb{A}_r \in (\mathbb{R}^{2 \times 2})^{N \times N}$ la matrice définie ainsi :

$$\mathbb{A}_r = \begin{pmatrix} 0 & \mathbb{A}_{\omega,a} & \mathbb{A}_{\omega,c} \\ 0 & \mathbb{A}_{\nu,a} & \mathbb{A}_{\nu,c} \\ 0 & 0 & -\mathbb{I}_{c,c} \end{pmatrix}.$$

Posons $\underline{\mathbf{L}}_0 = (\underline{\mathbf{L}}_{\omega}, \underline{\mathbf{L}}_{\nu}, \underline{\mathbf{Z}}_c)^T \in (\mathbb{R}^2)^N$, où $\underline{\mathbf{Z}}_c$ est le vecteur nul de $(\mathbb{R}^2)^{N_c}$.
 $\tilde{\underline{\mathbf{E}}}$ est alors solution du système matriciel :

$$\boxed{\mathbb{A}_0 \tilde{\underline{\mathbf{E}}} = \underline{\mathbf{L}}_0 - \mathbb{A}_r \underline{\varepsilon}.} \quad (4.43)$$

Nous ne détaillerons pas le calcul de $\underline{\mathbf{L}}_0$ car il est similaire à celui de $\underline{\mathbf{L}}$, effectué dans le paragraphe (4.6.3).

Une fois que $\tilde{\underline{\mathbf{E}}}$ est calculé, on peut écrire l'approximation calculée sur les arêtes dans la base canonique $(\mathbf{e}_1, \mathbf{e}_2)$. Pour cela, à chaque point M_i , $i \in I_a$, on associe $\mathbb{O}_i \in \mathbb{R}^{2 \times 2}$, la matrice de changement de base telle que :

$$\mathbb{O}_i = \begin{pmatrix} \nu_2(M_i) & \nu_1(M_i) \\ -\nu_1(M_i) & \nu_2(M_i) \end{pmatrix}.$$

Soit $\mathbb{O} \in (\mathbb{R}^{2 \times 2})^{N_a \times N_a}$ le matrice diagonale par blocs telle que :

$$\mathbb{O} = \begin{pmatrix} \mathbb{O}_{N_{\omega}+1} & 0 & 0 \\ 0 & \ddots & 0 \\ 0 & 0 & \mathbb{O}_{N-N_c} \end{pmatrix}.$$

Le vecteur correspondant aux composantes dans la base canonique $(\mathbf{e}_1, \mathbf{e}_2)$ s'écrit alors : $\mathbb{O} \tilde{\underline{\mathbf{E}}}_a$.

Remarque 4.10 *En pratique, il est plus simple de créer la matrice et le second membre dans la base cartésienne complète, et de procéder à une élimination a posteriori, en faisant un changement de base avant l'élimination, et un autre après pour revenir au système cartésien.*

Soit $l \in I_a$. On extrait la double ligne $\underline{\mathbf{L}} \in (\mathbb{R}^2)^N$ correspondant à $i = l$, $\alpha = 1$ et 2. En multipliant à gauche cette double ligne par $\mathbb{O}_l$, on l'exprime dans la base locale $(\boldsymbol{\tau}_l, \boldsymbol{\nu}_l)$. On procède à l'élimination de la première ligne de $\mathbb{O}_l \underline{\mathbf{L}}$, puis on multiplie à gauche par $\mathbb{O}_l^T$ pour revenir à la base canonique.

Ensuite, on extrait la double colonne $\underline{\mathbb{C}} \in (\mathbb{R}^2)^N$ correspondant à $j = l$, $\alpha = 1$ et 2. En multipliant à droite cette double colonne par $\mathbb{O}_l^T$, on l'exprime dans la base locale $(\boldsymbol{\tau}_l, \boldsymbol{\nu}_l)$. On procède à l'élimination de la première colonne de $\underline{\mathbb{C}}\mathbb{O}_l^T$. Afin de rendre la matrice inversible, on impose que $(\underline{\mathbb{C}}\mathbb{O}_l^T)^{2,2} = 1$, puis on multiplie à droite par $\mathbb{O}_l$ pour revenir à la base canonique.

4.8.2 Relèvement de la CL

Nous venons de voir qu'il n'est pas nécessaire de calculer un relèvement explicite de $\mathbf{e}_{\lambda,h}$ pour calculer $\mathbf{E}_h$. Nous allons détailler le calcul des composantes de $\underline{\varepsilon}$. Tout d'abord, décomposons $\underline{\varepsilon}$ en une partie dépendant de e et une partie dépendant des $\lambda_h^l, \mathbf{x}_l^P$:

$$\underline{\varepsilon} = \underline{\mathbf{e}} - \sum_{l=1}^{N_{cr}} \lambda_h^l \underline{\mathbf{x}}_{\partial\omega}^l,$$

où : $\underline{\mathbf{x}}_{\partial\omega}^l = (\underline{\mathbf{z}}_\omega, \underline{\mathbf{x}}_\tau^l, \underline{\mathbf{x}}_c^l)^T$, et $\underline{\mathbf{e}} = (\underline{\mathbf{z}}_\omega, \underline{\mathbf{e}}_\tau, \underline{\mathbf{e}}_c)^T$, avec :

- $\underline{\mathbf{x}}_\tau^l = (\mathbf{x}_l^P(M_{N_\omega+1}) \cdot \boldsymbol{\tau}_{N_\omega+1}, 0, \dots, \mathbf{x}_l^P(M_{N-N_c}) \cdot \boldsymbol{\tau}_{N-N_c}, 0)^T$;
- $\underline{\mathbf{x}}_c^l = (\mathbf{x}_l^P(M_{N-N_c+1}) \cdot \mathbf{e}_1, \mathbf{x}_l^P(M_{N-N_c+1}) \cdot \mathbf{e}_2, \dots, \mathbf{x}_l^P(M_N) \cdot \mathbf{e}_1, \mathbf{x}_l^P(M_N) \cdot \mathbf{e}_2)^T$;
- $\underline{\mathbf{e}}_\tau = -(e(M_{N_\omega+1}), 0, \dots, e(M_{N-N_c}), 0)^T$.

Physiquement, la donnée de $\mathbf{E} \cdot \boldsymbol{\tau}_{|\partial\omega}$ est une valeur associée aux composantes du bord, et non aux sommets. Cependant, dans notre approximation, nous associons la valeur $\mathbf{E}_h \cdot \boldsymbol{\tau}_{|\partial\omega}$ aux sommets du bord car nous utilisons des éléments finis nodaux. Nous allons consacrer le reste du paragraphe à la détermination de $\underline{\mathbf{e}}_c$.

Lorsque $\mathbf{E} \in \mathbf{H}^1(\omega)$, $\mathbf{E} \cdot \boldsymbol{\tau}_{|A_k} \in H^{1/2}(A_k)$, $\forall k \in \{1, \dots, K\}$, mais on n'a pas $\mathbf{E} \cdot \boldsymbol{\tau}_{|\partial\omega} \in H^{1/2}(\partial\omega)$. La fonction e n'est donc pas nécessairement continue entre deux arêtes successives : il n'y a pas unicité de la valeur de e aux points $(M_i)_{i \in I_c}$. Or, comme nous approchons la champ électrique par des éléments finis nodaux, nous devons attribuer une valeur au champ électrique approché aux coins du maillage pour résoudre cette difficulté. Soit M le coin du maillage égal à l'intersection $\partial A_1 \cap \partial A_2$. La valeur attribuée à $\mathbf{E}_h$ en M est donnée par la résolution du système :

$$\begin{cases} \mathbf{E}_h(M) \cdot \boldsymbol{\tau}_{|A_1} &= \lim_{M_1 \rightarrow M, M_1 \in A_1} e(M_1), \\ \mathbf{E}_h(M) \cdot \boldsymbol{\tau}_{|A_2} &= \lim_{M_2 \rightarrow M, M_2 \in A_2} e(M_2). \end{cases} \quad (4.44)$$

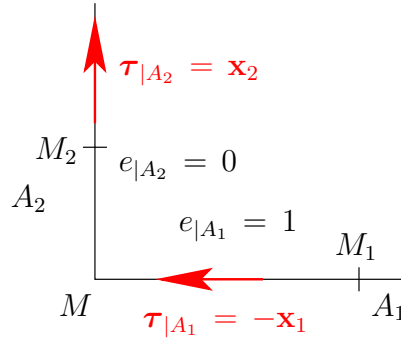
Ci-dessus, on a supposé pour simplifier que la donnée e est continue par arête, ce qui est vrai si pour tout k , il existe $\varepsilon_k > 0$ tel que $e|_{A_k} \in H^{1/2+\varepsilon}(A_k)$, car pour tout $\varepsilon > 0$, $H^{1/2+\varepsilon}(A_k) \subset C^0(\overline{A_k})$. Considérons l'exemple suivant (figure 4.2), avec $e|_{A_1} = 1$ et $e|_{A_2} = 0$.

On a alors :

$$\begin{cases} E_2(M) &= \mathbf{E}_h(M) \cdot \boldsymbol{\tau}_{|A_2} = 0, \\ -E_1(M) &= \mathbf{E}_h(M) \cdot \boldsymbol{\tau}_{|A_1} = 1. \end{cases}$$

D'où : $\mathbf{E}_h(M) = (-1, 0)^T$. Nous sommes ainsi en mesure de déterminer $\underline{\mathbf{e}}_c$. Il s'agit tout simplement de résoudre le système (4.44) pour chaque coin $(M_i)_{i \in I_c}$ du maillage, afin de déterminer $\mathbf{E}_{h,1}(M_i)$ et $\mathbf{E}_{h,2}(M_i)$. On a alors :

$$\underline{\mathbf{e}}_c = (\mathbf{E}_{h,1}(M_{N-N_c+1}), \mathbf{E}_{h,2}(M_{N-N_c+1}), \dots, \mathbf{E}_{h,1}(M_N), \mathbf{E}_{h,2}(M_N))^T.$$

FIG. 4.2 – Valeur de $\mathbf{E}_h$ en un sommet du domaine.

4.9 Complément singulier orthogonal : discrétisation

Pour la méthode du complément singulier orthogonale, l'approximation du champ électrique s'écrit : $\mathbf{E}_h = \hat{\mathbf{E}}_h + \sum_{l=1}^{N_{cr}} c_h^l \mathbf{x}_l^S$. Les inconnues sont : $\hat{\mathbf{E}}_h$ ainsi que les N_{cr} coefficients c_h^l , obtenus à l'aide des coefficients $\beta_h^{l,m}$ et λ_h^l . Dans le paragraphe 4.4.3, nous avons montré comment calculer les $\beta_h^{l,m}$, et dans le paragraphe 4.4.4, nous avons détaillé le calcul des λ_h^l .

Considérons $\hat{\mathbf{E}}_h = \hat{\mathbf{E}}_h^0 + \mathbf{e}_h$, avec $\mathbf{e}_h = \Pi_k(\mathbf{e})$, son approximation. Nous allons nous ramener au calcul direct de $\hat{\mathbf{E}}_h$, de la même façon qu'on se ramène au calcul de $\tilde{\mathbf{E}}_h$ dans la λ -approche.

4.9.1 Champ électrique : partie régulière

Soit $\mathbb{X}_h$ est la matrice approchée de $\mathbb{X} \in \mathbb{R}^{N_{cr} \times N_{cr}}$, construite à l'aide des $\beta_h^{l,m}$, et $\boldsymbol{\lambda}_h \in \mathbb{R}^{N_{cr}}$, contenant les λ_h^l , est l'approximation du vecteur $\boldsymbol{\lambda}$. La formulation variationnelle (2.57)-(2.58) discrétisée dans $\mathbf{X}_{\mathbf{E},k}^{0,R}$ s'écrit :

Trouver $\hat{\mathbf{E}}_h^0 \in \mathbf{X}_{\mathbf{E},k}^{0,R}$ et $\mathbf{c}_h \in \mathbb{R}^{N_{cr}}$ tels que :

$$\begin{aligned} \forall i \in I_\omega, \forall \alpha \in \{1, 2\}, \quad \mathcal{A}_0(\hat{\mathbf{E}}_h^0, \mathbf{v}_{i,\alpha}) &= \mathcal{L}_0(\mathbf{v}_{i,\alpha}) - \mathcal{A}_0(\mathbf{e}_h, \mathbf{v}_{i,\alpha}); \\ \forall i \in I_a, \quad \mathcal{A}_0(\hat{\mathbf{E}}_h^0, v_i \boldsymbol{\nu}_i) &= \mathcal{L}_0(v_i \boldsymbol{\nu}_i) - \mathcal{A}_0(\mathbf{e}_h, v_i \boldsymbol{\nu}_i), \\ \mathbb{X}_h \mathbf{c}_h &= \boldsymbol{\lambda}_h, \end{aligned} \tag{4.45}$$

Cela revient au problème suivant :

Trouver $\hat{\mathbf{E}}_h \in \mathbf{X}_k$ et $\mathbf{c}_h \in \mathbb{R}^{N_{cr}}$ tels que :

$$\begin{aligned} \forall i \in I_\omega, \forall \alpha \in \{1, 2\}, \quad \mathcal{A}_0(\hat{\mathbf{E}}_h, \mathbf{v}_{i,\alpha}) &= \mathcal{L}_0(\mathbf{v}_{i,\alpha}); \\ \forall i \in I_a, \quad \hat{\mathbf{E}}_h(M_i) \cdot \boldsymbol{\tau}_i &= \mathbf{e}_h(M_i) \cdot \boldsymbol{\tau}_i, \\ \forall i \in I_c, \quad \hat{\mathbf{E}}_h(M_i) &= \mathbf{e}_h(M_i), \\ \mathbb{X}_h \mathbf{c}_h &= \boldsymbol{\lambda}_h. \end{aligned} \tag{4.46}$$

Soit $\hat{\underline{\mathbf{E}}} \in (\mathbb{R}^2)^N$ le vecteur associé à la décomposition de $\hat{\mathbf{E}}_h$ dans la base B . Décomposons $\hat{\underline{\mathbf{E}}}$ de la façon suivante : $\hat{\underline{\mathbf{E}}} = (\hat{\underline{\mathbf{E}}}_\omega, \hat{\underline{\mathbf{E}}}_a, \hat{\underline{\mathbf{E}}}_c)^T$, où :

- $\widehat{\underline{\mathbf{E}}}_a = \widehat{\underline{\mathbf{E}}}_\nu + \underline{\mathbf{e}}_\tau$, avec : $\widehat{\underline{\mathbf{E}}}_\nu = (0, \widehat{\mathbf{E}}_h(M_{N_\omega+1}) \cdot \nu_{N_\omega+1}, \dots, 0, \widehat{\mathbf{E}}_h(M_{N-N_c}) \cdot \nu_{N-N_c})^T$,

- $\widehat{\underline{\mathbf{E}}}_c = \underline{\mathbf{e}}_c$,

- $\underline{\mathbf{e}}_\tau$ et $\underline{\mathbf{e}}_c$ sont déterminés comme $\underline{\varepsilon}_\tau$ et $\underline{\varepsilon}_c$ dans le paragraphe 4.8.1.

Soit $\mathbb{A}' = \begin{pmatrix} \mathbb{A}_0 & 0 \\ 0 & \pi \mathbb{X}_h \end{pmatrix} \in \mathbb{R}^{(2N+N_{cr}) \times (2N+N_{cr})}$ et $\underline{\mathbf{L}}' = \begin{pmatrix} \underline{\mathbf{L}}_0 - \mathbb{A}_r \underline{\mathbf{e}} \\ \pi \lambda_h \end{pmatrix} \in \mathbb{R}^{2N+N_{cr}}$. Posons $\underline{\mathbf{E}}' = \begin{pmatrix} \widehat{\underline{\mathbf{E}}} \\ \mathbf{c}_h \end{pmatrix} \in \mathbb{R}^{2N+N_{cr}}$. Les équations (4.46) se mettent sous la forme suivante :

$$\boxed{\mathbb{A}' \widehat{\underline{\mathbf{E}}} = \underline{\mathbf{L}}'} \quad (4.47)$$

$\mathbb{A}_0, \mathbb{A}_r, \underline{\mathbf{L}}_0$ ont été construits dans le paragraphe 4.8.1. Pour obtenir $\widehat{\mathbf{E}}_h$, il ne nous reste plus qu'à calculer les $\mathbf{x}_{l,h}^S$. Pour la méthode des potentiels, il faut évaluer les $\tilde{\varphi}_{D,l}$ et $\tilde{\varphi}_{N,l}$. Pour calculer les $\mathbf{x}_{l,h}^S$ par la λ -approche, il suffit de procéder comme il est expliqué précédemment, en prenant simplement : $f_h = s_{N,l,h}$ et $s_{D,l,h}$.

4.9.2 Champ électrique : partie singulière

D'après le théorème 2.43, on sait que : $\mathbf{x}_l^S = -\mathbf{grad} \varphi_{D,l} + \mathbf{rot} \varphi_{N,l}$, ce qui s'écrit aussi de la façon suivante :

$$\mathbf{x}_l^S = -\mathbf{grad} \tilde{\varphi}_{D,l} + \mathbf{rot} \tilde{\varphi}_{N,l} + \sum_{m=1}^{N_{cr}} \beta_h^{l,m} \mathbf{x}_m^P.$$

On peut donc approcher les $\mathbf{x}_l^S$ par dérivation des potentiels de la façon suivante :

$$\mathbf{x}_{l,h}^S = -\mathbf{grad} (\tilde{\varphi}_{D,l})_h + \mathbf{rot} (\tilde{\varphi}_{N,l})_h + \sum_{l=1}^{N_{cr}} \beta_h^{l,m} \Pi_k(\mathbf{x}_m^P).$$

Comme pour le calcul de $\mathbf{E}_h$ par les potentiels ϕ_D et ϕ_N (paragraphe 4.3.3), on utilise une projection $P_{k-1}-P_k$ pour obtenir une approximation continue de $\mathbf{x}_{l,h}^S$.

$\tilde{\mathbf{x}}_l$ est alors solution de :

$$\boxed{\tilde{\mathbf{x}}_l = -\mathbb{G}_D \tilde{\phi}_{D,l} + \mathbb{R}_N \tilde{\phi}_{N,l}} \quad (4.48)$$

On peut d'autre part calculer les $\mathbf{x}_l^S$ par la λ -approche : soit $\tilde{\mathbf{x}}_l \in (\mathbb{R}^2)^N$ représentant $\tilde{\mathbf{x}}_{l,h}$ dans la base B . D'après (2.64) et le paragraphe 4.8.2, $\tilde{\mathbf{x}}_l$ est solution de :

$$\boxed{\mathbb{A}_0 \tilde{\mathbf{x}}_l = -\mathbb{A}_r \left(\sum_{m=1}^{N_{cr}} \beta_h^{l,m} \mathbf{x}_{\partial\omega}^m \right)} \quad (4.49)$$

4.9.3 Conclusion

Théorème 4.11 *La λ -approche et la MCSO discrétisées donnent tout calcul fait le même résultat :*

$$\hat{\underline{\mathbf{E}}} + \sum_{l=1}^{N_{cr}} c_{l,h} \tilde{\underline{\mathbf{x}}}_l = \tilde{\underline{\mathbf{E}}}.$$

DÉMONSTRATION. L'addition de la première équation de (4.47), qui nous donne : $\mathbb{A}_0 \underline{\mathbf{E}} = \underline{\mathbf{L}}_0 - \mathbb{A}_r \underline{\mathbf{e}}$ et de (4.49) donne :

$$\mathbb{A}_0 \left(\hat{\underline{\mathbf{E}}} + \sum_{l=1}^{N_{cr}} c_{l,h} \tilde{\underline{\mathbf{x}}}_l \right) = \underline{\mathbf{L}}_0 - \mathbb{A}_r \left(\underline{\mathbf{e}} + \sum_{l=1}^{N_{cr}} c_{l,h} \sum_{m=1}^{N_{cr}} \beta_h^{l,m} \underline{\mathbf{x}}_{\partial\omega}^m \right),$$

ce qui devient à l'aide de la seconde équation de (4.47) : $\mathbb{A}_0 \left(\hat{\underline{\mathbf{E}}} + \sum_{l=1}^{N_{cr}} c_{l,h} \tilde{\underline{\mathbf{x}}}_l \right) = \underline{\mathbf{L}}_0 - \mathbb{A}_r \underline{\underline{\mathbf{e}}}$. D'où le résultat. □

4.10 Régularisation à poids : discrétisation

$\mathbf{X}_{\mathbf{E},\gamma}^0 \cap \mathbf{H}^1(\omega)$ étant dense dans $\mathbf{X}_{\mathbf{E},\gamma}^0$, on peut considérer comme espace d'approximation de $\mathbf{X}_{\mathbf{E},\gamma}^0$ l'espace : $\mathbf{X}_{\mathbf{E},k}^{0,R}$, décrit dans le paragraphe 4.7.1. La formulation variationnelle (2.69) s'écrit : Trouver $\mathbf{E}_{\gamma,h}^0 \in \mathbf{X}_{\mathbf{E},k}^{0,R}$ tel que :

$$\begin{aligned} \forall i \in I_\omega, \forall \alpha \in \{1, 2\} \quad \mathcal{A}_\gamma(\mathbf{E}_{\gamma,h}^0, v_i \mathbf{e}_\alpha) &= \mathcal{L}_\gamma(v_i \mathbf{e}_\alpha) - \mathcal{A}_\gamma(\mathbf{e}, v_i \mathbf{e}_\alpha), \\ \forall i \in I_a, \quad \mathcal{A}_\gamma(\mathbf{E}_{\gamma,h}^0, v_i \boldsymbol{\nu}_i) &= \mathcal{L}_\gamma(v_i \boldsymbol{\nu}_i) - \mathcal{A}_\gamma(\mathbf{e}, v_i \boldsymbol{\nu}_i). \end{aligned} \tag{4.50}$$

Pour exprimer ces équations sous forme matricielle, on procède de la même façon que dans le paragraphe 4.7.1 par pseudo-élimination des degrés de liberté connus. Soit $\mathbb{A}_{0,\gamma}$ (*resp.* $\mathbb{A}_{r,\gamma}$) la matrice obtenue en remplaçant l'opérateur $\mathcal{A}_0$ par $\mathcal{A}_\gamma$ dans $\mathbb{A}_0$ (*resp.* $\mathbb{A}_r$). On se ramène au calcul direct de $\mathbf{E}_{\gamma,h} = \mathbf{E}_{\gamma,h}^0 + \mathbf{e}_h$, avec comme relèvement discret de la condition tangentielle au bord e le vecteur $\underline{\mathbf{e}}$. On construit le vecteur $\underline{\mathbf{E}}_\gamma$ de la même façon que le vecteur $\hat{\underline{\mathbf{E}}}$; et le vecteur $\underline{\mathbf{L}}_\gamma$ de la même façon que $\underline{\mathbf{L}}_0$, en remplaçant l'opérateur $\mathcal{L}_0$ par $\mathcal{L}_\gamma$. Il s'agit alors de résoudre le système linéaire :

$$\mathbb{A}_\gamma \underline{\mathbf{E}}_\gamma = \underline{\mathbf{L}}_\gamma + \mathbb{A}_{r,\gamma} \underline{\mathbf{e}}. \tag{4.51}$$

Il faut construire la matrice $\mathbb{D}iv_\gamma$, qui correspond au produit scalaire dans V_γ^0 des divergences des vecteurs de la base B_0 . Cette construction est similaire à la construction de $\mathbb{D}iv$, on utilise un schéma d'intégration numérique ad hoc.

4.11 Analyse d'erreur

Nous allons maintenant étudier la convergence des méthodes de calcul direct de $\mathbf{E}$, c'est à dire l'erreur que l'on commet en fonction du pas du maillage h . On appelle *taux de convergence* la paramètre τ tel que : $\|\mathbf{E} - \mathbf{E}_h\|_{\mathbf{X}} \leq C h^\tau$, où C est une constante.

4.11.1 Méthode avec CL naturelles : convergence

Nous n'avons pas encore obtenu de résultat quant à l'analyse d'erreur pour le calcul de $\mathbf{E}_h$ dans $\mathbf{X}_k$. Cependant, nous avons observé numériquement que cette erreur est gouvernée par l'erreur sur la condition aux limites. Comme nous utilisons des éléments finis continus, nous ne pouvons pas approcher correctement $\mathbf{E} \cdot \boldsymbol{\tau}|_{\partial\omega}$ au voisinage des coins. En effet, cette quantité dépend de $\mathbf{E} \cdot \boldsymbol{\nu}|_{\partial\omega}$, par continuité le long des arêtes. Or nous ne contrôlons pas la quantité $\mathbf{E} \cdot \boldsymbol{\nu}|_{\partial\omega}$, qui est très élevée au(x) voisinage(s) du (des) coin(s) rentrant(s).

4.11.2 Méthode du complément singulier : convergence

Nous avons établi que pour le cas statique, la λ -approche et la méthode complément singulier orthogonale sont identiques d'un point de vue problème continu (section 2.4), formulation variationnelle (proposition 2.47) et problème discret (théorème 4.11). Nous nous contentons donc d'étudier la convergence de la λ -approche.

L'analyse de la convergence de la méthode du complément singulier a été réalisée par C. Hazard et S. Lohrengel dans [68]. Elle utilise le lemme de Céa 4.1, et l'équivalence des normes de $\mathbf{H}^1$ et de $\mathbf{X}_E^0$ pour les vecteurs de $\mathbf{X}_E^{0,R}$ (théorème 2.25). La preuve de convergence que nous présentons ici est légèrement différente car nous n'utilisons pas de fonction de troncature, mais nous obtenons le même taux de convergence, de l'ordre $2\alpha - 1 - \epsilon$. Nous allons montrer le théorème suivant :

Théorème 4.12 *On suppose qu'il existe $\epsilon > 0$ tel que $2\alpha - 1 - \epsilon > 0$ et $f, g \in H^{2\alpha-1-\epsilon}(\omega)$, $e \in H^{2\alpha-1/2-\epsilon}(\partial\omega)$. Alors il existe une constante $C_\epsilon > 0$ qui ne dépend que ϵ telle que :*

$$\|\mathbf{E} - \mathbf{E}_h\|_{\mathbf{X}_E} \leq C_\epsilon h^{2\alpha-1-\epsilon}. \quad (4.52)$$

Pour prouver ce théorème, nous avons besoin de la proposition suivante :

Proposition 4.13 *On a l'estimation d'erreur suivante sur le calcul des λ^i :*

Il existe une constante C_λ ne dépendant que du domaine et des données telle que : $\forall i \in \{1, \dots, N_{cr}\} :$

$$|\lambda^i - \lambda_h^i| < C_\lambda h^{2\alpha}.$$

DÉMONSTRATION. La preuve de cette proposition est faite par P. Ciarlet, Jr. et J. He dans [41] pour obtenir un taux de $2\alpha - \epsilon$, et grâce aux résultats de M. Amara et M. Moussaoui dans [4], on obtient le taux optimal. □

Pour tout $i \in \{1, \dots, N_{cr}\}$, $\mathbf{x}_{i,h}^P$ désigne l'approximation numérique de $\mathbf{x}_i^P$ choisie pour la représentation. En pratique, on prend souvent : $\mathbf{x}_{i,h}^P = \Pi_k(\mathbf{x}_i^P)$. Décomposons $\mathbf{E}$ et $\mathbf{E}_h$ ainsi :

$$\mathbf{E} = \tilde{\mathbf{E}}^0 + \mathbf{e}_\lambda + \sum_{i=1}^{N_{cr}} \lambda^i \mathbf{x}_i^P, \quad \mathbf{E}_h = \tilde{\mathbf{E}}_h^0 + \mathbf{e}_{\lambda,h} + \sum_{i=1}^{N_{cr}} \lambda_h^i \mathbf{x}_{i,h}^P.$$

Posons $\mathbf{z}^0 = \tilde{\mathbf{E}}^0 - \tilde{\mathbf{E}}_h^0$, $\mathbf{z}_\lambda = \mathbf{e}_\lambda - \mathbf{e}_{\lambda,h}$ et $\mathbf{z}^P = \sum_{i=1}^{N_{cr}} (\lambda^i \mathbf{x}_i^P - \lambda_h^i \mathbf{x}_{i,h}^P)$. D'après l'inégalité de Cauchy-Schwarz, on a donc :

$$\|\mathbf{E} - \mathbf{E}_h\|_{\mathbf{X}_E}^2 \leq 3 (\|\mathbf{z}^0\|_{\mathbf{X}_E}^2 + \|\mathbf{z}_\lambda\|_{\mathbf{X}_E}^2 + \|\mathbf{z}^P\|_{\mathbf{X}_E}^2).$$

Pour montrer le théorème 4.12, il s'agit alors d'estimer les trois termes du second membre de l'inégalité ci-dessus.

Lemme 4.14 *On suppose qu'il existe $\epsilon > 0$ tel que $2\alpha - 1 - \epsilon > 0$ et $f, g \in H^{2\alpha-1-\epsilon}(\omega)$, $e \in H^{2\alpha-1/2-\epsilon}(\partial\omega)$. Alors il existe une constante $C'_\epsilon > 0$ qui ne dépend que ϵ telle que :*

$$\|\mathbf{z}^0\|_{\mathbf{X}_E^0} = \|\tilde{\mathbf{E}}^0 - \tilde{\mathbf{E}}_h^0\|_{\mathbf{X}_E^0} \leq C'_\epsilon h^{2\alpha-1-\epsilon} \|\tilde{\mathbf{E}}^0\|_{\mathbf{H}^{2\alpha-\epsilon}}. \quad (4.53)$$

DÉMONSTRATION. La preuve de cette proposition, que nous reprenons de [68] se fait en trois étapes :

- D'après le lemme de Céa 4.1, comme $\mathcal{A}_0$ est coercitive, il existe une constante $C_0 > 0$ telle que :

$$\|\tilde{\mathbf{E}}^0 - \tilde{\mathbf{E}}_h^0\|_{\mathbf{X}_E^0} \leq C_0 \inf_{\mathbf{F}_h \in \mathbf{X}_{E,k}^{0,R}} \|\tilde{\mathbf{E}}^0 - \mathbf{F}_h\|_{\mathbf{X}_E^0}, \quad (4.54)$$

- D'après le théorème 2.25, il existe une constante $C_1 > 0$ telle que :

$$\|\tilde{\mathbf{E}}^0 - \tilde{\mathbf{E}}_h^0\|_{\mathbf{X}_E^0} \leq C_0 C_1 \inf_{\mathbf{F}_h \in \mathbf{X}_{E,k}^{0,R}} \|\tilde{\mathbf{E}}^0 - \mathbf{F}_h\|_{\mathbf{H}^1}. \quad (4.55)$$

- Or, d'après le théorème 2.31, pour le nombre ϵ considéré, $\tilde{\mathbf{E}}^0 \in \mathbf{H}^{2\alpha-\epsilon}(\omega)$. Comme $2\alpha - \epsilon > 1$, l'analyse d'erreur standard des éléments finis s'applique [33, 66] :

$$\inf_{\mathbf{F}_h \in \mathbf{X}_{E,k}^{0,R}} \|\tilde{\mathbf{E}}^0 - \mathbf{F}_h\|_{\mathbf{H}^1} \leq C''_\epsilon h^{2\alpha-1-\epsilon} \|\tilde{\mathbf{E}}^0\|_{\mathbf{H}^{2\alpha-\epsilon}}.$$

En prenant, $C'_\epsilon = C_0 C_1 C''_\epsilon$, on obtient (4.53). □

Lemme 4.15 *Il existe une constante $C_{\partial\omega} > 0$ telle que :*

$$\|\mathbf{z}_\lambda\|_{\mathbf{X}_E} \leq C_{\partial\omega} h^{2\alpha-\epsilon}. \quad (4.56)$$

DÉMONSTRATION. Notons que pour $e \in H^{2\alpha-1/2-\epsilon}(\partial\omega)$ avec ϵ tel que $2\alpha - 1 - \epsilon > 0$, il existe un relèvement $\mathbf{e} \in \mathbf{H}^{2\alpha-1-\epsilon}(\omega)$ de e , et comme $\mathbf{H}^{2\alpha-1-\epsilon}(\omega) \subset C^0(\overline{\omega})$, $\Pi_k(\mathbf{e})$ est bien défini.

Pour tout $i \in \{1, \dots, N_{cr}\}$, on pose $\mathbf{x}_i^R \in \mathbf{H}^{1+s}(\omega)$, $s > 0$, un relèvement régulier de $\mathbf{x}_i^P \cdot \boldsymbol{\tau}|_{\partial\omega}$. Comme $\mathbf{H}^{1+s}(\omega) \subset C^0(\overline{\omega})$, $\Pi_k(\mathbf{x}_i^R)$ est bien défini.

On a : $\mathbf{e}_\lambda = \mathbf{e} - \sum_{i=1}^{N_{cr}} \lambda^i \mathbf{x}_i^R$, et $\mathbf{e}_{\lambda,h} = \Pi_k(\mathbf{e}) - \sum_{i=1}^{N_{cr}} \lambda_h^i \mathbf{x}_{i,h}^R$. Décomposons $\mathbf{z}_\lambda$ de la façon suivante :

$$\mathbf{z}_\lambda = \mathbf{e} - \Pi_k(\mathbf{e}) + \sum_{i=1}^{N_{cr}} (\lambda_h^i - \lambda^i) \Pi_k(\mathbf{x}_i^R) + \sum_{i=1}^{N_{cr}} \lambda^i (\Pi_k(\mathbf{x}_i^R) - \mathbf{x}_i^R).$$

D'après l'inégalité de Cauchy-Schwarz, en décomposant cette somme en $2N_{cr} + 1$ termes, on a l'inégalité suivante :

$$\begin{aligned} \|\mathbf{z}_\lambda\|_{\mathbf{X}_E}^2 \leq (2N_{cr} + 1) & \left(\|\mathbf{e} - \Pi_k(\mathbf{e})\|_{\mathbf{X}_E}^2 + \sum_{i=1}^{N_{cr}} |\lambda_h^i - \lambda^i|^2 \|\Pi_k(\mathbf{x}_i^R)\|_{\mathbf{X}_E}^2 \right. \\ & \left. + \sum_{i=1}^{N_{cr}} |\lambda^i|^2 \|\mathbf{x}_i^R - \Pi_k(\mathbf{x}_i^R)\|_{\mathbf{X}_E}^2 \right). \end{aligned}$$

Évaluons les trois ensembles de termes du second membre de cette inégalité.

- Comme $\mathbf{e} \in \mathbf{H}^{2\alpha-\epsilon}(\omega)$, il existe une constante $C_{\mathbf{e},\epsilon} > 0$ telle que :

$$\|\mathbf{e} - \Pi_k(\mathbf{e})\|_{\mathbf{X}_E} \leq C_{\mathbf{e},\epsilon} h^{2\alpha-\epsilon} \|\mathbf{e}\|_{\mathbf{H}^{2\alpha-\epsilon}}.$$

- Comme Π_k est continue de $C^0(\overline{\omega})$ dans $\mathbf{H}^1(\omega)$, on a pour tout i :

$$\|\Pi_k(\mathbf{x}_i^R)\|_{\mathbf{X}_E} \leq \|\Pi_k\| \|\mathbf{x}_i^R\|_{\mathbf{H}^{1+s}},$$

avec : $\|\Pi_k\| = \sup_{\mathbf{F} \in \mathbf{X}_{E,k}^{0,R} : \|\mathbf{F}\|_{\mathbf{X}_E^0} = 1} \|\Pi_k(\mathbf{F})\|_{\mathbf{X}_E}.$

En utilisant ce résultat ainsi que celui de la proposition 4.13, on arrive à :

$$\sum_{i=1}^{N_{cr}} |\lambda^i - \lambda_h^i|^2 \|\Pi_k(\mathbf{x}_i^R)\|_{\mathbf{X}_E}^2 \leq C_\lambda'^2 (h^{2\alpha})^2,$$

avec $C_\lambda'^2 = N_{cr} C_\lambda^2 \|\Pi_k\|^2 \max_i (\|\mathbf{x}_i^R\|_{\mathbf{H}^{1+s}}^2).$

- Comme $\forall i, \mathbf{x}_i^R \in \mathbf{H}^{1+s}(\omega)$, il existe une constante $C_s' > 0$ telle que $\forall i$:

$$\|\mathbf{x}_i^R - \Pi_k(\mathbf{x}_i^R)\|_{\mathbf{X}_E} \leq C_s' h^k \|\mathbf{x}_i^R\|_{\mathbf{H}^{1+s}},$$

d'où :

$$\sum_{i=1}^{N_{cr}} |\lambda^i|^2 \|\mathbf{x}_i^R - \Pi_k(\mathbf{x}_i^R)\|_{\mathbf{X}_E}^2 \leq C_s^2 h^{2k},$$

avec $C_s^2 = N_{cr} \max_i (|\lambda^i|^2) C_s'^2 \|\mathbf{x}_i^R\|_{\mathbf{H}^{1+s}}^2.$

Finalement, en prenant $C_{\partial\omega} = \max(C_{\mathbf{e},\epsilon} \|\mathbf{e}\|_{\mathbf{H}^{2\alpha-\epsilon}}, C_\lambda', C_s,)$

$$\begin{aligned} \|\mathbf{z}_\lambda\|_{\mathbf{X}_E} & \leq C_{\partial\omega} h^{\min(2\alpha-\epsilon, 2\alpha, k)}, \\ & \leq C_{\partial\omega} h^{2\alpha-\epsilon}, \end{aligned}$$

où $C_{\partial\omega}$ dépend de $\|\mathbf{e}\|_{\mathbf{H}^{2\alpha-\epsilon}}$ et des $\|\mathbf{x}_i^R\|_{\mathbf{H}^{1,s}}.$

□

- Enfin, réécrivons $\mathbf{z}^P$ de la façon suivante :

$$\mathbf{z}^P = \sum_{i=1}^{N_{cr}} ((\lambda^i - \lambda_h^i) \mathbf{x}_i^P - \lambda_h^i (\mathbf{x}_i^P - \mathbf{x}_{i,h}^P)),$$

d'où :

$$\begin{aligned} \|\mathbf{z}^P\|_{\mathbf{X}_E} &\leq \sum_{i=1}^{N_{cr}} (|\lambda^i - \lambda_h^i| \|\mathbf{x}_i^P\|_{\mathbf{X}_E} + |\lambda_h^i| \|\mathbf{x}_i^P - \mathbf{x}_{i,h}^P\|_{\mathbf{X}_E}) , \\ &\leq C'_P h^{2\alpha} + \sum_{i=1}^{N_{cr}} |\lambda_h^i| \|\mathbf{x}_i^P - \mathbf{x}_{i,h}^P\|_{\mathbf{X}_E} , \end{aligned}$$

où C_P est une constante qui dépend des $\|\mathbf{x}_i^P\|_{\mathbf{X}_E}$ et de C_λ . Or, les erreurs $\|\mathbf{x}_i^P - \mathbf{x}_{i,h}^P\|_{\mathbf{X}_E}$ sont des erreurs de représentation et dépendent du nombre de quadratures p utilisé pour faire les intégrations numériques. D'où : $\|\mathbf{x}_i^P - \mathbf{x}_{i,h}^P\|_{\mathbf{X}_E} = O(h^{\phi(p)})$. Cette erreur est négligeable par rapport à celle faite sur les λ^i . D'où :

$$\|\mathbf{z}^P\|_{\mathbf{X}_E} \leq C_P h^{2\alpha} .$$

On est alors en mesure de montrer le théorème 4.12 :

DÉMONSTRATION DU THÉORÈME 4.12. Nous avons établi que l'erreur $\|\mathbf{z}^0\|_{\mathbf{X}_E}$ est en $h^{2\alpha-1-\epsilon}$. Les erreurs $\|\mathbf{z}^P\|_{\mathbf{X}_E}$ et $\|\mathbf{z}_\lambda\|_{\mathbf{X}_E}$ sont plus faibles. D'où, en regroupant les résultats des lemmes 4.14 et 4.15, on obtient la majoration (4.52). $\square$

Nous allons maintenant étudier l'erreur en norme $\mathbf{L}^2(\omega)$. Pour montrer le théorème 4.17 ci-dessous, obtenu en utilisant l'astuce d'Aubin-Nitsche, on a besoin du lemme suivant, prouvé dans la section 14.4 de la partie IV :

Lemme 4.16 Soit $\mathbf{f} \in \mathbf{L}^2(\omega)$. Considérons le problème variationnelle suivant :

Trouver $\mathbf{G} \in \mathbf{X}_E^0$ tel que :

$$\forall \mathbf{F} \in \mathbf{X}_E^0, \mathcal{A}_0(\mathbf{G}, \mathbf{F}) = (\mathbf{f}, \mathbf{F})_0 . \quad (4.57)$$

D'après le théorème de Lax-Milgram, ce problème admet une unique solution qui dépend continûment de $\mathbf{f}$. Soit $\mathbf{G}_h$ l'approximation de $\mathbf{G}$ par la λ -approche. On a alors l'approximation d'erreur suivante sur le calcul de $\mathbf{G}_h$: $\forall \epsilon > 0$ tel que $2\alpha - 1 - \epsilon > 0$ et $2(\alpha - 1) - \epsilon > 0$:

$$\|\mathbf{G} - \mathbf{G}_h\|_{\mathbf{X}_E} \leq C_G \|\mathbf{f}\|_0 h^{2\alpha-1-\epsilon} \text{ pour } \alpha \leq 3/4 ,$$

$$\|\mathbf{G} - \mathbf{G}_h\|_{\mathbf{X}_E} \leq C_G \|\mathbf{f}\|_0 h^{2(1-\alpha)-\epsilon} \text{ pour } \alpha \geq 3/4 .$$

Théorème 4.17 On suppose qu'il existe $\epsilon > 0$ tel que $2(2\alpha - 1) - \epsilon > 0$, $1 - \epsilon > 0$, $f, g \in H^{2\alpha-1-\epsilon}(\omega)$ et $e \in H^{2\alpha-1/2-\epsilon}(\partial\omega)$. Alors, il existe une constante $C_\epsilon^0 > 0$ qui ne dépend que de ϵ telle que :

$$\|\mathbf{E} - \mathbf{E}_h\|_0 \leq C_\epsilon^0 h^{2(2\alpha-1)-\epsilon} \text{ pour } \alpha \leq 3/4 , \quad (4.58)$$

$$\|\mathbf{E} - \mathbf{E}_h\|_0 \leq C_\epsilon^0 h^{1-\epsilon} \text{ pour } \alpha \geq 3/4 .$$

DÉMONSTRATION.

• Considérons tout d'abord le problème suivant :

Trouver $\mathbf{G} \in \mathbf{X}_E^0$ tel que :

$$\forall \mathbf{F} \in \mathbf{X}_E^0, \mathcal{A}_0(\mathbf{G}, \mathbf{F}) = (\mathbf{E} - \mathbf{E}_h, \mathbf{F})_0 . \quad (4.59)$$

Comme $(\mathbf{E} - \mathbf{E}_h) \in \mathbf{L}^2(\omega)$, on est dans le cadre du lemme 4.16. Soit $\epsilon > 0$ tel que $2\alpha - 1 - \epsilon > 0$ et $2(1 - \alpha) - \epsilon > 0$. L'approximation de $\mathbf{G}$ par la λ -approche est telle que :

$$\begin{aligned} \|\mathbf{G} - \mathbf{G}_h\|_{\mathbf{X}_E} &\leq C_G \|\mathbf{E} - \mathbf{E}_h\|_0 h^{2\alpha-1-\epsilon} \text{ pour } \alpha \leq 3/4 , \\ \|\mathbf{G} - \mathbf{G}_h\|_{\mathbf{X}_E} &\leq C_G \|\mathbf{E} - \mathbf{E}_h\|_0 h^{2(1-\alpha)-\epsilon} \text{ pour } \alpha \geq 3/4 . \end{aligned} \quad (4.60)$$

• Maintenant, nous allons majorer $\|\mathbf{E} - \mathbf{E}_h\|_0$ à l'aide de $\|\mathbf{G} - \mathbf{G}_h\|_{\mathbf{X}_E^0}$. Notons d'abord qu'en prenant comme fonction test $\mathbf{G}_h$ dans (2.43) et dans (4.41)-(4.42), on a l'égalité suivante :

$$(\mathbf{E}, \mathbf{G}_h)_{\mathbf{X}_E^0} = (\mathbf{E}_h, \mathbf{G}_h)_{\mathbf{X}_E^0}, \quad (4.61)$$

ce qui nous permet d'obtenir que : $(\mathbf{E} - \mathbf{E}_h, \mathbf{G}_h)_{\mathbf{X}_E^0} = 0$. C'est la propriété d'orthogonalité de Galerkin. Décomposons $\|\mathbf{E} - \mathbf{E}_h\|_0^2$ ainsi :

$$\begin{aligned} \|\mathbf{E} - \mathbf{E}_h\|_0^2 &= \mathcal{A}_0(\mathbf{G}, \mathbf{E} - \mathbf{E}_h) = \mathcal{A}_0(\mathbf{E} - \mathbf{E}_h, \mathbf{G}), \\ &= (\mathbf{E} - \mathbf{E}_h, \mathbf{G})_{\mathbf{X}_E^0} - (\mathbf{E} - \mathbf{E}_h, \mathbf{G}_h)_{\mathbf{X}_E^0}, \\ &= (\mathbf{E} - \mathbf{E}_h, \mathbf{G} - \mathbf{G}_h)_{\mathbf{X}_E^0}, \\ &\leq M \|\mathbf{E} - \mathbf{E}_h\|_{\mathbf{X}_E^0} \|\mathbf{G} - \mathbf{G}_h\|_{\mathbf{X}_E^0}, \text{ où } M \text{ est la constante de coercivité de } \mathcal{A}_0; \\ &\leq M C_\epsilon h^{2\alpha-1-\epsilon} \|\mathbf{G} - \mathbf{G}_h\|_{\mathbf{X}_E^0}, \text{ d'après le théorème 4.12,} \\ &\leq \begin{cases} M C_\epsilon C_G h^{2(2\alpha-1)-\epsilon} \|\mathbf{E} - \mathbf{E}_h\|_0, & \text{pour } \alpha \leq 3/4, \\ M C_\epsilon C_G h^{2\alpha-1+2(1-\alpha)-\epsilon} \|\mathbf{E} - \mathbf{E}_h\|_0, & \text{pour } \alpha \leq 3/4, \end{cases} \quad \text{d'après (4.60).} \end{aligned}$$

Posons $C_\epsilon^0 = M C_\epsilon C_G$. En divisant la dernière inégalité par $\|\mathbf{E} - \mathbf{E}_h\|_0$, on obtient (4.58). $\square$

Remarque 4.18 *Si nécessaire (en particulier, dans le cas où α est proche de $1/2$), il est possible d'améliorer la vitesse de convergence : en raffinant localement le maillage, ou en calculant explicitement les termes singuliers suivants du champ [86], ou encore en calculant avec précision la solution au voisinage des coins de $\partial\omega$ à l'aide d'opérateurs Dirichlet-Neumann [8].*

4.11.3 Régularisation à poids : convergence

L'analyse d'erreur faite dans [53] repose sur la décomposition de $\mathbf{E}$ en une partie régulière et le gradient d'une fonction dont le Laplacien est dans V_γ^0 : $\mathbf{E} = \mathbf{z}_0 + \mathbf{grad} \phi_\gamma$, où $\mathbf{z}_0 \in \mathbf{H}^1(\omega)$ et $\phi_\gamma \in V_\gamma^2 \cap H_0^1(\omega)$. Les hypothèses sur le champ électrique sont plus fortes que celles de notre problème, car dans [53], il s'agit des équations de Maxwell harmoniques : $\mathbf{E}$ est à divergence nulle et $\text{rot } \mathbf{E} \in H^1(\omega)$. D'autre part, les hypothèses sur l'espace de discrétisation nécessitent de prendre des éléments finis d'ordre supérieur ou égal à 2. Dans ces conditions (et avec les bons poids γ), on a le résultat suivant :

Proposition 4.19 *Soit $\varepsilon > 0$ donné. Il existe une constante $C_\varepsilon > 0$ qui ne dépend que ε telle que :*

$$\|\mathbf{E}_\gamma - \mathbf{E}_{\gamma,h}\|_{\mathbf{X}_{E,\gamma}^0} \leq C_\varepsilon h^{\alpha+\gamma-1-\varepsilon}. \quad (4.62)$$

Nous avons toutefois constaté dans les expériences numériques que la méthode converge avec la même vitesse pour les éléments finis P_1 de Lagrange, et pour des données moins régulières (par exemple, $\text{div } \mathbf{E} = s_D$ et/ou $\text{rot } \mathbf{E} = s_N$). Notons qu'il faut travailler avec un maillage assez régulier, car les expériences numériques montrent que cette méthode est plus sensible que les autres à la qualité du maillage.

La vitesse de convergence de la méthode à poids est donc meilleure que celle de la méthode du complément singulier, d'autant meilleure que γ est proche de 1. Cependant, la norme de $\mathbf{X}_{\mathbf{E},\gamma}^0$ est plus faible que la norme $\mathbf{X}_{\mathbf{E}}^0$: dans $\mathbf{X}_{\mathbf{E},\gamma}^0$, $\operatorname{div} \mathbf{E}_{\gamma,h}$ converge vers g en norme $L^2_\gamma(\omega)$, mais ne converge pas vers g en norme $L^2(\omega)$.

4.11.4 Conclusion

D'après [44], si $\operatorname{rot}(\operatorname{rot} \mathbf{E}) \in \mathbf{L}^2(\omega)$ et si il existe $0 < \epsilon < \alpha$ tel que $\mathbf{E} \in \mathbf{H}^{\alpha-\epsilon}(\omega)$, alors $\mathbf{E}_{Nd}$, l'approximation de $\mathbf{E}$ par les éléments finis de Nédélec de [88] est telle que :

$$\|\operatorname{rot} \mathbf{E} - \operatorname{rot} \mathbf{E}_{Nd}\|_0 \leq C_\epsilon h^{\alpha-\epsilon} \text{ et } \|\mathbf{E} - \mathbf{E}_{Nd}\|_0 \leq C'_\epsilon h^{\alpha-\epsilon},$$

où $C_\epsilon > 0$ et $C'_\epsilon > 0$ sont des constantes.

D'après [70], avec les mêmes hypothèses sur $\mathbf{E}$ que celles ci-dessus, $\mathbf{E}_{GD}$, l'approximation de $\mathbf{E}$ par les éléments finis mixtes de Galerkin discontinus est telle que :

$$\|\mathbf{E} - \mathbf{E}_{GD}\|_{GD} \leq C_{GD,\epsilon} h^{\alpha-\epsilon}$$

où $C_{GD,\epsilon} > 0$ est une constante. $\|\cdot\|_{GD}$ est une norme qui contient la norme $\mathbf{L}^2(\omega)$, la somme des normes $L^2(T_l)$ du rotationnel et la somme des norme des sauts de discontinuité à travers les arêtes. Sur la figure 4.3, on a représenté en bleu le taux de convergence de ces deux méthodes en fonction de α .

Ainsi, si $\operatorname{rot}(\operatorname{rot} \mathbf{E}) \in \mathbf{L}^2(\omega)$, le taux de convergence en norme $\mathbf{X}_{\mathbf{E},\gamma}^0$ de la méthode à poids correspond à celui des deux méthodes ci-dessus, à $(1 - \gamma)$ près. Pour la λ -approche, bien que le taux de convergence en norme $\mathbf{X}_{\mathbf{E}}$ soit plus faible, nous avons d'une part une hypothèse plus faible sur $\operatorname{rot} \mathbf{E}$, et d'autre part, le taux de convergence en norme $\mathbf{L}^2(\omega)$ est meilleur si $\alpha > 2/3$, c'est-à-dire si l'angle du coin rentrant est plus petit que $3\pi/2$ (voir la figure 4.3).

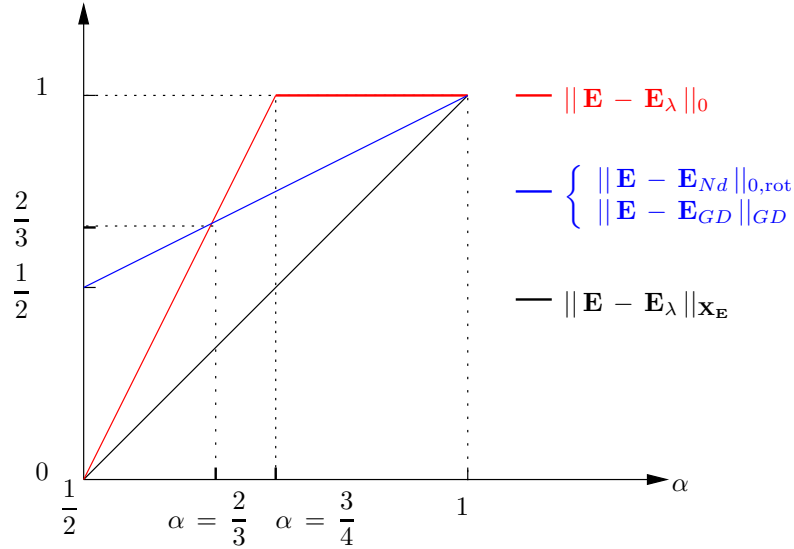


FIG. 4.3 – Comportement des erreurs.

Chapitre 5

Le problème statique 2D mixte discret

5.1 L'élément fini mixte de Taylor-Hood P_2 - P_1

Définition 5.1 On parle d'élément fini de Taylor-Hood P_2 - P_1 lorsque le multiplicateur de Lagrange et le champ électrique sont évalués sur le même maillage, le multiplicateur de Lagrange étant approché par les éléments finis de Lagrange continus en P_1 et le champ électrique est approché composante par composante par les éléments finis de Lagrange continus en P_2 . L'espace d'approximation du champ électrique est alors :

$$\mathbf{X}_2 = \{ \mathbf{v} \in C^0(\overline{\omega})^2 : \mathbf{v}|_{T_h} \in P_2^2(T_h), \forall T_h \in \mathcal{T}_h \},$$

et celui du multiplicateur de Lagrange est alors :

$$V_1 = \{ u \in C^0(\overline{\omega}) : u|_{T_h} \in P_1(T_h), \forall T_h \in \mathcal{T}_h \}.$$

On notera N_1 le nombre de degrés de liberté P_1 , c'est-à-dire le nombre de sommets de la triangulation du maillage ; et N le nombre de degrés de liberté P_2 , c'est-à-dire le nombre de sommets plus le nombre d'arêtes de la triangulation du maillage.

Considérons S_{i_1} , $i_1 \in \{1, \dots, N_1\}$ les sommets des triangles T_l . Soit N_ω^1 le nombre de sommets intérieurs à ω , et $N_{\partial\omega}^1$ le nombre de sommets du bord $\partial\omega$. On ordonne ces points P_1 ainsi : $\forall i_1 \in I_\omega^1$, $S_{i_1} \in \omega$, et $\forall i_1 \in I_{\partial\omega}^1$, $S_{i_1} \in \partial\omega$, où on a défini les ensembles d'indices suivants :

$$I_\omega^1 = \{1, \dots, N_\omega^1\}, I_{\partial\omega}^1 = \{N_\omega^1 + 1, \dots, N_\omega^1 + N_{\partial\omega}^1\} \text{ et } I_1 = I_\omega^1 \cup I_{\partial\omega}^1.$$

Soient u_{i_1} , $i_1 \in I_1$ les fonctions de base P_1 , qui génèrent V_1 .

Le champ électrique discretisé est recherché dans l'espace $(V_2)^2$. On ordonnera les points de discrétisation P_2 (sommets et arêtes des triangles) M_i , $i \in \{1, \dots, N\}$ comme c'est indiqué dans la section 4.2. Soit $i \in I$ tel que M_i soit un sommet de la triangulation. Alors il existe $i_1 \in I_1$ tel que $S_{i_1} = M_i$. On considérera v_i , $i \in I$ les fonctions de base P_2 qui génèrent V_2 . On notera p_h l'approximation du multiplicateur de Lagrange et $\mathbf{E}_{(\gamma),h}$ l'approximation du champ électrique. $p_{(\gamma),h}$ est décomposé ainsi :

$$p_{(\gamma),h} = \sum_{i_1 \in I_1} p_{(\gamma),h}(S_{i_1}) u_{i_1}.$$

On notera $\underline{p}_{(\gamma)} \in \mathbb{R}^{N_1}$ le vecteur associé à $p_{(\gamma),h} : \underline{p}_{(\gamma)} = (p_{(\gamma),h}(S_1), \dots, p_{(\gamma),h}(S_{N_1}))^T$.

Remarque 5.2 On peut aussi utiliser les éléments finis de Taylor-Hood P_2 -iso- P_1 : le multiplicateur de Lagrange est alors évalué sur un maillage grossier $\mathcal{T}_h$ en P_1 et les composantes du champ

électrique sont évaluées sur un maillage fin $\mathcal{T}_{2h}$, obtenu par subdivision des triangles du maillage grossier en P_1 . Les espaces d'approximation sont les suivants :

$$\mathbf{X}'_2 = \{ \mathbf{v} \in C^0(\bar{\omega})^2 : \mathbf{v}|_{T_h} \in P_1^2(T_h), \forall T_h \in \mathcal{T}_h \},$$

pour le champ électrique, et

$$V_1 = \{ u \in C^0(\bar{\omega}) : u|_{T_{2h}} \in P_1(T_{2h}), \forall T_{2h} \in \mathcal{T}_{2h} \},$$

pour le multiplicateur de Lagrange. La complexité de l'implémentation est similaire à celle des éléments finis P_2 - P_1 . Les éléments finis P_2 - P_1 convergent mieux si la solution est régulière.

5.2 Condition inf-sup discrète

5.2.1 Méthode avec CL naturelles

La preuve de la condition inf-sup discrète uniforme de la formulation mixte (3.5) pour l'élément fini de Taylor-Hood P_2 -iso- P_1 a été faite par P. Ciarlet, Jr. et V. Girault dans [39] en 3D, à l'aide d'un relèvement discret de la condition aux limites normale (voir aussi [19] pour le problème de Stokes).

Il s'agit alors de montrer la condition inf-sup pour le couple d'espaces $((V_2^0)^3, (V_1 \cap L_0^2(\Omega)))$. Pour cela, dans le cas d'un maillage quelconque, les auteurs se ramènent aux inégalités locales suivantes :

$$(q_h, \operatorname{div} \mathbf{v}_h)_{0,T} \geq c \rho_T^2 |q_h|_{1,T}^2 \text{ et } |\mathbf{v}_h|_{1,T} \leq c' h_T |q_h|_{1,T}.$$

Grâce à la technique de [106], qui prouve la condition inf-sup discrète pour le couple d'espace $((V_2^0)^2, (V_1 \cap L_0^2(\omega)))$ pour le problème de Stokes, on obtient la condition inf-sup globale.

Pour le couple d'espaces $(\mathbf{X}_2, V_1)$, la preuve de la condition inf-sup se montre de la même façon. D'où :

Proposition 5.3 *Il existe une constante $\kappa_{\mathbf{E}}^* > 0$ indépendante de h telle que :*

$$\inf_{u_h \in V_1} \sup_{\mathbf{F}_h \in \mathbf{X}_{\mathbf{E},2}} \frac{\mathcal{B}_{\mathbf{E}}(\mathbf{F}_h, u_h)}{\|\mathbf{F}_h\|_{\mathbf{X}_{\mathbf{E}}} \|u_h\|_0} \geq \kappa_{\mathbf{E}}^*.$$

5.2.2 Méthode avec CL essentielles

$\mathbf{X}_{\mathbf{E},2}^{0,R}$ est un sous-espace de $\mathbf{X}_2$, et $\forall (\mathbf{F}_h, u_h) \in \mathbf{X}_{\mathbf{E},2}^{0,R} \times V_1$, $\mathcal{B}_0(\mathbf{F}_h, u_h) = \mathcal{B}_{\mathbf{E}}(\mathbf{F}_h, u_h)$, on a donc aussi une condition inf-sup discrète pour le couple d'espaces $(\mathbf{X}_{\mathbf{E},2}^{0,R}, V_1)$:

Proposition 5.4 *Il existe une constante $\kappa_0^* > 0$ indépendante de h telle que :*

$$\inf_{u_h \in V_1} \sup_{\mathbf{F}_h \in \mathbf{X}_{\mathbf{E},2}^{0,R}} \frac{\mathcal{B}_0(\mathbf{F}_h, u_h)}{\|\mathbf{F}_h\|_{\mathbf{X}_{\mathbf{E}}^0} \|u_h\|_0} \geq \kappa_0^*.$$

5.2.3 Régularisation à poids

Nous n'avons pas obtenu de résultat théorique sur la condition inf-sup discrète dans $\mathbf{X}_{\mathbf{E},\gamma}^0$, cependant les résultats numériques suggèrent que cette condition n'est pas vérifiée de façon uniforme : sur des maillages homogénéisés, le nombre d'itérations de l'algorithme d'Uzawa se comporte linéairement en fonction du logarithme du pas maillage, ce qui nous laisse penser que κ_{γ}^* dépend du

logarithme du pas du maillage. Comme on peut montrer que $\|p_h\|_{0,\gamma}$ reste petit (voir la section 5.6), la méthode converge tout de même assez rapidement.

Le problème de la condition inf-sup discrète dans des espaces à poids a été étudié pour le problème de Stokes. Dans [64], V. Girault et al. obtiennent une condition inf-sup discrète avec un poids qui n'est jamais nul (il est de la forme $(r^2 + \kappa^2 h^2)^{1/2}$). Dans [14], Z. Belhachmi et al. obtiennent une condition inf-sup discrète pour le problème de Stokes axisymétrique.

5.3 Méthode avec CL naturelles mixte : discrétisation

La formulation variationnelle (3.5) discrétisée par les éléments finis de Taylor-Hood P_2 - P_1 s'écrit de la façon suivante :

Trouver le couple $(p_h, \mathbf{E}_h) \in V_1 \times \mathbf{X}_2$ tel que :

$$\begin{aligned} \forall i \in I, \forall \alpha \in \{1, 2\} \quad \mathcal{A}_{\mathbf{E}}(\mathbf{E}_h, v_i \mathbf{e}_\alpha) + \mathcal{B}_{\mathbf{E}}(v_i \mathbf{e}_\alpha, p_h) &= \mathcal{L}_{\mathbf{E}}(v_i \mathbf{e}_\alpha), \\ \forall i_1 \in I_1 \quad \sum_{j \in I} \mathcal{B}_{\mathbf{E}}(\mathbf{E}_h, u_{i_1}) &= \mathcal{G}_{\mathbf{E}}(v_{i_1}). \end{aligned} \quad (5.1)$$

En décomposant p_h dans (5.1), on obtient :

$$\forall i \in I, \forall \alpha \in \{1, 2\} \quad \mathcal{A}_{\mathbf{E}}(\mathbf{E}_h, v_i \mathbf{e}_\alpha) + \sum_{j_1 \in I_1} p_h(S_{j_1}) \mathcal{B}_{\mathbf{E}}(v_i \mathbf{e}_\alpha, u_{i_1}) = \mathcal{L}_{\mathbf{E}}(v_i \mathbf{e}_\alpha),$$

$$\forall i_1 \in I_1 \quad \mathcal{B}_{\mathbf{E}}(\mathbf{E}_h, u_{i_1}) = \mathcal{G}_{\mathbf{E}}(u_{i_1}).$$

Soit $\underline{\mathbf{E}} \in (\mathbb{R}^2)^N$, la représentation de $\mathcal{E}_h$ dans $\text{vect}(v_i \mathbf{e}_\alpha)_{i \in I, \alpha \in \{1,2\}}$. Nous allons écrire les équations ci-dessus de façon matricielle. Pour cela, nous devons construire les matrices associées à $\mathcal{B}_{\mathbf{E}}$ et $\mathcal{G}_{\mathbf{E}}$.

- Soit $\mathbb{C}_{\mathbf{E}} \in (\mathbb{R}^{1 \times 2})^{N_1 \times N}$ la matrice composée des sous blocs $\mathbb{C}_{\mathbf{E}}^{i_1, j} \in \mathbb{R}^2$ tels que :

$$\forall i_1 \in I_1, \forall j \in I, \mathbb{C}_{\mathbf{E}}^{i_1, j} = \begin{pmatrix} \mathcal{B}_{\mathbf{E}}(v_j \mathbf{e}_1, u_{i_1}) & \mathcal{B}_{\mathbf{E}}(v_j \mathbf{e}_2, u_{i_1}) \end{pmatrix} = \begin{pmatrix} (\partial_1 v_j, u_{i_1})_0 & (\partial_2 v_j, u_{i_1})_0 \end{pmatrix}.$$

- Soit $\underline{\mathbf{G}} \in \mathbb{R}^{N_1}$ le vecteur tel que :

$$\forall i_1 \in I_1, \underline{\mathbf{G}}^{i_1} = \mathcal{G}_{\mathbf{E}}(u_{i_1}).$$

Les équations (5.1) se mettent sous la forme :

$$\boxed{\begin{aligned} \mathbb{A}_{\mathbf{E}} \underline{\mathbf{E}} + \mathbb{C}_{\mathbf{E}}^T \underline{\mathbf{p}} &= \underline{\mathbf{L}}, \\ \mathbb{C}_{\mathbf{E}} \underline{\mathbf{E}} &= \underline{\mathbf{G}}. \end{aligned}} \quad (5.2)$$

L'algorithme de résolution de (5.2) est le suivant :

- Résoudre $\mathbb{A}_{\mathbf{E}} \underline{\mathbf{E}}_0 = \underline{\mathbf{L}}$;
- Résoudre $(\mathbb{C}_{\mathbf{E}} \mathbb{A}_{\mathbf{E}}^{-1} \mathbb{C}_{\mathbf{E}}^T) \underline{\mathbf{p}} = \mathbb{C}_{\mathbf{E}} \underline{\mathbf{E}}_0 - \underline{\mathbf{G}}$;
- Résoudre $\mathbb{A}_{\mathbf{E}} \underline{\mathbf{E}} = \underline{\mathbf{L}} - \mathbb{C}_{\mathbf{E}}^T \underline{\mathbf{p}}$.

5.3.1 Matrice mixte

La matrice mixte $\mathbb{C}_{\mathbf{E}} \in (\mathbb{R}^{1 \times 2})^{N_1 \times 2N}$ est composée des sous blocs $\mathbb{C}_{\mathbf{E}}^{i_1, j} \in \mathbb{R}^{1 \times 2}$ tels que :

$\forall i_1 \in I_1, \forall j \in I,$

$$\mathbb{C}_{\mathbf{E}}^{i_1, j} = \begin{pmatrix} (\partial_1 v_j, u_{i_1})_0 & (\partial_2 v_j, u_{i_1})_0 \end{pmatrix} = \sum_{T_l \mid S_{i_1}, M_j \in T_l} \left(\int_{T_l} \partial_1 v_j u_{i_1} d\omega \quad \int_{T_l} \partial_2 v_j u_{i_1} d\omega \right).$$

Pour chaque triangle T_l , on doit donc calculer les fonctions de base P_1 , notées u_{i_1} et les fonctions de base P_2 , notées v_j . L'intégration se fera par un schéma d'intégration numérique.

5.3.2 Second membre

Le vecteur du second membre $\underline{\mathbf{G}} \in \mathbb{R}^{N_1}$ est tel que : $\underline{\mathbf{G}}^{i_1} = (g, u_{i_1})_0$. Ce calcul peut être effectué directement, à l'aide d'un schéma d'intégration numérique, ou en utilisant la fonction d'interpolation P_1 de g : $g_h = \Pi_1(g)$. Soit $\mathbb{M}_1 \in \mathbb{R}^{N_1 \times N_1}$ la matrice de masse P_1 telle que :

$\forall i_1, j_1 \in I_1,$

$$\mathbb{M}_1^{i_1, j_1} = (u_{i_1}, u_{j_1})_0 = \sum_{T_l \mid S_{i_1}, S_{j_1} \in T_l} \int_{T_l} u_{i_1} u_{j_1} d\omega.$$

On en déduit dans le second cas : $\underline{\mathbf{G}} = \mathbb{M}_1 \underline{\mathbf{g}}$.

5.4 Complément singulier orthogonal mixte : discrétisation

La formulation variationnelle (3.7) discrétisée par les éléments finis de Taylor-Hood P_2 - P_1 s'écrit de la façon suivante :

Trouver le triplet $(p_h, \hat{\mathbf{E}}_h, \mathbf{c}_h) \in V_1 \times \mathbf{X}_2 \times \mathbb{R}^{N_{cr}}$ tel que :

$$\begin{aligned} \forall i \in I_\omega, \forall \alpha \in \{1, 2\}, \quad & \mathcal{A}_0(\hat{\mathbf{E}}_h, v_i \mathbf{e}_\alpha) + \mathcal{B}_0(v_i \mathbf{e}_\alpha, p_h) = \mathcal{L}_0(v_i \mathbf{e}_\alpha), \\ \forall i \in I_a, \quad & \mathcal{A}_0(\hat{\mathbf{E}}_h, v_i \boldsymbol{\nu}_i) + \mathcal{B}_0(v_i \boldsymbol{\nu}_i, p_h) = \mathcal{L}_0(v_i \boldsymbol{\nu}_i), \\ \forall i \in I_a, \quad & \hat{\mathbf{E}}_h(M_i) \cdot \boldsymbol{\tau}_i = \mathbf{e}_h(M_i) \cdot \boldsymbol{\tau}_i, \\ \forall i \in I_c, \quad & \hat{\mathbf{E}}_h(M_i) = \mathbf{e}_h(M_i), \\ \forall m \in \{1, \dots, N_{cr}\}, \quad & \pi \sum_{l=1}^{N_{cr}} c_{l,h} \beta_h^{l,m} + \mathcal{B}_0(\mathbf{x}_m^S, p_h) = \pi \lambda_h^m, \\ \forall i_1 \in I_1, \quad & \mathcal{B}_0(\hat{\mathbf{E}}_h, u_{i_1}) + \sum_{l=1}^{N_{cr}} c_{l,h} \mathcal{B}_0(\mathbf{x}_l^S, u_{i_1}) = \mathcal{G}_{\mathbf{E}}(u_{i_1}). \end{aligned} \tag{5.3}$$

Rappelons que $\operatorname{div} \mathbf{x}_l^S = s_D^l$. Ainsi, on approche $\mathcal{B}_0(\mathbf{x}_l^S, u_{i_1})$ par $(s_{D,h}, u_{i_1})_0$. En décomposant p_h dans (5.3), on obtient alors :

$$\forall i \in I_\omega, \forall \alpha \in \{1, 2\}, \quad \mathcal{A}_0(\widehat{\mathbf{E}}_h(M_j), v_i \mathbf{e}_\alpha) + \sum_{i_1=1}^{N_1} p_h(S_{i_1}) \mathcal{B}_0(v_i \mathbf{e}_\alpha, u_{i_1}) = \mathcal{L}_0(v_i \mathbf{e}_\alpha),$$

$$\forall i \in I_a, \quad \mathcal{A}_0(\widehat{\mathbf{E}}_h, v_i \boldsymbol{\nu}_i) + \sum_{i_1=1}^{N_1} p_h(S_{i_1}) \mathcal{B}_0(v_i \boldsymbol{\nu}_i, u_{i_1}) = \mathcal{L}_0(v_i \boldsymbol{\nu}_i),$$

$$\forall i \in I_a, \quad \widehat{\mathbf{E}}_h \cdot \boldsymbol{\tau}_i = \mathbf{e}_h(M_i) \cdot \boldsymbol{\tau}_i,$$

$$\forall i \in I_c, \quad \widehat{\mathbf{E}}_h(M_i) = \mathbf{e}_h(M_i),$$

$$\forall m \in \{1, \dots, N_{cr}\}, \quad \pi \sum_{l=1}^{N_{cr}} c_{l,h} \beta_h^{l,m} + (s_{D,h}^m, p_h)_0 = \pi \lambda_h^m,$$

$$\forall i_1 \in I_1, \quad \mathcal{B}_0(\widehat{\mathbf{E}}_h, u_{i_1}) + \sum_{l=1}^{N_{cr}} c_{l,h} (s_{D,h}^l, u_{i_1})_0 = \mathcal{G}_{\mathbf{E}}(u_{i_1}).$$

Soit $\widehat{\mathbf{E}} \in (\mathbb{R}^2)^N$, la représentation de $\widehat{\mathbf{E}}_h$ dans la base B (4.40). Nous allons écrire les équations ci-dessus de façon matricielle. Pour cela, nous devons construire les matrices associées à $(s_{D,h}^l, \cdot)_0$ et $\mathcal{B}_0$.

- Soit $\mathbb{S}_D \in \mathbb{R}^{N_1 \times N_{cr}}$ la matrice composée de N_{cr} vecteurs colonnes de $\mathbb{R}^N$ tels que $\forall l \in \{1, \dots, N_{cr}\}$:

$$\forall i_1 \in I_1, \mathbb{S}_D^{i_1, l} = (s_{D,h}^l, u_{i_1})_0.$$

- Soit $\mathbb{C}_0 \in (\mathbb{R}^{1 \times 2})^{N_1 \times 2N}$ la matrice composée des sous-blocs $\mathbb{C}_0^{i_1, j} \in \mathbb{R}^{1 \times 2}$ tels que :

$$\forall i_1 \in I_1, \forall j \in I_\omega \quad \mathbb{C}_0^{i_1, j} = \mathbb{C}_{\mathbf{E}}^{i_1, j},$$

$$\forall i_1 \in I_1, \forall j \in I_c \quad \mathbb{C}_0^{i_1, j} = \begin{pmatrix} 0 & 0 \end{pmatrix},$$

$$\forall i_1 \in I_1, \forall j \in I_a \quad \mathbb{C}_0^{i_1, j} = \begin{pmatrix} 0 & \mathcal{B}_0(v_j \boldsymbol{\nu}_j, u_{i_1}) \end{pmatrix} = \begin{pmatrix} 0 & (\partial_{\nu_j} v_j, u_{i_1})_0 \end{pmatrix}.$$

- Soit $\mathbb{C}_r \in (\mathbb{R}^{1 \times 2})^{N_1 \times 2N}$ la matrice composée des sous-blocs $\mathbb{C}_r^{i_1, j} \in \mathbb{R}^{1 \times 2}$ tels que :

$$\forall i_1 \in I_1, \forall j \in I_\omega \quad \mathbb{C}_r^{i_1, j} = \begin{pmatrix} 0 & 0 \end{pmatrix},$$

$$\forall i_1 \in I_1, \forall j \in I_c \quad \mathbb{C}_r^{i_1, j} = \mathbb{C}_{\mathbf{E}}^{i_1, j},$$

$$\forall i_1 \in I_1, \forall j \in I_a \quad \mathbb{C}_r^{i_1, j} = \begin{pmatrix} \mathcal{B}_0(v_j \boldsymbol{\tau}_j, u_{i_1}) & 0 \end{pmatrix} = \begin{pmatrix} (\partial_{\tau_j} v_j, u_{i_1})_0 & 0 \end{pmatrix}.$$

Il s'agit de résoudre le système matriciel suivant :

$$\mathbb{A}_0 \widehat{\mathbf{E}} + \mathbb{C}_0^T \underline{\mathbf{p}} = \underline{\mathbf{L}}_0 - \mathbb{A}_r \underline{\mathbf{e}},$$

$$\pi \mathbb{X}_h \mathbf{c}_h + \mathbb{S}_D^T \underline{\mathbf{p}} = \pi \boldsymbol{\lambda}_h,$$

$$\mathbb{C}_0 \widehat{\mathbf{E}} + \mathbb{S}_D \mathbf{c}_h = \underline{\mathbf{G}} - \mathbb{C}_r \underline{\mathbf{e}}.$$

Écrivons ces équations sous la forme condensée suivante :

$$\begin{pmatrix} \mathbb{A}_0 & 0 \\ 0 & \pi \mathbb{X}_h \end{pmatrix} \begin{pmatrix} \hat{\underline{\mathbf{E}}} \\ \mathbf{c}_h \end{pmatrix} + \begin{pmatrix} \mathbb{C}_0^T \\ \mathbb{S}_D^T \end{pmatrix} \underline{\mathbf{p}} = \begin{pmatrix} \underline{\mathbf{L}}_0 - \mathbb{A}_r \underline{\mathbf{e}} \\ \pi \boldsymbol{\lambda}_h \end{pmatrix} \quad (5.4)$$

$$\begin{pmatrix} \mathbb{C}_0 & \mathbb{S}_D \end{pmatrix} \begin{pmatrix} \hat{\underline{\mathbf{E}}} \\ \mathbf{c}_h \end{pmatrix} = \underline{\mathbf{G}} - \mathbb{C}_r \underline{\mathbf{e}}.$$

Soit $\mathbb{C}' = \begin{pmatrix} \mathbb{C}_0 & \mathbb{S}_D \end{pmatrix} \in \mathbb{R}^{(N_1+N_{cr}) \times (2N+N_{cr})}$. Le système (5.4) se réécrit ainsi (avec $\mathbb{A}'$, $\underline{\mathbf{E}}'$ et $\underline{\mathbf{L}}'$ décrits au paragraphe 4.9.1) :

$$\begin{aligned} \mathbb{A}' \underline{\mathbf{E}}' + \mathbb{C}'^T \underline{\mathbf{p}} &= \underline{\mathbf{L}}', \\ \mathbb{C}' \underline{\mathbf{E}}' &= \underline{\mathbf{G}} - \mathbb{C}_r \underline{\mathbf{e}}. \end{aligned}$$

(5.5)

L'algorithme de résolution de (5.5) est le suivant :

- Résoudre $\mathbb{A}' \underline{\mathbf{E}}' = \underline{\mathbf{L}}'$, c'est-à-dire : $\underline{\mathbf{E}}'_0 = \begin{pmatrix} \hat{\underline{\mathbf{E}}}_0 \\ \mathbf{c}_h^0 \end{pmatrix}$, avec : $\mathbb{A}_0 \hat{\underline{\mathbf{E}}} = \underline{\mathbf{L}}_0 - \mathbb{A}_r \underline{\mathbf{e}}$ et $\mathbb{X}_h \mathbf{c}_h^0 = \boldsymbol{\lambda}_h$;
- Résoudre $(\mathbb{C}' (\mathbb{A}')^{-1} \mathbb{C}'^T) \underline{\mathbf{p}} = \mathbb{C}' \underline{\mathbf{E}}'_0 - (\underline{\mathbf{G}} - \mathbb{C}_r \underline{\mathbf{e}})$;
- Résoudre $\mathbb{A}_0 \hat{\underline{\mathbf{E}}} = \underline{\mathbf{L}}_0 - \mathbb{A}_r \underline{\mathbf{e}} - \mathbb{C}_0^T \underline{\mathbf{p}}$ et $\pi \mathbb{X}_h \mathbf{c}_h = \pi \boldsymbol{\lambda}_h - \mathbb{S}_D^T \underline{\mathbf{p}}$.

5.4.1 Matrices mixtes

Les calculs de $\mathbb{C}_0$ et $\mathbb{C}_r$ sont similaires à ceux de $\mathbb{C}_{\mathbf{E}}$. On a :

$$\forall i_1 \in \mathbf{I}_1, \forall j \in \mathbf{I}_a \quad \mathbb{C}_0^{i_1, j} = \begin{pmatrix} 0 & \sum_{T_l \mid S_{i_1}, M_j \in T_l} \int_{T_l} \partial_{\nu_j} v_j u_{i_1} d\omega \end{pmatrix},$$

$$\forall i_1 \in \mathbf{I}_1, \forall j \in \mathbf{I}_a \quad \mathbb{C}_r^{i_1, j} = \begin{pmatrix} \sum_{T_l \mid S_{i_1}, M_j \in T_l} \int_{T_l} \partial_{\tau_j} v_j u_{i_1} d\omega & 0 \end{pmatrix}.$$

5.4.2 Second membre

Nous avons remarqué que : $\mathcal{B}_0(\mathbf{x}_l^S, u_{i_1}) = (s_{D,l}, u_{i_1})_0$. Ainsi, $\mathbb{S}_D \in \mathbb{R}^{N_1 \times N_{cr}}$ est telle que :

$$\mathbb{S}_D^{i_1, l} = (s_{D,l,h}, u_{i_1})_0 = (\tilde{s}_{D,l,h}, u_{i_1})_0 + (s_{D,l,h}^P, u_{i_1})_0.$$

Le premier terme est tel que : $(s_{D,l,h}^P, u_{i_1})_0 = (\mathbb{M}_1 \underline{\mathbf{s}}_l)_{i_1}$, où $\underline{\mathbf{s}}_{D,l}$ représente l'approximation par les éléments finis P_1 de Lagrange de $\tilde{s}_{D,l}$. Le second terme, qui dépend de $s_{D,l}^P$ est calculé directement, à l'aide d'un schéma d'intégration numérique (section 15.2).

5.5 Régularisation à poids mixte : discrétisation

La formulation variationnelle (3.10) discrétisée par les éléments finis de Taylor-Hood P_2 - P_1 s'écrit de la façon suivante : Trouver le couple $(p_{\gamma,h}, \mathbf{E}_{\gamma,h}) \in \mathbf{V}_1 \times \mathbf{X}_2$ tel que :

$$\begin{aligned}
 \forall i \in \mathbf{I}_\omega, \forall \alpha \in \{1, 2\} \quad & \mathcal{A}_\gamma(\mathbf{E}_{\gamma,h}, v_i \mathbf{e}_\alpha) + \mathcal{B}_\gamma(p_{\gamma,h}, u_{j_1}) = \mathcal{L}_\gamma(v_i \mathbf{e}_\alpha), \\
 \forall i \in \mathbf{I}_a, \quad & \mathcal{A}_\gamma(\mathbf{E}_{\gamma,h}, v_i \boldsymbol{\nu}_i) + \mathcal{B}_\gamma(p_{\gamma,h}, u_{j_1}) = \mathcal{L}_\gamma(v_i \boldsymbol{\nu}_i), \\
 \forall i \in \mathbf{I}_a, \quad & \mathbf{E}_{\gamma,h}(M_i) \cdot \boldsymbol{\tau}_i = \mathbf{e}_h(M_i) \cdot \boldsymbol{\tau}_i, \\
 \forall i \in \mathbf{I}_c, \quad & \mathbf{E}_{\gamma,h}(M_i) = \mathbf{e}_h(M_i), \\
 \forall i_1 \in \mathbf{I}_1, \quad & \mathcal{B}_\gamma(\mathbf{E}_{\gamma,h}, u_{i_1}) = \mathcal{G}_\gamma(u_{i_1}).
 \end{aligned} \tag{5.6}$$

En décomposant $p_{\gamma,h}$ dans (5.6), on obtient alors :

$$\begin{aligned}
 \forall i \in \mathbf{I}_\omega, \forall \alpha \in \{1, 2\}, \quad & \mathcal{A}_\gamma(\mathbf{E}_{\gamma,h}, v_i \mathbf{e}_\alpha) + \sum_{i_1=1}^{N_1} p_{\gamma,h}(S_{i_1}) \mathcal{B}_\gamma(v_i \mathbf{e}_\alpha, u_{i_1}) = \mathcal{L}_\gamma(v_i \mathbf{e}_\alpha), \\
 \forall i \in \mathbf{I}_a, \quad & \mathcal{A}_\gamma(\mathbf{E}_{\gamma,h}, v_i \boldsymbol{\nu}_i) + \sum_{i_1=1}^{N_1} p_{\gamma,h}(S_{i_1}) \mathcal{B}_\gamma(v_i \boldsymbol{\nu}_i, u_{i_1}) = \mathcal{L}_\gamma(v_i \boldsymbol{\nu}_i), \\
 \forall i \in \mathbf{I}_a, \quad & \mathbf{E}_{\gamma,h}(M_i) \cdot \boldsymbol{\tau}_i = \mathbf{e}_h(M_i) \cdot \boldsymbol{\tau}_i, \\
 \forall i \in \mathbf{I}_c, \quad & \mathbf{E}_{\gamma,h}(M_i) = \mathbf{e}_h(M_i), \\
 \forall i_1 \in \mathbf{I}_1 \quad & \mathcal{B}_\gamma(\mathbf{E}_{\gamma,h}, u_{i_1}) = \mathcal{G}_\gamma(u_{i_1}).
 \end{aligned}$$

Soit $\underline{\mathbf{E}}_\gamma \in (\mathbb{R}^2)^N$, la représentation de $\mathbf{E}_{\gamma,h}$ dans B (4.40). Nous allons écrire les équations ci-dessus de façon matricielle. Pour cela, nous devons construire les matrices associées à $\mathcal{B}_\gamma$ et $\mathcal{G}_\gamma$.

- Soit $\mathbb{C}_\gamma \in (\mathbb{R}^{1 \times 2})^{N_1 \times 2N}$ la matrice composée des sous-blocs $\mathbb{C}_\gamma^{i_1, j} \in \mathbb{R}^{1 \times 2}$ tels que :

$$\begin{aligned}
 \forall i_1 \in \mathbf{I}_1, \forall j \in \mathbf{I}_\omega \quad & \mathbb{C}_\gamma^{i_1, j} = \begin{pmatrix} \mathcal{B}_\gamma(v_j \mathbf{e}_1, u_{i_1}) & \mathcal{B}_\gamma(v_j \mathbf{e}_2, u_{i_1}) \\ (\partial_1 v_j, u_{i_1})_{0, \gamma} & (\partial_2 v_j, u_{i_1})_{0, \gamma} \end{pmatrix},
 \end{aligned}$$

$$\forall i_1 \in \mathbf{I}_1, \forall j \in \mathbf{I}_c \quad \mathbb{C}_\gamma^{i_1, j} = \begin{pmatrix} 0 & 0 \end{pmatrix},$$

$$\forall i_1 \in \mathbf{I}_1, \forall j \in \mathbf{I}_a \quad \mathbb{C}_\gamma^{i_1, j} = \begin{pmatrix} 0 & \mathcal{B}_\gamma(v_j \boldsymbol{\nu}_j, u_{i_1}) \\ 0 & (\partial_{\nu_j} v_j, u_{i_1}) \end{pmatrix}.$$

- Soit $\mathbb{C}_{r, \gamma} \in (\mathbb{R}^{1 \times 2})^{N_1 \times 2N}$ la matrice composée des sous blocs $\mathbb{C}_{r, \gamma}^{i_1, j} \in \mathbb{R}^{1 \times 2}$ tels que :

$$\forall i_1 \in \mathbf{I}_1, \forall j \in \mathbf{I}_\omega \quad \mathbb{C}_{r, \gamma}^{i_1, j} = \begin{pmatrix} 0 & 0 \end{pmatrix},$$

$$\begin{aligned}
 \forall i_1 \in \mathbf{I}_1, \forall j \in \mathbf{I}_c \quad & \mathbb{C}_{r, \gamma}^{i_1, j} = \begin{pmatrix} \mathcal{B}_\gamma(v_j \mathbf{e}_1, u_{i_1}) & \mathcal{B}_\gamma(v_j \mathbf{e}_2, u_{i_1}) \\ (\partial_1 v_j, u_{i_1})_{0, \gamma} & (\partial_2 v_j, u_{i_1})_{0, \gamma} \end{pmatrix}, \\
 \forall i_1 \in \mathbf{I}_1, \forall j \in \mathbf{I}_a \quad & \mathbb{C}_{r, \gamma}^{i_1, j} = \begin{pmatrix} \mathcal{B}_\gamma(v_j \boldsymbol{\tau}_j, u_{i_1}) & 0 \\ (\partial_{\tau_j} v_j, u_{i_1})_{0, \gamma} & 0 \end{pmatrix}.
 \end{aligned}$$

- Soit enfin $\underline{\mathbf{G}}_\gamma \in \mathbb{R}^{N_1}$ le vecteur tel que :

$$\forall i_1 \in I_1, \underline{\mathbf{G}}_\gamma^{i_1} = \mathcal{G}_\gamma(u_{i_1}) = (g, u_{i_1})_{0,\gamma}.$$

Les équations (5.6) se mettent sous la forme :

$$\boxed{\begin{aligned} \mathbb{A}_\gamma \underline{\mathbf{E}}_\gamma + \mathbb{C}_\gamma^T \underline{\mathbf{p}}_\gamma &= \underline{\mathbf{L}}_\gamma - \mathbb{A}_{r,\gamma} \underline{\mathbf{e}}, \\ \mathbb{C}_\gamma \underline{\mathbf{E}}_\gamma &= \underline{\mathbf{G}}_\gamma - \mathbb{C}_{r,\gamma} \underline{\mathbf{e}}. \end{aligned}} \quad (5.7)$$

L'algorithme de résolution de (5.7) est le suivant :

- Résoudre $\mathbb{A}_\gamma \underline{\mathbf{E}}_{\gamma,0} = \underline{\mathbf{L}}_\gamma - \mathbb{A}_{r,\gamma} \underline{\mathbf{e}}$;
- Résoudre $(\mathbb{C}_\gamma \mathbb{A}_\gamma^{-1} \mathbb{C}_\gamma^T) \underline{\mathbf{p}}_\gamma = \mathbb{C}_\gamma \underline{\mathbf{E}}_{\gamma,0} - (\underline{\mathbf{G}}_\gamma - \mathbb{C}_{r,\gamma} \underline{\mathbf{e}})$;
- Résoudre $\mathbb{A}_\gamma \underline{\mathbf{E}}_\gamma = \underline{\mathbf{L}}_\gamma - \mathbb{A}_{r,\gamma} \underline{\mathbf{e}} - \mathbb{C}_\gamma^T \underline{\mathbf{p}}_\gamma$.

5.5.1 Matrices mixtes

$\mathbb{C}_\gamma \in \mathbb{R}^{N_1 \times 2N}$ est composée des sous blocs $\mathbb{C}_\gamma^{i_1,j} \in \mathbb{R}^2$ tels que :

$$\forall i_1 \in I_1, \forall j \in I_\omega \quad \mathbb{C}_\gamma^{i_1,j} = \begin{pmatrix} \sum_{T_l | S_{i_1}, M_j \in T_l} \int_{T_l} r^{2\gamma} \partial_1 v_j u_{i_1} d\omega & \sum_{T_l | S_{i_1}, M_j \in T_l} \int_{T_l} r^{2\gamma} \partial_2 v_j u_{i_1} d\omega \end{pmatrix},$$

$$\forall i_1 \in I_1, \forall j \in I_c \quad \mathbb{C}_\gamma^{i_1,j} = \begin{pmatrix} 0 & 0 \end{pmatrix},$$

$$\forall i_1 \in I_1, \forall j \in I_a \quad \mathbb{C}_\gamma^{i_1,j} = \begin{pmatrix} 0 & \sum_{T_l | S_{i_1}, M_j \in T_l} \int_{T_l} r^{2\gamma} \partial_{\nu_j} v_j u_{i_1} d\omega \end{pmatrix}.$$

$\mathbb{C}_{r,\gamma} \in \mathbb{R}^{N_1 \times 2N}$ est composée des sous blocs $\mathbb{C}_{r,\gamma}^{i_1,j} \in \mathbb{R}^{1 \times 2}$ tels que :

$$\forall i_1 \in I_1, \forall j \in I_\omega \quad \mathbb{C}_{r,\gamma}^{i_1,j} = \begin{pmatrix} 0 & 0 \end{pmatrix},$$

$$\forall i_1 \in I_1, \forall j \in I_c \quad \mathbb{C}_{r,\gamma}^{i_1,j} = \begin{pmatrix} \sum_{T_l | S_{i_1}, M_j \in T_l} \int_{T_l} r^{2\gamma} \partial_1 v_j u_{i_1} d\omega & \sum_{T_l | S_{i_1}, M_j \in T_l} \int_{T_l} r^{2\gamma} \partial_2 v_j u_{i_1} d\omega \end{pmatrix},$$

$$\forall i_1 \in I_1, \forall j \in I_a \quad \mathbb{C}_{r,\gamma}^{i_1,j} = \begin{pmatrix} \sum_{T_l | S_{i_1}, M_j \in T_l} \int_{T_l} r^{2\gamma} \partial_{\tau_j} v_j u_{i_1} d\omega & 0 \end{pmatrix}.$$

5.5.2 Second membre

$\underline{\mathbf{G}}_\gamma \in \mathbb{R}^{N_1}$ est tel que : $\underline{\mathbf{G}}_\gamma^{i_1} = (g, u_{i_1})_{0,\gamma}$. Ce calcul peut être approché directement, à l'aide d'un schéma d'intégration numérique, ou en utilisant la fonction d'interpolation P_1 de g . Soit

$\mathbb{M}_{\gamma,1} \in \mathbb{R}^{N_1 \times N_1}$ la matrice de masse pondérée P_1 telle que : $\forall i_1, j_1 \in I_1$,

$$\mathbb{M}_{\gamma,1}^{i_1,j_1} = (u_{i_1}, u_{j_1})_{0,\gamma} = \sum_{T_l | S_{i_1}, S_{j_1} \in T_l} \int_{T_l} r^{2\gamma} u_{i_1} u_{j_1} d\omega.$$

Dans le second cas, en déduit : $\underline{\mathbf{G}}_\gamma = \mathbb{M}_{\gamma,1} \underline{\mathbf{g}}$.

5.6 Analyse d'erreur

Rappelons que dans le cas quasi-électrostatique, le multiplicateur de Lagrange est nul : $p = 0$.

Lemme 5.5 *Pour les trois méthodes étudiées, on a l'estimation d'erreur suivante sur le multiplicateur de Lagrange : Il existe une constante C_p qui ne dépend que du domaine et des données telle que, pour h suffisamment petit :*

$$\|p_h\|_Q \leq C_p h^\tau, \quad (5.8)$$

où τ est le taux de convergence du calcul de $\mathbf{E}$ sans multiplicateur de Lagrange.

Ainsi, même si p_h ne s'annule pas, il prend des petites valeurs. La démonstration suivante est due à P. Ciarlet, Jr. [35].

DÉMONSTRATION. Le couple d'espaces $(\mathbf{X}, Q)$ peut représenter $(\mathbf{X}_\mathbf{E}, L^2(\omega))$, $(\mathbf{X}_\mathbf{E}^0, L^2(\omega))$ ou $(\mathbf{X}_{\mathbf{E},\gamma}^0, L_\gamma^2(\omega))$. Soit $(\mathbf{X}_h, Q_h)$ le couple d'espaces discretisés respectif. Φ représente l'espace Φ_D ou Φ_γ (méthode à poids). Soit $\mathcal{L}$ la forme linéaire continue suivante : $\forall \mathbf{F} \in \mathbf{X}$,

$$\mathcal{L}(\mathbf{F}) = \begin{cases} \mathcal{L}_\mathbf{E}(\mathbf{F}) & \text{si } \mathbf{X} = \mathbf{X}_\mathbf{E}, \\ \mathcal{L}_0(\mathbf{F}) - \mathcal{A}_0(\mathbf{e}, \mathbf{F}) & \text{si } \mathbf{X} = \mathbf{X}_\mathbf{E}^0, \\ \mathcal{L}_\gamma(\mathbf{F}) - \mathcal{A}_\gamma(\mathbf{e}, \mathbf{F}) & \text{si } \mathbf{X} = \mathbf{X}_{\mathbf{E},\gamma}^0. \end{cases}$$

g_X vaut g pour $\mathbf{X} = \mathbf{X}_\mathbf{E}$, et g^0 sinon. Considérons les problèmes suivants :

Le problème direct continu :

Trouver $\mathbf{E} \in \mathbf{X}$ tel que :

$$(\mathbf{E}, \mathbf{F})_\mathbf{X} = \mathcal{L}(\mathbf{F}), \quad \forall \mathbf{F} \in \mathbf{X}. \quad (5.9)$$

Le problème mixte continu :

Trouver $(\mathbf{E}, p) \in (\mathbf{X}, Q)$ tel que :

$$(\mathbf{E}, \mathbf{F})_\mathbf{X} + (p, \operatorname{div} \mathbf{F})_Q = \mathcal{L}(\mathbf{F}), \quad \forall \mathbf{F} \in \mathbf{X}, \quad (5.10)$$

$$(\operatorname{div} \mathbf{E}, q)_Q = (g_X, q)_Q, \quad \forall q \in Q. \quad (5.11)$$

Le problème mixte discrétisé :

Trouver $(\mathbf{E}_h, p_h) \in (\mathbf{X}_h, Q_h)$ tel que :

$$(\mathbf{E}_h, \mathbf{F}_h)_\mathbf{X} + (p_h, \operatorname{div} \mathbf{F}_h)_Q = \mathcal{L}(\mathbf{F}_h), \quad \forall \mathbf{F}_h \in \mathbf{X}_h, \quad (5.12)$$

$$(\operatorname{div} \mathbf{E}_h, q_h)_Q = (g_X, q_h)_Q, \quad \forall q_h \in Q_h. \quad (5.13)$$

Soit $\phi \in \Phi$ tel que $\Delta\phi = p_h$ dans ω . Considérons $\mathbf{v}^*$ tel que $\mathbf{v}^* = \mathbf{grad} \phi$. On a donc $\mathbf{v}^* \in \mathbf{X}$, tel que :

$$\begin{cases} \operatorname{div} \mathbf{v}^* &= p_h, \text{ dans } \omega, \\ \operatorname{rot} \mathbf{v}^* &= 0, \text{ dans } \omega, \\ \mathbf{v}^* \cdot \boldsymbol{\tau}|_{\partial\omega} &= 0, \text{ sur } \partial\omega. \end{cases}$$

$\mathbf{v}^*$ est donc solution du problème :

Trouver $\mathbf{v}^* \in \mathbf{X}$ tel que :

$$(\mathbf{v}^*, \mathbf{F})_{\mathbf{X}} = (p_h, \operatorname{div} \mathbf{F})_Q, \forall \mathbf{F} \in \mathbf{X}. \quad (5.14)$$

• Soit $\mathbf{v}_h^*$ la solution du problème discrétisé correspondant :

Trouver $\mathbf{v}_h^* \in \mathbf{X}_h$ tel que :

$$(\mathbf{v}_h^*, \mathbf{F}_h)_{\mathbf{X}} = (p_h, \operatorname{div} \mathbf{F}_h)_Q, \forall \mathbf{F}_h \in \mathbf{X}_h. \quad (5.15)$$

On a donc :

$$\begin{aligned} \|\mathbf{v}^* - \mathbf{v}_h^*\|_{\mathbf{X}} &\leq \inf_{\mathbf{F}_h \in \mathbf{X}_h} \|\mathbf{v}^* - \mathbf{v}_h^*\|_{\mathbf{X}_h}, \text{ d'après le lemme de Céa ??,} \\ &\leq C \|p_h\|_Q h^\tau, \end{aligned} \quad (5.16)$$

d'après l'analyse d'erreur pour le problème discret, où C est une constante qui ne dépend que du domaine ω .

• Choisissons de prendre $\mathbf{F}_h = \mathbf{E}_h$ dans (5.15), et $q_h = p_h$ dans (5.13). On a alors :

$$(\mathbf{v}_h^*, \mathbf{E}_h)_{\mathbf{X}} = (p_h, \operatorname{div} \mathbf{E}_h)_Q = (g_X, p_h)_Q = (g_X, \operatorname{div} \mathbf{v}^*)_Q.$$

Lorsqu'on injecte $\mathbf{F}_h = \mathbf{v}_h^*$ dans (5.12) on obtient :

$$\begin{aligned} (\mathbf{E}_h, \mathbf{v}_h^*)_{\mathbf{X}} + (p_h, \operatorname{div} \mathbf{v}_h^*)_Q &= (f, \operatorname{rot} \mathbf{v}_h^*)_0 + (g, \operatorname{div} \mathbf{v}_h^*)_0 + \int_{\partial\omega} e \mathbf{v}_h^* \cdot \boldsymbol{\tau} \, d\sigma, \text{ si } \mathbf{X} = \mathbf{X}_{\mathbf{E}}, \\ &= (f^0, \operatorname{rot} \mathbf{v}_h^*)_0 + (g^0, \operatorname{div} \mathbf{v}_h^*)_Q, \text{ sinon.} \end{aligned}$$

En remplaçant le terme $(\mathbf{E}_h, \mathbf{v}_h^*)_{\mathbf{X}}$ par $(g_X, \operatorname{div} \mathbf{v}^*)_Q$ et en le transférant au second membre, on en déduit :

$$\begin{aligned} (p_h, \operatorname{div} \mathbf{v}_h^*)_Q &= (f, \operatorname{rot} \mathbf{v}_h^*)_0 + (g, \operatorname{div} (\mathbf{v}_h^* - \mathbf{v}^*))_0 + \int_{\partial\omega} e \mathbf{v}_h^* \cdot \boldsymbol{\tau} \, d\sigma, \text{ si } \mathbf{X} = \mathbf{X}_{\mathbf{E}}, \\ &= (f^0, \operatorname{rot} \mathbf{v}_h^*)_0 + (g^0, \operatorname{div} (\mathbf{v}_h^* - \mathbf{v}^*))_Q, \text{ sinon.} \end{aligned}$$

• Comme par définition de $\mathbf{v}^*$, $\operatorname{rot} \mathbf{v}^* = 0$, et $\mathbf{v}^* \cdot \boldsymbol{\tau}|_{\partial\omega} = 0$, on a :

$$\begin{aligned} (p_h, \operatorname{div} \mathbf{v}_h^*)_Q &= (f, \operatorname{rot} (\mathbf{v}_h^* - \mathbf{v}^*))_0 + (g, \operatorname{div} (\mathbf{v}_h^* - \mathbf{v}^*))_0 \\ &\quad + \int_{\partial\omega} e (\mathbf{v}_h^* - \mathbf{v}^*) \cdot \boldsymbol{\tau} \, d\sigma, \text{ si } \mathbf{X} = \mathbf{X}_{\mathbf{E}}, \\ &= (f^0, \operatorname{rot} (\mathbf{v}_h^* - \mathbf{v}^*))_0 + (g^0, \operatorname{div} (\mathbf{v}_h^* - \mathbf{v}^*))_Q, \text{ sinon,} \end{aligned}$$

ce qui se réécrit : $(p_h, \operatorname{div} \mathbf{v}_h^*)_Q = \mathcal{L}(\mathbf{v}_h^* - \mathbf{v}^*)$. Or, on a :

$$\begin{aligned} \|p_h\|_Q^2 &= (p_h, \operatorname{div} \mathbf{v}^*)_Q = (p_h, \operatorname{div} (\mathbf{v}^* - \mathbf{v}_h^*))_Q + (p_h, \operatorname{div} \mathbf{v}_h^*)_Q \\ &= (p_h, \operatorname{div} (\mathbf{v}^* - \mathbf{v}_h^*))_Q + \mathcal{L}(\mathbf{v}_h^* - \mathbf{v}^*). \end{aligned}$$

On en déduit les inégalités suivantes (avec : $\|\mathcal{L}\|_{\mathbf{X}'} = \sup_{\mathbf{F} \in \mathbf{X}} |\mathcal{L}(\mathbf{F})|$) :

$$\begin{aligned} \|p_h\|_Q^2 &\leq \|p_h\|_Q \|\operatorname{div} (\mathbf{v}^* - \mathbf{v}_h^*)\|_Q + \|\mathcal{L}\|_{\mathbf{X}'} \|\mathbf{v}^* - \mathbf{v}_h^*\|_{\mathbf{X}}, \\ &\quad \text{d'après l'inégalité de Cauchy-Schwarz,} \\ &\leq (\|p_h\|_Q + \|\mathcal{L}\|_{\mathbf{X}'}) \|\mathbf{v}^* - \mathbf{v}_h^*\|_{\mathbf{X}} \\ &\leq C \|p_h\|_Q h^\tau (\|p_h\|_Q + \|\mathcal{L}\|_{\mathbf{X}'}), \text{ d'après (5.16).} \end{aligned}$$

D'où, en simplifiant par $\|p_h\|_Q$, on a : $\|p_h\|_Q \leq C h^\tau (\|p_h\|_Q + \|\mathcal{L}\|_{\mathbf{X}'})$.

En regroupant les $\|p_h\|_Q$ dans la partie gauche de l'inégalité, on en déduit, pour h suffisamment petit, tel que $C h^\tau < 1/2$:

$$\begin{aligned} \|p_h\|_Q &\leq \frac{1}{2} \|p_h\|_Q + C h^\tau \|\mathcal{L}\|_{\mathbf{X}'}, \text{ soit :} \\ \|p_h\|_Q &\leq 2 C h^\tau \|\mathcal{L}\|_{\mathbf{X}'}. \end{aligned}$$

D'où le résultat, avec $C_p = 2 \|\mathcal{L}\|_{\mathbf{X}'}$. □

Notons provisoirement $\mathbf{E}_h^* \in \mathbf{X}_h$ la solution du problème discrétisé direct :

$$\forall \mathbf{F}_h \in \mathbf{X}_h, \mathcal{A}(\mathbf{E}_h^*, \mathbf{F}_h) = \mathcal{L}(\mathbf{F}_h).$$

Soit $(\mathbf{E}_h, p_h) \in \mathbf{X}_h \times Q_h$ tel que :

$$\begin{aligned} \mathcal{A}(\mathbf{E}_h, \mathbf{F}_h) + \mathcal{B}(\mathbf{F}_h, p_h) &= \mathcal{L}(\mathbf{F}_h), \forall \mathbf{F}_h \in \mathbf{X}, \\ \mathcal{B}(\mathbf{E}_h, q_h) &= \mathcal{G}(q_h), \forall q_h \in Q. \end{aligned}$$

Lemme 5.6 *On a l'estimation suivante :*

$$\|\mathbf{E}_h - \mathbf{E}_h^*\|_{\mathbf{X}} \leq \|p_h\|_Q. \quad (5.17)$$

DÉMONSTRATION. On a : $\mathcal{A}(\mathbf{E}_h - \mathbf{E}_h^*, \mathbf{F}_h) + \mathcal{B}(\mathbf{F}_h, p_h) = 0$. Prenons $\mathbf{F}_h = \mathbf{E}_h - \mathbf{E}_h^*$. D'après l'inégalité de Cauchy-Schwarz, on obtient :

$$\|\mathbf{E}_h - \mathbf{E}_h^*\|_{\mathbf{X}}^2 \leq \|p_h\|_Q \|\operatorname{div}(\mathbf{E}_h - \mathbf{E}_h^*)\|_Q \leq \|p_h\|_Q \|\mathbf{E}_h - \mathbf{E}_h^*\|_{\mathbf{X}}.$$

□

Théorème 5.7 *Pour les trois méthodes étudiées, on a les mêmes estimations d'erreurs avec ou sans multiplicateur de Lagrange.*

DÉMONSTRATION. D'après l'inégalité triangulaire, on a en effet :

$$\begin{aligned} \|\mathbf{E} - \mathbf{E}_h\|_{\mathbf{X}} &= \|(\mathbf{E} - \mathbf{E}_h^*) + (\mathbf{E}_h^* - \mathbf{E}_h)\|_{\mathbf{X}} \\ &\leq \|\mathbf{E} - \mathbf{E}_h^*\|_{\mathbf{X}} + \|\mathbf{E}_h^* - \mathbf{E}_h\|_{\mathbf{X}}. \end{aligned}$$

D'où, d'après la section (5.5) : $\|\mathbf{E} - \mathbf{E}_h\|_{\mathbf{X}} \leq C h^\tau + C_p h^\tau = (C + C_p) h^\tau$. □

5.7 Optimalité de l'algorithme d'Uzawa

La matrice $\mathbb{C}\mathbb{A}^{-1}\mathbb{C}^T$ est symétrique et définie positive car $\mathcal{A}$ est une forme bilinéaire coercitive sur $\mathbf{X} \times \mathbf{X}$, et $\mathbb{C}$ est de rang maximal. On peut donc appliquer donc *l'algorithme du gradient conjugué* (voir la section 15.4) pour calculer $\underline{\mathbf{p}}$.

Afin de diminuer le nombre d'itérations de l'algorithme, qui est lié au nombre de conditionnement de $\mathbb{C}\mathbb{A}^{-1}\mathbb{C}^T$, on utilise comme matrice de préconditionnement la matrice de masse $\mathbb{M}$. Celle-ci

est équivalente à $\mathbb{C} \mathbb{A}^{-1} \mathbb{C}^T$ si la condition inf-sup discrète est uniforme (comme prouvé dans la section 15.7). Dans ce cas, le nombre d'itérations est indépendant de h . Pour réduire le coût d'une itération, on remplace la matrice de masse pleine par la matrice de masse réduite.

Les résultats illustrent le fait que cet algorithme est optimal pour résoudre le problème mixte lorsque la condition inf-sup discrète uniforme est vérifiée (espaces $\mathbf{X}_{\mathbf{E}}$ et $\mathbf{X}_{\mathbf{E}}^0$). Ainsi, pour l'espace $\mathbf{X}_{\mathbf{E},\gamma}^0$, l'algorithme d'Uzawa n'est pas optimal car la condition inf-sup discrète dépend du pas du maillage. Notons pour finir que dans la mesure où d'une part la dépendance numérique est faible (voir les expériences numériques du chapitre 6), et d'autre part, le multiplicateur de Lagrange est petit, ce n'est pas pénalisant, si toutefois on préconditionne le système avec la matrice de masse à poids réduite.

Chapitre 6

Résultats numériques du problème statique $2D$

6.1 Introduction

Les méthodes pour la résolution du problème quasi-électrostatique (2.1)-(2.3) ont été programmées en Matlab. Nous avons choisis ce logiciel de calcul scientifique car la syntaxe est matricielle, et on dispose de plusieurs fonctions spécialisées pour le calcul (inversion de matrices, recherche de valeurs propres...) et la représentation.

Le domaine d'étude ω pour les tests numériques est un polygone non-convexe en forme de L (figure 6.1). Il ne comporte donc qu'un seul coin rentrant à angle droit, tel que $\alpha = 2/3$ (voir les sections 4.11 et 5.6). Nous avons travaillé avec cinq différents maillages, générés par un maillage initial, raffiné cinq fois mais non- homogénéisé. Les caractéristiques de ces maillages sont les suivantes :

Maillage	1	2	3	4	5
h , pas du maillage	2.00×10^{-1}	1.00×10^{-1}	5.00×10^{-2}	2.50×10^{-2}	1.25×10^{-2}
Nombre de triangles	616	2 464	9 856	39 42	158 720
Nombre de noeuds P_1	351	1 317	5 097	20 049	79 521
Nombre de noeuds P_2	1 317	5 097	20 049	79 521	—

TAB. 6.1 – Caractéristiques des maillages.

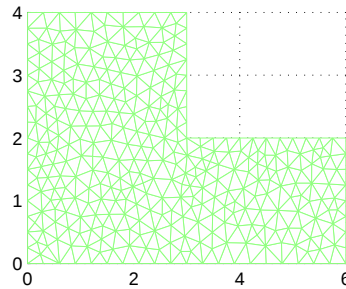


FIG. 6.1 – Représentation du domaine d'étude, et du maillage 1.

Nous utiliserons les notations suivantes :

- $\mathbf{E}_\phi$, l'approximation du champ électrique calculée à l'aide des potentiels ϕ_D et ϕ_N ,
- $\mathbf{E}_{nat}$, l'approximation du champ électrique calculée dans $\mathbf{X}_{\mathbf{E}}$,
- $\mathbf{E}_\lambda$, l'approximation du champ électrique calculée par la λ -approche, et λ_h , l'approximation de λ ,
- $\mathbf{E}_\perp$, l'approximation du champ électrique calculée dans $\mathbf{X}_{\mathbf{E}}^{0,R}$,
- $\mathbf{E}_0$, l'approximation du champ électrique calculée dans $\mathbf{X}_{\mathbf{E}}^0$,
- $\mathbf{E}_\gamma$, l'approximation du champ électrique calculée dans $\mathbf{X}_{\mathbf{E},\gamma}^0$.

$\mathbf{X}$ désigne l'un ou l'autre des espaces considérés, et $\mathbf{E}_h$ est de façon générale l'une ou l'autre des approximations associées. On appelle $\mathbf{E}$ le champ électrique solution exacte du problème. On notera la méthode avec condition aux limites naturelles : MCLN et la méthode de régularisation à poids : MRP.

Pour la MRP, nous avons pris $\gamma = 0.99$. Ainsi, le taux de convergence théorique (avec les bonnes hypothèses sur les données) de cette méthode est d'environ 0.66. Pour la λ -approche ou la MCSO, le taux de convergence théorique (avec les bonnes hypothèses sur les données) est de l'ordre de 0.33.

Afin d'étudier le taux de convergence des méthodes, nous calculons le taux de décroissance de l'erreur en norme $\mathbf{X}$ ou $\mathbf{L}^2(\omega)$ entre deux maillages consécutifs :

$$\text{Pour } i \in \{1, \dots, 4\}, \tau_i = - \frac{\log_{10}(\|\mathbf{E}_{h_{i+1}} - \mathbf{E}\|_{\mathbf{X}} \text{ ou } 0) - \log_{10}(\|\mathbf{E}_{h_i} - \mathbf{E}\|_{\mathbf{X}} \text{ ou } 0)}{\log_{10}(h_{i+1}) - \log_{10}(h_i)}$$

Nous utiliserons la notation suivante : $\tau_{meth,norm}$ les taux de convergence de la méthode *meth* en norme *norm*. Notons que pour étudier l'erreur en norme $\mathbf{X}$, nous n'avons besoin que des données et de $\mathbf{E}_h$. En revanche, pour calculer l'erreur en norme $\mathbf{L}^2(\omega)$, il faut disposer d'une approximation discrète du champ électrique exact, ce qui n'est pas le cas en général.

Dans la section 6.2, on présente les résultats du calcul de $\mathbf{E}_h$ par les éléments finis P_1 de Lagrange continus composante par composante. Dans la section 6.3, on présente les résultats du calcul de $\mathbf{E}$ par les éléments finis mixtes P_1 - P_2 de Lagrange continus composante par composante.

6.2 Calcul direct

Afin de tester les méthodes directes en P_1 , nous choisissons les trois cas-tests suivants :

- **Cas test 1** : g sinusoïdale en x et y , $f = 0$, et $e = 0$. On connaît alors la solution analytique de (2.1)-(2.3), qui est de plus régulière ; on a donc $\lambda_D = 0$.
- **Cas test 2** : $g = s_D$, $f = 0$ et $e = 0$. La solution de (2.1)-(2.3) est $-\mathbf{grad} \varphi_D = -\mathbf{grad} \tilde{\varphi}_D + \lambda_D \mathbf{x}^P$. On ne connaît pas $\mathbf{grad} \tilde{\varphi}_D$ exactement, mais on peut l'approcher en calculant $\tilde{\varphi}_D$ par les éléments finis P_1 de Lagrange. $-\mathbf{grad} \tilde{\varphi}_D$ est alors $(P_0)^2$, c'est-à-dire constant par triangle.
- **Cas test 3** : $f = 0$, $g = 0$ et $e = \mathbf{x}^P \cdot \boldsymbol{\tau}_{|\partial\omega}$. Dans ce cas, la condition aux limites tangentielle n'est pas nulle. La solution de (2.1)-(2.3) est connue et vaut $\mathbf{E} = \mathbf{x}^P$.

6.2.1 Cas régulier

On résout (2.1)-(2.3) avec : $f = 0$, $g = 2\pi \sin \pi x \sin \pi y$ dans ω , $e = 0$ sur la frontière $\partial\omega$. Le champ $\mathbf{E}$ solution est alors : $\mathbf{E} = - \begin{pmatrix} \cos \pi x \sin \pi y \\ \sin \pi x \cos \pi y \end{pmatrix}$. Dans ce paragraphe, $\mathbf{E}_0$ représente l'approximation de $\mathbf{E}$ calculé dans $\mathbf{X}_{\mathbf{E}}^{0,R}$ sans complément singulier.

• Dans les tableaux 6.2, on a représenté à gauche les erreurs en norme $\mathbf{L}^2(\Omega)$ et à droite les taux de convergence de ces erreurs. Le calcul de l'erreur est fait à l'aide d'un schéma d'intégration numérique à sept points par triangle, sachant qu'on connaît la solution exacte en tous points. On observe que la méthode qui converge le moins bien est le calcul du champ par les potentiels. En effet, $\mathbf{E}_\phi$ est P_0 alors que les autres approximations sont P_1 . Comme λ est nul, la méthode du complément singulier revient à effectuer le calcul du champ quasi-électrostatique dans $\mathbf{X}_{\mathbf{E}}^{0,R}$. La différence (inférieure à 0.01%) est due à l'erreur sur le calcul de λ : λ_h n'est pas exactement nul. Notons de plus qu'on obtient exactement le même résultat numérique avec la λ -approche et la MCSO.

• Dans les tableaux 6.3, on a représenté à gauche les erreurs en norme $\mathbf{X}$ et à droite les taux de convergence de ces erreurs. De même, le calcul de l'erreur est fait à l'aide d'un schéma d'intégration numérique à sept points par triangle, sachant qu'on connaît g en tous points. Pour les méthodes de calcul direct, les taux de convergence sont de l'ordre de 1, ce qui correspond au taux attendu lorsqu'on approche un champ régulier avec les éléments finis de Lagrange P_1 . Le calcul du champ dans $\mathbf{X}_{\mathbf{E}}^0$ (avec ou sans complément singulier) converge mieux que le calcul dans $\mathbf{X}_{\mathbf{E}}$, car on tient compte explicitement de la condition limite (comme on le verra dans les tableaux suivants). Notons qu'on a un facteur deux entre le taux de convergence de la λ -approche en norme $\mathbf{X}_{\mathbf{E}}$ et celui en norme $\mathbf{L}^2(\Omega)$, comme c'est attendu par la théorie (voir le paragraphe 4.11.2). Pour la MRP, on remarque qu'on a quasiment la même erreur en norme $\mathbf{X}_{\mathbf{E}}^0$ et en norme $\mathbf{X}_{\mathbf{E},\gamma}^0$, toujours car la donnée et la solution sont régulières. Il semble qu'on ait aussi un facteur deux entre le taux de convergence de la MRP en norme $\mathbf{X}_{\mathbf{E}}$ et celui en norme $\mathbf{L}^2(\Omega)$.

• Dans les tableaux 6.4, on a représenté à gauche les erreurs sur la composante tangentielle sur $\partial\omega$, en norme $L^2(\partial\omega)$ et à droite les taux de convergence de ces erreurs. La quantité $\|\mathbf{E}_\lambda \cdot \boldsymbol{\tau}\|_{0,\partial\omega}$ correspond ici à l'erreur d'interpolation entre $\lambda \mathbf{x}^P \cdot \boldsymbol{\tau}|_{\partial\omega}$ et $\lambda_h \pi_h(\mathbf{x}^P \cdot \boldsymbol{\tau}|_{\partial\omega})$.

Maillage	1	2	3	4	5
$\ \mathbf{E}_\phi - \mathbf{E}\ _0$	3.0%	0.9%	0.3%	0.0%	0.0%
$\ \mathbf{E}_{nat} - \mathbf{E}\ _0$	4.0%	1.4%	0.4%	0.1%	0.0%
$\ \mathbf{E}_\lambda - \mathbf{E}\ _0$	2.5%	0.6%	0.2%	0.0%	0.0%
$\ \mathbf{E}_0 - \mathbf{E}\ _0$	2.4%	0.6%	0.2%	0.0%	0.0%
$\ \mathbf{E}_\gamma - \mathbf{E}\ _0$	1.2%	0.3%	0.1%	0.0%	0.0%

τ_i	1/2	2/3	3/4	4/5
$\tau_{\phi,0}$	1.80	1.75	1.68	1.62
$\tau_{nat,0}$	1.55	1.82	1.94	1.93
$\tau_{\lambda,0}$	1.99	1.99	2.00	2.01
$\tau_{\gamma,0}$	1.93	1.94	1.94	1.96

TAB. 6.2 – Cas régulier : Erreurs $\mathbf{L}^2(\omega)$ et taux de convergence.

Maillage	1	2	3	4	5
$\ \mathbf{E}_\phi - \mathbf{E}\ _{\mathbf{X}_E}$	37.0%	20.3%	11.7%	7.1%	4.5%
$\ \mathbf{E}_{nat} - \mathbf{E}\ _{\mathbf{X}_E}$	31.4%	15.8%	7.9%	4.0%	2.0%
$\ \mathbf{E}_\lambda - \mathbf{E}\ _{\mathbf{X}_E}$	32.3%	16.0%	8.0%	4.0%	2.0%
$\ \mathbf{E}_\perp - \mathbf{E}\ _{\mathbf{X}_E}$	32.3%	16.0%	8.0%	4.0%	2.0%
$\ \mathbf{E}_0 - \mathbf{E}\ _{\mathbf{X}_E}$	32.3%	16.0%	7.9%	4.0%	2.0%
$\ \mathbf{E}_\gamma - \mathbf{E}\ _{\mathbf{X}_{E,\gamma}^0}$	30.2%	14.8%	7.4%	3.7%	1.8%
$\ \mathbf{E}_\gamma - \mathbf{E}\ _{\mathbf{X}_E^0}$	33.0%	16.1%	8.0%	4.0%	2.0%

τ_i	1/2	2/3	3/4	4/5
$\tau_{\phi, \mathbf{X}_E}$	0.87	0.79	0.72	0.64
$\tau_{nat, \mathbf{X}_E}$	0.99	0.99	1.00	1.00
$\tau_{\lambda, \mathbf{X}_E}$	1.01	1.01	1.00	1.00
$\tau_{\gamma, \mathbf{X}_{E,\gamma}^0}$	1.02	1.01	1.00	1.00
$\tau_{\gamma, \mathbf{X}_E^0}$	1.02	1.01	1.00	1.00

TAB. 6.3 – Cas régulier : Erreurs $\mathbf{X}$ et taux de convergence.

Maillage	1	2	3	4	5
$\ \mathbf{E}_\phi - \mathbf{E}\ _{0, \partial\omega}$	2.6%	0.1%	0.0%	0.0%	0.0%
$\ \mathbf{E}_{nat} - \mathbf{E}\ _{0, \partial\omega}$	5.6%	2.0%	0.6%	0.2%	0.0%
$\ \mathbf{E}_\lambda - \mathbf{E}\ _{0, \partial\omega}$	0.0%	0.0%	0.0%	0.0%	0.0%

τ_i	1/2	2/3	3/4	4/5
$\tau_{\phi, \partial\omega}$	1.47	1.47	1.48	1.50
$\tau_{nat, \partial\omega}$	1.36	1.70	1.82	1.85
$\tau_{\lambda, \partial\omega}$	4.12	4.03	4.01	4.00

TAB. 6.4 – Cas régulier : Erreurs $L^2(\partial\omega)$ et taux de convergence.

6.2.2 Premier cas singulier

On résout (2.1)-(2.3) avec cette fois : $f = 0$, $g = s_D$, $e = 0$.

Le champ $\mathbf{E}$ solution est alors : $\mathbf{E} = -\mathbf{grad} \varphi_D = -\mathbf{grad} \tilde{\varphi}_D + \lambda \mathbf{x}^P$. On ne connaît donc pas exactement le champ exact, aussi on prend comme approximation du champ exact pour calculer l'erreur en norme $\mathbf{L}^2(\omega)$: $\mathbf{E}_D = -\mathbf{grad} \tilde{\varphi}_{D,h} + \lambda_h \mathbf{x}^P$, où $\tilde{\varphi}_{D,h}$ approché par les éléments finis P_2 , et $-\mathbf{grad} \tilde{\varphi}_{D,h}$ est obtenu par projection P_1 - P_2 . Rappelons que dans notre situation, on a $\alpha = 2/3$, et donc $1 - \alpha = 2\alpha - 1$. Ainsi, $s_D \in H^{2\alpha-1-\varepsilon}(\omega)$ pour tous $0 < \varepsilon < 2\alpha - 1$ (d'après [67], thm. 1.2.18). On est dans le cadre des hypothèses du paragraphe 4.11.2 : le taux de convergence attendu pour la λ -approche est bien de l'ordre de 0.33. En revanche, on ne peut pas savoir quel est le taux de convergence attendu pour la méthode à poids car on n'a pas l'hypothèse $\text{div } \mathbf{E} = 0$ requise pour appliquer la proposition 4.19. De même que dans le cas régulier, la λ -approche et la MCSO donnent les mêmes résultats numériques.

- Dans les tableaux 6.5, on a représenté à gauche les erreurs en norme $\mathbf{L}^2(\Omega)$, calculées à l'aide d'un schéma d'intégration numérique à sept points (pour approcher au mieux $\mathbf{x}^P$), et à droite les taux de convergence de ces erreurs. On remarque que cette fois-ci, la donnée n'étant pas régulière,

il est crucial d'utiliser le complément singulier pour approcher correctement $\mathbf{E}$ dans $\mathbf{X}_{\mathbf{E},h}^0$. Sinon, il n'y a pas convergence vers la solution exacte (voir la quatrième ligne : $\|\mathbf{E}_0 - \mathbf{E}_D\|_0$), comme c'est annoncé par la théorie. En fait, $\mathbf{E}_0$ converge vers une solution régulière autre que $\mathbf{E}$.

- Dans les tableaux 6.6, on a représenté à gauche les erreurs en norme $\mathbf{X}$ et à droite les taux de convergence de ces erreurs. De même, le calcul de l'erreur est fait à l'aide d'un schéma d'intégration numérique à sept points par triangle (pour approcher au mieux s_D^P). Pour la λ -approche, on a un taux de convergence de l'ordre de 0.33, ce qui correspond à la théorie. Le taux de convergence en norme $\mathbf{L}^2(\omega)$ est de l'ordre de 0.72, ce qui est meilleur que deux fois meilleur que le taux attendu, de l'ordre de 0.66, mais comme nous ne disposons pas de la solution analytique, cela correspond au taux de convergence entre deux calculs différents. Notons que, contrairement au cas où la solution exacte est régulière, la MRP ne converge pas en norme $\mathbf{X}_{\mathbf{E}}^0$.

- Détaillons les erreurs en norme $\mathbf{X}$. Dans les tableaux 6.7, on donne à gauche les erreurs de $\|\operatorname{div} \mathbf{E}_{\mathbf{X}} - s_{D,h}\|_{0(\gamma)}$, de $\|\operatorname{rot} \mathbf{E}_{\mathbf{X}}\|_0$ et de $\|\mathbf{E}_{\mathbf{X}} \cdot \boldsymbol{\tau}\|_{0,\partial\omega}$; et à droite les taux de convergence correspondants. On remarque que pour les trois méthodes, l'erreur est moins forte sur le rotationnel que sur la divergence. Cela est dû au fait que nous avons choisi une solution à rotationnel nul. De même que dans l'exemple du précédent paragraphe, pour la méthode du complément singulier, l'erreur $\|\mathbf{E}_{\lambda} \cdot \boldsymbol{\tau}\|_{0,\partial\omega}$ est simplement une erreur d'interpolation. On remarque en revanche que pour la MCLN, l'erreur sur la composante tangentielle à $\partial\omega$ est médiocre et surtout ne décroît pas. En fait, on ne contrôle pas la composante normale au bord de $\mathbf{E}_{nat}$, qui est singulière au voisinage du coin, comme c'est illustré sur la figure 6.2. En effet, sur cette figure, on a représenté $|\mathbf{E}_{nat} \cdot \boldsymbol{\tau}|_{\partial\omega}$ arête par arête sur le bord $\partial\omega$. On remarque qu'au voisinage du coin rentrant (sur A^0 et A^Θ), $|\mathbf{E}_{nat} \cdot \boldsymbol{\tau}|_{\partial\omega}$ augmente brusquement. Cela est dû au fait que par continuité des éléments finis P_1 , $|\mathbf{E}_{nat} \cdot \boldsymbol{\tau}|_{A^0}$ tend vers $|\mathbf{E}_{nat} \cdot \boldsymbol{\nu}|_{A^\Theta}$ et $|\mathbf{E}_{nat} \cdot \boldsymbol{\tau}|_{A^\Theta}$ tend vers $|\mathbf{E}_{nat} \cdot \boldsymbol{\nu}|_{A^0}$ au voisinage du coin rentrant : la valeur de $\mathbf{E}_{nat} \cdot \boldsymbol{\tau}$ sur A^0 tend en son extrémité vers la valeur de $-\mathbf{E}_{nat} \cdot \boldsymbol{\nu}$ sur A^Θ . De même, la valeur de $\mathbf{E}_{nat} \cdot \boldsymbol{\tau}$ sur A^Θ tend en son extrémité vers la valeur de $\mathbf{E}_{nat} \cdot \boldsymbol{\nu}$ sur A^0 . Or, cette quantité prend des valeurs plus élevées que la moyenne.

Maillage	1	2	3	4	5
$\ \mathbf{E}_\phi - \mathbf{E}_D\ _0$	10.3%	6.2%	3.8%	2.3%	1.4%
$\ \mathbf{E}_{nat} - \mathbf{E}_D\ _0$	51.9%	49.1%	45.5%	41.7%	37.8%
$\ \mathbf{E}_\lambda - \mathbf{E}_D\ _0$	12.8%	7.6%	4.6%	2.8%	1.7%
$\ \mathbf{E}_0 - \mathbf{E}_D\ _0$	70.9%	73.3%	74.7%	75.5%	—
$\ \mathbf{E}_\gamma - \mathbf{E}_D\ _0$	56.8%	45.2%	34.5%	26.7%	21%

τ_i	1/2	2/3	3/4	4/5
$\tau_{\phi,0}$	0.73	0.72	0.70	0.69
$\tau_{nat,0}$	0.24	0.22	0.20	0.17
$\tau_{\lambda,0}$	0.75	0.73	0.72	0.70
$\tau_{\gamma,0}$	0.33	0.39	0.37	0.35

TAB. 6.5 – Cas singulier 1 : Erreurs $\mathbf{L}^2(\omega)$ et taux de convergence.

Maillage	1	2	3	4	5
$\ \mathbf{E}_{nat} - \mathbf{E}_D\ _{\mathbf{X}_E}$	44.2%	37.5%	32.3%	28.1%	24.9%
$\ \mathbf{E}_\lambda - \mathbf{E}_D\ _{\mathbf{X}_E}$	39.7%	31.4%	24.9%	19.8%	15.7%
$\ \mathbf{E}_\gamma - \mathbf{E}_D\ _{\mathbf{X}_{E,\gamma}^0}$	45.3%	35.3%	28.6%	23.6%	19.6%
$\ \mathbf{E}_\gamma - \mathbf{E}_D\ _{\mathbf{X}_E^0}$	76.3%	93.8%	—	—	—

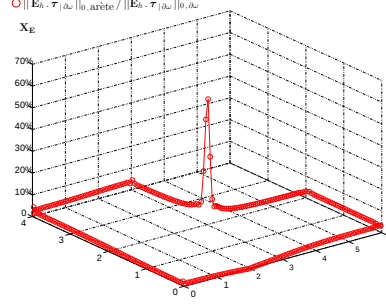
τ_i	1/2	2/3	3/4	4/5
$\tau_{nat,\mathbf{X}_E}$	0.24	0.22	0.20	0.17
$\tau_{\lambda,\mathbf{X}_E}$	0.34	0.33	0.33	0.33
$\tau_{\gamma,\mathbf{X}_{E,\gamma}^0}$	0.34	0.30	0.28	0.27

TAB. 6.6 – Cas singulier 1 : Erreurs $\mathbf{X}$ et taux de convergence.

Maillage	1	2	3	4	5
$\ \text{rot } \mathbf{E}_{nat}\ _0$	14.4%	13.3%	12.0%	10.7%	9.6%
$\ \text{div } \mathbf{E}_{nat} - s_D\ _0$	37.8%	30.2%	24.2%	19.5%	15.7%
$\ \mathbf{E}_{nat} \cdot \boldsymbol{\tau} _{\partial\omega}\ _{0,\partial\omega}$	17.7%	17.9%	17.7%	17.3%	16.7%
$\ \text{rot } \mathbf{E}_\lambda\ _0$	13.5%	10.5%	8.2%	6.4%	5.0%
$\ \text{div } \mathbf{E}_\lambda - s_D\ _0$	37.3%	29.6%	23.6%	18.7%	14.8%
$\ \mathbf{E}_\lambda \cdot \boldsymbol{\tau} _{\partial\omega}\ _{0,\partial\omega}$	0.0%	0.0%	0.0%	0.0%	0.0%
$\ \text{rot } \mathbf{E}_\gamma\ _0$	25.7%	21.9%	17.7%	14.4%	11.7%
$\ \text{div } \mathbf{E}_\gamma - s_D\ _{0,\gamma}$	37.3%	27.7%	22.4%	18.7%	15.8%

τ_i	1/2	2/3	3/4	4/5
$\tau_{nat,rot}$	0.11	0.15	0.17	0.16
$\tau_{nat,div}$	0.32	0.32	0.31	0.31
$\tau_{nat,\partial\omega}$	−0.02	0.02	0.03	0.05
$\tau_{\lambda,rot}$	0.36	0.36	0.35	0.35
$\tau_{\lambda,div}$	0.33	0.33	0.34	0.34
$\tau_{\lambda,\partial\omega}$	2.00	2.00	2.00	2.01
$\tau_{\gamma,rot}$	0.23	0.31	0.30	0.30
τ_{γ,div_γ}	0.43	0.31	0.26	0.24

TAB. 6.7 – Cas singulier 1 : Erreurs rot , div , $\mathbf{E} \cdot \boldsymbol{\tau}|_{\partial\omega}$, et taux de convergence.

FIG. 6.2 – Cas singulier 1 : $|\mathbf{E}_{nat} \cdot \boldsymbol{\tau}|_{\partial\omega}$ arête par arête.

6.2.3 Second cas singulier

On résout (2.1)-(2.3) avec : $f = 0, g = 0, e = \mathbf{x}^P \cdot \boldsymbol{\tau}$. Le champ $\mathbf{E}$ solution est alors : $\mathbf{E} = \mathbf{x}^P$ (donc $\lambda = 1$). On ne donne que l'erreur en norme $\mathbf{L}^2(\omega)$ et en norme $\mathbf{X}_E$ puisque nous venons d'analyser l'erreur au bord au premier cas-test singulier et que les comportements sont similaires entre les deux cas singuliers. Ce test permet de vérifier la robustesse du code : il montre qu'il fonctionne correctement pour des problèmes avec condition aux limites non homogènes, et qu'il donne de bons résultats même avec des données très singulières.

- Dans les tableaux 6.8, on a représenté à gauche les erreurs en norme $\mathbf{L}^2(\omega)$, calculées à l'aide d'un schéma d'intégration numérique à sept points (pour approcher au mieux $\mathbf{x}^P$).

- Dans les tableaux 6.9, on a représenté à gauche les erreurs en norme $\mathbf{X}$ et à droite les taux de convergence de ces erreurs. De même, le calcul de l'erreur est fait à l'aide d'un schéma d'intégration numérique à sept points par triangle. Pour la λ -approche, on a une erreur très faible car elle ne porte pratiquement que sur la composante tangentielle à $\partial\omega$.

- Dans le tableau 6.10, nous présentons les résultats des calculs proposés au paragraphe 2.4.3. Rappelons que $\mathbf{x}^P \cdot \boldsymbol{\tau}$ est nul sur les arêtes du coin rentrant, et régulier ailleurs. Nous pouvons donc utiliser le théorème 2.45. Soit $\mathbf{e}$ un relèvement régulier de e .

On note λ_{Gr} l'approximation de λ à la Grisvard : $\lambda_{Gr} = -((\text{rot } \mathbf{e}_h, s_{N,h})_0 + (\text{div } \mathbf{e}_h, s_{D,h})_0) / \pi$; et l'approximation de λ à la Nazarov-Plamenevsky : $\lambda_{NP} = \int_{\gamma_0} e_h s_{N,h} d\sigma / \pi$. Pour le calcul de λ_{Gr} , $\mathbf{e}_h$ est construit de sorte que :

$$\begin{cases} \mathbf{e}_h(S_i) &= 0, \forall i \in I_\omega, \\ \mathbf{e}_h \cdot \boldsymbol{\nu}_i &= 0, \forall i \in I_a, \\ \mathbf{e}_h \cdot \boldsymbol{\tau}_i &= \mathbf{x}^P(S_i) \cdot \boldsymbol{\tau}_i, \forall i \in I_a, \\ \mathbf{e}_h(S_i) &= \mathbf{x}^P(S_i), \forall i \in I_c. \end{cases}$$

Les parties régulières des fonctions singulières duales $\tilde{s}_D$ et $\tilde{s}_N$ sont approchées avec les éléments finis P_1 . Nous avons utilisé un schéma d'intégration à sept points par triangle pour le calcul de λ_{Gr} , et à quatre points par arête pour le calcul de λ_{NP} . Le premier schéma est exact pour les polynômes de degré cinq et le second pour les polynômes de degré trois. Ainsi, comme on peut l'observer sur le tableau 6.10, le calcul de λ_{Gr} est plus précis. On constate néanmoins que les deux calculs donnent de très bons résultats, avec un erreur inférieure à 5% dès le premier maillage. Pour les deux calculs, le taux de convergence est de l'ordre de 1.

Maillage	1	2	3	4	5
$\ \mathbf{E}_{nat} - \mathbf{x}^P\ _0$	69.6%	65.6%	55.2%	49.5%	44.5%
$\ \mathbf{E}_\lambda - \mathbf{x}^P\ _0$	0.4%	0.1%	0.0%	0.0%	0.0%
$\ \mathbf{E}_\gamma - \mathbf{x}^P\ _0$	54.7%	42.1%	32.1%	24.6%	19.0%

τ_i	1/2	2/3	3/4	4/5
$\tau_{nat,0}$	0.08	0.11	0.13	0.15
$\tau_{\lambda,0}$	1.38	1.36	1.35	1.34
$\tau_{\gamma,0}$	0.38	0.39	0.38	0.37

TAB. 6.8 – Cas singulier 2 : Erreurs $\mathbf{L}^2(\omega)$ et taux de convergence.

Maillage	1	2	3	4	5
$\ \mathbf{E}_{nat} - \mathbf{x}^P\ _{\mathbf{X}_E}$	31.4%	30.2%	28.9%	27.5%	26.1%
$\ \mathbf{E}_\lambda - \mathbf{x}^P\ _{\mathbf{X}_E}$	0.2%	0.1%	0.0%	0.0%	0.0%
$\ \mathbf{E}_\gamma - \mathbf{x}^P\ _{\mathbf{X}_{E,\gamma}^0}$	43.4%	36.1%	29.5%	23.9%	19.2%

τ_i	1/2	2/3	3/4	4/5
$\tau_{nat,\mathbf{X}_E}$	0.06	0.06	0.07	0.08
$\tau_{\lambda,\mathbf{X}_E}$	1.38	1.36	1.35	1.34
$\tau_{\gamma,\mathbf{X}_{E,\gamma}}$	0.27	0.29	0.30	0.32

TAB. 6.9 – Cas singulier 2 : Erreurs $\mathbf{X}$ et taux de convergence.

Maillage	1	2	3	4	5
λ_{Gr}	0.996	0.998	0.999	1.000	1.000
λ_{NP}	0.956	0.978	0.989	0.995	0.997

TAB. 6.10 – Cas singulier 2 : Deux calculs de λ .

6.3 Utilisation du multiplicateur de Lagrange

Dans cette section, nous donnons des résultats de l'approximation du champ quasi-électrostatique par les éléments finis de Taylor-Hood P_2 - P_1 , avec les données suivantes : $f = 0$, $g = 1$ dans ω et $e = 0$. Pour les trois méthodes, nous avons évalué le multiplicateur avec l'algorithme du gradient conjugué avec un critère d'arrêt égal à 10^{-5} de trois façons différents :

- Sans préconditionner la matrice $\mathbb{C}\mathbb{A}^{-1}\mathbb{C}^T$ (paragraphe 15.4.1),
- En préconditionnant cette matrice par la matrice de masse $L_{2(\gamma)}(\omega)$ P_1 (paragraphe 15.4.2),
- En préconditionnant cette matrice par la matrice de masse $L_{2(\gamma)}(\omega)$ P_1 réduite (sections 15.5, 15.6, et 15.7, partie IV).

Rappelons que si la condition inf-sup discrète est uniforme, la matrice $\mathbb{C}\mathbb{A}^{-1}\mathbb{C}^T$ est équivalente à la matrice de masse $\mathbb{M}$, elle-même équivalente à la matrice de masse diagonale $\tilde{\mathbb{M}}$. Ces trois méthodes de calcul donnent le même résultat pour p_h , on aura donc le même résultat pour l'approximation

$\mathbf{E}_h$, mais les algorithmes de résolution ne convergent pas à la même vitesse et n'ont pas le même coût calcul. Notons que comme p_h est petit, l'approximation de $\mathbf{E}_h$ correspond à celle en P_2 .

- Dans le tableau 6.11, on donne les erreurs d'approximation de $\mathbf{E}$ en P_2 - P_1 et en P_1 pour la MCLN, la MCSO et la MRP, ainsi que les taux de convergence correspondants. On remarque que c'est la MCSO qui donne le résultat le plus précis. Notons que bien que le taux de convergence de la MCLN soit faible, cette méthode donne approximation correcte sur le maillage le plus grossier. Pour la MRP, en P_2 - P_1 , on a un taux de convergence proche du taux théorique attendu. En revanche, en P_1 , le taux de convergence décroît de 0.70 à 0.45. Ainsi, l'utilisation d'éléments finis P_2 tend à stabiliser le taux de convergence. Le taux de convergence de la MCSO est supérieur à 0.33, que ce soit en P_2 - P_1 ou en P_1 : il y a un phénomène de super-convergence, certainement dû au fait que la solution est plus régulière que $\mathbf{H}^{2\alpha-1-\epsilon}(\omega)$. Pour les trois méthodes, notons que l'erreur en P_1 sur le maillage $i+1$ est moins bonne que l'erreur en P_2 sur le maillage i , pour un coût calcul similaire. En effet, les éléments finis P_2 permettent d'approcher avec plus de précision la partie régulière de la solution.

- Dans le tableau 6.12, on donne le nombre d'itérations N_{it} et le nombre de conditionnement κ de $\mathbb{C}\mathbf{A}^{-1}\mathbb{C}^T$ (κ est défini en 15.1, partie IV) pour les trois méthodes directes en P_2 - P_1 . Ce calcul n'est pas rédhibitoire pour la MCSO et la MCLN, mais on peut diminuer le nombre d'itérations. Pour la MRP, le calcul est beaucoup trop coûteux : le nombre d'itérations est très élevé et croît très rapidement.

- Dans le tableau 6.13, on donne le nombre d'itérations et le nombre de conditionnement de $\mathbb{M}^{-1}\mathbb{C}\mathbf{A}^{-1}\mathbb{C}^T$, où $\mathbb{M}$ est la matrice de masse exacte. Le nombre d'itérations est peu élevé, et surtout, il est indépendant de h pour la MCLN et la MCSO. Pour la MRP, $\log_{10}(\kappa)$ croît linéairement en fonction de $\log_{10}(h)$, avec une pente d'ordre 1. Pour les trois méthodes, le nombre d'itérations est raisonnable, néanmoins, chaque itération reste coûteuse car il faut inverser la matrice de masse $\mathbb{M}$.

Dans le tableau 6.14, on donne le nombre d'itérations et le nombre de conditionnement de $\tilde{\mathbb{M}}^{-1/2}\mathbb{C}\mathbf{A}^{-1}\mathbb{C}^T\tilde{\mathbb{M}}^{-1/2}$, où $\tilde{\mathbb{M}}$ est la matrice de masse réduite. Le nombre d'itérations est peu élevé, et le coût d'une itération est le même que sans préconditionnement. Ce préconditionnement est donc bon compromis.

Maillage	1	2	3	4	5
$\ \mathbf{E}_{nat} - \mathbf{E}\ _{\mathbf{X}_{\mathbf{E}}, P_2-P_1}$	11.2%	10.6%	10.0%	9.3%	—
$\ \mathbf{E}_{nat} - \mathbf{E}\ _{\mathbf{X}_{\mathbf{E}}, P_1}$	16.1%	13.2%	11.7%	10.8%	10.1%
$\ \mathbf{E}_{\perp} - \mathbf{E}\ _{\mathbf{X}_{\mathbf{E}}, P_2-P_1}$	3.7%	2.0%	1.2%	0.8%	—
$\ \mathbf{E}_{\perp} - \mathbf{E}\ _{\mathbf{X}_{\mathbf{E}}, P_1}$	13.9%	7.8%	4.3%	2.4%	1.4%
$\ \mathbf{E}_{\gamma} - \mathbf{E}\ _{\mathbf{X}_{\mathbf{E},\gamma}^0, P_2-P_1}$	7.3%	4.6%	2.9%	1.8%	—
$\ \mathbf{E}_{\gamma} - \mathbf{E}\ _{\mathbf{X}_{\mathbf{E},\gamma}^0, P_1}$	19.5%	12.0%	7.7%	5.3%	4.9%

τ_i	1/2	2/3	3/4	4/5
$\tau_{nat, \mathbf{X}_{\mathbf{E}}, P_2-P_1}$	0.07	0.09	0.09	—
$\tau_{nat, \mathbf{X}_{\mathbf{E}}, P_1}$	0.28	0.17	0.12	0.10
$\tau_{\perp, \mathbf{X}_{\mathbf{E}}, P_2-P_1}$	0.86	0.73	0.65	—
$\tau_{\perp, \mathbf{X}_{\mathbf{E}}, P_1}$	0.83	0.86	0.84	0.78
$\tau_{\gamma, \mathbf{X}_{\mathbf{E},\gamma}^0, P_2-P_1}$	0.67	0.67	0.69	—
$\tau_{\gamma, \mathbf{X}_{\mathbf{E},\gamma}^0, P_1}$	0.70	0.64	0.54	0.45

TAB. 6.11 – Problème mixte : Erreurs $\mathbf{X}$ et taux de convergence.

Maillage	1	2	3	4
MCLN, N_{it}	21	22	22	23
MCLN, κ	25	28	30	31
MCSO N_{it}	37	37	40	41
MCSO κ	120	138	151	157
MRP N_{it}	137	364	> 600	—
MRP κ	$> 10^4$	$> 10^5$	$> 10^6$	—

TAB. 6.12 – Problème mixte, sans préconditionnement.

Maillage	1	2	3	4
MCLN, N_{it}	4	3	3	3
MCLN, κ	1	1	1	1
MCSO N_{it}	7	7	7	7
MCSO κ	5	5	5	5
MRP N_{it}	9	12	15	19
MRP κ	6	12	26	56

TAB. 6.13 – Problème mixte, préconditionnement avec la matrice de masse exacte.

Maillage	1	2	3	4
MCLN, N_{it}	15	15	14	14
MCLN, κ	8	8	9	9
MCSO N_{it}	25	26	27	27
MCSO κ	37	41	43	45
MRP N_{it}	48	82	155	271
MRP κ	317	1 890	8 900	34 400

TAB. 6.14 – Problème mixte, préconditionnement avec la matrice de masse réduite.

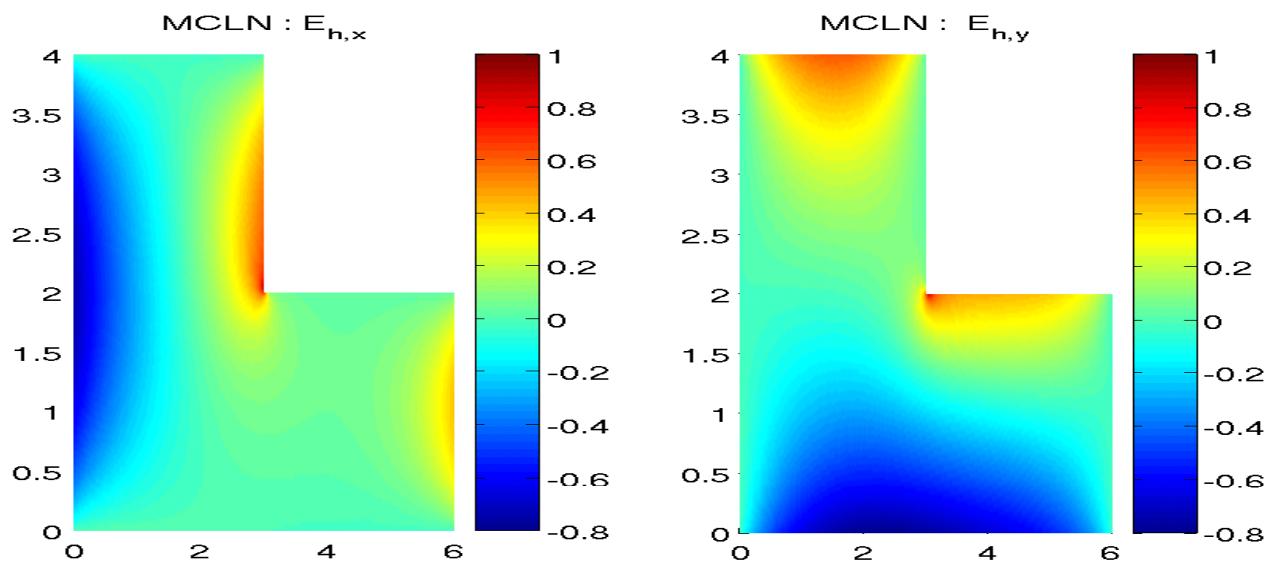
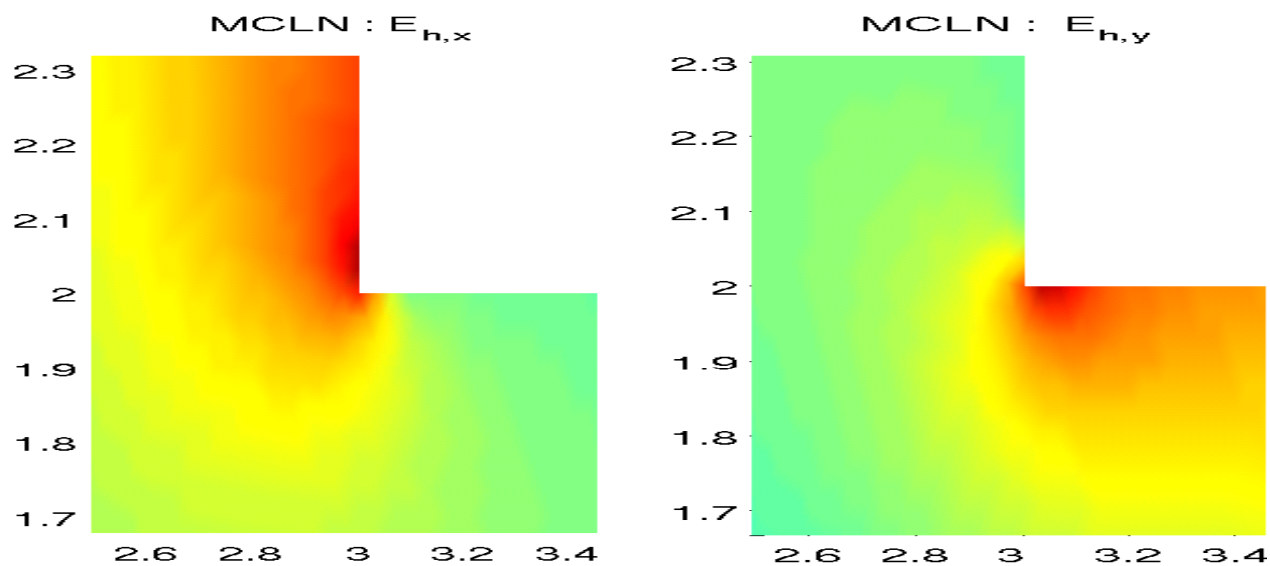
6.4 Allure du champ quasi-électrostatique

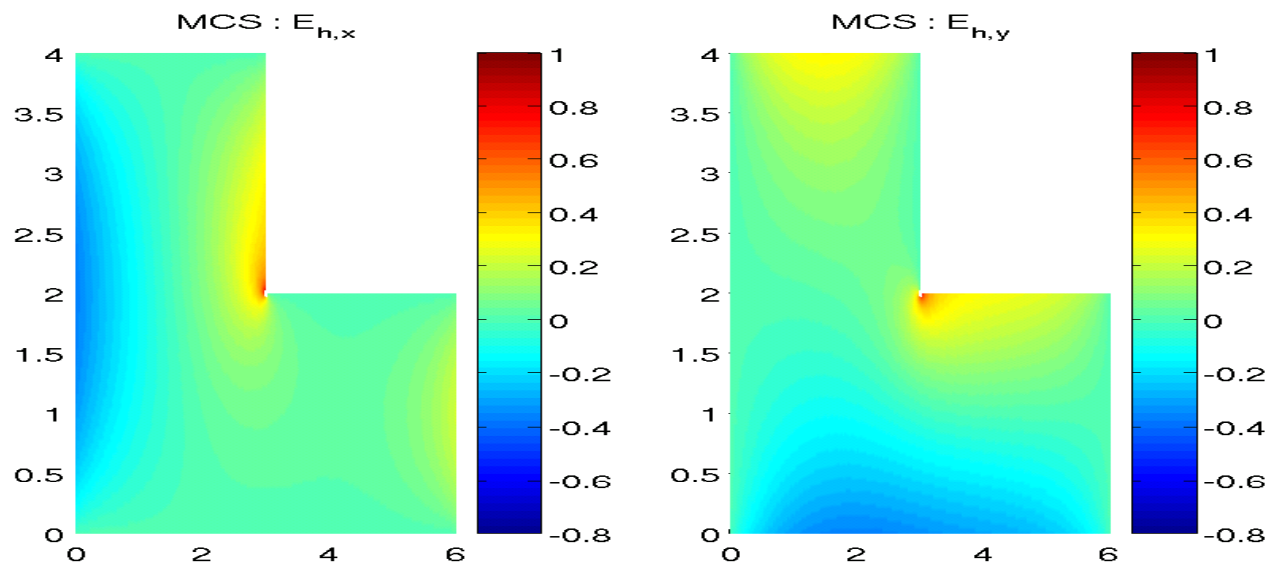
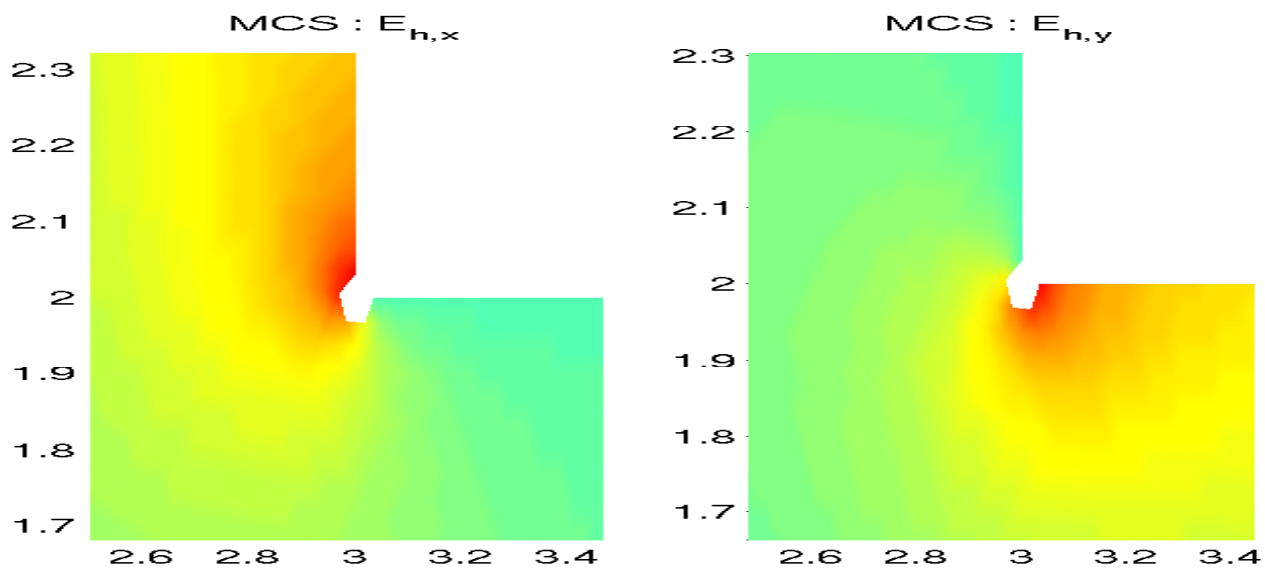
Sur les figures qui suivent, on a représenté les valeurs normalisées des composantes de $\mathbf{E}_{\mathbf{x}}$, approché avec les éléments finis de Taylor-Hood P_2 - P_1 sur le quatrième maillage. On remarquera pour les trois méthodes, le champ électrique a la même allure, il s'enroule autour du coin rentrant. Les valeurs maximales des composantes de $\mathbf{E}_h$ sont de l'ordre de 1.3 pour la MCLN et la MRP, alors que pour la MCSO, elles sont de l'ordre de 1.8, car on ajoute explicitement le champ singulier, qui prend de très grandes valeurs au voisinage du coin rentrant.

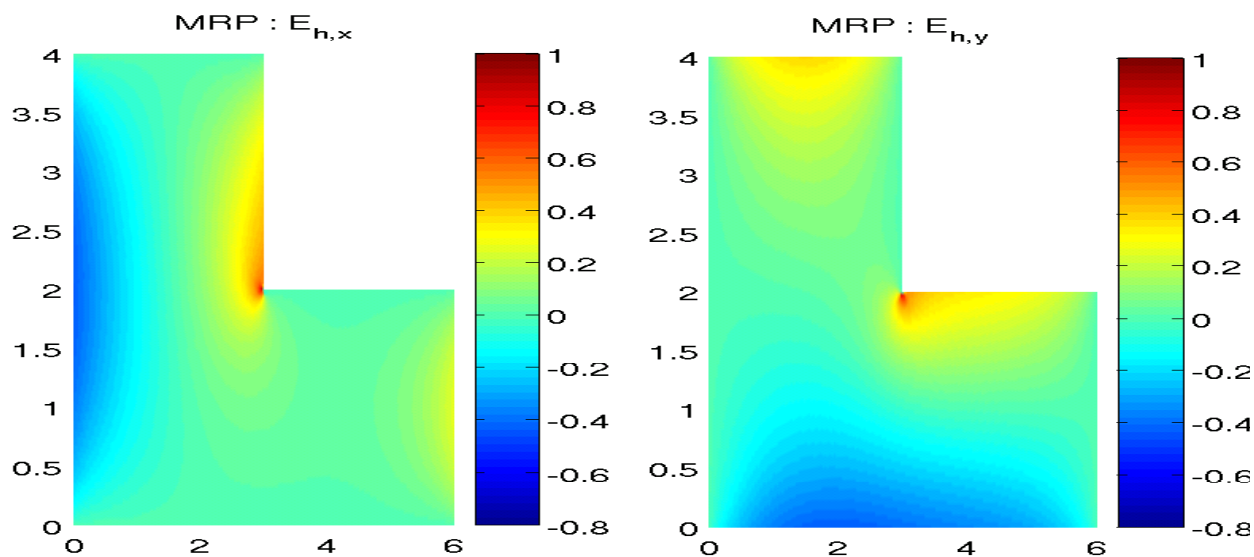
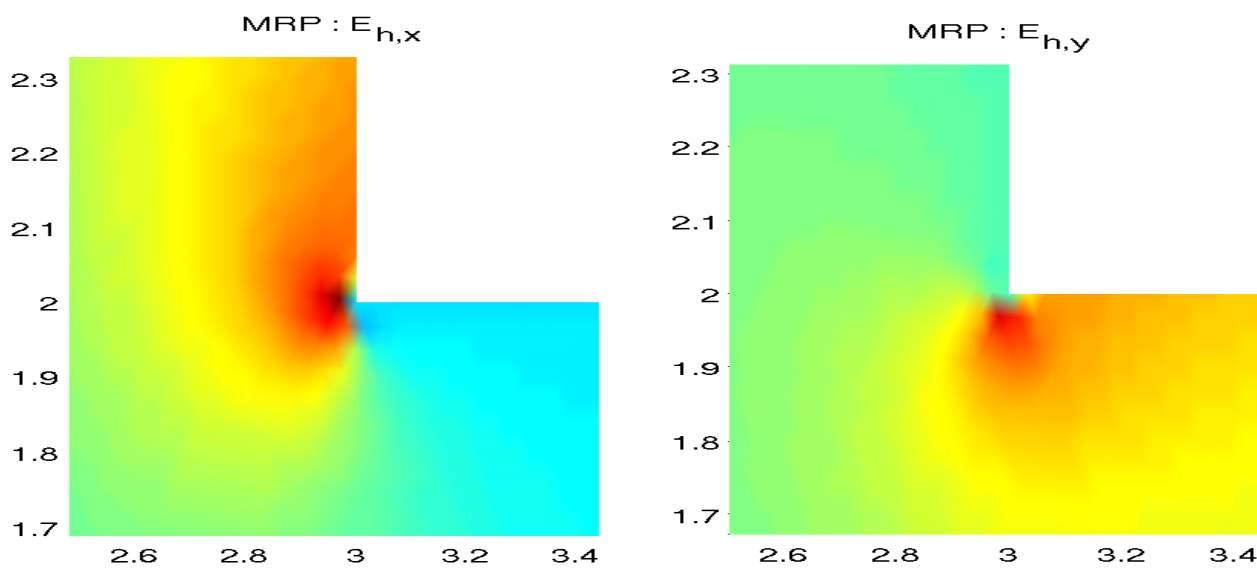
- Sur les figures 6.3 et 6.4, $\mathbf{E}_{\mathbf{x}} = \mathbf{E}_{nat}$. On observe comme prévu que le champ électrique est très intense au voisinage du coin rentrant. On peut voir sur le zoom que la condition aux limites tangentielle n'est pas exactement nulle au voisinage du coin rentrant, car elle est imposée faiblement.

- Sur la figure 6.5, $\mathbf{E}_{\mathbf{x}} = \mathbf{E}_{\perp}$. De même, le champ électrique est très intense au voisinage du coin rentrant. L'approximation du champ électrique est infinie au coin rentrant, on ne peut donc pas représenter le champ électrique en ce point. C'est pourquoi, pour la représentation, on a tronqué les triangles qui touchent le coin rentrant, comme on peut l'observer sur le zoom. On observe cette fois-ci que la condition aux limites tangentielle est bien nulle partout.

- Sur la figure 6.8, $\mathbf{E}_{\mathbf{x}} = \mathbf{E}_{\gamma}$. Encore une fois, on constate que le champ électrique est très intense au voisinage du coin rentrant. Par rapport aux figures précédentes, on remarque un léger déplacement de la valeur maximale, qui est en fait un artefact de calcul. Comme pour la MCSO, il est clair que la condition aux limites tangentielle est nulle (elle est en fait exactement nulle pour cette approche).

FIG. 6.3 – MCLN : amplitudes relatives de $E_{nat,x}$ et $E_{nat,y}$.FIG. 6.4 – $E_{nat,x}$ et $E_{nat,y}$: zoom au voisinage du coin rentrant.

FIG. 6.5 – MCSO : amplitudes relatives de $E_{\perp,x}$ et $E_{\perp,y}$.FIG. 6.6 – $E_{\perp,x}$ et $E_{\perp,y}$: zoom au voisinage du coin rentrant.

FIG. 6.7 – MRP : amplitudes relatives de $E_{\gamma,x}$ et $E_{\gamma,y}$.FIG. 6.8 – $E_{\gamma,x}$ et $E_{\gamma,y}$: zoom au voisinage du coin rentrant.

6.5 Conclusions

- Pour conclure, rappelons les principaux résultats observés pour la MCLN, la MCSO et la MRP :
- Dans le cas d'une solution régulière, ces méthodes sont équivalentes (erreur en norme $\mathbf{X}$ de l'ordre de 30% sur le premier maillage et de l'ordre de 2% sur le dernier).
 - La méthode directe la plus simple à programmer est la MCLN (pas de traitement particulier pour les noeuds du bord ni pour les singularités). La méthode directe la plus complexe à programmer est la λ -approche en P_1 et la MCSO en P_2-P_1 .
 - L'approximation de $\mathbf{E}$ dans $\mathbf{X}_{\mathbf{E},k}^{0,R}$ diverge si $\mathbf{E}_0 \notin \mathbf{H}^1(\omega)$,
 - La λ -approche et la MCSO donnent exactement le même résultat numérique en P_1 ,
 - Pour la MRP, $\|\operatorname{div} \mathbf{E}_\gamma - g\|_0$ diverge si la solution exacte $\mathbf{E} \notin \mathbf{H}^1(\omega)$,
 - Les trois méthodes sont robustes : elles convergent même avec des données peu régulières (paragraphe 6.2.3),
 - Pour la résolution du problème mixte, le préconditionnement de $\mathbb{CA}^{-1}\mathbb{C}^T$ avec la matrice de masse réduite permet de réduire nombre d'itérations sans coût supplémentaire,
 - Expérimentalement, la méthode qui est la plus performante en terme de précision est la λ -approche pour le P_1 et la MCSO pour le P_2-P_1 ,
 - Pour la MCLN, l'approximation médiocre de la composante tangentielle à $\partial\omega$ n'est pas rédhibitoire, car pour des données peu singulières (comme celles de la section 6.3), l'erreur de convergence est de l'ordre de 16% en P_1 et 11% en P_2-P_1 dès le premier maillage.

Nous ne présentons pas tous les résultats obtenus, néanmoins nous signalons au lecteur que la MRP dépend fortement du maillage. Notons d'une part que l'utilisation de maillages homogénéisés stabilise le taux de convergence en P_1 ; d'autre part, si le maillage est trop raffiné localement autour d'un noeud autre que le coin rentrant, le champ électrique peut prendre de grandes valeurs au voisinage de ce noeud.

Troisième partie

Le problème tridimensionnel

Chapitre 7

Notations et résultats préliminaires

3D

7.1 Notations relatives au domaine d'étude

Le domaine d'étude $\Omega \subset \mathbb{R}^3$ est un polyèdre de frontière lipschitzienne $\partial\Omega$. On suppose en particulier que Ω est simplement connexe et que $\partial\Omega$ est connexe. Dans le cas général, on renvoie le lecteur à [5].

La frontière $\partial\Omega$ est composée de K faces ouvertes F_k , $k \in \{1, \dots, K\}$: $\partial\Omega = \cup_k \overline{F_k}$. La normale unitaire sortante à F_k est notée $\boldsymbol{\nu}_k$.

On appellera $A_{k,l}$ l'arête partagée entre les faces F_k et F_l , et $\Theta_{k,l} \in]0, \pi[\cap]\pi, 2\pi[$ l'angle dièdre entre ces faces. Soit $\boldsymbol{\tau}_{k,l}$ le vecteur parallèle à $A_{k,l}$, et $\boldsymbol{\tau}_k = \boldsymbol{\tau}_{k,l} \times \boldsymbol{\nu}_k$. Ainsi, le couple $(\boldsymbol{\tau}_k, \boldsymbol{\tau}_{k,l})$

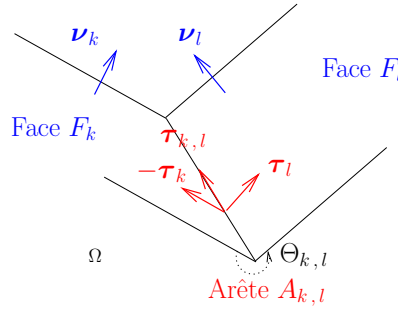


FIG. 7.1 – Vecteurs normaux et tangentiels à F_k et F_l .

forme une base orthonormée dans le plan généré par F_k ; et $(\boldsymbol{\tau}_k, \boldsymbol{\tau}_{k,l}, \boldsymbol{\nu}_k)$ forme une base orthonormée de $\mathbb{R}^3$ (voir la figure 7.1). On utilisera les notations suivantes :

$(\mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3) = (\mathbf{e}_x, \mathbf{e}_y, \mathbf{e}_z)$: base orthonormale canonique de $\mathbb{R}^3$.

$(x_1, x_2, x_3) = (x, y, z)$: coordonnées d'un point de $\mathbb{R}^3$.

$\mathbf{u} = (u_1, u_2, u_3) = (u_x, u_y, u_z)$: composantes d'un champ de vecteurs de $\mathbb{R}^3$.

$\mathbf{u} \cdot \mathbf{v} = u_1 v_1 + u_2 v_2 + u_3 v_3 \in \mathbb{R}$: produit scalaire 3D entre $\mathbf{u}$ et $\mathbf{v}$.

$\mathbf{u} \times \mathbf{v} = (u_2 v_3 - u_3 v_2, u_3 v_1 - u_1 v_3, u_1 v_2 - u_2 v_1) \in \mathbb{R}^3$: produit vectoriel 3D entre $\mathbf{u}$ et $\mathbf{v}$.

$\mathbf{u} \times^2 \mathbf{v} = u_1 v_2 - u_2 v_1 \in \mathbb{R}$: produit vectoriel 2D entre $\mathbf{u}$ et $\mathbf{v}$.

Pour tous $k \in \{1, \dots, K\}$, on note :

$\boldsymbol{\nu}|_{F_k} = (\nu_1, \nu_2, \nu_3)$: vecteur unitaire sortant normal à F_k .

$\mathbf{u}_\nu = \mathbf{u} \cdot \boldsymbol{\nu}|_{F_k}$: composante normale à F_k du vecteur $\mathbf{u}$.

$\mathbf{u}_T = \boldsymbol{\nu} \times (\mathbf{u} \times \boldsymbol{\nu})|_{F_k}$: composante tangentielle à F_k du vecteur $\mathbf{u}$.

Notons que sur F_k :

$$\mathbf{u} = \mathbf{u}_T + \mathbf{u}_\nu \boldsymbol{\nu}.$$

Par abus de notation, on généralisera ces définitions à $\partial\Omega$.

Propriétés 7.1 Soient $\mathbf{u}, \mathbf{v}, \mathbf{w}$ des vecteurs de $\mathbb{R}^3$. On a alors les propriétés vectorielles suivantes :

$$\mathbf{u} \cdot (\mathbf{v} \times \mathbf{w}) = (\mathbf{w} \times \mathbf{u}) \cdot \mathbf{v}, \quad (7.1)$$

$$\mathbf{u} \times (\mathbf{v} \times \mathbf{w}) = (\mathbf{u} \cdot \mathbf{w}) \mathbf{v} - (\mathbf{u} \cdot \mathbf{v}) \mathbf{w}. \quad (7.2)$$

On suppose que Ω présente N_{ar} arêtes rentrantes $(A_e)_{e \in \{1, \dots, N_{ar}\}}$, c'est-à-dire des arêtes dont les angles dièdres, notés $(\Theta_e)_{e \in \{1, \dots, N_{ar}\}}$ sont compris strictement entre π et 2π . On note $(\alpha_e)_{e \in \{1, \dots, N_{ar}\}}$, et on pose :

$$\boxed{\alpha = \min_e \alpha_e} \quad (7.3)$$

Définition 7.2 Soit $E = \cup_e \overline{A_e}$, la fermeture de l'ensemble des arêtes rentrantes de $\partial\Omega$. On pose $d_0(\mathbf{x}) = d(\mathbf{x}, E)$.

7.2 Opérateurs tridimensionnels

On utilisera les notations suivantes :

t : variable temporelle.

$\partial_t(\cdot) = \frac{\partial(\cdot)}{\partial t}$: dérivée partielle par rapport au temps.

$\partial_t^m(\cdot) = \frac{\partial^m(\cdot)}{\partial t^m}$: dérivée partielle $m^{\text{ième}}$ par rapport au temps.

$\partial_i(\cdot)$: la dérivée partielle suivant x_i .

$\partial_i^m(\cdot) = \frac{\partial^m(\cdot)}{\partial x_i^m}$: la dérivée partielle $m^{\text{ième}}$ suivant x_i .

$$\partial_{\mathbf{x}}^m v = \sum_{m_1, m_2, m_3 \geq 0 \mid m_1 + m_2 + m_3 = m} \frac{\partial^m v}{\partial x_1^{m_1} \partial x_2^{m_2} \partial x_3^{m_3}}.$$

$$\mathbf{grad} v = (\partial_1 v, \partial_2 v, \partial_3 v).$$

$$\mathbf{grad}_{\partial\Omega} v = \boldsymbol{\nu} \times (\mathbf{grad} v \times \boldsymbol{\nu})|_{\partial\Omega} : \text{gradient surfacique de } v \text{ sur } \partial\Omega, \text{ tangentiel à } \partial\Omega.$$

$$\mathbf{rot}_{\partial\Omega} v = (\mathbf{grad} v \times \boldsymbol{\nu})|_{\partial\Omega} : \text{rotationnel surfacique de } v \text{ sur } \partial\Omega, \text{ tangentiel à } \partial\Omega.$$

$$\partial_\nu v|_{\partial\Omega} = \mathbf{grad} v \cdot \boldsymbol{\nu}|_{\partial\Omega} : \text{dérivée de } v \text{ sur } \partial\Omega \text{ dans la direction normale à } \partial\Omega.$$

$$\mathbf{rot}_{\partial\Omega} \mathbf{u} : \text{l'opérateur adjoint de } \mathbf{rot}_{\partial\Omega}.$$

$$\mathbf{div}_{\partial\Omega} \mathbf{u} : \text{l'opérateur adjoint de } -\mathbf{grad}_{\partial\Omega}.$$

$$\Delta v = \partial_1^2 v + \partial_2^2 v + \partial_3^2 v : \text{le laplacien de } v.$$

$$\mathbf{rot} \mathbf{u} = (\partial_2 u_3 - \partial_3 u_2, \partial_3 u_1 - \partial_1 u_3, \partial_1 u_2 - \partial_2 u_1).$$

$$\mathbf{div} \mathbf{u} = \partial_1 u_1 + \partial_2 u_2 + \partial_3 u_3.$$

$$\Delta \mathbf{u} = (\Delta u_1, \Delta u_2, \Delta u_3) : \text{le laplacien vectoriel de } \mathbf{u}.$$

$$\mathbf{grad} : \mathbf{u} = \begin{pmatrix} \partial_1 u_1 & \partial_2 u_1 & \partial_3 u_1 \\ \partial_1 u_2 & \partial_2 u_2 & \partial_3 u_2 \\ \partial_1 u_3 & \partial_2 u_3 & \partial_3 u_3 \end{pmatrix} : \text{le gradient matriciel de } \mathbf{u}.$$

Relations entre opérateurs tridimensionnels

$$\mathbf{div} \mathbf{rot} (.) = 0, \quad (7.4)$$

$$\mathbf{rot} \mathbf{grad} (.) = 0, \quad (7.5)$$

$$\mathbf{div} \mathbf{grad} (.) = \Delta (.), \quad (7.6)$$

$$-\mathbf{rot} \mathbf{rot} (.) + \mathbf{grad} \mathbf{div} (.) = \Delta (.), \quad (7.7)$$

$$\mathbf{div} (\mathbf{v} \times \mathbf{w}) = \mathbf{w} \cdot \mathbf{rot} \mathbf{v} - \mathbf{v} \cdot \mathbf{rot} \mathbf{w}. \quad (7.8)$$

Propriété 7.3 La relation suivante est montrée par É. Heintz dans [69] :
Soient $\mathbf{u}$ et $\mathbf{v}$ réguliers. On a alors :

$$\begin{aligned} (\mathbf{div} \mathbf{u}, \mathbf{div} \mathbf{v})_0 + (\mathbf{rot} \mathbf{u}, \mathbf{rot} \mathbf{v})_0 &= (\mathbf{grad} : \mathbf{u}, \mathbf{grad} : \mathbf{v})_0 \\ &+ \sum_{\alpha=1}^3 (\mathbf{grad} u_\alpha \times \boldsymbol{\nu}, \mathbf{e}_\alpha \times \mathbf{v})_{0, \partial\Omega}. \end{aligned}$$

Propriété 7.4 La relation suivante est montrée par A. Buffa et P. Ciarlet, Jr. dans [29] :
Soit $\mathbf{u} \in \mathcal{H}(\mathbf{rot}, \Omega)$. On a alors : $\mathbf{div}_{\partial\Omega}(\mathbf{u} \times \boldsymbol{\nu})|_{\partial\Omega} = \mathbf{rot} \mathbf{u} \cdot \boldsymbol{\nu}|_{\partial\Omega}$.

7.3 Espaces de Hilbert usuels et leurs normes associées

De façon générale, les espaces désignés par des lettres majuscules en caractère calligraphiques sont des espaces de champs vectoriels, et les espaces désignés par des lettres majuscules italiques en caractère normal sont des espaces de champs scalaires. Soit $H(\cdot)$ un espace de champs scalaires.

On notera $\mathcal{H}(\cdot) = H(\cdot)^3$, l'espace de champs vectoriels correspondant.

$d\Omega$ désigne la mesure de l'ouvert Ω et $d\Sigma$ désigne la mesure de l'ouvert $\partial\Omega$.

$\mathcal{D}(\Omega)$ est l'espace des fonctions $C^\infty(\Omega)$, à support compact dans Ω . On appelle $\mathcal{D}'(\Omega)$ son dual, l'espace des distributions.

Pour la définition de la dualité, on renvoie le lecteur à la section 1.6.

7.3.1 Espaces de champs scalaires

$$L^2(\Omega) = \left\{ u \text{ mesurable sur } \Omega : \int_{\Omega} u^2 d\Omega < \infty \right\}, \quad \|u\|_0 = \left(\int_{\Omega} u^2 d\Omega \right)^{1/2}.$$

$$L^2_{loc}(\Omega) = \left\{ u \text{ mesurable sur } \Omega : \forall \Omega_c \mid \overline{\Omega_c} \subset \Omega, \int_{\Omega_c} u^2 d\Omega < \infty \right\}.$$

$$L^2_0(\Omega) = \left\{ u \in L^2(\Omega) : \int_{\Omega} u d\Omega = 0 \right\}.$$

$$H^1(\Omega) = \{ u \in L^2(\Omega) : \mathbf{grad} u \in \mathcal{L}^2(\Omega) \}, \quad \|u\|_{H^1} = (\|u\|_0^2 + \|\mathbf{grad} u\|_0^2)^{1/2}.$$

$$H^1_0(\Omega) = \{ u \in H^1(\Omega) : u|_{\partial\Omega} = 0 \}, \quad \|u\|_{H^1_0} = \|\mathbf{grad} u\|_0.$$

$$H^1_{\Delta}(\Omega) = \{ u \in H^1(\Omega) : \Delta u \in L^2(\Omega) \}, \quad \|u\|_{H^1_{\Delta}} = (\|u\|_{H^1}^2 + \|\Delta u\|_0^2)^{1/2}.$$

$$H^2(\Omega) = \{ u \in H^1(\Omega) : \mathbf{grad} u \in \mathcal{H}^1(\Omega) \}, \quad \|u\|_{H^2} = (\|u\|_0^2 + \|\mathbf{grad} u\|_{\mathcal{H}^1}^2)^{1/2}.$$

L'espace $\mathcal{H}^1(\Omega)$ et sa norme sont définis dans le paragraphe 7.3.2 ci-dessous. L'équivalence entre la norme du graphe et la semi-norme dans $H^1_0(\Omega)$ est due à l'inégalité de Poincaré.

On désigne par $H^{-1}(\Omega)$ le dual de $H^1_0(\Omega)$.

$\Phi_{\mathcal{N}}$ et $\Phi_{\mathcal{D}}$ sont les espaces des solutions du Laplacien :

$$\Phi_{\mathcal{N}} = \{ u \in H^1(\Omega) \cap L^2_0(\Omega) : \Delta u \in L^2(\Omega), \partial_{\nu} u|_{\partial\Omega} = 0 \},$$

$$\Phi_{\mathcal{D}} = \{ u \in H^1_0(\Omega) : \Delta u \in L^2(\Omega) \}.$$

Dans $\Phi_{\mathcal{N}}$ et $\Phi_{\mathcal{D}}$, la semi-norme est équivalente à la norme du graphe : $\|u\|_{\Phi_{\mathcal{D},\mathcal{N}}} = \|u\|_{\Phi} := \|\Delta u\|_0$ (inégalité de Poincaré-Friedrichs pour $\Phi_{\mathcal{D}}$ et inégalité de Poincaré-Wirtinger et théorème de Lax-Milgram pour $\Phi_{\mathcal{N}}$).

7.3.2 Espaces de champs vectoriels

$$\mathcal{L}^2(\Omega) = L^2(\Omega)^3, \quad \|\mathbf{u}\|_0 = (\|u_1\|_0^2 + \|u_2\|_0^2 + \|u_3\|_0^2)^{1/2}.$$

$$\mathcal{H}^1(\Omega) = H^1(\Omega)^3, \quad \|\mathbf{u}\|_{\mathcal{H}^1} = (\|\mathbf{u}\|_0^2 + \|\mathbf{grad} : \mathbf{u}\|_0^2)^{1/2}.$$

$$\mathcal{H}(\mathbf{rot}, \Omega) = \{\mathbf{u} \in \mathcal{L}^2(\Omega) : \mathbf{rot} \mathbf{u} \in \mathcal{L}^2(\Omega)\}, \quad \|\mathbf{u}\|_{\mathcal{H}(\mathbf{rot})} = (\|\mathbf{u}\|_0^2 + \|\mathbf{rot} \mathbf{u}\|_0^2)^{1/2}.$$

$$\mathcal{H}_0(\mathbf{rot}, \Omega) = \{\mathbf{u} \in \mathcal{H}(\mathbf{rot}, \Omega) : \mathbf{u} \times \boldsymbol{\nu}_{|\partial\Omega} = 0\}.$$

$$\mathcal{H}(\mathbf{rot}^0, \Omega) = \{\mathbf{u} \in \mathcal{L}^2(\Omega) : \mathbf{rot} \mathbf{u} = 0\}, \quad \|\mathbf{u}\|_{\mathcal{H}(\mathbf{rot}^0)} = \|\mathbf{u}\|_0.$$

$$\mathcal{H}_0(\mathbf{rot}^0, \Omega) = \{\mathbf{u} \in \mathcal{L}^2(\Omega) : \mathbf{rot} \mathbf{u} = 0, \mathbf{u} \times \boldsymbol{\nu}_{|\partial\Omega} = 0\}.$$

$$\mathcal{H}(\text{div}, \Omega) = \{\mathbf{u} \in \mathcal{L}^2(\Omega) : \text{div} \mathbf{u} \in L^2(\Omega)\}, \quad \|\mathbf{u}\|_{\mathcal{H}(\text{div})} = (\|\mathbf{u}\|_0^2 + \|\text{div} \mathbf{u}\|_0^2)^{1/2}.$$

$$\mathcal{H}_0(\text{div}, \Omega) = \{\mathbf{u} \in \mathcal{H}(\text{div}, \Omega) : u_\nu = 0\}.$$

$$\mathcal{H}(\text{div}^0, \Omega) = \{\mathbf{u} \in \mathcal{L}^2(\Omega) : \text{div} \mathbf{u} = 0\}, \quad \|\mathbf{u}\|_{\mathcal{H}(\text{div}^0)} = \|\mathbf{u}\|_0.$$

$$\mathcal{H}_0(\text{div}^0, \Omega) = \{\mathbf{u} \in \mathcal{L}^2(\Omega) : \text{div} \mathbf{u} = 0, u_\nu = 0\}.$$

$\mathcal{X}_{\mathcal{E}}$ et $\mathcal{X}_{\mathcal{H}}$ sont les espaces des solutions des équations de Maxwell quasi-statique :

$$\mathcal{X}_{\mathcal{E}} = \{\mathbf{u} \in \mathcal{H}(\mathbf{rot}, \Omega) \cap \mathcal{H}(\text{div}, \Omega) : \mathbf{u} \times \boldsymbol{\nu}_{|\partial\Omega} \in \mathcal{L}_t^2(\partial\Omega)\}.$$

$$\mathcal{X}_{\mathcal{H}} = \{\mathbf{u} \in \mathcal{H}(\mathbf{rot}, \Omega) \cap \mathcal{H}(\text{div}, \Omega) : u_\nu \in L^2(\partial\Omega)\}.$$

Voir le paragraphe 7.3.3 pour une définition des espaces de trace.

D'après [49], ces espaces sont identiques algébriquement et topologiquement.

Dans $\mathcal{X}_{\mathcal{E}}$ et $\mathcal{X}_{\mathcal{H}}$, la semi-norme est équivalente à la norme du graphe (voir la section 8.1), d'où :

$$\|\mathbf{u}\|_{\mathcal{X}_{\mathcal{E}}} = (\|\mathbf{rot} \mathbf{u}\|_0^2 + \|\text{div} \mathbf{u}\|_0^2 + \|\mathbf{u} \times \boldsymbol{\nu}_{|\partial\Omega}\|_0^2)^{1/2}, \quad \|\mathbf{u}\|_{\mathcal{X}_{\mathcal{H}}} = (\|\mathbf{rot} \mathbf{u}\|_0^2 + \|\text{div} \mathbf{u}\|_0^2 + \|u_\nu\|_0^2)^{1/2}.$$

On considère les sous-espaces de $\mathcal{X}_{\mathcal{E}}$ suivants :

$$\mathcal{X}_{\mathcal{E}}^0 = \{\mathbf{u} \in \mathcal{X}_{\mathcal{E}} : \mathbf{u} \times \boldsymbol{\nu}_{|\partial\Omega} = 0\}, \quad \|\mathbf{u}\|_{\mathcal{X}_{\mathcal{E}}^0} = (\|\mathbf{rot} \mathbf{u}\|_0^2 + \|\text{div} \mathbf{u}\|_0^2)^{1/2}.$$

$$\mathcal{V}_{\mathcal{E}} = \{\mathbf{u} \in \mathcal{X}_{\mathcal{E}}^0 : \text{div} \mathbf{u} = 0\}, \quad \|\mathbf{u}\|_{\mathcal{V}_{\mathcal{E}}} = \|\mathbf{rot} \mathbf{u}\|_0.$$

$$\mathcal{W}_{\mathcal{E}} = \{\mathbf{u} \in \mathcal{X}_{\mathcal{E}}^0 : \mathbf{rot} \mathbf{u} = 0\}, \quad \|\mathbf{u}\|_{\mathcal{W}_{\mathcal{E}}} = \|\text{div} \mathbf{u}\|_0.$$

On considère les sous-espaces de $\mathcal{X}_{\mathcal{H}}$ suivants :

$$\mathcal{X}_{\mathcal{H}}^0 = \{\mathbf{u} \in \mathcal{X}_{\mathcal{H}} : u_\nu = 0\}, \quad \|\mathbf{u}\|_{\mathcal{X}_{\mathcal{H}}^0} = (\|\mathbf{rot} \mathbf{u}\|_0^2 + \|\text{div} \mathbf{u}\|_0^2)^{1/2}.$$

$$\mathcal{V}_{\mathcal{H}} = \{\mathbf{u} \in \mathcal{X}_{\mathcal{H}}^0 : \text{div} \mathbf{u} = 0\}, \quad \|\mathbf{u}\|_{\mathcal{V}_{\mathcal{H}}} = \|\mathbf{rot} \mathbf{u}\|_0.$$

$$\mathcal{W}_{\mathcal{H}} = \{\mathbf{u} \in \mathcal{X}_{\mathcal{H}}^0 : \mathbf{rot} \mathbf{u} = 0\}, \quad \|\mathbf{u}\|_{\mathcal{W}_{\mathcal{H}}} = \|\text{div} \mathbf{u}\|_0.$$

7.3.3 Espaces des traces

$$L^2(\partial\Omega) = \left\{ u \text{ mesurable sur } \partial\Omega : \int_{\partial\Omega} u^2 d\Sigma < \infty \right\}, \quad \|u\|_{0,\partial\Omega} = \left(\int_{\partial\Omega} u^2 d\Sigma \right)^{1/2}.$$

$$\mathcal{L}_t^2(\partial\Omega) = \left\{ \mathbf{u} \in L^2(\partial\Omega)^3 : \forall k \in \{1, \dots, K\}, \mathbf{u} \cdot \boldsymbol{\nu}_{|F_k} = 0 \right\}, \quad \|\mathbf{u}\|_{0,\partial\Omega} = \left(\int_{\partial\Omega} |\mathbf{u}|^2 d\Sigma \right)^{1/2}.$$

$\forall u \in H^1(\Omega), u|_{\partial\Omega} \in H^{1/2}(\partial\Omega)$, où :

$$H^{1/2}(\partial\Omega) = \left\{ u \in L^2(\partial\Omega) : \int_{\partial\Omega} \int_{\partial\Omega} \frac{(u(\mathbf{x}) - u(\mathbf{y}))^2}{\|\mathbf{x} - \mathbf{y}\|^3} d\Sigma(\mathbf{x}) d\Sigma(\mathbf{y}) < \infty \right\}.$$

Proposition 7.5 *L'application : $H^1(\Omega) \rightarrow H^{1/2}(\partial\Omega)$, $u \mapsto u|_{\partial\Omega}$ est surjective et continue.*

$H^{-1/2}(\partial\Omega)$ est le dual de $H^{1/2}(\partial\Omega)$.

Les espaces suivants sont caractérisés par A. Buffa et P. Ciarlet, Jr. dans [29].

$$\mathcal{H}_-^{1/2}(\partial\Omega) = \left\{ \mathbf{u} \in \mathcal{L}_t^2(\partial\Omega) : \forall k \in \{1, \dots, K\}, \mathbf{u}|_{F_k} \in \mathcal{H}^{1/2}(F_k) \right\}.$$

$$\mathcal{H}_\perp^{1/2}(\partial\Omega) = \left\{ \mathbf{v} \times \boldsymbol{\nu}_{|\partial\Omega}, \mathbf{v} \in \mathcal{H}^1(\Omega) \right\}.$$

$$\mathcal{H}_\parallel^{1/2}(\partial\Omega) = \left\{ \mathbf{v}_T, \mathbf{v} \in \mathcal{H}^1(\Omega) \right\}.$$

$$\mathcal{H}_\perp^{-1/2}(\mathbf{rot}_{\partial\Omega}, \partial\Omega) = \left\{ \mathbf{v}_T, \mathbf{v} \in \mathcal{H}(\mathbf{rot}, \partial\Omega) \right\}.$$

$$\mathcal{H}_\parallel^{-1/2}(\text{div}_{\partial\Omega}, \partial\Omega) = \left\{ \mathbf{v} \times \boldsymbol{\nu}_{|\partial\Omega}, \mathbf{v} \in \mathcal{H}(\mathbf{rot}, \partial\Omega) \right\}.$$

D'après [30], $\mathcal{H}_\perp^{-1/2}(\mathbf{rot}_{\partial\Omega}, \partial\Omega)' = (\mathcal{H}_\parallel^{-1/2}(\text{div}_{\partial\Omega}, \partial\Omega))'$, ce qui permet de définir une formule d'intégration par parties pour les éléments de $\mathcal{H}(\mathbf{rot}, \Omega)$ (voir l'équation (7.15)).

Proposition 7.6 *L'application suivante :*

$$\begin{aligned} \mathcal{X}_{\mathcal{E}} &\rightarrow \mathcal{L}_t^2(\partial\Omega) \cap \mathcal{H}_\perp^{-1/2}(\mathbf{rot}_{\partial\Omega}, \partial\Omega) \\ \mathbf{v} &\mapsto \mathbf{v}_T \end{aligned}$$

est surjective et continue.

DÉMONSTRATION. Soit $\mathbf{v} \in \mathcal{L}_t^2(\partial\Omega) \cap \mathcal{H}_\parallel^{-1/2}(\text{div}_{\partial\Omega}, \partial\Omega)$. Il existe donc $\mathbf{u} \in \mathcal{H}(\mathbf{rot}, \Omega)$ tel que $\mathbf{u} \times \boldsymbol{\nu}_{|\partial\Omega} = \mathbf{v}$. D'après le théorème de Lax-Milgram, il existe un unique $q \in H_0^1(\Omega)$ tel que : $\Delta q = \text{div } \mathbf{u}$ dans $H^{-1}(\Omega)$. Soit $\mathcal{F} = \mathbf{u} - \mathbf{grad } q \in \mathcal{L}^2(\Omega)$. On a : $\text{div } \mathcal{F} = 0 \in L^2(\Omega)$, et $\mathbf{rot } \mathcal{F} = \mathbf{rot } \mathbf{u} \in \mathcal{L}^2(\Omega)$. De plus, $\mathcal{F} \times \boldsymbol{\nu}_{|\partial\Omega} = \mathbf{u} \times \boldsymbol{\nu}_{|\partial\Omega} = \mathbf{v} \in \mathcal{L}_t^2(\partial\Omega)$. Finalement, $\mathcal{F} \in \mathcal{X}_{\mathcal{E}}$, et $\mathcal{F} \times \boldsymbol{\nu}_{|\partial\Omega} = \mathbf{v}$. On en conclut que pour tout $\mathbf{v} \in \mathcal{L}_t^2(\partial\Omega) \cap \mathcal{H}_\parallel^{-1/2}(\text{div}_{\partial\Omega}, \partial\Omega)$, il existe $\mathcal{F} \in \mathcal{X}_{\mathcal{E}}$ tel que $\mathcal{F} \times \boldsymbol{\nu}_{|\partial\Omega} = \mathbf{v}_{|\partial\Omega}$. La continuité est immédiate. $\square$

Ci-dessus, on a utilisé le fait que pour tout élément $q \in H_0^1(\Omega)$, $\mathbf{grad } q \times \boldsymbol{\nu}_{|\partial\Omega} = 0$ (voir [65]).

7.4 Espaces à poids

Pour les espaces à poids dans les domaines tridimensionnels, plusieurs définitions sont possibles. Nous choisissons celle de [36], plutôt que celle de [53] car elle est adaptée aux domaines polyédriques, pour lesquels on peut confondre les distances aux arêtes rentrantes et les distances aux coins rentrants, ce qui n'est pas possible si le domaine présente des pointes coniques à bord régulier. Pour ce cas particulier, et pour une définition plus précise, nous renvoyons le lecteur à [53].

$$V_\gamma^l(\Omega) = \left\{ u \in L_{loc}^2(\Omega) : \sum_{|j| \leq l} \int_\Omega d_0^{2(\gamma-l+|j|)} |\partial_{\mathbf{x}}^j u|^2 d\Omega < \infty \right\}.$$

On a en particulier :

$$V_\gamma^0(\Omega) = \{ u \in L_{loc}^2(\Omega) : d_0^\gamma u \in L^2(\Omega) \}, \quad \|u\|_{0,\gamma} = \|d_0^\gamma u\|_0,$$

$$V_\gamma^1(\Omega) = \{ u \in L_{loc}^2(\Omega) : d_0^{\gamma-1} u \in L^2(\Omega), d_0^\gamma \mathbf{grad} u \in \mathcal{L}^2(\Omega) \}, \quad \|u\|_{V_\gamma^1(\Omega)} = (\|d_0^{\gamma-1} u\|_0 + \|\mathbf{grad} u\|_{0,\gamma})^{1/2},$$

$$V_0^1(\Omega) = \{ u \in L_{loc}^2(\Omega) : d_0^{-1} u \in L^2(\Omega), \mathbf{grad} u \in \mathcal{L}^2(\Omega) \}, \quad \|u\|_{V_0^1(\Omega)} = (\|d_0^{-1} u\|_0 + \|\mathbf{grad} u\|_0)^{1/2}.$$

On utilisera aussi la notation : $L_\gamma^2(\Omega) := V_\gamma^0(\Omega)$, et le produit scalaire de $L_\gamma^2(\Omega)$ ainsi que la norme associée seront notés :

$$(u, v)_{0,\gamma} := \int_\Omega d_0^{2\gamma} u v d\Omega, \quad \|u\|_{0,\gamma} := \left(\int_\Omega d_0^{2\gamma} u^2 d\Omega \right)^{1/2}.$$

Pour $0 < \gamma < 1$, on définit de plus :

$$H_{0,\gamma}^1(\Omega) = \{ \phi \in V_\gamma^1(\Omega) : \phi|_{\partial\Omega} = 0 \}, \quad \|\phi\|_{H_{0,\gamma}^1} = \|\mathbf{grad} \phi\|_{0,\gamma}.$$

$$\Phi_\gamma = \{ \phi \in H_0^1(\Omega) : \Delta \phi \in L_\gamma^2(\Omega) \}, \quad \|\phi\|_{\Phi,\gamma} = \|\Delta \phi\|_{0,\gamma}.$$

$$\mathcal{H}(\text{div}_\gamma, \Omega) = \{ \mathbf{u} \in \mathcal{L}^2(\Omega) : \text{div } \mathbf{u} \in L_\gamma^2(\Omega) \}, \quad \|\mathbf{u}\|_{\mathcal{H}(\text{div}_\gamma)} = (\|\mathbf{u}\|_0^2 + \|\text{div } \mathbf{u}\|_{0,\gamma}^2)^{1/2}.$$

$$\mathcal{H}_0(\text{div}_\gamma, \Omega) = \{ \mathbf{u} \in \mathcal{H}(\text{div}_\gamma, \Omega) : \mathbf{u}_\nu = 0 \}.$$

$$\mathcal{X}_{\mathcal{E},\gamma}^0 = \{ \mathbf{u} \in \mathcal{H}_0(\mathbf{rot}, \Omega) \cap \mathcal{H}(\text{div}_\gamma, \Omega) \}, \quad \|\mathbf{u}\|_{\mathcal{X}_{\mathcal{E},\gamma}^0} = (\|\mathbf{rot} \mathbf{u}\|_0^2 + \|\text{div } \mathbf{u}\|_{0,\gamma}^2)^{1/2}.$$

$$\mathcal{X}_{\mathcal{H},\gamma}^0 = \{ \mathbf{u} \in \mathcal{H}(\mathbf{rot}, \Omega) \cap \mathcal{H}_0(\text{div}_\gamma, \Omega) \}, \quad \|\mathbf{u}\|_{\mathcal{X}_{\mathcal{H},\gamma}^0} = (\|\mathbf{u}\|_0^2 + \|\mathbf{rot} \mathbf{u}\|_0^2 + \|\text{div } \mathbf{u}\|_{0,\gamma}^2)^{1/2}.$$

Notons que $H_{0,\gamma}^1(\Omega)$ est la fermeture des fonctions $C_0^\infty(\Omega)$ dans $V_\gamma^1(\Omega)$. En particulier, $H_0^1(\Omega)$ est la fermeture des fonctions $C_0^\infty(\Omega)$ dans $V_0^1(\Omega)$. Pour $\mathcal{X}_{\mathcal{E},\gamma}^0$, l'équivalence entre la semi-norme et la norme du graphe est prouvée dans la section 8.5.

Proposition 7.7 Dans $H_{0,\gamma}^1(\Omega)$, la norme du graphe $\|u\|_{V_\gamma^1}$ la semi-norme définie par : $|u|_{1,\gamma} := \|\mathbf{grad} d_0^\gamma u\|_0$ et sont équivalentes.

DÉMONSTRATION. Faisons la preuve par contradiction. Supposons qu'il existe $(q_n)_n \in H_{0,\gamma}^1(\Omega)$ une suite telle que : $\forall n, \|d_0^{\gamma-1} q_n\|_0 = 1$ ou $\|d_0^\gamma \mathbf{grad} q_n\|_0 = 1$; et $\|\mathbf{grad} d_0^\gamma q_n\|_0 \rightarrow 0$.

- Supposons que $\|d_0^{\gamma-1} q_n\|_0 = 1$.

Comme $\|\mathbf{grad} d_0^\gamma q_n\|_0 \rightarrow 0$ et que $d_0^\gamma q_n|_{\partial\Omega} = 0$, alors $d_0^\gamma q_n \in H_0^1(\Omega)$. D'où $d_0^\gamma q_n \rightarrow 0$ dans $H_0^1(\Omega)$. Or d'après [53], l'injection de $V_0^1(\Omega)$ dans $V_{-1}^0(\Omega)$ est continue, d'où, comme $H_0^1(\Omega) \subset V_0^1(\Omega)$: $d_0^{-1}(d_0^\gamma q_n) \rightarrow 0$ dans $L^2(\Omega)$, ce qui contredit l'hypothèse initiale.

• Supposons que $\|d_0^\gamma \mathbf{grad} q_n\|_0 = 1$.

Prenons le cas d'une unique arête rentrante, le long de l'axe $z = 0$. Posons $d_0 = r$, la distance orthogonale à cette arête, telle que : $r^2 = x_1^2 + x_2^2$. Pour $i \in \{1, 2\}$: $\partial_i(r^\gamma q_n) = \gamma x_i r^{\gamma-2} q_n + r^\gamma \partial_i q_n$, c'est-à-dire $r^\gamma \partial_i q_n = \partial_i(r^\gamma q_n) - \gamma x_i r^{\gamma-2} q_n$. Or $\|x_i r^{\gamma-2} q_n\|_0 \leq \|r^{\gamma-1} q_n\|_0 \rightarrow 0$ et $\|\partial_i(r^\gamma q_n)\|_0 \rightarrow 0$ par hypothèse. De plus, $\partial_3(r^\gamma q_n) = r^\gamma \partial_3 q_n$. D'où $r^\gamma \partial_i q_n \rightarrow 0$ pour $\forall i$, et $\|r^\gamma \mathbf{grad} q_n\|_0 \rightarrow 0$, ce qui contredit l'hypothèse initiale. $\square$

7.5 Formules d'intégration par parties classiques dans un ouvert

7.5.1 Formules de Green

$$\forall u \in \mathcal{D}(\overline{\Omega}), \forall v \in \mathcal{D}(\overline{\Omega}), \quad \int_{\Omega} \mathbf{grad} u \cdot \mathbf{grad} v \, d\Omega + \int_{\Omega} \Delta u v \, d\Omega = \int_{\partial\Omega} \partial_\nu u v \, d\Sigma. \quad (7.9)$$

$$\forall \mathbf{u} \in \mathcal{D}(\overline{\Omega})^3, \forall v \in \mathcal{D}(\overline{\Omega}), \quad \int_{\Omega} \mathbf{u} \cdot \mathbf{grad} v \, d\Omega + \int_{\Omega} \operatorname{div} \mathbf{u} v \, d\Omega = \int_{\partial\Omega} u_\nu v \, d\Sigma. \quad (7.10)$$

$$\forall \mathbf{u} \in \mathcal{D}(\overline{\Omega})^3, \forall \mathbf{v} \in \mathcal{D}(\overline{\Omega})^3, \quad \int_{\Omega} \mathbf{rot} \mathbf{u} \cdot \mathbf{v} \, d\Omega - \int_{\Omega} \mathbf{u} \cdot \mathbf{rot} \mathbf{v} \, d\Omega = \int_{\partial\Omega} \mathbf{u} \cdot (\mathbf{v} \times \boldsymbol{\nu}) \, d\Sigma. \quad (7.11)$$

7.5.2 Généralisation des formules de Green

Les preuves de ces formules se trouvent dans [65]. L'argument principal est la densité des fonctions régulières dans les espaces de Hilbert considérés. Ici, $H = H^{1/2}(\partial\Omega)$.

$$\forall u \in H_\Delta^1(\Omega), \forall v \in H^1(\Omega), \quad \int_{\Omega} \mathbf{grad} u \cdot \mathbf{grad} v \, d\Omega + \int_{\Omega} \Delta u v \, d\Omega = \langle \partial_\nu u, v \rangle_{H', H}. \quad (7.12)$$

$$\forall \mathbf{u} \in \mathcal{H}(\operatorname{div}, \Omega), \forall v \in H^1(\Omega), \quad \int_{\Omega} \mathbf{u} \cdot \mathbf{grad} v \, d\Omega + \int_{\Omega} \operatorname{div} \mathbf{u} v \, d\Omega = \langle u_\nu, v \rangle_{H', H}. \quad (7.13)$$

$$\forall \mathbf{u} \in \mathcal{H}(\mathbf{rot}, \Omega), \forall \mathbf{v} \in \mathcal{H}^1(\Omega), \quad \int_{\Omega} \mathbf{rot} \mathbf{u} \cdot \mathbf{v} \, d\Omega - \int_{\Omega} \mathbf{u} \cdot \mathbf{rot} \mathbf{v} \, d\Omega = \langle \mathbf{u}, \mathbf{v} \times \boldsymbol{\nu} \rangle_{\mathcal{H}', \mathcal{H}}. \quad (7.14)$$

Généralisation de (7.14), prouvée par A. Buffa et P. Ciarlet, Jr. [29] : $\forall \mathbf{u}, \mathbf{v} \in \mathcal{H}(\mathbf{rot}, \Omega)$

$$\int_{\Omega} \mathbf{rot} \mathbf{u} \cdot \mathbf{v} \, d\Omega - \int_{\Omega} \mathbf{u} \cdot \mathbf{rot} \mathbf{v} \, d\Omega = \langle \boldsymbol{\nu} \times (\mathbf{u} \times \boldsymbol{\nu}), \mathbf{v} \times \boldsymbol{\nu} \rangle_{\mathcal{H}', \mathcal{H}}, \quad (7.15)$$

avec $\mathcal{H} = \mathcal{H}_{\parallel}^{-1/2}(\operatorname{div}_{\partial\Omega}, \partial\Omega)$.

7.6 Espaces fonctionnels du problème en temps

Lorsqu'on résout les équations de Maxwell stationnaires, le champ électrique dépend à la fois du temps t et de la variable d'espace $\mathbf{x}$. On distingue ces deux variables. En général, on s'intéresse aux valeurs prises par le champ en un instant donné : on étudie par exemple $\mathbf{x} \rightarrow \mathcal{E}(\mathbf{x}, t_0)$, où t_0 est fixé. On peut se contenter de définir alors deux types d'espaces fonctionnels et une classe de distributions à valeurs vectorielles.

Soit $T > 0$ un temps final donné, $m \in \mathbb{N}$, X un espace de Banach et H un espace de Hilbert, de variable $\mathbf{x}$ tous deux.

Définitions 7.8 • $C^m(0, T; X)$ désigne l'ensemble des fonctions de classe C^m sur $]0, T[$ à valeurs dans X . C'est un espace de Banach muni de la norme :

$$\|u\|_{C^m(0, T; X)} := \sum_{k=0}^m \sup_{t \in [0, T]} \|\partial_t^k u(t)\|_X. \quad (7.16)$$

• L'espace $L^2(0, T; X)$ est l'ensemble des fonctions mesurables de carré intégrable sur $]0, T[$ à valeurs dans X . C'est un espace de Banach, muni de la norme :

$$\|u\|_{L^2(0, T; X)} := \left(\int_0^T \|u(\cdot, t)\|_X^2 dt \right)^{1/2}. \quad (7.17)$$

Si $X = H$, l'espace $L^2(0, T; X)$ est un espace de Hilbert muni du produit scalaire :

$$(u, v)_{L^2(0, T; X)} := \int_0^T (u(\cdot, t), v(\cdot, t))_H dt. \quad (7.18)$$

• L'espace des distributions sur $]0, T[$ à valeurs dans X , noté $\mathcal{D}'(]0, T[; X)$ est l'ensemble des applications linéaires et continues de $\mathcal{D}(]0, T[)$ dans X .

Par convention, on écrira que $\partial_t^m v \in L^2(0, T; X)$, ou de façon équivalente $v \in H^m(0, T; X)$.

Chapitre 8

Le problème statique 3D direct continu

8.1 Introduction

Le domaine Ω représente l'intérieur *vide* d'un conducteur parfait, qu'on borne si besoin est par une frontière artificielle Γ_A *ne coupant pas de singularités géométriques*. La frontière $\partial\Omega$ est dans ce cas composée de deux parties : $\partial\Omega = \overline{\Gamma_C} \cup \overline{\Gamma_A}$, où Γ_C est le conducteur parfait.

Dans cette configuration, le champ électrique bidimensionnel satisfait $\mathcal{E} \times \boldsymbol{\nu} = 0$ sur Γ_C et $\mathcal{E} \times \boldsymbol{\nu} = \mathbf{e}^* \times \boldsymbol{\nu} + c(\mathcal{B} \times \boldsymbol{\nu}) \times \boldsymbol{\nu}$ sur Γ_A , avec :

- Pour une condition aux limites d'onde entrante : $\mathbf{e}^* \times \boldsymbol{\nu} = \mathbf{e}_i \times \boldsymbol{\nu}$, où $\mathbf{e}_i$ est une donnée,
- Pour une condition aux limites absorbante : $\mathbf{e}^* \times \boldsymbol{\nu} = 0$.

Posons $\mathcal{E} \times \boldsymbol{\nu} = \mathbf{e} \times \boldsymbol{\nu}$. On suppose que $\mathbf{e} \times \boldsymbol{\nu}$ est connu (c'est-à-dire que $(\boldsymbol{\nu} \times \mathcal{B}) \times \boldsymbol{\nu}|_{\Gamma_A}$ est connu). Considérons $\mathcal{V}_{\Gamma_A}$, un voisinage de Γ_A ne contenant pas de singularité.

Localement, $\mathcal{E}|_{\mathcal{V}_{\Gamma_A}} \in \mathcal{H}^1(\mathcal{V}_{\Gamma_A})$ [36], (rem. 1, p. 562), d'après [5], (thm. 2.9 p. 829 et thm. 2.12 p. 830). De plus, $\mathbf{e} \times \boldsymbol{\nu}|_{\Gamma_C} = 0$, et d'après [36] (prop. 2.7 p. 15), $\mathbf{e} \times \boldsymbol{\nu}|_{\Gamma_A} \in \mathcal{H}_\perp^{1/2}(\Gamma_A)$. D'où, en prolongeant $\mathbf{e} \times \boldsymbol{\nu}|_{\Gamma_A}$ par 0 sur Γ_C , on a : $\mathbf{e} \times \boldsymbol{\nu}|_{\partial\Omega} \in \mathcal{H}_\perp^{1/2}(\partial\Omega)$. Ainsi, dans la suite de cette partie, on considère toujours que :

$\mathcal{E} \times \boldsymbol{\nu}$ sur $\partial\Omega$ est donné par $\mathbf{e} \times \boldsymbol{\nu} \in \mathcal{H}_\perp^{1/2}(\partial\Omega)$,
s'annulant au voisinage des coins et des arêtes rentrants de Ω .

Notons que comme $\mathcal{E} \in \mathcal{H}(\mathbf{rot}, \Omega)$, on a aussi $\mathbf{e} \times \boldsymbol{\nu}|_{\partial\Omega} \in \mathcal{H}_\parallel^{-1/2}(\text{div}_{\partial\Omega}, \partial\Omega)$. Cette hypothèse n'est en aucun cas restrictive, et correspond à une réalité de la modélisation. En effet, comme on l'a mentionné plus haut, on peut toujours placer la frontière artificielle de façon à ne pas couper les singularités géométriques. On en déduit immédiatement la proposition suivante :

Proposition 8.1 Γ_A *ne coupant pas de singularité géométrique*, il existe un relèvement régulier $\mathcal{E}^r \in \mathcal{H}^1(\Omega)$ de $\mathbf{e} \times \boldsymbol{\nu}$ tel que : $\mathcal{E}^r \times \boldsymbol{\nu}|_{\partial\Omega} = \mathbf{e} \times \boldsymbol{\nu}|_{\partial\Omega}$ sur $\partial\Omega$.

Au contraire du cas bidimensionnel, on ne peut pas facilement construire de relèvement comme décrit dans la démonstration de la proposition 2.2.

Nous étudions le problème quasi-électrostatique tridimensionnel, qui s'écrit mathématiquement de la façon suivante :

Trouver $\mathcal{E} \in \mathcal{X}_{\mathcal{E}}$ tel que :

$$\mathbf{rot} \mathcal{E} = \mathbf{f} \quad \text{dans } \Omega \quad \mathbf{f} \in \mathcal{L}^2(\Omega), \quad (8.1)$$

$$\operatorname{div} \mathcal{E} = g \quad \text{dans } \Omega \quad g \in L^2(\Omega), \quad (8.2)$$

$$\mathcal{E} \times \boldsymbol{\nu} = \mathbf{e} \times \boldsymbol{\nu} \quad \text{sur } \partial\Omega \quad \mathbf{e} \times \boldsymbol{\nu} \in \mathcal{L}_t^2(\partial\Omega), \quad (8.3)$$

où $\mathbf{f} = \partial_t \mathcal{B}$ et $g = \rho/\varepsilon_0$ sont connus. Lorsque le problème est statique, $\mathbf{f} = 0$.

En vertu de la relation (7.4), $\operatorname{div} \mathbf{f} = 0$, d'où $\mathbf{f} \in \mathcal{H}(\operatorname{div}^0, \Omega)$. D'autre part, on a la relation de compatibilité suivante entre $\mathbf{e}$ et $\mathbf{f}$ (propriété 7.4).

$$\mathbf{f} \cdot \boldsymbol{\nu}|_{\partial\Omega} = \mathbf{rot} \mathcal{E} \cdot \boldsymbol{\nu}|_{\partial\Omega} = \operatorname{div}_{\partial\Omega}(\mathbf{e} \times \boldsymbol{\nu}|_{\partial\Omega}). \quad (8.4)$$

Définitions 8.2 On appelle la norme du graphe (ou norme naturelle) de $\mathcal{X}_{\mathcal{E}}$ la quantité suivante :

$$\|\mathbf{v}\|_{0, \mathbf{rot}, \operatorname{div}, \mathcal{L}_t^2(\partial\Omega)} := \left(\|\mathbf{v}\|_0^2 + \|\mathbf{rot} \mathbf{v}\|_0^2 + \|\operatorname{div} \mathbf{v}\|_0^2 + \int_{\partial\Omega} |\mathbf{v} \times \boldsymbol{\nu}|^2 d\Sigma \right)^{1/2}.$$

On appelle la semi-norme de $\mathcal{X}_{\mathcal{E}}$ la quantité suivante :

$$|\mathbf{v}|_{\mathbf{rot}, \operatorname{div}, \mathcal{L}_t^2(\partial\Omega)} := \left(\|\mathbf{rot} \mathbf{v}\|_0^2 + \|\operatorname{div} \mathbf{v}\|_0^2 + \int_{\partial\Omega} |\mathbf{v} \times \boldsymbol{\nu}|^2 d\Sigma \right)^{1/2}.$$

P. Fernandez et G. Gilardi ont montré dans [61] le théorème suivant :

Théorème 8.3 L'injection de $\mathcal{X}_{\mathcal{E}}$ dans $\mathcal{L}^2(\Omega)$ est compacte.

Ce théorème permet de montrer le lemme suivant :

Lemme 8.4 Il existe une constante $C > 0$ ne dépendant que de Ω telle que :

$$\forall \mathbf{v} \in \mathcal{X}_{\mathcal{E}}, \quad \|\mathbf{v}\|_0 \leq C |\mathbf{v}|_{\mathbf{rot}, \operatorname{div}, \mathcal{L}_t^2(\partial\Omega)}. \quad (8.5)$$

DÉMONSTRATION. Raisonnons par l'absurde. Supposons que $\forall n \in \mathbb{N}^*$, il existe $\mathbf{v}_n \in \mathcal{X}_{\mathcal{E}}$ tel que :

$$\|\mathbf{v}_n\|_0 = 1 \text{ et } |\mathbf{v}_n|_{\mathbf{rot}, \operatorname{div}, \mathcal{L}_t^2(\partial\Omega)} \leq \frac{1}{n}. \quad (8.6)$$

La suite $(\mathbf{v}_n)_{n \in \mathbb{N}^*}$ est bornée dans $\mathcal{X}_{\mathcal{E}}$. Comme l'injection de $\mathcal{X}_{\mathcal{E}}$ dans $\mathcal{L}^2(\Omega)$ est compacte, il existe une sous-suite extraite de $(\mathbf{v}_n)_{n \in \mathbb{N}^*}$, que l'on assimile à $(\mathbf{v}_n)_{n \in \mathbb{N}^*}$, qui converge fortement dans $\mathcal{L}^2(\Omega)$ vers $\mathbf{v}$. On en déduit que $(\operatorname{div} \mathbf{v}_n)_{n \in \mathbb{N}^*}$ et $(\mathbf{rot} \mathbf{v}_n)_{n \in \mathbb{N}^*}$ convergent au sens des distributions vers $\operatorname{div} \mathbf{v}$ et $\mathbf{rot} \mathbf{v}$ respectivement. D'après (8.6), on obtient la convergence forte dans $\mathcal{L}^2(\Omega)$ et la valeur de la limite :

$$\begin{cases} \mathbf{rot} \mathbf{v}_n & \rightarrow \mathbf{rot} \mathbf{v} = 0 \text{ dans } \mathcal{L}^2(\Omega), \\ \operatorname{div} \mathbf{v}_n & \rightarrow \operatorname{div} \mathbf{v} = 0 \text{ dans } L^2(\Omega). \end{cases}$$

On a alors $\mathbf{v}_n \rightarrow \mathbf{v}$ dans $\mathcal{H}(\mathbf{rot}^0, \Omega) \cap \mathcal{H}(\operatorname{div}^0, \Omega)$. On en déduit la convergence de la trace : $\mathbf{v}_n \times \boldsymbol{\nu}|_{\partial\Omega} \rightarrow \mathbf{v} \times \boldsymbol{\nu}|_{\partial\Omega}$ dans $\mathcal{H}_{\parallel}^{-1/2}(\operatorname{div}_{\partial\Omega}, \partial\Omega)$. D'après (8.6), on a $\mathbf{v} \times \boldsymbol{\nu}|_{\partial\Omega} = 0$ au sens des distributions.

Ω est simplement connexe et $\mathbf{v} \in \mathcal{L}^2(\Omega)$ est tel que $\mathbf{rot} \mathbf{v} = 0$ dans Ω . On en déduit qu'il existe un unique $p \in H^1(\Omega)$ tel que $\mathbf{v} = \mathbf{grad} p$ [65] (thm. 2.9, p. 31). Comme $\mathbf{grad} p \times \boldsymbol{\nu}|_{\partial\Omega} = 0$, p est constant sur chaque composante connexe de $\partial\Omega$ [65] (rem. 1.3, p. 35). Posons a sa valeur. p satisfait alors le problème de Dirichlet homogène suivant :

Trouver $p \in H^1(\Omega)$ tel que : $\begin{cases} \Delta p = 0 & \text{dans } \Omega, \\ p|_{\partial\Omega} = a & \text{sur } \partial\Omega. \end{cases}$

D'après le théorème de Lax-Milgram, ce problème admet une unique solution : $p = a$ dans Ω . On en déduit que $\mathbf{v} = \mathbf{grad} p = 0$. Comme $\mathbf{v}_n \rightarrow \mathbf{v}$ dans $\mathcal{L}^2(\Omega)$, on a nécessairement $\|\mathbf{v}\|_0 = 1$, et donc $\mathbf{v} \neq 0$, ce qui contredit la conclusion précédente. $\square$

On en déduit alors le théorème qui suit :

Théorème 8.5 *Dans $\mathcal{X}_{\mathcal{E}}$, la semi-norme est équivalente à la norme du graphe : la semi-norme définit une norme sur $\mathcal{X}_{\mathcal{E}}$.*

DÉMONSTRATION. du théorème 8.5 L'inégalité (8.5) permet de montrer que :

$$C' \|\mathbf{v}\|_{0,\mathbf{rot},\text{div},\mathcal{L}_t^2(\partial\Omega)} \leq |\mathbf{v}|_{\mathbf{rot},\text{div},\mathcal{L}_t^2(\partial\Omega)} \leq \|\mathbf{v}\|_{0,\mathbf{rot},\text{div},\mathcal{L}_t^2(\partial\Omega)}, \text{ avec : } C' = 1/\sqrt{C^2 + 1}.$$

Les normes $\|\mathbf{v}\|_{0,\mathbf{rot},\text{div},\mathcal{L}_t^2(\partial\Omega)}$ et $|\mathbf{v}|_{\mathbf{rot},\text{div},\mathcal{L}_t^2(\partial\Omega)}$ sont donc équivalentes. $\square$

Définitions 8.6 *On peut définir la norme de $\mathcal{X}_{\mathcal{E}}$ par :*

$$\|\mathbf{v}\|_{\mathcal{X}_{\mathcal{E}}} = \left(\|\mathbf{rot} \mathbf{v}\|_0^2 + \|\text{div} \mathbf{v}\|_0^2 + \int_{\partial\Omega} |\mathbf{v} \times \boldsymbol{\nu}|^2 d\Sigma \right)^{1/2}.$$

Le produit scalaire dans $\mathcal{X}_{\mathcal{E}}$ étant ainsi défini :

$$\begin{aligned} (\mathbf{u}, \mathbf{v})_{\mathcal{X}_{\mathcal{E}}} &= (\mathbf{rot} \mathbf{u}, \mathbf{rot} \mathbf{v})_0 + (\text{div} \mathbf{u}, \text{div} \mathbf{v})_0 + \int_{\partial\Omega} (\mathbf{u} \times \boldsymbol{\nu}) \cdot (\mathbf{v} \times \boldsymbol{\nu}) d\Sigma \\ &= \int_{\Omega} \mathbf{rot} \mathbf{u} \cdot \mathbf{rot} \mathbf{v} d\Omega + \int_{\Omega} \text{div} \mathbf{u} \text{div} \mathbf{v} d\Omega + \int_{\partial\Omega} (\mathbf{u} \times \boldsymbol{\nu}) \cdot (\mathbf{v} \times \boldsymbol{\nu}) d\Sigma. \end{aligned}$$

$\mathcal{X}_{\mathcal{E}}^0$ étant un sous-espace de $\mathcal{X}_{\mathcal{E}}$, sa norme est définie par : $\|\mathbf{v}\|_{\mathcal{X}_{\mathcal{E}}} = (\|\mathbf{rot} \mathbf{v}\|_0^2 + \|\text{div} \mathbf{v}\|_0^2)^{1/2}$. Le produit scalaire dans $\mathcal{X}_{\mathcal{E}}^0$ est défini de la façon suivante :

$$(\mathbf{u}, \mathbf{v})_{\mathcal{X}_{\mathcal{E}}^0} = (\mathbf{rot} \mathbf{u}, \mathbf{rot} \mathbf{v})_0 + (\text{div} \mathbf{u}, \text{div} \mathbf{v})_0 = \int_{\Omega} \mathbf{rot} \mathbf{u} \cdot \mathbf{rot} \mathbf{v} d\Omega + \int_{\Omega} \text{div} \mathbf{u} \text{div} \mathbf{v} d\Omega.$$

Remarque 8.7 *Lorsque Ω n'est pas connexe, il apparaît un terme supplémentaire dans la norme : c'est la charge sur chaque composante connexe Γ_k , égale à $\int_{\Gamma_k} \mathbf{u} \cdot \boldsymbol{\nu} d\Sigma$ $\int_{\Gamma_k} \mathbf{v} \cdot \boldsymbol{\nu} d\Sigma$. Cela correspond à la valeur du potentiel électrostatique à la surface de chaque conducteur parfait [36].*

Posons $\mathcal{E}^0 = \mathcal{E} - \mathcal{E}^r$. $\mathcal{E}^0 \in \mathcal{X}_{\mathcal{E}}^0$ satisfait les équations suivantes :

$$\mathbf{rot} \mathcal{E}^0 = \mathbf{f}^0 \quad \text{dans } \Omega \quad \mathbf{f}^0 \in \mathcal{L}^2(\Omega), \quad (8.7)$$

$$\text{div} \mathcal{E}^0 = g^0 \quad \text{dans } \Omega \quad g^0 \in L^2(\Omega), \quad (8.8)$$

avec $\mathbf{f}^0 = \mathbf{f} - \mathbf{rot} \mathcal{E}^r$ et $g^0 = g - \text{div} \mathcal{E}^r$. La relation de compatibilité (8.4) devient : $\mathbf{f}^0 \cdot \boldsymbol{\nu}|_{\partial\Omega} = 0$, d'où :

$$\mathbf{f}^0 \in \mathcal{H}_0(\text{div}^0, \Omega). \quad (8.9)$$

Lorsque $\mathbf{f} = 0$ et $\mathbf{e} = 0$, le problème (8.1)-(8.3) peut être résolu par calcul du potentiel statique, $\phi_0 \in H_0^1(\Omega)$ tel que : $-\Delta\phi_0 = g$. Le développement de cette méthode fera l'objet des sections 8.2 (problème continu) et 10.2 (problème discrétisé).

D'autre part, le problème (8.1)-(8.3) peut se résoudre numériquement par les éléments finis de Lagrange continus P_k , en calculant $\mathcal{E}$ directement dans $\mathcal{X}_{\mathcal{E}}$. Cette méthode sera détaillée dans les sections 8.3 (problème continu) et 10.3 (problème discrétisé).

Il est aussi possible de calculer $\mathcal{E}^0$ par les éléments finis Lagrange continus P_k dans $\mathcal{X}_{\mathcal{E}}^0$ si Ω est convexe : voir les sections 8.4 (problème continu) et 10.4 (problème discrétisé). Lorsque Ω n'est pas convexe, le calcul de $\mathcal{E}^0$ par les éléments finis de Lagrange continus dans $\mathcal{X}_{\mathcal{E}}^0$ ne converge pas car on ne capte pas les parties singulières du champ électrique. En revanche, il converge dans l'espace à poids $\mathcal{X}_{\mathcal{E},\gamma}^0$, comme c'est détaillé dans les sections 8.5 (problème continu) et 10.5 (problème discrétisé). On peut adapter la méthode du complément singulier dans le cas d'un domaine prismatique, ce qui fait l'objet des sections 8.6 (problème continu) et 10.6 (problème discret).

8.2 Champ électrostatique 3D

Le potentiel ϕ_0 est solution du problème de Dirichlet suivant : Trouver $\phi_0 \in H_0^1(\Omega)$ tel que :

$$-\Delta\phi_0 = g \text{ dans } \Omega. \quad (8.10)$$

Soit $\mathcal{E} = -\mathbf{grad} \phi_0 \in \mathcal{L}^2(\Omega)$. On a : $\text{div } \mathcal{E} = g$, $\mathbf{rot} \mathcal{E} = 0$, et $\mathcal{E} \times \boldsymbol{\nu}_{|\partial\Omega} = -\mathbf{grad} \phi_0 \times \boldsymbol{\nu}_{|\partial\Omega} = 0$: $\mathcal{E}$ est solution de (8.1)-(8.3) avec $\mathbf{f} = 0$ et $\mathbf{e} \times \boldsymbol{\nu}_{|\partial\Omega} = 0$. Le potentiel ϕ_0 est calculé par les éléments finis continus P_k de Lagrange, et $\mathcal{E}$ est approché par l'élément fini *discontinu* P_{k-1} composante par composante.

Proposition 8.8 *Le problème (8.10) est équivalent au problème variationnel suivant (FVD) : Trouver $\phi_0 \in H_0^1(\Omega)$ tel que :*

$$\forall u \in H_0^1(\Omega), (\mathbf{grad} \phi_0, \mathbf{grad} u)_0 = \int_{\Omega} g u \, d\Omega \quad (8.11)$$

D'après le théorème de Lax-Milgram, le problème (FVD) admet une unique solution $\phi_0 \in H_0^1(\Omega)$.

Remarque 8.9 *Notons que ϕ_0 appartient par construction à $\Phi_{\mathcal{D}}$, ce qui prouve que :*

$$\mathcal{H}_0(\mathbf{rot}^0, \Omega) \cap \mathcal{H}(\text{div}, \Omega) = \mathbf{grad} \Phi_{\mathcal{D}}.$$

Contrairement au cas bidimensionnel, l'espace singulier de $\Phi_{\mathcal{D}}$, $\Phi_{\mathcal{D}}^S$ tel que : $\Phi_{\mathcal{D}} = \Phi_{\mathcal{D}}^S \oplus (\Phi_{\mathcal{D}} \cap H^2(\Omega))$ n'est pas de dimension finie. Ainsi, on ne peut pas décomposer ϕ_0 en une partie $H^2(\Omega)$ et une partie exacte, écrite sous la forme d'une somme finie.

Comme l'approximation du champ électrique par le potentiel électrostatique est calculée par dérivation d'éléments continus P_k , il y a conformité uniquement dans $\mathcal{H}(\mathbf{rot}, \Omega)$ et pas dans $\mathcal{X}_{\mathcal{E}}$.

8.3 Champ électrique 3D : CL naturelles

Étudions la résolution du problème (8.1)-(8.3) dans $\mathcal{X}_{\mathcal{E}}$.

Proposition 8.10 *Le problème (8.1)-(8.3) est équivalent au problème variationnel suivant : Trouver $\mathcal{E} \in \mathcal{X}_{\mathcal{E}}$ tel que :*

$$\forall \mathcal{F} \in \mathcal{X}_{\mathcal{E}}, (\mathcal{E}, \mathcal{F})_{\mathcal{X}_{\mathcal{E}}} = \mathcal{L}_{\mathcal{E}}(\mathcal{F}), \quad (8.12)$$

où $\mathcal{L}_{\mathcal{E}}$ est la forme linéaire suivante :

$$\begin{aligned} \mathcal{L}_{\mathcal{E}} : \mathcal{X}_{\mathcal{E}} &\rightarrow \mathbb{R} \\ \mathcal{F} &\mapsto (\mathbf{f}, \mathbf{rot} \mathcal{F})_0 + (g, \text{div } \mathcal{F})_0 + \int_{\partial\Omega} (\mathbf{e} \times \boldsymbol{\nu}) \cdot (\mathcal{F} \times \boldsymbol{\nu}) \, d\Sigma. \end{aligned}$$

Comme $\mathcal{L}_{\mathcal{E}} \in \mathcal{X}'_{\mathcal{E}}$, d'après le théorème 1.7, p. 1.7 le problème (8.12) admet une solution unique $\mathcal{E} \in \mathcal{X}_{\mathcal{E}}$, qui dépend continûment des données $\mathbf{f}$, g et $\mathbf{e} \times \boldsymbol{\nu}_{|\partial\Omega}$.

DÉMONSTRATION. Il est clair que si $\mathcal{E}$ satisfait (8.1)-(8.3), alors $\mathcal{E}$ satisfait (8.12).

Montrons la réciproque. Soit $h \in L^2(\Omega)$. Il existe une unique fonction $\phi \in H_0^1(\Omega)$ telle que $\Delta\phi = h$. Posons $\mathcal{F} = \mathbf{grad} \phi$. On a alors : $\mathcal{F} \in \mathcal{L}^2(\Omega)$, $\mathbf{rot} \mathcal{F} = 0 \in \mathcal{L}^2(\Omega)$, $\text{div} \mathcal{F} = h \in L^2(\Omega)$, $\mathcal{F} \times \boldsymbol{\nu} = \mathbf{grad}_{\partial\Omega} \phi \times \boldsymbol{\nu} = 0 \in \mathcal{L}_t^2(\partial\Omega)$, donc $\mathcal{F} \in \mathcal{X}_{\mathcal{E}}$.

Lorsqu'on injecte $\mathcal{F}$ dans (8.12), on obtient : $(\text{div} \mathcal{E}, h)_0 = (g, h)_0$. Comme $\text{div} \mathcal{E}$ et g appartiennent à $L^2(\Omega)$, on en déduit : $\text{div} \mathcal{E} = g$ dans $L^2(\Omega)$. L'équation (8.12) se réduit donc à :

Trouver $\mathcal{E} \in \mathcal{X}_{\mathcal{E}}$ tel que :

$$\forall \mathcal{F} \in \mathcal{X}_{\mathcal{E}} \quad (\mathbf{rot} \mathcal{E}, \mathbf{rot} \mathcal{F})_0 + \int_{\partial\Omega} (\mathcal{E} \times \boldsymbol{\nu}) \cdot (\mathcal{F} \times \boldsymbol{\nu}) d\Sigma = (\mathbf{f}, \mathbf{rot} \mathcal{F})_0 + \int_{\partial\Omega} (\mathbf{e} \times \boldsymbol{\nu}) \cdot (\mathcal{F} \times \boldsymbol{\nu}) d\Sigma.$$

D'après la proposition 7.6, il existe $\mathcal{E}^* \in \mathcal{X}_{\mathcal{E}}$ tel que : $\mathcal{E} \times \boldsymbol{\nu}^* = \mathbf{e} \times \boldsymbol{\nu} \in \mathcal{L}_t^2(\partial\Omega) \cap \mathcal{H}_{\parallel}^{-1/2}(\text{div}_{\partial\Omega}, \partial\Omega)$.

Posons $\mathcal{E}' = \mathcal{E} - \mathcal{E}^* \in \mathcal{X}_{\mathcal{E}}^0$. On obtient que $\mathcal{E}' \in \mathcal{X}_{\mathcal{E}}^0$ satisfait :

$$\forall \mathcal{F} \in \mathcal{X}_{\mathcal{E}}, (\mathbf{rot} \mathcal{E}', \mathbf{rot} \mathcal{F})_0 = (\mathbf{f} - \mathbf{rot} \mathcal{E}^*, \mathbf{rot} \mathcal{F})_0. \quad (8.13)$$

Posons $\mathbf{f}' = \mathbf{f} - \mathbf{rot} \mathcal{E}^*$. On remarque que : $\mathbf{f}' \in \mathcal{L}^2(\Omega)$, $\text{div} \mathbf{f}' = 0 \in L^2(\Omega)$, et d'après (8.4), on a :

$$\mathbf{f}' \cdot \boldsymbol{\nu}_{|\partial\Omega} = \mathbf{f} \cdot \boldsymbol{\nu}_{|\partial\Omega} - \mathbf{rot} \mathcal{E}^* \cdot \boldsymbol{\nu}_{|\partial\Omega} = \text{div}_{\partial\Omega}(\mathbf{e} \times \boldsymbol{\nu}_{|\partial\Omega}) - \text{div}_{\partial\Omega}(\mathbf{e} \times \boldsymbol{\nu}_{|\partial\Omega}) = 0.$$

On en déduit que $\mathbf{f}' \in \mathcal{H}_0(\text{div}^0, \Omega)$. D'après [65] (thm. 3.4 p. 45), il existe $\mathcal{F}^* \in \mathcal{X}_{\mathcal{E}}^0$ tel que : $\mathbf{f}' = \mathbf{rot} \mathcal{F}^*$. Ainsi, (8.13) devient : $(\mathbf{rot} \mathcal{E}', \mathbf{rot} \mathcal{F})_0 = (\mathbf{rot} \mathcal{F}^*, \mathbf{rot} \mathcal{F})_0$, ce qui se réécrit : $(\mathbf{rot}(\mathcal{E}' - \mathcal{F}^*), \mathbf{rot} \mathcal{F})_0 = 0$. Choisissons $\mathcal{F} = \mathcal{E}' - \mathcal{F}^* = \mathcal{E} - (\mathcal{E}^* + \mathcal{F}^*)$.

D'où : $\|\mathbf{rot}(\mathcal{E} - (\mathcal{E}^* + \mathcal{F}^*))\|_0^2 = 0$. De cette égalité, on déduit que : $\mathbf{rot} \mathcal{E} = \mathbf{rot}(\mathcal{E}^* + \mathcal{F}^*) = \mathbf{rot} \mathcal{E}^* + \mathbf{f}' = \mathbf{f}$ dans $\mathcal{L}^2(\Omega)$. Finalement, (8.12) devient :

Trouver $\mathcal{E} \in \mathcal{X}_{\mathcal{E}}$ tel que :

$$\forall \mathcal{F} \in \mathcal{X}_{\mathcal{E}}, \int_{\partial\Omega} (\mathcal{E} \times \boldsymbol{\nu}) \cdot (\mathcal{F} \times \boldsymbol{\nu}) d\Sigma = \int_{\partial\Omega} (\mathbf{e} \times \boldsymbol{\nu}) \cdot (\mathcal{F} \times \boldsymbol{\nu}) d\Sigma,$$

ce qui se réécrit : $\forall \mathcal{F} \in \mathcal{X}_{\mathcal{E}}, \int_{\partial\Omega} (\mathcal{E} \times \boldsymbol{\nu} - \mathbf{e} \times \boldsymbol{\nu}) \cdot (\mathcal{F} \times \boldsymbol{\nu}) d\Sigma = 0$. En prenant $\mathcal{F} = \mathcal{E} - \mathcal{E}^*$, on a

alors : $\int_{\partial\Omega} |\mathcal{E} \times \boldsymbol{\nu} - \mathbf{e} \times \boldsymbol{\nu}|^2 d\Sigma = 0$, d'où : $\mathcal{E} \times \boldsymbol{\nu}_{|\partial\Omega} = \mathbf{e} \times \boldsymbol{\nu}_{|\partial\Omega}$ dans $\mathcal{L}_t^2(\partial\Omega)$. □

Nous avons donc montré que $\mathcal{E}$ satisfait les équations (8.1)-(8.3). Les équations (8.1)-(8.3) et (8.12) sont donc équivalentes sous les hypothèses :

$$\mathbf{f} \in \mathcal{H}(\text{div}^0, \Omega), g \in L^2(\Omega), \mathbf{e} \times \boldsymbol{\nu}_{|\partial\Omega} \in \mathcal{L}_t^2(\partial\Omega) \cap \mathcal{H}_{\parallel}^{-1/2}(\text{div}_{\partial\Omega}, \Omega);$$

et la relation de compatibilité : $\mathbf{f} \cdot \boldsymbol{\nu}_{|\partial\Omega} = \text{div}_{\partial\Omega}(\mathbf{e} \times \boldsymbol{\nu}_{|\partial\Omega})$.

Notons pour finir que, comme $\mathcal{X}_{\mathcal{E}} \cap \mathcal{H}^1(\Omega)$ est dense dans $\mathcal{X}_{\mathcal{E}}$ [40, 51], on peut approcher la solution de (8.12) par les éléments finis de Lagrange continus P_k .

Remarque 8.11 *Au contraire de la méthode des potentiels, cette méthode s'applique au cas dépendant du temps.*

8.4 Champ électrique 3D : CL essentielles

La condition aux limites étant naturelle dans la formulation (8.12), on s'intéresse également à la résolution de (8.7)-(8.8), dans l'espace $\mathcal{X}_\varepsilon^0$ dans lequel la condition aux limites est prise en compte de façon essentielle.

Proposition 8.12 *Le problème (8.7)-(8.8) est équivalent au problème variationnel suivant : Trouver $\mathcal{E}^0 \in \mathcal{X}_\varepsilon^0$ tel que :*

$$\forall \mathcal{F} \in \mathcal{X}_\varepsilon^0, (\mathcal{E}, \mathcal{F})_{\mathcal{X}_\varepsilon^0} = \mathcal{L}_0(\mathcal{F}) - (\mathcal{E}^r, \mathcal{F})_{\mathcal{X}_\varepsilon^0}, \quad (8.14)$$

où $\mathcal{L}_0$ est la forme linéaire suivante :

$$\begin{aligned} \mathcal{L}_0 : \mathcal{X}_\varepsilon^0 &\rightarrow \mathbb{R} \\ \mathcal{F} &\mapsto (\mathbf{f}^0, \mathbf{rot} \mathcal{F})_0 + (g^0, \operatorname{div} \mathcal{F})_0. \end{aligned}$$

Comme $\mathcal{L}_0 \in (\mathcal{X}_\varepsilon^0)'$, d'après le théorème 1.7, le problème (8.14) admet une solution unique $\mathcal{E}^0 \in \mathcal{X}_\varepsilon^0$, qui dépend continûment des données $\mathbf{f}^0$ et g^0 .

DÉMONSTRATION. $\mathcal{X}_\varepsilon^0$ étant un sous-espace de $\mathcal{X}_\varepsilon$, la démonstration est similaire à celle de la proposition 8.10. □

Lorsque Ω est convexe, $\mathcal{X}_\varepsilon^{0,R} = \mathcal{X}_\varepsilon^0$. La solution de (8.14) approchée par les éléments finis de Lagrange continus P_k est fautive lorsque Ω présente des singularités géométriques. Contrairement au cas bidimensionnel, l'espace $\mathcal{X}_\varepsilon^{0,R}$ est de codimension infinie, ainsi, on ne peut pas généraliser la méthode du complément singulier.

8.5 Champ électrique 3D : régularisation à poids

Dans cette section, nous reprenons les résultats de M. Costabel et M. Dauge dans [53], que nous avons résumés pour le problème bidimensionnel dans la section 2.5 de la partie II. Rappelons que lorsque Ω est non-convexe, l'espace $\mathcal{X}_\varepsilon^0 \cap \mathcal{H}^1(\Omega)$ est fermé et strictement inclus dans $\mathcal{X}_\varepsilon^0$, ce qui implique que la solution du problème discrétisé par les éléments finis P_k dans l'espace $\mathcal{X}_\varepsilon^0$ ne converge pas vers la bonne solution. La méthode à poids consiste à déterminer un espace $\mathcal{X}_{\varepsilon,\gamma}^0$ plus gros que $\mathcal{X}_\varepsilon^0$ de façon à avoir la densité de $\mathcal{X}_{\varepsilon,\gamma}^0 \cap \mathcal{H}^1(\Omega)$ dans $\mathcal{X}_{\varepsilon,\gamma}^0$. Cet espace est construit en considérant que la divergence du champ électrique appartient à un espace de type L^2 à poids, où le poids dépend de la distance aux singularités.

La principale différence avec le cas bidimensionnel est qu'il existe des interactions entre les singularités des coins et celles des arêtes, mais dans un polyèdre, un coin étant une intersection d'arêtes, on peut toujours définir le poids à l'aide d'un produit de distances aux arêtes.

Notons que $\mathcal{X}_{\varepsilon,\gamma}^0 \cap \mathcal{H}^1(\Omega) = \mathcal{X}_\varepsilon^{0,R}$. Comme pour le cas bidimensionnel (thm. 2.60, p. 79), on peut décomposer les éléments de $\mathcal{X}_{\varepsilon,\gamma}^0$ en une partie $\mathcal{H}^1$ -régulière et une partie qui n'est pas $\mathcal{H}^1(\Omega)$.

Théorème 8.13 *L'espace $\mathcal{X}_{\varepsilon,\gamma}^0$ se décompose de la façon suivante :*

$$\mathcal{X}_{\varepsilon,\gamma}^0 = \mathcal{X}_\varepsilon^{0,R} \oplus \mathbf{grad} \Phi_\gamma.$$

Ainsi, pour avoir l'injection compacte de $\mathcal{L}^2(\Omega)$ dans $\mathcal{X}_{\varepsilon,\gamma}^0$, on cherche à avoir l'injection compacte de $\mathcal{L}^2(\Omega)$ dans $\mathbf{grad} \Phi_\gamma$, et donc celle de $H^1(\Omega)$ dans Φ_γ , ce qui est réalisé lorsque V_γ^0 s'injecte compactement dans $H^{-1}(\Omega)$. Or, comme pour le cas bidimensionnel (prop. 2.62, p. 80), on a le théorème suivant :

Théorème 8.14 *L'injection de V_γ^0 dans $H^{-1}(\Omega)$ est compacte si et seulement si $\gamma < 1$.*

On obtient alors l'injection compacte de $\mathcal{X}_{\varepsilon,\gamma}^0$ dans $\mathcal{L}^2(\Omega)$, ce qui nous permet d'obtenir l'équivalence entre la norme du graphe et la semi-norme dans $\mathcal{X}_{\varepsilon,\gamma}^0$:

Proposition 8.15 *La norme du graphe dans $\mathcal{X}_{\varepsilon,\gamma}^0$ est équivalente à la semi-norme si et seulement si $\gamma < 1$. La norme dans $\mathcal{X}_{\varepsilon,\gamma}^0$ est alors définie ainsi :*

$$\|\mathbf{u}\|_{\mathcal{X}_{\varepsilon,\gamma}^0} = \left(\|\operatorname{rot} \mathbf{u}\|_0^2 + \|\operatorname{div} \mathbf{u}\|_{0,\gamma}^2 \right)^{1/2}. \quad (8.15)$$

La preuve se fait par l'absurde, et est similaire à celle de la proposition 2.56 (p. 78).

Afin d'obtenir la convergence des éléments finis de Galerkin dans $\mathcal{X}_{\varepsilon,\gamma}^0$, on cherche la condition sur γ de façon à obtenir la densité de $\mathcal{X}_{\varepsilon}^{0,R}$ dans $\mathcal{X}_{\varepsilon,\gamma}^0$, c'est-à-dire la densité de $\mathcal{X}_{\varepsilon}^{0,R}$ dans $\mathbf{grad} \Phi_\gamma$. Comme dans le cas bidimensionnel (prop. 2.63, p. 80), on a le résultat suivant :

Proposition 8.16 *Quelque soit γ , les fonctions $C^\infty(\overline{\Omega})$ à trace nulle sur $\partial\Omega$ sont denses dans $V_\gamma^2 \cap H_0^1(\Omega)$.*

Le théorème fondamental suivant (thm. 2.64, p. 80 pour le 2D) nous donne les valeurs de γ telles que $\Phi_\gamma \subset V_\gamma^2$:

Théorème 8.17 *Pour tout γ tel que : $1 - \alpha < \gamma \leq 1$, l'opérateur Δ est un isomorphisme de $V_\gamma^2 \cap H_0^1(\Omega)$ dans V_γ^0 . De plus, $H^2(\Omega) \cap H_0^1(\Omega)$ est dense dans Φ_γ .*

Ainsi, pour $1 - \alpha < \gamma \leq 1$, les fonctions $C^\infty(\overline{\Omega})$ à trace nulle sur $\partial\Omega$ sont denses dans Φ_γ , et par conséquent, $\mathcal{X}_{\varepsilon}^{0,R}$ est dense dans $\mathbf{grad} \Phi_\gamma$. En pratique, cette condition permet d'un point de vue numérique de capter les singularités du champ électrique. Le choix optimal de γ pour lequel on a l'injection compacte de $\mathcal{L}^2(\Omega)$ dans $\mathcal{X}_{\varepsilon,\gamma}^0$ et la densité des fonctions $\mathcal{H}^1$ -régulières dans $\mathcal{X}_{\varepsilon,\gamma}^0$ est donc :

$$1 - \alpha < \gamma < 1 \quad (8.16)$$

Proposition 8.18 *Le problème (8.7)-(8.8) accompagné de l'hypothèse (8.9) est équivalent à la formulation variationnelle suivante :*

Trouver $\mathcal{E}^0 \in \mathcal{X}_{\varepsilon,\gamma}^0$ tel que :

$$\forall \mathcal{F} \in \mathcal{X}_{\varepsilon,\gamma}^0, (\mathcal{E}^0, \mathcal{F})_{\mathcal{X}_{\varepsilon,\gamma}^0} = \mathcal{L}_\gamma(\mathcal{F}) - (\mathcal{E}^r, \mathcal{F})_{\mathcal{X}_{\varepsilon,\gamma}^0}, \quad (8.17)$$

où $\mathcal{L}_\gamma$ est la forme linéaire continue suivante :

$$\begin{aligned} \mathcal{L}_\gamma : \mathcal{X}_{\varepsilon,\gamma}^0 &\rightarrow \mathbb{R} \\ \mathcal{F} &\mapsto (g, \operatorname{div} \mathcal{F})_{0,\gamma} + (\mathbf{f}, \operatorname{rot} \mathcal{F})_0. \end{aligned}$$

DÉMONSTRATION. Soit $h \in V_\gamma^0$. Considérons $\phi \in \Phi_\gamma$ tel que $\Delta\phi = h$. Le vecteur $\mathcal{F} = \mathbf{grad} \phi$ est dans $\mathcal{X}_{\varepsilon,\gamma}^0$. En l'injectant (8.17), on obtient :

$$(\operatorname{div} \mathcal{E}^0, h)_{0,\gamma} = (g - \operatorname{div} \mathcal{E}^r, h)_{0,\gamma}$$

On en déduit que : $\operatorname{div} \mathcal{E}^0 = g - \operatorname{div} \mathcal{E}^r := g^0$ au sens de V_γ^0 . L'équation (8.17) se réduit alors à : Trouver $\mathcal{E}^0 \in \mathcal{X}_{\mathcal{E},\gamma}^0$ tel que :

$$\forall \mathcal{F} \in \mathcal{X}_{\mathcal{E},\gamma}^0, (\mathbf{rot} \mathcal{E}^0, \mathbf{rot} \mathcal{F})_0 = (\mathbf{f} - \mathbf{rot} \mathcal{E}^r, \mathbf{rot} \mathcal{F})_0. \quad (8.18)$$

Comme $\mathbf{f} - \mathbf{rot} \mathcal{E}^r := \mathbf{f}^0 \in \mathcal{H}_0(\operatorname{div}^0, \Omega)$, d'après [65] (thm. 3.4 p. 45), il existe $\mathcal{F}^* \in \mathcal{X}_{\mathcal{E}}^0$ tel que : $\mathbf{f}^0 = \mathbf{rot} \mathcal{F}^*$. D'où, (8.18) se réécrit :

Trouver $\mathcal{E}^0 \in \mathcal{X}_{\mathcal{E},\gamma}^0$ tel que :

$$\forall \mathcal{F} \in \mathcal{X}_{\mathcal{E},\gamma}^0, (\mathbf{rot} (\mathcal{E}^0 - \mathcal{F}^*), \mathbf{rot} \mathcal{F})_0 = 0. \quad (8.19)$$

Comme $\mathcal{X}_{\mathcal{E}}^0 \subset \mathcal{X}_{\mathcal{E},\gamma}^0$, on peut injecter $\mathcal{E}^0 - \mathcal{F}^* \in \mathcal{X}_{\mathcal{E},\gamma}^0$ dans (8.19), ce qui donne : $\|\mathbf{rot} \mathcal{E}^0 - \mathbf{f}^0\|_0 = 0$. On en déduit que $\mathbf{rot} \mathcal{E}^0 = \mathbf{f}^0$ au sens $\mathcal{L}^2(\Omega)$. D'où le résultat. $\square$

Soit $\gamma < 1$. Alors $(\cdot, \cdot)_{\mathcal{X}_{\mathcal{E},\gamma}^0}$ est une forme bilinéaire coercitive sur $\mathcal{X}_{\mathcal{E},\gamma}^0 \times \mathcal{X}_{\mathcal{E},\gamma}^0$. D'après le théorème de Lax-Milgram, il existe alors une unique solution à la formulation variationnelle (8.17), qui dépend continûment des données $\mathbf{f}^0$ et g^0 .

Soit γ satisfaisant (8.16). $\mathcal{X}_{\mathcal{E},\gamma}^0 \cap \mathcal{H}^1(\Omega)$ est alors dense dans $\mathcal{X}_{\mathcal{E},\gamma}^0$, on peut donc approcher la solution de (8.17) par les éléments finis de Lagrange continus P_k .

8.6 Champ électrique $2D_{\frac{1}{2}}$: le complément singulier

Nous avons vu dans la section 8.4 que dans les domaines tridimensionnels, l'espace des singularités électromagnétiques n'est pas de dimension finie. On ne peut donc pas généraliser la méthode du complément singulier. Cependant, dans les domaines prismatique, c'est-à-dire invariants selon une direction, on peut décomposer le champ électromagnétique de façon à résoudre une série de problèmes $2D$, pour lesquels on peut appliquer la méthode du complément singulier. Ceci peut être assez avantageux pour diminuer le coût de calcul et garder une bonne précision.

8.6.1 Introduction

Dans ce paragraphe, nous reprenons l'approche de l'article [38], adaptée au problème (8.1)-(8.3), et permettant de résoudre les équations de Maxwell dans un filtre à stubs (voir la figure 2.3). Ω est un prisme droit : $\Omega = \omega \times]0, L[$, où ω est un polygone bidimensionnel du plan $(O, \mathbf{x}, \mathbf{y})$, possédant N_{ar} coins rentrants, d'angles $(\Theta_e)_{e \in \{1, \dots, N_{ar}\}}$ supérieurs à π . Les singularités géométriques de Ω sont donc les N_{ar} arêtes rentrantes d'angles dièdres intérieurs $(\Theta_e)_{e \in \{1, \dots, N_{ar}\}}$.

Soit $\mathcal{F} \in \mathbb{R}^3$. On considère $\mathbf{F} \in \mathbb{R}^2$ et $F_z \in \mathbb{R}$ tels que : $\mathcal{F} = \mathbf{F} + F_z \mathbf{e}_z$, où : $\mathbf{F} = F_x \mathbf{e}_x + F_y \mathbf{e}_y$. Le bord $\partial\omega$ de ω est composé de γ_C et γ_A : $\partial\omega = \overline{\gamma_C} \cup \overline{\gamma_A}$.

Le bord $\partial\Omega$ de Ω est composé de Γ_C , le conducteur parfait et de Γ_A , la frontière artificielle qui borne le domaine. Γ_C et Γ_A sont telles que :

$$\Gamma_C = (\gamma_C \times]0, L[) \cup \{\mathbf{x} \in \partial\Omega : z = 0 \text{ ou } L\} \text{ et } \Gamma_A = \gamma_A \times]0, L[.$$

Comme $\mathcal{E} \times \boldsymbol{\nu}|_{\Gamma_C} = 0$, on a : $\mathbf{E}_z|_{\Gamma_C} = 0$. Γ_z désignera la réunion des bases $\Gamma_C \cap \{z = 0\}$ et $\Gamma_C \cap \{z = L\}$. Sur Γ_z , on a $\boldsymbol{\nu} = \pm \mathbf{e}_z$, d'où, d'après la condition aux limites d'un conducteur parfait (1.28), on a : $\mathbf{E} = 0$ sur Γ_z , c'est-à-dire $E_x = E_y = 0$ sur Γ_z .

Dans cette configuration d'après [18, 42], on a la propriété suivante :

Proposition 8.19 *Pour tout $\mathcal{F} \in \mathcal{X}_{\mathcal{E}}^0$, $F_z \in H_0^1(\Omega)$, et $\partial_z F_x, \partial_z F_y \in L^2(\Omega)$.*

Définitions 8.20 • $\mathcal{X}_{\mathcal{E}}^{A,0}$ désigne le sous-espace de $\mathcal{X}_{\mathcal{E}}$ suivant :

$$\mathcal{X}_{\mathcal{E}}^{A,0} = \{ \mathcal{F} \in \mathcal{X}_{\mathcal{E}} : \mathcal{F} \times \boldsymbol{\nu}|_{\Gamma_C} = 0, F_z|_{\partial\Omega} = 0 \}. \quad (8.20)$$

$\mathcal{X}_{\mathcal{E}}^{A,0}$ est muni du produit scalaire suivant : $\forall \mathcal{F}, \mathcal{F}' \in \mathcal{X}_{\mathcal{E}}^{A,0}$,

$$(\mathcal{F}, \mathcal{F}')_{\mathcal{X}_{\mathcal{E}}^{A,0}} = (\mathcal{F}, \mathcal{F}')_{\mathcal{X}_{\mathcal{E}}^0} + \int_{\Gamma_A} (\mathcal{F} \times \boldsymbol{\nu}) \cdot (\mathcal{F}' \times \boldsymbol{\nu}) d\Sigma,$$

dont la norme associée est : $\|\mathcal{F}\|_{\mathcal{X}_{\mathcal{E}}^{A,0}} = \left(\|\mathcal{F}\|_{\mathcal{X}_{\mathcal{E}}^0}^2 + \int_{\Gamma_A} |\mathcal{F} \times \boldsymbol{\nu}|^2 d\Sigma \right)^{1/2}$.

• On définit aussi $\mathbf{X}_{\mathbf{E}}^A$, le sous-espace de $\mathbf{X}_{\mathbf{E}}$ sur ω suivant :

$$\mathbf{X}_{\mathbf{E}}^A = \{ \mathbf{u} \in \mathbf{X}_{\mathbf{E}} : u_{\tau}|_{\gamma_C} = 0 \}, \quad (8.21)$$

où $u_{\tau} = \mathbf{u} \cdot \boldsymbol{\tau}|_{\partial\omega}$, avec $\boldsymbol{\tau}$ le vecteur tangentiel à la frontière polygonale $\partial\omega$. $\mathbf{X}_{\mathbf{E}}^A$ est muni du produit scalaire suivant : $\forall \mathbf{u}, \mathbf{v} \in \mathbf{X}_{\mathbf{E}}^A$,

$$(\mathbf{u}, \mathbf{v})_{\mathbf{X}_{\mathbf{E}}^A} = (\mathbf{u}, \mathbf{v})_{\mathbf{X}_{\mathbf{E}}^0} + \int_{\gamma_A} u_{\tau} v_{\tau} d\sigma,$$

dont la norme associée est : $\|\mathbf{u}\|_{\mathbf{X}_{\mathbf{E}}^A} = \left(\|\mathbf{u}\|_{\mathbf{X}_{\mathbf{E}}^0}^2 + \int_{\gamma_A} u_{\tau}^2 d\sigma \right)^{1/2}$.

Rappelons que $\mathbf{X}_{\mathbf{E}}$ est défini au paragraphe 1.3.2 de la partie II). Notons que : $\mathcal{X}_{\mathcal{E}}^0 \subset \mathcal{X}_{\mathcal{E}}^{A,0} \subset \mathcal{X}_{\mathcal{E}}$ et $\mathbf{X}_{\mathbf{E}}^0 \subset \mathbf{X}_{\mathbf{E}}^A \subset \mathbf{X}_{\mathbf{E}}$. On peut décomposer $\mathbf{X}_{\mathbf{E}}^A$ de la façon suivante :

$$\mathbf{X}_{\mathbf{E}}^A = \{ \mathbf{v} := \mathbf{u}_0 + \mathbf{e}, \text{ avec } \mathbf{u}_0 \in \mathbf{X}_{\mathbf{E}}^0, \text{ et } \mathbf{e} \in \mathbf{H}^1(\omega) \text{ tel que } e_{\tau}|_{\gamma_C} = 0 \}.$$

On en déduit donc que $\mathbf{X}_{\mathbf{E}}^A \cap \mathbf{H}^1(\Omega)$ est un sous-espace fermé de $\mathbf{X}_{\mathbf{E}}^A$. Ainsi, $\mathcal{X}_{\mathcal{E}}^{A,0} \cap \mathcal{H}^1(\Omega)$ est un sous-espace fermé de $\mathcal{X}_{\mathcal{E}}^{A,0}$.

Proposition 8.21 Pour tout $\mathcal{F} \in \mathcal{X}_{\mathcal{E}}^{A,0}$, $F_z \in H_0^1(\Omega)$, et $\partial_z F_x, \partial_z F_y \in L^2(\Omega)$.

DÉMONSTRATION. Soit $\mathcal{F} \in \mathcal{X}_{\mathcal{E}}^{A,0}$. D'après la proposition 8.1, il existe un relèvement régulier de $\mathcal{F} \times \boldsymbol{\nu}|_{\Gamma_A}$, $\mathcal{F}^r \in \mathcal{H}^1(\Omega)$, tel que : $\mathcal{F}^0 = \mathcal{F} - \mathcal{F}^r \in \mathcal{X}_{\mathcal{E}}^0$. Appliquons la proposition 8.19 à $\mathcal{F}^0$: on a $F^0 = F_z - F_z^r \in H^1(\Omega)$, et $\partial_z F_x - \partial_z F_x^r, \partial_z F_y - \partial_z F_y^r \in L^2(\Omega)$. Comme $F_z^r \in H_0^1(\Omega)$, et que $\partial_z F_x^r$ et $\partial_z F_y^r \in L^2(\Omega)$, on en déduit 8.21. □

Ainsi, seules les composantes F_x et F_y d'un champ électrique $\mathcal{F}$ de $\mathcal{X}_{\mathcal{E}}^0$ ou $\mathcal{X}_{\mathcal{E}}^{A,0}$ peuvent être singulières. Comme $\mathcal{E} \in \mathcal{L}^2(\Omega)$ et que par ailleurs : $E_x = E_y = 0$ sur Γ_z , il semble logique de décomposer $\mathcal{E}$ en une série de Fourier de la façon suivante :

$$\mathcal{E}(x, y, z) = \sum_{k=0}^{\infty} \left(\mathbf{E}^k(x, y) \sin\left(\frac{k\pi}{L}z\right) + E_z^k(x, y) \cos\left(\frac{k\pi}{L}z\right) \mathbf{e}_z \right), \quad (8.22)$$

Tous les vecteurs de $\mathcal{X}_{\mathcal{E}}^0$ ou $\mathcal{X}_{\mathcal{E}}^{A,0}$ seront décomposés de la même façon. De même, les données du problème sont décomposées en série de Fourier :

$$\begin{aligned} g(x, y, z) &= \sum_{k=0}^{\infty} g^k(x, y) \sin\left(\frac{k\pi}{L}z\right). \\ \mathbf{f}(x, y, z) &= \sum_{k=0}^{\infty} \left(\mathbf{f}^k(x, y) \cos\left(\frac{k\pi}{L}z\right) + f_z^k(x, y) \sin\left(\frac{k\pi}{L}z\right) \mathbf{e}_z \right), \\ \mathbf{e}(x, y, z) &= \sum_{k=0}^{\infty} \left(\mathbf{e}^k(x, y) \sin\left(\frac{k\pi}{L}z\right) + e_z^k(x, y) \cos\left(\frac{k\pi}{L}z\right) \mathbf{e}_z \right). \end{aligned}$$

Comme $\mathbf{f} \cdot \boldsymbol{\nu}_{|\Gamma_z} = 0$, on décompose $\mathbf{f} \cdot \mathbf{z}$ en série de sinus selon la direction $\mathbf{e}_z$, et $\mathbf{f}$ en série de cosinus dans le plan $(\mathbf{e}_x, \mathbf{e}_y)$. Comme $\mathbf{e} \times \boldsymbol{\nu}_{|\partial\Omega}$ est la trace tangentielle d'un vecteur de $\mathcal{X}_{\mathcal{E}}^{A,0} \cap \mathcal{H}^1(\Omega)$, on décompose $\mathbf{e}$ de la même façon que $\mathcal{E}$.

Remarque 8.22 On peut indifféremment choisir de développer g en série de cosinus ou de sinus car on a simplement $g \in L^2(\Omega)$. Le choix d'un développement en série de sinus permet de simplifier certains calculs.

Soit $\mathcal{F} \in \mathcal{X}_{\mathcal{E}}^0$ (resp. $\mathcal{X}_{\mathcal{E}}^{A,0}$), décomposé selon 8.22. On a :

$$\mathbf{rot} \mathcal{F} = \sum_{k=0}^{\infty} \begin{pmatrix} \left(\partial_y F_z^k - \frac{k\pi}{L} F_y^k \right) \cos\left(\frac{k\pi}{L}z\right) \\ \left(\frac{k\pi}{L} F_x^k - \partial_x F_z^k \right) \cos\left(\frac{k\pi}{L}z\right) \\ \mathbf{rot} \mathbf{F}^k \sin\left(\frac{k\pi}{L}z\right) \end{pmatrix},$$

et :

$$\operatorname{div} \mathcal{F} = \sum_{k=0}^{\infty} \left(\operatorname{div} \mathbf{F}^k - \frac{k\pi}{L} F_z^k \right) \sin\left(\frac{k\pi}{L}z\right).$$

On a alors pour tout $\mathcal{F} \in \mathcal{X}_{\mathcal{E}}^0$ (resp. $\mathcal{X}_{\mathcal{E}}^{A,0}$), $\mathbf{F}^k \in \mathbf{X}_{\mathbf{E}}^0$ (resp. $\mathbf{X}_{\mathbf{E}}^A$) et $F_z^k \in H_0^1(\omega)$.

Pour $(\mathbf{F}, F_z) \in \mathbf{X}_{\mathcal{E}}^{0,A} \times H_0^1(\omega)$, on pose :

$$\mathcal{F}^k = \mathbf{F}(x, y) \sin\left(\frac{k\pi}{L}z\right) + F_z(x, y) \cos\left(\frac{k\pi}{L}z\right) \mathbf{e}_z. \quad (8.23)$$

Comme $\int_0^L \cos\left(\frac{k\pi}{L}z\right) \cos\left(\frac{l\pi}{L}z\right) dz = \frac{L}{2} \delta_{kl}$ (et de même pour le produit des sinus), on a :

$$(\mathcal{E}, \mathcal{F}^k)_0 = \frac{L}{2} \left((\mathbf{E}^k, \mathbf{F})_{0,\omega} + (\mathbf{E}_z^k, F_z)_{0,\omega} \right),$$

où $(\cdot, \cdot)_{0,\omega}$ désigne le produit scalaire $\mathbf{L}^2$ ou L^2 dans le domaine bidimensionnel ω . Par un calcul direct on a :

$$\begin{aligned} (\mathbf{rot} \mathcal{E}, \mathbf{rot} \mathcal{F}^k)_0 &= \frac{L}{2} \left((\mathbf{rot} \mathbf{E}^k, \mathbf{rot} \mathbf{F})_{0,\omega} + (\mathbf{grad}_\omega \mathbf{E}_z^k, \mathbf{grad}_\omega \mathbf{F}_z)_{0,\omega} + \frac{k^2 \pi^2}{L^2} (\mathbf{E}^k, \mathbf{F})_{0,\omega} \right. \\ &\quad \left. - \frac{k \pi}{L} \left((\mathbf{grad}_\omega \mathbf{E}_z^k, \mathbf{F})_{0,\omega} + (\mathbf{grad}_\omega \mathbf{F}_z, \mathbf{E}^k)_{0,\omega} \right) \right), \end{aligned}$$

où $\mathbf{grad}_\omega$ est le gradient bidimensionnel dans ω . De plus, on a aussi :

$$\begin{aligned} (\mathbf{div} \mathcal{E}, \mathbf{div} \mathcal{F}^k)_0 &= \frac{L}{2} \left((\mathbf{div} \mathbf{E}^k, \mathbf{div} \mathbf{F})_{0,\omega} + \frac{k^2 \pi^2}{L^2} (\mathbf{E}_z^k, \mathbf{F}_z)_{0,\omega} \right. \\ &\quad \left. - \frac{k \pi}{L} \left((\mathbf{div}_\omega \mathbf{E}^k, \mathbf{F}_z)_{0,\omega} + (\mathbf{div}_\omega \mathbf{F}^k, \mathbf{E}_z^k)_{0,\omega} \right) \right), \end{aligned}$$

où $\mathbf{div}_\omega$ est la divergence bidimensionnelle dans ω . Ainsi, le produit scalaire entre $\mathcal{E}^0 \in \mathcal{X}_\mathcal{E}^0$ et $\mathcal{F}^k \in \mathcal{X}_\mathcal{E}^0$ s'écrit alors (à l'aide de la formule d'intégration par parties (1.14) dans ω) :

$$\begin{aligned} (\mathcal{E}^0, \mathcal{F}^k)_{\mathcal{X}_\mathcal{E}^0} &= \frac{L}{2} \left((\mathbf{E}^{0,k}, \mathbf{F})_{\mathbf{X}_\mathbf{E}^0} + (\mathbf{grad}_\omega \mathbf{E}_z^{0,k}, \mathbf{grad}_\omega \mathbf{F}_z)_{0,\omega} \right. \\ &\quad \left. + \frac{k^2 \pi^2}{L^2} \left((\mathbf{E}^{0,k}, \mathbf{F})_{0,\omega} + (\mathbf{E}_z^{0,k}, \mathbf{F}_z)_{0,\omega} \right) \right). \end{aligned} \quad (8.24)$$

Notons d'autre part que pour $\mathcal{E} \in \mathcal{X}_\mathcal{E}$ et $\mathcal{F}^k \in \mathcal{X}_\mathcal{E}^{A,0}$, on a :

$$(\mathcal{E} \times \boldsymbol{\nu}, \mathcal{F}^k \times \boldsymbol{\nu})_{0,\Gamma_A} = \frac{L}{2} \left(\int_{\gamma_A} \mathbf{E}^k \cdot \boldsymbol{\tau} \mathbf{F} \cdot \boldsymbol{\tau} d\sigma \right).$$

On obtient alors le produit scalaire entre $\mathcal{E} \in \mathcal{X}_\mathcal{E}$ et $\mathcal{F}^k \in \mathcal{X}_\mathcal{E}^{A,0}$, (à l'aide de (1.14)) :

$$(\mathcal{E}, \mathcal{F}^k)_{\mathcal{X}_\mathcal{E}^{A,0}} = (\mathcal{E}^0, \mathcal{F}^k)_{\mathcal{X}_\mathcal{E}^0} + \frac{L}{2} \left(\int_{\gamma_A} \mathbf{E}^k \cdot \boldsymbol{\tau} \mathbf{F} \cdot \boldsymbol{\tau} d\sigma - \frac{k \pi}{L} \int_{\gamma_A} \mathbf{F} \cdot \boldsymbol{\nu} \mathbf{E}_z^k d\sigma \right). \quad (8.25)$$

Comme $\mathbf{E}_{z|\gamma_A}^k = \mathbf{e}_{z|\gamma_A}^k$ est une donnée du problème, on pourra le passer au second membre la formulation variationnelle. Étudions les seconds membres des formulations variationnelles. On a pour $\mathcal{F}^k \in \mathcal{X}_\mathcal{E}^0$ (*resp.* $\mathcal{X}_\mathcal{E}^{A,0}$) :

$$\begin{aligned} (g, \mathbf{div} \mathcal{F}^k)_0 &= \frac{L}{2} \left((g^k, \mathbf{div} \mathbf{F})_{0,\omega} - \frac{k \pi}{L} (g^k, \mathbf{F}_z)_{0,\omega} \right), \\ (\mathbf{f}, \mathbf{rot} \mathcal{F}^k)_0 &= \frac{L}{2} \left((\mathbf{f}^k, \mathbf{rot} \mathbf{F}_z)_{0,\omega} + (\mathbf{f}_z^k, \mathbf{rot} \mathbf{F})_{0,\omega} - \frac{k \pi}{L} \int_\omega \mathbf{f}^k \stackrel{2}{\times} \mathbf{F} d\omega \right), \\ \int_{\Gamma_A} (\mathbf{e} \times \boldsymbol{\nu}) \cdot (\mathcal{F}^k \times \boldsymbol{\nu}) d\Sigma &= \frac{L}{2} \left(\int_{\gamma_A} \mathbf{e}^k \cdot \boldsymbol{\tau} \mathbf{F} \cdot \boldsymbol{\tau} d\sigma \right). \end{aligned}$$

8.6.2 Cas prismatique : CL essentielles

Décomposons $\mathcal{E}^0$ et $\mathcal{E}^r$ selon (8.22). Comme les $\mathbf{E}^{0,k}$ appartiennent à $\mathbf{X}_{\mathbf{E}}^0$, on peut les décomposer selon la méthode du complément singulier orthogonal (voir le paragraphe 2.4.4) :

$$\mathbf{E}^{0,k} = \widehat{\mathbf{E}}^{0,k} + \sum_{i=1}^{N_e} c_i^k \mathbf{x}_i^S, \text{ avec } \widehat{\mathbf{E}}^{0,k} \in \mathbf{X}_{\mathbf{E}}^{0,R}, \text{ et } \forall i, c_i^k \in \mathbb{R}.$$

La formulation variationnelle (8.14) reste vraie pour tout vecteur $\mathcal{F}^k$, $k \in \mathbb{N}$, de la forme (8.23), en particulier lorsque $\mathbf{F} = 0$ ou $\mathbf{F}_z = 0$. On obtient alors une série de problèmes posés dans ω et dépendants de k .

Soient $\mathcal{A}_0^k$ et a_D^k les formes bilinéaires symétriques et coercitives suivantes :

$$\begin{aligned} \mathcal{A}_0^k : \mathbf{X}_{\mathbf{E}}^A \times \mathbf{X}_{\mathbf{E}}^A &\rightarrow \mathbb{R} \\ (\mathbf{E}, \mathbf{F}) &\mapsto (\mathbf{E}, \mathbf{F})_{\mathbf{X}_{\mathbf{E}}^0} + \frac{k^2 \pi^2}{L^2} (\mathbf{E}, \mathbf{F})_{0,\omega}, \end{aligned}$$

$$\begin{aligned} a_D^k : H_0^1(\omega) \times H_0^1(\omega) &\rightarrow \mathbb{R} \\ (u, v) &\mapsto (\mathbf{grad}_{\omega} u, \mathbf{grad}_{\omega} v)_0 + \frac{k^2 \pi^2}{L^2} (u, v)_{0,\omega}. \end{aligned}$$

Soient $\mathcal{L}_0^k$ et l_D^k les formes linéaires suivantes :

$$\begin{aligned} \mathcal{L}_0^k : \mathbf{X}_{\mathbf{E}}^0 &\rightarrow \mathbb{R} \\ \mathbf{F} &\mapsto (g^k, \operatorname{div} \mathbf{F})_{0,\omega} + (\mathbf{f}_z^k, \operatorname{rot} \mathbf{F})_{0,\omega} - \frac{k \pi}{L} \int_{\omega} \mathbf{f}^k \times \mathbf{F} \, d\omega, \end{aligned}$$

$$\begin{aligned} l_D^k : H_0^1(\omega) &\rightarrow \mathbb{R} \\ v &\mapsto (\mathbf{f}^k, \mathbf{rot} v)_{0,\omega} - \frac{k \pi}{L} (g^k, v)_{0,\omega}. \end{aligned}$$

Soient $\lambda^{k,j}$ les coefficients de singularité suivants :

$$\lambda^{k,j} = \left((g^k - \operatorname{div} \mathbf{E}^{r,k}, s_{D,j})_{0,\omega} + (\mathbf{f}_z^k - \operatorname{rot} \mathbf{E}^{r,k}, s_{N,j})_{0,\omega} \right) / \pi - \frac{k}{L} \int_{\omega} \mathbf{f}^k \times \mathbf{x}_j^S \, d\omega,$$

où $s_{D,j}$ et $s_{N,j}$ sont les fonctions singulières primales du Laplacien (partie II, paragraphe 2.2.3).

La proposition 8.12 se réécrit alors :

Proposition 8.23 *Le problème (8.7)-(8.8) est équivalent à résoudre la suite de problèmes variationnels suivants :*

Pour tout $k \in \mathbb{N}$, trouver $(\widehat{\mathbf{E}}^{0,k}, \mathbf{E}_z^{0,k}) \in \mathbf{X}_{\mathbf{E}}^{0,R} \times H_0^1(\omega)$ et $(c_i^k)_{i \in \{1, \dots, N_{ar}\}} \in \mathbb{R}^{N_{ar}}$ tels que :

D'une part :

$$\forall \mathbf{F} \in \mathbf{X}_{\mathbf{E}}^{0,R}, \mathcal{A}_0^k(\widehat{\mathbf{E}}^{0,k}, \mathbf{F}) + \frac{k^2 \pi^2}{L^2} \sum_{i=1}^{N_e} c_i^k (\mathbf{x}_i^S, \mathbf{F})_{0,\omega} = \mathcal{L}_0^k(\mathbf{F}) - \mathcal{A}_0^k(\mathbf{E}^{r,k}, \mathbf{F}), \quad (8.26)$$

$$\forall j \in \{1, \dots, N_{ar}\}, \frac{k^2 \pi^2}{L^2} (\widehat{\mathbf{E}}^{0,k}, \mathbf{x}_j^S)_{0,\omega} + \sum_{i=1}^{N_{ar}} c_i^k \mathcal{A}_0^k(\mathbf{x}_i^S, \mathbf{x}_j^S) = \pi \lambda^{k,j}. \quad (8.27)$$

D'autre part :

$$\forall v \in H_0^1(\omega), a_D^k(\mathbf{E}_z^{0,k}, v) = l_D^k(v) - a_D^k(\mathbf{E}_z^{r,k}, v). \quad (8.28)$$

Nous obtenons donc des systèmes d'équations découplées en $((\widehat{\mathbf{E}}^{0,k}, \mathbf{c}^k), \mathbf{E}_z^{0,k})$, ce qui facilite la résolution. On pourra au choix écrire un système linéaire en dimension trois ou deux systèmes linéaires, l'un en dimension deux (voir le calcul discret du champ électrique bidimensionnel $\mathbf{E}$ dans le paragraphe 10.6.1), et l'autre en dimension un (voir le calcul discret du champ électrique scalaire E_z dans le paragraphe 10.6.3). Pour $k = 0$, le système d'équations $(\widehat{\mathbf{E}}^{0,k}, \mathbf{c}^k)$ est en plus découplé car la décomposition de $\mathbf{X}_{\mathbf{E}}^0$ choisie est orthogonale.

8.6.3 Cas prismatique : CL presque essentielles

L'espace $\mathbf{X}_{\mathbf{E}}^A$ se décompose en une partie régulière et une partie singulière [7] :

Proposition 8.24 *Soit $\mathbf{X}_{\mathbf{E}}^{A,R} = \mathbf{X}_{\mathbf{E}}^A \cap \mathbf{H}^1(\omega)$ l'espace régularisé de $\mathbf{X}_{\mathbf{E}}^A$. On a alors la décomposition suivante :*

$$\mathbf{X}_{\mathbf{E}}^A = \mathbf{X}_{\mathbf{E}}^{A,R} \oplus \mathbf{X}_{\mathbf{E}}^{0,S}. \quad (8.29)$$

DÉMONSTRATION. Nous reprenons la preuve de [7]. Soit ω_A un voisinage de la frontière artificielle γ_A . Soit η_A une fonction de troncature régulière, qui s'annule en dehors de ω_A , et qui vaut 1 dans un voisinage bidimensionnel $\mathcal{V}_A \subset \omega_A$ de γ_A . Par construction, on a :

$$\mathbf{E} = \eta_A \mathbf{E} + (1 - \eta_A) \mathbf{E}, \text{ avec : } \eta_A \mathbf{E} \in \mathbf{H}^1(\omega) \text{ et } (1 - \eta_A) \mathbf{E} \in \mathbf{X}_{\mathbf{E}}^0.$$

A priori, $\eta_A \mathbf{E}|_{\gamma_A} \neq 0$, alors que $\eta_A \mathbf{E}|_{\gamma_C} = 0$. Le champ électrique se décompose donc en une somme de deux termes, l'un appartenant à $\mathbf{X}_{\mathbf{E}}^{A,R}$ et l'autre à $\mathbf{X}_{\mathbf{E}}^0$, d'où :

$$\mathbf{X}_{\mathbf{E}}^A = \mathbf{X}_{\mathbf{E}}^{A,R} \oplus \mathbf{X}_{\mathbf{E}}^0 = \mathbf{X}_{\mathbf{E}}^{A,R} \oplus \mathbf{X}_{\mathbf{E}}^{0,R} \oplus \mathbf{X}_{\mathbf{E}}^{0,R}.$$

Comme $\mathbf{X}_{\mathbf{E}}^{0,R}$ est un sous-espace de $\mathbf{X}_{\mathbf{E}}^{A,R}$, on obtient (8.29). □

Pour tout $k \in \mathbb{N}$, on décompose alors $\mathbf{E}^k \in \mathbf{X}_{\mathbf{E}}^A$ de la façon suivante : $\mathbf{E}^k = \widehat{\mathbf{E}}^k + \sum_{i=1}^{N_{ar}} c_{A,i}^k \mathbf{x}_i^S$, où $\widehat{\mathbf{E}}^k \in \mathbf{X}_{\mathbf{E}}^{A,R}$ et $\forall i, c_{A,i}^k \in \mathbb{R}$. Décomposons $\mathbf{E}_z^k$ en $\mathbf{E}_z^{0,k} + \mathbf{E}_z^{r,k}$.

Soit $\mathcal{A}_A^k$ la forme bilinéaire symétrique coercitive suivante :

$$\begin{aligned} \mathcal{A}_A^k : \mathbf{X}_{\mathbf{E}}^A \times \mathbf{X}_{\mathbf{E}}^A &\rightarrow \mathbb{R} \\ (\mathbf{E}, \mathbf{F}) &\mapsto \mathcal{A}_0^k(\mathbf{E}, \mathbf{F}) + \int_{\gamma_A} \mathbf{E} \cdot \boldsymbol{\tau} \mathbf{F} \cdot \boldsymbol{\tau} \, d\sigma. \end{aligned}$$

Soit $\mathcal{L}_A^k$ la forme linéaire suivante :

$$\begin{aligned} \mathcal{L}_A^k : \mathbf{X}_{\mathbf{E}}^A &\rightarrow \mathbb{R} \\ \mathbf{F} &\mapsto \mathcal{L}_0^k(\mathbf{F}) + \int_{\gamma_A} \mathbf{e}^k \cdot \boldsymbol{\tau} \mathbf{F} \cdot \boldsymbol{\tau} \, d\sigma + \frac{k\pi}{L} \int_{\gamma_A} \mathbf{F} \cdot \boldsymbol{\nu} \mathbf{e}_z^k \, d\sigma. \end{aligned}$$

Soient $\lambda_A^{k,j}$ les coefficients de singularité suivants :

$$\lambda_A^{k,j} = \left((g^k, s_{D,j})_{0,\omega} + (f_z^k, s_{N,j})_{0,\omega} \right) / \pi - \frac{k}{L} \int_{\omega} \mathbf{f}^k \times \mathbf{x}_j^S \, d\omega + \frac{k}{L} \int_{\gamma_A} \mathbf{x}_j^S \cdot \boldsymbol{\nu} \mathbf{e}_z^k \, d\sigma.$$

La proposition 8.10 se réécrit alors :

Proposition 8.25 *Le problème (8.1)-(8.3) est équivalent la suite de problèmes variationnels suivants :*

Pour tout $k \in \mathbb{N}$, trouver $(\hat{\mathbf{E}}^k, \mathbf{E}_z^{0,k}) \in \mathbf{X}_{\mathbf{E}}^{A,R} \times H_0^1(\omega)$ et $(c_{A,i}^k)_{i \in \{1, \dots, N_{ar}\}} \in \mathbb{R}^{N_{ar}}$ tels que :
D'une part,

$$\forall \mathbf{F} \in \mathbf{X}_{\mathbf{E}}^{A,R}, \mathcal{A}_A^k(\hat{\mathbf{E}}^k, \mathbf{F}) + \sum_{i=1}^{N_{ar}} c_{A,i}^k \mathcal{A}_A^k(\mathbf{x}_i^S, \mathbf{F}) = \mathcal{L}_A^k(\mathbf{F}), \quad (8.30)$$

$$\forall j \in \{1, \dots, N_{ar}\}, \mathcal{A}_A^k(\hat{\mathbf{E}}^k, \mathbf{x}_j^S) + \sum_{i=1}^{N_{ar}} c_{A,i}^k \mathcal{A}_A^k(\mathbf{x}_i^S, \mathbf{x}_j^S) = \pi \lambda_A^{k,j}. \quad (8.31)$$

D'autre part, $\mathbf{E}_z^{0,k}$ satisfait (8.28)

De même que pour le calcul dans $\mathbf{X}_{\mathbf{E}}^0 \times H_0^1(\omega)$, on a un découplage du système d'équations en $(\hat{\mathbf{E}}^k, \mathbf{c}^k)$ pour la partie transverse et en $\mathbf{E}_z^k$ pour la partie longitudinale. Pour $k = 0$, le système d'équations $(\hat{\mathbf{E}}^{0,k}, \mathbf{c}^k)$ n'est pas découplé car la décomposition de $\mathbf{X}_{\mathbf{E}}^A$ choisie n'est pas orthogonale.

Chapitre 9

Le problème statique 3D mixte continu

Dans ce chapitre, nous reprenons le formalisme du chapitre 3 de la partie II. Les preuves de la condition inf-sup sont les mêmes que dans le cas bidimensionnel. Dans la section 3.1, p. 85, nous avons fait un rappel sur les formulations mixtes.

Nous rappelons à la section 17.2 de la partie IV les espaces fonctionnels des différentes méthodes.

9.1 Formulation mixte 3D : CL naturelles

Considérons le problème (3.1) avec $\mathbf{X} = \mathcal{X}_{\mathcal{E}}$. On a alors : $\mathcal{A}(\cdot, \cdot) = (\cdot, \cdot)_{\mathcal{X}_{\mathcal{E}}}$ et $\mathcal{L} = \mathcal{L}_{\mathcal{E}}$. Comme $\operatorname{div} \mathcal{E} \in L^2(\Omega)$, on prend $Q = L^2(\Omega)$, et l'on notera $\mathcal{B}_{\mathcal{E}}$ la forme bilinéaire, et $\mathcal{G}_{\mathcal{E}}$ la forme linéaire associées :

$$\begin{aligned} \mathcal{B}_{\mathcal{E}} : \mathcal{X}_{\mathcal{E}} \times L^2(\Omega) &\rightarrow \mathbb{R} \\ (\mathcal{F}, q) &\mapsto (\operatorname{div} \mathcal{F}, q)_0 ; \end{aligned} \quad (9.1)$$

$$\begin{aligned} \mathcal{G}_{\mathcal{E}} : L^2(\Omega) &\rightarrow \mathbb{R} \\ q &\mapsto (g, q)_0 . \end{aligned} \quad (9.2)$$

Proposition 9.1 *Il existe une constante $\kappa_{\mathcal{E}} \geq 1$ telle que :*

$$\inf_{q \in L^2(\Omega)} \sup_{\mathcal{F} \in \mathcal{X}_{\mathcal{E}}} \frac{\mathcal{B}_{\mathcal{E}}(\mathcal{F}, q)}{\|\mathcal{F}\|_{\mathcal{X}_{\mathcal{E}}} \|q\|_0} \geq \kappa_{\mathcal{E}}.$$

DÉMONSTRATION. Soit $q \in L^2(\Omega)$. Considérons $\phi \in H_0^1(\Omega)$ tel que $-\Delta \phi = q$, et $\mathbf{F} = -\mathbf{grad} \phi$ (voir la section 8.2). Alors $\mathcal{F} \in L^2(\Omega)$ est à rotationnel nul, à divergence dans $L^2(\Omega)$, et vérifie $\mathcal{F} \times \boldsymbol{\nu}_{|\partial\Omega} = 0$. On en déduit que $\mathcal{F} \in \mathcal{X}_{\mathcal{E}}$ et que : $\mathcal{B}_{\mathcal{E}}(\mathcal{F}, q) = \|\operatorname{div} \mathcal{F}\|_0^2 = \|\mathcal{F}\|_{\mathcal{X}_{\mathcal{E}}}^2$, et aussi : $\mathcal{B}_{\mathcal{E}}(\mathcal{F}, q) = \|q\|_0^2$, d'où : $\mathcal{B}_{\mathcal{E}}(\mathcal{F}, q) = \|q\|_0 \|\mathcal{F}\|_{\mathcal{X}_{\mathcal{E}}}$. □

Théorème 9.2 *Le problème (8.1)-(8.3) est équivalent à la formulation suivante : Trouver $(\mathcal{E}, p) \in \mathcal{X}_{\mathcal{E}} \times L^2(\Omega)$ tels que :*

$$\begin{aligned} (\mathcal{E}, \mathcal{F})_{\mathcal{X}_{\mathcal{E}}} + \mathcal{B}_{\mathcal{E}}(\mathcal{F}, p) &= \mathcal{L}_{\mathcal{E}}(\mathcal{F}) \quad \forall \mathcal{F} \in \mathcal{X}_{\mathcal{E}}, \\ \mathcal{B}_{\mathcal{E}}(\mathcal{E}, q) &= \mathcal{G}_{\mathcal{E}}(q) \quad \forall q \in L^2(\Omega). \end{aligned} \quad (9.3)$$

DÉMONSTRATION. La condition inf-sup est satisfaite, le problème est donc bien posé et admet une unique solution. Montrons que (9.3) est équivalent à (8.12), c'est-à-dire que $p = 0$. Considérons $\mathcal{F} = -\mathbf{grad} \phi$, où : $\phi \in H_0^1(\Omega)$ est solution de $-\Delta \phi = p$. La première équation de (9.3) devient alors : $(\operatorname{div} \mathcal{E}, p)_0 + (p, p)_0 = (g, p)_0$. D'après la seconde équation de (9.3), $(\operatorname{div} \mathbf{E}, p)_0 = (g, p)_0$, puisque $p \in L^2(\Omega)$. D'où : $\|p\|_0^2 = 0$, et $p = 0$ au sens $L^2(\Omega)$. $\square$

Remarque 9.3 Dans le cas dépendant du temps ou dans le cas harmonique, les conditions aux limites tangentielles naturelles ne sont plus satisfaites quand on traite seulement la divergence comme contrainte : il faut donc un multiplicateur de Lagrange sur la divergence et sur la composante tangentielle au bord du champ électrique. La condition inf-sup sur ce multiplicateur de Lagrange supplémentaire est vérifiée pour l'induction magnétique, mais pas pour le champ électrique [36].

9.2 Formulation mixte 3D : CL essentielles

Considérons le problème (3.1) avec $\mathbf{X} = \mathcal{X}_{\mathcal{E}}^0$. On a alors : $\mathcal{A}(\cdot, \cdot) = (\cdot, \cdot)_{\mathcal{X}_{\mathcal{E}}^0}$ et $\mathcal{L} = \mathcal{L}_0$. Comme $\operatorname{div} \mathcal{E}^0 \in L^2(\Omega)$, on prend $Q = L^2(\Omega)$. On en déduit le théorème suivant :

Théorème 9.4 Le problème (8.7)-(8.8) est équivalent à la formulation variationnelle mixte suivante :

Trouver $(\mathcal{E}^0, p) \in \mathcal{X}_{\mathcal{E}}^0 \times L^2(\Omega)$ tels que :

$$\begin{aligned} (\mathcal{E}^0, \mathcal{F})_{\mathcal{X}_{\mathcal{E}}^0} + \mathcal{B}_{\mathcal{E}}(\mathcal{F}, p) &= \mathcal{L}_0(\mathcal{F}) - \mathcal{A}_0(\mathcal{E}^r, \mathcal{F}) \quad \forall \mathcal{F} \in \mathcal{X}_{\mathcal{E}}^0, \\ \mathcal{B}_{\mathcal{E}}(\mathcal{E}^0, q) &= \mathcal{G}_{\mathcal{E}}(q) - (\operatorname{div} \mathcal{E}^r, q)_0 \quad \forall q \in L^2(\Omega). \end{aligned} \quad (9.4)$$

DÉMONSTRATION. $\mathcal{X}_{\mathcal{E}}^0$ est un sous espace de $\mathcal{X}_{\mathcal{E}}$, donc la condition inf-sup est vérifiée automatiquement, puisque le champ $\mathcal{F}$ construit appartient toujours à $\mathcal{X}_{\mathcal{E}}^0$. L'équivalence avec le problème (8.7)-(8.8) se montre de la même façon que pour le théorème 9.2. $\square$

9.3 Formulation mixte 3D : régularisation à poids

Considérons le problème (3.1) avec $\mathbf{X} = \mathcal{X}_{\mathcal{E}, \gamma}^0$. On a alors : $\mathcal{A}(\cdot, \cdot) = (\cdot, \cdot)_{\mathcal{X}_{\mathcal{E}, \gamma}^0}$ et $\mathcal{L} = \mathcal{L}_{\gamma}$. Comme $\operatorname{div} \mathcal{E} \in L_{\gamma}^2(\Omega)$, on prend $Q = L_{\gamma}^2(\Omega)$. On note $\mathcal{B}_{\gamma}$ la forme bilinéaire, et $\mathcal{G}_{\gamma}$ la forme linéaire associées :

$$\begin{aligned} \mathcal{B}_{\gamma} : \mathcal{X}_{\mathcal{E}, \gamma}^0 \times L_{\gamma}^2(\Omega) &\rightarrow \mathbb{R} \\ (\mathcal{F}, q) &\mapsto (\operatorname{div} \mathcal{F}, q)_{0, \gamma}; \end{aligned} \quad (9.5)$$

$$\begin{aligned} \mathcal{G}_{\gamma} : L_{\gamma}^2(\Omega) &\rightarrow \mathbb{R} \\ q &\mapsto (g, q)_{0, \gamma}. \end{aligned} \quad (9.6)$$

Proposition 9.5 Il existe une constante $\kappa_{\gamma} \geq 1$ telle que :

$$\inf_{q \in L_{\gamma}^2(\Omega)} \sup_{\mathcal{F} \in \mathcal{X}_{\mathcal{E}, \gamma}^0} \frac{\mathcal{B}_{\gamma}(\mathcal{F}, q)}{\|\mathbf{F}\|_{\mathcal{X}_{\mathcal{E}, \gamma}^0} \|q\|_0} \geq \kappa_{\gamma}.$$

DÉMONSTRATION. Soit $q \in L^2_\gamma(\Omega)$. Considérons $\phi \in H^1_0(\Omega)$ tel que $-\Delta\phi = q$. Comme $L^2_\gamma(\Omega) \subset H^{-1}(\Omega)$, il existe une unique solution ϕ . Posons $\mathcal{F} = -\mathbf{grad} \phi$. Alors $\mathcal{F} \in \mathcal{L}^2(\Omega)$ est à rotationnel nul, à divergence dans $L^2_\gamma(\Omega)$, et vérifie $\mathcal{F} \times \boldsymbol{\nu}_{|\partial\Omega} = 0$. On en déduit que $\mathcal{F} \in \mathcal{X}^0_{\mathcal{E},\gamma}$ et que : $\mathcal{B}_\gamma(\mathcal{F}, q) = \|\operatorname{div} \mathcal{F}\|_{0,\gamma} = \|\mathcal{F}\|_{\mathcal{X}^0_{\mathcal{E},\gamma}}$.

□

Théorème 9.6 *Le problème (8.7)-(8.8) est équivalent à la formulation variationnelle mixte suivante :*

Trouver $(\mathcal{E}^0, p) \in \mathcal{X}^0_{\mathcal{E},\gamma} \times L^2_\gamma(\Omega)$ tels que :

$$\begin{aligned} (\mathcal{E}^0, \mathcal{F})_{\mathcal{X}^0_{\mathcal{E},\gamma}} + \mathcal{B}_\gamma(\mathcal{F}, p) &= \mathcal{L}_\gamma(\mathcal{F}) - \mathcal{A}_\gamma(\mathcal{E}^r, \mathcal{F}) \quad \forall \mathbf{F} \in \mathcal{X}^0_{\mathcal{E},\gamma}, \\ \mathcal{B}_\gamma(\mathcal{E}^0, q) &= \mathcal{G}_\gamma(q) - (\operatorname{div} \mathcal{E}^r, q)_{0,\gamma} \quad \forall q \in L^2(\Omega). \end{aligned} \tag{9.7}$$

DÉMONSTRATION. La condition inf-sup est satisfaite. Le problème (9.7) est donc bien posé et admet un unique couple $(\mathcal{E}, p) \in \mathcal{X}_{\mathcal{E},\gamma} \times L^2_\gamma(\Omega)$ solution. Montrons que $p = 0$, et donc que (9.7) est équivalente à (8.17). Prenons $\mathcal{F} = -\mathbf{grad} \phi$, où : $\phi \in H^1_0(\Omega)$ est la solution de $-\Delta\phi = p$. La première équation de (9.7) devient alors :

$$(\operatorname{div} \mathcal{E}, p)_{0,\gamma} + (p, p)_{0,\gamma} = (g, p)_{0,\gamma}.$$

D'après la seconde équation de (9.7), $(\operatorname{div} \mathcal{E}, p)_{0,\gamma} = (g, p)_{0,\gamma}$, puisque $p \in L^2_\gamma(\Omega)$. D'où : $\|p\|_{0,\gamma}^2 = 0$, et on a bien $p = 0$ au sens $L^2_\gamma(\Omega)$.

□

9.4 Formulation mixte $2D\frac{1}{2}$: le complément singulier

Lorsqu'on introduit un multiplicateur de Lagrange sur la divergence dans le système d'équations (8.26)-(8.28), des termes de couplage entre le système d'équation pour le champ transverse et l'équation pour le champ longitudinal apparaissent. En effet, comme $p \in L^2(\Omega)$, on peut le décomposer de la façon suivante :

$$p = \sum_{k=0}^{\infty} p^k \sin\left(\frac{k\pi}{L}z\right).$$

Le terme $\mathcal{B}_\mathcal{E}(\mathcal{F}, p)$ de (9.4) devient alors :

- Pour $\mathcal{F} = \mathbf{F}^k$: $\mathcal{B}_\mathcal{E}(\mathbf{F}^k, p) = \frac{L}{2} (p^k, \operatorname{div} \mathbf{F})_{0,\omega}$,
- Pour $\mathcal{F} = \mathbf{F}^k_z \mathbf{e}_3$: $\mathcal{B}_\mathcal{E}(\mathbf{F}^k_z, p) = \frac{L}{2} \left(-\frac{k\pi}{L} (p^k, \mathbf{F}_z)_{0,\omega} \right)$.

Posons $q^k = q \sin\left(\frac{k\pi}{L}z\right)$, $q \in L^2(\omega)$. Le terme $\mathcal{B}_\mathcal{E}(\mathcal{E}^0, q^k)$ devient :

$$\mathcal{B}_\mathcal{E}(\mathcal{E}, q^k) = \frac{L}{2} \left((q, \operatorname{div} \mathbf{E}^k)_{0,\omega} - \frac{k\pi}{L} (q, \mathbf{E}^k_z)_{0,\omega} \right).$$

Ainsi, de la formulation variationnelle mixte (9.4), on déduit le théorème suivant :

Théorème 9.7 *Le problème (8.7)-(8.8) est équivalent à résoudre la suite de problèmes variationnels suivants : Pour tout $k \in \mathbb{N}$, trouver $(\widehat{\mathbf{E}}^{0,k}, \mathbf{E}_z^{0,k}, p^k) \in \mathbf{X}_{\mathbf{E}}^{0,R} \times H_0^1(\omega) \times L^2(\omega)$ et $(c_i^k)_{i \in \{1, \dots, N_{ar}\}} \in \mathbb{R}^{N_{ar}}$ tels que :*

D'une part : $\forall \mathbf{F} \in \mathbf{X}_{\mathbf{E}}^{0,R}$,

$$\mathcal{A}_0^k(\widehat{\mathbf{E}}^{0,k}, \mathbf{F}) + \frac{k^2 \pi^2}{L^2} \sum_{i=1}^{N_{ar}} c_i^k (\mathbf{x}_i^S, \mathbf{F})_{0,\omega} + (p^k, \operatorname{div} \mathbf{F})_{0,\omega} = \mathcal{L}_0^k(\mathbf{F}) - \mathcal{A}_0^k(\mathbf{E}^{r,k}, \mathbf{F}), \quad (9.8)$$

ainsi que : $\forall j \in \{1, \dots, N_{ar}\}$,

$$\frac{k^2 \pi^2}{L^2} (\widehat{\mathbf{E}}^{0,k}, \mathbf{x}_j^S)_{0,\omega} + \sum_{i=1}^{N_{ar}} c_i^k \mathcal{A}_0^k(\mathbf{x}_i^S, \mathbf{x}_j^S) + (p^k, s_{D,j})_{0,\omega} = \pi \lambda^{k,j}, \quad (9.9)$$

et enfin : $\forall q \in L^2(\omega)$,

$$\begin{aligned} (\operatorname{div} \widehat{\mathbf{E}}^{0,k}, q)_{0,\omega} + \sum_{i=1}^{N_{ar}} c_i^k (s_{D,i}, q)_{0,\omega} - \frac{k \pi}{L} (\mathbf{E}_z^{0,k}, q)_{0,\omega} &= (g^k, q)_{0,\omega} - (\operatorname{div} \mathbf{E}^{r,k}, q)_{0,\omega} \\ &\quad + \frac{k \pi}{L} (\mathbf{E}_z^{r,k}, q)_{0,\omega}. \end{aligned} \quad (9.10)$$

D'autre part : $\forall v \in H_0^1(\omega)$,

$$a_D^k(\mathbf{E}_z^{0,k}, v) - \frac{k \pi}{L} (p^k, v)_{0,\omega} = l_D^k(v) - a_D^k(\mathbf{E}_z^{r,k}, v). \quad (9.11)$$

En procédant de la même façon pour $\mathcal{X}_{\mathcal{E}}^A$, on obtient le théorème suivant :

Théorème 9.8 *Le problème (8.7)-(8.8) est équivalent à résoudre la suite de problèmes variationnels suivants : Pour tout $k \in \mathbb{N}$, trouver $(\widehat{\mathbf{E}}^k, \mathbf{E}_z^{0,k}, p^k) \in \mathbf{X}_{\mathbf{E}}^{A,R} \times H_0^1(\omega) \times L^2(\omega)$ et $(c_{A,i}^k)_{i \in \{1, \dots, N_{ar}\}} \in \mathbb{R}^{N_{ar}}$ tels que :*

D'une part : $\forall \mathbf{F} \in \mathbf{X}_{\mathbf{E}}^{A,R}$,

$$\mathcal{A}_A^k(\widehat{\mathbf{E}}^k, \mathbf{F}) + \sum_{i=1}^{N_{ar}} c_{A,i}^k \mathcal{A}_A^k(\mathbf{x}_i^S, \mathbf{F}) + (p^k, \operatorname{div} \mathbf{F})_{0,\omega} = \mathcal{L}_A^k(\mathbf{F}), \quad (9.12)$$

ainsi que $\forall j \in \{1, \dots, N_{ar}\}$,

$$\mathcal{A}_A^k(\mathbf{x}_j^S, \mathbf{F}) + \sum_{i=1}^{N_{ar}} c_{A,i}^k \mathcal{A}_A^k(\mathbf{x}_i^S, \mathbf{x}_j^S) + (p^k, s_{D,j})_{0,\omega} = \pi \lambda_A^{k,j}, \quad (9.13)$$

et enfin : $\forall q \in L^2(\omega)$,

$$\begin{aligned} (\operatorname{div} \widehat{\mathbf{E}}^k, q)_{0,\omega} + \sum_{i=1}^{N_{ar}} c_{A,i}^k (s_{D,i}, q)_{0,\omega} - \frac{k \pi}{L} (\mathbf{E}_z^{0,k}, q)_{0,\omega} &= (g^k, q)_{0,\omega} \\ &\quad + \frac{k \pi}{L} (\mathbf{E}_z^{r,k}, q)_{0,\omega}. \end{aligned} \quad (9.14)$$

D'autre part $(\mathbf{E}_z^{0,k}, p^k)$ satisfait (9.11).

Nous avons donc obtenu des systèmes d'équations en $((\widehat{\mathbf{E}}^{0,k}, \mathbf{c}^k), \mathbf{E}_z^{0,k}, p^k)$ ou $((\widehat{\mathbf{E}}^k, \mathbf{c}_A^k), \mathbf{E}_z^{0,k}, p^k)$ couplées.

Chapitre 10

Le problème statique 3D direct discret

10.1 Discrétisation du domaine d'étude

On construit un maillage du domaine Ω constitué de L tétraèdres $\{\mathcal{T}_l, l = 1, \dots, L\}$, d'intérieurs non vides, tels que $\cup_l \mathcal{T}_l = \overline{\Omega}$, et définissant un maillage conforme :

$$\forall l, l' \in \{1, \dots, L\} \mathcal{T}_l \cap \mathcal{T}_{l'} = \begin{cases} \text{soit } \emptyset, \\ \text{soit un sommet commun,} \\ \text{soit une arête commune.} \\ \text{soit une face commune.} \end{cases}$$

On note $h = \max_l h_l$, où h_l est le rayon de la sphère circonscrite au tétraèdre $\mathcal{T}_l$.

N_Ω désignera le nombre de points de discrétisation intérieurs à Ω , $N_{\partial\Omega}$, le nombre de points de discrétisation sur le bord $\partial\Omega$, et $N = N_\Omega + N_{\partial\Omega}$ le nombre total de points.

Ainsi, on ordonne les points de discrétisation $(M_i)_{i=1, N}$ de la façon suivante :

- $\forall i \in I_\Omega, M_i \in \Omega$, et
- $\forall i \in I_{\partial\Omega}, M_i \in \partial\Omega$,

où on a défini les ensembles d'indices suivants :

$$I_\Omega = \{1, \dots, N_\Omega\}, I_{\partial\Omega} = \{N_\Omega + 1, \dots, N_\Omega + N_{\partial\Omega}\} \text{ et } I = I_\Omega \cup I_{\partial\Omega}.$$

On considère :

$$V_k = \{v_h \in C^0(\overline{\Omega}) \mid \forall l \in \{1, \dots, L\} u_h|_{\mathcal{T}_l} \in P_k\}, \quad (10.1)$$

où P_k est l'ensemble des fonctions continues, polynômiales de degré k .

- Si $k = 1$, $u_h \in V_1$ est une fonction continue, affine par tétraèdre. Les points de discrétisation de maillage sont les sommets des tétraèdres. Il y a donc quatre degrés de liberté par tétraèdre. u_h est complètement définie par ses valeurs aux sommets des tétraèdres $\mathcal{T}_l$.

- Si $k = 2$, $u_h \in V_2$ est une fonction continue, quadratique par tétraèdre. Les points de discrétisation du maillage sont les sommets des tétraèdres et les milieux des arêtes des tétraèdres. Il y a donc dix degrés de liberté par tétraèdre. u_h est complètement définie par ses valeurs aux sommets et aux milieux des arêtes de la tétraédration.

V_k est un sous-espace de dimension finie de $H^1(\Omega)$. Il est engendré par les fonctions continues, polynômiales par tétraèdre $(v_i)_{i=1, N}$ telles que $v_i(M_j) = \delta_{ij}$. Pour tout i , v_i a pour support les

tétraèdres $\mathcal{T}_i$ contenant M_i . On introduit l'opérateur d'interpolation Π_k tel que :

$$\begin{aligned} \Pi_k : H^1(\Omega) \cap C^0(\overline{\Omega}) &\rightarrow V_k \\ u &\mapsto \sum_{i=1}^N u(M_i) v_i. \end{aligned}$$

Enfin, les F_q , $q \in \{1, \dots, N_F\}$ désigneront les faces du bord (qui sont des triangles).

Soit π_k l'opérateur d'interpolation au bord :

$$\begin{aligned} \pi_k : H^{1/2}(\partial\Omega) \cap C^0(\partial\Omega) &\rightarrow V_k \\ u &\mapsto \sum_{i \in I_{\partial\Omega}} u(M_i) v_i. \end{aligned}$$

Remarque 10.1 *Pour présenter les calculs, on choisit de ne pas passer par les éléments finis de référence car on ne fait que du P_1 ou du P_2 . Ces éléments sont présentés dans la section 15.3 de la partie IV.*

10.2 Électrostatique : discrétisation

10.2.1 Le Laplacien avec CL de Dirichlet

Nous allons discrétiser (8.11), qui donne une solution faible de (8.10). Il s'agit d'un problème de Dirichlet homogène, nous allons donc introduire V_k^0 le sous-espace de $H_0^1(\Omega)$, de dimension N_Ω :

$$V_k^0 := V_k \cap H_0^1(\Omega) = \{u_h \in V_k : u_h|_{\partial\Omega} = 0\}.$$

Les $(v_i)_{i \in I_\Omega}$ forment une base de V_k^0 . Soit Π_k^0 l'opérateur d'interpolation associé à cet espace :

$$\begin{aligned} \Pi_k^0 : H_0^1(\Omega) \cap C^0(\overline{\Omega}) &\rightarrow V_k^0 \\ u &\mapsto \sum_{i \in I_\Omega} u(M_i) v_i. \end{aligned} \tag{10.2}$$

Soit $\phi_{0,h} \in V_k^0$ l'approximation de ϕ_0 , telle que : $\phi_{0,h} = \sum_{j \in I_\Omega} \phi_{0,h}(M_j) v_j$.

La formulation variationnelle (8.11) discrétisée dans V_k^0 s'écrit :
Trouver $\phi_{0,h} \in V_k^0$ tel que :

$$\forall i \in I_\omega, (\mathbf{grad} \phi_{0,h}, \mathbf{grad} v_i)_0 = \int_\Omega g v_i d\Omega, \tag{10.3}$$

où : $\phi_{0,h} = \sum_{j \in I_\omega} \phi_{0,h}(M_j) v_j$.

Ce problème linéaire est de dimension $\mathbb{R}^{N_\Omega}$. Il est résolu *en pratique* dans $\mathbb{R}^N$: on procède par pseudo-élimination des degrés de liberté liés au bord $\partial\Omega$, en modifiant le système linéaire (10.3) de la façon suivante :

Trouver $\phi_{0,h} \in V_k$ tel que :

$$\begin{aligned} \forall i \in I_\Omega, \quad & (\mathbf{grad} \phi_{0,h}, \mathbf{grad} v_i)_0 = \int_\Omega g v_i d\Omega, \\ \forall i \in I_{\partial\Omega}, \quad & \phi_{0,h}(M_i) = 0. \end{aligned}$$

Nous allons exprimer ce problème linéaire sous forme matricielle.

Le second membre peut être calculé de 2 façons :

- Si la fonction d'interpolation $g_h := \Pi_k(g)$ existe, on calcule le second membre par un produit matrice-vecteur.

- Sinon, on construit directement le vecteur $\underline{G} \in \mathbb{R}^N$, de composantes :

$$\begin{aligned} \underline{G}_i &= \int_{\Omega} g v_i \, d\Omega \text{ si } i \in I_{\Omega}, \\ &= 0 \text{ si } i \in I_{\partial\Omega}; \end{aligned}$$

Supposons que g_h existe. On écrit alors l'équation (10.3) sous la forme :

Trouver $\phi_{0,h} \in V_k^0$ tel que :

$$\forall i \in I_{\Omega}, \sum_{j \in I_{\Omega}} \phi_{0,h}(M_j) (\mathbf{grad} v_j, \mathbf{grad} v_i)_0 = \sum_{j \in I_{\Omega}} g(M_j) (v_j, v_i)_0, \quad (10.4)$$

avec : $(\mathbf{grad} v_j, \mathbf{grad} v_i)_0 = (\partial_1 v_j, \partial_1 v_i)_0 + (\partial_2 v_j, \partial_2 v_i)_0 + (\partial_3 v_j, \partial_3 v_i)_0$.

Considérons $\mathbb{K}_0 \in \mathbb{R}^{N \times N}$, appelée *matrice de raideur interne*, la matrice d'éléments :

$$\begin{aligned} \mathbb{K}_0^{i,j} &= (\mathbf{grad} v_j, \mathbf{grad} v_i)_0 \text{ si } i \text{ et } j \in I_{\Omega}, \\ &= \delta_{ij} \text{ si } i \text{ ou } j \in I_{\partial\Omega}. \end{aligned} \quad (10.5)$$

$\forall i, j, \mathbb{K}_0^{i,j} = \mathbb{K}_0^{j,i}$: $\mathbb{K}_0$ est symétrique. Par ailleurs, $\mathbb{K}_0$ est formée de 2 sous-blocs :

$$\mathbb{K}_0 = \begin{pmatrix} \mathbb{K}_{\Omega} & 0 \\ 0 & \mathbb{I}_{\partial\Omega} \end{pmatrix},$$

où $\mathbb{K}_{\Omega} \in \mathbb{R}^{N_{\Omega} \times N_{\Omega}}$ est telle que : $\forall i, j \in I_{\Omega}, \mathbb{K}_{\Omega}^{i,j} = \mathbb{K}_0^{i,j}$; et $\mathbb{I}_{\partial\Omega} \in \mathbb{R}^{N_{\partial\Omega} \times N_{\partial\Omega}}$ est la matrice identité. On remarque que pour $i, j \in I_{\Omega}$, si M_i et M_j n'appartiennent pas à un même tétraèdre, l'élément $\mathbb{K}_0^{i,j}$ nul : la matrice est donc creuse. Lorsque M_i et M_j appartiennent à un même tétraèdre, on a :

$$\mathbb{K}_0^{i,j} = \sum_{T_i | M_i, M_j \in T_i} \int_{T_i} (\partial_1 v_j \partial_1 v_i + \partial_2 v_j \partial_2 v_i + \partial_3 v_j \partial_3 v_i) \, d\Omega.$$

Considérons la matrice $\mathbb{M}_0 \in \mathbb{R}^{N \times N}$, appelée *matrice de masse interne*, la matrice éléments :

$$\begin{aligned} \mathbb{M}_0^{i,j} &= (v_j, v_i)_0 \text{ si } i \in I_{\Omega}, \forall j \in I, \\ &= 0 \text{ si } i \in I_{\partial\Omega}, \forall j \in I. \end{aligned} \quad (10.6)$$

$\mathbb{M}_0$ est formée de deux sous-blocs :

$$\mathbb{M}_0 = \begin{pmatrix} \mathbb{M}_{\Omega} & \mathbb{M}_{\Omega, \partial\Omega} \\ 0 & 0 \end{pmatrix},$$

où $\mathbb{M}_{\Omega} \in \mathbb{R}^{N_{\Omega} \times N_{\Omega}}$ et $\mathbb{M}_{\Omega, \partial\Omega} \in \mathbb{R}^{N_{\Omega} \times N_{\partial\Omega}}$. On remarque que le sous-bloc $\mathbb{M}_{\Omega}$ est symétrique. De même que pour $\mathbb{K}_0$, lorsque M_i et M_j n'appartiennent pas à un même triangle, l'élément $\mathbb{M}_0^{i,j}$ est nul : $\mathbb{M}_0$ est une matrice creuse. Comme précédemment, $\forall i \in I_{\Omega}, \forall j \in I$,

$$\mathbb{M}_0^{i,j} = \sum_{T_i | M_i, M_j \in T_i} \int_{T_i} v_j v_i \, d\Omega.$$

Supposons que $\Pi_k g$ existe. Soit $\underline{\mathbf{g}} \in \mathbb{R}^N$ tel que : $\underline{\mathbf{g}} = (g(M_1), \dots, g(M_N))^T$.

Posons $\underline{\phi}_0 \in \mathbb{R}^N$ tel que $\underline{\phi}_0 = (\phi_{0,h}(M_1), \dots, \phi_{0,h}(M_{N_\Omega}), 0, \dots, 0)^T$.

On peut alors écrire le problème (10.4) sous la forme matricielle suivante :

$$\mathbb{K}_0 \underline{\phi}_0 = \mathbb{M}_0 \underline{\mathbf{g}}.$$

10.2.2 Dérivation du champ électrique

Une fois que l'on a évalué $\phi_{0,h}$, il faut calculer $\mathcal{E}_h$, l'approximation de $\mathcal{E}$, composante par composante. Pour cela, il faut calculer $\mathbf{grad} \phi_{D,h}$. $\phi_{0,h}$ est P_k -continu, donc par dérivation, $\mathcal{E}_h \in (P_{k-1})^3$ -discontinu. Par exemple, si on a utilisé les éléments finis P_1 pour calculer le potentiel électrostatique, ceux-ci sont continus, affines par triangles. Ainsi, le champ $\mathcal{E}_h$ calculé est discontinu, constant par triangle. On a :

$$\mathcal{E}_h = -\mathbf{grad} \phi_{0,h}, \text{ avec } \phi_{0,h} \text{ approximation de } \phi_0, \text{ solution de 10.4.} \quad (10.7)$$

Pour avoir une représentation continue de $\mathcal{E}_h$, on utilise une projection $(P_{k-1})^3 \rightarrow (P_k)^3$ [6, 63], c'est-à-dire qu'au lieu de calculer directement les dérivées du potentiel, on résout :

$$(\mathcal{E}_h, v_i \mathbf{e}_\alpha)_0 = -(\mathbf{grad} \phi_{0,h}, v_i \mathbf{e}_\alpha)_0, \forall i \in I, \forall \alpha \in \{1, 2, 3\}. \quad (10.8)$$

Soit $\mathbb{G}_0 \in (\mathbb{R}^{3 \times 1})^{N \times N}$ la matrice telle que : $\forall i, j \in I$,

$$\mathbb{G}_0^{i,j} = \begin{pmatrix} (v_i, \partial_1 v_j)_0 \\ (v_i, \partial_2 v_j)_0 \\ (v_i, \partial_3 v_j)_0 \end{pmatrix}.$$

Soit $\underline{\mathbf{E}}$ la représentation de $\mathcal{E}_h$ suivante (voir la section 10.3 suivante) :

$$\underline{\mathbf{E}} = (E_{h,1}(M_1), E_{h,2}(M_1), E_{h,3}(M_1), \dots, E_{h,1}(M_N), E_{h,2}(M_N), E_{h,3}(M_N))^T.$$

On a alors :

$$\underline{\mathbf{E}} = -\mathbb{G}_0 \underline{\phi}_0. \quad (10.9)$$

10.3 Méthode avec CL naturelles : discrétisation

$\mathcal{X}_\mathcal{E} \cap \mathcal{H}^1(\Omega)$ est dense dans $\mathcal{X}_\mathcal{E}$ [40, 51]. On définit alors une discrétisation de $\mathcal{X}_\mathcal{E} \cap \mathcal{H}^1(\Omega)$ afin de résoudre les problèmes posés dans la section 8.3 à l'aide des éléments finis continus de Lagrange P_k . Une bonne définition de cet espace discrétisé est, a priori, l'espace $\mathcal{X}_k$:

$$\mathcal{X}_k = \{ \mathbf{v}_h \in C^0(\overline{\Omega})^3 \mid \mathbf{v}_h|_{\mathcal{T}_l} \in P_k(\mathcal{T}_l)^3, \forall l \in \{1, \dots, L\} \}.$$

Par construction, $\mathcal{X}_k \subset \mathcal{H}^1(\Omega)$. Lorsque $k = 1$, les fonctions $\mathbf{v}_h$ de $\mathcal{X}_k$ sont telles que : $\mathbf{v}_h|_{\mathcal{T}_l} = \mathbf{a}|_{\mathcal{T}_l} + \mathbf{b}|_{\mathcal{T}_l}(\mathbf{x})$, où $\mathbf{a}|_{\mathcal{T}_l} \in \mathbb{R}^3$; et $\mathbf{b}|_{\mathcal{T}_l}(\mathbf{x})$ est une application linéaire de $\mathbb{R}^3$ dans $\mathbb{R}^3$. On peut donc

considérer comme espace d'approximation de $\mathcal{X}_{\mathcal{E}} \cap \mathcal{H}^1(\Omega)$ l'espace : $\mathcal{X}_k = (V_k)^3$, pour lequel deux approximations sont possibles :

- D'une part, on peut considérer $\mathcal{X}_k$ comme un espace de dimension finie $3N$ sur $\mathbb{R}$.

Soit $(\mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3)^T$ une base orthonormale directe de $\mathbb{R}^3$. Les $(v_i \mathbf{e}_\alpha)_{i \in \{1, \dots, N\}, \alpha \in \{1, 2, 3\}}$ forment une base de $\mathcal{X}_k$, ce qui va nous permettre de discrétiser la formulation variationnelle (8.12). On cherche $\mathcal{E}_h$, l'approximation de $\mathcal{E}$ sous la forme :

$$\mathcal{E}_h = \sum_{j \in I} \sum_{\beta=1}^3 E_{h,\beta}(M_j) v_j \mathbf{e}_\beta.$$

- D'autre part, on peut considérer $\mathcal{X}_k$ comme un espace de dimension finie N sur $\mathbb{R}^3$, ce qui consiste à chercher les N valeurs *physiques* de $\mathcal{E}_h$ aux points de la discrétisation, c'est-à-dire :

$$\mathcal{E}_h = \sum_{j \in I} \mathcal{E}_h(M_j) v_j,$$

avec $\mathcal{E}_h(M_j) \in \mathbb{R}^3$. Bien sûr, on a : $\mathcal{E}_h(M_j) = E_{h,1}(M_j) \mathbf{e}_1 + E_{h,2}(M_j) \mathbf{e}_2 + E_{h,3}(M_j) \mathbf{e}_3$ pour tout j , et les deux choix d'approximation coïncident.

En pratique, on utilise la première représentation, qui consiste à considérer des inconnues scalaires. Nous allons la détailler dans ce qui suit.

La formulation variationnelle (8.12) discrétisée dans $\mathcal{X}_k$ s'écrit :

Trouver $\mathcal{E}_h \in \mathcal{X}_k$ tel que :

$$\forall i \in I, \forall \alpha \in \{1, 2, 3\}, (\mathcal{E}_h, \mathbf{v}_{i,\alpha})_{\mathcal{X}_{\mathcal{E}}} = \mathcal{L}_{\mathcal{E}}(\mathbf{v}_{i,\alpha}). \quad (10.10)$$

En décomposant $\mathcal{E}_h$ selon la base canonique de $\mathcal{X}_k$ dans (10.10) on obtient :

$$\forall i \in I, \forall \alpha \in \{1, 2, 3\}, \sum_{j \in I} \sum_{\beta=1}^3 E_{h,\beta}(M_j) (\mathbf{v}_{j,\beta}, \mathbf{v}_{i,\alpha})_{\mathcal{X}_{\mathcal{E}}} = \mathcal{L}_{\mathcal{E}}(\mathbf{v}_{i,\alpha}).$$

Soit $\mathbb{A}_{\mathcal{E}} \in (\mathbb{R}^{3 \times 3})^{N \times N}$ la matrice ainsi construite, définie par les sous-blocs $\mathbb{A}_{\mathcal{E}}^{i,j} \in \mathbb{R}^{3 \times 3}$ tels que :

$$\mathbb{A}_{\mathcal{E}}^{i,j} = \begin{pmatrix} (v_j \mathbf{e}_1, v_i \mathbf{e}_1)_{\mathcal{X}_{\mathcal{E}}} & (v_j \mathbf{e}_2, v_i \mathbf{e}_1)_{\mathcal{X}_{\mathcal{E}}} & (v_j \mathbf{e}_3, v_i \mathbf{e}_1)_{\mathcal{X}_{\mathcal{E}}} \\ (v_j \mathbf{e}_1, v_i \mathbf{e}_2)_{\mathcal{X}_{\mathcal{E}}} & (v_j \mathbf{e}_2, v_i \mathbf{e}_2)_{\mathcal{X}_{\mathcal{E}}} & (v_j \mathbf{e}_3, v_i \mathbf{e}_2)_{\mathcal{X}_{\mathcal{E}}} \\ (v_j \mathbf{e}_1, v_i \mathbf{e}_3)_{\mathcal{X}_{\mathcal{E}}} & (v_j \mathbf{e}_2, v_i \mathbf{e}_3)_{\mathcal{X}_{\mathcal{E}}} & (v_j \mathbf{e}_3, v_i \mathbf{e}_3)_{\mathcal{X}_{\mathcal{E}}} \end{pmatrix}.$$

Ces sous-blocs agissent sur les vecteurs de $\mathbb{R}^3$ suivants : $\mathcal{E}_h(M_j) = \begin{pmatrix} E_{h,1}(M_j) \\ E_{h,2}(M_j) \\ E_{h,3}(M_j) \end{pmatrix}$.

Par souci de consistance, on appelle $\underline{\mathbf{L}} \in (\mathbb{R}^3)^N$ le vecteur composé des valeurs $\mathcal{L}_{\mathcal{E}}(v_i \mathbf{v}_\alpha)$, $i \in I$, $\alpha \in \{1, 2, 3\}$:

$$\underline{\mathbf{L}} = (\mathcal{L}_{\mathcal{E}}(v_1 \mathbf{e}_1), \mathcal{L}_{\mathcal{E}}(v_1 \mathbf{e}_2), \mathcal{L}_{\mathcal{E}}(v_1 \mathbf{e}_3), \dots, \mathcal{L}_{\mathcal{E}}(v_N \mathbf{e}_1), \mathcal{L}_{\mathcal{E}}(v_N \mathbf{e}_2), \mathcal{L}_{\mathcal{E}}(v_N \mathbf{e}_3))^T.$$

Selon la même idée, posons :

$$\underline{\mathbf{E}} = (E_{h,1}(M_1), E_{h,2}(M_1), E_{h,3}(M_1), \dots, E_{h,1}(M_N), E_{h,2}(M_N), E_{h,3}(M_N))^T \in (\mathbb{R}^3)^N.$$

On a aussi $\underline{\mathbf{E}} = \begin{pmatrix} \mathcal{E}_h(M_1) \\ \vdots \\ \mathcal{E}_h(M_N) \end{pmatrix} \in (\mathbb{R}^3)^N$. Le problème (10.10) s'écrit alors : $\mathbb{A}_{\mathcal{E}} \underline{\mathbf{E}} = \underline{\mathbf{L}}$.

Remarque 10.2 Dans les sections 10.4 et 10.5, on utilisera une autre base de décomposition de $\mathcal{X}_k$ car il faudra alors éliminer explicitement les valeurs tangentielles.

Nous allons étudier les composantes de $\mathbb{A}_{\mathcal{E}}$ et $\underline{\mathbb{L}}$.

Décomposons $\mathbb{A}_{\mathcal{E}}$ en 3 matrices : $\mathbb{A}_{\mathcal{E}} = \mathbb{D}\text{iv} + \mathbb{R}\text{ot} + \mathbb{B}$, avec $\forall i, j \in \mathbf{I}$:

$$\mathbb{D}\text{iv}^{i,j} = \begin{pmatrix} (\text{div}(v_j \mathbf{e}_1), \text{div}(v_i \mathbf{e}_1))_0 & (\text{div}(v_j \mathbf{e}_2), \text{div}(v_i \mathbf{e}_1))_0 & (\text{div}(v_j \mathbf{e}_3), \text{div}(v_i \mathbf{e}_1))_0 \\ (\text{div}(v_j \mathbf{e}_1), \text{div}(v_i \mathbf{e}_2))_0 & (\text{div}(v_j \mathbf{e}_2), \text{div}(v_i \mathbf{e}_2))_0 & (\text{div}(v_j \mathbf{e}_3), \text{div}(v_i \mathbf{e}_2))_0 \\ (\text{div}(v_j \mathbf{e}_1), \text{div}(v_i \mathbf{e}_3))_0 & (\text{div}(v_j \mathbf{e}_2), \text{div}(v_i \mathbf{e}_3))_0 & (\text{div}(v_j \mathbf{e}_3), \text{div}(v_i \mathbf{e}_3))_0 \end{pmatrix},$$

$$\mathbb{R}\text{ot}^{i,j} = \begin{pmatrix} (\mathbf{rot}(v_j \mathbf{e}_1), \mathbf{rot}(v_i \mathbf{e}_1))_0 & (\mathbf{rot}(v_j \mathbf{e}_2), \mathbf{rot}(v_i \mathbf{e}_1))_0 & (\mathbf{rot}(v_j \mathbf{e}_3), \mathbf{rot}(v_i \mathbf{e}_1))_0 \\ (\mathbf{rot}(v_j \mathbf{e}_1), \mathbf{rot}(v_i \mathbf{e}_2))_0 & (\mathbf{rot}(v_j \mathbf{e}_2), \mathbf{rot}(v_i \mathbf{e}_2))_0 & (\mathbf{rot}(v_j \mathbf{e}_3), \mathbf{rot}(v_i \mathbf{e}_2))_0 \\ (\mathbf{rot}(v_j \mathbf{e}_1), \mathbf{rot}(v_i \mathbf{e}_3))_0 & (\mathbf{rot}(v_j \mathbf{e}_2), \mathbf{rot}(v_i \mathbf{e}_3))_0 & (\mathbf{rot}(v_j \mathbf{e}_3), \mathbf{rot}(v_i \mathbf{e}_3))_0 \end{pmatrix},$$

$$\mathbb{B}^{i,j} = \begin{pmatrix} (v_j \mathbf{e}_{1,T}, v_i \mathbf{e}_{1,T})_{0,\partial\Omega} & (v_j \mathbf{e}_{2,T}, v_i \mathbf{e}_{1,T})_{0,\partial\Omega} & (v_j \mathbf{e}_{3,T}, v_i \mathbf{e}_{1,T})_{0,\partial\Omega} \\ (v_j \mathbf{e}_{1,T}, v_i \mathbf{e}_{2,T})_{0,\partial\Omega} & (v_j \mathbf{e}_{2,T}, v_i \mathbf{e}_{2,T})_{0,\partial\Omega} & (v_j \mathbf{e}_{3,T}, v_i \mathbf{e}_{2,T})_{0,\partial\Omega} \\ (v_j \mathbf{e}_{1,T}, v_i \mathbf{e}_{3,T})_{0,\partial\Omega} & (v_j \mathbf{e}_{2,T}, v_i \mathbf{e}_{3,T})_{0,\partial\Omega} & (v_j \mathbf{e}_{3,T}, v_i \mathbf{e}_{3,T})_{0,\partial\Omega} \end{pmatrix}.$$

Par calcul, on remarque que :

$$(\mathbf{v}_{i,\alpha} \times \boldsymbol{\nu}) \cdot (\mathbf{v}_{j,\beta} \times \boldsymbol{\nu}) = v_i v_j (\delta_{\alpha\beta} - \nu_\alpha \nu_\beta).$$

La matrice $\mathbb{B}$ s'écrit alors :

$$\mathbb{B}^{i,j} = \begin{pmatrix} \int_{\partial\Omega} (1 - \nu_1^2) v_j v_i \, d\Sigma & - \int_{\partial\Omega} \nu_1 \nu_2 v_j v_i \, d\Sigma & - \int_{\partial\Omega} \nu_1 \nu_3 v_j v_i \, d\Sigma \\ - \int_{\partial\Omega} \nu_1 \nu_2 v_j v_i \, d\Sigma & \int_{\partial\Omega} (1 - \nu_2^2) v_j v_i \, d\Sigma & - \int_{\partial\Omega} \nu_2 \nu_3 v_j v_i \, d\Sigma \\ - \int_{\partial\Omega} \nu_1 \nu_3 v_j v_i \, d\Sigma & - \int_{\partial\Omega} \nu_2 \nu_3 v_j v_i \, d\Sigma & \int_{\partial\Omega} (1 - \nu_3^2) v_j v_i \, d\Sigma \end{pmatrix}, \forall i, j \in \mathbf{I}_{\partial\Omega},$$

$$= 0 \text{ sinon.}$$

Notons que : $\mathbb{B}^{i,j} = \mathbb{B}^{j,i}$. $\mathbb{B}$ est donc symétrique par blocs 3×3 .

10.3.1 Matrice de raideur interne

Les sous-blocs $\mathbb{D}iv^{i,j}$ et $\mathbb{R}ot^{i,j}$ sont non nuls seulement si M_i et M_j sont les sommets d'un même tétraèdre. D'où : $\forall i, j \in I$,

$$\mathbb{D}iv^{i,j} = \sum_{\mathcal{T}_l \mid M_i, M_j \in \mathcal{T}_l} \begin{pmatrix} \int_{\mathcal{T}_l} \partial_1 v_j \partial_1 v_i \, d\Omega & \int_{\mathcal{T}_l} \partial_2 v_j \partial_1 v_i \, d\Omega & \int_{\mathcal{T}_l} \partial_3 v_j \partial_1 v_i \, d\Omega \\ \int_{\mathcal{T}_l} \partial_1 v_j \partial_2 v_i \, d\Omega & \int_{\mathcal{T}_l} \partial_2 v_j \partial_2 v_i \, d\Omega & \int_{\mathcal{T}_l} \partial_3 v_j \partial_2 v_i \, d\Omega \\ \int_{\mathcal{T}_l} \partial_1 v_j \partial_3 v_i \, d\Omega & \int_{\mathcal{T}_l} \partial_2 v_j \partial_3 v_i \, d\Omega & \int_{\mathcal{T}_l} \partial_3 v_j \partial_3 v_i \, d\Omega \end{pmatrix},$$

et

$$\mathbb{R}ot^{i,j} = \sum_{\mathcal{T}_l \mid M_i, M_j \in \mathcal{T}_l} \begin{pmatrix} c_{i,j}^{2,3} & - \int_{\mathcal{T}_l} \partial_1 v_j \partial_2 v_i \, d\Omega & - \int_{\mathcal{T}_l} \partial_1 v_j \partial_3 v_i \, d\Omega \\ - \int_{\mathcal{T}_l} \partial_2 v_j \partial_1 v_i \, d\Omega & c_{i,j}^{1,3} & - \int_{\mathcal{T}_l} \partial_2 v_j \partial_3 v_i \, d\Omega \\ - \int_{\mathcal{T}_l} \partial_3 v_j \partial_1 v_i \, d\Omega & - \int_{\mathcal{T}_l} \partial_3 v_j \partial_2 v_i \, d\Omega & c_{i,j}^{1,2} \end{pmatrix},$$

avec $\forall \alpha, \beta \in \{1, 2, 3\}$, $c_{i,j}^{\alpha,\beta} = \int_{\mathcal{T}_l} (\partial_\alpha v_j \partial_\alpha v_i + \partial_\beta v_j \partial_\beta v_i) \, d\Omega$.

Notons que la somme de ces deux blocs est telle :

$$(\mathbb{R}ot + \mathbb{D}iv)^{i,j} = \sum_{\mathcal{T}_l \mid M_i, M_j \in \mathcal{T}_l} \begin{pmatrix} c_{i,j} & a_{i,j} & b_{i,j} \\ -a_{i,j} & c_{i,j} & d_{i,j} \\ -b_{i,j} & -d_{i,j} & c_{i,j} \end{pmatrix},$$

avec :

$$\begin{aligned} a_{i,j} &= \int_{\Omega} (\partial_2 v_j \partial_1 v_i - \partial_1 v_j \partial_2 v_i) \, d\Omega \\ c_{i,j} &= \int_{\Omega} (\mathbf{grad} v_i, \mathbf{grad} v_j) \, d\Omega \\ b_{i,j} &= \int_{\Omega} (\partial_3 v_j \partial_1 v_i - \partial_1 v_j \partial_3 v_i) \, d\Omega \\ d_{i,j} &= \int_{\Omega} (\partial_3 v_j \partial_2 v_i - \partial_2 v_j \partial_3 v_i) \, d\Omega. \end{aligned}$$

$\mathbb{R}ot + \mathbb{D}iv$ est donc anti-symétrique par blocs 3×3 .

10.3.2 Matrice de masse du bord

Le bloc $\mathbb{B}^{i,j}$ est non nul seulement si M_i et M_j sont les sommets d'une même face du bord $\partial\Omega$. On a donc : $\forall i$ ou $j \notin I_{\partial\Omega}$, $\mathbb{B}^{i,j} = 0$ et :

$\forall i, j \in \mathcal{I}_{\partial\Omega},$

$$\begin{aligned} \mathbb{B}^{i,j} &= \sum_{F_q \mid M_i, M_j \in F_q} \int_{F_q} v_j v_i \, d\Sigma \begin{pmatrix} 1 - \nu_{q,1}^2 & -\nu_{q,1} \nu_{q,2} & -\nu_{q,1} \nu_{q,3} \\ -\nu_{q,1} \nu_{q,2} & 1 - \nu_{q,2}^2 & -\nu_{q,2} \nu_{q,3} \\ -\nu_{q,1} \nu_{q,3} & -\nu_{q,2} \nu_{q,3} & 1 - \nu_{q,3}^2 \end{pmatrix} \\ &= \sum_{F_q \mid M_i, M_j \in F_q} \left(\int_{F_q} v_j v_i \, d\Sigma (\mathbb{I}_3 - \boldsymbol{\nu}_q \cdot \boldsymbol{\nu}_q^T) \right), \end{aligned}$$

où $\mathbb{I}_3 \in \mathbb{R}^{3 \times 3}$ est la matrice identité de $\mathbb{R}^3$ et $\boldsymbol{\nu}_q$ est le vecteur normal à la face F_q .

10.3.3 Second membre

Le second membre peut-être calculé directement, à l'aide de schémas d'intégration numériques, ou avec les fonctions d'interpolation $g_h := \Pi_k(g)$, $\mathbf{f}_h := \Pi_k(f)$ et $\mathbf{e}_h = \pi_k(e)$, à condition qu'on puisse définir leur valeur en chaque point de discrétisation (voir la section 10.2). Nous allons développer cette approche. Ceci permet d'écrire le second membre sous forme d'une somme de produits matrice-vecteur, en utilisant $\underline{\mathbf{g}} \in \mathbb{R}^N$, $\underline{\mathbf{f}} \in (\mathbb{R}^3)^N$ et $\underline{\mathbf{e}} \in (\mathbb{R}^3)^N$:

$$\underline{\mathbf{g}} = (g_h(M_1), \dots, g_h(M_N))^T.$$

$$\underline{\mathbf{f}} = (f_{h,1}(M_1), f_{h,2}(M_1), f_{h,3}(M_1), \dots, f_{h,1}(M_N), f_{h,2}(M_N), f_{h,3}(M_N))^T.$$

$$\underline{\mathbf{e}} = (0, \dots, 0, e_{h,1}(M_{N_{\partial\Omega}}), e_{h,2}(M_{N_{\partial\Omega}}), e_{h,3}(M_{N_{\partial\Omega}}), \dots, e_{h,1}(M_N), e_{h,2}(M_N), e_{h,3}(M_N))^T$$

Soit $\mathbb{L}_g \in (\mathbb{R}^3)^{N \times N}$ la matrice formée des $N \times N$ sous-blocs $(\mathbb{L}_g)_{i,j} \in \mathbb{R}^3$ tels que :

$$\forall i, j \in \mathcal{I}, \mathbb{L}_g^{i,j} = \sum_{\mathcal{T}_l \mid M_i, M_j \in \mathcal{T}_l} \begin{pmatrix} \int_{\mathcal{T}_l} v_j \partial_1 v_i \, d\Omega \\ \int_{\mathcal{T}_l} v_j \partial_2 v_i \, d\Omega \\ \int_{\mathcal{T}_l} v_j \partial_3 v_i \, d\Omega \end{pmatrix}.$$

De même, on considère $\mathbb{L}_f \in (\mathbb{R}^{3 \times 3})^{N \times N}$ la matrice formée des $N \times N$ sous-blocs $\mathbb{L}_f^{i,j} \in \mathbb{R}^{3 \times 3}$ tels que :

$$\forall i, j \in \mathcal{I}, \mathbb{L}_f^{i,j} = \sum_{\mathcal{T}_l \mid M_i, M_j \in \mathcal{T}_l} \begin{pmatrix} 0 & \int_{\mathcal{T}_l} v_j \partial_3 v_i \, d\Omega & -\int_{\mathcal{T}_l} v_j \partial_2 v_i \, d\Omega \\ -\int_{\mathcal{T}_l} v_j \partial_3 v_i \, d\Omega & 0 & \int_{\mathcal{T}_l} v_j \partial_1 v_i \, d\Omega \\ \int_{\mathcal{T}_l} v_j \partial_2 v_i \, d\Omega & -\int_{\mathcal{T}_l} v_j \partial_1 v_i \, d\Omega & 0 \end{pmatrix}.$$

Posons enfin $\mathbb{L}_e \in (\mathbb{R}^{3 \times 3})^{N \times N}$ la matrice définie par les $N \times N$ sous-blocs :

$$\begin{aligned} \mathbb{L}_e^{i,j} &= \sum_{F_q \mid M_i, M_j \in F_q} \int_{F_q} v_j v_i d\Sigma (\mathbb{I}_3 - \boldsymbol{\nu}_q \cdot \boldsymbol{\nu}_q^T), \text{ si } i \text{ et } j \in I_{\partial\Omega}, \\ &= 0 \text{ sinon.} \end{aligned}$$

On peut donc calculer $\underline{\mathbb{L}}$ sous la forme matricielle suivante : $\underline{\mathbb{L}} = \mathbb{L}_g \underline{\mathbf{g}} + \mathbb{L}_f \underline{\mathbf{f}} + \mathbb{L}_e \underline{\mathbf{e}}$. Pour approcher le champ électrique, on résout donc le système $3N \times 3N$ suivant :

$$\begin{aligned} \mathbb{A}_{\mathcal{E}} \underline{\mathbf{E}} &= \mathbb{L}_g \underline{\mathbf{g}} + \mathbb{L}_f \underline{\mathbf{f}} + \mathbb{L}_e \underline{\mathbf{e}}, \\ \text{avec : } \mathbb{A}_{\mathcal{E}} &= \mathbb{D}\text{iv} + \mathbb{R}\text{ot} + \mathbb{B}. \end{aligned}$$

(10.11)

10.4 Méthode avec CL essentielles : discrétisation

$\mathcal{X}_{\mathcal{E}}^0 \cap \mathcal{H}^1(\Omega)$ est dense dans $\mathcal{X}_{\mathcal{E}}^0$, seulement si Ω est convexe. Ainsi, la méthode décrite ici converge dans la cas où Ω présente des singularités géométriques seulement si les données sont régulières. Nous la présentons de façon à expliciter l'élimination des conditions aux limites tangentielles, nécessaire pour la discrétisation de la méthode à poids.

Soit N_c le nombre de sommets du domaine polyédrique Ω . Chaque sommet est un point de discrétisation du bord, quelque soit le maillage considéré. On appelle $\mathcal{C}$ l'ensemble de ces sommets. Soit N_a le nombre de points de discrétisation du bord se trouvant sur une arête de $\partial\Omega$ et n'étant pas un coin. On appelle $\mathcal{A}$ l'ensemble de ces sommets. On pose $N_f = N_{\partial\Omega} - N_c - N_a$, le nombre de points de discrétisation du bord qui ne trouvent pas sur des arêtes. On ordonne les points du bord $\partial\Omega$ de la façon suivante :

- $\forall i \in I_f, M_i \in \partial\Omega \setminus (\mathcal{C} \cup \mathcal{A})$,
- $\forall i \in I_a, M_i \in \mathcal{A}$,
- $\forall i \in I_c, M_i \in \mathcal{C}$,

où on a décomposé $I_{\partial\Omega}$ en : $I_f = \{N_{\Omega} + 1, \dots, N_{\Omega} + N_f\}$, $I_a = \{N_{\Omega} + N_f + 1, \dots, N_{\Omega} + N_f + N_a\}$ et $I_c = \{N - N_c + 1, \dots, N\}$.

On remarque que : $N_{\Omega} + N_f + N_a + N_c = N$.

On considère comme espace de discrétisation de $\mathcal{X}_{\mathcal{E}}^{0,R}$ le sous-espace de $\mathcal{X}_k$, conforme dans $\mathcal{X}_{\mathcal{E}}^0$:

$$\mathcal{X}_{\mathcal{E},k}^{0,R} = \{ \mathbf{v} \in \mathcal{X}_k \mid \mathbf{v} \times \boldsymbol{\nu}|_{\partial\Omega} = 0 \text{ sur } \partial\Omega \}.$$

Par construction, $\mathcal{X}_{\mathcal{E},k}^{0,R} \subset \mathcal{H}^1(\Omega)$. La discrétisation de (8.14) se fait de la même façon que celle de (8.12), mais il faut prendre en compte la condition aux limites tangentielle dans les matrices $\mathbb{D}\text{iv}$ et $\mathbb{R}\text{ot}$, et dans les termes du second membre.

10.4.1 Élimination des conditions aux limites essentielles

Propriété 10.3 *L'espace $\mathcal{X}_{\mathcal{E},k}^{0,R}$ est de dimension finie $3N - 2N_f - 3N_a - 3N_c$.*

DÉMONSTRATION. Soit $\mathbf{u} \in \mathcal{X}_{\mathcal{E},k}^{0,R}$. Comme $\mathcal{X}_{\mathcal{E},k}^{0,R}$ est un sous-espace de $\mathcal{X}_{\mathcal{E},k}$, il est de dimension inférieure ou égale à $3N$ et on peut écrire $\mathbf{u} = \sum_{i=1}^N \sum_{\alpha=1}^3 u_{\alpha}(M_i) \mathbf{v}_{i,\alpha}$.

• Considérons un point M_i , $i \in I_a$, se trouvant sur l'arête $A_{k,l}$ (figure 7.1, p. 153). On a :

$$\begin{cases} \mathbf{u}(M_i) \cdot \boldsymbol{\tau}_{k,l} &= 0, \\ \mathbf{u}(M_i) \cdot \boldsymbol{\tau}_l &= 0, \\ \mathbf{u}(M_i) \cdot \boldsymbol{\tau}_k &= 0. \end{cases}$$

Comme les vecteurs $\boldsymbol{\tau}_{k,l}$, $\boldsymbol{\tau}_k$ et $\boldsymbol{\tau}_l$ ne sont pas colinéaires, on a nécessairement $\mathbf{u}(M_i) = 0$. Ceci étant valable pour tous les points d'arêtes, on peut écrire : $\mathbf{u} = \sum_{i \in I_a} \sum_{\alpha=1}^3 u_{\alpha}(M_i) \mathbf{v}_{i,\alpha}$. Pour chaque

vecteur de $\mathcal{X}_{\mathcal{E},k}^{0,R}$, nous éliminons alors les trois degrés de liberté liés à chaque point d'arête. D'où : $\dim(\mathcal{X}_{\mathcal{E},k}^{0,R}) \leq 3N - 3N_a$.

• Considérons un coin M_i , $i \in I_c$. M_i est l'intersection d'au moins deux arêtes, d'où $\mathbf{u}(M_i) = 0$.

Ceci étant valable pour tous les coins, on peut écrire : $\mathbf{u} = \sum_{i=1}^{N-N_a-N_c} \sum_{\alpha=1}^3 u_{\alpha}(M_i) \mathbf{v}_{i,\alpha}$. Pour chaque

vecteur de $\mathcal{X}_{\mathcal{E},k}^{0,R}$, nous éliminons alors les trois degrés de liberté liés à chaque coin. Il y en a N_c , d'où : $\dim(\mathcal{X}_{\mathcal{E},k}^{0,R}) \leq 3N - 3N_a - 3N_c$.

• Considérons un point M_i , $i \in I_f$, se trouvant sur la face F_k de Ω . Comme $\mathbf{u}(M_i) \times \boldsymbol{\nu}|_{F_k} = 0$, on

peut écrire $\mathbf{u}(M_i)|_{F_k} = u_{\nu}(M_i) \boldsymbol{\nu}|_{F_k}$. D'où : $\mathbf{u} = \sum_{i=1}^{N_{\Omega}} \sum_{\alpha=1}^3 u_{\alpha}(M_i) \mathbf{v}_{i,\alpha} + \sum_{i=1+N_{\Omega}}^{N_{\Omega}+N_f} u_{\nu}(M_i) v_i \boldsymbol{\nu}_i$.

Notons qu'il n'y a pas d'ambiguïté sur $\boldsymbol{\nu}(M_i)$ lorsque $i \in I_f$ car alors M_i n'est ni sur une arête ni sur un coin. Ainsi, pour chaque vecteur de $\mathcal{X}_{\mathcal{E},k}^{0,R}$, nous éliminons deux degrés de liberté par point du bord qui n'est pas ni sur une arête ni sur un coin. D'où : $\dim(\mathcal{X}_{\mathcal{E},k}^{0,R}) \leq 3N - 2N_f - 3N_a - 3N_c$.

Pour conclure, la famille $B_0 := (\mathbf{v}_{i,\alpha})_{i \in I_{\omega}, \alpha \in \{1,2,3\}} \cup (v_i \boldsymbol{\nu}_i)_{i \in I_f}$ est contenue dans $\mathcal{X}_{\mathcal{E},k}^{0,R}$, d'où : $\dim(\mathcal{X}_{\mathcal{E},k}^{0,R}) \geq 3N - 2N_f - 3N_a - 3N_c$. □

Nous avons donc montré la propriété 10.3, et au passage que B_0 engendre $\mathcal{X}_{\mathcal{E},k}^{0,R}$. Par la suite, on utilisera les notations suivantes : $I_{ac} = I_a \cup I_c$, et $N_{ac} = N_a + N_c$. Pour $M_i \in I_f$, se trouvant sur la face F_k , on appellera $(\boldsymbol{\tau}_k^1, \boldsymbol{\tau}_k^2)$ la base du plan généré par la face F_k , telle que $(\boldsymbol{\tau}_k^1, \boldsymbol{\tau}_k^2, \boldsymbol{\nu}_k)$ soit une base orthonormée directe. On utilisera donc deux types de base locale de $\mathbb{R}^3$ selon l'emplacement de M_i :

- $\forall i \in I_{\Omega}$, $M_i \in \Omega$, on travaille dans la base canonique $(\mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3)$,
- $\forall i \in I_f$, $M_i \in \partial\Omega \setminus (\mathcal{A} \cup \mathcal{C})$, on travaille dans la base locale $(\boldsymbol{\tau}_k^1, \boldsymbol{\tau}_k^2, \boldsymbol{\nu}_k)$,
- $\forall i \in I_{ac}$, $M_i \in \mathcal{A} \cup \mathcal{C}$, afin de lever l'ambiguïté sur le vecteur normal, on travaille dans la base canonique $(\mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3)$.

On appelle B la base de $\mathcal{X}_k$ ainsi formée :

$$B = (\mathbf{v}_{i,\alpha})_{i \in I_{\Omega} \cup I_{ac}, \alpha \in \{1,2,3\}} \cup (v_i \boldsymbol{\tau}_i^1)_{i \in I_f} \cup (v_i \boldsymbol{\tau}_i^2)_{i \in I_f} \cup (v_i \boldsymbol{\nu}_i)_{i \in I_f}. \quad (10.12)$$

Soit $\mathbb{A} \in (\mathbb{R}^{3 \times 3})^{N \times N}$ la décomposition de la matrice $\mathbb{D}iv + \mathbb{R}ot$ dans B , qu'on écrit de la façon

suivante :

$$\mathbb{A} = \begin{pmatrix} \mathbb{A}_{\Omega, \Omega} & \mathbb{A}_{\Omega, f} & \mathbb{A}_{\Omega, ac} \\ \mathbb{A}_{f, \Omega} & \mathbb{A}_{f, f} & \mathbb{A}_{f, ac} \\ \mathbb{A}_{ac, \Omega} & \mathbb{A}_{ac, f} & \mathbb{A}_{ac, ac} \end{pmatrix}.$$

Décrivons les sous-matrices de $\mathbb{A}$.

Les sous-blocs $\mathbb{A}^{i,j} \in \mathbb{R}^3$ tels que i et $j \in I_{\Omega} \cup I_{ac}$ sont définis dans la base $(\mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3)^T$ et ont déjà été calculés dans le paragraphe 10.3.1. Ils sont contenus dans les sous-matrices suivantes :

- $\mathbb{A}_{\Omega, \Omega} = (\text{Div} + \text{Rot})^{i \in I_{\Omega}, j \in I_{\Omega}} \in (\mathbb{R}^{3 \times 3})^{N_{\Omega} \times N_{\Omega}}$;
- $\mathbb{A}_{ac, ac} = (\text{Div} + \text{Rot})^{i \in I_{ac}, j \in I_{ac}} \in (\mathbb{R}^{3 \times 3})^{N_{ac} \times N_{ac}}$;
- $\mathbb{A}_{\Omega, ac} = (\text{Div} + \text{Rot})^{i \in I_{\Omega}, j \in I_{ac}} \in (\mathbb{R}^{3 \times 3})^{N_{\Omega} \times N_{ac}}$;
- $\mathbb{A}_{ac, \Omega} = (\text{Div} + \text{Rot})^{i \in I_{ac}, j \in I_{\Omega}} \in (\mathbb{R}^{3 \times 3})^{N_{ac} \times N_{\Omega}}$.

La sous-matrice $\mathbb{A}_{f, f} \in (\mathbb{R}^{3 \times 3})^{N_f \times N_f}$ est composée des $N_f \times N_f$ sous-blocs $\mathbb{A}_{f, f}^{i, j} \in \mathbb{R}^{3 \times 3}$, tels que : $\forall i \in I_f, \forall j \in I_f$:

$$\mathbb{A}_{f, f}^{i, j} = \begin{pmatrix} (v_j \boldsymbol{\tau}_j^1, v_i \boldsymbol{\tau}_i^1)_{\mathcal{X}_{\varepsilon}^0} & (v_j \boldsymbol{\tau}_j^2, v_i \boldsymbol{\tau}_i^1)_{\mathcal{X}_{\varepsilon}^0} & (v_j \boldsymbol{\nu}_j, v_i \boldsymbol{\tau}_i^1)_{\mathcal{X}_{\varepsilon}^0} \\ (v_j \boldsymbol{\tau}_j^1, v_i \boldsymbol{\tau}_i^2)_{\mathcal{X}_{\varepsilon}^0} & (v_j \boldsymbol{\tau}_j^2, v_i \boldsymbol{\tau}_i^2)_{\mathcal{X}_{\varepsilon}^0} & (v_j \boldsymbol{\nu}_j, v_i \boldsymbol{\tau}_i^2)_{\mathcal{X}_{\varepsilon}^0} \\ (v_j \boldsymbol{\tau}_j^1, v_i \boldsymbol{\nu}_i)_{\mathcal{X}_{\varepsilon}^0} & (v_j \boldsymbol{\tau}_j^2, v_i \boldsymbol{\nu}_i)_{\mathcal{X}_{\varepsilon}^0} & (v_j \boldsymbol{\nu}_j, v_i \boldsymbol{\nu}_i)_{\mathcal{X}_{\varepsilon}^0} \end{pmatrix}.$$

Ces sous-blocs agissent sur les vecteurs de $\mathbb{R}^3$ décomposés dans la base locale $(\boldsymbol{\tau}_j^1, \boldsymbol{\tau}_j^2, \boldsymbol{\nu}_j)^T$:

$$\mathcal{E}_h(M_j) = \begin{pmatrix} \mathbf{E}_{\tau, 1}(M_j) \\ \mathbf{E}_{\tau, 2}(M_j) \\ \mathbf{E}_{\nu}(M_j) \end{pmatrix}, \text{ où } \mathbf{E}_{\tau, \alpha}(M_j) = \mathbf{E}_h(M_j) \cdot \boldsymbol{\tau}_j^{\alpha}, \alpha = 1, 2, \text{ et } \mathbf{E}_{\nu}(M_j) = \mathcal{E}_h(M_j) \cdot \boldsymbol{\nu}_j.$$

La sous-matrice $\mathbb{A}_{\Omega, f} \in (\mathbb{R}^{3 \times 3})^{N_{\Omega} \times N_f}$ est composée des $N_{\Omega} \times N_f$ sous-blocs de $\mathbb{R}^{3 \times 3}$ suivants : $\forall i \in I_{\Omega}, \forall j \in I_f$,

$$\mathbb{A}_{\Omega, f}^{i, j} = \begin{pmatrix} (v_j \boldsymbol{\tau}_j^1, v_i \mathbf{e}_1)_{\mathcal{X}_{\varepsilon}^0} & (v_j \boldsymbol{\tau}_j^2, v_i \mathbf{e}_1)_{\mathcal{X}_{\varepsilon}^0} & (v_j \boldsymbol{\nu}_j, v_i \mathbf{e}_1)_{\mathcal{X}_{\varepsilon}^0} \\ (v_j \boldsymbol{\tau}_j^1, v_i \mathbf{e}_2)_{\mathcal{X}_{\varepsilon}^0} & (v_j \boldsymbol{\tau}_j^2, v_i \mathbf{e}_2)_{\mathcal{X}_{\varepsilon}^0} & (v_j \boldsymbol{\nu}_j, v_i \mathbf{e}_2)_{\mathcal{X}_{\varepsilon}^0} \\ (v_j \boldsymbol{\tau}_j^1, v_i \mathbf{e}_3)_{\mathcal{X}_{\varepsilon}^0} & (v_j \boldsymbol{\tau}_j^2, v_i \mathbf{e}_3)_{\mathcal{X}_{\varepsilon}^0} & (v_j \boldsymbol{\nu}_j, v_i \mathbf{e}_3)_{\mathcal{X}_{\varepsilon}^0} \end{pmatrix}.$$

Symétriquement, la sous-matrice $\mathbb{A}_{f, \Omega} \in (\mathbb{R}^{3 \times 3})^{N_f \times N_{\Omega}}$ est telle que :

$$\forall i \in I_f, \forall j \in I_{\Omega}, \mathbb{A}_{f, \Omega}^{i, j} = \mathbb{A}_{\Omega, f}^{j, i}.$$

De même, la sous-matrice $\mathbb{A}_{f, ac} \in (\mathbb{R}^{3 \times 3})^{N_f \times N_{ac}}$ est composée des $N_f \times N_{ac}$ sous-blocs de

$\mathbb{R}^{3 \times 3}$ suivants : $\forall i \in I_f, \forall j \in I_{ac}$,

$$\mathbb{A}_{f,ac}^{i,j} = \begin{pmatrix} (v_j \mathbf{e}_1, v_i \boldsymbol{\tau}_i^1)_{\mathcal{X}_\varepsilon^0} & (v_j \mathbf{e}_2, v_i \boldsymbol{\tau}_i^1)_{\mathcal{X}_\varepsilon^0} & (v_j \mathbf{e}_3, v_i \boldsymbol{\tau}_i^1)_{\mathcal{X}_\varepsilon^0} \\ (v_j \mathbf{e}_1, v_i \boldsymbol{\tau}_i^2)_{\mathcal{X}_\varepsilon^0} & (v_j \mathbf{e}_2, v_i \boldsymbol{\tau}_i^2)_{\mathcal{X}_\varepsilon^0} & (v_j \mathbf{e}_3, v_i \boldsymbol{\tau}_i^2)_{\mathcal{X}_\varepsilon^0} \\ (v_j \mathbf{e}_1, v_i \boldsymbol{\nu}_i)_{\mathcal{X}_\varepsilon^0} & (v_j \mathbf{e}_2, v_i \boldsymbol{\nu}_i)_{\mathcal{X}_\varepsilon^0} & (v_j \mathbf{e}_3, v_i \boldsymbol{\nu}_i)_{\mathcal{X}_\varepsilon^0} \end{pmatrix},$$

et symétriquement, la sous-matrice $\mathbb{A}_{ac,f} \in (\mathbb{R}^{3 \times 3})^{N_f \times N_\Omega}$ est telle que :

$$\forall i \in I_f, \forall j \in I_\Omega, \mathbb{A}_{ac,f}^{i,j} = \mathbb{A}_{f,ac}^{j,i}.$$

Nous ne détaillerons pas les calculs des sous-matrices de $\mathbb{A}$, car ils sont similaires aux calculs des éléments de $\mathbb{R}ot$ et $\mathbb{D}iv$, expliqués dans le paragraphe 10.3.1.

Soit $\mathcal{F} \in \mathcal{X}_k$. Afin de définir le produit matrice-vecteur $\mathbb{A}\underline{\mathcal{F}}$, où $\underline{\mathcal{F}}$ est le vecteur de $(\mathbb{R}^3)^N$ associé à $\mathcal{F}$, on décompose $\mathcal{F}$ dans B :

$$\mathcal{F} = \sum_{j \in I_\Omega \cup I_{ac}} \sum_{\beta=1}^3 F_\beta(M_j) \mathbf{v}_{j,\beta} + \sum_{j \in I_f} (F_{\tau,1}(M_j) v_j \boldsymbol{\tau}_j^1 + F_{\tau,2}(M_j) v_j \boldsymbol{\tau}_j^2 + F_\nu(M_j) v_j \boldsymbol{\nu}_j),$$

où $\forall i \in I_f, \forall \alpha \in \{1,2\}, F_{\tau,\alpha}(M_j) = \mathcal{F}(M_i) \cdot \boldsymbol{\tau}_i^\alpha$ et $F_\nu(M_j) = \mathcal{F}(M_i) \cdot \boldsymbol{\nu}_i$.

On considère alors $\underline{\mathcal{F}} = (\underline{\mathcal{F}}_\Omega, \underline{\mathcal{F}}_f, \underline{\mathcal{F}}_{ac})^T$ le vecteur de $\mathbb{R}^{3N}$ associé, dont les sous-composantes sont :

- $\underline{\mathcal{F}}_\Omega = (F_1(M_1), F_2(M_1), F_3(M_1), \dots, F_1(M_{N_\Omega}), F_2(M_{N_\Omega}), F_3(M_{N_\Omega}))^T \in (\mathbb{R}^3)^{N_\Omega}$,
- $\underline{\mathcal{F}}_f = (F_{\tau,1}(M_{N_\Omega+1}), F_{\tau,2}(M_{N_\Omega+1}), F_\nu(M_{N_\Omega+1}), \dots, F_{\tau,1}(M_{N-N_{ac}}), F_{\tau,2}(M_{N-N_{ac}}), F_\nu(M_{N-N_{ac}}))^T \in (\mathbb{R}^3)^{N_f}$,
- $\underline{\mathcal{F}}_{ac} = (F_1(M_{N-N_{ac}+1}), F_2(M_{N-N_{ac}+1}), F_3(M_{N-N_{ac}+1}), \dots, F_1(M_N), F_2(M_N), F_3(M_N))^T \in (\mathbb{R}^3)^{N_{ac}}$.

Le produit matrice-vecteur $\mathbb{A}\underline{\mathcal{F}}$ est ainsi bien défini. Pour construire la matrice de la formulation variationnelle (8.14) discrétisée dans $\mathcal{X}_{\varepsilon,k}^{0,R}$, on peut alors procéder à l'élimination des conditions aux limites essentielles dans $\mathbb{A}$. Cela correspond à éliminer d'une part les lignes et les colonnes de $\mathbb{A}$ dont les éléments dépendent de $\boldsymbol{\tau}^{1,2}$ pour les sommets situés à l'intérieur des faces du bord ($i \in I_f$) ; et d'autre part toutes les lignes et les colonnes pour les sommets situés sur les arêtes de $\partial\Omega$ ($i \in I_{ac}$). On appelle $\mathbb{A}_0$ la matrice ainsi obtenue.

Soit $\mathbb{A}_{\Omega,\nu} \in (\mathbb{R}^{3 \times 3})^{N_\Omega \times N_f}$ la matrice $\mathbb{A}_{\Omega,f}$ dont on a éliminé les colonnes dépendant de $\boldsymbol{\tau}^{1,2}$: $\forall i \in I_\Omega, \forall j \in I_f$,

$$\mathbb{A}_{\Omega,\nu}^{i,j} = \begin{pmatrix} 0 & 0 & (v_j \boldsymbol{\nu}_j, v_i \mathbf{e}_1)_{\mathcal{X}_\varepsilon^0} \\ 0 & 0 & (v_j \boldsymbol{\nu}_j, v_i \mathbf{e}_2)_{\mathcal{X}_\varepsilon^0} \\ 0 & 0 & (v_j \boldsymbol{\nu}_j, v_i \mathbf{e}_3)_{\mathcal{X}_\varepsilon^0} \end{pmatrix}.$$

Symétriquement, $\mathbb{A}_{\nu,\Omega} \in (\mathbb{R}^{3 \times 3})^{N_f \times N_\Omega}$ est la matrice $\mathbb{A}_{f,\Omega}$ dont on a éliminé les lignes qui dépendent de $\boldsymbol{\tau}^{1,2}$: $\mathbb{A}_{\nu,\Omega}^{i,j} = \mathbb{A}_{\Omega,\nu}^{j,i}$.

On considère $\mathbb{A}_{\nu,\nu}$ la matrice $\mathbb{A}_{f,f}$ dont on a éliminé les termes dépendant de $\boldsymbol{\tau}^{1,2}$, sauf les termes diagonaux, pour lesquels on a imposé la valeur 1 :

$$\mathbb{A}_{\nu,\nu}^{i,j} = \begin{pmatrix} \delta_{ij} & 0 & 0 \\ 0 & \delta_{ij} & 0 \\ 0 & 0 & (v_j \boldsymbol{\nu}_j, v_i \boldsymbol{\nu}_i)_{\mathcal{X}_\varepsilon^0} \end{pmatrix}.$$

En effet, $\mathbb{A}_{f,f}$ correspond à un bloc diagonal de $\mathbb{A}$. En imposant la valeur 1 sur les termes diagonaux éliminés, on s'assure que la matrice ainsi modifiée soit inversible.

On en déduit que la structure par blocs de la matrice $\mathbb{A}_0 \in (\mathbb{R}^{3 \times 3})^{N \times N}$ est la suivante :

$$\mathbb{A}_0 = \begin{pmatrix} \mathbb{A}_{\Omega,\Omega} & \mathbb{A}_{\Omega,\nu} & 0 \\ \mathbb{A}_{\nu,\Omega} & \mathbb{A}_{\nu,\nu} & 0 \\ 0 & 0 & \mathbb{I}_{ac,ac} \end{pmatrix},$$

où $\mathbb{I}_{ac,ac}$ est la matrice identité de $(\mathbb{R}^{3 \times 3})^{N_{ac} \times N_{ac}}$.

Proposition 10.4 *La matrice $\mathbb{A}_0$ est inversible.*

DÉMONSTRATION. Les sous-blocs diagonaux de $\mathbb{A}_{\omega,\omega}$ et $\mathbb{A}_{\nu,\nu}$ sont strictement positifs. On a :

$$\mathbb{A}_{\Omega,\Omega}^{i,i} = \begin{pmatrix} \|\mathbf{grad} v_i\|_0^2 & 0 & 0 \\ 0 & \|\mathbf{grad} v_i\|_0^2 & 0 \\ 0 & 0 & \|\mathbf{grad} v_i\|_0^2 \end{pmatrix}, \quad \text{et } \mathbb{A}_{\nu,\nu}^{i,i} = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & \|\mathbf{grad} v_i\|_0^2 \end{pmatrix}.$$

Le sous-bloc diagonal $\mathbb{I}_{ac,ac}$ garantit donc l'inversibilité de $\mathbb{A}_0$. □

10.4.2 Champ électrique

Tout se passe comme dans le cas bidimensionnel (voir la section 4.7 et le paragraphe 4.7.1), aussi nous nous contentons de donner les éléments matriciels, sans détailler la discrétisation.

Le problème variationnel (8.14) discrétisé dans $\mathcal{X}_{\mathcal{E},k}^{0,R}$ s'écrit :

Trouver $\mathcal{E}_h^0 \in \mathcal{X}_{\mathcal{E},k}^{0,R}$ tel que :

$$\begin{aligned} \forall i \in I_\Omega, \forall \alpha \in \{1, 2, 3\}, \quad (\mathcal{E}_h^0, v_i \mathbf{e}_\alpha)_{\mathcal{X}_\mathcal{E}^0} &= \mathcal{L}_0(v_i \mathbf{e}_\alpha) - (\mathcal{E}^r, v_i \mathbf{e}_\alpha)_{\mathcal{X}_\mathcal{E}^0}, \\ \forall i \in I_f, \quad (\mathcal{E}_h^0, v_i \boldsymbol{\nu}_i)_{\mathcal{X}_\mathcal{E}^0} &= \mathcal{L}_0(v_i \boldsymbol{\nu}_i) - (\mathcal{E}^r, v_i \boldsymbol{\nu}_i)_{\mathcal{X}_\mathcal{E}^0}. \end{aligned}$$

En regroupant $\mathcal{E}_h^0$ et $\mathcal{E}^r$, on peut réécrire ces équations sous la forme :

Trouver $\mathcal{E}_h \in \mathcal{X}_k$ tel que :

$$\begin{aligned} \forall i \in I_\Omega, \forall \alpha \in \{1, 2, 3\}, \quad (\mathcal{E}_h, \mathbf{v}_{i,\alpha})_{\mathcal{X}_\mathcal{E}^0} &= \mathcal{L}_0(\mathbf{v}_{i,\alpha}), \\ \forall i \in I_f, \quad (\mathcal{E}_h, v_i \boldsymbol{\nu}_i)_{\mathcal{X}_\mathcal{E}^0} &= \mathcal{L}_0(v_i \boldsymbol{\nu}_i), \\ \forall i \in I_f, \quad \mathcal{E}_h(M_i) \times \boldsymbol{\nu}_i &= \mathbf{e}(M_i) \times \boldsymbol{\nu}_i, \\ \forall i \in I_{ac}, \quad \mathcal{E}_h(M_i) &= \mathbf{e}(M_i). \end{aligned} \tag{10.13}$$

La dernière égalité a lieu au sens du paragraphe 4.8.2 de la partie II, et est détaillée dans le paragraphe 10.4.3.

Soit $\underline{\mathbf{E}} = (\underline{\mathbf{E}}_\Omega, \underline{\mathbf{E}}_f, \underline{\mathbf{E}}_{ac})^T$ le vecteur de $\mathbb{R}^{3N}$ associé dont les composantes sont égales à celles de $\mathcal{E}_h$ dans B . Considérons le relèvement discret suivant : $\underline{\mathbf{E}}^r = (\underline{\mathbf{Z}}_\Omega, \underline{\mathbf{E}}_r, \underline{\mathbf{E}}_{ac})^T$, où :

- $\underline{\mathbf{Z}}_\Omega$ est le vecteur nul de $(\mathbb{R}^3)^{N_\Omega}$;

- $\underline{\mathbf{E}}_r^r = (\mathbf{e}(M_{N_\Omega+1}).\boldsymbol{\tau}_{N_\Omega+1}^1, \mathbf{e}(M_{N_\Omega+1}).\boldsymbol{\tau}_{N_\Omega+1}^2, 0, \dots, \mathbf{e}(M_{N_\Omega+N_f}).\boldsymbol{\tau}_{N_\Omega+N_f}^1, \mathbf{e}(M_{N_\Omega+N_f}).\boldsymbol{\tau}_{N_\Omega+N_f}^2, 0)^T \in$

$(\mathbb{R}^3)^{N_f}$, est le vecteur contenant les valeurs des composantes tangentielles de $\mathcal{E}_h$ sur $\partial\Omega \setminus (\mathcal{A} \cup \mathcal{C})$;

- $\underline{\mathbf{E}}_{ac}^r = (e_1(M_{N_\Omega+N_f+1}), e_2(M_{N_\Omega+N_f+1}), e_3(M_{N_\Omega+N_f+1}), \dots, e_1(M_N), e_2(M_N), e_3(M_N))^T \in (\mathbb{R}^3)^{N_{ac}}$ est le vecteur contenant les valeurs des composantes de $\mathbf{e}$ sur les arêtes et aux coins du maillage.

Soit $\mathbb{A}_{\nu,f} \in (\mathbb{R}^{3 \times 3})^{N_f \times N_f}$ la matrice $\mathbb{A}_{f,f}$ dont on a éliminé les lignes dépendant de $\boldsymbol{\tau}^{1,2}$:
 $\forall i, j \in I_f$,

$$\mathbb{A}_{\nu,f}^{i,j} = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & 0 \\ (v_j \boldsymbol{\tau}_j^1, v_i \boldsymbol{\nu}_i)_{\mathcal{X}_\varepsilon^0} & (v_j \boldsymbol{\tau}_j^2, v_i \boldsymbol{\nu}_i)_{\mathcal{X}_\varepsilon^0} & (v_j \boldsymbol{\nu}_j, v_i \boldsymbol{\nu}_i)_{\mathcal{X}_\varepsilon^0} \end{pmatrix}.$$

Soit $\mathbb{A}_{\nu,ac} \in (\mathbb{R}^{3 \times 3})^{N_f \times N_{ac}}$, la matrice $\mathbb{A}_{f,ac}$ dont on a éliminé les lignes dépendant de $\boldsymbol{\tau}^{1,2}$:
 $\forall i \in I_f, \forall j \in I_{ac}$,

$$\mathbb{A}_{\nu,ac}^{i,j} = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & 0 \\ (v_j \mathbf{e}_1, v_i \boldsymbol{\nu}_i)_{\mathcal{X}_\varepsilon^0} & (v_j \mathbf{e}_2, v_i \boldsymbol{\nu}_i)_{\mathcal{X}_\varepsilon^0} & (v_j \mathbf{e}_3, v_i \boldsymbol{\nu}_i)_{\mathcal{X}_\varepsilon^0} \end{pmatrix}.$$

Soit $\underline{\mathbf{L}}_\Omega \in (\mathbb{R}^3)^{N_\Omega}$ la restriction de $\underline{\mathbf{L}}$ aux $i \in I_\Omega$, et $\underline{\mathbf{L}}_\nu \in (\mathbb{R}^3)^{N_f}$ le vecteur tel que :

$$\underline{\mathbf{L}}_\nu = (0, 0, \mathcal{L}_0(v_{N_\Omega+1} \boldsymbol{\nu}_{N_\Omega+1}), \dots, 0, 0, \mathcal{L}_0(v_{N_\Omega+N_f} \boldsymbol{\nu}_{N_\Omega+N_f})).$$

Soit $\mathbb{A}_r \in (\mathbb{R}^{3 \times 3})^{N \times N}$ la matrice définie ainsi :

$$\mathbb{A}_r = \begin{pmatrix} 0 & \mathbb{A}_{\Omega,f} & \mathbb{A}_{\Omega,ac} \\ 0 & \mathbb{A}_{\nu,f} & \mathbb{A}_{\nu,ac} \\ 0 & 0 & -\mathbb{I}_{ac,ac} \end{pmatrix}.$$

Posons $\underline{\mathbf{L}}_0 = (\underline{\mathbf{L}}_\Omega, \underline{\mathbf{L}}_\nu, \underline{\mathbf{Z}}_{ac})^T \in (\mathbb{R}^3)^N$, où $\underline{\mathbf{Z}}_{ac}$ est le vecteur nul de $(\mathbb{R}^3)^{N_{ac}}$. On montre de la même façon que dans le paragraphe 4.7.1 que les équations (10.13) se réécrivent sous la forme du système matriciel suivant :

$$\mathbb{A}_0 \underline{\mathbf{E}} = \underline{\mathbf{L}}_0 - \mathbb{A}_r \underline{\mathbf{E}}^r.$$

Le calcul de $\underline{\mathbf{L}}_0$ est similaire à celui de $\underline{\mathbf{L}}$, effectué dans le paragraphe 10.3.3. Une fois que $\underline{\mathbf{E}}$ est calculé, on peut écrire l'approximation calculée sur l'intérieur des faces du bord, $\underline{\mathbf{E}}_f$ dans la base canonique $(\mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3)$.

Pour cela, à chaque point M_i , $i \in I_f$, on associe $\mathbb{O}_{f,i} \in \mathbb{R}^{3 \times 3}$, la matrice de changement de base telle que :

$$\mathbb{O}_{f,i} = \begin{pmatrix} \tau_1^1(M_i) & \tau_1^2(M_i) & \nu_1(M_i) \\ \tau_2^1(M_i) & \tau_2^2(M_i) & \nu_2(M_i) \\ \tau_3^1(M_i) & \tau_3^2(M_i) & \nu_3(M_i) \end{pmatrix},$$

avec $\boldsymbol{\tau}^1 = \sum_{\alpha=1}^3 \tau_\alpha^1 \mathbf{e}_\alpha$ et $\boldsymbol{\tau}^2 = \sum_{\alpha=1}^3 \tau_\alpha^2 \mathbf{e}_\alpha$. Soit $\mathbb{O}_f \in (\mathbb{R}^3)^{N_f \times N_f}$ la matrice diagonale par bloc définie par :

$$\mathbb{O}_f = \begin{pmatrix} \mathbb{O}_{f, N_\Omega + 1} & 0 & 0 \\ 0 & \ddots & 0 \\ 0 & 0 & \mathbb{O}_{f, N - N_{ac}} \end{pmatrix}.$$

Le vecteur correspondant aux composantes dans la base $(\mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3)^T$ s'écrit alors : $\mathbb{O}_f \underline{\mathbf{E}}_f$.

Remarque 10.5 Comme dans le cas bidimensionnel, en pratique, on crée la matrice et le second membre dans la base cartésienne complète, et on procède à une élimination a posteriori, en faisant un changement de base avant l'élimination, et un autre après pour revenir au système cartésien. Peu importe la base $(\boldsymbol{\tau}^1, \boldsymbol{\tau}^2)^T$ du bord choisie.

10.4.3 Relèvement de la CL

Comme en 2D, il n'est pas nécessaire de calculer un relèvement explicite de $\mathbf{e}$ pour calculer $\mathcal{E}_h$. Nous détaillons ici comment déterminer $\underline{\mathbf{E}}_{ac}^r$ dans le cas où $\partial\Omega = \overline{\Gamma_A} \cup \overline{\Gamma_C}$ (voir section 8.1). Pour les points M_i qui se trouvent sur les arêtes ou les coins de Γ_C , on a toujours $(\underline{\mathbf{E}}_{ac}^r)_i = (0, 0, 0)^T$. En pratique, Γ_A est un plan, et les points qui sont sur la frontière $\overline{\Gamma_A} \cap \overline{\Gamma_C}$ sont géométriquement soit des arêtes, soit des coins (voir par exemple la figure 2.3 de la partie I). On les considère comme étant des points de Γ_C . Prenons $M_i \in \overline{\Gamma_A} \cap \overline{\Gamma_C}$. Si M_i est géométriquement sur une arête, alors $i \in I_f$: M_i est interprété comme étant un point d'une face de Γ_C . Si M_i est géométriquement un coin, alors $i \in I_a$: M_i est interprété comme étant un point d'une arête de Γ_C , d'où : $(\underline{\mathbf{E}}_{ac}^r)_i = (0, 0, 0)^T$.

10.5 Régularisation à poids : discrétisation

Pour les valeurs de γ ad hoc, $\mathcal{X}_{\mathcal{E}, \gamma}^0 \cap \mathcal{H}^1(\Omega)$ est dense dans $\mathcal{X}_{\mathcal{E}, \gamma}^0$, on peut donc considérer comme espace d'approximation de $\mathcal{X}_{\mathcal{E}, \gamma}^0$ l'espace : $\mathcal{X}_{\mathcal{E}, k}^{0, R}$, décrit dans le paragraphe 10.4.2. La formulation variationnelle (8.17) s'écrit :

Trouver $\mathcal{E}_{\gamma, h}^0 \in \mathcal{X}_{\mathcal{E}, k}^{0, R}$ tel que :

$$\begin{aligned} \forall i \in I_\Omega, \forall \alpha \in \{1, 2\} \quad & (\mathcal{E}_h^0, v_i \mathbf{e}_\alpha)_{\mathcal{X}_{\mathcal{E}, \gamma}^0} = \mathcal{L}_\gamma(v_i \mathbf{e}_\alpha) - (\mathcal{E}^r, v_i \mathbf{e}_\alpha)_{\mathcal{X}_{\mathcal{E}, \gamma}^0}, \\ \forall i \in I_f, \quad & (\mathcal{E}_h^0, v_i \boldsymbol{\nu}_i)_{\mathcal{X}_{\mathcal{E}, \gamma}^0} = \mathcal{L}_\gamma(v_i \boldsymbol{\nu}_i) - (\mathcal{E}^r, v_i \boldsymbol{\nu}_i)_{\mathcal{X}_{\mathcal{E}, \gamma}^0}. \end{aligned} \quad (10.14)$$

Pour exprimer ces équations sous forme matricielle, on procède de la même façon que dans le paragraphe 10.4.2 par pseudo-élimination des degrés de liberté connus. Soit $\mathbb{A}_{0, \gamma}$ (*resp.* $\mathbb{A}_{r, \gamma}$) la matrice obtenue en remplaçant l'opérateur $\mathcal{A}_0$ par $\mathcal{A}_\gamma$ dans $\mathbb{A}_0$ (*resp.* $\mathbb{A}_r$). On se ramène au calcul direct de $\mathcal{E}_{\gamma, h} = \mathcal{E}_{\gamma, h}^0 + \mathcal{E}_h^r$, avec comme relèvement discret de la condition tangentielle au bord $\mathbf{e}$ le vecteur $\underline{\mathbf{E}}^r$. On construit le vecteur $\underline{\mathbf{L}}_\gamma$ de la même façon que $\underline{\mathbf{L}}_0$, en remplaçant l'opérateur $\mathcal{L}_0$ par $\mathcal{L}_\gamma$. Il s'agit alors de résoudre le système linéaire :

$$\mathbb{A}_\gamma \underline{\mathbf{E}}_\gamma = \underline{\mathbf{L}}_\gamma + \mathbb{A}_{r, \gamma} \underline{\mathbf{E}}^r.$$

Il faut construire la matrice $\mathbb{D}\text{iv}_\gamma$, qui correspond au produit scalaire dans $L_\gamma^2(\Omega)$ des divergences des vecteurs de la base B_0 . Cette construction est similaire à la construction de $\mathbb{D}\text{iv}$. On utilise un schéma d'intégration numérique à quinze points (voir le paragraphe 15.3.5), partie IV).

Remarque 10.6 D'après A. Buffa et al. [31, 32], le champ électrique est moins singulier dans la direction des arêtes rentrantes qu'orthogonalement à ces directions. Ainsi, on peut imaginer d'utiliser des éléments finis anisotropes, c'est-à-dire qu'en dehors d'un voisinage des coins rentrants, les arêtes des tétraèdres parallèles aux arêtes de Ω sont plus longues que les autres. En pratique, cela permet de réduire le nombre de points de discrétisation. Dans [32], les auteurs étudient la convergence des éléments finis d'arêtes sur un maillage anisotrope pour le problème harmonique.

10.6 Cas prismatique : discrétisation

Dans le cas où $\Omega = \omega \times]0, L[$, (avec ω polygone de $\mathbb{R}^2$ et $L > 0$) est un prisme droit, nous avons vu dans la section 8.6 qu'on pouvait réécrire le problème sous forme d'une suite de problèmes posés dans ω . Les systèmes d'équations obtenus sont découplés en :

- $((\hat{\mathbf{E}}^{0,k}, \mathbf{c}^k), \mathbf{E}_z^{0,k})$ si on traite le problème avec des conditions aux limites essentielles partout, ou en :

- $((\hat{\mathbf{E}}^k, \mathbf{c}_A^k), \mathbf{E}_z^{0,k})$, si on traite le problème avec des conditions aux limites essentielles seulement sur γ_C pour la partie longitudinale du champ.

On se ramène alors pour chaque $k \in \mathbb{N}$ à la résolution d'un système matriciel issu d'une discrétisation dans $\mathbf{X}_{\mathbf{E},k}^{0,R}$ et d'un système matriciel issu d'une discrétisation dans $V_k^0(\omega)$. Nous ne détaillerons pas la formation des matrices de discrétisation des formulations variationnelles, puisque cela a été fait dans le chapitre 4 de partie II, qui traite le problème bidimensionnel. Nous reprenons les notations relatives à cette partie. Rappelons que le nombre de coins rentrants du problème bidimensionnel est égal au nombre d'arêtes rentrantes du problème tridimensionnel.

10.6.1 Champ électrique transverse : CL essentielles

Dans ce paragraphe, on s'intéresse à la discrétisation de la formulation variationnelle (8.26)-(8.27) dans $\mathbf{X}_{\mathbf{E},k}^{0,R}$. On introduit de nouvelles matrices de $(\mathbb{R}^{2 \times 2})^{N \times N}$ et de nouveaux vecteurs de $(\mathbb{R}^2)^N$, obtenus par discrétisation (8.26)-(8.27) dans la base B de $\mathbb{R}^2$ (définie en (4.40), partie II, section 4.7).

Soit $\mathbb{M}_0 \in (\mathbb{R}^{2 \times 2})^{N \times N}$ la matrice de masse de l'espace $\mathbf{X}_{\mathbf{E},k}^{0,R}$. Cette matrice se forme comme la matrice $\mathbb{A}_0$ du problème bidimensionnel, en remplaçant le produit scalaire dans $\mathbf{X}_{\mathbf{E}}^0 : \mathcal{A}_0(\cdot, \cdot)$ par le produit $L^2(\omega) : (\cdot, \cdot)_{0,\omega}$. On posera de même $\mathbb{M}_r \in (\mathbb{R}^{2 \times 2})^{N \times N}$ la matrice de masse appliquée au relèvement discret, de la même forme que $\mathbb{A}_r$.

On appelle alors $\mathbb{A}_0^k = \mathbb{A}_0 + \frac{k^2 \pi^2}{L^2} \mathbb{M}_0$, la matrice dont les coefficients correspondent aux produits scalaires $\mathcal{A}_0^k(\cdot, \cdot)$ dans la base B_0 . De même, $\mathbb{A}_r^k = \mathbb{A}_r + \frac{k^2 \pi^2}{L^2} \mathbb{M}_r$. $\underline{\mathbf{e}}^k$ désignera la représentation de $\mathbf{e}_h^k$ décomposé dans la base B .

Soit $\mathbb{L}_k \in (\mathbb{R}^{1 \times 2})^{N_{ar} \times N}$, la matrice représentant les valeurs de $\mathcal{A}_0^k(\cdot, \mathbf{x}_{l,h}^S)$ dans la base B_0 : $\forall l \in \{1, \dots, N_{ar}\}$,

$$\begin{aligned} \mathbb{L}_k^{l,j} &= \left(\mathcal{A}_0^k(v_j \mathbf{e}_1, \mathbf{x}_{l,h}^S) \quad \mathcal{A}_0^k(v_j \mathbf{e}_2, \mathbf{x}_{l,h}^S) \right), \forall j \in I_\omega, \\ &= \left(0 \quad \mathcal{A}_0^k(v_j \boldsymbol{\nu}_j, \mathbf{x}_{l,h}^S) \right), \forall j \in I_a, \\ &= \left(0 \quad 0 \right), \forall j \in I_c. \end{aligned}$$

Pour calculer cette matrice, il faut d'abord approcher les $\mathbf{x}_l^S$, comme c'est expliqué dans le paragraphe 4.9.2 de la partie II.

Soit $\underline{\mathbf{L}}_k \in (\mathbb{R}^2)^N$, le vecteur représentant les valeurs approchées de $\mathcal{L}_0^k(\cdot)$ dans la base B_0 :

$$\begin{aligned} \underline{\mathbf{L}}_k^i &= \begin{pmatrix} (g_h^k, \partial_1 v_i)_{0,\omega} - (\mathbf{f}_{z,h}^k, \partial_2 v_i)_{0,\omega} + \frac{k\pi}{L} (\mathbf{f}_{y,h}^k, v_i)_{0,\omega} \\ (g_h^k, \partial_2 v_i)_{0,\omega} + (\mathbf{f}_{z,h}^k, \partial_1 v_i)_{0,\omega} - \frac{k\pi}{L} (\mathbf{f}_{x,h}^k, v_i)_{0,\omega} \end{pmatrix} \forall i \in \mathbf{I}_\omega, \\ &= \begin{pmatrix} (g_h^k, \partial_\tau v_i)_{0,\omega} - (\mathbf{f}_{z,h}^k, \partial_\nu v_i)_{0,\omega} \\ (g_h^k, \partial_\nu v_i)_{0,\omega} + (\mathbf{f}_{z,h}^k, \partial_\tau v_i)_{0,\omega} - \frac{k\pi}{L} (\mathbf{f}_h^k \cdot \boldsymbol{\tau}_i, v_i)_{0,\omega} \end{pmatrix} \forall i \in \mathbf{I}_a, \\ &= \begin{pmatrix} 0 \\ 0 \end{pmatrix} \forall i \in \mathbf{I}_c. \end{aligned}$$

Soit $\mathbb{X}_h^k \in \mathbb{R}^{N_{ar} \times N_{ar}}$ la matrice des produits : $\mathcal{A}_0^k(\mathbf{x}_{l,h}^S, \mathbf{x}_{m,h}^S)$. Pour le calcul des coefficients diagonaux, le calcul de $(\mathbf{x}_{l,h}^S, \mathbf{x}_{l,h}^S)_{0,\omega}$ se décompose de la même façon que pour le calcul des $\beta^{l,l}$, expliqué au paragraphe 4.4.3. On a : $\int_{\omega_{R_l}} |\mathbf{x}_l^P|^2 d\omega = \pi \alpha_l R_l^{2\alpha_l}$.

Soit $\underline{\lambda}^k \in \mathbb{R}^{N_{ar}}$, le vecteur contenant les approximations des $\lambda^{k,j}$.

Soit $\widehat{\underline{\mathbf{E}}}^k \in \mathbb{R}^{2N}$ la représentation de $\widehat{\mathbf{E}}^k$ dans la base B (4.40).

Les équations (8.26)-(8.27) se mettent alors sous la forme discrétisée suivante :

Pout tout $k \in \mathbb{N}$, trouver $\underline{\mathbf{E}}_k \in (\mathbb{R}^2)^N$ et $\mathbf{c}_{h,k} \in \mathbb{R}^{N_{ar}}$ tels que :

$$\begin{aligned} \mathbb{A}_0^k \widehat{\underline{\mathbf{E}}}^k + \mathbb{L}_k^T \mathbf{c}_h^k &= \underline{\mathbf{L}}_k - \mathbb{A}_r^k \underline{\mathbf{e}}^k, \\ \mathbb{L}_k \widehat{\underline{\mathbf{E}}}^k + \pi \mathbb{X}_h^k \mathbf{c}_h^k &= \pi \underline{\lambda}^k, \end{aligned} \tag{10.15}$$

Soit $\mathbb{A}_0'^k = \begin{pmatrix} \mathbb{A}_0^k & \mathbb{L}_k^T \\ \mathbb{L}_k & \pi \mathbb{X}_h^k \end{pmatrix} \in \mathbb{R}^{(2N+N_{cr}) \times (2N+N_{cr})}$, et $\underline{\mathbf{L}}'^k = \begin{pmatrix} \underline{\mathbf{L}}_k - \mathbb{A}_r^k \underline{\mathbf{e}}^k \\ \pi \underline{\lambda}^k \end{pmatrix} \in \mathbb{R}^{(2N+N_{cr})}$. Posons : $\underline{\mathbf{E}}'^k = \begin{pmatrix} \widehat{\underline{\mathbf{E}}}^k \\ \mathbf{c}_h^k \end{pmatrix} \in \mathbb{R}^{(2N+N_{cr})}$. Les équations (10.15) se mettent sous la forme suivante :

$$\mathbb{A}_0'^k \underline{\mathbf{E}}'^k = \underline{\mathbf{L}}'^k.$$

L'algorithme de résolution de (10.15) est le suivant :

- Résoudre $\mathbb{A}_0^k \underline{\mathbf{E}}_0 = \underline{\mathbf{L}}_k - \mathbb{A}_r^k \underline{\mathbf{e}}^k$;
- Résoudre $(\mathbb{L}_k (\mathbb{A}_0^k)^{-1} \mathbb{L}_k^T - \pi \mathbb{X}_h^k) \mathbf{c}_h^k = \mathbb{L}_k \underline{\mathbf{E}}_0 - \pi \underline{\lambda}^k$;
- Résoudre $\mathbb{A}_0^k \widehat{\underline{\mathbf{E}}}^k = \underline{\mathbf{L}}_k - \mathbb{A}_r^k \underline{\mathbf{e}}^k - \mathbb{L}_k^T \mathbf{c}_h^k := \mathbb{A}_0^k \underline{\mathbf{E}}_0 - \mathbb{L}_k^T \mathbf{c}_h^k$.

10.6.2 Champ électrique transverse : CL presque essentielles

Dans ce paragraphe, on s'intéresse à la discrétisation de la formulation variationnelle (8.30)-(8.31). Dans un premier temps, nous allons déterminer une base de $\mathbf{X}_{\mathbf{E},k}^{A,R}$ l'espace discrétisé de

$\mathbf{X}_{\mathbf{E}}^{A,R}$ par les éléments finis de Lagrange P_k :

$$\mathbf{X}_{\mathbf{E},k}^{A,R} = \{ \mathbf{v} \in \mathbf{X}_k \mid \mathbf{v} \cdot \boldsymbol{\tau}|_{\gamma_C} = 0 \text{ sur } \gamma_C \}. \quad (10.16)$$

Par construction, $\mathbf{X}_{\mathbf{E},k}^{A,R} \in \mathbf{H}^1(\Omega)$ et est conforme dans $\mathbf{X}_{\mathbf{E}}^{A,R}$. On suppose pour simplifier les notations que γ_A ne contient pas de coin (en pratique, on choisit pour γ_A un plan).

Les coins appartenant à $\overline{\gamma_C} \cap \overline{\gamma_A}$ sont maintenant considérés comme des points des arêtes de γ_C . On suppose qu'il y en a N_{AC} , indicés par I_{AC} . Il reste donc $N_c - N_{AC}$ coins et on a $N_a + N_{AC}$ points d'arêtes. On restreint l'ensemble d'indices I_c à $I_{c'} = I_c \setminus I_{AC}$, et on complète alors l'ensemble des indices I_a , en $I_a \cup I_{AC}$. Soit I_A l'ensemble des indices des N_A points du bord appartenant à γ_A . Pour ces points, on ne procédera pas à la pseudo-élimination. On complète alors l'ensemble des indices I_ω , en $I_o = I_\omega \cup I_A$, et on restreint l'ensemble des indices $I_a \cup I_{AC}$ à $I_{a'} := (I_a \setminus I_A) \cup I_{AC}$. On en déduit que $\mathbf{X}_{\mathbf{E},k}^{A,R}$ est de dimension $2N - N_{a'} - 2N_{c'}$.

Soit B_A la base de $\mathbf{X}_{\mathbf{E},k}^{A,R}$:

$$B_A := (\mathbf{v}_i, \alpha)_{i \in I_o, \alpha \in \{1,2\}} \cup (v_i, \boldsymbol{\nu}_i)_{i \in I_{a'}}. \quad (10.17)$$

On reprend la structure par bloc détaillée dans le paragraphe 4.7.1 de la partie II, en remplaçant I_ω par I_o , I_a par $I_{a'}$ et I_c par $I_{c'}$. On reprend la construction de la matrice $\mathbb{A}_0$. Considérons :

$$\mathbb{A}_A = \begin{pmatrix} \mathbb{A}_{o,o} & \mathbb{A}_{o,\nu} & 0 \\ \mathbb{A}_{\nu,o} & \mathbb{A}_{\nu,\nu} & 0 \\ 0 & 0 & \mathbb{I}_{c',c'} \end{pmatrix}$$

avec :

- $\mathbb{A}_{o,o} \in (\mathbb{R}^{2 \times 2})^{N_o \times N_o}$, telle que : $\forall, j \in I_o$:

$$\mathbb{A}_{o,o}^{i,j} = \begin{pmatrix} (v_j \mathbf{e}_1, v_i \mathbf{e}_1)_{\mathbf{X}_{\mathbf{E}}^A} & (v_j \mathbf{e}_2, v_i \mathbf{e}_1)_{\mathbf{X}_{\mathbf{E}}^A} \\ (v_j \mathbf{e}_1, v_i \mathbf{e}_2)_{\mathbf{X}_{\mathbf{E}}^A} & (v_j \mathbf{e}_2, v_i \mathbf{e}_2)_{\mathbf{X}_{\mathbf{E}}^A} \end{pmatrix}.$$

- $\mathbb{A}_{\nu,\nu} \in (\mathbb{R}^{2 \times 2})^{N_{f'} \times N_{f'}}$, telle que : $\forall, j \in I_{f'}$:

$$\mathbb{A}_{\nu,\nu}^{i,j} = \begin{pmatrix} (v_j \boldsymbol{\tau}_j, v_i \boldsymbol{\tau}_i)_{\mathbf{X}_{\mathbf{E}}^A} & (v_j \boldsymbol{\nu}_j, v_i \boldsymbol{\tau}_i)_{\mathbf{X}_{\mathbf{E}}^A} \\ (v_j \boldsymbol{\tau}_j, v_i \boldsymbol{\nu}_i)_{\mathbf{X}_{\mathbf{E}}^A} & (v_j \boldsymbol{\nu}_j, v_i \boldsymbol{\nu}_i)_{\mathbf{X}_{\mathbf{E}}^A} \end{pmatrix}.$$

- $\mathbb{A}_{o,\nu} \in (\mathbb{R}^{2 \times 2})^{N_o \times N_{a'}}$, et $\mathbb{A}_{\nu,o} \in (\mathbb{R}^{2 \times 2})^{N_{a'} \times N_o}$ telles que :

$\forall i \in I_o, \forall j \in I_{a'}$,

$$\mathbb{A}_{o,\nu}^{i,j} = \begin{pmatrix} 0 & (v_j \boldsymbol{\nu}_j, v_i \mathbf{e}_1)_{\mathbf{X}_{\mathbf{E}}^A} \\ 0 & (v_j \boldsymbol{\nu}_j, v_i \mathbf{e}_2)_{\mathbf{X}_{\mathbf{E}}^A} \end{pmatrix} \text{ et } \mathbb{A}_{\nu,o}^{j,i} = \begin{pmatrix} 0 & 0 \\ (v_i \mathbf{e}_1, v_j \boldsymbol{\nu}_j)_{\mathbf{X}_{\mathbf{E}}^A} & (v_i \mathbf{e}_2, v_j \boldsymbol{\nu}_j)_{\mathbf{X}_{\mathbf{E}}^A} \end{pmatrix}.$$

$\mathbb{I}_{c',c'}$ étant la matrice identité de $(\mathbb{R}^{2 \times 2})^{N_{c'} \times N_{c'}}$.

Soit $\mathbb{M}_A$ la matrice construite de la même façon, en remplaçant le produit scalaire $(\cdot, \cdot)_{\mathbf{x}_{\mathbf{E}}^A}$ par

le produit scalaire $L^2(\omega)$. La matrice $\mathbb{A}_A^k := \mathbb{A}_A + \frac{k^2 \pi^2}{L^2} \mathbb{M}_A$ représente alors le produit scalaire $\mathcal{A}_A^k(\cdot, \cdot)$ dans la base B_A .

Soit $\mathbb{L}_{A,k} \in (\mathbb{R}^2)^{N_{ar} \times N}$, la matrice représentant $\mathcal{A}_A^k(\cdot, \mathbf{x}_{l,h}^S)$ dans la base B_A . $\mathbb{L}_{A,k}$ est telle que : $\forall l \in \{1, \dots, N_{ar}\}$,

$$\begin{aligned} \mathbb{L}_{A,k}^{l,i} &= \left(\mathcal{A}_A^k(v_i \mathbf{e}_1, \mathbf{x}_{l,h}^S) \quad \mathcal{A}_A^k(v_i \mathbf{e}_2, \mathbf{x}_{l,h}^S) \right), \forall i \in \mathbb{I}_0, \\ &= \left(0 \quad \mathcal{A}_A^k(v_i \boldsymbol{\nu}_j, \mathbf{x}_{l,h}^S) \right), \forall i \in \mathbb{I}_{a'}, \\ &= \left(0 \quad 0 \right), \forall i \in \mathbb{I}_{c'}. \end{aligned}$$

Soit $\underline{\mathbb{L}}_{A,k} \in (\mathbb{R}^2)^N$, le vecteur représentant les approximations de $\mathcal{L}_A^k(\cdot)$ dans la base B_A , $\mathcal{L}_A^k$ est tel que :

$$\begin{aligned} \underline{\mathbb{L}}_{A,k} &= \begin{pmatrix} (g_h^k, \partial_1 v_i)_{0,\omega} - (f_h^k, \partial_2 v_i)_{0,\omega} \\ (g_h^k, \partial_2 v_i)_{0,\omega} + (f_h^k, \partial_1 v_i)_{0,\omega} \end{pmatrix}, \forall i \in \mathbb{I}_\omega \cup \mathbb{I}_{a'}, \\ &= \begin{pmatrix} (g_h^k, \partial_1 v_i)_{0,\omega} - (f_h^k, \partial_2 v_i)_{0,\omega} + \int_{\gamma_A} \mathbf{e}_h^k \cdot \boldsymbol{\tau} v_i \tau_1 d\sigma - \frac{k\pi}{L} \int_{\gamma_A} \mathbf{e}_{z,h}^k v_i \tau_2 d\sigma \\ (g_h^k, \partial_2 v_i)_{0,\omega} + (f_h^k, \partial_1 v_i)_{0,\omega} + \int_{\gamma_A} \mathbf{e}_h^k \cdot \boldsymbol{\tau} v_i \tau_2 d\sigma + \frac{k\pi}{L} \int_{\gamma_A} \mathbf{e}_{z,h}^k v_i \tau_1 d\sigma \end{pmatrix}, \forall i \in \mathbb{I}_A, \\ &= \begin{pmatrix} 0 \\ 0 \end{pmatrix}, \forall i \in \mathbb{I}_{c'}. \end{aligned}$$

Soit $\mathbb{X}_{A,h}^k \in (\mathbb{R}^{1 \times 2})^{N_{ar} \times N}$, telle que : $\forall l, k \in \{1, \dots, N_{ar}\}$,

$$\pi(\mathbb{X}_{A,h}^k)^{l,m} = \mathcal{A}_A^k(\mathbf{x}_{l,h}^S, \mathbf{x}_{m,h}^S).$$

Cette matrice se calcule de la même façon que $\mathbb{X}_h$.

Soit $\underline{\lambda}_A^k \in \mathbb{R}^{N_{ar}}$, le vecteur contenant les approximations des $\lambda_A^{k,j}$.

Soit $\widehat{\underline{\mathbb{E}}}_A^k \in \mathbb{R}^{2N}$ la représentation de $\widehat{\mathbf{E}}^k$ dans la base B_A .

Les équations (8.30)-(8.30) se mettent alors sous la forme discrétisée suivante :

Pout tout $k \in \mathbb{N}$, trouver $\widehat{\underline{\mathbb{E}}}_A^k \in (\mathbb{R}^2)^N$ et $\mathbf{c}_{A,k} \in \mathbb{R}^N$ tels que :

$$\boxed{\begin{aligned} \mathbb{A}_A^k \widehat{\underline{\mathbb{E}}}_A^k + \mathbb{L}_{A,k}^T \mathbf{c}_{A,k} &= \underline{\mathbb{L}}_A^k. \\ \mathbb{L}_{A,k} \widehat{\underline{\mathbb{E}}}_A^k + \mathbb{X}_{A,h}^k \mathbf{c}_{A,k} &= \pi \underline{\lambda}_A^k. \end{aligned}} \quad (10.18)$$

Soit $\mathbb{A}'_A = \begin{pmatrix} \mathbb{A}_A^k & \mathbb{L}_{A,k}^T \\ \mathbb{L}_{A,k} & \pi \mathbb{X}_{A,h}^k \end{pmatrix} \in \mathbb{R}^{(2N+N_{cr}) \times (2N+N_{cr})}$, et $\underline{\mathbb{L}}'_A = \begin{pmatrix} \underline{\mathbb{L}}_A^k - \mathbb{A}_A^k \underline{\mathbb{E}}^k \\ \pi \underline{\lambda}_A^k \end{pmatrix} \in \mathbb{R}^{(2N+N_{cr})}$.

Posons : $\underline{\mathbf{E}}_A^{'k} = \begin{pmatrix} \widehat{\underline{\mathbf{E}}}_A^k \\ \mathbf{c}_{A,h}^k \end{pmatrix} \in \mathbb{R}^{(2N+N_{cr})}$. Les équations (10.15) se mettent sous la forme suivante :

$$\mathbb{A}_A^{'k} \underline{\mathbf{E}}_A^{'k} = \underline{\mathbf{L}}_A^{'k}.$$

L'algorithme de résolution de (10.18) est le suivant :

- Résoudre $\mathbb{A}_A^k \underline{\mathbf{E}}_A = \underline{\mathbf{L}}_{A,k}$;
- Résoudre $(\mathbb{L}_{A,k} (\mathbb{A}_A^k)^{-1} \mathbb{L}_{A,k}^T - \pi \mathbb{X}_{A,h}^k) \mathbf{c}_{A,h}^k = \mathbb{L}_{A,k} \underline{\mathbf{E}}_A - \pi \underline{\lambda}_A^k$;
- Résoudre $\mathbb{A}_A^k \widehat{\underline{\mathbf{E}}}_A^k = \underline{\mathbf{L}}_{A,k} - \mathbb{L}_{A,k}^T \mathbf{c}_{A,h}^k := \mathbb{A}_A^k \underline{\mathbf{E}}_A - \mathbb{L}_{A,k}^T \mathbf{c}_{A,h}^k$.

10.6.3 Champ électrique longitudinal

Qu'on choisisse de discrétiser le problème prismatique dans $\mathcal{X}_\varepsilon^0$ ou dans $\mathcal{X}_\varepsilon^A$, le champ électrique longitudinal satisfait la formulation variationnelle (8.28). Nous allons discrétiser cette formulation dans l'espace V_k^0 (équation (4.9), partie II, paragraphe 4.3.1).

On introduit les matrices et les vecteurs suivants :

Soit $\mathbb{K}_D^k \in \mathbb{R}^{N \times N}$ la matrice représentant le produit scalaire $a_D^k(\cdot, \cdot)$ dans la base discrète de $V_k^0(\omega)$, telle que :

$$\begin{aligned} (\mathbb{K}_D^k)^{i,j} &= a_D^k(v_i, v_j), \text{ si } i \text{ et } j \in I_\omega, \\ &= \delta_{ij}, \text{ si } i \text{ ou } j \in I_{\partial\omega} \end{aligned}$$

Soit $\mathbb{K}_D^{P,k} \in \mathbb{R}^{N \times N}$ telle que :

$$\mathbb{K}_D^{P,k} = \begin{pmatrix} 0 & \mathbb{K}_{\omega, \partial\omega}^k \\ 0 & -\mathbb{I}_{\partial\omega} \end{pmatrix},$$

où $\mathbb{I}_{\partial\omega}$ est la matrice identité de $\mathbb{R}^{N_{\partial\omega} \times N_{\partial\omega}}$; et $\mathbb{K}_{\omega, \partial\omega}^k \in \mathbb{R}^{N_\omega \times N_{\partial\omega}}$ est telle que :

$$\mathbb{K}_{\omega, \partial\omega}^{i,j} = a_D^k(v_j, v_i) \text{ si } i \in I_\omega, \text{ et } j \in I_{\partial\omega}.$$

Soit $\underline{\mathbf{l}}_k \in \mathbb{R}^N$ le vecteur tel que :

$$\begin{aligned} \underline{\mathbf{l}}_k^i &= (f_{h,x}^k, \partial_1 v_i)_{0,\omega} - (f_{h,y}^k, \partial_2 v_i)_{0,\omega} - \frac{k\pi}{L} (g_h, \partial_1 v_i)_{0,\omega}, \forall i \in I_\omega, \\ &= 0 \text{ sinon.} \end{aligned}$$

Soit $\underline{\mathbf{e}}_z^k$ la représentation de $\mathbf{e}_{h,z}^k$ dans la base cartésienne.

L'équation (8.28) se met sous la forme discrétisée suivante :

Pout tout $k \in \mathbb{N}$, trouver $\underline{\mathbf{E}}_z^k \in \mathbb{R}^N$ tel que :

$$\boxed{\mathbb{K}_D^k \underline{\mathbf{E}}_z^k = \underline{\mathbf{l}}_k - \mathbb{K}_D^{P,k} \underline{\mathbf{e}}_z^k.} \quad (10.19)$$

Chapitre 11

Le problème statique 3D mixte discret

11.1 L'élément fini mixte de Taylor-Hood P_2 - P_1

Comme pour le cas bidimensionnel, on discrétise le problèmes mixte par les éléments finis de Taylor-Hood P_2 - P_1 (section 5.1, p. 123). Le champ électrique est approché en P_2 et le multiplicatuer de Lagrange en P_1 sur un même maillage. L'espace d'approximation du champ électrique est alors :

$$\mathcal{X}_2 = \{ \mathbf{v} \in C^0(\overline{\Omega})^2 : \mathbf{v}|_{T_l} \in P_2^3(T_l), \forall l \in \{1, \dots, L\} \},$$

et celui du multiplicateur de Lagrange est alors :

$$V_1 = \{ u \in C^0(\overline{\Omega}) : \mathbf{v}|_{T_l} \in P_1(T_l), \forall l \in \{1, \dots, L\} \}.$$

On notera N_1 le nombre de degrés de liberté P_1 , c'est-à-dire le nombre de sommets de la tétraédrisation du maillage; et N le nombre de degrés de liberté P_2 , c'est-à-dire le nombre de sommets plus le nombre d'arêtes de la triangulation du maillage.

Considérons $S_{i_1}, i_1 \in \{1, \dots, N_1\}$ les sommets des triangles T_l . Soit N_Ω^1 le nombre de sommets intérieurs à Ω , et $N_{\partial\Omega}^1$ le nombre de sommets du bord $\partial\Omega$. On ordonne ces points P_1 ainsi : $\forall i_1 \in I_\Omega^1, S_{i_1} \in \Omega$, et $\forall i_1 \in I_{\partial\Omega}^1, S_{i_1} \in \partial\Omega$, où on a défini les ensembles d'indices suivants :

$$I_\Omega^1 = \{1, \dots, N_\Omega^1\}, I_{\partial\Omega} = \{N_\Omega^1 + 1, \dots, N_\Omega^1 + N_{\partial\Omega}^1\} \text{ et } I_1 = I_\Omega^1 \cup I_{\partial\Omega}.$$

Soient $u_{i_1}, i_1 \in I_1$ les fonctions de base P_1 , qui génèrent V_1 .

Le champ électrique discret est recherché dans l'espace $(V_2)^3$. On ordonnera les points de discrétisation P_2 (sommets et arêtes des tétraèdres) $M_i, i \in \{1, \dots, N\}$. Soit $i \in I$ tel que M_i soit un sommet de la tétraédrisation. Alors il existe $i_1 \in I_1$ tel que $S_{i_1} = M_i$. On considérera $v_i, i \in I$ les fonctions de base P_2 qui génèrent V_2 . On notera $\mathcal{E}_h$ et p_h et les approximations du champ électrique et du multiplicateur de Lagrange dans $\mathcal{X}_\mathcal{E} \times L^2(\Omega)$ ou $\mathcal{X}_\mathcal{E}^0 \times L^2(\Omega)$, et $\mathcal{E}_{\gamma,h}$ et $p_{\gamma,h}$ et les approximations dans $\mathcal{X}_{\mathcal{E},\gamma}^0 \times L_\gamma^2(\Omega)$. Le multiplicateur de Lagrange discrétisé est décomposés ainsi :

$$p_h = \sum_{i_1 \in I_1} p_h(S_{i_1}) u_{i_1}, p_{\gamma,h} = \sum_{i_1 \in I_1} p_{\gamma,h}(S_{i_1}) u_{i_1}.$$

On notera $\underline{p}$ et $\underline{p}_\gamma \in \mathbb{R}^{N_1}$ les vecteurs associés à p_h et p_h : $\underline{p} = (p_h(S_1), \dots, p_h(S_{N_1}))^T$ et $\underline{p}_\gamma = (p_{\gamma,h}(S_1), \dots, p_{\gamma,h}(S_{N_1}))^T$.

11.2 Méthode avec CL naturelles mixte : discrétisation

La preuve de la condition inf-sup discrète uniforme de la formulation mixte (9.3) pour l'élément fini de Taylor-Hood P_2 -iso- P_1 a été faite par P. Ciarlet, Jr. et V. Girault dans [39]. De même, pour le couple d'espace $(\mathcal{X}_2, V_1)$, la condition inf-sup discrète est uniforme.

Proposition 11.1 *Il existe une constante $\kappa_{\mathcal{E}}^* > 0$, indépendante de h telle que :*

$$\inf_{u_h \in V_1} \sup_{\mathcal{F}_h \in \mathcal{X}_2} \frac{\mathcal{B}(\mathcal{F}_h, u_h)}{\|\mathcal{F}_h\|_{\mathcal{X}_{\mathcal{E}}} \|u_h\|_0} \geq \kappa_{\mathcal{E}}^*.$$

La formulation variationnelle (9.3) discrétisée par les éléments finis de Taylor-Hood P_2 - P_1 s'écrit de la façon suivante :

Trouver le couple $(p_h, \mathcal{E}_h) \in V_1 \times \mathcal{X}_2$ tel que :

$$\begin{aligned} \forall i \in I, \forall \alpha \in \{1, 3\} \quad (\mathcal{E}_h, v_i \mathbf{e}_{\alpha})_{\mathcal{X}_{\mathcal{E}}} + \mathcal{B}_{\mathcal{E}}(v_i \mathbf{e}_{\alpha}, p_h) &= \mathcal{L}_{\mathcal{E}}(v_i \mathbf{e}_{\alpha}), \\ \forall i_1 \in I_1 \quad \mathcal{B}_{\mathcal{E}}(\mathcal{E}_h, u_{i_1}) &= \mathcal{G}_{\mathcal{E}}(u_{i_1}). \end{aligned} \quad (11.1)$$

En décomposant p_h dans (11.1), on obtient :

$$\begin{aligned} \forall i \in I, \forall \alpha \in \{1, 3\} \quad (\mathcal{E}_h, v_i \mathbf{e}_{\alpha})_{\mathcal{X}_{\mathcal{E}}} + \sum_{j_1 \in I_1} p_h(S_{j_1}) \mathcal{B}_{\mathcal{E}}(v_i \mathbf{e}_{\alpha}, u_{i_1}) &= \mathcal{L}_{\mathcal{E}}(v_i \mathbf{e}_{\alpha}), \\ \forall i_1 \in I_1 \quad \mathcal{B}_{\mathcal{E}}(\mathcal{E}_h, u_{i_1}) &= \mathcal{G}_{\mathcal{E}}(u_{i_1}). \end{aligned}$$

Soit $\underline{\mathbf{E}} \in (\mathbb{R}^3)^N$, la représentation de $\mathcal{E}_h$ dans $\text{vect}(v_i \mathbf{e}_{\alpha})_{i \in I, \alpha \in \{1,2,3\}}$. Nous allons écrire les équations ci-dessus de façon matricielle. Pour cela, nous devons construire les matrices associées à $\mathcal{B}_{\mathcal{E}}$ et $\mathcal{G}_{\mathcal{E}}$.

- Soit $\mathbb{C}_{\mathcal{E}} \in (\mathbb{R}^{1 \times 3})^{N_1 \times N}$ la matrice composée des sous blocs $\mathbb{C}_{\mathcal{E}}^{i_1, j} \in \mathbb{R}^3$ tels que :

$$\forall i_1 \in I_1, \forall j \in I, \mathbb{C}_{\mathcal{E}}^{i_1, j} = \left((\partial_1 v_j, u_{i_1})_0 \quad (\partial_2 v_j, u_{i_1}) \quad (\partial_3 v_j, u_{i_1}) \right).$$

- Soit $\underline{\mathbf{G}} \in \mathbb{R}^{N_1}$ le vecteur tel que :

$$\forall i_1 \in I_1, \underline{\mathbf{G}}^{i_1} = \mathcal{G}_{\mathcal{E}}(u_{i_1}).$$

Les équations (11.1) se mettent sous la forme :

$$\boxed{\begin{aligned} \mathbb{A}_{\mathcal{E}} \underline{\mathbf{E}} + \mathbb{C}_{\mathcal{E}}^T \underline{\mathbf{p}} &= \underline{\mathbf{L}}, \\ \mathbb{C}_{\mathcal{E}} \underline{\mathbf{E}} &= \underline{\mathbf{G}}. \end{aligned}} \quad (11.2)$$

L'algorithme de résolution de (11.2) est le suivant :

- Résoudre $\mathbb{A}_{\mathcal{E}} \underline{\mathbf{E}}_0 = \underline{\mathbf{L}}$;
- Résoudre $(\mathbb{C}_{\mathcal{E}} \mathbb{A}_{\mathcal{E}}^{-1} \mathbb{C}_{\mathcal{E}}^T) \underline{\mathbf{p}} = \mathbb{C}_{\mathcal{E}} \underline{\mathbf{E}}_0 - \underline{\mathbf{G}}$;
- Résoudre $\mathbb{A}_{\mathcal{E}} \underline{\mathbf{E}} = \underline{\mathbf{L}} - \mathbb{C}_{\mathcal{E}}^T \underline{\mathbf{p}}$.

Les calculs de $\mathbb{C}_{\mathcal{E}}$ et $\underline{\mathbf{G}}$ sont similaires aux calculs de la section 5.3 (partie II).

11.3 Méthode avec CL essentielles mixte : discrétisation

$\mathcal{X}_{\varepsilon,2}^{0,R}$ est un sous-espace de $\mathcal{X}_2$, et $\forall (\mathcal{F}_h, u_h) \in \mathcal{X}_{\varepsilon,2}^{0,R} \times V_1$, $\mathcal{B}_0(\mathcal{F}_h, u_h) = \mathcal{B}_{\mathcal{E}}(\mathcal{F}_h, u_h)$, on a donc aussi une condition inf-sup discrète pour le couple d'espaces $(\mathcal{X}_{\varepsilon,2}^{0,R}, V_1)$:

Proposition 11.2 *Il existe une constante $\kappa_0^* > 0$ indépendante de h telle que :*

$$\inf_{u_h \in V_1} \sup_{\mathcal{F}_h \in \mathcal{X}_{\varepsilon,2}^{0,R}} \frac{\mathcal{B}_0(\mathcal{F}_h, u_h)}{\|\mathcal{F}_h\|_{\mathcal{X}_{\mathcal{E}}^0} \|u_h\|_0} \geq \kappa_0^*.$$

La formulation variationnelle (9.4) discrétisée par les éléments finis de Taylor-Hood P_2 - P_1 s'écrit de la façon suivante :

Trouver le couple $(p_h, \mathcal{E}_h) \in V_1 \times \mathcal{X}_2$ tel que :

$$\begin{aligned} \forall i \in I_{\Omega}, \forall \alpha \in \{1, 3\}, \quad & (\mathcal{E}_h, v_i \mathbf{e}_{\alpha})_{\mathcal{X}_{\mathcal{E}}^0} + \mathcal{B}_0(v_i \mathbf{e}_{\alpha}, p_h) = \mathcal{L}_0(v_i \mathbf{e}_{\alpha}), \\ \forall i \in I_f, \quad & (\mathcal{E}_h, v_i \boldsymbol{\nu}_i)_{\mathcal{X}_{\mathcal{E}}^0} + \mathcal{B}_0(v_i \boldsymbol{\nu}_i, p_h) = \mathcal{L}_0(v_i \boldsymbol{\nu}_i), \\ \forall i \in I_f, \quad & \mathcal{E}_h(M_i) \times \boldsymbol{\nu}_i = \mathbf{e}_h(M_i) \times \boldsymbol{\nu}_i, \\ \forall i \in I_{ac}, \quad & \mathcal{E}_h(M_i) = \mathbf{e}_h(M_i), \\ \forall i_1 \in I_1, \quad & \mathcal{B}_0(\mathcal{E}, u_{i_1}) = \mathcal{G}_{\mathcal{E}}(u_{i_1}). \end{aligned} \tag{11.3}$$

En décomposant p_h dans (11.3), on obtient alors :

$$\begin{aligned} \forall i \in I_{\Omega}, \forall \alpha \in \{1, 2, 3\}, \quad & (\mathcal{E}_h, v_i \mathbf{e}_{\alpha})_{\mathcal{X}_{\mathcal{E}}^0} + \sum_{i_1=1}^{N_1} p_h(S_{i_1}) \mathcal{B}_0(v_i \mathbf{e}_{\alpha}, u_{i_1}) = \mathcal{L}_0(v_i \mathbf{e}_{\alpha}), \\ \forall i \in I_f, \quad & (\mathcal{E}_h, v_i \boldsymbol{\nu}_i)_{\mathcal{X}_{\mathcal{E}}^0} + \sum_{i_1=1}^{N_1} p_h(S_{i_1}) \mathcal{B}_0(v_i \boldsymbol{\nu}_i, u_{i_1}) = \mathcal{L}_0(v_i \boldsymbol{\nu}_i), \\ \forall i \in I_f, \quad & \mathcal{E}_h(M_i) \times \boldsymbol{\nu}_i = \mathbf{e}_h(M_i) \times \boldsymbol{\nu}_i, \\ \forall i \in I_{ac}, \quad & \mathcal{E}_h(M_i) = \mathbf{e}_h(M_i), \\ \forall i_1 \in I_1, \quad & \sum_{j=1}^N \mathcal{B}_0(\mathcal{E}_h, u_{i_1}) = \mathcal{G}_{\mathcal{E}}(u_{i_1}). \end{aligned}$$

Soit $\underline{\mathbf{E}}^R \in (\mathbb{R}^3)^N$, la représentation de $\mathcal{E}_h$ dans B (10.12). Nous allons écrire les équations ci-dessus de façon matricielle. Pour cela, nous devons construire les matrices associées à $\mathcal{B}_0$.

• Soit $\mathbb{C}_0 \in (\mathbb{R}^{1 \times 3})^{N_1 \times 3N}$ la matrice composée des sous-blocs $\mathbb{C}_0^{i_1, j} \in \mathbb{R}^{1 \times 3}$ tels que :

$$\begin{aligned} \forall i_1 \in I_1, \forall j \in I_{\Omega}, \quad & \mathbb{C}_0^{i_1, j} = \mathbb{C}_{\mathcal{E}}^{i_1, j}, \\ \forall i_1 \in I_1, \forall j \in I_{ac}, \quad & \mathbb{C}_0^{i_1, j} = \begin{pmatrix} 0 & 0 & 0 \end{pmatrix}, \\ \forall i_1 \in I_1, \forall j \in I_f, \quad & \mathbb{C}_0^{i_1, j} = \begin{pmatrix} 0 & 0 & (\partial_{\nu_j} v_j, u_{i_1})_0 \end{pmatrix}. \end{aligned}$$

- Soit $\mathbb{C}_r$ la matrice composée des sous-blocs $\mathbb{C}_r^{i_1, j} \in \mathbb{R}^{1 \times 2}$ tels que :

$$\begin{aligned} \forall i_1 \in I_1, \forall j \in I_\omega, \quad \mathbb{C}_r^{i_1, j} &= \begin{pmatrix} 0 & 0 & 0 \end{pmatrix}, \\ \forall i_1 \in I_1, \forall j \in I_{ac}, \quad \mathbb{C}_r^{i_1, j} &= \mathbb{C}_{\mathcal{E}}^{i_1, j}, \\ \forall i_1 \in I_1, \forall j \in I_f, \quad \mathbb{C}_r^{i_1, j} &= \begin{pmatrix} (\partial_{\tau_j^1} v_j, u_{i_1})_0 & (\partial_{\tau_j^2} v_j, u_{i_1})_0 & 0 \end{pmatrix}. \end{aligned}$$

Il s'agit de résoudre le système matriciel suivant :

$$\boxed{\begin{aligned} \mathbb{A}_0 \underline{\mathbf{E}}^R + \mathbb{C}_0^T \underline{\mathbf{p}} &= \underline{\mathbf{L}}_0 - \mathbb{A}_r \underline{\mathbf{E}}^r, \\ \mathbb{C}_0 \underline{\mathbf{E}}^R &= \underline{\mathbf{G}} - \mathbb{C}_r \underline{\mathbf{E}}^r. \end{aligned}} \quad (11.4)$$

L'algorithme de résolution de (11.4) est le suivant :

- Résoudre $\mathbb{A}_0 \underline{\mathbf{E}}' = \underline{\mathbf{L}}_0$;
- Résoudre $(\mathbb{C} \mathbb{A}_0^{-1} \mathbb{C}^T) \underline{\mathbf{p}} = \mathbb{C}_0 \underline{\mathbf{E}}' - (\underline{\mathbf{G}} - \mathbb{C}_r \underline{\mathbf{E}}^r)$;
- Résoudre $\mathbb{A}_0 \underline{\mathbf{E}}^R = \underline{\mathbf{L}}_0 - \mathbb{A}_r \underline{\mathbf{E}}^r - \mathbb{C}_0^T \underline{\mathbf{p}}$.

Les calculs de $\mathbb{C}_0$ et $\mathbb{C}_r$ sont similaires à ceux de la section 5.4 (partie II).

11.4 Régularisation à poids mixte : discrétisation

Comme pour le cas bidimensionnel, nous n'avons pas obtenu de résultat théorique sur la condition inf-sup discrète dans $\mathcal{X}_{\mathcal{E}, \gamma}^0$. La formulation variationnelle (9.7) discrétisée par les éléments finis de Taylor-Hood P_2 - P_1 s'écrit de la façon suivante : Trouver le couple $(p_{\gamma, h}, \mathcal{E}_{\gamma, h}) \in V_1 \times \mathcal{X}_2$ tel que :

$$\begin{aligned} \forall i \in I_\Omega, \forall \alpha \in \{1, 3\} \quad & (\mathcal{E}_{\gamma, h}, v_i \mathbf{e}_\alpha)_{\mathcal{X}_{\mathcal{E}, \gamma}^0} + \mathcal{B}_\gamma(p_{\gamma, h}, u_{j_1}) = \mathcal{L}_\gamma(v_i \mathbf{e}_\alpha), \\ \forall i \in I_f, \quad & (\mathcal{E}_{\gamma, h}, v_i \boldsymbol{\nu}_i)_{\mathcal{X}_{\mathcal{E}, \gamma}^0} + \mathcal{B}_\gamma(v_i \boldsymbol{\nu}_i, p_{\gamma, h}) = \mathcal{L}_\gamma(v_i \boldsymbol{\nu}_i), \\ \forall i \in I_f, \quad & \mathcal{E}_{\gamma, h}(M_i) \times \boldsymbol{\nu}_i = \mathbf{e}_h(M_i) \times \boldsymbol{\nu}_i, \\ \forall i \in I_{ac}, \quad & \mathcal{E}_{\gamma, h}(M_i) = \mathbf{e}_h(M_i), \\ \forall i_1 \in I_1 \quad & \mathcal{B}_\gamma(\mathcal{E}_{\gamma, h}, u_{i_1}) = \mathcal{G}_\gamma(u_{i_1}). \end{aligned} \quad (11.5)$$

En décomposant $p_{\gamma,h}$ dans (11.5), on obtient alors :

$$\begin{aligned} \forall i \in I_{\Omega}, \forall \alpha \in \{1, 2, 3\}, \quad (\mathcal{E}_{\gamma,h}, v_i \mathbf{e}_{\alpha})_{\mathcal{X}_{\mathcal{E},\gamma}^0} + \sum_{i_1=1}^{N_1} p_{\gamma,h}(S_{i_1}) \mathcal{B}_{\gamma}(v_i \mathbf{e}_{\alpha}, u_{i_1}) &= \mathcal{L}_{\gamma}(v_i \mathbf{e}_{\alpha}), \\ \forall i \in I_f, \quad (\mathcal{E}_{\gamma,h}, v_i \boldsymbol{\nu}_i)_{\mathcal{X}_{\mathcal{E},\gamma}^0} + \sum_{i_1=1}^{N_1} p_{\gamma,h}(S_{i_1}) \mathcal{B}_{\gamma}(v_i \boldsymbol{\nu}_i, u_{i_1}) &= \mathcal{L}_{\gamma}(v_i \boldsymbol{\nu}_i), \\ \forall i \in I_f, \quad \mathcal{E}_{\gamma,h}(M_i) \times \boldsymbol{\nu}_i &= \mathbf{e}_h(M_i) \times \boldsymbol{\nu}_i, \\ \forall i \in I_{ac}, \quad \mathcal{E}_{\gamma,h}(M_i) &= \mathbf{e}_h(M_i). \\ \forall i_1 \in I_1, \quad \mathcal{B}_{\gamma}(\mathcal{E}_{\gamma,h}, u_{i_1}) &= \mathcal{G}_{\gamma}(u_{i_1}). \end{aligned}$$

Soit $\underline{\mathbf{E}}_{\gamma} \in (\mathbb{R}^3)^N$, la représentation de $\mathcal{E}_{\gamma,h}$ dans B (10.12). Nous allons écrire les équations ci-dessus de façon matricielle. Pour cela, nous devons construire les matrices associées à $\mathcal{B}_{\gamma}$ et $\mathcal{G}_{\gamma}$.

- Soit $\mathbb{C}_{\gamma} \in (\mathbb{R}^{1 \times 3})^{N_1 \times 3N}$ la matrice composée des sous-blocs $\mathbb{C}_{\gamma}^{i_1,j} \in \mathbb{R}^{1 \times 3}$ tels que :

$$\forall i_1 \in I_1, \forall j \in I_{\Omega} \quad \mathbb{C}_{\gamma}^{i_1,j} = \begin{pmatrix} (\partial_1 v_j, u_{i_1})_{0,\gamma} & (\partial_2 v_j, u_{i_1})_{0,\gamma} & (\partial_3 v_j, u_{i_1})_{0,\gamma} \end{pmatrix},$$

$$\forall i_1 \in I_1, \forall j \in I_{ac} \quad \mathbb{C}_{\gamma}^{i_1,j} = \begin{pmatrix} 0 & 0 & 0 \end{pmatrix},$$

$$\forall i_1 \in I_1, \forall j \in I_f \quad \mathbb{C}_{\gamma}^{i_1,j} = \begin{pmatrix} 0 & 0 & (\partial_{\nu_j} v_j, u_{i_1})_{0,\gamma} \end{pmatrix}.$$

- Soit $\mathbb{C}_{r,\gamma}$ la matrice composée des sous blocs $\mathbb{C}_{r,\gamma}^{i_1,j} \in \mathbb{R}^{1 \times 3}$ tels que :

$$\forall i_1 \in I_1, \forall j \in I_{\Omega} \quad \mathbb{C}_{r,\gamma}^{i_1,j} = \begin{pmatrix} 0 & 0 & 0 \end{pmatrix},$$

$$\forall i_1 \in I_1, \forall j \in I_{ac} \quad \mathbb{C}_{r,\gamma}^{i_1,j} = \begin{pmatrix} (\partial_1 v_j, u_{i_1})_{0,\gamma} & (\partial_2 v_j, u_{i_1})_{0,\gamma} & (\partial_3 v_j, u_{i_1})_{0,\gamma} \end{pmatrix},$$

$$\forall i_1 \in I_1, \forall j \in I_f \quad \mathbb{C}_{r,\gamma}^{i_1,j} = \begin{pmatrix} (\partial_{\tau_j^1} v_j, u_{i_1})_{0,\gamma} & (\partial_{\tau_j^2} v_j, u_{i_1})_{0,\gamma} & 0 \end{pmatrix}.$$

- Soit enfin $\underline{\mathbf{G}}_{\gamma} \in \mathbb{R}^{N_1}$ le vecteur tel que :

$$\forall i_1 \in I_1, \underline{\mathbf{G}}_{\gamma}^{i_1} = \mathcal{G}_{\gamma}(u_{i_1}) = (g, u_{i_1})_{0,\gamma}.$$

Les équations (11.5) se mettent sous la forme :

$$\boxed{\begin{aligned} \mathbb{A}_{\gamma} \underline{\mathbf{E}}_{\gamma} + \mathbb{C}_{\gamma}^T \underline{\mathbf{p}}_{\gamma} &= \underline{\mathbf{L}}_{\gamma} - \mathbb{A}_{r,\gamma} \underline{\mathbf{E}}^r, \\ \mathbb{C}_{\gamma} \underline{\mathbf{E}}_{\gamma} &= \underline{\mathbf{G}}_{\gamma} - \mathbb{C}_{r,\gamma} \underline{\mathbf{e}}. \end{aligned}} \quad (11.6)$$

L'algorithme de résolution de (11.6) est le suivant :

- Résoudre $\mathbb{A}_{\gamma} \underline{\mathbf{E}}_{\gamma,0} = \underline{\mathbf{L}}_{\gamma} - \mathbb{A}_{r,\gamma} \underline{\mathbf{E}}^r$;

- Résoudre $(\mathbb{C}_\gamma \mathbb{A}_\gamma^{-1} \mathbb{C}_\gamma^T) \underline{\mathbf{p}} = \mathbb{C}_\gamma \underline{\mathbf{E}}_{\gamma,0} - (\underline{\mathbf{G}}_\gamma - \mathbb{C}_{r,\gamma} \underline{\mathbf{E}}^r)$;

- Résoudre $\mathbb{A}_\gamma \underline{\mathbf{E}}_\gamma = \underline{\mathbf{L}}_\gamma - \mathbb{A}_{r,\gamma} \underline{\mathbf{e}} - \mathbb{C}_\gamma^T \underline{\mathbf{p}}$.

Les calculs de $\mathbb{C}_\gamma$ et $\mathbb{C}_{r,\gamma}$ sont similaires à ceux de la section 5.5 (partie II).

11.5 Cas prismatique mixte : discrétisation

Nous reprenons dans cette section les notations de la section 5.4 de la partie II. Ainsi, ici, N_1 (*resp.*, N) est le nombre de points de discrétisation P_1 (*resp.*, P_2) de ω , etc. La partie longitudinale du champ électrique E_z est discrétisé dans V_2^0 , avec :

$$V_2^0 = \{ u \in C^0(\bar{\omega}) : u|_{T_h} \in P_2(T_h), \forall T_h \in \mathcal{T}_h(\omega) \text{ et } u|_{\partial\omega} = 0 \}.$$

On introduit la matrice mixte suivante, pour la partie longitudinale du champ électrique E_z , correspondant au produit $(p^k, v)_{0,\omega}$ discrétisé dans $V_1 \times V_2^0$:

Soit $\mathbb{M}_k \in \mathbb{R}^{N_1 \times N}$ telle que :

$$\forall i_1 \in I_1, \forall j \in I_\omega, \quad \mathbb{M}_k^{i_1,j} = \frac{k\pi}{L} (u_{i_1}, v_j)_{0,\omega},$$

$$\forall i_1 \in I_1, \forall j \in I_{\partial\omega}, \quad \mathbb{M}_k^{i_1,j} = 0.$$

Soit $\mathbb{M}_{k,r} \in \mathbb{R}^{N_1 \times N}$ telle que :

$$\forall i_1 \in I_1, \forall j \in I_\omega, \quad \mathbb{M}_{k,r}^{i_1,j} = 0,$$

$$\forall i_1 \in I_1, \forall j \in I_{\partial\omega}, \quad \mathbb{M}_{k,r}^{i_1,j} = \frac{k\pi}{L} (u_{i_1}, v_j)_{0,\omega}.$$

Soit $\underline{\mathbf{G}}_k \in \mathbb{R}^{N_1}$ tel que : $\forall i_1 \in I_1, \underline{\mathbf{G}}_k^{i_1} = (g_h^k, u_{i_1})_{0,\omega}$.

Soit $\underline{\mathbf{p}}^k \in \mathbb{R}^{N_1}$ la représentation discrétisée de p_h , l'approximation de p dans V_1 .

Soit $\underline{\mathbf{E}}_z^k \in \mathbb{R}^N$ la représentation discrétisée de $E_{z,h}$, l'approximation de E_z dans V_2 .

11.5.1 CL essentielles

Soit $\hat{\underline{\mathbf{E}}}^k \in (\mathbb{R}^2)^N$ la représentation discrétisée de $\hat{\mathbf{E}}_h$, l'approximation de $\hat{\mathbf{E}}$ dans la base B (4.40) de $\mathbf{X}_{\mathbf{E},2}$. La discrétisation du système de formulations variationnelles (9.8)-(9.11) se met sous la forme :

Pour tout $k \in \mathbb{N}$, trouver $(\hat{\underline{\mathbf{E}}}^k, \underline{\mathbf{c}}_h^k, \underline{\mathbf{E}}_z^k, \underline{\mathbf{p}}_k) \in (\mathbb{R}^2)^N \times \mathbb{R}^{N_{ar}} \times \mathbb{R}^N \times \mathbb{R}^{N_1}$ tel que :

$$\mathbb{A}_0^k \hat{\underline{\mathbf{E}}}^k + \mathbb{L}_k^T \underline{\mathbf{c}}_h^k + \mathbb{C}_0^T \underline{\mathbf{p}} = \underline{\mathbf{L}}_k - \mathbb{A}_r^k \underline{\mathbf{e}}^k,$$

$$\mathbb{L}_k \hat{\underline{\mathbf{E}}}^k + \pi \mathbb{X}_h^k \underline{\mathbf{c}}_h^k + \mathbb{S}_D^T \underline{\mathbf{p}} = \pi \underline{\lambda}^k,$$

$$\mathbb{C}_0 \hat{\underline{\mathbf{E}}}^k + \mathbb{S}_D \underline{\mathbf{c}}_h^k - \mathbb{M}_k \underline{\mathbf{E}}_z^k = \underline{\mathbf{G}}^k - \mathbb{C}_{r,\gamma} \underline{\mathbf{e}}^k + \mathbb{M}_{k,r} \underline{\mathbf{e}}_z^k,$$

$$\mathbb{K}_D^k \underline{\mathbf{E}}_z^k - \mathbb{M}_k^T \underline{\mathbf{p}} = \underline{\mathbf{l}}_k - \mathbb{K}_D^{P,k} \underline{\mathbf{e}}_z^k,$$

Posons :

$$\underline{l}'_k = \underline{l}_k - \mathbb{K}_D^{P,k} \underline{e}_z^k, \text{ et } \underline{G}'^k = \underline{G}^k - \mathbb{C}_r \underline{e}^k + \mathbb{M}_{r,k} \underline{e}_z^k.$$

Les équations ci-dessus se mettent aussi sous la forme condensée suivante (en reprenant les notations du paragraphe 10.6.1 de la partie II) :

$$\boxed{\begin{aligned} \begin{pmatrix} \mathbb{A}'^k & 0 \\ 0 & \mathbb{K}_D^k \end{pmatrix} \begin{pmatrix} \underline{E}'^k \\ \underline{E}_z^k \end{pmatrix} + \begin{pmatrix} \mathbb{C}'^T \\ -\mathbb{M}_k^T \end{pmatrix} \underline{p}^k &= \begin{pmatrix} \underline{L}'_k \\ \underline{l}'_k \end{pmatrix}, \\ \begin{pmatrix} \mathbb{C}' & -\mathbb{M}_k \end{pmatrix} \begin{pmatrix} \underline{E}'^k \\ \underline{E}_z^k \end{pmatrix} &= \underline{G}'^k. \end{aligned}} \quad (11.7)$$

L'algorithme de résolution de (11.7) est le suivant :

- Résoudre $\mathbb{A}'^k \underline{E}_0 = \underline{L}'_k$ (voir l'algorithme de résolution de (10.15)) et $\mathbb{K}_D^k \underline{E}_{0,z} = \underline{l}'_k$;
- Résoudre $\begin{pmatrix} \mathbb{C}' & -\mathbb{M}_k \end{pmatrix} \begin{pmatrix} (\mathbb{A}'^k)^{-1} & 0 \\ 0 & (\mathbb{K}_D^k)^{-1} \end{pmatrix} \begin{pmatrix} \mathbb{C}'^T \\ -\mathbb{M}_k^T \end{pmatrix} \underline{p} = \underline{G}'^k$;
- Résoudre $\mathbb{A}'^k \underline{E}'^k = \underline{L}'^k - \mathbb{C}'^T \underline{p}$, et $\mathbb{K}_D^k \underline{E}_z^k = \underline{l}'_k + \mathbb{M}_k^T \underline{p}$.

Pour la deuxième étape, calcule une approximation de $\underline{p}$ par l'algorithme du gradient conjugué (section 15.4, partie IV). Notons qu'à chaque étape, on devra utiliser l'algorithme de résolution de (10.15).

11.5.2 CL presque essentielles

Soient $\mathbb{C}_A \in (\mathbb{R}^{1 \times 2})^{N_1 \times N}$ la matrice mixte suivante :

$$\begin{aligned} \forall i_1 \in I_1, \forall j \in I_o \quad \mathbb{C}_A^{i_1,j} &= \mathbb{C}_{\mathbf{E}}^{i_1,j}, \\ \forall i_1 \in I_1, \forall j \in I_{c'} \quad \mathbb{C}_A^{i_1,j} &= \begin{pmatrix} 0 & 0 \end{pmatrix}, \\ \forall i_1 \in I_1, \forall j \in I_{a'} \quad \mathbb{C}_A^{i_1,j} &= \begin{pmatrix} 0 & (\partial_{\nu_j} v_j, u_{i_1})_0 \end{pmatrix}. \end{aligned}$$

On pose : $\mathbb{C}'_A := \begin{pmatrix} \mathbb{C}_A & \mathbb{S}_D \end{pmatrix} \in \mathbb{R}^{N_1 \times (2N + N_{ar})}$.

Soit $\hat{\underline{E}}_A^k \in (\mathbb{R}^2)^N$ la représentation discrétisée de $\hat{\mathbf{E}}_h$, l'approximation de $\hat{\mathbf{E}}$ dans la base B_A (10.17) de $\mathbf{X}_{\mathbf{E},2}$. La discrétisation du système de formulations variationnelles (9.12)-(9.14), accompagnées de (9.11), se met sous la forme :

Pour tout $k \in \mathbb{N}$, trouver $(\hat{\underline{\mathbf{E}}}^k, \underline{\mathbf{c}}_{A,h}^k, \underline{\mathbf{E}}_z^k, \underline{\mathbf{p}}_k) \in (\mathbb{R}^2)^N \times \mathbb{R}^{N_{ar}} \times \mathbb{R}^N \times \mathbb{R}^{N_1}$ tel que :

$$\begin{aligned}\mathbb{A}_A^k \hat{\underline{\mathbf{E}}}^k + \mathbb{L}_{A,k}^T \underline{\mathbf{c}}_{A,h}^k + \mathbb{C}_A^T \underline{\mathbf{p}} &= \underline{\mathbf{L}}_{A,k}, \\ \mathbb{L}_{A,k} \hat{\underline{\mathbf{E}}}^k + \pi \mathbb{X}_{A,h}^k \underline{\mathbf{c}}_{A,h}^k + \mathbb{S}_D^T \underline{\mathbf{p}} &= \pi \underline{\lambda}_A^k, \\ \mathbb{C}_A \hat{\underline{\mathbf{E}}}^k + \mathbb{S}_D \underline{\mathbf{c}}_{A,h}^k - \mathbb{M}_k \underline{\mathbf{E}}_z^k &= \underline{\mathbf{G}}^k + \mathbb{M}_{k,r} \underline{\mathbf{e}}_z^k, \\ \mathbb{K}_D^k \underline{\mathbf{E}}_z^k - \mathbb{M}_k^T \underline{\mathbf{p}} &= \underline{\mathbf{l}}_k - \mathbb{K}_D^{P,k} \underline{\mathbf{e}}_z^k,\end{aligned}$$

Posons :

$$\underline{\mathbf{l}}'_{A,k} = \underline{\mathbf{l}}'_k - \mathbb{K}_D^{P,k} \underline{\mathbf{e}}_z^k, \text{ et } \underline{\mathbf{G}}'_A = \underline{\mathbf{G}}^k + \mathbb{M}_{r,k} \underline{\mathbf{e}}_z^k.$$

Les équations ci-dessus se mettent aussi sous la forme condensée suivante (en reprenant les notations du paragraphe 10.6.2 de la partie II) :

$$\boxed{\begin{aligned}\begin{pmatrix} \mathbb{A}'_A & 0 \\ 0 & \mathbb{K}_D^k \end{pmatrix} \begin{pmatrix} \underline{\mathbf{E}}'_A \\ \underline{\mathbf{E}}_z^k \end{pmatrix} + \begin{pmatrix} \mathbb{C}'_A \\ -\mathbb{M}_k^T \end{pmatrix} \underline{\mathbf{p}}^k &= \begin{pmatrix} \underline{\mathbf{L}}'_{A,k} \\ \underline{\mathbf{l}}'_{A,k} \end{pmatrix}, \\ \begin{pmatrix} \mathbb{C}'_A & -\mathbb{M}_k \end{pmatrix} \begin{pmatrix} \underline{\mathbf{E}}'_A \\ \underline{\mathbf{E}}_z^k \end{pmatrix} &= \underline{\mathbf{G}}'_A.\end{aligned}} \quad (11.8)$$

L'algorithme de résolution de (11.8) est le suivant :

- Résoudre $\mathbb{A}'_A \underline{\mathbf{E}}_A = \underline{\mathbf{L}}'_{A,k}$ (voir l'algorithme de résolution de (10.18)) et $\mathbb{K}_D^k \underline{\mathbf{E}}_{0,z} = \underline{\mathbf{l}}'_k$;
- Résoudre $\begin{pmatrix} \mathbb{C}'_A & -\mathbb{M}_k \end{pmatrix} \begin{pmatrix} (\mathbb{A}'_A)^{-1} & 0 \\ 0 & (\mathbb{K}_D^k)^{-1} \end{pmatrix} \begin{pmatrix} \mathbb{C}'_A \\ -\mathbb{M}_k^T \end{pmatrix} \underline{\mathbf{p}} = \underline{\mathbf{G}}'_A$;
- Résoudre $\mathbb{A}'_A \underline{\mathbf{E}}'_A = \underline{\mathbf{L}}'_A - \mathbb{C}'_A^T \underline{\mathbf{p}}$, et $\mathbb{K}_D^k \underline{\mathbf{E}}_z^k = \underline{\mathbf{l}}'_k + \mathbb{M}_k^T \underline{\mathbf{p}}$.

Pour la deuxième étape, calcule une approximation de $\underline{\mathbf{p}}$ par l'algorithme du gradient conjugué (section 15.4, partie IV). Notons qu'à chaque étape, on devra utiliser l'algorithme de résolution de (10.18).

Chapitre 12

Le problème temporel 3D

12.1 Introduction

Les premières expériences numériques de résolution des équations de Maxwell instationnaires tridimensionnelles par des éléments finis continus ont été réalisées en 1992 par É. Heintzé [69] (voir aussi [26]). É. Heintzé a codé la méthode avec conditions aux limites essentielles. Ainsi, l'approximation du champ électromagnétique n'était correcte que pour des domaines convexes.

Nous reprenons le travail d'É. Heintzé, et nous le généralisons au cas de domaines non-convexe, pour lesquels on résout les équations de Maxwell instationnaires avec la méthode de régularisation à poids.

Dans ce chapitre, Ω représente l'intérieur d'un conducteur parfait Γ_C , borné éventuellement par une frontière artificielle Γ_A , exactement comme c'est expliqué dans la section 8.1. Dans le vide, les équations de Maxwell s'écrivent :

$$\varepsilon_0 \partial_t \mathcal{E} - \mathbf{rot} \mathcal{H} = -\mathcal{J}, \quad (12.1)$$

$$\mu_0 \partial_t \mathcal{H} + \mathbf{rot} \mathcal{E} = 0, \quad (12.2)$$

$$\operatorname{div}(\varepsilon_0 \mathcal{E}) = \rho, \quad (12.3)$$

$$\operatorname{div}(\mu_0 \mathcal{H}) = 0. \quad (12.4)$$

$\mathcal{E}$ et $\mathcal{H}$ sont le champ électrique et magnétique, ρ est la densité de charges, et $\mathcal{J}$ est le vecteur densité de courant. Ces quantités dépendent de la variable d'espace $\mathbf{x}$ et de la variable de temps t . Rappelons que ε_0 est la permittivité diélectrique du vide, μ_0 la perméabilité magnétique du vide. Les quantités ε_0 et μ_0 sont telles que $\varepsilon_0 \mu_0 c^2 = 1$, où $c \approx 3.10^8 \text{m.s}^{-1}$ est la vitesse de propagation des ondes électromagnétiques dans le vide.

On a les hypothèses de régularité minimale suivantes sur ρ et $\mathcal{J}$, obtenues en analysant les équations de Maxwell, et en notant que l'énergie électromagnétique est finie (voir par exemple [6], chap. 2) :

$$\partial_t \mathcal{J} \in L^2(0, T; \mathcal{L}^2(\Omega)),$$

$$\partial_t \rho \in C^0(0, T; H^{-1}(\Omega)), \quad (12.5)$$

$$\mathcal{J} \text{ et } \rho \text{ satisfont (1.6) : } \operatorname{div} \mathcal{J} + \partial_t \rho = 0.$$

Notons que les hypothèses de régularité sur $\mathcal{J}$ et ρ sont liées par l'équation de la charge (1.6). On suppose de plus que la valeur initiale de $\mathcal{J}$, notée $\mathcal{J}(\cdot, 0)$, existe. On a également une condition de

conducteur parfait :

$$\mathcal{E} \times \boldsymbol{\nu} = 0 \text{ sur } \Gamma_C. \quad (12.6)$$

Remarquons que (12.2) et (12.6) impliquent que :

$$\mu_0 \partial_t (\mathcal{H} \cdot \boldsymbol{\nu}) = 0 \text{ sur } \Gamma_C. \quad (12.7)$$

Rappelons que l'équation de conservation de la charge (1.6) est une conséquence des équations (12.1) et (12.3).

De plus, on connaît $\mathcal{E}$ et $\mathcal{H}$ au temps initial :

$$\mathcal{E}(\cdot, 0) = \mathcal{E}_0, \quad (12.8)$$

$$\mathcal{H}(\cdot, 0) = \mathcal{H}_0, \quad (12.9)$$

où le couple $(\mathcal{E}_0, \mathcal{H}_0)$ ne dépend que de $\mathbf{x}$.

On suppose en général que $\mu_0 \mathcal{H}_0 \cdot \boldsymbol{\nu} = 0$ sur Γ_C , de sorte que :

$$\mathcal{H} \cdot \boldsymbol{\nu} = 0 \text{ sur } \Gamma_C. \quad (12.10)$$

Sur la frontière artificielle Γ_A , on impose une condition aux limites de Silver-Müller, qui s'écrit de deux façons équivalentes :

$$\left(\mathcal{E} - \sqrt{\frac{\mu_0}{\varepsilon_0}} \mathcal{H} \times \boldsymbol{\nu} \right) \times \boldsymbol{\nu} = \left(\mathbf{e} - \sqrt{\frac{\mu_0}{\varepsilon_0}} \mathbf{h} \times \boldsymbol{\nu} \right) \times \boldsymbol{\nu} := \mathbf{e}^* \times \boldsymbol{\nu} \text{ sur } \Gamma_A \text{ ou bien :} \quad (12.11)$$

$$\left(\sqrt{\frac{\mu_0}{\varepsilon_0}} \mathcal{H} + \mathcal{E} \times \boldsymbol{\nu} \right) \times \boldsymbol{\nu} = \left(\sqrt{\frac{\mu_0}{\varepsilon_0}} \mathbf{h} + \mathbf{e} \times \boldsymbol{\nu} \right) \times \boldsymbol{\nu} := \sqrt{\frac{\mu_0}{\varepsilon_0}} \mathbf{h}^* \times \boldsymbol{\nu} \text{ sur } \Gamma_A. \quad (12.12)$$

où $\mathbf{e}^* \in \mathcal{H}^1(\Omega)$. On décompose Γ_A en Γ_A^i et Γ_A^a . Sur Γ_A^i , on modélise des ondes planes incidentes, et sur Γ_A^a , on impose une condition aux limites absorbantes, de sorte que : $\mathbf{e}^*|_{\Gamma_A^a} = 0$ et $\mathbf{h}^*|_{\Gamma_A^a} = 0$ par définition, alors que $\mathbf{e}^*|_{\Gamma_A^i}$ et $\mathbf{h}^*|_{\Gamma_A^i} = 0$ sont liés aux ondes planes.

On définit les espaces vectoriels suivants :

$$\mathcal{H}_A(\mathbf{rot}, \Omega) := \{ \mathbf{u} \in \mathcal{H}(\mathbf{rot}, \Omega) : \mathbf{u} \times \boldsymbol{\nu}_{\partial\Omega} \in \mathcal{L}_t^2(\partial\Omega), \mathbf{u} \times \boldsymbol{\nu}_{\Gamma_C} = 0 \},$$

$$\mathcal{H}_A(\mathbf{div}_{(\gamma)}, \Omega) := \{ \mathbf{u} \in \mathcal{H}(\mathbf{div}_{(\gamma)}, \Omega) : \mathbf{u} \cdot \boldsymbol{\nu}_{\partial\Omega} \in L^2(\partial\Omega), \mathbf{u} \cdot \boldsymbol{\nu}_{\Gamma_C} = 0 \}.$$

On pose alors : $\mathcal{X}_{\mathcal{E}(\cdot, \gamma)}^A = \mathcal{H}_A(\mathbf{rot}, \Omega) \cap \mathcal{H}(\mathbf{div}_{(\gamma)}, \Omega)$, et $\mathcal{X}_{\mathcal{H}(\cdot, \gamma)}^A = \mathcal{H}(\mathbf{rot}, \Omega) \cap \mathcal{H}_A(\mathbf{div}_{(\gamma)}, \Omega)$. Notons que $\mathcal{X}_{\mathcal{E}}^{A, R} := \mathcal{X}_{\mathcal{E}}^A \cap \mathcal{H}^1(\Omega)$ (resp. $\mathcal{X}_{\mathcal{H}}^{A, R} := \mathcal{X}_{\mathcal{H}}^A \cap \mathcal{H}^1(\Omega)$) est un sous-espace fermé de $\mathcal{X}_{\mathcal{E}}^A$ (resp. $\mathcal{X}_{\mathcal{H}}^A$). On a les hypothèses minimales de régularité suivantes sur le champ électromagnétique initial $(\mathcal{E}_0, \mathcal{H}_0)$:

$$\begin{aligned} \mathcal{E}_0 &\in \mathcal{H}_A(\mathbf{rot}, \Omega) \text{ tel que : } \operatorname{div}(\varepsilon_0 \mathcal{E}_0) = \rho, \\ \mathcal{H}_0 &\in \mathcal{X}_{\mathcal{H}}^A \text{ tel que : } \operatorname{div}(\mu_0 \mathcal{E}_0) = 0, \\ \left(\mathcal{E}_0 - \sqrt{\frac{\mu_0}{\varepsilon_0}} \mathcal{H}_0 \times \boldsymbol{\nu} \right) \times \boldsymbol{\nu} &= \mathbf{e}_0^* \times \boldsymbol{\nu} \text{ sur } \Gamma_A. \end{aligned}$$

(12.13)

Les hypothèses minimales de régularité sur $\mathcal{E}$ et $\mathcal{H}$ sont alors :

$$\boxed{\begin{aligned} \mathcal{E} &\in L^2(0, T; \mathcal{H}_A(\mathbf{rot}, \Omega)); \partial_t \mathcal{E} \in L^2(0, T; \mathcal{L}^2(\Omega)), \\ \mathcal{H} &\in L^2(0, T; \mathcal{X}_{\mathcal{H}}^A); \partial_t \mathcal{H} \in L^2(0, T; \mathcal{L}^2(\Omega)). \end{aligned}} \quad (12.14)$$

À partir des équations (12.1)-(12.4) de premier ordre en temps couplées en $\mathcal{E}$ et $\mathcal{H}$, on obtient des équations de second ordre en temps découplées, à partir desquelles on construit les formulations variationnelles. L'existence et l'unicité des solutions découle alors du théorème suivant, dû à J.-L. Lions et E. Magenes [78] (thm. 8.2, p. 296).

Théorème 12.1

- Soient V et H deux espaces de Hilbert tels que $H' \equiv H$ et V soit dense dans H .
- Soit a une forme bilinéaire, continue et symétrique sur V . On suppose qu'il existe $\lambda > 0$ tel que $a(\cdot, \cdot) + \lambda(\cdot, \cdot)_H$ soit coercitive sur V , c'est-à-dire :

$$\exists \alpha > 0 \mid \forall v \in V, a(v, v) + \lambda \|v\|_H^2 \geq \alpha \|v\|_V^2. \quad (12.15)$$

On définit l'opérateur $A \in \mathcal{L}(V, V')$ par : $\langle Au, v \rangle_{V', V} := a(u, v)$.

- Soit $f \in L^2(0, T; H)$, $(u_0, u_1) \in V \times H$.

On considère le problème suivant : Trouver u tel que :

$$\begin{cases} u''(t) + Au(t) = f(t), \text{ dans } V', \text{ pour } t \in]0, T[, \\ \text{avec les données de Cauchy : } u(0) = u_0 \text{ et } u'(0) = u_1. \end{cases} \quad (12.16)$$

Le problème (12.16) est équivalent à la formulation variationnelle suivante : Trouver u tel que :

$$\begin{cases} \langle u''(t), v \rangle_{V', V} + a(u(t), v) = (f(t), v)_H, \forall v \in V, \text{ pour } t \in]0, T[, \\ \text{avec les données de Cauchy : } u(0) = u_0 \text{ et } u'(0) = u_1. \end{cases} \quad (12.17)$$

Il existe une unique solution aux problèmes (12.16) et (12.17). De plus, l'application suivante η est continue :

$$\begin{aligned} \eta : L^2(0, T; H) \times V \times H &\rightarrow C^0(0, T; V) \times C^0(0, T; H) \\ (f, u_0, u_1) &\mapsto (u, u') \end{aligned}$$

Les preuves de la section qui suit sont dues à P. Ciarlet, Jr. [37].

12.2 Champ électrique : formulations variationnelles

12.2.1 Introduction

Dans un premier temps, on suppose que $\partial\Omega = \Gamma_C$: il n'y a pas de frontière absorbante. On suppose que $\mathcal{J}$ et ρ satisfont les hypothèses de régularité minimale (12.5). Le champ électromagnétique est alors tel que :

$$\mathcal{E} \in L^2(0, T; \mathcal{H}_0(\mathbf{rot}, \Omega)) \text{ et } \mathcal{H} \in L^2(0, T; \mathcal{X}_{\mathcal{H}}^0).$$

On s'intéresse à l'évolution du champ électromagnétique $(\mathcal{E}, \mathcal{H})$ solution du problème (12.1)-(12.4) accompagné des conditions initiales (12.8)-(12.9), et des conditions aux limites (12.6) et (12.10) dans l'intervalle de temps $]0, T[$ ($T \in \mathbb{R}_+^*$) dans Ω . On note : $\Omega_T = \Omega \times]0, T[$. Par la suite, nous allons détailler l'étude du champ électrique, celle du champ magnétique étant similaire.

Proposition 12.2 Dans le problème de premier ordre en temps (12.1)-(12.4), on peut remplacer (12.1) par les deux conditions suivantes :

$$\partial_t^2 \mathcal{E} + c^2 \mathbf{rot}(\mathbf{rot} \mathcal{E}) = -\partial_t \mathcal{J} / \varepsilon_0, \text{ dans } \Omega, \text{ pour } t \in]0, T[, \quad (12.18)$$

$$\partial_t \mathcal{E}(\cdot, 0) = \mathcal{E}_1 := \frac{1}{\varepsilon_0} (\mathbf{rot} \mathcal{H}_0 - \mathcal{J}(\cdot, 0)), \text{ dans } \Omega, \text{ à } t = 0, \quad (12.19)$$

DÉMONSTRATION. En dérivant (12.1) par rapport au temps, on trouve (12.18), dans $\mathcal{D}'([0, T[\times \Omega)^3$:

$$\mu_0 \mathbf{rot} \partial_t \mathcal{H} - c^2 \partial_t^2 \mathcal{E} = -\partial_t \mathcal{J} / \varepsilon_0 \stackrel{(12.2)}{\Rightarrow} -\mathbf{rot}(\mathbf{rot} \mathcal{E}) - \frac{1}{c^2} \partial_t^2 \mathcal{E} = \mu_0 \partial_t \mathcal{J}.$$

En considérant (12.1) à $t = 0$, on obtient (12.19). Considérons la variable auxiliaire :

$$\mathcal{U} := \mu_0 \mathbf{rot} \mathcal{H} - c^2 \partial_t \mathcal{E} - \mathcal{J} / \varepsilon_0.$$

Pour montrer la réciproque, on considère la variable auxiliaire : $\mathcal{U} := \mu_0 \mathbf{rot} \mathcal{H} - c^2 \partial_t \mathcal{E} - \mathcal{J} / \varepsilon_0$. Les relations (12.2) et (12.18) donnent :

$$\partial_t \mathcal{U} = \mu_0 \mathbf{rot} \partial_t \mathcal{H} - c^2 \partial_t^2 \mathcal{E} - \partial_t \mathcal{J} / \varepsilon_0 \stackrel{(12.2)}{=} \partial_t^2 \mathcal{E} + c^2 \mathbf{rot}(\mathbf{rot} \mathcal{E}) + \partial_t \mathcal{J} / \varepsilon_0 \stackrel{(12.18)}{=} 0$$

$\mathcal{U}$ est donc une constante. Par ailleurs, on a : $\mathcal{U}(\cdot, 0) = \mathbf{rot} \mathcal{H}_0 / \varepsilon_0 - \mathcal{E}_1 - \mathcal{J}(\cdot, 0) / \varepsilon_0 = 0$. Ainsi, $\mathcal{U} \equiv 0$, ce qu'il fallait démontrer. □

On cherche alors à résoudre le problème (PE) suivant : Soit $\mathcal{J}$ et ρ satisfaisant l'équation de la charge. Trouver $\mathcal{E}$ tel que :

$$\partial_t^2 \mathcal{E} + c^2 \mathbf{rot}(\mathbf{rot} \mathcal{E}) = -\frac{1}{\varepsilon_0} \partial_t \mathcal{J}, \text{ dans } \Omega, \text{ pour } t \in]0, T[, \quad (12.20)$$

$$\operatorname{div} \mathcal{E} = \frac{\rho}{\varepsilon_0}, \text{ dans } \Omega, \text{ pour } t \in]0, T[, \quad (12.21)$$

$$\mathcal{E} \times \boldsymbol{\nu}_{\partial\Omega} = 0, \text{ sur } \partial\Omega, \text{ pour } t \in]0, T[, \quad (12.22)$$

$$\mathcal{E}(\cdot, 0) = \mathcal{E}_0, \partial_t \mathcal{E}(\cdot, 0) = \mathcal{E}_1 := c^2 (\mathbf{rot} \mathcal{B}_0 - \mu_0 \mathcal{J}(\cdot, 0)), \text{ dans } \Omega, \text{ à } t = 0, \quad (12.23)$$

$$\operatorname{div} \mathcal{E}_0 = \frac{\rho(\cdot, 0)}{\varepsilon_0}, \operatorname{div} \mathcal{E}_1 = -\frac{1}{\varepsilon_0} \operatorname{div} \mathcal{J}(\cdot, 0), \text{ dans } \Omega \text{ à } t = 0. \quad (12.24)$$

Dans les paragraphes qui suivent, nous utilisons la convention suivante : $L^2_{(\gamma)}(\Omega)$ désigne au choix $L^2(\Omega)$ ou $L^2_\gamma(\Omega)$; et $(\cdot, \cdot)_{0(\gamma)}$ le produit scalaire correspondant au choix. $\mathcal{X}^0_{\mathcal{E}(\cdot, \gamma)}$ désigne au choix $\mathcal{X}^0_{\mathcal{E}}$ ou $\mathcal{X}^0_{\mathcal{E}, \gamma}$. $H^1_{0(\cdot, \gamma)}(\Omega)$ désigne au choix $H^1_0(\Omega)$ ou $H^1_{0, \gamma}(\Omega)$. Bien sûr, $L^2(\Omega)$ et $H^1_0(\Omega)$ sont utilisés avec $\mathcal{X}^0_{\mathcal{E}}$; et $L^2_\gamma(\Omega)$ et $H^1_{0, \gamma}(\Omega)$ sont utilisés avec $\mathcal{X}^0_{\mathcal{E}, \gamma}$.

12.2.2 Formulation variationnelle classique

Comme $\mathbf{rot} \mathcal{E} \in \mathcal{L}^2(\Omega)$ avec $\mathcal{E} \in \mathcal{H}_0(\mathbf{rot}, \Omega)$, on a : $\mathbf{rot}(\mathbf{rot} \mathcal{E}) \in \mathcal{H}_0(\mathbf{rot}, \Omega)'$. De plus, comme $\partial_t \mathcal{J} \in \mathcal{L}^2(\Omega)$, par différence (voir l'équation (12.20)), on a : $\partial_t^2 \mathcal{E} \in \mathcal{H}_0(\mathbf{rot}, \Omega)'$.

En multipliant (12.20) par une fonction test $\mathcal{F} \in \mathcal{H}_0(\mathbf{rot}, \Omega)$, et en intégrant sur Ω à l'aide de la formule d'intégrations par parties (7.15), on obtient alors formellement (au sens du chapitre 3 de [78]) la formulation variationnelle en $\mathcal{E}$ (FV) suivante :

Trouver $\mathcal{E}(., t) \in \mathcal{H}_0(\mathbf{rot}, \Omega)$ tel que :

$$\langle \partial_t^2 \mathcal{E}, \mathcal{F} \rangle_{\mathcal{H}', \mathcal{H}} + c^2 (\mathbf{rot} \mathcal{E}, \mathbf{rot} \mathcal{F})_0 = -\frac{1}{\varepsilon_0} (\partial_t \mathcal{J}, \mathcal{F})_0, \quad \forall \mathcal{F} \in \mathcal{H}_0(\mathbf{rot}, \Omega), \quad (12.25)$$

$$(12.26)$$

$$\mathcal{E}(., 0) = \mathcal{E}_0, \partial_t \mathcal{E}(., 0) = \mathcal{E}_1, \quad (12.27)$$

où ici, $\mathcal{H} = \mathcal{H}_0(\mathbf{rot}, \Omega)$ et $\langle ., . \rangle_{\mathcal{H}', \mathcal{H}}$ désigne le produit de dualité entre $\mathcal{H}$ et son dual $\mathcal{H}'$. On suppose que $\partial_t \mathcal{J} \in L^2(0, T; L^2(\Omega))$ pour que le second membre de cette formulation soit bien défini. La solution du problème (PE) est alors solution de (FV).

Notons que le premier terme est souvent écrit sous la forme $\frac{d^2}{dt^2}(\mathcal{E}, \mathcal{F})_0$.

Proposition 12.3 *Supposons que $\partial_t \mathcal{J} \in L^2(0, T; \mathcal{L}^2(\Omega))$, et $(\mathcal{E}_0, \mathcal{E}_1) \in \mathcal{H}_0(\mathbf{rot}, \Omega) \times \mathcal{L}^2(\Omega)$. La formulation variationnelle (FV) admet une unique solution telle que : $(\mathcal{E}, \partial_t \mathcal{E}) \in C^0(0, T; \mathcal{H}_0(\mathbf{rot}, \Omega)) \times C^0(0, T; \mathcal{L}^2(\Omega))$.*

DÉMONSTRATION. Appliquons le théorème de Lions-Magenes 12.1, avec $V = \mathcal{H}_0(\mathbf{rot}, \Omega)$ et $H = \mathcal{L}^2(\Omega)$. V est dense dans H car V contient $\mathcal{D}(\Omega)$. La forme bilinéaire symétrique a est telle que : $a(\mathbf{u}, \mathbf{v}) = (\mathbf{rot} \mathbf{u}, \mathbf{rot} \mathbf{v})_0, \forall \mathbf{u}, \mathbf{v} \in V$, et satisfait (12.15) avec $\lambda = \alpha = 1$. Par définition, la donnée $-\partial_t \mathcal{J} \in L^2(0, T; H)$, et par définition $\mathcal{E}_0 \in V$ et $\mathcal{E}_1 \in H$. D'où l'existence et l'unicité d'une solution à (FV) telle que $(\mathcal{E}, \mathcal{E}') \in C^0(0, T; V) \cap C^0(0, T; H)$. □

Montrons que la formulation variationnelle (FV) est équivalente au problème (PE), c'est-à-dire que $\mathcal{E}$ solution de (FV) satisfait (12.21).

Théorème 12.4 *Supposons que $\mathcal{J}$ et ρ satisfassent les hypothèses de régularité minimale (12.5), et que $(\mathcal{E}_0, \mathcal{E}_1) \in \mathcal{H}_0(\mathbf{rot}, \Omega) \times \mathcal{L}^2(\Omega)$. Le problème (PE) est équivalent à la formulation variationnelle (FV).*

DÉMONSTRATION. Montrons que la solution $\mathcal{E}$ de (FV) vérifie la relation de Coulomb. D'après la proposition 12.3, $\partial_t \mathcal{E} \in C^0(0, T; \mathcal{L}^2(\Omega))$, d'où : $\text{div}(\partial_t \mathcal{E}) \in C^0(0, T; H^{-1}(\Omega))$. Posons :

$$V = \text{div} \mathcal{E} - \rho/\varepsilon_0,$$

une distribution de $\mathcal{D}'(\Omega \times]0, T[)$. On a la relation suivante, vraie au sens des distributions :

$$\begin{aligned} \partial_t^2 V &= \text{div}(\partial_t^2 \mathcal{E}) - \partial_t^2 \rho/\varepsilon_0, \\ &= \text{div}(-\partial_t \mathcal{J}/\varepsilon_0 - c^2 \mathbf{rot} \mathbf{rot} \mathcal{E} - \partial_t^2 \rho/\varepsilon_0), \text{ d'après (12.20),} \\ &= -1/\varepsilon_0 \partial_t(\text{div} \mathcal{J} + \partial_t \rho) = 0. \end{aligned}$$

D'après [101] (pp. 54-55), et sachant que $\partial_t V = \text{div}(\partial_t \mathcal{E}) - \partial_t \rho/\varepsilon_0 \in C^0(0, T; H^{-1}(\Omega))$ (ce qui signifie que $\partial_t V(., 0)$ a bien un sens), on a alors : $\partial_t V = \partial_t V(., 0) \equiv (\text{div} \mathcal{E}_1 - \partial_t \rho(., 0)/\varepsilon_0) = 0$, d'après (12.24). Ainsi, on obtient : $V = V(., 0) \equiv \text{div} \mathcal{E}_0 - \rho_0/\varepsilon_0 = 0$, d'après (12.27) et (12.23). □

D'après la proposition 12.3 et le théorème 12.4, avec les hypothèses (12.5), la formulation (FV) admet une unique solution égale à la solution de (PE). Pour résoudre (PE), on peut donc discrétiser la formulation variationnelle (FV), par exemple avec des éléments finis d'arêtes, conformes dans $\mathcal{H}_0(\mathbf{rot}, \Omega)$. On déduit du théorème 12.4 le corollaire suivant :

Corollaire 12.5 *Supposons que les hypothèses du théorème 12.4 sont vérifiées, et que de plus $\rho \in C^0(0, T; L^2_{(\gamma)}(\Omega))$. Il existe alors une unique solution de (PE) ou (FV) telle que :*
 $(\mathcal{E}, \partial_t \mathcal{E}) \in C^0(0, T; \mathcal{X}^0_{\mathcal{E}(\cdot, \gamma)}) \times C^0(0, T; \mathcal{L}^2(\Omega))$.

12.2.3 Formulation variationnelle augmentée

Supposons que $\rho \in L^2(0, T; L^2_{(\gamma)}(\Omega))$. Afin de discrétiser (FV) avec des éléments finis conformes dans $\mathcal{X}^0_{\mathcal{E}, \gamma}$, on a ajouté un terme en $c^2 (\operatorname{div} \mathcal{E}, \operatorname{div} \mathcal{F})_{0(\gamma)}$ dans le membre principal de (12.25) et $c^2 (\rho, \operatorname{div} \mathcal{F})_{0(\gamma)}/\varepsilon_0$ au membre de gauche. On considère alors la formulation variationnelle augmentée (FVA) suivante :

Trouver $\mathcal{E}(\cdot, t) \in \mathcal{X}^0_{\mathcal{E}(\cdot, \gamma)}$ tel que :

$$\langle \partial_t^2 \mathcal{E}, \mathcal{F} \rangle_{\mathcal{H}', \mathcal{H}} + c^2 (\mathcal{E}, \mathcal{F})_{\mathcal{X}^0_{\mathcal{E}(\cdot, \gamma)}} = -\frac{1}{\varepsilon_0} (\partial_t \mathcal{J}, \mathcal{F})_0 + \frac{c^2}{\varepsilon_0} (\rho, \operatorname{div} \mathcal{F})_{0(\gamma)}, \quad \forall \mathcal{F} \in \mathcal{X}^0_{\mathcal{E}(\cdot, \gamma)} \quad (12.28)$$

$$\mathcal{E}(\cdot, 0) = \mathcal{E}_0, \quad \partial_t \mathcal{E}(\cdot, 0) = \mathcal{E}_1. \quad (12.29)$$

On suppose que $\partial_t \mathcal{J} \in L^2(0, T; L^2(\Omega))$ et $\rho \in L^2(0, T; L^2_{(\gamma)}(\Omega))$ pour que le second membre de cette formulation soit bien défini. La solution du problème (PE) est alors solution de (FVA).

Pour appliquer le théorème de Lions-Magenes 12.1, on devrait modifier le second membre de (12.28), de façon à ce qu'il soit égal à $(f(t), \mathcal{F})_0$, où $f \in L^2(0, T; \mathcal{L}^2(\Omega))$. Ceci étant, on peut utiliser la version légèrement modifiée de ce résultat, due à C. Bernardi et F. Ben Belgacem [15] (thm. 2.3, p. 1502). Dans ce cas, la donnée d'un second membre de la forme $(\rho(t), \operatorname{div} \mathcal{F})_{0(\gamma)}$ est suffisante pour conclure.

Proposition 12.6 *Supposons que $\partial_t \mathcal{J} \in L^2(0, T; \mathcal{L}^2(\Omega))$, que $\rho \in L^2(0, T; L^2_{(\gamma)}(\Omega))$; et que $(\mathcal{E}_0, \mathcal{E}_1) \in \mathcal{X}^0_{\mathcal{E}(\cdot, \gamma)} \times \mathcal{L}^2(\Omega)$. La formulation variationnelle (FV) admet une unique solution telle que : $(\mathcal{E}, \partial_t \mathcal{E}) \in C^0(0, T; \mathcal{X}^0_{\mathcal{E}(\cdot, \gamma)}) \cap C^0(0, T; \mathcal{L}^2(\Omega))$.*

On veut maintenant montrer que la formulation variationnelle (FVA) est équivalente au problème (PE). Pour cela, il faut établir que $\mathcal{E}$, solution de (FVA) satisfait (12.21). À cette fin, il faut simplement prouver que $\operatorname{div} \mathcal{E} = \rho/\varepsilon_0$, pour la solution de (FVA). On note que $V = \operatorname{div} \mathcal{E} - \rho/\varepsilon_0$ appartient cette fois à $L^2(0, T; L^2_{(\gamma)}(\Omega))$. Par ailleurs, si on raisonne au sens des distributions dans $\mathcal{D}'(\Omega) \times]0, T[$, on a maintenant $\partial_t^2 V - c^2 \Delta V = 0$, ce qui ne permet plus de conclure. Pour arriver à nos fins, passons par le résultat suivant :

Théorème 12.7 *Supposons que $\partial_t \rho \in C^0(0, T; L^2_{(\gamma)}(\Omega))$ et $\partial_t^2 \mathcal{J} \in L^2(0, T; \mathcal{L}^2(\Omega))$ satisfont l'équation de conservation de la charge. Le problème (PE) est équivalent à (FVA).*

DÉMONSTRATION. Posons $V = \operatorname{div} \mathcal{E} - \rho/\varepsilon_0$. Par hypothèse, $V \in C^0(0, T; L^2_{(\gamma)}(\Omega))$, et $V(\cdot, 0)$ est nul. Qui plus est, comme $\partial_t^2 \mathcal{J}$ et $\partial_t \rho$ sont de régularités suffisantes, la formulation variationnelle (FVA) avec second membre égal à :

$$-\frac{1}{\varepsilon_0} (\partial_t^2 \mathcal{J}, \mathcal{F})_0 + \frac{c^2}{\varepsilon_0} (\partial_t \rho, \operatorname{div} \mathcal{F})_{0(\gamma)}$$

admet une solution unique, qui coïncide donc avec $\partial_t \mathcal{E}$. En conséquence, $\partial_t \mathcal{E} \in C^0(0, T; \mathcal{X}^0_{\mathcal{E}(\cdot, \gamma)})$, et $\partial_t V \in C^0(0, T; L^2_{(\gamma)}(\Omega))$, et $\partial_t V(\cdot, 0) = 0$. On construit une fonction test ad hoc pour la

formulation variationnelle (FVA) en résolvant le problème adjoint ci-dessous :

Trouver $\phi(., t) \in H_{0(\cdot, \gamma)}^1(\Omega)$ tel que :

$$\langle \partial_t^2 \phi, u \rangle + c^2 (\mathbf{grad}((d_0^\gamma) \phi), \mathbf{grad}((d_0^\gamma) u))_0 = ((d_0^\gamma) V, u)_0, \forall u \in H_{0(\cdot, \gamma)}^1(\Omega), t \in]0, T[,$$

$$\text{avec : } \phi(T) = \partial_t \phi(T) = 0.$$

(12.30)

Par construction (voir [78]), $\phi \in C^0(0, T; H_{0(\cdot, \gamma)}^1(\Omega))$ et $\partial_t \phi \in C^0(0, T; L^2(\Omega))$. Comme $\partial_t \phi$ est la solution unique de (12.30) avec second membre égal à $\partial_t V$, on a : $\partial_t \phi \in C^0(0, T; H_{0(\cdot, \gamma)}^1(\Omega))$ et $\partial_t^2 \phi \in C^0(0, T; L^2(\Omega))$. En particulier, $\Delta(d_0^\gamma) \phi \in C^0(0, T; L_{(\gamma)}^2(\Omega))$. On considère alors $\mathcal{F} = \mathbf{grad}(d_0^\gamma) \phi \in C^0(0, T; \mathcal{X}_{\mathcal{E}(\cdot, \gamma)}^0)$, que l'on peut utiliser dans la formulation variationnelle (FVA), intégrée sur $]0, T[$. Tous calculs faits, et après une double intégration par parties en temps, on arrive à :

$$\int_0^T \|V(., t)\|_{0(\cdot, \gamma)}^2 dt = 0.$$

□

D'après la proposition 12.6 et le théorème 12.7, avec les bonnes hypothèses sur les données, la formulation (FVA) admet une unique solution égale à la solution de (PE). Pour résoudre (PE), on peut donc discrétiser la formulation variationnelle (FVA), par exemple avec des éléments finis de Lagrange P_k , conformes dans $\mathcal{H}_0(\mathbf{rot}, \Omega) \cap \mathcal{H}(\text{div}_{(\gamma)}, \Omega)$.

12.2.4 Formulation variationnelle augmentée mixte

Rappelons que lorsqu'on souhaite résoudre le système couplé Maxwell-Vlasov, afin de satisfaire numériquement l'équation de la charge, on renforce la condition de divergence par un multiplicateur de Lagrange. Considérons dans un premier temps la formulation variationnelle augmentée "à demi" mixte (FVA $_{\frac{1}{2}}^M$) suivante :

Trouver $\mathcal{E}(., t) \in \mathcal{X}_{\mathcal{E}(\cdot, \gamma)}^0$ tel que :

$$\langle \partial_t^2 \mathcal{E}, \mathcal{F} \rangle_{\mathcal{H}', \mathcal{H}} + c^2 (\mathcal{E}, \mathcal{F})_{\mathcal{X}_{\mathcal{E}(\cdot, \gamma)}^0} = -\frac{1}{\varepsilon_0} (\partial_t \mathcal{J}, \mathcal{F})_0 + \frac{c^2}{\varepsilon_0} (\rho, \text{div } \mathcal{F})_{0(\cdot, \gamma)}, \quad \forall \mathcal{F} \in \mathcal{X}_{\mathcal{E}(\cdot, \gamma)}^0 \quad (12.31)$$

$$(\text{div } \mathcal{E}, q)_{0(\cdot, \gamma)} = \frac{1}{\varepsilon_0} (\rho, q)_{0(\cdot, \gamma)}, \quad \forall q \in L_{(\gamma)}^2(\Omega), \quad (12.32)$$

$$\mathcal{E}(., 0) = \mathcal{E}_0, \partial_t \mathcal{E}(., 0) = \mathcal{E}_1. \quad (12.33)$$

La solution de (FVA $_{\frac{1}{2}}^M$) satisfait (FVA) et l'équation de Coulomb. Si de plus on suppose que $\partial_t \rho \in C^0(0, T; L_{(\gamma)}^2(\Omega))$, alors la solution de (FVA) satisfait (FVA $_{\frac{1}{2}}^M$) : (FVA $_{\frac{1}{2}}^M$) est équivalente à (FVA), qui est équivalente à (PE). La formulation variationnelle (FVA $_{\frac{1}{2}}^M$) va nous permettre de montrer que la formulation variationnelle augmentée mixte (FVAM) suivante est équivalente à (PE) :

Trouver $(\mathcal{E}(\cdot, t), p(\cdot, t)) \in \mathcal{X}_{\mathcal{E}(\cdot, \gamma)}^0 \times L_{(\gamma)}^2(\Omega)$ tel que :

$$\begin{aligned} & \langle \partial_t^2 \mathcal{E}, \mathcal{F} \rangle_{\mathcal{H}', \mathcal{H}} + c^2 (\mathcal{E}, \mathcal{F})_{\mathcal{X}_{\mathcal{E}(\cdot, \gamma)}^0} + (p, \operatorname{div} \mathcal{F})_{0(\cdot, \gamma)} \\ &= -\frac{1}{\varepsilon_0} (\partial_t \mathcal{J}, \mathcal{F})_0 + \frac{c^2}{\varepsilon_0} (\rho, \operatorname{div} \mathcal{F})_{0(\cdot, \gamma)}, \quad \forall \mathcal{F} \in \mathcal{X}_{\mathcal{E}(\cdot, \gamma)}^0, \end{aligned} \quad (12.34)$$

$$(\operatorname{div} \mathcal{E}, q)_{0(\cdot, \gamma)} = \frac{1}{\varepsilon_0} (\rho, q)_{0(\cdot, \gamma)} \quad \forall q \in L_{(\gamma)}^2(\Omega), \quad (12.35)$$

$$\mathcal{E}(\cdot, 0) = \mathcal{E}_0, \quad \partial_t \mathcal{E}(\cdot, 0) = \mathcal{E}_1. \quad (12.36)$$

Pour prouver que les formulations (FVA $\frac{1}{2}$ M) et (FVAM) sont équivalentes, il faut et il suffit de montrer que $p = 0$. Notons que p est solution de (FVAM) si et seulement si $p \in C^0(0, T; L_{(\gamma)}^2(\Omega))$. Montrons tout d'abord le lemme suivant :

Lemme 12.8 *Soit $\mathcal{F} \in \mathcal{H}_0(\mathbf{rot}, \Omega)'$. Alors $\operatorname{div} \mathcal{F} \in \mathcal{H}^{-1}(\Omega)$.*

DÉMONSTRATION. Considérons $\mathcal{F} \in \mathcal{H}_0(\mathbf{rot}, \Omega)'$. Pour tout $\eta \in \mathcal{D}(\Omega)$, on a $\mathbf{grad} \eta \in \mathcal{H}$, d'où :

$$\langle \operatorname{div} \mathcal{F}, \eta \rangle_{\mathcal{D}', \mathcal{D}} = - \langle \mathcal{F}, \mathbf{grad} \eta \rangle_{\mathcal{H}', \mathcal{H}} \Rightarrow | \langle \operatorname{div} \mathcal{F}, \eta \rangle_{\mathcal{D}', \mathcal{D}} | \leq \| \mathcal{F} \|_{\mathcal{H}'} \| \eta \|_{H^1}.$$

Donc le produit $\langle \operatorname{div} \mathcal{F}, \eta \rangle_{\mathcal{D}', \mathcal{D}}$ est borné, continûment en fonction de $\| \eta \|_{H^1}$. Comme c'est valable pour tout $\eta \in H_0^1(\Omega)$, par densité de $H_0^1(\Omega)$ dans $\mathcal{D}(\Omega)$ et par prolongement, $\operatorname{div} \mathcal{F} \in H^{-1}(\Omega)$. □

Ce lemme nous permet de montrer la proposition suivante :

Proposition 12.9 *Supposons que $\partial_t^2 \rho \in L^2(0, T; L_{(\gamma)}^2(\Omega))$. Alors $p = 0$ dans (12.34).*

DÉMONSTRATION. Comme $p \in C^0(0, T; L_{(\gamma)}^2(\Omega))$, il existe un unique $\phi(\cdot, t) \in C^0(0, T; H_0^1(\Omega))$ tel que : $\Delta \phi = p$. Soit $\mathbf{v} = \mathbf{grad} \phi$. On a donc $\mathbf{v} \in C^0(0, T; \mathcal{X}_{\mathcal{E}(\cdot, \gamma)}^0)$ tel que : $\mathbf{rot} \mathbf{v} = 0$ et $\operatorname{div} \mathbf{v} = p$ (voir la section 8.2). D'après [6] (p. 44), on sait que $\partial_t^2 \mathcal{E} \in L^2(0, T; \mathcal{H}_0(\mathbf{rot}, \Omega)')$. En injectant $\mathbf{v}$ dans (12.34) intégrée sur $]0, T[$, on obtient, en reprenant la notation de (FV) :

$$\begin{aligned} & \int_0^T \left(\langle \partial_t^2 \mathcal{E}, \mathbf{grad} \phi \rangle_{\mathcal{H}', \mathcal{H}} + c^2 (\operatorname{div} \mathcal{E}, p)_{0(\cdot, \gamma)} + \| p \|_{0(\cdot, \gamma)}^2 \right) dt \\ &= \int_0^T \left(-\frac{1}{\varepsilon_0} (\partial_t \mathcal{J}, \mathbf{grad} \phi)_0 + \frac{c^2}{\varepsilon_0} (\rho, p)_{0(\cdot, \gamma)} \right) dt. \end{aligned}$$

Or, d'après (12.35), $\operatorname{div} \mathcal{E} = \rho/\varepsilon_0$. D'où :

$$\int_0^T \left(\langle \partial_t^2 \mathcal{E}, \mathbf{grad} \phi \rangle_{\mathcal{H}', \mathcal{H}} + \| p \|_{0(\cdot, \gamma)}^2 \right) dt = \int_0^T \left(-\frac{1}{\varepsilon_0} (\partial_t \mathcal{J}, \mathbf{grad} \phi)_0 \right) dt.$$

D'après le lemme 12.8, comme $\partial_t^2 \mathcal{E}(\cdot, t) \in \mathcal{H}'$, pour tout t , on a également $\operatorname{div}(\partial_t^2 \mathcal{E}(\cdot, t)) \in H^{-1}(\Omega)$. Comme $\phi(\cdot, t) \in H_0^1(\Omega)$, on peut alors écrire :

$$\int_0^T \left(\langle \operatorname{div}(\partial_t^2 \mathcal{E}), \phi \rangle_{H^{-1}, H_0^1} + \| p \|_{0(\cdot, \gamma)}^2 \right) dt = \int_0^T \left(-\frac{1}{\varepsilon_0} \langle \operatorname{div}(\partial_t \mathcal{J}), \phi \rangle_{H^{-1}, H_0^1} \right) dt.$$

Or, on a : $\operatorname{div}(\partial_t \mathcal{J}) = \partial_t \operatorname{div} \mathcal{J} = -\partial_t^2 \rho \in L^2(0, T; L^2_{(\gamma)}(\Omega))$ d'après (1.6). De plus, comme $\operatorname{div} \mathcal{E} = \rho/\varepsilon_0$, alors : $\operatorname{div}(\partial_t^2 \mathcal{E}) = \partial_t^2 \rho/\varepsilon_0 \in L^2(0, T; L^2_{(\gamma)}(\Omega))$. Finalement, on aboutit à :

$$\int_0^T \|p(\cdot, t)\|_{0(\cdot, \gamma)}^2 dt = 0.$$

□

On en déduit que lorsque $\partial_t^2 \rho \in L^2(0, T; L^2_{(\gamma)}(\Omega))$, (FVAM) est équivalent à (FVA $\frac{1}{2}$ M).

Proposition 12.10 *Supposons que $\partial_t \mathcal{J} \in L^2(0, T; \mathcal{L}^2(\Omega))$ et $\partial_t \rho \in C^0(0, T; L^2_{(\gamma)}(\Omega))$ satisfont l'équation de la charge. Soit $(\mathcal{E}_0, \mathcal{E}_1) \in \mathcal{X}_{\mathcal{E}(\cdot, \gamma)}^0 \times \mathcal{L}^2(\Omega)$. La formulation variationnelle (FVAM) admet une unique solution telle que $p \equiv 0$ et $(\mathcal{E}, \partial_t \mathcal{E}) \in C^0(0, T; \mathcal{X}_{\mathcal{E}(\cdot, \gamma)}^0) \times C^0(0, T; \mathcal{H}(\operatorname{div}_{(\gamma)}, \Omega))$.*

DÉMONSTRATION. • Pour l'existence, notons qu'avec les hypothèses sur les données et les conditions initiales, d'après la proposition 12.3, le théorème 12.4, et le corollaire 12.5, il existe une unique solution à (FV) telle que $(\mathcal{E}, \partial_t \mathcal{E}) \in C^0(0, T; \mathcal{X}_{\mathcal{E}(\cdot, \gamma)}^0) \times C^0(0, T; \mathcal{H}(\operatorname{div}_{(\gamma)}, \Omega))$. On remarque alors que $(\mathcal{E}, 0)$ est aussi solution de (FVAM).

• Montrons l'unicité : soient deux solutions de (FVAM) $(\mathcal{E}_1, p_1)$ et $(\mathcal{E}_2, p_2)$. Alors le couple $(\mathcal{E}, p) = (\mathcal{E}_1 - \mathcal{E}_2, p_1 - p_2)$ vérifie :

$$\begin{cases} \partial_t^2 (\mathcal{E}, \mathcal{F})_0 + c^2 (\operatorname{rot} \mathcal{E}, \operatorname{rot} \mathcal{F})_0 + c^2 (\operatorname{div} \mathcal{E}, \operatorname{div} \mathcal{F})_{0(\gamma)} + (p, \operatorname{div} \mathcal{F})_{0(\gamma)} = 0, & \forall \mathcal{F} \in \mathcal{X}_{\mathcal{E}(\cdot, \gamma)}^0, \\ (\operatorname{div} \mathcal{E}, q)_{0(\gamma)} = 0, & \forall q \in L^2_{(\gamma)}(\Omega). \end{cases}$$

D'après la proposition 12.9, comme $0 \in H^2$, on a $p \equiv 0$. D'après la seconde ligne, $(\operatorname{div} \mathcal{E}, \operatorname{div} \mathcal{F})_{0(\gamma)}$ est nul. Il reste donc :

$$\begin{cases} \langle \partial_t^2 \mathcal{E}, \mathcal{F} \rangle_{\mathcal{H}', \mathcal{H}} + c^2 (\operatorname{rot} \mathcal{E}, \operatorname{rot} \mathcal{F})_0 = 0, & \forall \mathcal{F} \in \mathcal{X}_{\mathcal{E}(\cdot, \gamma)}^0, \\ \operatorname{div} \mathcal{E} = 0, \end{cases}$$

avec $\mathcal{H} = \mathcal{H}_0(\operatorname{rot}, \Omega)$. Pour se ramener à (FV) avec comme donnée $\mathcal{J} = 0$, on procède comme suit : Soit $\mathbf{w} \in \mathcal{H}_0(\operatorname{rot}, \Omega)$. Il existe un unique $\phi_w \in H_0^1(\Omega)$ tel que : $\Delta \phi_w = \operatorname{div} \mathbf{w}$. On a alors $\mathbf{v} = \mathbf{w} - \operatorname{grad} \phi_w \in \mathcal{X}_{\mathcal{E}(\cdot, \gamma)}^0$, puisque $\operatorname{div} \mathbf{v} = 0$, et $\operatorname{rot} \mathbf{v} = \operatorname{rot} \mathbf{w}$. En prenant $\mathcal{F} = \mathbf{v}$, on trouve alors :

$$\begin{aligned} \langle \partial_t^2 \mathcal{E}, \mathbf{w} \rangle_{\mathcal{H}', \mathcal{H}} - \langle \partial_t^2 \mathcal{E}, \operatorname{grad} \phi_w \rangle_{\mathcal{H}', \mathcal{H}} + c^2 (\operatorname{rot} \mathcal{E}, \operatorname{rot} \mathbf{v})_0 &= 0, \\ \langle \partial_t^2 \mathcal{E}, \mathbf{w} \rangle_{\mathcal{H}', \mathcal{H}} - \langle \operatorname{div}(\partial_t^2 \mathcal{E}), \phi_w \rangle_{H^{-1}, H_0^1} + c^2 (\operatorname{rot} \mathcal{E}, \operatorname{rot} \mathbf{w})_0 &= 0, \text{ d'après le lemme 12.8,} \\ \langle \partial_t^2 \mathcal{E}, \mathbf{w} \rangle_{\mathcal{H}', \mathcal{H}} + c^2 (\operatorname{rot} \mathcal{E}, \operatorname{rot} \mathbf{w})_0 &= 0, \text{ comme } \operatorname{div} \mathcal{E} = 0. \end{aligned}$$

La dernière ligne correspond à (FV) puisqu'elle est valable pour tous $\mathbf{w} \in \mathcal{H}_0(\operatorname{rot}, \Omega)$.

□

Théorème 12.11 *Supposons que $\partial_t \mathcal{J} \in L^2(0, T; \mathcal{L}^2(\Omega))$ et $\partial_t^2 \rho \in L^2(0, T; L^2_{(\gamma)}(\Omega))$ satisfont l'équation de la charge. Soit $(\mathcal{E}_0, \mathcal{E}_1) \in \mathcal{X}_{\mathcal{E}(\cdot, \gamma)}^0 \times \mathcal{L}^2(\Omega)$. Alors (FVAM) est équivalente à (PE).*

DÉMONSTRATION. Avec ces hypothèses, d'après la proposition 12.9, (FVAM) est équivalent à (FVA $\frac{1}{2}$ M), qui est équivalent à (PE) d'après le théorème 12.7.

□

Remarque 12.12 Voir aussi [13], pour une résolution directe de la (FVAM), lorsque la relation de conservation de la charge n'est pas vérifiée.

D'après la proposition 12.10 et le théorème 12.11, avec les bonnes hypothèses sur les données, (FVAM) est équivalent à (PE). On peut donc approcher la solution de (PE) en discrétisant (FVAM), par exemple avec les éléments finis de Taylor-Hood P_{k+1} - P_k .

Considérons alors les hypothèses suivantes sur les données et sur les conditions initiales :

$$\begin{aligned}
 & \partial_t \mathcal{J} \in L^2(0, T; \mathcal{L}^2(\Omega)), \\
 & \partial_t^2 \rho \in L^2(0, T; L^2_{(\gamma)}(\Omega)), \\
 & \mathcal{J} \text{ et } \rho \text{ satisfont (1.6) : } \operatorname{div} \mathcal{J} + \partial_t \rho = 0, \\
 & (\mathcal{E}_0, \mathcal{E}_1) \in \mathcal{X}_{\mathcal{E}(\cdot, \gamma)}^0 \times \mathcal{L}^2(\Omega), .
 \end{aligned} \tag{12.37}$$

On a trois méthodes de calcul du champ électrique :

- Résoudre (FV) dans $\mathcal{H}_0(\mathbf{rot}, \Omega)$ (proposition 12.3, théorème 12.4),
- Résoudre (FVA) dans $\mathcal{X}_{\mathcal{E}(\cdot, \gamma)}^0$ (proposition 12.6, théorème 12.7),
- Résoudre (FVAM) dans $\mathcal{X}_{\mathcal{E}(\cdot, \gamma)}^0 \times L^2_{(\gamma)}(\Omega)$ (proposition 12.10, théorème 12.11).

12.2.5 Existence d'une frontière artificielle

Supposons maintenant qu'il existe une frontière artificielle Γ_A . Nous allons introduire la condition de Silver-Müller (12.12) dans la formulation variationnelle mixte augmentée. En dérivant cette condition par rapport au temps, on a :

$$\sqrt{\frac{\mu_0}{\varepsilon_0}} \partial_t (\mathcal{H} \times \boldsymbol{\nu}) = -\partial_t ((\mathcal{E} \times \boldsymbol{\nu}) \times \boldsymbol{\nu}) + \sqrt{\frac{\mu_0}{\varepsilon_0}} \partial_t (\mathbf{h}^* \times \boldsymbol{\nu}) \text{ sur } \Gamma_A.$$

On élimine ensuite le terme $\partial_t (\mathcal{H} \times \boldsymbol{\nu})$ en prenant la trace tangentielle (i.e. $\cdot \times \boldsymbol{\nu}|_{\Gamma_A}$) de l'équation de Maxwell-Faraday (12.2), soit :

$$\sqrt{\frac{\mu_0}{\varepsilon_0}} \partial_t (\mathcal{H} \times \boldsymbol{\nu})|_{\Gamma_A} + c \mathbf{rot} (\mathcal{E} \times \boldsymbol{\nu})|_{\Gamma_A} = 0.$$

On obtient alors comme condition aux limites pour $\mathcal{E}$ sur Γ_A , l'expression :

$$\partial_t ((\mathcal{E} \times \boldsymbol{\nu}) \times \boldsymbol{\nu}) - c \mathbf{rot} (\mathcal{E} \times \boldsymbol{\nu}) = \sqrt{\frac{\mu_0}{\varepsilon_0}} \partial_t (\mathbf{h}^* \times \boldsymbol{\nu}).$$

En conséquence, on doit ajouter deux termes supplémentaires dans la formulation variationnelle. Ces termes proviennent de la formule d'intégration par parties sur le terme en $\mathbf{rot} \mathbf{rot}$:

$$c^2 \langle \mathbf{rot} (\mathbf{rot} \mathcal{E}), \mathcal{F} \rangle = c^2 \int_{\Omega} \mathbf{rot} \mathcal{E} \cdot \mathbf{rot} \mathcal{F} \, d\Omega - c^2 \int_{\partial\Omega} (\mathbf{rot} \mathcal{E} \times \boldsymbol{\nu}) \cdot \mathcal{F} \, d\Sigma.$$

Or, on a :

$$\begin{aligned}
-c^2 \int_{\partial\Omega} (\mathbf{rot} \mathcal{E} \times \boldsymbol{\nu}) \cdot \mathcal{F} \, d\Sigma &= -c^2 \int_{\Gamma_C} (\mathbf{rot} \mathcal{E} \times \boldsymbol{\nu}) \cdot \mathcal{F} \, d\Sigma - c^2 \int_{\Gamma_A} (\mathbf{rot} \mathcal{E} \times \boldsymbol{\nu}) \cdot \mathcal{F} \, d\Sigma, \\
&= c^2 \int_{\Gamma_C} \mathbf{rot} \mathcal{E} \cdot (\mathcal{F} \times \boldsymbol{\nu}) \, d\Sigma - c^2 \int_{\Gamma_A} (\mathbf{rot} \mathcal{E} \times \boldsymbol{\nu}) \cdot \mathcal{F} \, d\Sigma, \\
&= -c^2 \int_{\Gamma_A} (\mathbf{rot} \mathcal{E} \times \boldsymbol{\nu}) \cdot \mathcal{F} \, d\Sigma, \text{ si } \mathcal{F} \in \mathcal{H}_A(\mathbf{rot}, \Omega), \\
&= -c \int_{\Gamma_A} ((\partial_t \mathcal{E} \times \boldsymbol{\nu}) \times \boldsymbol{\nu}) \cdot \mathcal{F} \, d\Sigma + \frac{1}{\varepsilon_0} \int_{\Gamma_A} \partial_t (\mathbf{h}^* \times \boldsymbol{\nu}) \cdot \mathcal{F} \, d\Sigma, \\
&= c \int_{\Gamma_A} (\partial_t \mathcal{E} \times \boldsymbol{\nu}) \cdot (\mathcal{F} \times \boldsymbol{\nu}) \, d\Sigma + \frac{1}{\varepsilon_0} \int_{\Gamma_A} \partial_t (\mathbf{h}^* \times \boldsymbol{\nu}) \cdot \mathcal{F} \, d\Sigma.
\end{aligned}$$

La formulation variationnelle (FVAM) se réécrit alors :

Trouver $(\mathcal{E}(\cdot, t), p(\cdot, t)) \in \mathcal{X}_{\mathcal{E}(\cdot, \gamma)}^A \times L_{(\gamma)}^2(\Omega)$ tel que :

$$\begin{aligned}
\frac{d^2}{dt^2}(\mathcal{E}, \mathcal{F})_0 + c \int_{\Gamma_A} (\partial_t \mathcal{E} \times \boldsymbol{\nu}) \cdot (\mathcal{F} \times \boldsymbol{\nu}) \, d\Sigma + c^2 (\mathcal{E}, \mathcal{F})_{\mathcal{X}_{\mathcal{E}(\cdot, \gamma)}^0} \\
+ (p, \operatorname{div} \mathcal{F})_{0(\cdot, \gamma)} = -\frac{1}{\varepsilon_0} (\partial_t \mathcal{J}, \mathcal{F})_0 + \frac{c^2}{\varepsilon_0} (\rho, \operatorname{div} \mathcal{F})_{0(\cdot, \gamma)} \\
- \frac{1}{\varepsilon_0} \int_{\Gamma_A} \partial_t (\mathbf{h}^* \times \boldsymbol{\nu}) \cdot \mathcal{F} \, d\Sigma, \quad \forall \mathcal{F} \in \mathcal{X}_{\mathcal{E}(\cdot, \gamma)} \quad (12.38)
\end{aligned}$$

$$(\operatorname{div} \mathcal{E}, q)_{0(\cdot, \gamma)} = \frac{1}{\varepsilon_0} (\rho, q)_{0(\cdot, \gamma)}, \quad \forall q \in L_{(\gamma)}^2(\Omega), \quad (12.39)$$

$$\mathcal{E}(\cdot, 0) = \mathcal{E}_0, \quad \partial_t \mathcal{E}(\cdot, 0) = \mathcal{E}_1. \quad (12.40)$$

12.3 Champ magnétique : formulation variationnelle

De même que pour le champ électrique, on peut modifier le système (12.1)-(12.4), de façon à n'avoir à résoudre qu'un problème dépendant de $\mathcal{H}$. Rappelons que le produit scalaire dans $\mathcal{X}_{\mathcal{E}(\cdot, \gamma)}^0$ et dans $\mathcal{X}_{\mathcal{H}(\cdot, \gamma)}^0$ est le même. La proposition suivante se montre de la même façon que la proposition (12.2).

Proposition 12.13 *Dans le problème de premier ordre en temps (12.1)-(12.4), on peut remplacer (12.2) par les trois conditions suivantes :*

$$\partial_t^2 \mathcal{H} + c^2 \mathbf{rot} ((\mathbf{rot} \mathcal{H} - \mathcal{J})) = 0, \quad \text{dans } \Omega, \text{ pour } t \in]0, T[, \quad (12.41)$$

$$\varepsilon_0^{-1} (\mathbf{rot} \mathcal{H} - \mathcal{J}) \times \boldsymbol{\nu} = 0, \quad \text{sur } \partial\Omega, \text{ pour } t \in]0, T[, \quad (12.42)$$

$$\partial_t \mathcal{H}(\cdot, 0) = \mathcal{H}_1 := -\mu_0^{-1} \mathbf{rot} \mathcal{E}_0, \quad \text{dans } \Omega, \text{ à } t = 0. \quad (12.43)$$

La formulation variationnelle mixte augmentée correspondant à ce problème s'écrit :

Trouver $(\mathcal{H}(\cdot, t), p(\cdot, t)) \in \mathcal{X}_{\mathcal{H}(\cdot, \gamma)}^A \times L_{(\gamma)}^2(\Omega)$ tel que :

$$\begin{aligned} & \frac{d^2}{dt^2}(\mathcal{H}, \mathcal{F})_0 + c \int_{\Gamma_A} (\partial_t \mathcal{H} \times \boldsymbol{\nu}) \cdot (\mathcal{F} \times \boldsymbol{\nu}) d\Sigma + c^2 (\mathcal{H}, \mathcal{F})_{\mathcal{X}_{\mathcal{H}(\cdot, \gamma)}^0} \\ & + (p, \operatorname{div} \mathcal{F})_{0(\cdot, \gamma)} = c^2 (\mathcal{J}, \mathbf{rot} \mathcal{F})_0 + \frac{1}{\mu_0} \int_{\Gamma_A} \partial_t (\mathbf{e}^* \times \boldsymbol{\nu}) \cdot \mathcal{F} d\Sigma, \quad \forall \mathcal{F} \in \mathcal{X}_{\mathcal{H}(\cdot, \gamma)} \end{aligned} \quad (12.44)$$

$$(\operatorname{div} \mathcal{H}, q)_{0(\cdot, \gamma)} = 0, \quad \forall q \in L_{(\gamma)}^2(\Omega), \quad (12.45)$$

$$\mathcal{H}(\cdot, 0) = \mathcal{H}_0, \partial_t \mathcal{H}(\cdot, 0) = \mathcal{H}_1. \quad (12.46)$$

À notre connaissance, il n'existe pas de travaux publiés établissant que $\mathcal{X}_{\mathcal{H}, \gamma}^0 \cap \mathcal{H}^1(\Omega)$ soit dense dans $\mathcal{X}_{\mathcal{H}, \gamma}^0$. Néanmoins d'après M. Dauge (communication privée), ce résultat est vrai sous des hypothèses sur γ semblables à celles qui permettent de retrouver la densité de $\mathcal{X}_{\mathcal{E}, \gamma}^0 \cap \mathcal{H}^1(\Omega)$ dans $\mathcal{X}_{\mathcal{E}, \gamma}^0$.

12.4 Semi-discrétisation en temps

Lorsque nous avons étudié le problème quasi-statique, nous avons obtenu une discrétisation en espace des équations de Maxwell. Il ne nous reste qu'à étudier la discrétisation en temps. Pour cela, on utilise un schéma explicite, centré de second ordre en temps. Soit T le temps final et N_t tel que : $T = N_t \Delta t$, où $N_t \in \mathbb{N}^*$ et $\Delta t > 0$ est le pas de temps choisi. Posons $t_n = n \Delta t$, $t_{n+1/2} = (n+1/2) \Delta t$, et $\mathcal{U}^n := \mathcal{U}(t_n)$, $\mathcal{U}^{n+1/2} := \mathcal{U}(t_{n+1/2})$. La dérivée seconde de $\mathcal{U}$ est approchée par un schéma de Newmark :

$$\frac{d^2}{dt^2} \mathcal{U}(t_n) = \frac{\mathcal{U}^{n+1} - 2\mathcal{U}^n + \mathcal{U}^{n-1}}{\Delta t^2} + \mathcal{O}(\Delta t^2).$$

La dérivée première est approchée par un schéma saute-mouton :

$$\frac{d}{dt} \mathcal{U}(t_n) = \frac{\mathcal{U}^{n+1} - \mathcal{U}^{n-1}}{2\Delta t} + \mathcal{O}(\Delta t^2).$$

Le champ électrique est défini aux instants entiers t_n et le champ magnétique est défini aux instants demi-entiers $t_{n+1/2}$. On peut alors écrire les schémas semi-discrétisés en temps suivants :

- Pour le champ électrique :

Pour tout $n \in \{0, \dots, N_t\}$, trouver $(\mathcal{E}^{n+1}, p^{n+1}) \in \mathcal{X}_{\mathcal{E}(\cdot, \gamma)}^A \times L_{(\gamma)}^2(\Omega)$ tel que : $\forall \mathcal{F} \in \mathcal{X}_{\mathcal{E}(\cdot, \gamma)}^A$,

$$\begin{aligned} & \frac{1}{\Delta t^2} (\mathcal{E}^{n+1} - 2\mathcal{E}^n + \mathcal{E}^{n-1}, \mathcal{F})_0 + \frac{c}{2\Delta t} \int_{\Gamma_A} ((\mathcal{E}^{n+1} - \mathcal{E}^{n-1}) \times \boldsymbol{\nu}) \cdot (\mathcal{F} \times \boldsymbol{\nu}) d\Sigma \\ & + c^2 (\mathcal{E}^n, \mathcal{F})_{\mathcal{X}_{\mathcal{E}(\cdot, \gamma)}^0} + (p^{n+1}, \operatorname{div} \mathcal{F})_{0(\cdot, \gamma)} \\ & = -\frac{1}{\varepsilon_0 \Delta t} (\mathcal{J}^{n+1/2} - \mathcal{J}^{n-1/2}, \mathcal{F})_0 + \frac{c^2}{\varepsilon_0} (\rho^n, \operatorname{div} \mathcal{F})_{0(\cdot, \gamma)} \\ & - \frac{1}{\Delta t} \int_{\Gamma_A} \left[\left(\frac{1}{\varepsilon_0} (\mathbf{h}^{n+1/2} - \mathbf{h}^{n-1/2}) + \left(\frac{\mathbf{e}^{n+1} - \mathbf{e}^{n-1}}{2} \times \boldsymbol{\nu} \right) \right) \times \boldsymbol{\nu} \right] \cdot \mathcal{F} d\Sigma \\ & := \mathcal{L}^n(\mathcal{F}), \end{aligned} \quad (12.47)$$

et $\forall q \in L^2_{(\gamma)}(\Omega)$,

$$(\operatorname{div} \mathcal{E}^{n+1}, q)_{0(\cdot, \gamma)} = \frac{1}{\varepsilon_0} (\rho^{n+1}, q)_{0(\cdot, \gamma)}, \quad (12.48)$$

avec les données initiales :

$$\mathcal{E}^0 = \mathcal{E}_0, \mathcal{E}^{-1} = \mathcal{E}_0 - \Delta t \mathcal{E}_1, p^0 = 0. \quad (12.49)$$

• Pour le champ magnétique :

Pour tout $n \in \{0, \dots, N_t\}$, trouver $(\mathcal{H}^{n+3/2}, p^{n+3/2}) \in \mathcal{X}_{\mathcal{H}(\cdot, \gamma)} \times L^2_{(\gamma)}(\Omega)$ tel que : $\forall \mathcal{F} \in \mathcal{X}_{\mathcal{E}(\cdot, \gamma)}^A$,

$$\begin{aligned} & \frac{1}{\Delta t^2} (\mathcal{H}^{n+3/2} - 2\mathcal{H}^{n+1/2} + \mathcal{H}^{n-1/2}, \mathcal{F})_0 + \frac{c}{2\Delta t} \int_{\Gamma_A} ((\mathcal{H}^{n+1/2} + \mathcal{H}^{n-1/2}) \times \boldsymbol{\nu}) \cdot (\mathcal{F} \times \boldsymbol{\nu}) \, d\Sigma \\ & + c^2 (\mathcal{H}^{n+1/2}, \mathcal{F})_{\mathcal{X}_{\mathcal{E}(\cdot, \gamma)}^0} + (p^{n+1/2}, \operatorname{div} \mathcal{F})_{0(\cdot, \gamma)} = c^2 (\mathcal{J}^{n+1/2}, \mathbf{rot} \mathcal{F})_0 \\ & - \frac{1}{\mu_0 \Delta t} \int_{\Gamma_A} \left[\left(\frac{1}{\varepsilon_0} (\mathbf{e}^{n+1} - \mathbf{e}^n) + \sqrt{\frac{\mu_0}{\varepsilon_0}} \left(\frac{\mathbf{h}^{n+1/2} - \mathbf{h}^{n-1/2}}{2} \times \boldsymbol{\nu} \right) \right) \times \boldsymbol{\nu} \right] \cdot \mathcal{F} \, d\Sigma \\ & := \mathcal{M}^{n+1/2}(\mathcal{F}), \end{aligned} \quad (12.50)$$

et, $\forall q \in L^2_{(\gamma)}(\Omega)$,

$$(\operatorname{div} \mathcal{H}^{n+3/2}, q)_{0(\cdot, \gamma)} = 0, \quad (12.51)$$

avec les conditions initiales :

$$\mathcal{H}^{1/2} = \mathcal{H}_0 - \frac{\Delta t}{2} \mathcal{H}_1, \mathcal{H}^{-1/2} = \mathcal{H}_0 + \frac{\Delta t}{2} \mathcal{H}_1, p^0 = 0. \quad (12.52)$$

12.5 Discrétisation complète

Nous allons maintenant discrétiser les schémas en espace, afin d'avoir à résoudre un système matriciel à chaque pas de temps. Comme nous avons beaucoup détaillé la discrétisation du problème statique (chapitres 10 et 11 pour le champ électrique, section 16.3 de la partie IV pour le champ magnétique), nous ne donnons ici que de brèves indications pour la construction des matrices.

12.5.1 Élimination des CL presque essentielles

Si Ω n'est pas convexe, $\mathcal{X}_{\mathcal{E}}^{A,R}$ (*resp.* $\mathcal{X}_{\mathcal{H}}^{A,R}$) n'est pas dense dans $\mathcal{X}_{\mathcal{E}}^A$ (*resp.* $\mathcal{X}_{\mathcal{H}}^A$). En revanche, pour $1 - \alpha < \gamma < 1$, $\mathcal{X}_{\mathcal{E}}^{A,R}$ est dense dans $\mathcal{X}_{\mathcal{E},\gamma}^A$. On n'a pas de preuve que $\mathcal{X}_{\mathcal{H}}^{A,R}$ soit dense dans $\mathcal{X}_{\mathcal{H},\gamma}^A$, néanmoins, on suppose que c'est vrai pour $1 - \alpha < \gamma < 1$ (cf M. Dauge, communication privée). Pour discrétiser $\mathcal{X}_{\mathcal{E}}^{A,R}$, on procède de la même façon que pour discrétiser $\mathcal{X}_{\mathcal{E}}^{0,R}$, ce qui est expliqué au paragraphe 10.4. Pour simplifier les notations, on ne différencie par le cas où Ω est convexe (calcul dans $\mathcal{X}_{\mathcal{E}}^0$ et $\mathcal{X}_{\mathcal{H}}^0$) du cas où Ω n'est pas convexe (calcul dans $\mathcal{X}_{\mathcal{E},\gamma}^0$ et $\mathcal{X}_{\mathcal{H},\gamma}^0$).

Soit $\mathcal{X}_{\mathcal{E},k}^{A,R}$ (*resp.* $\mathcal{X}_{\mathcal{H},k}^{A,R}$) l'espace $\mathcal{X}_{\mathcal{E}}^{A,R}$ (*resp.* $\mathcal{X}_{\mathcal{H}}^{A,R}$) discrétisé par les éléments finis de Lagrange P_k :

$$\begin{aligned} \mathcal{X}_{\mathcal{E},k}^{A,R} &:= \{ \mathbf{v}_h \in \mathcal{X}_k \mid \mathbf{v}_h \times \boldsymbol{\nu}|_{\Gamma_C} = 0 \}, \\ \mathcal{X}_{\mathcal{H},k}^{A,R} &:= \{ \mathbf{v}_h \in \mathcal{X}_k \mid \mathbf{v}_h \cdot \boldsymbol{\nu}|_{\Gamma_C} = 0 \}. \end{aligned} \quad (12.53)$$

Par construction, $\mathcal{X}_{\mathcal{E},k}^{A,R} \in \mathbf{H}^1(\Omega)$ (*resp.* $\mathcal{X}_{\mathcal{H},k}^{A,R} \in \mathbf{H}^1(\Omega)$) et est conforme dans $\mathcal{X}_{\mathcal{E}}^{A,R}$ (*resp.* $\mathcal{X}_{\mathcal{H}}^{A,R}$). On va expliquer comment discrétiser la formulation variationnelle augmentée mixte (12.38)-(12.39) (*resp.* (12.44)-(12.45)) par les éléments finis de Taylor-Hood dans $\mathcal{X}_{\mathcal{E},k+1}^{A,R} \times V_k$ (*resp.* $\mathcal{X}_{\mathcal{H},k+1}^{A,R} \times V_k$).

On suppose que Γ_A est un plan, et que la frontière $\overline{\Gamma_C} \cup \overline{\Gamma_A}$ ne contient géométriquement que des points d'arêtes ou de coins.

- Soit I_A l'ensemble des indices des N_A points du bord appartenant à Γ_A . Pour ces points, on ne procédera pas à la pseudo-élimination. On complète alors l'ensemble des indices I_Ω en $I_o := I_\Omega \cup I_A$, et on restreint l'ensemble des indices I_f à $I_f \setminus I_A$.
- Les coins appartenant à $\overline{\Gamma_C} \cup \overline{\Gamma_A}$ sont interprétés comme étant des points d'arêtes de Γ_C . On suppose qu'il y en a $N_{a,AC}$, indicés par $I_{a,AC}$. Il reste donc $N_{c'} := N_c - N_{a,AC}$ coins et on a $N_a + N_{a,AC}$ points d'arêtes. On restreint l'ensemble des indices I_c à $I_{c'} := I_c \setminus I_{a,AC}$, et on complète l'ensemble des indices I_a par $I_a \cup I_{a,AC}$.
- Les points d'arêtes de $\overline{\Gamma_C} \cup \overline{\Gamma_A}$ sont interprétés comme étant des points de face de Γ_C . On suppose qu'il y en a $N_{f,AC}$, indicés par $I_{f,AC}$. Il reste donc $N_{a'} := N_a + N_{a,AC} - N_{f,AC}$. On restreint l'ensemble des indices $I_a \cup I_{a,AC}$: $I_{a'} := (I_a \cup I_{a,AC}) \setminus I_{f,AC}$, et on complète l'ensemble des indices $I_f \setminus I_A$ par $I_{f'} := (I_f \setminus I_A) \cup I_{f,AC}$. Par la suite, on pose : $N_{ac'} := N_{a'} + N_{c'}$ et $I_{ac'} := I_{a'} \cup I_{c'}$. On en déduit que l'espace $\mathcal{X}_{\mathcal{E},k}^{A,R}$, l'espace $\mathcal{X}_{\mathcal{E}}^{A,R}$ discrétisé par les éléments finis de Lagrange P_k est de dimension finie $3N - 2N_{f'} - 3N_{ac'}$ (voir le paragraphe 10.4.1) De même, l'espace $\mathcal{X}_{\mathcal{H},k}^{A,R}$, l'espace $\mathcal{X}_{\mathcal{H}}^{A,R}$ discrétisé par les éléments finis de Lagrange P_k est de dimension finie $3N - N_{f'} - 2N_{a'} - 3N_{c'}$ (voir le paragraphe 16.3.2 de la partie IV).

Soient B_A , et B'_A les bases de $\mathcal{X}_k$ suivantes :

$$\begin{aligned} B_A &:= (\mathbf{v}_i, \alpha)_{i \in I_o \cup I_{ac'}, \alpha \in \{1,2,3\}} \cup (v_i \boldsymbol{\tau}_i^\alpha)_{i \in I_{f'}, \alpha \in \{1,2\}} \cup (v_i \boldsymbol{\nu}_i)_{i \in I_{f'}} \\ B'_A &:= (\mathbf{v}_i, \alpha)_{i \in I_o \cup I_{c'}, \alpha \in \{1,2,3\}} \cup (v_i \boldsymbol{\tau}_i^\alpha)_{i \in I_{f'}, \alpha \in \{1,2\}} \cup (v_i \boldsymbol{\nu}_i)_{i \in I_{f'}} \\ &\quad \cup (v_i \boldsymbol{\tau}_i)_{i \in I_{a'}} \cup (v_i \boldsymbol{\nu}_i^\alpha)_{i \in I_{a'}, \alpha \in \{1,2\}}. \end{aligned} \quad (12.54)$$

On décomposera les éléments de $\mathcal{X}_{\mathcal{E},k}^{A,R}$ dans B_A et ceux de $\mathcal{X}_{\mathcal{E},k}^{A,R}$ dans B'_A . Pour la signification de $(v_i \boldsymbol{\tau}_i)_{i \in I_{a'}} \cup (v_i \boldsymbol{\nu}_i^\alpha)_{i \in I_{a'}, \alpha \in \{1,2\}}$, on revoit le lecteur au paragraphe 16.3.2 de la partie IV : $\boldsymbol{\tau}_i$ est un vecteur directeur de l'arête sur laquelle se trouve M_i , et on crée les vecteurs $\boldsymbol{\nu}_i^\alpha$ de sorte que $(\boldsymbol{\tau}_i, \boldsymbol{\nu}_i^1, \boldsymbol{\nu}_i^2)$ forme une base orthonormée directe.

On reprend la structure par bloc, détaillée au paragraphe 10.4.1 pour le champ électrique et au paragraphe 16.3.2 de la partie IV pour le champ magnétique, en remplaçant I_Ω par I_o , I_f par $I_{f'}$, I_{ac} par $I_{ac'}$, I_a par $I_{a'}$, I_c par $I_{c'}$. Soit $\mathbb{A}_A \in (\mathbb{R}^{3 \times 3})^{N \times N}$ (*resp.* $\mathbb{A}'_A \in (\mathbb{R}^{3 \times 3})^{N \times N}$) la matrice des produits scalaires $\mathcal{X}_{\mathcal{E}(\gamma)}^0$ des éléments de la base B_A .

$$\mathbb{A}_A = \begin{pmatrix} \mathbb{A}_{o,o} & \mathbb{A}_{o,\nu} & 0 \\ \mathbb{A}_{\nu,o} & \mathbb{A}_{\nu,\nu} & 0 \\ 0 & 0 & \mathbb{I}_{ac',ac'} \end{pmatrix}, \quad \mathbb{A}'_A = \begin{pmatrix} \mathbb{A}'_{o,o} & \mathbb{A}'_{o,\tau} & \mathbb{A}'_{o,\tau_a} & 0 \\ \mathbb{A}'_{\tau,o} & \mathbb{A}'_{\tau,\tau} & \mathbb{A}'_{\tau,\tau_a} & 0 \\ \mathbb{A}'_{\tau_a,o} & \mathbb{A}'_{\tau_a,\tau} & \mathbb{A}'_{\tau_a,\tau_a} & 0 \\ 0 & 0 & 0 & \mathbb{I}_{c',c'} \end{pmatrix},$$

où $\mathbb{I}_{ac',ac'}$ est la matrice identité de $(\mathbb{R}^{3 \times 3})^{N_{ac'} \times N_{ac'}}$, et $\mathbb{I}_{c',c'}$ est la matrice identité de $(\mathbb{R}^{3 \times 3})^{N_{c'} \times N_{c'}}$. Nous ne détaillons pas les calculs de ces matrices car ils sont similaires à ceux des matrices $\mathbb{A}_0$ et $\mathbb{A}'_0$. On crée exactement de la même façon $\mathbb{M}_A \in (\mathbb{R}^{3 \times 3})^{N \times N}$ et $\mathbb{M}'_A \in (\mathbb{R}^{3 \times 3})^{N \times N}$, les matrices de

masse des problèmes électrique et magnétique, en remplaçant le produit scalaire dans $\mathcal{X}_{\mathcal{E}(\gamma)}^0$ par le produit scalaire $\mathcal{L}^2(\Omega)$.

Construisons $\mathbb{M}_{\Gamma_A} \in (\mathbb{R}^{3 \times 3})^{N \times N}$, la matrice qui correspond aux produits scalaires $\mathcal{L}_t^2(\Gamma_A)$ des composantes tangentielles sur Γ_A des fonctions de base $(\mathbf{v}_{i,\alpha})_{i \in I_A, \alpha \in \{1,2,3\}}$. On a :

$$\begin{aligned} \mathbb{M}_{\Gamma_A}^{i,j} &= \sum_{F_q \mid M_i, M_j \in F_q} \left(\int_{F_q} v_i v_j d\Sigma (\mathbb{I}_3 - \boldsymbol{\nu}_q \cdot \boldsymbol{\nu}_q^T) \right), \forall i, j \in I_A, \\ &= \delta_{ij} \mathbb{I}_3 \text{ sinon.} \end{aligned}$$

Construisons $\mathbb{C}_A$ et $\mathbb{C}'_A \in (\mathbb{R}^{1 \times 3})^{N_1 \times 3N}$ les matrices des produits mixtes :

- Soit $\mathbb{C}_A \in (\mathbb{R}^{1 \times 3})^{N_1 \times 3N}$ la matrice composée des sous-blocs $\mathbb{C}_A^{i_1,j} \in \mathbb{R}^{1 \times 3}$ tels que :

$$\forall i_1 \in I_1, \forall j \in I_o \quad \mathbb{C}_A^{i_1,j} = \begin{pmatrix} (\partial_1 v_j, u_{i_1})_{0(\gamma)} & (\partial_2 v_j, u_{i_1})_{0(\gamma)} & (\partial_3 v_j, u_{i_1})_{0(\gamma)} \end{pmatrix},$$

$$\forall i_1 \in I_1, \forall j \in I_{ac'} \quad \mathbb{C}_A^{i_1,j} = \begin{pmatrix} 0 & 0 & 0 \end{pmatrix},$$

$$\forall i_1 \in I_1, \forall j \in I_{f'} \quad \mathbb{C}_A^{i_1,j} = \begin{pmatrix} 0 & 0 & (\partial_{\nu_j} v_j, u_{i_1})_{0(\gamma)} \end{pmatrix}.$$

- Soit $\mathbb{C}'_A \in (\mathbb{R}^{1 \times 3})^{N_1 \times 3N}$ la matrice composée des sous-blocs $(\mathbb{C}'_A)^{i_1,j} \in \mathbb{R}^{1 \times 3}$ tels que :

$$\forall i_1 \in I_1, \forall j \in I_o \cup I_{c'} \quad (\mathbb{C}'_A)^{i_1,j} = \mathbb{C}_A^{i_1,j},$$

$$\forall i_1 \in I_1, \forall j \in I_{f'} \quad (\mathbb{C}'_A)^{i_1,j} = \begin{pmatrix} (\partial_{\tau_j^1} v_j, u_{i_1})_{0(\gamma)} & (\partial_{\tau_j^2} v_j, u_{i_1})_{0(\gamma)} & 0 \end{pmatrix},$$

$$\forall i_1 \in I_1, \forall j \in I_{a'} \quad (\mathbb{C}'_A)^{i_1,j} = \begin{pmatrix} (\partial_{\tau_j} v_j, u_{i_1})_{0(\gamma)} & 0 & 0 \end{pmatrix}.$$

12.5.2 Formation des second-membres

Nous avons formé toutes les matrices du membre principal de la discrétisation de (12.38)-(12.39) (*resp.* (12.44)-(12.45)). Il ne nous reste plus qu'à former les second-membres.

Soit $n \in \{0, \dots, N_t\}$. On appelle $\underline{\mathbf{L}}^n \in \mathbb{R}^{3N}$ le vecteur tel que :

- $\forall i \in I_o, \underline{\mathbf{L}}_i^n = (\mathcal{L}^n(v_i \mathbf{e}_1), \mathcal{L}^n(v_i \mathbf{e}_2), \mathcal{L}^n(v_i \mathbf{e}_3))^T$,
- $\forall i \in I_{f'}, \underline{\mathbf{L}}_i^n = (0, 0, \mathcal{L}^n(v_i \boldsymbol{\nu}_i))^T$,
- $\forall i \in I_{ac'}, \underline{\mathbf{L}}_i^n = (0, 0, 0)^T$.

Soit $n \in \{0, \dots, N_t\}$. On appelle $\underline{\mathbf{R}}^{n+1} \in \mathbb{R}^{N_1}$ le vecteur tel que :

- $\forall i \in I_1, \underline{\mathbf{R}}^{n+1} = (\rho^{n+1}, u_{i_1})_{0(\gamma)} / \varepsilon_0$.

Soit $n \in \{0, \dots, N_t\}$. On appelle $\underline{\mathbf{M}}^{n+1/2} \in \mathbb{R}^{3N}$ le vecteur tel que :

- $\forall i \in I_o, \underline{\mathbf{M}}_i^{n+1/2} = (\mathcal{M}^{n+1/2}(v_i \mathbf{e}_1), \mathcal{M}^{n+1/2}(v_i \mathbf{e}_2), \mathcal{M}^{n+1/2}(v_i \mathbf{e}_3))^T$,
- $\forall i \in I_{f'}, \underline{\mathbf{M}}_i^{n+1/2} = (\mathcal{M}^{n+1/2}(v_i \boldsymbol{\tau}_i^1), \mathcal{M}^{n+1/2}(v_i \boldsymbol{\tau}_i^2), 0)^T$,
- $\forall i \in I_{a'}, \underline{\mathbf{M}}_i^{n+1/2} = (\mathcal{M}^{n+1/2}(v_i \boldsymbol{\tau}_i), 0, 0)^T$,
- $\forall i \in I_{c'}, \underline{\mathbf{M}}_i^{n+1/2} = (0, 0, 0)^T$.

12.5.3 Champ électromagnétique

Soit $\underline{\mathbf{E}}^0$ et $\underline{\mathbf{E}}^{-1}$ les décompositions de $\Pi_{k+1}(\mathcal{E}^0)$ et $\Pi_{k+1}(\mathcal{E}^{-1})$ dans la base B_A . De même, on considère $\underline{\mathbf{H}}^{1/2}$ et $\underline{\mathbf{H}}^{-1/2}$ les décompositions de $\Pi_{k+1}(\mathcal{H}^{1/2})$ et $\Pi_{k+1}(\mathcal{H}^{-1/2})$ dans la base B'_A . Pour

$n \in \{1, \dots, N_t + 1\}$, on désigne par $\underline{\mathbf{E}}^n$ (*resp.* $\underline{\mathbf{H}}^{n+1/2}$) la décomposition de $\mathcal{E}_h^n$ (*resp.* $\mathcal{H}_h^{n+1/2}$) dans la base B_A (*resp.* B'_A), on désigne par $\underline{\mathbf{p}}^n \in \mathbb{R}^{N_1}$ la décomposition de p_h (pour le champ électrique ou magnétique) dans la base de V_k .

La FVAM (12.38)-(12.39) discrétisée par un schéma explicite en temps et par les éléments finis de Taylor-Hood P_{k+1} - P_k dans $\mathcal{X}_{\mathcal{E},k}^{0,R} \times V_k$ s'écrit alors : Pour tout $n \in \{0, \dots, N\}$, trouver $(\underline{\mathbf{E}}^{n+1} \times \underline{\mathbf{p}}^{n+1}) \in \mathcal{X}_{\mathcal{E},k+1}^{0,R} \times V_k$ tel que :

$$\begin{aligned} (\mathbb{M}_A + \frac{c \Delta t}{2} \mathbb{M}_{\Gamma_A}) \underline{\mathbf{E}}^{n+1} + \mathbb{C}_A^T \underline{\mathbf{p}}^{n+1} &= \underline{\mathbf{L}}'^n, \\ \mathbb{C}_A \underline{\mathbf{E}}^{n+1} &= \underline{\mathbf{R}}^{n+1}, \end{aligned} \quad (12.55)$$

avec : $\underline{\mathbf{L}}'^n := \Delta t^2 \underline{\mathbf{L}}^n + (2\mathbb{M}_A - c^2 \Delta t^2 \mathbb{A}_A) \underline{\mathbf{E}}^n - (\mathbb{M}_A - \frac{c \Delta t}{2} \mathbb{M}_{\Gamma_A}) \underline{\mathbf{E}}^{n-1}$.

L'algorithme de résolution à chaque pas de temps est le suivant :

- Résoudre $(\mathbb{M}_A + \frac{c \Delta t}{2} \mathbb{M}_{\Gamma_A}) \underline{\mathbf{E}}' = \underline{\mathbf{L}}'^n$;
- Résoudre $(\mathbb{C}_A (\mathbb{M}_A + \frac{c \Delta t}{2} \mathbb{M}_{\Gamma_A})^{-1} \mathbb{C}_A^T) \underline{\mathbf{p}}^{n+1} = \mathbb{C}_A \underline{\mathbf{E}}' - \underline{\mathbf{R}}^{n+1}$;
- Résoudre $(\mathbb{M}_A + \frac{c \Delta t}{2} \mathbb{M}_{\Gamma_A}) \underline{\mathbf{E}}^{n+1} = \underline{\mathbf{L}}'^n - \mathbb{C}_A^T \underline{\mathbf{p}}^{n+1}$.

La FVAM (12.44)-(12.45) discrétisée par un schéma explicite en temps et par les éléments finis de Taylor-Hood P_{k+1} - P_k dans $\mathcal{X}_{\mathcal{H},k+1}^{0,R} \times V_k$ s'écrit alors : Pour tout $n \in \{0, \dots, N\}$, trouver $(\underline{\mathbf{H}}^{n+1} \times \underline{\mathbf{p}}^{n+1}) \in \mathcal{X}_{\mathcal{E},k}^{0,R} \times V_k$ tel que :

$$\begin{aligned} (\mathbb{M}'_A + \frac{c \Delta t}{2} \mathbb{M}'_{\Gamma_A}) \underline{\mathbf{H}}^{n+1/2} + \mathbb{C}'_A^T \underline{\mathbf{p}}^{n+1} &= \underline{\mathbf{M}}'^{n+1/2}, \\ \mathbb{C}'_A \underline{\mathbf{H}}^{n+1} &= 0, \end{aligned} \quad (12.56)$$

avec : $\underline{\mathbf{M}}'^{n+1/2} := \Delta t^2 \underline{\mathbf{M}}^{n+1/2} + (2\mathbb{M}'_A - c^2 \Delta t^2 \mathbb{A}'_A) \underline{\mathbf{H}}^{n+1/2} - (\mathbb{M}'_A - \frac{c \Delta t}{2} \mathbb{M}'_{\Gamma_A}) \underline{\mathbf{H}}^{n-1/2}$.

L'algorithme de résolution à chaque pas de temps est le suivant :

- Résoudre $(\mathbb{M}'_A + \frac{c \Delta t}{2} \mathbb{M}'_{\Gamma_A}) \underline{\mathbf{H}}' = \underline{\mathbf{M}}'^{n+1/2}$;
- Résoudre $(\mathbb{C}'_A (\mathbb{M}'_A + \frac{c \Delta t}{2} \mathbb{M}'_{\Gamma_A})^{-1} \mathbb{C}'_A^T) \underline{\mathbf{p}}^{n+1} = \mathbb{C}'_A \underline{\mathbf{H}}'$;
- Résoudre $(\mathbb{M}'_A + \frac{c \Delta t}{2} \mathbb{M}'_{\Gamma_A}) \underline{\mathbf{H}}^{n+3/2} = \underline{\mathbf{M}}'^{n+1/2} - \mathbb{C}'_A^T \underline{\mathbf{p}}^{n+1}$.

Afin d'avoir un algorithme de calcul rapide, il est intéressant de modifier les matrices de masses (qui sont creuses) pour obtenir des matrices diagonales. En P_1 (la discrétisation se fait alors sans multiplicateur de Lagrange), on a la formule de quadrature suivante, exacte pour f un polynôme P_1 sur $\mathcal{T}_l$:

$$\int_{\mathcal{T}_l} f \partial \Omega = \frac{|\mathcal{T}_l|}{4} \sum_{i=1}^4 f(S_i),$$

où les S_i sont les sommets de $\mathcal{T}_l$. D'après la section 15.5 de la partie IV, on peut construire des matrices de masse diagonales $\tilde{\mathbb{M}}_A$ et $\tilde{\mathbb{M}}_{\Gamma_A}$ pour le champ électrique (*resp.* $\tilde{\mathbb{M}}'_A$ et $\tilde{\mathbb{M}}'_{\Gamma_A}$ pour le champ magnétique) équivalentes à $\mathbb{M}_A$ et $\mathbb{M}_{\Gamma_A}$ (*resp.* $\mathbb{M}'_A$ et $\mathbb{M}'_{\Gamma_A}$). En revanche, pour les éléments finis P_2 , la seule formule de quadrature exacte pour les polynômes P_2 faisant intervenir les noeuds P_2 est la suivante : pour f un polynôme P_2 sur $\mathcal{T}_l$:

$$\int_{\mathcal{T}_l} f \partial\Omega = \frac{|\mathcal{T}_l|}{6} \sum_{i'=1}^6 f(S_{i'}),$$

où les $S_{i'}$ sont les milieux des arêtes de $\mathcal{T}_l$. Ainsi, dans cette formule de quadrature, les poids associés aux sommets de $\mathcal{T}_l$ sont nuls. Ceci a pour conséquence que la matrice diagonale associée à cette formule a des zéros sur la diagonale, et n'est donc pas inversible. Dans [45], G. Cohen propose alors une nouvelle famille d'éléments finis : les éléments $\tilde{P}_2$ tels que : $P_2 \subset \tilde{P}_2 \subset P_4$ (voir le paragraphe 15.3.4 de la partie IV), qui contiennent vingt-trois éléments par tétraèdre (voir la figure 15.3), et dont la formule de quadrature exacte pour un polynôme P_4 sur $\mathcal{T}_l$ est telle que :

$$\int_{\mathcal{T}_l} f \partial\Omega = \sum_{i=1}^{23} \omega_i f(M_i),$$

où les $(M_i)_{i=1,\dots,23}$ sont les points $\tilde{P}_2$ du tétraèdre, et les poids ω_i associés sont tous strictement positifs. Il est important de noter que les formules de quadrature en $\tilde{P}_1$ et en $\tilde{P}_2$ préservent l'ordre de convergence.

12.5.4 Étude de la stabilité

Comme on utilise un schéma explicite, on doit imposer une condition de stabilité (condition de type CFL) sur le pas de temps. Pour obtenir cette condition, on dérive des identités d'énergie à partir des formulations variationnelles discrètes. Dans le cas où ρ et $\mathcal{J}$ sont nuls, on peut réécrire la formulation variationnelle augmentée (FVA) totalement discrétisée sous la forme :

Pour tout n , trouver $\mathcal{E}_h^{n+1} \in \mathcal{X}_{\mathcal{E},k}^{0,R}$ tel que : $\forall \mathcal{F}_h \in \mathcal{X}_{\mathcal{E},k}^{0,R}$,

$$\frac{(\mathcal{E}_h^{n+1} - 2\mathcal{E}_h^n + \mathcal{E}_h^{n-1}, \mathcal{F}_h)_0}{\Delta t^2} + c^2(\mathcal{E}_h^n, \mathcal{F}_h)_{\mathcal{X}_{\mathcal{E},(\gamma)}^0} = 0.$$

Considérons $\mathcal{F}_h = \frac{\mathcal{E}_h^{n+1} - \mathcal{E}_h^{n-1}}{2\Delta t}$, qui correspond à la discrétisation de $\partial_t \mathcal{E}$, qu'on met sous la forme :

$$\frac{\mathcal{E}_h^{n+1} - \mathcal{E}_h^n}{2\Delta t} + \frac{\mathcal{E}_h^n - \mathcal{E}_h^{n-1}}{2\Delta t}.$$

On utilise alors l'identité :

$$(\mathcal{E}_h^{n+1} - 2\mathcal{E}_h^n + \mathcal{E}_h^{n-1}, \mathcal{E}_h^{n+1} - \mathcal{E}_h^n)_0 = \|\mathcal{E}_h^{n+1} - \mathcal{E}_h^n\|_0^2 - \|\mathcal{E}_h^n - \mathcal{E}_h^{n-1}\|_0^2,$$

de sorte que la FVA totalement discrétisée se réécrit :

$$\frac{1}{2\Delta t} \left(\left\| \frac{\mathcal{E}_h^{n+1} - \mathcal{E}_h^n}{\Delta t} \right\|_0^2 - \left\| \frac{\mathcal{E}_h^n - \mathcal{E}_h^{n-1}}{\Delta t} \right\|_0^2 + c^2(\mathcal{E}_h^{n+1}, \mathcal{E}_h^n)_{\mathcal{X}_{\mathcal{E},(\gamma)}^0} - c^2(\mathcal{E}_h^n, \mathcal{E}_h^{n-1})_{\mathcal{X}_{\mathcal{E},(\gamma)}^0} \right) = 0. \quad (12.57)$$

On définit la quantité discrète W^{n+1} , qui a la dimension d'une énergie, par :

$$W^{n+1} = \frac{1}{2} \left\| \frac{\mathcal{E}_h^{n+1} - \mathcal{E}_h^n}{\Delta t} \right\|_0^2 + \frac{c^2}{2} (\mathcal{E}_h^{n+1}, \mathcal{E}_h^n)_{\mathcal{X}_{\mathcal{E},(\gamma)}^0}, \quad (12.58)$$

de sorte que l'identité (12.57) s'écrit :

$$\frac{1}{\Delta t} (W^{n+1} - W^n) = 0,$$

ce qui traduit la conservation d'une énergie. Mais on n'a pas garanti la positivité de l'énergie. Considérons la propriété suivante, valable pour toute forme bilinéaire symétrique $\mathcal{A}$:

$$\begin{aligned} \mathcal{A}(\mathbf{u}, \mathbf{v}) &= \frac{1}{4} (\mathcal{A}(\mathbf{u} + \mathbf{v}, \mathbf{u} + \mathbf{v}) - \mathcal{A}(\mathbf{u} - \mathbf{v}, \mathbf{u} - \mathbf{v})) \\ &= \mathcal{A}\left(\frac{\mathbf{u} + \mathbf{v}}{2}, \frac{\mathbf{u} + \mathbf{v}}{2}\right) - \frac{\Delta t^2}{4} \mathcal{A}\left(\frac{\mathbf{u} - \mathbf{v}}{\Delta t}, \frac{\mathbf{u} - \mathbf{v}}{\Delta t}\right), \end{aligned}$$

appliquée pour le produit scalaire dans $\mathcal{X}_{\mathcal{E},(\gamma)}^0$ pour $\mathbf{u} = \mathcal{E}_h^{n+1}$ et $\mathbf{v} = \mathcal{E}_h^n$. On a alors :

$$2W^{n+1} = \left\| \frac{\mathcal{E}_h^{n+1} - \mathcal{E}_h^n}{\Delta t} \right\|_0^2 + c^2 \left\| \frac{\mathcal{E}_h^{n+1} - \mathcal{E}_h^n}{2} \right\|_{\mathcal{X}_{\mathcal{E},(\gamma)}^0}^2 - \frac{c^2 \Delta t^2}{4} \left\| \frac{\mathcal{E}_h^{n+1} - \mathcal{E}_h^n}{\Delta t} \right\|_{\mathcal{X}_{\mathcal{E},(\gamma)}^0}^2.$$

Considérons λ_{max} , la plus grande valeur propre du problème aux valeurs propres suivants : Trouver $(\lambda, \mathcal{E}_h) \in \mathbb{R}^+ \times \mathcal{X}_{\mathcal{E},k}^{0,R}$ tel que : $\forall \mathcal{F}_h \in \mathcal{X}_{\mathcal{E},k}^{0,R}$,

$$(\mathcal{E}_h, \mathcal{F}_h)_{\mathcal{X}_{\mathcal{E},(\gamma)}^0} = \lambda (\mathcal{E}_h, \mathcal{F}_h)_0.$$

λ_{max} est caractérisée par :

$$\lambda_{max} = \max_{\mathcal{F}_h \in \mathcal{X}_{\mathcal{E},k}^{0,R}} \frac{\|\mathcal{F}_h\|_{\mathcal{X}_{\mathcal{E},(\gamma)}^0}^2}{\|\mathcal{F}_h\|_0^2}.$$

On en déduit l'inégalité suivante :

$$\|\mathcal{E}_h^{n+1} - \mathcal{E}_h^n\|_{\mathcal{X}_{\mathcal{E},(\gamma)}^0}^2 \leq \lambda_{max} \|\mathcal{E}_h^{n+1} - \mathcal{E}_h^n\|_0^2,$$

ce que nous permet d'obtenir que W^{n+1} satisfait l'inégalité :

$$2W^{n+1} \leq \left(1 - c^2 \frac{\Delta t^2}{4} \lambda_{max}\right) \left\| \frac{\mathcal{E}_h^{n+1} - \mathcal{E}_h^n}{\Delta t} \right\|_0^2 + c^2 \left\| \frac{\mathcal{E}_h^{n+1} - \mathcal{E}_h^n}{2} \right\|_{\mathcal{X}_{\mathcal{E},(\gamma)}^0}^2.$$

Ainsi, la condition suffisante de positivité de W^{n+1} est :

$$1 - c^2 \frac{\Delta t^2}{4} \lambda_{max} \geq 0.$$

En pratique, on ne calcule pas λ . On peut montrer que :

$$\lambda_{max} \leq \frac{K^2}{r_h^2},$$

où $K > 0$ est une constante qui dépend de la géométrie du maillage ainsi que du degré des éléments finis; et r_h est la valeur minimale des rayons des sphères inscrites des tétraèdres. On choisit alors le pas de temps Δt de sorte que la condition suffisante de stabilité suivante soit vérifiée :

$$\Delta t < \frac{2 r_h}{c K} . \quad (12.59)$$

C'est une stabilité de type L^2 .

12.6 Présentation du code de calcul

Le code de calcul, implémenté en Fortran 77 a été intégralement réalisé par l'auteur, aidée de P. Ciarlet, Jr. pour certains aspects algorithmiques.

Afin d'optimiser le coût mémoire, nous avons utilisé des structures de stockage de type matrices creuses. Il faut construire les fonctions de multiplications matrice-matrice et matrice-vecteur correspondantes.

Deux types de solveurs ont été choisis :

- Pour les petits systèmes (au plus, des matrices carrées de taille 30×30), nous utilisons l'algorithme du pivot de Gauss,
- Pour les gros systèmes, nous utilisons l'algorithme du gradient conjugué, éventuellement préconditionné.

Pour construire les matrices éléments finis, nous procédons en suivant ces étapes :

- Pour le tétraèdre de référence, on crée les coordonnées barycentriques des degrés de liberté et les fonctions de base associées.
- Dans une boucle sur les tétraèdres, on calcule les fonctions de base P_1 , ce qui nous permet d'obtenir les valeurs des fonctions de base locale et/ou de leurs gradients aux points d'intégration numérique. On remplit ensuite les éléments des matrices.

Une fois les matrices éléments finis créées, il faut éliminer les degrés de liberté du bord supplémentaires, afin de satisfaire la condition aux limites $\mathcal{E} \times \boldsymbol{\nu}|_{\Gamma_C} = 0$ ou $\mathcal{B} \cdot \boldsymbol{\nu}|_{\Gamma_C} = 0$. Pour cela, on procède en sélectionnant pour chaque noeud les 3 lignes et les 3 colonnes qui lui correspondent. Puis on fait un changement de base locale, comme c'est expliqué dans le paragraphe 10.4.1.

Sur la figure 12.1, nous avons représenté schématiquement l'arborescence du code du calcul.

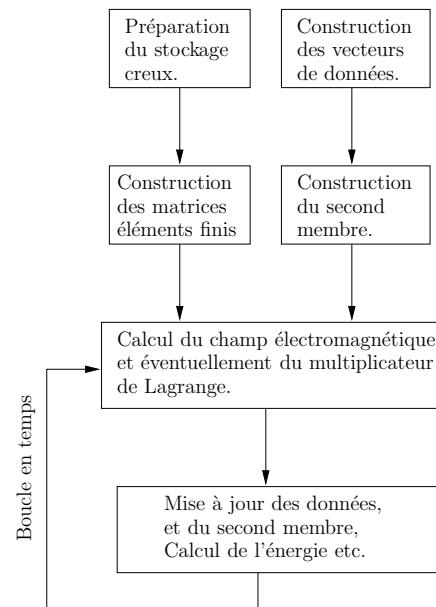


FIG. 12.1 – Schéma du code de calcul.

Chapitre 13

Problème temporel 3D : résultats numériques

13.1 Méthode avec CL essentielles

Pour tester le code, on commence par résoudre le problème académique suivant : une cavité résonnante cubique de longueur $L = 1$ m, c'est-à-dire un domaine Ω en forme de cube dont la frontière $\partial\Omega$ est un conducteur parfait (ce test est proposé par É. Heintzé dans [69]) : Trouver $\mathcal{E} \in \mathcal{X}_{\mathcal{E}}^0$ et $\mathcal{H} \in \mathcal{X}_{\mathcal{H}}^0$ tels que :

$$\begin{aligned} \forall \mathcal{F} \in \mathcal{X}_{\mathcal{E}}^0, \quad \partial_t^2(\mathcal{E}, \mathcal{F})_0 + c^2(\mathcal{E}, \mathcal{F})_{\mathcal{X}_{\mathcal{E}}^0} &= 0, \\ \forall \mathcal{F} \in \mathcal{X}_{\mathcal{H}}^0, \quad \partial_t^2(\mathcal{H}, \mathcal{F})_0 + c^2(\mathcal{H}, \mathcal{F})_{\mathcal{X}_{\mathcal{H}}^0} &= 0, \end{aligned} \tag{13.1}$$

où c est la vitesse de la lumière dans le vide, avec les conditions initiales à divergence et rotationnel nuls ci-dessous :

$$\mathcal{E}_0 = \begin{pmatrix} \cos(\pi x) \sin(\pi y) \sin(-2\pi z) \\ \sin(\pi x) \cos(\pi y) \sin(-2\pi z) \\ \sin(\pi x) \sin(\pi y) \cos(-2\pi z) \end{pmatrix}, \quad \partial_t \mathcal{E}_0 = 0,$$

et :

$$\mathcal{H}_0 = \frac{3\pi}{\mu_0 \omega} \sin(\omega t) \begin{pmatrix} -\sin(\pi x) \cos(\pi y) \cos(-2\pi z) \\ \cos(\pi x) \sin(\pi y) \cos(-2\pi z) \\ 0 \end{pmatrix}, \quad \partial_t \mathcal{H}_0 = 0,$$

où $\omega \approx 2.3$ GHz. Ces conditions initiales vérifient bien les conditions aux limites (12.6) et (12.10) requises, ainsi que les équations de Maxwell dans le vide. La solution exacte est alors la suivante :

$$\mathcal{E}(t) = \cos(\omega t) \begin{pmatrix} \cos(\pi x) \sin(\pi y) \sin(-2\pi z) \\ \sin(\pi x) \cos(\pi y) \sin(-2\pi z) \\ \sin(\pi x) \sin(\pi y) \cos(-2\pi z) \end{pmatrix}, \quad \mathcal{H}(t) = \frac{3\pi}{\mu_0 \omega} \sin(\omega t) \begin{pmatrix} -\sin(\pi x) \cos(\pi y) \cos(-2\pi z) \\ \cos(\pi x) \sin(\pi y) \cos(-2\pi z) \\ 0 \end{pmatrix}.$$

Chaque composante du champ électromagnétique est un mode propre de l'équation des ondes. La période d'oscillations spatiale est de deux mètres dans les directions x et y ; et de un mètre dans la direction z . Le calcul est fait pour un cube de 24 576 tétraèdres, soit 33 noeuds par côté, c'est-à-dire 33 noeuds par période d'oscillations dans les directions x et y ; et 16 noeuds par période d'oscillations dans la direction z , ce qui est supérieur à la limite minimale heuristique de 10 noeuds par période d'oscillations. Notons que les tétraèdres de ce cube sont construits par division de

sous-cubes identiques. On a $r_h \approx 1.4$ cm. Pour observer une oscillation complète, on choisit comme temps final $T_f = 4$ ns. Nous avons fait trois simulations :

- Approximation de la solution de 13.1 par les éléments finis de Lagrange P_1 , notée $(\mathcal{E}_1, \mathcal{B}_1)$.
- Approximation de la solution de 13.1 par les éléments finis de Lagrange $\tilde{P}_1$ (éléments finis P_1 en condensant la matrice de masse), notée $(\mathcal{E}_{\tilde{1}}, \mathcal{B}_{\tilde{1}})$.
- Approximation de la solution de 13.1 par les éléments finis de Lagrange P_2 , notée $(\mathcal{E}_2, \mathcal{B}_2)$.

L'induction magnétique est bien calculée directement, et non par dérivation du champ électrique. Pour les calculs en P_1 , le pas de temps est de l'ordre de 47 ps, ce qui permet de respecter la condition CFL (12.59) et d'avoir plus de 10 pas de temps de discrétisation par période d'oscillations temporelles. Nous avons observé lors des expériences numériques que la condition CFL était plus restrictive pour les éléments finis P_2 que pour les éléments finis P_1 . Ainsi, pour le calcul en P_2 , le pas de temps est de l'ordre de 4.7 ps.

Sur la figure 13.1, nous avons représenté les amplitudes relatives de $E_{1,y}$ (en bleu), $E_{\tilde{1},y}$ (en vert), $E_{2,y}$ (en noir) au point $(0.19, 0.12, 0.12)$ en fonction du temps. Nous les comparons à celle de E_y (en rouge), la composante selon l'axe y de la solution exacte. Notons qu'à cette échelle (précision relative à 0.1%), les courbes de E_y (rouge) et $E_{2,y}$ (noire) sont confondues. Pour la courbe de $E_{1,y}$ (bleue), nous constatons un léger déphasage et une plus grande amplitude par rapport à la courbe de E_y . Ce déphasage s'accroît au cours du temps. Il diminue lorsqu'on utilise les éléments finis $\tilde{P}_1$ (courbe verte), mais la différence d'amplitude persiste. Les résultats pour le champ magnétique sont très similaires. Nous avons aussi testé la conservation de l'énergie discrète : $\frac{W_{n+1} - W_0}{W_0}$.

L'évolution temporelle de cette quantité pour le calcul P_1 est donnée sur la figure 13.2 (à gauche pour $\mathcal{E}_1$ et à droite pour $\mathcal{B}_1$). Comme prévu, cette quantité est négligeable. Sur la figure 13.3, on a représenté l'évolution temporelle de la norme $L^2(\Omega)$ de la divergence (normalisée par W_0) de $\mathcal{E}_1$ (à gauche) et de $\mathcal{B}_1$ (à droite). Cette quantité oscille au cours du temps, mais reste inférieure à 1%. La qualité de cette approximation est liée au fait que notre maillage soit très régulier. Dans ce cas, il n'est pas nécessaire d'utiliser un multiplicateur de Lagrange pour avoir de bons résultats.

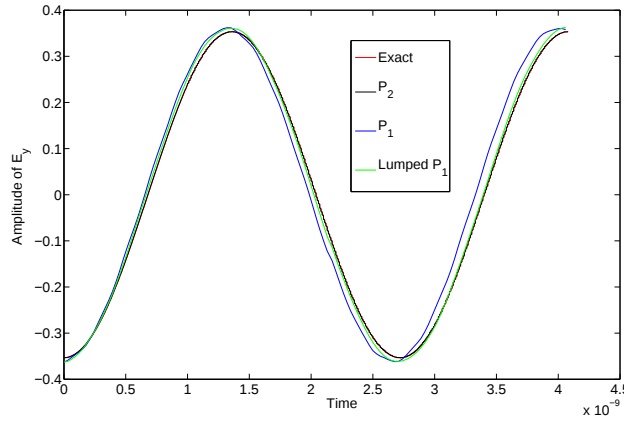


FIG. 13.1 – Amplitude relative de $E_{h,y}$.

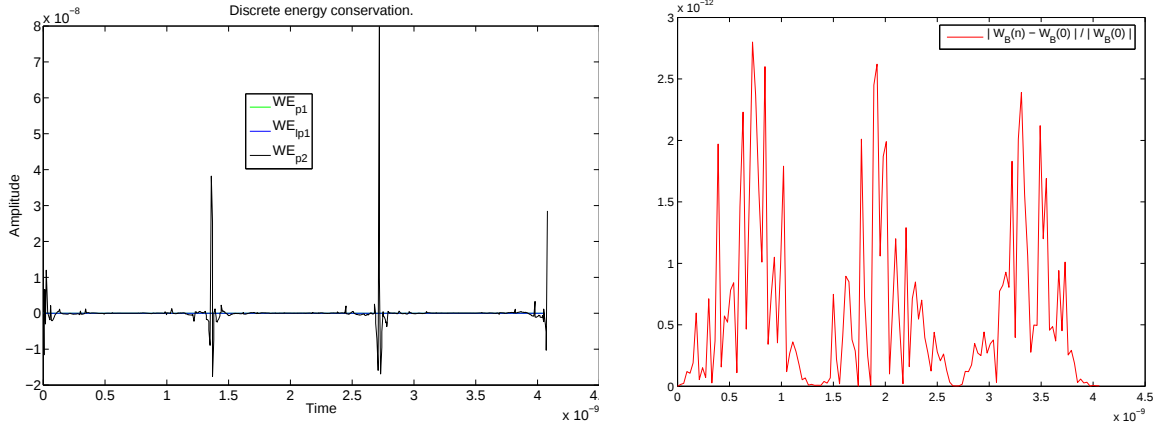


FIG. 13.2 – Évolution temporelle de la conservation de l'énergie discrète.

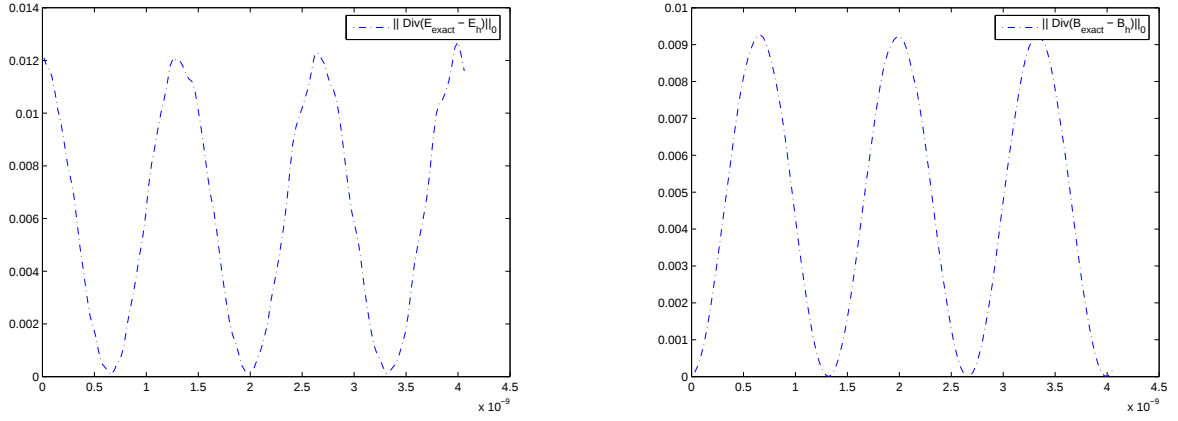


FIG. 13.3 – Évolution temporelle de la divergence.

13.2 Méthode de régularisation à poids

13.2.1 Introduction

Nous allons maintenant présenter un cas test dans un domaine 3D singulier, dont la géométrie est donnée sur la figure 13.4. Le domaine de calcul Ω est un prisme droit présentant une arête rentrante d'angle dièdre $3\pi/2$ (en rouge). Posons r la distance orthogonale à l'arête rentrante. Nous avons pris un poids $\gamma = 0.95$ si $r \leq 1$ et $\gamma = 0$ si $r > 1$ (dans la région où le champ est régulier). Sur la totalité de la frontière artificielle Γ_A (en vert), on impose une condition aux limites de Silver-Müller d'onde absorbante : $\mathbf{e}^* = 0$ et $\mathbf{h}^* = 0$. Une barre de courant traverse le prisme et va générer un champ électromagnétique. Le vecteur densité de courant et la densité de charges circulant dans la barre sont tels que :

$$\mathcal{J} = 10^{-5} \omega \sin\left(\frac{\pi z}{L}\right) \cos(\omega t) \mathbf{e}_z, \quad \rho = 10^{-5} \frac{\pi}{L} \cos\left(\frac{\pi z}{L}\right) \sin(\omega t),$$

avec $\omega \approx 2.5$ GHz. Le courant est d'intensité maximale au milieu du prisme, et il s'annule en ses extrémités. Notons que $\mathcal{J}$ et ρ ici choisis sont très réguliers.

Le champ électrique est solution de la formulation variationnelle augmentée suivante :

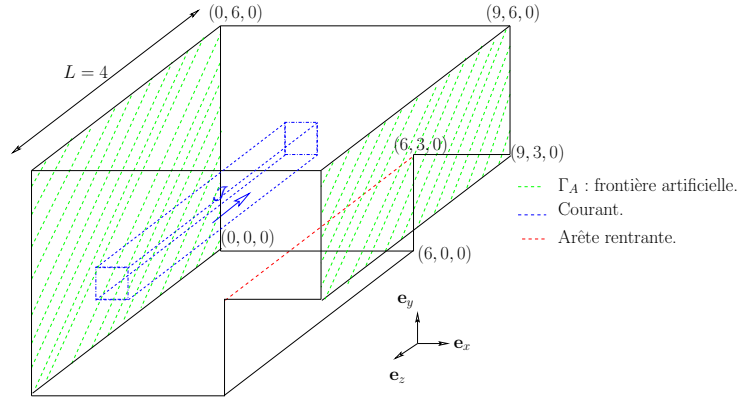


FIG. 13.4 – Cas singulier : géométrie du problème.

Trouver $\mathcal{E} \in \mathcal{X}_{\mathcal{E},\gamma}^A$ tel que $\forall \mathcal{F} \in \mathcal{X}_{\mathcal{E},\gamma}^A$:

$$\frac{d^2}{dt^2}(\mathcal{E}, \mathcal{F})_0 + c^2 (\mathcal{E}, \mathcal{F})_{\mathcal{X}_{\mathcal{E},\gamma}^0} + c \frac{d}{dt} \int_{\Gamma_A} (\mathcal{E} \times \boldsymbol{\nu}) \cdot (\mathcal{F} \times \boldsymbol{\nu}) d\Sigma = - \frac{d}{dt} (\mathcal{J}/\varepsilon_0, \mathcal{F})_0 + c^2 (\rho/\varepsilon_0, \text{div } \mathcal{F})_{0,\gamma},$$

Le champ magnétique est solution de la formulation variationnelle augmentée suivante :

Trouver $\mathcal{H} \in \mathcal{X}_{\mathcal{H},\gamma}^A$ tel que $\forall \mathcal{F} \in \mathcal{X}_{\mathcal{H},\gamma}^A$:

$$\begin{aligned} \frac{d^2}{dt^2}(\mathcal{H}, \mathcal{F})_0 + c^2 (\mathcal{H}, \mathcal{F})_{\mathcal{X}_{\gamma}^0} + c \frac{d}{dt} \int_{\Gamma_A} (\mathcal{H} \times \boldsymbol{\nu}) \cdot (\mathcal{F} \times \boldsymbol{\nu}) d\Sigma \\ = c^2 (\mathcal{J}, \text{rot } \mathcal{F})_0, \end{aligned}$$

Nous résolvons le problème avec les éléments finis de Lagrange $\tilde{P}_1$. Le maillage contient environ 685 000 tétraèdres, et on a $r_h \approx 3$ cm. Le pas de temps est de l'ordre de 40 ps, ce qui permet de respecter la condition CFL (12.59) et d'avoir plus de 10 pas de temps de discrétisation par période d'oscillations temporelles. Le temps d'observation final est de 25 ns, ce qui permet à l'onde générée par la barre de courant d'atteindre les extrémités du domaine (rappelons que le champ se déplace à la vitesse finie $c \approx 3.10^8$ m.s⁻¹). La longueur d'onde associée à la pulsation ω est de l'ordre de : $\lambda = 2\pi c/\omega \approx 75$ cm.

13.2.2 Évolution spatiale

Étudions l'évolution spatiale du champ. Sur les figures 13.5 à 13.10, on représente les valeurs des composantes du champ électrique et de l'induction magnétique dans le plan $z = 2.5$ m, orthogonal à la barre de courant, aux quatre temps suivants : $T_1 = 1$ ns (en haut à gauche), $T_2 = 8$ ns (en haut à droite), $T_3 = 15$ ns (en bas à gauche), $T_4 = 20$ ns (en bas à droite). Au temps T_1 , l'onde électromagnétique vient de se créer autour de la barre de courant, elle se développe et atteint la frontière artificielle Γ_A , en $x = 0$ au temps T_2 . On constate que la longueur d'onde est proche de 75 cm, comme attendu. Au temps T_3 , l'onde est arrivée sur l'arête rentrante, et au temps T_4 , elle a presque atteint l'extrémité de droite de la frontière artificielle Γ_A .

Sur les figures 13.5 et 13.7, on remarque aux temps T_3 et T_4 que les composantes x et y du champ électrique ont un comportement singulier au voisinage du coin rentrant, comme on peut s'y attendre. D'autre part, ce comportement ressemble à celui qu'on observe en 2D (figures 6.3 à 6.8) : les valeurs de la composante $E_{\tilde{\Gamma},x}$ se développent plus la direction x , et les valeurs de la composante

$E_{\tilde{1},x}$ dans la direction y (la partie transverse du champ s'enroule autour de l'arête rentrante). Notons que l'onde est réfléchiée lorsqu'elle atteint le conducteur parfait. On observe aussi des réflexions sur la frontière artificielle. Ce sont des réflexions parasites : il faudrait une condition limite d'ordre supérieur à l'ordre 1 car l'onde n'est pas plane. Sur la figure 13.9, on constate que les valeurs de la composante $E_{\tilde{1},y}$ sont plus basses que celles des deux autres composantes (d'environ un facteur deux), et que $E_{\tilde{1},z}$ est régulier partout (rappelons qu'on a la proposition 8.21). Cela s'explique de la façon suivante : z est la direction de l'arête rentrante, et il n'y a pas de singularité géométrique dans cette direction. Dans cette configuration d'après A. Buffa et al. [31], $E_{\tilde{1},z} \in H^1(\Omega)$.

Sur les figures 13.5 et 13.7, on remarque que le comportement singulier des composantes x et y de l'induction magnétique (*resp.* $B_{\tilde{1},x}$ et $B_{\tilde{1},y}$) existe mais est très discret. Comme le courant est dirigé selon l'axe z , on s'attend à ce que la composante z de l'induction $B_{\tilde{1},z}$ soit nulle. Sur la figure 13.10, on voit que les valeurs sont au plus d'intensité dix fois inférieure à celle des deux autres composantes. Pour le moment, on ne sait pas expliquer que cette composante ait un comportement singulier au voisinage du coin rentrant.

Quant à la forme du champ électromagnétique, on peut l'expliquer formellement avec la loi de Biot et Savart (magnétostatique, voir par exemple [100]). En effet, d'après cette loi, l'induction magnétique au point M créé par la barre de courant B_J est tel que :

$$\mathcal{B}(M) = \frac{\mu_0}{4\pi} \int \int \int_{B_J} \frac{\mathcal{J} \times \vec{P_M}}{|\vec{P_M}|^2} dV,$$

où P décrit le volume B_J . Ainsi, au point M , on a $\mathcal{B}(M) \propto \frac{\mathcal{J} \times P_M \vec{M}}{|P_M \vec{M}|^2}$, où P_M est la projection orthogonale de M sur l'axe de la barre. Soit (x, y, z) les coordonnées de M et (x', y', z) celles de P_M . On en déduit :

$$\mathcal{B}(M) \propto \frac{J_z}{|P_M \vec{M}|^2} \begin{pmatrix} -(y - y') \\ (x - x') \\ 0 \end{pmatrix}.$$

On en déduit alors que pour $y = y'$ (i.e. dans la direction verticale), $B_x \approx 0$ et que pour $x = x'$ (i.e. dans la direction horizontale), $B_y \approx 0$. Sur la figure 13.6, on observe en effet au temps T_1 que $B_{\tilde{1},x}$ est non nulle au dessus et en dessous de la barre de courant, et oscille dans la direction verticale. De même, sur la figure 13.8, on observe en effet au temps T_1 que $B_{\tilde{1},y}$ est non nulle à droite et à gauche de la barre de courant, et oscille dans la direction horizontale.

Le champ électrique est orthogonal à l'induction magnétique, c'est pourquoi on a un comportement orthogonal : sur la figure 13.5, on observe au temps T_1 que $E_{\tilde{1},x}$ est non nulle à droite et à gauche de la barre de courant, et oscille dans la direction horizontale. De même, sur la figure 13.7, on observe en effet au temps T_1 que $E_{\tilde{1},y}$ est non nulle au dessus et en dessous de la barre de courant, et oscille dans la direction verticale. Comme le courant est dans la direction z , $E_{\tilde{1},z}$ n'est pas nul.

Notons que si on n'utilise pas de poids, on obtient le même résultat tant que l'onde n'a pas atteint le coin rentrant. Une fois qu'il est atteint, les résultats diffèrent considérablement car sans poids, on ne capte pas les singularités du champ électromagnétique.

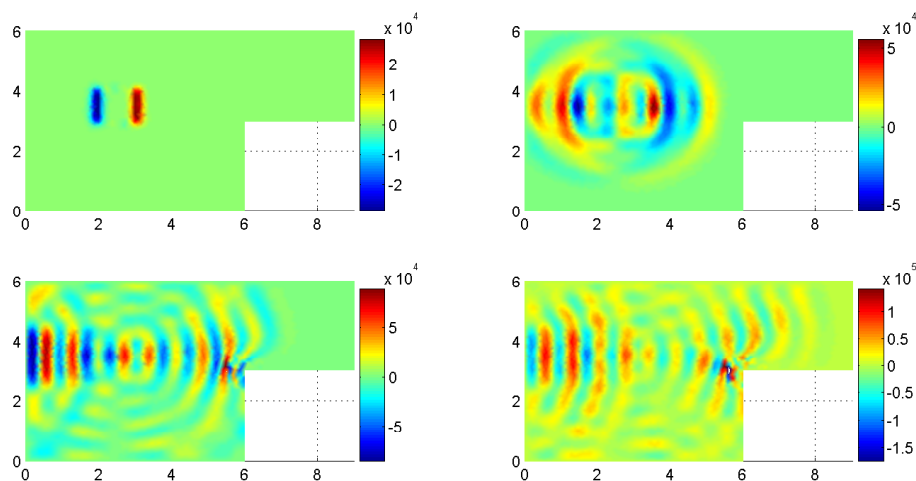


FIG. 13.5 – Composante $E_{1,x}$ aux temps T_i , $i = 1$ à 4, dans le plan $z = 2, 5$ m.

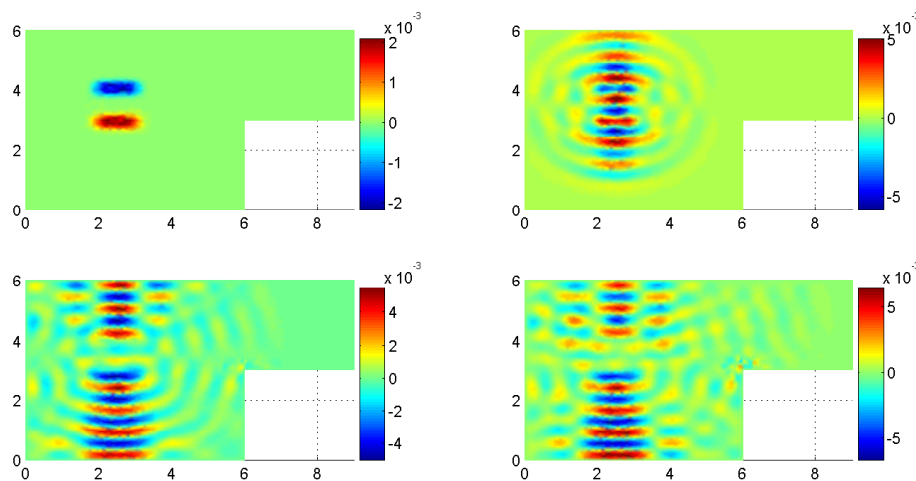


FIG. 13.6 – Composante $B_{1,x}$ aux temps T_i , $i = 1$ à 4, dans le plan $z = 2, 5$ m.

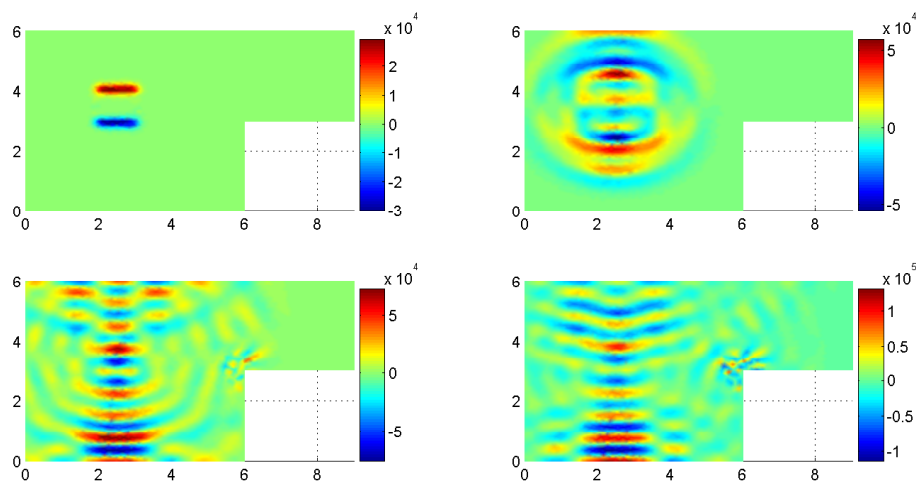


FIG. 13.7 – Composante $E_{1,y}$ aux temps T_i , $i = 1$ à 4 , dans le plan $z = 2, 5$ m.

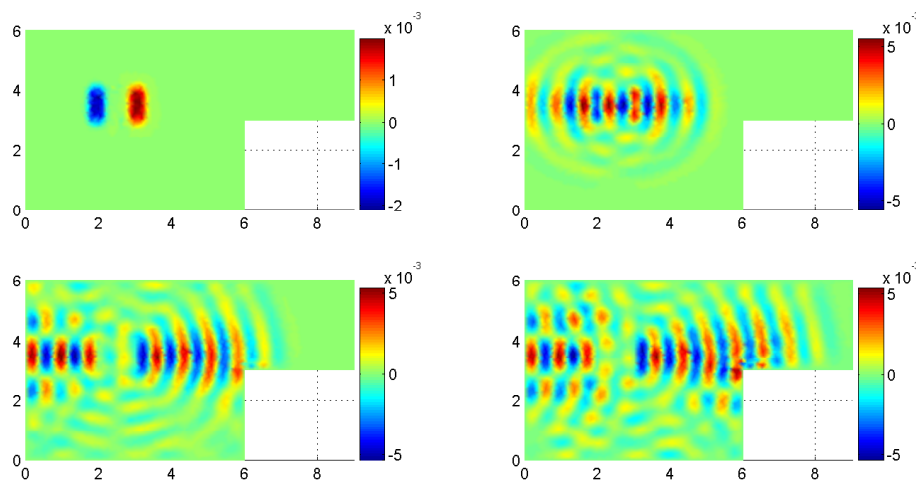


FIG. 13.8 – Composante $B_{1,y}$ aux temps T_i , $i = 1$ à 4 , dans le plan $z = 2, 5$ m.

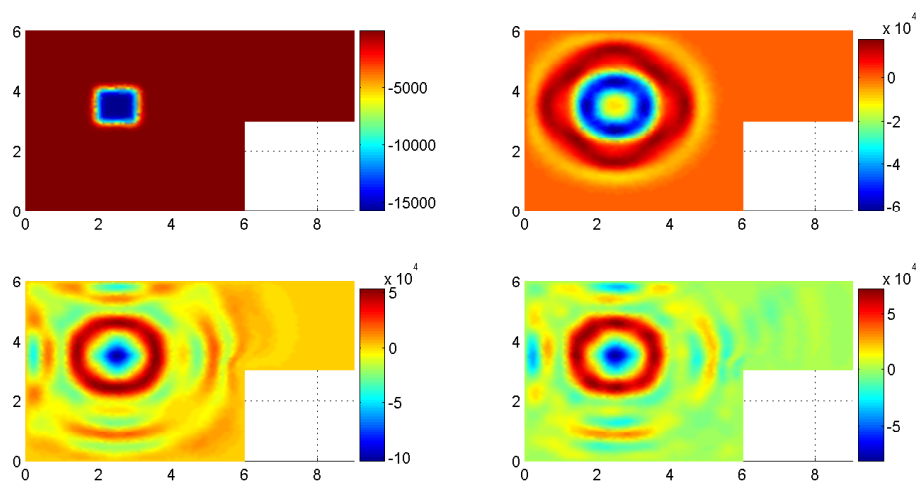


FIG. 13.9 – Composante $E_{1,z}$ aux temps T_i , $i = 1$ à 4, dans le plan $z = 2,5$ m.

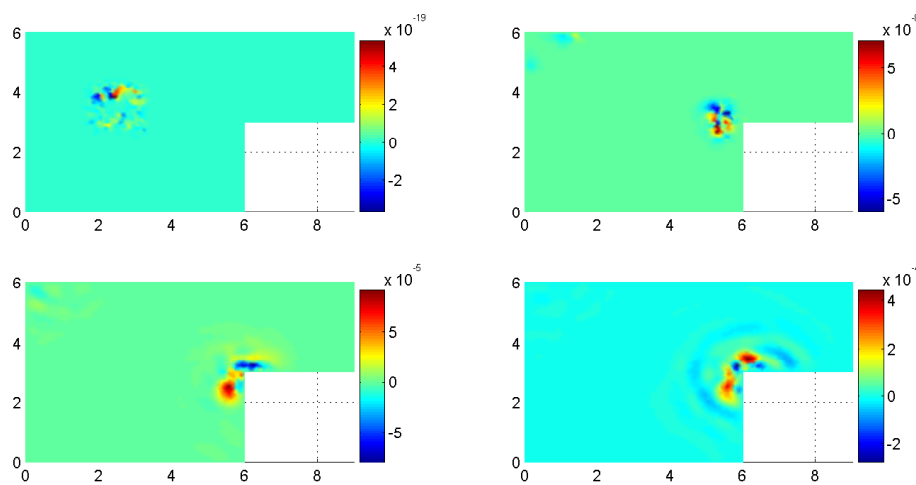


FIG. 13.10 – Composante $B_{1,z}$ aux temps T_i , $i = 1$ à 4, dans le plan $z = 2,5$ m.

13.2.3 Évolution temporelle

Sur la figure 13.12, on a représenté l'évolution temporelle de $E_{1,x}$ au niveau de quatre différents points de Ω , situés dans le plan $z = 2$: $M_1 = (1, 1, 2)$, $M_2 = (1, 5, 2)$, $M_3 = (8, 5.5, 2)$ (figure 13.11). Le champ pris au point considéré est nul tant que l'onde n'a pas atteint ce point. Puis on observe des oscillations forcées, de période de l'ordre de 2.5 ns, ce qui correspond à la période d'oscillations temporelles du courant ($T = 2\pi/\omega \approx 2.5$ ns). On remarque de plus des effets d'interférence entre les ondes créées et réfléchies : certaines oscillations se brisent.

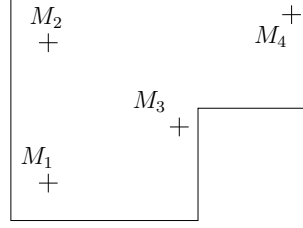


FIG. 13.11 – Localisation des points d'observation.

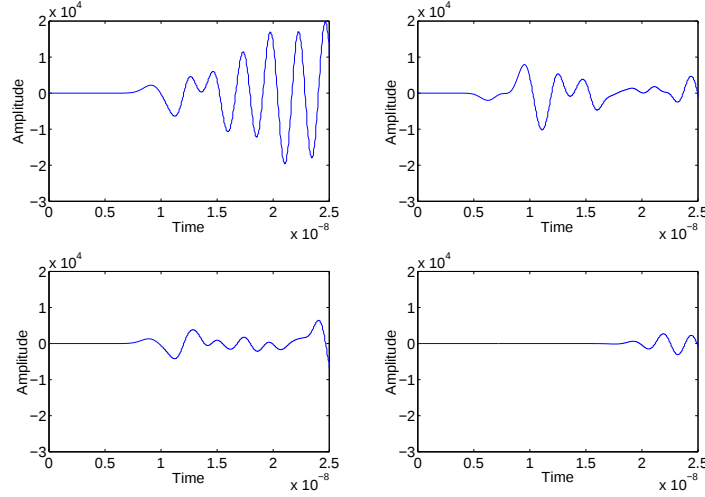


FIG. 13.12 – Composante $E_{1,x}$ aux points M_i , $i = 1$ à 4.

13.3 Conclusions sur les méthodes utilisés

Le calcul du champ électromagnétique instationnaire avec la méthode avec conditions aux limites essentielles est connue depuis les années 90. Cette méthode donne d'excellents résultats lorsque le domaine de calcul est régulier, mais cette méthode n'est pas envisageable pour des domaines singuliers, ce qui bien sûr restreint fortement son utilisation. La méthode de régularisation à poids, inventée par M. Costabel et M. Dauge [53], est en quelque sorte une généralisation de la première méthode aux domaines singuliers, puisque qu'elle coïncide avec celle-ci lorsqu'il n'y a pas de singularité géométrique. Bien que nous n'ayons pas de résultats comparatifs, la méthode de régularisation à poids semble donner des résultats pour le moins corrects en P_1 , pour le champ électrique aussi

bien que pour le champ magnétique. Facile à implémenter, cette méthode est un concurrent sérieux pour les méthodes éléments finis d'arêtes ou de Galerkin discontinus.

Conclusions et perspectives

Nous disposons de deux codes éléments finis continus pour résoudre les équations de Maxwell :

- Un code Matlab, qui propose trois méthodes pour résoudre le problème quasi-électrostatique dans des domaines singuliers par les éléments finis de Lagrange P_k ($k = 1$ ou 2) ou les éléments finis de Taylor-Hood P_2-P_1 . Les trois méthodes sont les suivantes :
 - La méthode avec CL naturelles, qui est simple à implémenter et efficace sur des maillages grossiers ;
 - La méthode du complément singulier, pour laquelle on propose une nouvelle décomposition et une nouvelle façon d'approcher les fonctions de base singulières. Cette méthode a donné lieu à l'article [72] (à lire accompagné du corrigendum [73]), et semble être la plus précise en $2D$ (ce qui peut s'expliquer par la connaissance explicite de la partie la plus singulière du champ) ;
 - La méthode de régularisation à poids, introduite par M. Costabel et M. Dauge pour le calcul de champ électrique en régime harmonique. Cette méthode n'avait pas encore été testée pour le problème quasi-électrostatique mixte. Elle donne de très bons résultats et est assez rapide à coder.
- Un code Fortran 77, qui résout les équations de Maxwell instationnaires dans des domaines réguliers par des éléments finis de Lagrange $P_1, \tilde{P}_1$ ou P_2 (ce qui existait déjà) ; et dans des domaines singuliers par des éléments finis de Lagrange $\tilde{P}_1$ (à notre connaissance, cela n'a jamais été fait). La programmation du $\tilde{P}^2$ (voir [45]) et $\tilde{P}^2-P_1$ est en cours.

D'un point de vue théorique, certains points sont à éclaircir, en particulier :

- Calcul de l'erreur d'approximation de la méthode avec conditions aux limites naturelles ;
- Fondements théoriques de la méthode de régularisation à poids pour le champ magnétique (notons que nous avons tout de même adapté cette méthode au calcul du champ magnétique instationnaire) ;
- Calcul de la condition inf-sup discrète pour la méthode de régularisation à poids mixte ;
- Calcul de l'erreur d'approximation du problème instationnaire (en s'aidant de [44]).

La suite naturelle du travail de thèse est l'exploitation des codes de calculs pour simuler des expériences physiques. Pour cela, il faut approfondir l'étude numérique du problème instationnaire :

- Nécessité du multiplicateur de Lagrange,
- Choix d'un schéma explicite ou implicite en temps,
- Améliorer la prise en compte de conditions aux limites absorbantes en utilisant des couches parfaitement adaptée (en anglais : P.M.L. pour perfectly matched layer) [16], afin de réduire les réflexions parasites,
- Intérêt de combiner l'adaptation du pas de temps et du maillage, comme c'est proposé dans [17] pour l'équation des ondes scalaire,
- Utilisation de polynômes d'ordre élevé, comme c'est fait dans [56] pour la méthode de régularisation à poids en $2D$.

Ainsi, il est utile de compléter l'étude mathématique et numérique des méthodes, afin d'améliorer la performance des codes. D'une part, on peut chercher à améliorer la vitesse de convergence de la méthode avec conditions aux limites naturelle, et de l'adapter au cas instationnaire. D'autre part, l'étude et la programmation de la méthode du complément singulier dans le cas prismatique est une voie qui promet des résultats compétitifs en terme de précision et de coût calcul. De plus, il est

important de comparer les expériences numériques avec d'autres méthodes telles que les volumes finis ou les éléments finis de Galerkin discontinus.

Notons que nous avons présenté l'utilisation des éléments finis de Taylor-Hood P_{k+1} - P_k sans que nous n'en ayons eu vraiment besoin. Mais nous rappelons que pour simuler des interactions champ-particule, il est souhaitable de coupler le code à un code d'interaction de particules, puisqu'on doit résoudre le système d'équations couplées Maxwell-Vlasov. Or dans ce cas, l'équation de continuité de la charge discrète n'est pas vérifiée, et sans multiplicateur de Lagrange, les calculs du champ électromagnétique divergent. De plus, on peut envisager de résoudre le problème harmonique avec les éléments finis de Taylor-Hood P_{k+1} - P_k , afin de tuer les valeurs propres parasites.

Quatrième partie

Annexe

Chapitre 14

Calculs complémentaires pour la MCS

14.1 Calcul de β_D et β_N

Les calculs exposés dans cette section ont été fournis par P. Ciarlet, Jr. [35].

14.1.1 Cas d'un unique coin rentrant

Considérons le cas où ω n'a qu'un seul coin rentrant O . Calculons $\beta_{D,N}$ de la proposition 2.13. Réécrivons $\varphi_{D,N}$ ainsi :

$$\varphi_{D,N} = (\tilde{\varphi}_{D,N} + \beta_{D,N} (1 - \eta) \varphi_{D,N}^P) + \beta_{D,N} \eta \varphi_{D,N}^P,$$

où η est la fonction de troncature définie en 1.1.

Par construction, la fonction $(1 - \eta)$ s'annule dans un voisinage du coin rentrant, de sorte que $(1 - \eta) \varphi_{D,N}^P$ est régulière et $(\tilde{\varphi}_{D,N} + \beta_{D,N} (1 - \eta) \varphi_{D,N}^P)$ est dans $H^2(\omega)$.

D'autre part, le support de η est tel que la trace de $\eta \varphi_D^P$ (*resp.* $\partial_\nu(\eta \varphi_N^P)$) s'annule sur $\partial\omega$. Comme ceci est valable aussi pour $\varphi_{D,N}$, on en déduit que :

$$(\tilde{\varphi}_{D,N} + \beta_{D,N} (1 - \eta) \varphi_{D,N}^P) \in \Phi_{D,N}^R.$$

Ainsi, le Laplacien de cette fonction est orthogonal à toute fonction de $S_{D,N}$, d'où :

$$\begin{aligned} \|s_{D,N}\|_0^2 &= - \int_{\omega} \Delta \varphi_{D,N} s_{D,N} d\omega = -\beta_{D,N} \int_{\omega} \Delta(\eta \varphi_{D,N}) s_{D,N} d\omega \\ &= -\beta_{D,N} \left(\int_{\omega} \Delta(\eta \varphi_{D,N}) \tilde{s}_{D,N} d\omega + \int_{\omega} \Delta(\eta \varphi_{D,N}) s_{D,N}^P d\omega \right). \end{aligned} \quad (14.1)$$

Soient I_1 et I_2 la première et la seconde intégrale du second-membre de 14.1. Pour I_1 , rappelons que $\tilde{s}_{D,N}$ est dans $H^1(\omega)$ et est à Laplacien L^2 : on peut donc appliquer l'intégration par parties (1.13) pour obtenir :

$$\begin{aligned} I_1 &= - \int_{\omega} \mathbf{grad}(\eta \varphi_{D,N}^P) \cdot \mathbf{grad} \tilde{s}_{D,N} d\omega + \int_{\partial\omega} \partial_\nu(\eta \varphi_{D,N}^P) \tilde{s}_{D,N} d\sigma \\ &= \int_{\omega} \eta \varphi_{D,N}^P \Delta \tilde{s}_{D,N} d\omega - \langle \partial_\nu \tilde{s}_{D,N}, \eta \varphi_{D,N}^P \rangle_{H',H} + \int_{\partial\omega} \partial_\nu(\eta \varphi_{D,N}^P) \tilde{s}_{D,N} d\sigma, \end{aligned}$$

avec $H = H^{1/2}(\partial\omega)$.

Le dernier terme est sous forme intégrale grâce à la propriété $\partial_\nu(\eta \varphi_{D,N}^P)|_{\partial\omega} \in L^2(\partial\omega)$.

Comme $\tilde{s}_{D,N}$ est harmonique et que $\eta \tilde{\varphi}_{D|\partial\omega} = 0$ (*resp.* $\partial_\nu \tilde{s}_{N|\partial\omega} = 0$), les deux premiers termes s'annulent. Soit γ_c la portion de $\partial\omega$ définie ainsi :

$$\gamma_c := \{(r, \theta) \in \omega : r \leq \epsilon \text{ et } \theta = 0 \text{ ou } \pi/\alpha\},$$

où ϵ est tel que le support de η est $\omega \cap \mathcal{B}(O, \epsilon)$, où $\mathcal{B}(O, \epsilon)$ est la boule de centre O et de rayon ϵ . Séparons le dernier terme en une intégrale sur γ_c et une intégrale sur $\partial\omega \setminus \gamma_c$.

La première partie s'annule car d'après la proposition 2.10, $\tilde{s}_{D|\gamma_c} = 0$ dans $H^{1/2}(\gamma_c)$ (*resp.* d'après la proposition 2.16, $\partial_\nu(\eta \varphi_N^P)|_{\gamma_c} = 0$ dans $L^2(\gamma_c)$) et le reste est aussi nul car le support de $\eta \varphi_D^P$ est inclus dans $\omega \cap \mathcal{B}(O, \epsilon)$.

Qu'en est-il de I_2 , qui contient la partie principale de $s_{D,N}$?

Considérons $\omega_\epsilon := \omega \setminus \mathcal{B}(O, \epsilon)$. Soit $\Sigma_\epsilon = \partial\omega_\epsilon \cap \omega$, et écrivons l'intégrale I_2 sous forme d'une limite :

$$I_2 = \lim_{\epsilon \rightarrow 0^+} I_2^\epsilon \text{ avec } I_2^\epsilon = \int_{\omega_\epsilon} \Delta(\eta \varphi_{D,N}^P) s_{D,N}^P d\omega.$$

Comme $s_{D,N}^P$ est régulier dans ω_ϵ , on peut appliquer l'intégration par parties (1.13). Comme $s_{D,N}^P$ est régulier sur $\partial\omega_\epsilon$, les crochets de dualité sont remplacés par des intégrales.

$$\begin{aligned} I_2^\epsilon &= \int_{\omega_\epsilon} \eta \varphi_{D,N}^P \Delta s_{D,N}^P d\omega + \int_{\partial\omega_\epsilon} (\partial_\nu(\eta \varphi_{D,N}^P) s_{D,N}^P - \partial_\nu s_{D,N}^P \eta \varphi_{D,N}^P) d\sigma \\ &= \int_{\partial\omega_\epsilon} (\partial_\nu(\eta \varphi_{D,N}^P) s_{D,N}^P - \partial_\nu s_{D,N}^P \eta \varphi_{D,N}^P) d\sigma, \text{ car } s_{D,N}^P \text{ est harmonique} \\ &= \int_{\Sigma_\epsilon} (\partial_\nu(\eta \varphi_{D,N}^P) s_{D,N}^P - \partial_\nu s_{D,N}^P \eta \varphi_{D,N}^P) d\sigma, \text{ grâce à la C. L. sur } \partial\omega_\epsilon \setminus \Sigma_\epsilon \\ &= \int_{\Sigma_\epsilon} (\partial_\nu \varphi_{D,N}^P s_{D,N}^P - \partial_\nu s_{D,N}^P \varphi_{D,N}^P) d\sigma, \text{ car } \eta \equiv 1 \text{ près du coin rentrant.} \end{aligned}$$

La dernière égalité est valable pour $\epsilon \leq \epsilon/2$.

On peut calculer explicitement cette intégrale en remplaçant $\varphi_{D,N}^P$ et $s_{D,N}^P$ par leurs expressions. Notons de plus que $\Sigma_\epsilon = \{(r, \theta) : r = \epsilon, \theta \in]0, \pi/\alpha[\}$, de sorte que $\partial_\nu \equiv -\partial_r$ sur l'interface. En faisant le changement de variable $\theta' = \alpha\theta$, on obtient pour $\|s_D\|_0^2$:

$$I_2^\epsilon = \int_{\theta=0}^{\pi/\alpha} (-\alpha \epsilon^{-1} \sin^2(\alpha\theta) - \alpha \epsilon^{-1} \sin^2(\alpha\theta)) \epsilon d\theta = -2 \int_{\theta'=0}^{\pi} \sin^2(\theta') d\theta' = -\pi.$$

On obtient le même résultat pour $\|s_N\|_0^2$ (il suffit de changer les sinus par des cosinus). D'où :

$$\beta_{D,N} = \frac{1}{\pi} \|s_{D,N}\|_0^2.$$

14.1.2 Cas de plusieurs coins rentrants

Considérons maintenant le cas où il existe plusieurs coins rentrants.

Montrons que $\beta_D^{i,j} = (s_{D,i}, s_{D,j})_0/\pi$. La preuve pour les $\beta_N^{i,j}$ est similaire.

Réécrivons les $\varphi_{D,i}$, $i \in \{1, \dots, N_{cr}\}$ ainsi :

$$\varphi_{D,i} = \left(\tilde{\varphi}_{D,i} + \sum_{j=1}^{N_{cr}} \beta_D^{i,j} (1 - \eta_j) \varphi_{D,j}^P \right) + \sum_{j=1}^{N_{cr}} \beta_D^{i,j} \eta_j \varphi_{D,j}^P.$$

où η_j est la fonction de troncature définie en 1.1.

Par définition des fonctions de troncature, la première partie est dans Φ_D^R . Ainsi son Laplacien est orthogonal à S_D et on a pour tout $j \in \{1, \dots, N_{cr}\}$:

$$\int_{\omega} s_{D,i} s_{D,j} d\omega = - \int_{\omega} \Delta \varphi_{D,i} s_{D,j} d\omega = - \sum_{k=1}^{N_{cr}} \beta_D^{i,k} \int_{\omega} \Delta(\eta_k \varphi_{D,k}^P) s_{D,j} d\omega. \quad (14.2)$$

Soient $\mathbb{P}_D, \mathbb{B}_D, \mathbb{D}_D \in \mathbb{R}^{N_{cr} \times N_{cr}}$ les matrices symétriques telles que : $\forall i, j \in \{1, \dots, N_{cr}\}$:

$$\mathbb{P}_D^{i,j} = \int_{\omega} s_{D,i} s_{D,j} d\omega, \quad \mathbb{B}_D^{i,j} = \beta_D^{i,j}, \quad \mathbb{D}_D^{i,j} = - \int_{\omega} \Delta(\eta_i \varphi_{D,i}^P) s_{D,j} d\omega.$$

L'équation (14.2) s'écrit donc sous la forme matricielle suivante : $\mathbb{P}_D = \mathbb{B}_D \mathbb{D}_D$.

Il s'agit alors de montrer que : $\mathbb{D}_D = \pi \mathbb{I}$, où $\mathbb{I} \in \mathbb{R}^{N_{cr} \times N_{cr}}$ est la matrice identité de $\mathbb{R}^{N_{cr} \times N_{cr}}$. Considérons d'abord les termes diagonaux. Pour j donné, on a :

$$\int_{\omega} \Delta(\eta_j \varphi_{D,j}^P) s_{D,j} d\omega = \int_{\omega} \Delta(\eta_j \varphi_{D,j}^P) \tilde{s}_{D,j} d\omega + \int_{\omega} \Delta(\eta_j \varphi_{D,j}^P) s_{D,j}^P d\omega = I_1 + I_2.$$

Reprenons la preuve donnée dans le paragraphe 14.1.1 précédent. En intégrant deux fois par parties I_1 , on obtient que $I_1 = 0$. Pour la seconde intégrale, on introduit ω_{ε}^i et Σ_{ε}^i , et on écrit que $I_2 = \lim_{\varepsilon \rightarrow 0^+} I_2^{\varepsilon}$, avec :

$$I_2^{\varepsilon} = \int_{\omega_{\varepsilon}^i} \Delta(\eta_i \varphi_{D,i}^P) s_{D,i}^P d\omega = \int_{\Sigma_{\varepsilon}^i} (\partial_{\nu} \varphi_{D,i}^P s_{D,i}^P - \partial_{\nu} s_{D,i}^P \varphi_{D,i}^P) d\sigma,$$

pour $\varepsilon \leq \varepsilon_i/2$. En prenant compte de la forme circulaire de Σ_{ε}^i et des expression explicites des parties principales, on trouve que $I_2^{\varepsilon} = -\pi$, d'où :

$$\mathbb{D}_D^{i,i} = \pi, \quad \forall i \in \{1, \dots, N_{cr}\}.$$

Considérons les coefficients non-diagonaux. Pour $i \neq j$, on a :

$$\int_{\omega} \Delta(\eta_i \varphi_{D,i}^P) s_{D,j} d\omega = \int_{\omega_{\varepsilon_i}} s_{D,j} d\omega,$$

où ω_{ε_i} est le support de η_i . Au voisinage du $i^{\text{ième}}$ coin rentrant, la fonction harmonique $s_{D,j}$ est dans $H^1(\omega_{\varepsilon_i})$. En utilisant les propriétés des éléments de s_D sur la frontière $\partial\omega$, on peut montrer que : $s_{D,j}|_{\partial\omega \cap \partial\omega_{\varepsilon_i}} = 0$ au sens $H^{1/2}(\partial\omega \cap \partial\omega_{\varepsilon_i})$. En intégrant deux fois par parties, on obtient alors :

$$\int_{\omega_{\varepsilon_i}} s_{D,j} d\omega = - \langle \partial_{\nu} s_{D,i}, \eta_i \varphi_{D,i}^P \rangle_{H,H'} + \int_{\partial\omega_{\varepsilon_i} \setminus \partial\omega} \partial_{\nu}(\eta_i \varphi_{D,i}^P) s_{D,i} d\sigma,$$

où ici, $H = H^{1/2}(\partial\omega_{\varepsilon_i})$. Mais les deux termes s'annulent, car par construction, $\eta_i \varphi_{D,i}^P \in H_0^1(\omega_{\varepsilon_i})$ et s'annule dans un voisinage de $\partial\omega_{\varepsilon_i} \setminus \partial\omega$, de sorte que $\partial_{\nu}(\eta_i \varphi_{D,i}^P)|_{\partial\omega_{\varepsilon_i} \setminus \partial\omega} = 0$. On conclut ainsi que :

$$\mathbb{D}_D^{i,j} = 0, \quad \forall i, j \in \{1, \dots, N_{cr}\}, \quad i \neq j.$$

Ce résultat permet d'étendre le découplage de M. Moussaoui [83] au cas où il existe plusieurs coins rentrants.

14.2 Calcul de λ_D et λ_N

Considérons le cas où ω n'a qu'un seul coin rentrant O . Montrons que le coefficients c_D de la décomposition orthogonale de ϕ_D^0 dans Φ_D est tel que :

$$c_D = \frac{\int_{\omega} g^0 s_D d\omega}{\|s_D\|_0^2}.$$

Rappelons que : $\phi_D^0 = \hat{\phi}_D + c_D \varphi_D$, avec $\hat{\phi}_D \in \Phi_D^R$.

On a donc : $\int_{\omega} \Delta \phi_D^0 s_D d\omega = \int_{\omega} \Delta \hat{\phi}_D s_D d\omega + c \int_{\omega} \Delta \varphi_D s_D d\omega$. Comme $\hat{\phi}_D \in \Phi_D^R$, la première intégrale du second membre s'annule. Comme $-\Delta \varphi_D = s_D$ et $-\Delta \phi_D^0 = g^0$, on en déduit :

$$\int_{\omega} g^0 s_D d\omega = c_D \int_{\omega} s_D^2 d\omega$$

d'où le résultat. Or, en décomposant φ_D en partie régulière et partie principale, on obtient :

$$\phi_D^0 = (\hat{\phi}_D + c \tilde{\varphi}_D) + c \beta_D \varphi_D^P,$$

d'où en posant $\lambda_D = c \beta_D$, et $\tilde{\phi}_D = \hat{\phi}_D + c \tilde{\varphi}_D$, on obtient que $\lambda_D = (g^0, s_D)_0/\pi$ et :

$$\phi_D^0 = \tilde{\phi}_D + \lambda_D \varphi_D^P.$$

Généralisons le résultat précédent au cas où il existe plusieurs coins rentrants. On a alors :

$$\phi_D^0 = \hat{\phi}_D + \sum_{i=1}^{N_{cr}} c_D^i \varphi_{D,i}, \text{ avec } \hat{\phi}_D \in \Phi_D^R. \text{ D'où, } \forall j \in \{1, \dots, N_{cr}\} :$$

$$\int_{\omega} \Delta \phi_D^0 s_D^j d\omega = - \int_{\omega} g^0 s_D^j d\omega = - \sum_{i=1}^{N_{cr}} c_D^i \int_{\omega} s_D^i s_D^j d\omega,$$

car $\Delta \hat{\phi}_D$ est orthogonal à s_D^j et $-\Delta \varphi_{D,i} = s_D^i$. Soient λ_D et $\mathbf{c}_D \in \mathbb{R}^{N_{cr}}$ les vecteurs tels que : $\lambda_D^i = (g^0, s_D^i)_0/\pi$ et $\mathbf{c}_D^i = c_D^i$. On doit alors résoudre le système suivant pour obtenir les c_D^i : $\mathbb{B}_D \mathbf{c}_D = \lambda_D$. D'où, $\forall j \in \{1, \dots, N_{cr}\}$:

$$\sum_{i=1}^{N_{cr}} c_D^i \beta_D^{i,j} = (g^0, s_D^j)_0/\pi := \lambda_D^j.$$

Pour le problème de Neumann, la démonstration est similaire.

14.3 Simplification du calcul de λ

14.3.1 Preuve du lemme 2.37

Soit $u \in H_0^1(\mathcal{B}_{\varepsilon})$. Posons $\phi(r, \theta) = r^{-\alpha/2} u(r, \theta)$. On rappelle que les espaces à poids $V_{\gamma}^l(\omega)$ sont définis dans la section 1.4.

D'après [95] (p. 46, thm 1.3), $u \in V_{-1}^0(\mathcal{B}_{\varepsilon})$. On en déduit que $\phi \in V_{\alpha/2-1}^0(\mathcal{B}_{\varepsilon})$. D'autre part, $\mathbf{grad} \phi = \begin{pmatrix} -\alpha/2 r^{-\alpha/2-1} u + r^{-\alpha/2} \partial_r u \\ r^{-\alpha/2-1} \partial_{\theta} u \end{pmatrix}$, en coordonnées polaires par rapport au coin rentrant.

Or $\partial_r u$ et $r^{-1} \partial_\theta u$ sont dans $L^2(\mathcal{B}_\varepsilon)$. D'où $\mathbf{grad} \phi \in V_{\alpha/2}^0(\mathcal{B}_\varepsilon)$. D'où ϕ appartient à l'espace de Sobolev à poids $V_{\alpha/2}^1(\mathcal{B}_\varepsilon)$.

D'après [95] (p. 45, thm. 1.1), on a l'injection continue $V_\gamma^l(\mathcal{B}_\varepsilon) \hookrightarrow W^{l,q}(\mathcal{B}_\varepsilon)$, pour $1 \leq q \leq \frac{4}{2\gamma+2}$. Ainsi ϕ est dans l'espace de Sobolev fractionnaire $W^{1,4/(\alpha+2)}(\mathcal{B}_\varepsilon)$. Déterminons l'espace de Hilbert auquel appartient ϕ .

Pour cela, utilisons un second théorème d'injection continue dû cette fois à P. Grisvard, [66] (p. 27, thm 1.4.4.1) : $W^{s,p}(\mathcal{B}_\varepsilon) \hookrightarrow W^{t,q}(\mathcal{B}_\varepsilon)$, pour $t \leq s$ et $q \geq p$ tels que : $s - 2/p = t - 2/q$. Rappelons que $H^t(\mathcal{B}_\varepsilon) = W^{t,2}(\mathcal{B}_\varepsilon)$. Il s'agit de trouver t tel que : $s = 1$, $p = 4/(\alpha+2)$ et $q = 2$. On obtient : $t = 1 - \alpha/2$. Posons : $\epsilon_\alpha = (1 - \alpha)/2 \in]0, 1/4[$. On en déduit : $W^{1,4/(\alpha+2)}(\mathcal{B}_\varepsilon) \hookrightarrow H^{1/2+\epsilon_\alpha}(\mathcal{B}_\varepsilon)$, ce qui permet de conclure.

Soit $u \in H^1(\mathcal{B}_\varepsilon)$, dont la trace sur $\partial\mathcal{B}_\varepsilon$ s'annule au voisinage du coin rentrant. En d'autres termes : il existe $\varepsilon_0 > 0$ tel que $\forall r \in]0, \varepsilon_0[$, $u(r, 0) = u(r, \pi/\alpha) = 0$. Soit ε_1 un réel tel que $0 < \varepsilon_1 < \varepsilon_0$. Soit $\rho \in C^\infty([0, \varepsilon])$ tel que : $\rho = 1$ sur $[0, \varepsilon_1/2]$ et $\rho = 0$ sur $[\varepsilon_1, \varepsilon]$. On pose $v = \rho u \in H_0^1(\mathcal{B}_\varepsilon)$. On a donc $u = v + (1 - \rho)u$. D'après ce qui précède, $r^{-\alpha/2} v \in H^{1/2+\epsilon_\alpha}(\mathcal{B}_\varepsilon)$. Par ailleurs, comme $r^{-\alpha/2} v \in C^\infty(\mathcal{B}_\varepsilon)$, $r^{-\alpha/2} (1 - \rho)u \in H^1(\mathcal{B}_\varepsilon)$, et ainsi : $r^{-\alpha/2} u \in H^{1/2+\epsilon_\alpha}(\mathcal{B}_\varepsilon)$.

14.3.2 Preuve du lemme 2.38

Soit $\varepsilon > 0$. Soit $\varepsilon_1 > \varepsilon$. Considérons $f \in H'_{\varepsilon_1}$ et $g \in H_{\varepsilon_1}$. On a :

$$\langle f, g \rangle_{H'_\varepsilon, H_\varepsilon} \leq \|f\|_{H'_\varepsilon} \|g\|_{H_\varepsilon}.$$

Or, d'après la définition de H_ε , $\lim_{\varepsilon \rightarrow 0} \|g\|_{H_\varepsilon} = 0$ car le support des intégrales sur $\mathcal{B}_\varepsilon$ tend vers 0.

Rappelons que la norme de H'_ε est la suivante :

$$\|f\|_{H'_\varepsilon} = \sup_{\phi \in H_\varepsilon} \frac{\langle f, \phi \rangle_{H'_\varepsilon, H_\varepsilon}}{\|\phi\|_{H_\varepsilon}}.$$

Pour tout $\phi \in H_\varepsilon$, on définit par $\tilde{\phi} \in H_{\varepsilon_1}$, la fonction égale à ϕ sur $\mathcal{B}_\varepsilon$ et nulle sur $\mathcal{B}_{\varepsilon_1} \setminus \mathcal{B}_\varepsilon$. On a donc :

$$\|f\|_{H'_\varepsilon} = \sup_{\tilde{\phi} \in H_{\varepsilon_1}} \frac{\langle f, \tilde{\phi} \rangle_{H'_{\varepsilon_1}, H_{\varepsilon_1}}}{\|\tilde{\phi}\|_{H_{\varepsilon_1}}} \leq \sup_{\psi \in H_{\varepsilon_1}} \frac{\langle f, \psi \rangle_{H'_{\varepsilon_1}, H_{\varepsilon_1}}}{\|\psi\|_{H_{\varepsilon_1}}} := \|f\|_{H'_{\varepsilon_1}}.$$

Ainsi, $\|f\|_{H'_\varepsilon}$ est bornée lorsque ε tend vers zéro, et on a bien $\lim_{\varepsilon \rightarrow 0} \langle f, g \rangle_{H'_\varepsilon, H_\varepsilon} = 0$.

14.4 Preuve du lemme 4.16

Nous allons prouver le lemme 4.16. Considérons le problème variationnel (4.57). On est dans le cadre du problème (3.1)-(3.2) posé dans [52]. Décomposons $\mathbf{G}$ en une partie régulière et une partie

singulière ainsi : $\mathbf{G} = \hat{\mathbf{G}} + \sum_{i=1}^{N_{cr}} c_G^i \mathbf{x}_i^S$, où $\hat{\mathbf{G}} \in \mathbf{X}_E^{0,R}$ et $\forall i \in \{1, \dots, N_{cr}\}$, $c_G^i \in \mathbb{R}$ est une constante.

En prenant comme fonctions tests les vecteurs $\mathbf{x}_j^S$, $j \in \{1, \dots, N_{cr}\}$ dans (4.57), on obtient que les c_G^i sont solutions du système linéaire suivant : $\pi \sum_{i=1}^{N_{cr}} c_G^i \beta^{i,j} = (\mathbf{f}, \mathbf{x}_j^S)_0$, ce qui nous permet d'écrire la décomposition suivante pour $\mathbf{G}$:

$$\mathbf{G} = \tilde{\mathbf{G}} + \sum_{i=1}^{N_{cr}} \lambda_G^i \mathbf{x}_i^P, \text{ avec } \tilde{\mathbf{G}} \in \mathbf{H}^1(\omega) \text{ et } \lambda_G^i = (\mathbf{f}, \mathbf{x}_i^S)_0 / \pi.$$

$\tilde{\mathbf{G}}$, qui n'est pas dans $\mathbf{X}_{\mathbf{E}}^0$ se met sous la forme : $\tilde{\mathbf{G}} = \tilde{\mathbf{G}}^0 - \sum_{i=1}^{N_{cr}} \lambda_G^i \mathbf{x}_i^R$, avec $\tilde{\mathbf{G}}^0 \in \mathbf{X}_{\mathbf{E}}^{0,R}$. D'après

M. Costabel et M. Dauge [52], $\tilde{\mathbf{G}}^0 \in \mathbf{H}^{2\alpha-\epsilon}(\omega)$, pour tout $\epsilon > 0$ tel que $2\alpha - 1 - \epsilon > 0$. Pour obtenir l'estimation d'erreur sur $\mathbf{G}$, on procède de la même façon que pour montrer le théorème 4.12. Ainsi, l'erreur sur la partie régulière $\tilde{\mathbf{G}}^0$ est en $h^{2\alpha-1-\epsilon}$, et on doit évaluer l'erreur sur le calcul des λ_G^i pour estimer l'erreur sur la partie singulière, et sur le relèvement. Détailler cela.

• Soient $\lambda_{G,h}^i$ les approximations des λ_G^i . Pour les calculer, on doit approcher les $\mathbf{x}_i^S$. On appelle $\mathbf{x}_{i,h}^S$ ces approximations. Montrons d'abord qu'il existe que $\forall \epsilon > 0$ tel que $1 - \alpha - \epsilon > 0$, il existe une constante $C_{\lambda_G, \epsilon} > 0$ telle que $\forall i \in \{1, \dots, N_{cr}\}$:

$$|\lambda_G^i - \lambda_{G,h}^i| < C_{\lambda_G, \epsilon} h^{1-\alpha-\epsilon}. \quad (14.3)$$

En effet, on a :

$$\begin{aligned} |\lambda_G^i - \lambda_{G,h}^i| &= |(\mathbf{f}, \mathbf{x}_i^S)_0 - (\mathbf{f}_h, \mathbf{x}_{i,h}^S)_0|/\pi, \\ &= |(\mathbf{f} - \mathbf{f}_h, \mathbf{x}_i^S)_0 + (\mathbf{f}_h, \mathbf{x}_i^S - \mathbf{x}_{i,h}^S)_0|/\pi, \\ &\leq (||\mathbf{x}_i^S||_0 ||\mathbf{f} - \mathbf{f}_h||_0 + ||\mathbf{f}_h||_0 ||\mathbf{x}_i^S - \mathbf{x}_{i,h}^S||_0)/\pi \text{ (inégalité de Cauchy-Schwarz)}, \\ &\leq C_f h + C_k ||\mathbf{x}_i^S - \mathbf{x}_{i,h}^S||_0, \\ &\leq C_f h + C_k C' ||\mathbf{x}_i^S - \mathbf{x}_{i,h}^S||_{\mathbf{X}_{\mathbf{E}}^0}, \text{ d'après C. Weber [107]}. \end{aligned}$$

où C_f , C_k et C' sont des constantes strictement positives. Quelle est l'approximation d'erreur sur le calcul des $\mathbf{x}_{i,h}^S$? Rappelons que $\forall i \in \{1, \dots, N_{cr}\}$, $\mathbf{x}_i^S = \tilde{\mathbf{x}}_i^S + \sum_{j=1}^{N_{cr}} \beta^{i,j} \mathbf{x}_j^P$, avec $\tilde{\mathbf{x}}_i^S \in \mathbf{H}^{2-\alpha-\epsilon}(\omega)$

(voir le paragraphe 2.4.6). D'après [4], on a l'approximation d'erreur suivante sur le calcul des $\beta^{i,j}$: $\forall i, j \in \{1, \dots, N_{cr}\}$,

$$|\beta^{i,j} - \beta_h^{i,j}| \leq C_\beta h^{2\alpha},$$

où C_β ne dépend que du domaine. Ainsi, en reprenant la démonstration du théorème 4.12 (on peut le faire pour $\varepsilon > 0$ tel que $2 - \alpha - \varepsilon > 1$), on obtient que pour tout $\varepsilon > 0$ tel que $1 - \alpha - \epsilon > 0$, il existe une constante $C_\epsilon > 0$ telle que :

$$||\mathbf{x}_i^S - \mathbf{x}_{i,h}^S||_{\mathbf{X}_{\mathbf{E}}^0} \leq C_\epsilon h^{1-\alpha-\epsilon}. \quad (14.4)$$

On en déduit l'estimation (14.3).

• Posons $\mathbf{z}_G = - \left(\sum_{i=1}^{N_{cr}} \lambda_{G,h}^i \Pi_k(\mathbf{x}_i^P) - \sum_{i=1}^{N_{cr}} \lambda_G^i \mathbf{x}_i^P \right)$. En reprenant la démonstration du lemme 4.15, on a l'estimation suivante : pour tout $\varepsilon > 0$ tel que $1 - \alpha - \epsilon > 0$, il existe une constante $C_{\partial\omega} > 0$ telle que :

$$||\mathbf{z}_G||_{\mathbf{X}_{\mathbf{E}}} \leq C_{\partial\omega} h^{1-\alpha-\epsilon}. \quad (14.5)$$

• Comme $\tilde{\mathbf{G}}^0 \in \mathbf{H}^{2\alpha-\epsilon}(\omega)$, on en déduit que l'erreur d'approximation de $\tilde{\mathbf{G}}^0$ par les éléments finis de Lagrange P_k la suivante : $\forall \epsilon > 0$ tel que $2\alpha - 1 - \epsilon > 0$, il existe une constante $C_{G,\epsilon}$ telle que :

$$||\tilde{\mathbf{G}}^0 - \tilde{\mathbf{G}}_h^0||_{\mathbf{X}_{\mathbf{E}}^0} < C_{G,\epsilon} h^{2\alpha-1-\epsilon}. \quad (14.6)$$

En regroupant les résultats (14.3), (14.5) et (14.6), on en déduit que $\forall \epsilon > 0$ tel que $2\alpha - 1 - \epsilon > 0$ et $\alpha - 1 - \epsilon > 0$, il existe une constante $C'_{G,\epsilon} > 0$ telle que :

$$\begin{aligned} \|\mathbf{G} - \mathbf{G}_h\|_{\mathbf{X}_E} &\leq C_G \|\mathbf{f}\|_0 h^{2\alpha-1-\epsilon} \text{ pour } \alpha \leq 2/3, \\ \|\mathbf{G} - \mathbf{G}_h\|_{\mathbf{X}_E} &\leq C_G \|\mathbf{f}\|_0 h^{1-\alpha-\epsilon} \text{ pour } \alpha \geq 2/3. \end{aligned} \quad (14.7)$$

• Recalculons maintenant l'erreur (14.3). $\forall i \in \{1, \dots, N_{cr}\}$, notons $\mathbf{G}_i^S$ la solution du problème : Trouver $\mathbf{G}_i^S \in \mathbf{X}_E^0$ tel que $\forall \mathbf{F} \in \mathbf{X}_E^0$:

$$\mathcal{A}_0(\mathbf{G}_i^S, \mathbf{F}) = (\mathbf{x}_i^S - \mathbf{x}_{i,h}^S, \mathbf{F})_0.$$

Soit $\mathbf{G}_{i,h}^S$ l'approximation de $\mathbf{G}_i^S$ par les éléments finis de Lagrange P_k . Comme $\mathbf{x}_i^S - \mathbf{x}_{i,h}^S \in \mathbf{L}^2(\omega)$, d'après (14.12), on a alors :

$$\begin{aligned} \|\mathbf{G}_i^S - \mathbf{G}_{i,h}^S\|_{\mathbf{X}_E} &\leq C_{G_i^S} \|\mathbf{x}_i^S - \mathbf{x}_{i,h}^S\|_0 h^{2\alpha-1-\epsilon} \text{ pour } \alpha \leq 2/3, \\ \|\mathbf{G}_i^S - \mathbf{G}_{i,h}^S\|_{\mathbf{X}_E} &\leq C_{G_i^S} \|\mathbf{x}_i^S - \mathbf{x}_{i,h}^S\|_0 h^{1-\alpha-\epsilon} \text{ pour } \alpha \geq 2/3. \end{aligned} \quad (14.8)$$

Or, on remarque que :

$$\begin{aligned} \|\mathbf{x}_i^S - \mathbf{x}_{i,h}^S\|_0^2 &= \mathcal{A}_0(\mathbf{G}_i^S, \mathbf{x}_i^S - \mathbf{x}_{i,h}^S), \\ &= (\mathbf{G}_i^S - \mathbf{G}_{i,h}^S, \mathbf{x}_i^S - \mathbf{x}_{i,h}^S)_{\mathbf{X}_E^0}, \\ &\leq M \|\mathbf{G}_i^S - \mathbf{G}_{i,h}^S\|_{\mathbf{X}_E^0} \|\mathbf{x}_i^S - \mathbf{x}_{i,h}^S\|_{\mathbf{X}_E^0}, \text{ où } M \text{ est la constante de coercivité de } \mathcal{A}_0; \\ &\leq \begin{cases} M C_\epsilon C_{G_i^S} h^{\alpha-\epsilon} \|\mathbf{x}_i^S - \mathbf{x}_{i,h}^S\|_0, & \text{pour } \alpha \leq 2/3, \\ M C_\epsilon C_{G_i^S} h^{2(1-\alpha)-\epsilon} \|\mathbf{x}_i^S - \mathbf{x}_{i,h}^S\|_0, & \text{pour } \alpha \geq 2/3, \end{cases} \quad \text{d'après (14.4) et (14.8).} \end{aligned}$$

Ainsi, $\forall \epsilon > 0$ tel que $2(1-\alpha) - \epsilon > 0$ et $\alpha - \epsilon > 0$, il existe une constante C'_ϵ telle que :

$$\|\mathbf{x}_i^S - \mathbf{x}_{i,h}^S\|_0 \leq \begin{cases} C'_\epsilon h^{\alpha-\epsilon}, & \text{pour } \alpha \leq 2/3, \\ C'_\epsilon h^{2(1-\alpha)-\epsilon}, & \text{pour } \alpha \geq 2/3. \end{cases} \quad (14.9)$$

Nous avons obtenu une meilleure estimation d'erreur que (14.4), qui nous permet d'obtenir que $\forall \epsilon > 0$ tel que $\alpha - \epsilon > 0$ et $2(1-\alpha) - \epsilon > 0$, il existe une constante $C'_{\lambda_G, \epsilon}$ telle que $\forall i \in \{1, \dots, N_{cr}\}$:

$$|\lambda_G^i - \lambda_{G,h}^i| < \begin{cases} C'_{\lambda_G, \epsilon} h^{\alpha-\epsilon}, & \text{pour } \alpha \leq 2/3, \\ C'_{\lambda_G, \epsilon} h^{2(1-\alpha)-\epsilon}, & \text{pour } \alpha \geq 2/3. \end{cases} \quad (14.10)$$

On obtient de cette façon une estimation sur l'erreur $\|\mathbf{z}_G\|_{\mathbf{X}_E}$ meilleure que (14.5) : $\forall \epsilon > 0$ tel que $\alpha - \epsilon > 0$ et $2(1-\alpha) - \epsilon > 0$, il existe une constante $C'_{\partial\omega} > 0$ telle que :

$$\|\mathbf{z}_G\|_{\mathbf{X}_E} \leq \begin{cases} C'_{\partial\omega} h^{\alpha-\epsilon}, & \text{pour } \alpha \leq 2/3, \\ C'_{\partial\omega} h^{2(1-\alpha)-\epsilon}, & \text{pour } \alpha \geq 2/3. \end{cases} \quad (14.11)$$

En regroupant les résultats (14.10), (14.11) et (14.6), on en déduit l'estimation suivante :

$$\begin{aligned} \|\mathbf{G} - \mathbf{G}_h\|_{\mathbf{X}_E} &\leq C_G \|\mathbf{f}\|_0 h^{2\alpha-1-\epsilon} \text{ pour } \alpha \leq 3/4, \\ \|\mathbf{G} - \mathbf{G}_h\|_{\mathbf{X}_E} &\leq C_G \|\mathbf{f}\|_0 h^{2(1-\alpha)-\epsilon} \text{ pour } \alpha \geq 3/4. \end{aligned} \tag{14.12}$$

Remarque 14.1 *L'estimation d'erreur en norme $\mathbf{L}^2(\omega)$ du calcul des $\tilde{\mathbf{x}}_i$, du même ordre que (14.9), est moins bonne que le résultat de E. Garcia [63], qui obtient une estimation en $O(h)$. Néanmoins, nous avons aussi une estimation en norme $\mathbf{X}_E^0$, ce qui n'est pas le cas dans [63].*

Chapitre 15

Calculs du problème discrétisé

15.1 Éléments finis P_k 2D

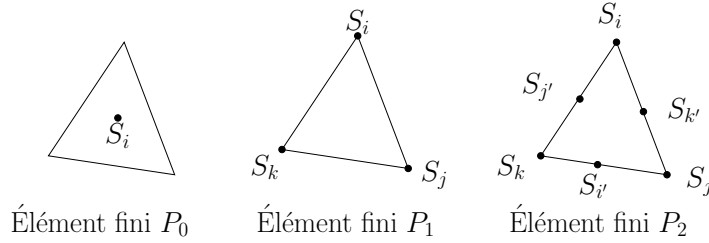


FIG. 15.1 – Éléments finis P_0 , P_1 , P_2 : point(s) de discrétisations par triangle.

15.1.1 Éléments finis P_0

Les points de discrétisation S_i , $i \in I$ sont les barycentres des triangles. Les fonctions de bases sont discontinues, constantes par triangles, telles que :

$$v_i|_{T_l} = \begin{cases} 1 & \text{si } S_i \in \overline{T_l}, \\ 0 & \text{sinon.} \end{cases}$$

15.1.2 Éléments finis P_1

Les points de discrétisation S_i , $i \in I$ sont les sommets des triangles. Il y a donc trois points de discrétisation par triangle. Les fonctions de bases sont continues, affines par triangles, telles que :

$$v_i|_{T_l}(x, y) = a_{i,l}x + b_{i,l}y + c_{i,l}, v_i(S_j) = \delta_{ij}, \forall j.$$

Les coefficients $a_{i,l}$, $b_{i,l}$, $c_{i,l}$ sont nuls si $S_i \notin \overline{T_l}$.

Soit T_l de sommets S_i , S_j , S_k . Soit $\mathbb{S}_3 = \begin{pmatrix} x_i & y_i & 1 \\ x_j & y_j & 1 \\ x_k & y_k & 1 \end{pmatrix}$. Soit $\mathbb{A}_3 = \begin{pmatrix} a_{i,l} & a_{j,l} & a_{k,l} \\ b_{i,l} & b_{j,l} & b_{k,l} \\ c_{i,l} & c_{j,l} & c_{k,l} \end{pmatrix}$.

Alors $\mathbb{S}_3 \mathbb{A}_3 = \mathbb{I}_3$, où $\mathbb{I}_3$ est la matrice identité d'ordre 3. Ainsi, on détermine les coefficients $a_{i,l}$, $b_{i,l}$, $c_{i,l}$ en inversant $\mathbb{S}_3$.

15.1.3 Éléments finis P_2

Les points de discrétisation S_i , $i \in I$ sont les sommets, ainsi que les milieux d'arêtes des triangles : il y a six points de discrétisation par triangle. Les fonctions de bases sont continues, quadratiques par triangles, telles que :

$$v_i|_{T_l}(x, y) = a_{i,l}x^2 + b_{i,l}y^2 + c_{i,l}xy + d_{i,l}x + e_{i,l}y + f_{i,l}, v_i(S_j) = \delta_{ij}, \forall j.$$

Les coefficients $a_{i,l}$, $b_{i,l}$, $c_{i,l}$, $e_{i,l}$, $f_{i,l}$ sont nuls si $S_i \notin \overline{T_l}$. Soit T_l de sommets S_i , S_j , S_k , et dont les milieux des arêtes sont $S_{i'}$, $S_{j'}$, $S_{k'}$. Pour obtenir les coefficients $a_{i,l}$, $b_{i,l}$, $c_{i,l}$, $d_{i,l}$, $e_{i,l}$, $f_{i,l}$ on inverse la matrice $\mathbb{S}_6 \in \mathbb{R}^{6 \times 6}$ telle que :

$$\mathbb{S}_6 = \begin{pmatrix} x_i^2 & y_i^2 & x_i y_i & x_i & y_i & 1 \\ x_j^2 & y_j^2 & x_j y_j & x_j & y_j & 1 \\ x_k^2 & y_k^2 & x_k y_k & x_k & y_k & 1 \\ x_{i'}^2 & y_{i'}^2 & x_{i'} y_{i'} & x_{i'} & y_{i'} & 1 \\ x_{j'}^2 & y_{j'}^2 & x_{j'} y_{j'} & x_{j'} & y_{j'} & 1 \\ x_{k'}^2 & y_{k'}^2 & x_{k'} y_{k'} & x_{k'} & y_{k'} & 1 \end{pmatrix}$$

15.2 Intégration numérique 2D

15.2.1 Schémas d'intégration numérique intérieure

Dans cette section, nous proposons trois schémas d'intégration numérique sur ω , pour un maillage triangulaire :

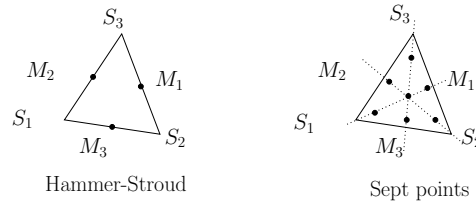


FIG. 15.2 – Localisation des points d'intégration numérique.

- Formule de quadrature exacte pour les polynômes P de degré un :

$$\int_{T_l} P d\omega = \frac{|T_l|}{3} [P(S_i) + P(S_j) + P(S_k)], \quad (15.1)$$

où S_i , S_j et S_k sont les sommets du triangle T_l (voir la figure 15.2).

Coordonnées barycentriques	Multiplicité	Poids $\bar{p}_k$
$(1; 0; 0)$	3	$\frac{1}{3} T_l $

- Schéma d'intégration d'Hammer-Stroud ([58], chap. 12, p. 780), exact pour les polynômes P de degré deux :

$$\int_{T_l} P \, d\omega = \frac{|T_l|}{3} [P(M_i) + P(M_j) + P(M_k)], \quad (15.2)$$

où M_i , M_j et M_k sont les milieux des arêtes du triangles (figure 15.2).

Coordonnées barycentriques	Multiplicité	Poids $\bar{p}_k$
$\left(\frac{1}{2}; \frac{1}{2}; 0\right)$	3	$\frac{1}{3} T_l $

- Schéma d'intégration à sept points (voir la figure 15.2), exact pour des polynômes de degré cinq ([58], chap. 12, p. 781) :

$$\int_{T_l} P \, d\omega = \sum_{k=1}^7 \bar{p}_k P(M_k^l). \quad (15.3)$$

Coordonnées barycentriques	Multiplicité	Poids $\bar{p}_k$
$\left(\frac{1}{3}; \frac{1}{3}; \frac{1}{3}\right)$	1	$\frac{9}{16} T_l $
$\left(\frac{6 - \sqrt{15}}{21}; \frac{6 - \sqrt{15}}{21}; \frac{9 + 2\sqrt{15}}{21}\right)$	3	$\frac{155 - \sqrt{15}}{1200} T_l $
$\left(\frac{6 + \sqrt{15}}{21}; \frac{6 + \sqrt{15}}{21}; \frac{9 - 2\sqrt{15}}{21}\right)$	3	$\frac{155 + \sqrt{15}}{1200} T_l $

15.2.2 Schémas d'intégration numérique sur la frontière

Dans ce paragraphe, nous proposons cette fois deux schémas d'intégrations numérique sur $\partial\omega$:

- Formule de quadrature sur les arêtes, exacte pour les polynômes P de degré un :

$$\int_{A_l} P \, d\sigma = \frac{|A_l|}{2} [P(S_i) + P(S_j)], \quad (15.4)$$

où S_i et S_j sont les extrémités de A_l .

Coordonnées barycentriques	Multiplicité	Poids $\bar{p}_k$
$(0, 1)$	1	$\frac{1}{2} A_l $

- Schéma d'intégration numérique à quatre points, exacte pour les polynômes P de degré trois :

$$\int_{A_l} P \, d\sigma = \sum_{k=1}^4 \bar{p}_k P(M_k^l). \quad (15.5)$$

Coordonnées barycentriques	Multiplicité	Poids $\bar{p}_k$
$(0, 1)$	2	$\frac{1}{8} A_l $
$\left(\frac{1}{3}, \frac{2}{3}\right)$	2	$\frac{3}{8} A_k $

15.3 Éléments finis P_k 3D

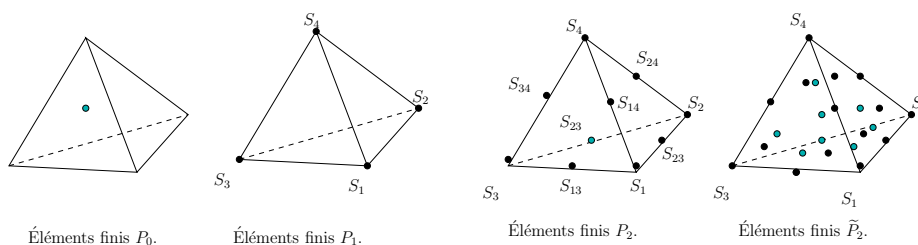


FIG. 15.3 – Éléments finis P_0 , P_1 , P_2 , $\tilde{P}_2$: point(s) de discrétisation par tétraèdre.

15.3.1 Éléments finis P_0

Les points de discrétisation S_i , $i \in I$ sont les barycentres des tétraèdres. Les fonctions de bases sont discontinues, constantes par tétraèdres, telles que :

$$v_i|_{\mathcal{T}_l} = \begin{cases} 1 & \text{si } S_i \in \bar{\mathcal{T}}_l, \\ 0 & \text{sinon.} \end{cases}$$

15.3.2 Éléments finis P_1

Les points de discrétisation S_i , $i \in I$ sont les sommets des tétraèdres. Il y a donc quatre points de discrétisation par tétraèdre. Les fonctions de bases sont continues, affines par triangles, telles

que :

$$v_i|_{T_l}(x, y, z) = a_{i,l}x + b_{i,l}y + c_{i,l}z + d_{i,l}, v_i(S_j) = \delta_{ij}, \forall j.$$

Les coefficients $a_{i,l}, b_{i,l}, c_{i,l}, d_{i,l}$ sont nuls si $S_i \notin \bar{T}_l$.

$$\text{Soit } T_l \text{ de sommets } S_i, S_j, S_k, S_m. \text{ Soit } \mathbb{S}_4 = \begin{pmatrix} x_i & y_i & z_i & 1 \\ x_j & y_j & z_j & 1 \\ x_k & y_k & z_k & 1 \\ x_m & y_m & z_m & 1 \end{pmatrix}.$$

$$\text{Soit } \mathbb{A}_4 = \begin{pmatrix} a_{i,l} & a_{j,l} & a_{k,l} & a_{m,l} \\ b_{i,l} & b_{j,l} & b_{k,l} & b_{m,l} \\ c_{i,l} & c_{j,l} & c_{k,l} & c_{m,l} \\ d_{i,l} & d_{j,l} & d_{k,l} & d_{m,l} \end{pmatrix}. \text{ Alors } \mathbb{S}_4 \mathbb{A}_4 = \mathbb{I}_4, \text{ où } \mathbb{I}_4 \text{ est la matrice identité d'ordre 4.}$$

Ainsi, on détermine les coefficients $a_{i,l}, b_{i,l}, c_{i,l}, d_{i,l}$ en inversant $\mathbb{S}_4$.

15.3.3 Éléments finis P_2

Les points de discrétisation $S_i, i \in I$ sont les sommets, ainsi que les milieux d'arêtes des tétraèdres : il y a dix points de discrétisation par tétraèdres. Les fonctions de bases sont continues, quadratiques par triangles, telles que :

$$v_i|_{T_l}(x, y, z) = a_{i,l}x^2 + b_{i,l}y^2 + c_{i,l}z^2 + d_{i,l}xy + e_{i,l}xz + f_{i,l}yz + g_{i,l}x + h_{i,l}y + i_{i,l}z + j_{i,l}, v_i(S_j) = \delta_{ij}, \forall j.$$

Les coefficients $a_{i,l}$, etc sont nuls si $S_i \notin \bar{T}_l$. Soit T_l de sommets S_i, S_j, S_k, S_m et dont les milieux des arêtes sont $S_{ij}, S_{ik}, S_{im}, S_{jk}, S_{jm}, S_{km}$. Pour obtenir les coefficients $a_{i,l}$ etc, on inverse la matrice $\mathbb{S}_{10} \in \mathbb{R}^{10 \times 10}$ telle que :

$$\mathbb{S}_6 = \begin{pmatrix} x_i^2 & y_i^2 & z_i^2 & x_i y_i & x_i z_i & y_i z_i & x_i & y_i & z_i & 1 \\ x_j^2 & y_j^2 & z_j^2 & x_j y_j & x_j z_j & y_j z_j & x_j & y_j & z_j & 1 \\ x_k^2 & y_k^2 & z_k^2 & x_k y_k & x_k z_k & y_k z_k & x_k & y_k & z_k & 1 \\ x_m^2 & y_m^2 & z_m^2 & x_m y_m & x_m z_m & y_m z_m & x_m & y_m & z_m & 1 \\ x_{ij}^2 & y_{ij}^2 & z_{ij}^2 & x_{ij} y_{ij} & x_{ij} z_{ij} & y_{ij} z_{ij} & x_{ij} & y_{ij} & z_{ij} & 1 \\ x_{ik}^2 & y_{ik}^2 & z_{ik}^2 & x_{ik} y_{ik} & x_{ik} z_{ik} & y_{ik} z_{ik} & x_{ik} & y_{ik} & z_{ik} & 1 \\ x_{im}^2 & y_{im}^2 & z_{im}^2 & x_{im} y_{im} & x_{im} z_{im} & y_{im} z_{im} & x_{im} & y_{im} & z_{im} & 1 \\ x_{jk}^2 & y_{jk}^2 & z_{jk}^2 & x_{jk} y_{jk} & x_{jk} z_{jk} & y_{jk} z_{jk} & x_{jk} & y_{jk} & z_{jk} & 1 \\ x_{jm}^2 & y_{jm}^2 & z_{jm}^2 & x_{jm} y_{jm} & x_{jm} z_{jm} & y_{jm} z_{jm} & x_{jm} & y_{jm} & z_{jm} & 1 \\ x_{km}^2 & y_{km}^2 & z_{km}^2 & x_{km} y_{km} & x_{km} z_{km} & y_{km} z_{km} & x_{km} & y_{km} & z_{km} & 1 \end{pmatrix}$$

15.3.4 Éléments finis $\tilde{P}_2$

Ces éléments finis sont donnés par G. Cohen dans [45]. Les points de discrétisation $S_i, i \in I$ sont les points de discrétisation P_2 , le barycentre du tétraèdre G_l , et les trois points par face suivants :

$$O\vec{G}_{im} = \zeta O\vec{S}_i + \frac{1}{2}(1 - \zeta)O\vec{S}_j + \frac{1}{2}(1 - \zeta)O\vec{S}_k,$$

$$O\vec{G}_{jm} = \zeta O\vec{S}_j + \frac{1}{2}(1 - \zeta)O\vec{S}_i + \frac{1}{2}(1 - \zeta_l)O\vec{S}_k,$$

$$O\vec{G}_{km} = \zeta O\vec{S}_k + \frac{1}{2}(1 - \zeta)O\vec{S}_i + \frac{1}{2}(1 - \zeta_l)O\vec{S}_j,$$

pour la face (S_i, S_j, S_k) , avec $\zeta = \frac{7 - \sqrt{13}}{18}$.

Les fonctions de base associées sont :

- Les fonctions usuelles P_2 pour les éléments P_2 ;
- La fonction bulle $\lambda_1 \lambda_2 \lambda_3 \lambda_4$, où $(\lambda_1, \lambda_2, \lambda_3, \lambda_4)$ sont les coordonnées barycentriques par rapport à (S_1, S_2, S_3, S_4) d'un point de $\mathcal{T}_l$;
- Pour le point G_{im} de la face (S_i, S_j, S_k) , la fonction bulle $\lambda_i^2 \lambda_j \lambda_k$, où $(\lambda_i, \lambda_j, \lambda_k)$ sont les coordonnées barycentriques par rapport à (S_i, S_j, S_k) d'un point de la face (S_i, S_j, S_k) .

On a alors la formule de quadrature suivante associée à ces points de discrétisation et exacte pour les polynômes de degré quatre :

Coordonnées barycentriques	Multiplicité	Poids $\bar{p}_k$
$(1; 0; 0; 0)$	4	$6 \frac{13 - 3\sqrt{13}}{10080} \mathcal{T}_l $
$\left(\frac{1}{2}; \frac{1}{2}; 0; 0\right)$	6	$6 \frac{4 - \sqrt{13}}{315} \mathcal{T}_l $
$\left(\frac{1}{4}; \frac{1}{4}; \frac{1}{4}; \frac{1}{4}\right)$	1	$6 \frac{16}{315} \mathcal{T}_l $
$\left(\zeta; \frac{1}{2}(1 - \zeta); \frac{1}{2}(1 - \zeta); 0\right)$	12	$6 \frac{29 + 17\sqrt{13}}{10080} \mathcal{T}_l $

15.3.5 Schémas d'intégration numérique 3D

Dans cette section, nous proposons trois schémas d'intégration numérique sur Ω , pour un maillage tétraédrique :

- Formule de quadrature exacte pour les polynômes P de degré un :

$$\int_{\mathcal{T}_l} P d\Omega = \frac{|\mathcal{T}_l|}{4} [P(S_i) + P(S_j) + P(S_k) + P(S_m)], \quad (15.6)$$

où S_i, S_j, S_k, S_m sont les sommets du triangles.

Coordonnées barycentriques	Multiplicité	Poids $\bar{p}_k$
$(1; 0; 0; 0)$	4	$\frac{1}{4} \mathcal{T}_l $

- Schéma d'intégration d'Hammer-Stroud ([58], chap. 12, p. 779), exact pour les polynômes P de degré deux :

Coordonnées barycentriques	Multiplicité	Poids $\bar{p}_k$
$\left(\frac{5 - \sqrt{5}}{20}; \frac{5 - \sqrt{5}}{20}; \frac{5 - \sqrt{5}}{20}; \frac{5 + 3\sqrt{5}}{20} \right)$	4	$\frac{1}{4} \mathcal{T}_l $

- Schéma d'intégration à quinze points, exact pour des polynômes P de degré cinq ([58], chap. 12, p. 779) :

Coordonnées barycentriques	Multiplicité	Poids
$(r; r; r)$	1	$A \mathcal{T}_l $
$(s_i; s_i; s_i; t_i), i = 1, 2$	8	$B_i \mathcal{T}_l $
$(u; u; v; v)$	6	$C \mathcal{T}_l $

avec : $r = 1/4$, $s_1 = (7 - \sqrt{15})/34$, $s_2 = (7 + \sqrt{15})/34$, $t_1 = (13 + 3\sqrt{15})/34$, $t_2 = (13 - 3\sqrt{15})/34$, $u = (10 - 2\sqrt{15})/40$, $v = (10 + 2\sqrt{15})/40$;
 $A = 16/135$, $B_1 = (2665 + 14\sqrt{15})/37800$, $B_2 = (2665 - 14\sqrt{15})/37800$ et enfin $C = 20/378$.

Pour les intégrations sur $\partial\Omega$, on se sert des schémas d'intégration numérique $2D$ vus au paragraphe 15.2.1.

Pour les schémas d'intégration $2D$ et $3D$ d'ordre supérieur, nous renvoyons le lecteur à l'ouvrage de P. Solin et al. [102].

15.4 Algorithme du gradient conjugué

Cette section est reprise de [34] (chap. 2). On souhaite résoudre le système linéaire :

$$\text{Trouver } \underline{x} \text{ solution de } \mathbb{K} \underline{x} = \underline{b}, \quad (15.7)$$

où $\mathbb{K} \in \mathbb{R}^{N \times N}$, $\underline{x} \in \mathbb{R}^N$ et $\underline{b} \in \mathbb{R}^N$. On notera $(\cdot | \cdot)$ le produit scalaire dans $\mathbb{R}^N \times \mathbb{R}^N$. Lorsque $\mathbb{K}$ est symétrique, définie-positive, on utilise à cette fin la méthode du gradient conjugué, préconditionné ou non.

15.4.1 Gradient conjugué non-préconditionné

L'algorithme du gradient conjugué (GC) est le suivant :

Soit $\varepsilon > 0$ donné.

Initialisation

$\underline{x}^0$ un vecteur quelconque,

$$\underline{r}^0 = \underline{b} - \mathbb{K} \underline{x}^0,$$

$$\underline{q}^0 = \underline{r}^0.$$

Itérer $k = 0, 1, \dots$

$$\alpha^k = \frac{(\underline{r}^k | \underline{r}^k)}{(\underline{q}^k | \mathbb{K} \underline{q}^k)}$$

$$\underline{x}^{k+1} = \underline{x}^k + \alpha^k \underline{q}^k$$

$$\underline{r}^{k+1} = \underline{r}^k - \alpha^k \mathbb{K} \underline{q}^k$$

$$\beta^k = \frac{(\underline{r}^{k+1} | \underline{r}^{k+1})}{(\underline{r}^k | \underline{r}^k)}$$

$$\underline{q}^{k+1} = \underline{r}^{k+1} + \beta^k \underline{q}^k.$$

jusqu'à $\frac{(\underline{r}^{k+1} | \underline{r}^{k+1})}{(\underline{r}^0 | \underline{r}^0)} < \varepsilon$.

Définitions 15.1 Pour une matrice symétrique, définie-positive, on note $\lambda_{\max}(\mathbb{K})$ sa plus grande valeur propre et $\lambda_{\min}(\mathbb{K})$ sa plus petite valeur propre.

On appelle $\kappa(\mathbb{K})$, le nombre de conditionnement de la matrice $\mathbb{K}$ le rapport entre $\lambda_{\max}(\mathbb{K})$ et $\lambda_{\min}(\mathbb{K})$:

$$\kappa(\mathbb{K}) = \frac{\lambda_{\max}}{\lambda_{\min}}.$$

Proposition 15.2 Posons : $\varepsilon(\underline{x}^k) = (\mathbb{K}(\underline{x}^k - \underline{x}) | \underline{x}^k - \underline{x})$. On a la majoration suivante :

$$\varepsilon(\underline{x}^k) \leq 2 \left(\frac{\sqrt{\kappa(\mathbb{K})} - 1}{\sqrt{\kappa(\mathbb{K})} + 1} \right)^k \varepsilon(\underline{x}^0). \quad (15.8)$$

15.4.2 Gradient conjugué preconditionné

La majoration (15.8) suggère que la méthode converge d'autant plus vite que $\kappa(\mathbb{K})$ est proche de 1. Afin de réduire le nombre d'itérations de la méthode, on multiplie $\mathbb{K}$ par l'inverse d'une matrice inversible $\mathbb{M}$, pour obtenir : $\mathbb{M}^{-1} \mathbb{K} \underline{x} = \mathbb{M}^{-1} \underline{b}$ qui est un système équivalent à (15.7). Pour appliquer l'algorithme du gradient conjugué, $\mathbb{M}$ doit être symétrique, définie-positive. Le système se réécrit en effet sous la forme :

$$\text{Trouver } \underline{y} \text{ solution de } \mathbb{M}^{-1/2} \mathbb{K} \mathbb{M}^{-1/2} \underline{y} = \mathbb{M}^{-1/2} \underline{b} \text{ et } \underline{x} \text{ solution de } \mathbb{M}^{1/2} \underline{x} = \underline{y}. \quad (15.9)$$

La matrice $\mathbb{M}^{-1/2}\mathbb{K}\mathbb{M}^{-1/2}$ étant symétrique, définie-positive, on peut donc utiliser la méthode du gradient conjugué pour résoudre le système en $\underline{y}$. En pratique, on réécrit l'algorithme en utilisant directement le vecteur $\underline{x}$ plutôt que $\underline{y}$ et la matrice $\mathbb{M}$ plutôt que $\mathbb{M}^{1/2}$. On obtient alors *l'algorithme du gradient conjugué préconditionné* (GCP). $\mathbb{M}$ est appelée *matrice de préconditionnement*.

Plus précisément, en appliquant l'algorithme du gradient conjugué à un système linéaire de matrice $\mathbb{M}^{-1/2}\mathbb{K}\mathbb{M}^{-1/2}$, on obtient l'algorithme suivant :

Soit $\varepsilon > 0$ donné.

Initialisation

$\underline{x}^0$ un vecteur quelconque,

$$\underline{r}^0 = \underline{b} - \mathbb{K}\underline{x}^0,$$

$$\mathbb{M}\underline{z}^0 = \underline{r}^0.$$

$$\underline{q}^0 = \underline{z}^0.$$

Boucler $k = 0, 1, \dots$

$$\alpha^k = \frac{(\underline{r}^k | \underline{z}^k)}{(\underline{q}^k | \mathbb{K}\underline{q}^k)}$$

$$\underline{x}^{k+1} = \underline{x}^k + \alpha^k \underline{q}^k$$

$$\underline{r}^{k+1} = \underline{r}^k - \alpha^k \mathbb{K}\underline{q}^k$$

$$\mathbb{M}\underline{z}^{k+1} = \underline{r}^{k+1}$$

$$\beta^k = \frac{(\underline{r}^{k+1} | \underline{z}^{k+1})}{(\underline{r}^k | \underline{z}^k)}$$

$$\underline{q}^{k+1} = \underline{z}^{k+1} + \beta^k \underline{q}^k.$$

jusqu'à $\frac{(\underline{r}^{k+1} | \underline{r}^{k+1})}{(\underline{r}^0 | \underline{r}^0)} < \varepsilon.$

Les matrices $\mathbb{M}^{-1/2}\mathbb{K}\mathbb{M}^{-1/2}$ et $\mathbb{M}^{-1}\mathbb{K}$ sont semblables. Elles ont donc mêmes valeurs propres et même nombre de conditionnement. Le taux de convergence de la méthode du GCP est gouverné par $\kappa(\mathbb{M}^{-1}\mathbb{K})$. Supposons que $\mathbb{M}$ soit facile à inverser et que $\kappa(\mathbb{M}^{-1}\mathbb{K})$ soit beaucoup plus petit que $\kappa(\mathbb{K})$. Alors le coût d'une itération de l'algorithme GCP ne sera pas beaucoup plus élevé que celui d'une itération de l'algorithme du GC. De plus, le taux de convergence sera meilleur. Ainsi, la méthode du GCP est plus efficace.

15.5 Réduction de la matrice de masse

Soit $\Omega \in \mathbb{R}^d$, $d = 2, 3$ un polygone en $2D$ ou un polyèdre en $3D$. Soit $(\mathcal{T}_h)_h$ un famille de triangulation ou une tétraédrisation de Ω . Soit V_h l'espace d'approximation de $H^1(\Omega)$ suivant :

$$V_h = \{v_h \in C^0(\bar{\Omega}), v_h|_K \in P^k(K), \forall K \in \mathcal{T}_h\}.$$

Soient $(S_i)_{i=1,\dots,N}$ les sommets de $\mathcal{T}_h$, et $(\phi_i)_{i=1,\dots,N}$ les fonctions de base associées, telles que : $\forall i, j, \phi_i(S_j) = \delta_{ij}$.

Soit $u_h \in V_h$, alors pour un point $M \in \Omega$, $u_h(M) = \sum_{i=1}^N u_h(S_i) \phi_i(M)$. Posons $u_h(S_i) = x_i$.

On peut représenter u_h sous la forme vectorielle suivante : $\underline{x} = (x_1, \dots, x_N)^T$.

15.5.1 Cas général

Calculons la norme L^2 de u_h :

$$\|u_h\|_0^2 = \left(\sum_{i=1}^N x_i \phi_i, \sum_{j=1}^N x_j \phi_j \right)_0 = \sum_{i=1}^N \sum_{j=1}^N x_i x_j (\phi_i, \phi_j)_0 = (\mathbb{M} \underline{x} | \underline{x}),$$

où $\mathbb{M} \in \mathbb{R}^{N \times N}$ est telle que : $\mathbb{M}^{i,j} = (\phi_i, \phi_j)_0$. $\mathbb{M}$ est appelée matrice de masse.

Considérons l'élément de référence $\widehat{K}$. Cet élément est transformé en l'élément K par la transformation affine suivante : $\widehat{\mathbf{x}} \mapsto \mathbf{x} = \mathbb{B}_K \widehat{\mathbf{x}} + \mathbf{b}_K$, où : $\widehat{\mathbf{x}} = \overrightarrow{O\widehat{M}}$, et $\mathbf{x} = \overrightarrow{OM}$, $\mathbb{B}_K \in \mathbb{R}^{d \times d}$, $\mathbf{b}_K \in \mathbb{R}^d$ (voir la figure 15.4). Soit $\widehat{u}_K$ la fonction telle que : $u_h|_K(M) = \widehat{u}_K(\widehat{M})$.

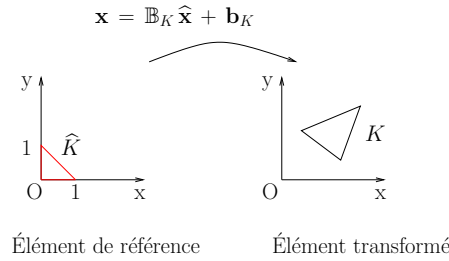


FIG. 15.4 – Formule de représentation en 2D.

On a alors :

$$(\mathbb{M} \underline{x} | \underline{x}) = \sum_{K \in \mathcal{T}_h} \int_K u_h(M)^2 d\mathbf{x} = \sum_{K \in \mathcal{T}_h} \int_{\widehat{K}} \widehat{u}_K(\widehat{M})^2 |\mathbb{B}_K| d\widehat{\mathbf{x}} = \sum_{K \in \mathcal{T}_h} |\mathbb{B}_K| \int_{\widehat{K}} \widehat{u}_K(\widehat{M})^2 d\widehat{\mathbf{x}}, \quad (15.10)$$

où $|\mathbb{B}_K| = \det(\mathbb{B}_K)$. Or, $\widehat{u}_K(\widehat{M})$ dépend linéairement des valeurs de u_h aux sommets de K ,

$(S_k^K)_{k=1,\dots,d+1}$. En effet, notons que : $u_h(M) = \sum_{k=1}^{d+1} \lambda_l^k u_h(S_l^K)$, où λ_l^k est la $k^{\text{ième}}$ coordonnée bary-

centrique de M dans K . Mais ces valeurs sont aussi les valeurs de $\widehat{u}_K$ aux sommets $(\widehat{S}_k)_{k=1,\dots,d+1}$: en d'autres termes, $\widehat{u}_K(\widehat{M})$ appartient à un espace vectoriel sur $\mathbb{R}$ de dimension d . Dans $\mathbb{R}^d$, toutes les normes sont équivalentes, et on a par exemple :

$$\int_{\widehat{K}} \widehat{u}_K(\widehat{M})^2 d\widehat{\mathbf{x}} \simeq \frac{1}{d+1} |\widehat{K}| \sum_{k=1}^{d+1} \widehat{u}_K(\widehat{S}_k)^2, \quad (15.11)$$

où $|\widehat{K}|$ l'aire ou le volume de l'élément de référence $\widehat{K}$.

Il existe donc un couple $(c_d, C_d) \in \mathbb{R}_+^2$ tel que :

$$\forall \widehat{u}_k, \frac{c_d}{d+1} |\widehat{K}| \sum_{k=1}^{d+1} \widehat{u}_k(\widehat{S}_k)^2 \leq \int_{\widehat{K}} \widehat{u}_K(\widehat{M})^2 d\widehat{\mathbf{x}} \leq \frac{C_d}{d+1} |\widehat{K}| \sum_{k=1}^{d+1} \widehat{u}_k(\widehat{S}_k)^2$$

Or, on remarque que : $|K| = \int_K d\mathbf{x} = \int_{\hat{K}} |\mathbb{B}_K| d\hat{\mathbf{x}} = |\mathbb{B}_K| |\hat{K}|$. On injecte (15.11) dans (15.10) pour obtenir :

$$\begin{aligned} (\mathbb{M}\underline{\mathbf{x}}|\underline{\mathbf{x}}) &\simeq \frac{1}{d+1} \sum_{K \in \mathcal{T}_h} |\mathbb{B}_K| |\hat{K}| \sum_{k=1}^{d+1} \hat{u}_K(\hat{S}_k)^2 = \frac{1}{d+1} \sum_{K \in \mathcal{T}_h} |K| \sum_{k=1}^{d+1} u_{h|K}(S_k^K)^2 \\ &\simeq \frac{1}{d+1} \sum_{i=1}^N \left(\sum_{K|S_i \in K} |K| \right) u_h(S_i)^2 = \frac{1}{d+1} \sum_{i=1}^N \left(\sum_{K|S_i \in K} |K| \right) \underline{\mathbf{x}}_i^2. \end{aligned} \quad (15.12)$$

Soit $\tilde{\mathbb{M}} \in \mathbb{R}^{N \times N}$ la matrice diagonale définie par :

$$\tilde{\mathbb{M}}^{i,i} = \frac{1}{d+1} \sum_{K|S_i \in K} |K|. \quad (15.13)$$

On déduit de (15.12) que :

$$(\mathbb{M}\underline{\mathbf{x}}|\underline{\mathbf{x}}) \simeq \frac{1}{d+1} \sum_{i=1}^N \tilde{\mathbb{M}}^{i,i} \underline{\mathbf{x}}_i^2 = \frac{1}{d+1} (\tilde{\mathbb{M}}\underline{\mathbf{x}}|\underline{\mathbf{x}}).$$

D'où le théorème suivant :

Théorème 15.3 *La matrice de masse pleine $\mathbb{M}$ est équivalente à la matrice de masse réduite $\tilde{\mathbb{M}}$, définie par (15.13).*

15.5.2 Triangulation ou tétraèdrisation régulière et quasi-uniforme

Soit h_K le rayon du cercle ou de la sphère circonscrit(e) à l'élément K de la triangulation ou de la tétraèdrisation $\mathcal{T}_h$, et ρ_K le rayon du cercle ou de la sphère inscrit(e). Un calcul direct permet de vérifier que $\rho_K^d \leq |\mathbb{B}_K| \leq h_K^d$. On rappelle que $h = \min_K h_K$. Les définitions suivantes sont données par P. G. Ciarlet dans [33].

Définition 15.4 (i) $(\mathcal{T}_h)_h$ est une famille de triangulations ou de tétraèdrisations régulière si : $\exists \sigma > 0 | \forall h, \forall K \in \mathcal{T}_h, h_K \leq \sigma \rho_K$.

(ii) $(\mathcal{T}_h)_h$ est quasi-uniforme : $\exists c > 0 | \forall h, \forall K \in \mathcal{T}_h, ch \leq h_K$.

Notons que (i) implique qu'il existe un nombre maximal d'éléments N_{max} auquel un sommet S_i appartient, et (ii) implique qu'on ne peut pas trop raffiner une zone de maillage par rapport à une autre.

Soit $\mathcal{T}_h$ satisfaisant l'hypothèse (i). Alors $\forall K \in \mathcal{T}_h, |\mathbb{B}_K| \simeq h_K^d$. Si de plus $\mathcal{T}_h$ satisfait (ii), alors $\forall K \in \mathcal{T}_h, |\mathbb{B}_K| \simeq h^d$ et $\forall i \in \{1, \dots, N\}, |\hat{K}| \sum_{K|S_i \in K} |\mathbb{B}_K| \simeq h^d$, puisque la somme comprend au

plus N_{max} triangles ou tétraèdres d'après (i). Reprenons les équations (15.12). On a :

$$\begin{aligned} (\mathbb{M}\underline{\mathbf{x}}|\underline{\mathbf{x}}) &\simeq \frac{1}{d+1} \sum_{i=1}^N \sum_{K|S_i \in K} |K| \underline{\mathbf{x}}_i^2 = \frac{1}{d+1} \sum_{i=1}^N |\hat{K}| \sum_{K|S_i \in K} |\mathbb{B}_K| \underline{\mathbf{x}}_i^2 \\ &\simeq \frac{h^d}{d+1} \sum_{i=1}^N \underline{\mathbf{x}}_i^2 = \frac{h^d}{d+1} (\underline{\mathbf{x}}|\underline{\mathbf{x}}). \end{aligned} \quad (15.14)$$

Soit $\mathbb{I} \in \mathbb{R}^{N \times N}$ la matrice identité. L'équation (15.14) nous permet d'énoncer le théorème suivant :

Théorème 15.5 *Soit $(\mathcal{T}_h)_h$ une famille de triangulations ou de tétraèdrisations régulière et quasi-uniforme. Alors la matrice de masse pleine $\mathbb{M}$ est équivalente à $\frac{h^d}{d+1} \mathbb{I}$.*

15.6 Réduction de la matrice de masse pondérée en $2D$

Peut-on trouver une matrice de diagonale équivalente à la matrice de masse pleine ? On ne peut pas appliquer systématiquement la relation (15.11) à $u_h(M)r(M)^\gamma$ car ce n'est pas une fonction affine. Il faut trouver une autre formule de quadrature. On a : $\int_K u_h(M)^2 r^{2\gamma} d\mathbf{x} = (\mathbb{M}_\gamma \underline{\mathbf{x}} | \underline{\mathbf{x}})$, où $(\mathbb{M}_\gamma)_{i,j} = (\phi_i, \phi_j)_{0,\gamma}$. D'autre part, on peut écrire : $\int_\omega u_h(M)^2 r^{2\gamma} d\mathbf{x} = \sum_{K \in \mathcal{T}_h} \int_K u_h(M)^2 r^{2\gamma} d\mathbf{x}$, avec des intégrales approchées par un schéma d'intégration numérique à n points :

$$\int_K u_h(M)^2 r^{2\gamma} d\mathbf{x} \simeq |K| \sum_{l=1}^n p_l u_h(M_l^K)^2 r(M_l^K)^{2\gamma}. \quad (15.15)$$

Posons pour tout k : $r^K = \min_l r(M_l^K)$ et $\mathcal{R}^K = \max_l r(M_l^K)$. On a alors :

$$|K| (r^K)^{2\gamma} \sum_{l=1}^n p_l u_h(M_l^K)^2 \leq \int_K u_h(M)^2 r^{2\gamma} d\mathbf{x} \leq |K| (\mathcal{R}^K)^{2\gamma} \sum_{l=1}^n p_l u_h(M_l^K)^2.$$

Numérotons les points du maillage de la façon suivante : S_N est le coin rentrant, pour $1 \leq i \leq N_0$, S_i n'est pas voisin de S_N , et pour $N_0 + 1 \leq i \leq N - 1$, S_i est voisin de S_N . Soit $N_c = N - N_0$, le nombre de sommets voisins du coin rentrant plus le coin reentrant. Soit $\mathcal{K}_c$ l'ensemble des éléments ayant S_N comme sommet, et $\mathcal{K}_0 = \mathcal{T}_h \setminus \mathcal{K}_c$ les autres éléments. Posons $\underline{\mathbf{x}}_0 = (x_1, \dots, x_{N-1}) \in \mathbb{R}^{N-1}$ et $\underline{\mathbf{x}}_c = (x_{N_0+1}, \dots, x_N) \in \mathbb{R}^{N_c}$.

15.6.1 Triangles intérieurs

Soit $K \in \mathcal{K}_0$. Notons que $r^K > 0$. D'autre part, pour une triangulation régulière, $\mathcal{R}^K \simeq r^K + h_K$, et $h_K \leq r^K$, d'où : $\mathcal{R}^K \leq 2r^K$. On obtient alors (par exemple pour le schéma d'intégration à trois points en $2D$ (15.1), $d = 2$ et $n = 3$) :

$$(\tilde{\mathbb{M}}_\gamma^0 \mathbf{x}_0 | \mathbf{x}_0) \leq \sum_{K \in \mathcal{K}_0} \int_K u_h(M)^2 r^{2\gamma} d\mathbf{x} \leq 4 (\tilde{\mathbb{M}}_\gamma^0 \mathbf{x}_0 | \mathbf{x}_0), \quad (15.16)$$

où : $\tilde{\mathbb{M}}_\gamma^0 \in \mathbb{R}^{(N-1) \times (N-1)}$ est la matrice diagonale telle que :

$$\forall i \in \{1, \dots, N-1\}, (\tilde{\mathbb{M}}_\gamma^0)_{i,i} = \sum_{K \in \mathcal{K}_0 | S_i \in K} |K| (r^K)^{2\gamma}.$$

La borne optimale est $2^{2\gamma}$, et comme $0 < \gamma < 1$, $2^{2\gamma} < 4$.

15.6.2 Triangles touchant le coin rentrant

Soit $K \in \mathcal{K}_c$. On ne peut plus choisir le schéma d'intégration à trois points (15.1) car il donne $r^K = 0$. Considérons le schéma d'intégration à trois points d'Hammer-Stroud en $2D$ (formule (15.2) et figure 15.2, $d = 2$ et $n = 3$). On a alors bien : $r^K > 0$. Pour tout l , M_l est le milieu de l'arête opposée au sommet S_l , $p_l = 1/3$, et :

$$\lambda_l^k = \begin{cases} 1/2 & \text{si } k \neq l, \\ 0 & \text{sinon.} \end{cases}$$

Posons provisoirement $u_h^K(S_i^K) = x_i$. On a donc :

$$\sum_{l=1}^n p_l u_h^K(M_l^K)^2 = \frac{|K|}{6} (x_1^2 + x_2^2 + x_3^2 + x_1 x_2 + x_1 x_3 + x_2 x_3). \quad (15.17)$$

Afin de trouver la matrice équivalente à $\mathbb{M}_\gamma$, il faut borner cette somme de produits.

Pour majorer (15.17), il suffit de remarquer que $2ab \leq a^2 + b^2$. Afin de minorer (15.17), notons que $a^2 + b^2 + c^2 + ac + ac + bc \geq (a^2 + b^2 + c^2)/2$, car $(a + b + c)^2 \geq 0$. D'où :

$$\frac{|K|}{12} \sum_{i=1}^3 u_h^K(S_i^K)^2 \leq \sum_{l=1}^n p_l u_h^K(M_l^K)^2 \leq \frac{|K|}{3} \sum_{k=1}^3 u_h^K(S_k^K)^2.$$

Pour les triangles qui touchent le coin rentrant, si la triangulation est régulière, alors : $r^K \simeq h_K$, et $\mathcal{R}^K \simeq 2r_K$, d'où :

$$\frac{1}{4} (\tilde{\mathbb{M}}_\gamma^c \mathbf{x}_c | \mathbf{x}_c) \leq \sum_{K \in \mathcal{K}_c} \int_K u_h(M)^2 r^{2\gamma} d\mathbf{x} \leq 4 (\tilde{\mathbb{M}}_\gamma^c \mathbf{x}_c | \mathbf{x}_c) \quad (15.18)$$

où : $\tilde{\mathbb{M}}_\gamma^c \in \mathbb{R}^{(N_c) \times (N_c)}$ est la matrice diagonale telle que :

$$\forall i \in \{N_c + 1, \dots, N\}, (\tilde{\mathbb{M}}_\gamma^c)_{i,i} = \frac{1}{d+1} \sum_{K \in \mathcal{K}_c | S_i \in K} |K| (h_K)^{2\gamma}.$$

15.6.3 Matrice équivalente

Soit $\tilde{\mathbb{M}}_\gamma \in \mathbb{R}^{N \times N}$ la matrice diagonale telle que :

$$\begin{cases} (\tilde{\mathbb{M}}_\gamma)_{i,i} = (\tilde{\mathbb{M}}_\gamma^0)_{i,i}, \forall i \in \{1, \dots, N_0\}, \\ (\tilde{\mathbb{M}}_\gamma)_{i,i} = (\tilde{\mathbb{M}}_\gamma^0)_{i,i} + (\tilde{\mathbb{M}}_\gamma^c)_{i,i}, \forall i \in \{N_0 + 1, \dots, N-1\}, \\ (\tilde{\mathbb{M}}_\gamma)_{i,i} = (\tilde{\mathbb{M}}_\gamma^c)_{i,i}, \text{ pour } i = N. \end{cases}$$

D'où : $\forall \underline{x} \in \mathbb{R}^N \quad \frac{1}{4} (\tilde{\mathbb{M}}_\gamma \underline{x} | \underline{x}) \leq (\mathbb{M}_\gamma \underline{x} | \underline{x}) \leq 4 (\tilde{\mathbb{M}}_\gamma \underline{x} | \underline{x})$. $\mathbb{M}_\gamma$ est équivalente à $\tilde{\mathbb{M}}_\gamma$.

15.7 Équivalence entre matrice de masse et matrice mixte

Soit $(\underline{E}, \underline{p}) \in \mathbb{R}^N \times \mathbb{R}^M$. On veut résoudre le problème suivant :

$$\begin{cases} \mathbb{A} \underline{E} + \mathbb{C}^T \underline{p} = \underline{f}, \\ \mathbb{C} \underline{E} = \underline{g}, \end{cases} \quad (15.19)$$

où $\mathbb{A} \in \mathbb{R}^{N \times N}$ est symétrique, définie positive ; $\mathbb{C} \in \mathbb{R}^{M \times N}$ est de rang maximal ; $\underline{f} \in \mathbb{R}^N$ et $\underline{g} \in \mathbb{R}^M$. Pour résoudre (15.19), il faut inverser $\mathbb{C} \mathbb{A}^{-1} \mathbb{C}^T \in \mathbb{R}^{M \times M}$.

Lorsque la condition inf-sup discrète est vérifiée, il existe une constante $\kappa_h^* > 0$ telle que :

$$\forall \underline{q} \in \mathbb{R}^M, \exists \underline{E} \in \mathbb{R}^N : \frac{(\mathbb{C}^T \underline{q} | \underline{E})^2}{(\mathbb{A} \underline{E} | \underline{E})} \geq \kappa_h^* (\mathbb{M} \underline{q} | \underline{q}), \quad (15.20)$$

Comme la forme bilinéaire $\mathcal{B}$ (3.2) est continue, il existe une constante $b_h > 0$ telle que :

$$\forall \underline{q} \in \mathbb{R}^M, \forall \underline{E} \in \mathbb{R}^N, (\mathbb{C}^T \underline{q} | \underline{E})^2 \leq b_h^2 (\mathbb{M} \underline{q} | \underline{q}) (\mathbb{A} \underline{E} | \underline{E}). \quad (15.21)$$

Théorème 15.6 *Lorsque la condition inf-sup discrète est vérifiée, $\kappa(\mathbb{M}^{-1}(\mathbb{C}\mathbb{A}^{-1}\mathbb{C}^T)) \leq \frac{b_h}{\kappa_h^{*2}}$.*

DÉMONSTRATION. Montrons le théorème 15.6. Considérons (15.20) et (15.21).

Effectuons les changements de variable : $\underline{\mathbf{q}}' = \mathbb{M}^{1/2}\underline{\mathbf{q}}$, et $\underline{\mathbf{F}}' = \mathbb{A}^{1/2}\underline{\mathbf{F}}$. On obtient :

$$-(\mathbb{M}\underline{\mathbf{q}}|\underline{\mathbf{q}}) = (\mathbb{M}^{1/2}\underline{\mathbf{q}}|\mathbb{M}^{-1/2}\underline{\mathbf{q}}) = (\underline{\mathbf{q}}'|\underline{\mathbf{q}}') = |\underline{\mathbf{q}}'|^2,$$

$$-(\mathbb{A}\underline{\mathbf{F}}|\underline{\mathbf{F}}) = (\mathbb{A}^{1/2}\underline{\mathbf{F}}|\mathbb{A}^{-1/2}\underline{\mathbf{F}}) = (\underline{\mathbf{F}}'|\underline{\mathbf{F}}') = |\underline{\mathbf{F}}'|^2,$$

$$-(\mathbb{C}^T\underline{\mathbf{q}}|\underline{\mathbf{F}}) = (\mathbb{C}^T\mathbb{M}^{-1/2}\underline{\mathbf{q}}|\mathbb{A}^{-1/2}\underline{\mathbf{F}}) = (\mathbb{A}^{-1/2}\mathbb{C}^T\mathbb{M}^{-1/2}\underline{\mathbf{q}}|\underline{\mathbf{F}}') = (\mathbb{L}\underline{\mathbf{q}}'|\underline{\mathbf{F}}'),$$

avec $\mathbb{L} = \mathbb{A}^{-1/2}\mathbb{C}^T\mathbb{M}^{-1/2} \in \mathbb{R}^{N \times M}$. On déduit de (15.20) que :

$$\forall \underline{\mathbf{q}}' \in \mathbb{R}^M, \exists \underline{\mathbf{F}}' \in \mathbb{R}^N : (\mathbb{L}\underline{\mathbf{q}}'|\underline{\mathbf{F}}')^2 \geq \kappa_h^{*2} |\underline{\mathbf{q}}'|^2 |\underline{\mathbf{F}}'|^2.$$

On déduit :

$$|||\mathbb{L}||| := \inf_{\underline{\mathbf{q}}' \in \mathbb{R}^M} \sup_{\underline{\mathbf{F}}' \in \mathbb{R}^N} \frac{(\mathbb{L}\underline{\mathbf{q}}'|\underline{\mathbf{F}}')}{|\underline{\mathbf{q}}'| |\underline{\mathbf{F}}'|} \geq \kappa_h^*, \quad (15.22)$$

où $|||\mathbb{L}|||$ est par définition la norme de $\mathbb{L}$ dans $\mathbb{R}^{N \times M}$.

En utilisant les mêmes changements de variables sur $\underline{\mathbf{F}}$ et $\underline{\mathbf{q}}$ on obtient de (15.21) que :

$$\forall \underline{\mathbf{q}}' \in \mathbb{R}^M, \forall \underline{\mathbf{F}}' \in \mathbb{R}^N : (\mathbb{L}\underline{\mathbf{q}}'|\underline{\mathbf{F}}') \leq b_h^2 |\underline{\mathbf{q}}'|^2 |\underline{\mathbf{F}}'|^2,$$

d'où :

$$|||\mathbb{L}||| \leq b_h. \quad (15.23)$$

Or, on a :

$$\begin{aligned} (\mathbb{C}\mathbb{A}^{-1}\mathbb{C}^T\underline{\mathbf{q}}|\underline{\mathbf{q}}) &= ((\mathbb{C}\mathbb{A}^{-1/2}\mathbb{A}^{-1/2}\mathbb{C}^T)\underline{\mathbf{q}}|\underline{\mathbf{q}}), \\ &= |\mathbb{A}^{-1/2}\mathbb{C}^T\underline{\mathbf{q}}|^2, \\ &= |\mathbb{L}\underline{\mathbf{q}}'|^2, \text{ par changement de variable.} \end{aligned}$$

On en déduit l'inégalité suivante, d'après (15.22) et (15.23) et comme $|\underline{\mathbf{q}}'|^2 = (\mathbb{M}\underline{\mathbf{q}}|\underline{\mathbf{q}})$:

$$\boxed{\kappa_h^{*2} (\mathbb{M}\underline{\mathbf{q}}|\underline{\mathbf{q}}) \leq (\mathbb{C}\mathbb{A}^{-1}\mathbb{C}^T\underline{\mathbf{q}}|\underline{\mathbf{q}}) \leq b_h^2 (\mathbb{M}\underline{\mathbf{q}}|\underline{\mathbf{q}}).} \quad (15.24)$$

□

Le corollaire suivant est immédiat.

Corollaire 15.7 *Si la condition inf-sup est uniforme et si la constante de continuité ne dépend pas du pas du maillage (κ_h^* et b_h indépendants de h), les matrices $\mathbb{M}$ et $\mathbb{C}\mathbb{A}^{-1}\mathbb{C}^T$ sont équivalentes.*

Si la condition inf-sup est uniforme, $\mathbb{M}^{-1}(\mathbb{C}\mathbb{A}^{-1}\mathbb{C}^T)$ est équivalente à une constante fois la matrice identité. Ainsi, le nombre de conditionnement de $\mathbb{M}^{-1}(\mathbb{C}\mathbb{A}^{-1}\mathbb{C}^T)$ est plus proche de un que celui de $\mathbb{C}\mathbb{A}^{-1}\mathbb{C}^T$, d'autant plus que κ_h^* et b_h sont proches de un. Afin de limiter le nombre d'itérations de l'algorithme d'Uzawa, il est donc judicieux de choisir $\mathbb{M}$ comme matrice de préconditionnement. Si le maillage le permet (voir la section 15.5), utiliser la matrice de masse réduite $\tilde{\mathbb{M}}$ permet de diminuer avantageusement le coût d'une itération.

Chapitre 16

Le champ magnétique

16.1 Le problème quasi-magnétostatique 2D

16.1.1 Introduction

Pour le problème bidimensionnel, les équations quasi-magnétostatique se résolvent de façon similaire aux équations quasi-électrostatique. La condition aux limites portant sur la composante normale de l'induction magnétique, le problème est tourné de 90° . Par rapport à l'espace $\mathbf{X}_{\mathbf{E}}^0$, les rôles de Φ_N et Φ_D sont inversés.

Le problème quasi-magnétostatique est le suivant :
Trouver $\mathbf{H} \in \mathbf{X}_{\mathbf{H}}$ tel que :

$$\operatorname{rot} \mathbf{H} = f_{\mathbf{H}} \text{ dans } \omega, f_{\mathbf{H}} \in L^2(\omega), \quad (16.1)$$

$$\operatorname{div} \mathbf{H} = g_{\mathbf{H}} \text{ dans } \omega, g_{\mathbf{H}} \in L^2(\omega), \quad (16.2)$$

$$H_\nu = h \text{ sur } \partial\omega, h \in L^2(\partial\omega), \quad (16.3)$$

Pour que les équations (16.1)-(16.3) soient bien posées, d'après l'intégration par parties (1.14), il faut que : $\int_{\omega} g_{\mathbf{H}} d\omega = \int_{\partial\omega} h d\sigma$. Supposons que h s'annule au voisinage des coins rentrants. Il existe alors un relèvement régulier $\mathbf{h} \in \mathbf{H}^1(\omega)$ de h ($h_\nu = h$). Posons $\mathbf{H} = \mathbf{H}^0 + \mathbf{h}$. $\mathbf{H}^0 \in \mathbf{X}_{\mathbf{H}}^0$ satisfait :

$$\operatorname{rot} \mathbf{H}^0 = f_{\mathbf{H}}^0 \in L^2(\omega), \quad (16.4)$$

$$\operatorname{div} \mathbf{H}^0 = g_{\mathbf{H}}^0 \in L^2(\omega), \quad (16.5)$$

où : $f_{\mathbf{H}}^0 = f_{\mathbf{H}} - \operatorname{rot} \mathbf{h}$ et $g_{\mathbf{H}}^0 = g_{\mathbf{H}} - \operatorname{div} \mathbf{h}$.

16.1.2 Le problème direct

On peut décomposer $\mathbf{H}$ et $\mathbf{H}^0$ comme somme de dérivées de potentiels :
 $\mathbf{H} = -\mathbf{grad} \psi_N + \mathbf{rot} \psi_D$ où ψ_N et ψ_D satisfont les problèmes de Neumann et de Dirichlet suivants :

$$\begin{cases} -\Delta \psi_N = g_{\mathbf{H}} \text{ dans } \omega. \\ \partial_\nu \psi_N = -h \text{ sur } \partial\omega, \end{cases} \quad \text{et} \quad \begin{cases} -\Delta \psi_D = f_{\mathbf{H}} \text{ dans } \omega. \\ \psi_D = 0 \text{ sur } \partial\omega. \end{cases} \quad (16.6)$$

$\mathbf{H}_0 = -\mathbf{grad} \psi_N^0 + \mathbf{rot} \psi_D^0$ où $\psi_N^0 \in \Phi_N$ et $\psi_D^0 \in \Phi_D$ satisfont les problèmes de Neumann et de Dirichlet homogènes suivants :

$$\begin{cases} -\Delta \psi_N^0 = g_{\mathbf{H}}^0 \text{ dans } \omega. \\ \partial_\nu \psi_N^0 = 0 \text{ sur } \partial\omega, \end{cases} \quad \text{et} \quad \begin{cases} -\Delta \psi_D^0 = f_{\mathbf{H}}^0 \text{ dans } \omega. \\ \psi_D^0 = 0 \text{ sur } \partial\omega. \end{cases} \quad (16.7)$$

On a donc la décomposition orthogonale suivante : $\mathbf{X}_H^0 = \mathbf{grad} \Phi_N \oplus^\perp \mathbf{rot} \Phi_D$.

On peut approcher le champ quasi-magnétostatique de plusieurs façons :

- Le champ quasi-magnétostatique peut donc être calculé par la méthode des potentiels décrite dans la section 2.2 de la même façon et avec la même précision que le champ quasi-électrostatique.

D'après la section 2.2, $\psi_N^0 = \tilde{\psi}_N + \sum_{i=1}^{N_{cr}} \mu_N^i \varphi_{N,i}^P$, et $\psi_D^0 = \tilde{\psi}_D + \sum_{i=1}^{N_{cr}} \mu_D^i \varphi_{D,i}^P$, où $\tilde{\psi}_{N,D} \in H^2(\omega)$, et : $\forall i \mu_N^i = (g_H^0, s_{N,i})_0 / \pi$, $\mu_D^i = (f_H^0, s_{D,i})_0 / \pi$.

- Comme $\mathbf{X}_H \cap \mathbf{H}^1(\omega)$ est dense dans $\mathbf{X}_H$, on peut approcher la solution de (16.1)-(16.3) par les éléments finis de Lagrange continus P_k .

Proposition 16.1 *Le problème (16.1)-(16.3) est équivalent au problème variationnel suivant : Trouver $\mathbf{H} \in \mathbf{X}_H$ tel que : $\forall \mathbf{F} \in \mathbf{X}_H$,*

$$(\mathbf{H}, \mathbf{F})_{\mathbf{X}_H} = (g_H, \operatorname{div} \mathbf{F})_0 + (f_H, \operatorname{rot} \mathbf{F})_0 + \int_{\partial\omega} h F_\nu d\sigma.$$

- On peut aussi appliquer la λ -approche au calcul du champ quasi-magnétostatique. On remarque en effet que : $\mathbf{grad} \varphi_{N,i}^P = \mathbf{rot} \varphi_{D,i}^P = \mathbf{y}_i^P$, avec $\mathbf{y}_i^P = \alpha_i r_i^{\alpha_i-1} \begin{pmatrix} \cos(\alpha_i \theta_i) \\ -\sin(\alpha_i \theta_i) \end{pmatrix}$, exprimé en coordonnées polaires par rapport au coin rentrant O_i .

On obtient alors la décomposition suivante : $\mathbf{H}^0 = \mathbf{H}^R + \sum_{i=1}^{N_{cr}} \mu^i \mathbf{y}_i^P$, $\mathbf{H}^R \in \mathbf{H}^1(\omega)$, et $\mu^i = \mu_D^i - \mu_N^i$.

Posons $\tilde{\mathbf{H}} = \mathbf{H}^R + \mathbf{h} \in \mathbf{H}^1(\omega)$. D'où : $\mathbf{H} = \tilde{\mathbf{H}} + \sum_{i=1}^{N_{cr}} \mu^i \mathbf{y}_i^P$. On remarque que : $\mathbf{y}_i^P \cdot \boldsymbol{\nu}|_{\partial\omega} \in L^2(\partial\omega)$ est nul sur les arêtes du coin rentrant O_i , et régulier ailleurs. Soit $\mathbf{h}_\mu$ un relèvement régulier de la condition aux limites normale : $\tilde{\mathbf{H}} \cdot \boldsymbol{\nu}|_{\partial\omega} = h - \sum_{i=1}^{N_{cr}} \mu^i \mathbf{y}_i^P \cdot \boldsymbol{\nu}|_{\partial\omega}$. Posons $\tilde{\mathbf{H}} = \tilde{\mathbf{H}}^0 + \mathbf{h}_\mu$. $\tilde{\mathbf{H}}^0$ satisfait :

$$\operatorname{rot} \tilde{\mathbf{H}}^0 = f_H - \operatorname{rot} \mathbf{h}_\mu \in L^2(\omega), \quad (16.8)$$

$$\operatorname{div} \tilde{\mathbf{H}}^0 = g_H - \operatorname{div} \mathbf{h}_\mu \in L^2(\omega). \quad (16.9)$$

Proposition 16.2 *Le problème (16.8)-(16.9) est équivalent à la formulation variationnelle suivante : Trouver $\tilde{\mathbf{H}}^0 \in \mathbf{X}_H^{0,R}$ tel que :*

$$\forall \mathbf{F} \in \mathbf{X}_H^{0,R}, \mathcal{A}_0(\tilde{\mathbf{H}}^0, \mathbf{F}) = \mathcal{M}_0(\mathbf{F}) - \mathcal{A}_0(\mathbf{h}_\mu, \mathbf{F}). \quad (16.10)$$

$\mathcal{A}_0$ restreinte à $\mathbf{X}_H^0 \times \mathbf{X}_H^0$ est une forme bilinéaire symétrique coercitive (c'est le produit scalaire dans $\mathbf{X}_H^0$). $\mathcal{M}_0$ est la forme linéaire continue suivante :

$$\begin{aligned} \mathcal{M}_0 : \mathbf{X}_H^0 &\rightarrow \mathbb{R} \\ \mathbf{F} &\mapsto (g_H, \operatorname{div} \mathbf{F})_0 + (f_H, \operatorname{rot} \mathbf{F})_0. \end{aligned}$$

Considérons le cas où il n'existe qu'un seul coin rentrant. Lorsque $h \in H^{1/2}(\partial\omega)$ s'annule au voisinage du coin rentrant, on peut montrer de la même façon que dans le paragraphe 2.4.3 que :

$$(\operatorname{div} \mathbf{h}, s_N)_0 - (\operatorname{rot} \mathbf{h}, s_D)_0 = \int_{\partial\omega} h s_N d\sigma. \quad (16.11)$$

On en déduit la formule généralisée et simplifiée suivante :

$$\mu = (f_{\mathbf{H}}, s_D)_0 - (g_{\mathbf{H}}, s_N)_0 + \int_{\partial\omega} h s_N d\sigma. \quad (16.12)$$

Étudions la décomposition orthogonale de $\mathbf{X}_{\mathbf{H}}^0$. Soit $\mathbf{X}_{\mathbf{H}}^{0,R} = \mathbf{X}_{\mathbf{H}}^0 \cap \mathbf{H}^1(\omega)$ l'espace régularisé de $\mathbf{X}_{\mathbf{H}}^0$. Lorsque ω est convexe, d'après [50], $\mathbf{X}_{\mathbf{H}}^{0,R} = \mathbf{X}_{\mathbf{H}}^0$. Au contraire, lorsqu'il existe un coin rentrant, $\mathbf{X}_{\mathbf{H}}^{0,R}$, est strictement inclus et est fermé dans $\mathbf{X}_{\mathbf{H}}^0$. Les éléments finis P_k de Lagrange ne permettent pas de capter les parties singulières (non $\mathbf{H}^1$). Soit $\mathbf{X}_{\mathbf{H}}^{0,S} = (\mathbf{X}_{\mathbf{H}}^{0,R})^{\perp_{\mathbf{X}_{\mathbf{H}}^0}}$. En procédant de la même façon que pour démontrer la théorème 2.43, on obtient une base de $\mathbf{X}_{\mathbf{H}}^{0,S}$:

Proposition 16.3 *On considère les vecteurs $\mathbf{y}_i^S = \mathbf{grad} \varphi_{N,i} + \mathbf{rot} \varphi_{D,i}$, $i \in \{1, \dots, N_{cr}\}$. Alors $\mathbf{X}_{\mathbf{H}}^{0,S}$ est généré par les $\mathbf{y}_i^S$: $\mathbf{X}_{\mathbf{H}}^{0,S} = \text{vect}(\mathbf{y}_1^S, \dots, \mathbf{y}_{N_{cr}}^S)$.*

D'après la λ -approche, les $\mathbf{y}_i^S$ s'écrivent ainsi : $\mathbf{y}_i^S = \tilde{\mathbf{y}}_i + \sum_{j=1}^{N_{cr}} \beta^{i,j} \mathbf{y}_j^P$, où $\tilde{\mathbf{y}}_i \in \mathbf{H}^1(\omega)$ satisfait :

$$\begin{aligned} \text{div } \tilde{\mathbf{y}}_i &= -s_{N,i} \text{ dans } \omega, \\ \text{rot } \tilde{\mathbf{y}}_i &= s_{D,i} \text{ dans } \omega, \\ \tilde{\mathbf{y}}_i \cdot \boldsymbol{\nu}|_{A_k} &= -\sum_{j=1}^{N_{cr}} \beta^{i,j} \mathbf{y}_j^P \cdot \boldsymbol{\nu}|_{A_k}, \forall k \in \{1, \dots, K\}. \end{aligned}$$

Pour tout i , on pose $\tilde{\mathbf{y}}_i = \tilde{\mathbf{y}}_i^0 + \mathbf{h}_i$, où $\tilde{\mathbf{y}}_i^0 \in \mathbf{X}_{\mathbf{H}}^{0,R}$ et où $\mathbf{h}_i$ est un relèvement régulier de $-\sum_{j=1}^{N_{cr}} \beta^{i,j} \mathbf{y}_j^P \cdot \boldsymbol{\nu}|_{\partial\omega}$. $\tilde{\mathbf{y}}_i^0$ satisfait alors :

$$\text{div } \tilde{\mathbf{y}}_i^0 = -s_{N,i} - \text{div } \mathbf{h}_i \text{ dans } \omega, \quad (16.13)$$

$$\text{rot } \tilde{\mathbf{y}}_i^0 = s_{D,i} - \text{rot } \mathbf{h}_i \text{ dans } \omega. \quad (16.14)$$

Proposition 16.4 *Fixons i . Le problème (16.13)-(16.14) est équivalent à la formulation variationnelle suivante : Trouver $\tilde{\mathbf{y}}_i^0 \in \mathbf{X}_{\mathbf{H}}^{0,R}$ tel que :*

$$\forall \mathbf{F} \in \mathbf{X}_{\mathbf{H}}^{0,R}, \mathcal{A}_0(\tilde{\mathbf{y}}_i^0, \mathbf{F}) = -\mathcal{A}_0(\mathbf{h}_i, \mathbf{F}). \quad (16.15)$$

Décomposons $\mathbf{H}^0$ de façon orthogonale : $\mathbf{H}^0 = \hat{\mathbf{H}}^0 + \sum_{i=1}^{N_{cr}} d_i \mathbf{y}_i^S$, où $\hat{\mathbf{H}}^0 \in \mathbf{X}_{\mathbf{H}}^{0,R}$ et $\forall i, d_i \in \mathbb{R}$.

Soient $\mathbf{d}$ et $\boldsymbol{\mu} \in \mathbb{R}^{N_{cr}}$ tels que : $\forall i \in \{1, \dots, N_{cr}\}$, $\mathbf{d}_i = d_i$ et $\boldsymbol{\mu}_i = \mu^i$.

Proposition 16.5 *Le problème (16.4)-(16.5) est équivalent à la formulation variationnelle suivante : Trouver $\hat{\mathbf{H}}^0 \in \mathbf{X}_{\mathbf{H}}^{0,R}$ et $\mathbf{d} \in \mathbb{R}$ tels que :*

$$\mathcal{A}_0(\hat{\mathbf{H}}^0, \mathbf{F}) = \mathcal{M}_0(\mathbf{F}) - \mathcal{A}_0(\mathbf{h}, \mathbf{F}), \forall \mathbf{F} \in \mathbf{X}_{\mathbf{H}}^{0,R}, \quad (16.16)$$

$$\mathbb{X} \mathbf{d} = \boldsymbol{\mu}. \quad (16.17)$$

16.2 Le problème quasi-magnétostatique 3D

16.2.1 Introduction

Le problème quasi-magnétostatique est le suivant :

Trouver $\mathcal{H} \in \mathcal{X}_{\mathcal{H}}$ tel que :

$$\mathbf{rot} \mathcal{H} = \mathbf{f}_{\mathcal{H}} \quad \text{dans } \Omega, \quad \mathbf{f}_{\mathcal{H}} \in \mathcal{H}(\operatorname{div}^0, \Omega), \quad (16.18)$$

$$\operatorname{div} \mathcal{H} = g_{\mathcal{H}} \quad \text{dans } \Omega, \quad g_{\mathcal{H}} \in L^2(\Omega) \quad (16.19)$$

$$H_{\nu} = h \quad \text{sur } \partial\Omega, \quad h \in L^2(\partial\Omega), \quad (16.20)$$

$\mathbf{f}_{\mathcal{H}}$ est à divergence nulle, et $h \in H^{1/2}(\partial\Omega)$ s'annule au voisinage des coins et des arêtes rentrants. Contrairement au cas bidimensionnel, il n'y a pas de relation de compatibilité entre $\mathbf{f}_{\mathcal{H}}$ et h . Notons que $\int_{\partial\Omega} h \, d\Sigma = \int_{\partial\Omega} \mathcal{H} \cdot \boldsymbol{\nu} \, d\Sigma = \int_{\Omega} g_{\mathcal{H}} \, d\Omega$, d'après la formule d'intégration par parties (7.13) avec $\mathbf{u} = \mathcal{H}$ et $v = 1$.

Définitions 16.6 On appelle la norme du graphe (ou norme naturelle) de $\mathcal{X}_{\mathcal{H}}$ la quantité suivante :

$$\|\mathbf{v}\|_{0, \mathbf{rot}, \operatorname{div}, L^2(\partial\Omega)} := \left(\|\mathbf{v}\|_0^2 + \|\mathbf{rot} \mathbf{v}\|_0^2 + \|\operatorname{div} \mathbf{v}\|_0^2 + \int_{\partial\Omega} |\mathbf{v} \cdot \boldsymbol{\nu}|^2 \, d\Sigma \right)^{1/2}.$$

On appelle la semi-norme de $\mathcal{X}_{\mathcal{H}}$ la quantité suivante :

$$|\mathbf{v}|_{\mathbf{rot}, \operatorname{div}, L^2(\partial\Omega)} := \left(\|\mathbf{rot} \mathbf{v}\|_0^2 + \|\operatorname{div} \mathbf{v}\|_0^2 + \int_{\partial\Omega} |\mathbf{v} \cdot \boldsymbol{\nu}|^2 \, d\Sigma \right)^{1/2}.$$

P. Fernandez et G. Gilardi ont montré dans [61] le théorème suivant :

Théorème 16.7 L'injection de $\mathcal{X}_{\mathcal{H}}$ dans $\mathcal{L}^2(\Omega)$ est compacte.

Ce théorème permet de montrer le lemme suivant, de la même façon que le lemme 8.4 :

Lemme 16.8 Il existe une constante $C > 0$ ne dépendant que de Ω telle que :

$$\forall \mathbf{v} \in \mathcal{X}_{\mathcal{H}}, \quad \|\mathbf{v}\|_0 \leq C |\mathbf{v}|_{\mathbf{rot}, \operatorname{div}, L^2(\partial\Omega)}. \quad (16.21)$$

On en déduit alors le théorème qui suit :

Théorème 16.9 Dans $\mathcal{X}_{\mathcal{H}}$, la semi-norme est équivalente à la norme du graphe : la semi-norme définit une norme sur $\mathcal{X}_{\mathcal{H}}$.

16.2.2 Le problème direct

En magnétostatique ($\operatorname{div} \mathcal{H} = 0$), si la densité de courant est nulle, on a $\mathbf{rot} \mathcal{H} = 0$, et on peut écrire le champ magnétique sous forme du gradient d'un *potentiel scalaire magnétique* : $\mathcal{H} = -\mathbf{grad} \phi_M$, où $\phi_M \in H^1(\Omega) \cap L_0^2(\Omega)$ satisfait le problème de Neuman suivant : $-\Delta \phi_M = 0$ dans Ω , et $\partial_{\nu} \phi_M = h$ sur $\partial\Omega$. On peut alors appliquer les méthodes et les résultats de l'électrostatique, en notant bien qu'il s'agit non plus de résoudre un problème de Dirichlet, mais un problème de Neumann.

Comme $\operatorname{div} \mathcal{H}$ est nul (absence de monopole magnétique), on peut écrire $\mathcal{H}$ sous la forme : $\mathbf{rot} \mathcal{A}$. $\mathcal{A}$ est appelé *potentiel vecteur*, et est défini à un gradient près car $\mathbf{rot} \mathbf{grad} = 0$. Si on choisit $\mathcal{A}$ tel que $\operatorname{div} \mathcal{A} = 0$ (c'est de choix de la *jauge de Coulomb*), $\mathcal{A} \in \mathcal{H}^1(\Omega)$ satisfait le problème de Neuman suivant : $-\Delta \mathcal{A} = \mathbf{f}_{\mathcal{H}}$ dans Ω et $\mathbf{rot} \mathcal{A} \cdot \mathbf{n} = h$ sur $\partial\Omega$.

Proposition 16.10 *Le problème (16.18)-(16.20) est équivalent au problème variationnel suivant : Trouver $\mathcal{H} \in \mathcal{X}_{\mathcal{H}}$ tel que :*

$$\forall \mathcal{F} \in \mathcal{X}_{\mathcal{H}}, (\mathcal{H}, \mathcal{F})_{\mathcal{X}_{\mathcal{H}}} = \mathcal{L}_{\mathcal{H}}(\mathcal{F}), \quad (16.22)$$

où $\mathcal{M}_{\mathcal{H}}$ est la forme linéaire suivante :

$$\begin{aligned} \mathcal{M}_{\mathcal{H}} : \mathcal{X}_{\mathcal{H}} &\rightarrow \mathbb{R} \\ \mathcal{F} &\mapsto (\mathbf{f}_{\mathcal{H}}, \mathbf{rot} \mathcal{F})_0 + (g_{\mathcal{H}}, \operatorname{div} \mathcal{F})_0 + \int_{\partial\Omega} h F_{\nu} d\Sigma. \end{aligned}$$

Comme $\mathcal{M}_{\mathcal{H}} \in \mathcal{X}'_{\mathcal{H}}$, d'après le théorème de Riesz-Fréchet, le problème (16.22) admet une solution unique $\mathcal{H} \in \mathcal{X}_{\mathcal{H}}$, qui dépend continûment des données $\mathbf{f}_{\mathcal{H}}$, $g_{\mathcal{H}}$ et h .

DÉMONSTRATION. Voir article [36]. □

Comme $\mathcal{X}_{\mathcal{H}} \cap \mathcal{H}^1(\Omega)$ est dense dans $\mathcal{X}_{\mathcal{H}}$ [40, 51], on peut approcher la solution de (8.12) par les éléments finis de Lagrange continus P_k .

Considérons $\mathcal{H}^r \in \mathcal{H}^1(\Omega)$ un relèvement régulier de h . On pose alors : $\mathcal{H} = \mathcal{H}^0 + \mathcal{H}^r$, où $\mathcal{H}^0 \in \mathcal{X}_{\mathcal{H}}^0$ satisfait le problème suivant :

$$\mathbf{rot} \mathcal{H} = \mathbf{f}_{\mathcal{H}} - \mathbf{rot} \mathcal{H}^r \quad \text{dans } \Omega, \quad , \quad (16.23)$$

$$\operatorname{div} \mathcal{H} = g_{\mathcal{H}} - \operatorname{div} \mathcal{H}^r \quad \text{dans } \Omega, \quad (16.24)$$

Proposition 16.11 *Le problème (16.23)-(16.24) est équivalent au problème variationnel suivant : Trouver $\mathcal{H}^0 \in \mathcal{X}_{\mathcal{H}}^0$ tel que :*

$$\forall \mathcal{F} \in \mathcal{X}_{\mathcal{H}}^0, (\mathcal{H}^0, \mathcal{F})_{\mathcal{X}_{\mathcal{H}}^0} = \mathcal{M}_0(\mathcal{F}) - (\mathcal{H}^r, \mathcal{F})_{\mathcal{X}_{\mathcal{H}}^0}, \quad (16.25)$$

Comme $\mathcal{M}_0 \in (\mathcal{X}_{\mathcal{H}}^0)'$, d'après le théorème 1.7, le problème (16.25) admet une solution unique $\mathcal{H}^0 \in \mathcal{X}_{\mathcal{H}}^0$, qui dépend continûment des données.

Lorsque Ω n'est pas convexe, $\mathcal{X}_{\mathcal{H}}^{0,R} := \mathcal{X}_{\mathcal{H}}^0 \cap \mathcal{H}^1(\Omega)$, l'espace régularisé de $\mathcal{X}_{\mathcal{H}}^0$ n'est pas dense dans $\mathcal{X}_{\mathcal{H}}^0$. $\mathcal{X}_{\mathcal{H}}^{0,R}$ est fermé et strictement inclus dans $\mathcal{X}_{\mathcal{H}}^0$.

Lorsque Ω est convexe, $\mathcal{X}_{\mathcal{H}}^{0,R} = \mathcal{X}_{\mathcal{H}}^0$. La solution de (8.14) approchée par les éléments finis de Lagrange continus P_k est fautive lorsque Ω présente des singularités géométriques.

La discrétisation des formulations variationnelles (16.22) et (16.25) est donnée dans le paragraphe 16.3.2 de la section suivante.

16.3 Champ quasi-magnétostatique : discrétisation 3D

Dans cette section, nous indiquons comment construire les matrices pour résoudre le problème quasi-magnétostatique 3D. Nous reprenons les notations du chapitre 10 de la partie III.

16.3.1 Méthode avec CL naturelles

Pour la méthode avec CL naturelles, l'espace de discrétisation de $\mathcal{X}_{\mathcal{H}}$ par les éléments finis P_k est $\mathcal{X}_k$. La matrice des produits scalaires $\mathcal{X}_{\mathcal{H}}$ des fonctions de $\mathcal{X}_k$ est représenté par $\mathbb{A}_{\mathcal{H}} = \text{Rot} + \text{Div} + \mathbb{B}_{\nu}$, où $\mathbb{B}_{\nu}$ est telle que : $\forall i$ ou $j \notin I_{\partial\Omega}$, $\mathbb{B}_{\nu}^{i,j} = 0$ et : $\forall i, j \in I_{\partial\Omega}$,

$$\mathbb{B}_{\nu}^{i,j} = \sum_{F_q \mid M_i, M_j \in F_q} \left(\int_{F_q} v_j v_i d\Sigma \boldsymbol{\nu}_q \cdot \boldsymbol{\nu}_q^T \right).$$

Posons $\mathbb{L}_h \in (\mathbb{R}^{3 \times 3})^{N \times N}$ la matrice définie par les $N \times N$ sous-blocs :

$$\begin{aligned} \mathbb{L}_h^{i,j} &= \sum_{F_q \mid M_i, M_j \in F_q} \int_{F_q} v_j v_i d\Sigma \boldsymbol{\nu}_q \cdot \boldsymbol{\nu}_q^T, \text{ si } i \text{ et } j \in I_{\partial\Omega}, \\ &= 0 \text{ sinon.} \end{aligned}$$

Soient $\mathbf{f}_{\mathcal{H}} \in \mathbb{R}^{3N}$ et $\mathbf{g}_{\mathcal{H}} \in \mathbb{R}^N$ les représentations de $\Pi_k \mathbf{f}_{\mathcal{H}}$ et $\Pi_k g_{\mathcal{H}}$ dans la base de $\mathcal{X}_k$, et $\mathbf{h}$ la représentation de $\mathcal{H}_h$, l'approximation de $\mathcal{H}$ dans $\mathcal{X}_k$. Soit $\mathbf{h} \in \mathbb{R}^N$, tel que $\mathbf{h}_i = 0$ si $M_i \in \Omega$, $\mathbf{h}(M_i) = h(M_i)$ sinon. La représentation discrétisée de la formulation variationnelle (16.22) se met sous la forme matricielle suivante :

$$\begin{aligned} \mathbb{A}_{\mathcal{H}} \mathbf{H} &= \mathbb{L}_{\mathbf{f}} \mathbf{f}_{\mathcal{H}} + \mathbb{L}_g \mathbf{g}_{\mathcal{H}} + \mathbb{L}_h \mathbf{h}, \\ \text{avec : } \mathbb{A}_{\mathcal{H}} &= \text{Div} + \text{Rot} + \mathbb{B}_{\nu}. \end{aligned} \tag{16.26}$$

Par la suite, on pose : $\mathbf{M} = \mathbb{L}_{\mathbf{f}} \mathbf{f}_{\mathcal{H}} + \mathbb{L}_g \mathbf{g}_{\mathcal{H}} + \mathbb{L}_h \mathbf{h}$.

16.3.2 Méthode avec CL essentielles

$\mathcal{X}_{\mathcal{H}}^0 \cap \mathcal{H}^1(\Omega)$ est dense dans $\mathcal{X}_{\mathcal{H}}^0$, seulement si Ω est convexe. Ainsi, la méthode décrite ici converge dans la cas où Ω présente des singularités géométriques seulement si les données sont régulières.

On considère comme espace de discrétisation de $\mathcal{X}_{\mathcal{H}}^{0,R}$ le sous-espace de $\mathcal{X}_k$, conforme dans $\mathcal{X}_{\mathcal{H}}^0$:

$$\mathcal{X}_{\mathcal{H},k}^{0,R} = \{ \mathbf{v} \in \mathcal{X}_k \mid \mathbf{v} \cdot \boldsymbol{\nu} \big|_{\partial\Omega} = 0 \text{ sur } \partial\Omega \}. \tag{16.27}$$

Par construction, $\mathcal{X}_{\mathcal{H},k}^{0,R} \subset \mathcal{H}^1(\Omega)$. La discrétisation de (16.25) se fait de la même façon que celle de (16.22), mais il faut prendre en compte la condition aux limites normales dans les matrices Div et Rot , et dans les termes du second membre.

Propriété 16.12 *L'espace $\mathcal{X}_{\mathcal{H},k}^{0,R}$ est de dimension finie $3N - N_f - 2N_a - 3N_c$.*

DÉMONSTRATION. Soit $\mathbf{u} \in \mathcal{X}_{\mathcal{H},k}^{0,R}$. Comme $\mathcal{X}_{\mathcal{H},k}^{0,R}$ est un sous-espace de $\mathcal{X}_k$, il est de dimension inférieure ou égale à $3N$ et on peut écrire $\mathbf{u} = \sum_{i=1}^N \sum_{\alpha=1}^3 u_{\alpha}(M_i) \mathbf{v}_{i,\alpha}$.

• Considérons un point M_i , $i \in I_a$, se trouvant sur l'arête $A_{k,l}$ (figure 7.1, p. 153). On a :

$$\begin{cases} \mathbf{u}(M_i) \cdot \boldsymbol{\nu}_l &= 0, \\ \mathbf{u}(M_i) \cdot \boldsymbol{\nu}_k &= 0. \end{cases}$$

Comme les vecteurs $\boldsymbol{\nu}_k$ et $\boldsymbol{\nu}_l$ ne sont pas colinéaires, $\mathbf{u}(M_i)$ est nécessairement colinéaire à $\boldsymbol{\tau}_{k,l}$.

Ceci étant valable pour tous les points d'arêtes, on peut écrire : $\mathbf{u} = \sum_{i \in I \setminus I_a} \sum_{\alpha=1}^3 u_{\alpha}(M_i) \mathbf{v}_{i,\alpha} +$

$\sum_{i \in I_a} u_{\tau}(M_i) v_i \boldsymbol{\tau}_i$, avec $\boldsymbol{\tau}_i = \boldsymbol{\tau}_{k,l}$, défini de façon unique; et $u_{\tau}(M_i) = \mathbf{u}(M_i) \cdot \boldsymbol{\tau}_i$. Pour chaque

vecteur de $\mathcal{X}_{\mathcal{H},k}^{0,R}$, nous éliminons alors deux degrés de liberté liés à chaque point d'arête. D'où : $\dim(\mathcal{X}_{\mathcal{H},k}^{0,R}) \leq 3N - 2N_a$.

• Considérons un coin M_i , $i \in I_c$. M_i est l'intersection d'au moins trois faces, dont les vecteurs normaux ne sont pas colinéaires ni tous les trois dans le même plan, d'où $\mathbf{u}(M_i) = 0$. Ceci étant

valable pour tous les coins, on peut écrire : $\mathbf{u} = \sum_{i=1}^{N-N_a} \sum_{\alpha=1}^3 u_{\alpha}(M_i) \mathbf{v}_{i,\alpha} + \sum_{i \in I_a} u_{\tau}(M_i) \boldsymbol{\tau}_i$. Pour

chaque vecteur de $\mathcal{X}_{\mathcal{H},k}^{0,R}$, nous éliminons alors les trois degrés de liberté liés à chaque coin. Il y en a N_c , d'où : $\dim(\mathcal{X}_{\mathcal{H},k}^{0,R}) \leq 3N - 2N_a - 3N_c$.

• Considérons un point M_i , $i \in I_f$, se trouvant sur la face F_k de $\partial\Omega$. Comme $\mathbf{u}(M_i) \cdot \boldsymbol{\nu}|_{F_k} = 0$, on peut écrire $\mathbf{u}(M_i)|_{F_k} = u_{\tau^1}(M_i) v_i \boldsymbol{\tau}_{F_k}^1 + u_{\tau^2}(M_i) v_i \boldsymbol{\tau}_{F_k}^2$.

D'où : $\mathbf{u} = \sum_{i \in I_{\Omega}} \sum_{\alpha=1}^3 u_{\alpha}(M_i) \mathbf{v}_{i,\alpha} + \sum_{i \in I_f} \sum_{\alpha=1}^2 u_{\tau^{\alpha}}(M_i) v_i \boldsymbol{\tau}_i^{\alpha} + \sum_{i \in I_a} u_{\tau}(M_i) v_i \boldsymbol{\tau}_i$. Ainsi, pour chaque

vecteur de $\mathcal{X}_{\mathcal{H},k}^{0,R}$, nous éliminons un degré de liberté par point du bord qui n'est pas ni sur une arête ni sur un coin. D'où : $\dim(\mathcal{X}_{\mathcal{H},k}^{0,R}) \leq 3N - N_f - 2N_a - 3N_c$. Pour conclure, la famille

$B'_0 := (\mathbf{v}_{i,\alpha})_{i \in I_{\Omega}, \alpha \in \{1,2,3\}} \cup (v_i \boldsymbol{\tau}_i^{\alpha})_{i \in I_f, \alpha \in \{1,2\}} \cup (v_i \boldsymbol{\tau}_i)_{i \in I_a}$ est contenue dans $\mathcal{X}_{\mathcal{H},k}^{0,R}$, d'où : $\dim(\mathcal{X}_{\mathcal{H},k}^{0,R}) \geq 3N - 2N_f - 3N_a - 3N_c$.

□

Nous avons donc montré la propriété 16.12, et au passage que B'_0 engendre $\mathcal{X}_{\mathcal{H},k}^{0,R}$. On utilisera donc deux types de base locale de $\mathbb{R}^3$ selon l'emplacement de M_i :

- $\forall i \in I_{\Omega}$, $M_i \in \Omega$, on travaille dans la base canonique $(\mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3)$,
- $\forall i \in I_f$, $M_i \in \partial\Omega(\mathcal{A} \cup \mathcal{C})$, on travaille dans la base locale $(\boldsymbol{\tau}_i^1, \boldsymbol{\tau}_i^2, \boldsymbol{\nu}_i)$,
- $\forall i \in I_a$, $M_i \in \mathcal{A}$, si M_i est sur une arête de F_k et de F_l , on travaille dans la base locale $(\boldsymbol{\tau}_i, \boldsymbol{\nu}_i^1, \boldsymbol{\nu}_i^2)$, où par exemple si M_i est sur l'arête $A_{k,l}$ et que F_k et F_l sont orthogonales : $\boldsymbol{\nu}_i^1 = \boldsymbol{\nu}_k$ et $\boldsymbol{\nu}_i^2 = \boldsymbol{\nu}_l$ (ou le contraire, de sorte que $(\boldsymbol{\tau}_i, \boldsymbol{\nu}_i^1, \boldsymbol{\nu}_i^2)$ forme une base orthonormée directe).
- $\forall i \in I_c$, $M_i \in \mathcal{C}$, afin de lever l'ambiguïté sur le vecteur normal, on travaille dans la base canonique $(\mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3)$.

On appelle B' la base de $\mathcal{X}_k$ ainsi formée :

$$B' = (\mathbf{v}_i, \alpha)_{i \in I_\Omega \cup I_c, \alpha \in \{1,2,3\}} \cup (v_i \boldsymbol{\tau}_i^\alpha)_{i \in I_f, \alpha \in \{1,2\}} \cup (v_i \boldsymbol{\nu}_i)_{i \in I_f} \cup (v_i \boldsymbol{\nu}_i^\alpha)_{i \in I_a, \alpha \in \{1,2\}} \cup (v_i \boldsymbol{\tau}_i)_{i \in I_a}. \quad (16.28)$$

Les fonctions de base de B (10.12) et de B' sont les mêmes pour les points M_i tels que $i \in I \setminus I_a$.

Soit $\mathbb{A}' \in (\mathbb{R}^{3 \times 3})^{N \times N}$ la décomposition de la matrice $\text{Div} + \text{Rot}$ dans B' , cette fois-ci, qu'on écrit de la façon suivante :

$$\mathbb{A}' = \begin{pmatrix} \mathbb{A}_{\Omega, \Omega} & \mathbb{A}_{\Omega, f} & \mathbb{A}_{\Omega, a} & \mathbb{A}_{\Omega, c} \\ \mathbb{A}_{f, \Omega} & \mathbb{A}_{f, f} & \mathbb{A}_{f, a} & \mathbb{A}_{f, c} \\ \mathbb{A}_{a, \Omega} & \mathbb{A}_{a, f} & \mathbb{A}_{a, a} & \mathbb{A}_{a, c} \\ \mathbb{A}_{c, \Omega} & \mathbb{A}_{c, f} & \mathbb{A}_{c, a} & \mathbb{A}_{c, c} \end{pmatrix}.$$

Nous ne décrivons ici que les sous-blocs qui n'ont pas été décrits au paragraphe 10.4.1 :

- $\mathbb{A}_{c, c} = \mathbb{A}_{ac, ac}^{i \in I_c, j \in I_c} \in (\mathbb{R}^{3 \times 3})^{N_c \times N_c}$;
- $\mathbb{A}_{\Omega, c} = \mathbb{A}_{\Omega, ac}^{i \in I_\Omega, j \in I_c} \in (\mathbb{R}^{3 \times 3})^{N_\Omega \times N_c}$;
- $\mathbb{A}_{c, \Omega} = \mathbb{A}_{ac, \Omega}^{i \in I_c, j \in I_\Omega} \in (\mathbb{R}^{3 \times 3})^{N_c \times N_\Omega}$.
- $\mathbb{A}_{f, c} = \mathbb{A}_{f, ac}^{i \in I_c, j \in I_f} \in (\mathbb{R}^{3 \times 3})^{N_c \times N_f}$.
- $\mathbb{A}_{c, f} = \mathbb{A}_{fac, f}^{i \in I_f, j \in I_c} \in (\mathbb{R}^{3 \times 3})^{N_f \times N_c}$.

La sous-matrice $\mathbb{A}_{a, a} \in (\mathbb{R}^{3 \times 3})^{N_a \times N_a}$ est composée des $N_a \times N_a$ sous-blocs $\mathbb{A}_{a, a}^{i, j} \in \mathbb{R}^{3 \times 3}$, tels que : $\forall i \in I_a, \forall j \in I_a$:

$$\mathbb{A}_{a, a}^{i, j} = \begin{pmatrix} (v_j \boldsymbol{\tau}_j, v_i \boldsymbol{\tau}_i)_{\mathcal{H}}^0 & (v_j \boldsymbol{\nu}_j^1, v_i \boldsymbol{\tau}_i)_{\mathcal{H}}^0 & (v_j \boldsymbol{\nu}_j^2, v_i \boldsymbol{\tau}_i)_{\mathcal{H}}^0 \\ (v_j \boldsymbol{\tau}_j, v_i \boldsymbol{\nu}_i^1)_{\mathcal{H}}^0 & (v_j \boldsymbol{\nu}_j^1, v_i \boldsymbol{\nu}_i^1)_{\mathcal{H}}^0 & (v_j \boldsymbol{\nu}_j^2, v_i \boldsymbol{\nu}_i^1)_{\mathcal{H}}^0 \\ (v_j \boldsymbol{\tau}_j, v_i \boldsymbol{\nu}_i^2)_{\mathcal{H}}^0 & (v_j \boldsymbol{\nu}_j^1, v_i \boldsymbol{\nu}_i^2)_{\mathcal{H}}^0 & (v_j \boldsymbol{\nu}_j^2, v_i \boldsymbol{\nu}_i^2)_{\mathcal{H}}^0 \end{pmatrix}.$$

Ces sous-blocs agissent sur les vecteurs de $\mathbb{R}^3$ décomposés dans la base locale $(\boldsymbol{\nu}_j^1, \boldsymbol{\nu}_j^2, \boldsymbol{\tau}_j)^T$:

$$\mathcal{H}_h(M_j) = \begin{pmatrix} H_\tau(M_j) \\ H_{\nu^1}(M_j) \\ H_{\nu^2}(M_j) \end{pmatrix}, \text{ où } H_{\nu^\alpha}(M_j) = \mathbf{H}_h(M_j) \cdot \boldsymbol{\nu}_j^\alpha, \alpha = 1, 2, \text{ et } H_\tau(M_j) = \mathcal{H}_h(M_j) \cdot \boldsymbol{\tau}_j.$$

La sous-matrice $\mathbb{A}_{\Omega, a} \in (\mathbb{R}^{3 \times 3})^{N_\Omega \times N_a}$ est composée des $N_\Omega \times N_a$ sous-blocs $\mathbb{A}_{\Omega, a}^{i, j} \in \mathbb{R}^{3 \times 3}$, tels

que : $\forall i \in I_\Omega, \forall j \in I_a$:

$$\mathbb{A}_{\Omega,a}^{i,j} = \begin{pmatrix} (v_j \boldsymbol{\tau}_j, v_i \mathbf{e}_1)_{\mathcal{H}}^0 & (v_j \boldsymbol{\nu}_j^1, v_i \mathbf{e}_1)_{\mathcal{H}}^0 & (v_j \boldsymbol{\nu}_j^2, v_i \mathbf{e}_1)_{\mathcal{H}}^0 \\ (v_j \boldsymbol{\tau}_j, v_i \mathbf{e}_2)_{\mathcal{H}}^0 & (v_j \boldsymbol{\nu}_j^1, v_i \mathbf{e}_2)_{\mathcal{H}}^0 & (v_j \boldsymbol{\nu}_j^2, v_i \mathbf{e}_2)_{\mathcal{H}}^0 \\ (v_j \boldsymbol{\tau}_j, v_i \mathbf{e}_3)_{\mathcal{H}}^0 & (v_j \boldsymbol{\nu}_j^1, v_i \mathbf{e}_3)_{\mathcal{H}}^0 & (v_j \boldsymbol{\nu}_j^2, v_i \mathbf{e}_3)_{\mathcal{H}}^0 \end{pmatrix}.$$

Symétriquement, on a : $\mathbb{A}_{a,\Omega} \in (\mathbb{R}^{3 \times 3})^{N_a \times N_\Omega}$ telle que : $\forall i \in I_a, \forall j \in I_\Omega, \mathbb{A}_{a,\Omega}^{i,j} = \mathbb{A}_{\Omega,a}^{j,i}$. Les sous-blocs $\mathbb{A}_{a,c} \in (\mathbb{R}^{3 \times 3})^{N_a \times N_c}$ et $\mathbb{A}_{c,a} \in (\mathbb{R}^{3 \times 3})^{N_c \times N_a}$ se forment de la même façon.

La sous-matrice $\mathbb{A}_{f,a} \in (\mathbb{R}^{3 \times 3})^{N_f \times N_a}$ est composée des $N_f \times N_a$ sous-blocs $\mathbb{A}_{f,a}^{i,j} \in \mathbb{R}^{3 \times 3}$, tels que : $\forall i \in I_f, \forall j \in I_a$:

$$\mathbb{A}_{f,a}^{i,j} = \begin{pmatrix} (v_j \boldsymbol{\tau}_j, v_i \boldsymbol{\tau}_i^1)_{\mathcal{H}}^0 & (v_j \boldsymbol{\nu}_j^1, v_i \boldsymbol{\tau}_i^1)_{\mathcal{H}}^0 & (v_j \boldsymbol{\nu}_j^2, v_i \boldsymbol{\tau}_i^1)_{\mathcal{H}}^0 \\ (v_j \boldsymbol{\tau}_j, v_i \boldsymbol{\tau}_i^2)_{\mathcal{H}}^0 & (v_j \boldsymbol{\nu}_j^1, v_i \boldsymbol{\tau}_i^2)_{\mathcal{H}}^0 & (v_j \boldsymbol{\nu}_j^2, v_i \boldsymbol{\tau}_i^2)_{\mathcal{H}}^0 \\ (v_j \boldsymbol{\tau}_j, v_i \boldsymbol{\nu}_i)_{\mathcal{H}}^0 & (v_j \boldsymbol{\nu}_j^1, v_i \boldsymbol{\nu}_i)_{\mathcal{H}}^0 & (v_j \boldsymbol{\nu}_j^2, v_i \boldsymbol{\nu}_i)_{\mathcal{H}}^0 \end{pmatrix}.$$

Symétriquement, on a : $\mathbb{A}_{a,f} \in (\mathbb{R}^{3 \times 3})^{N_a \times N_f}$ telle que : $\forall i \in I_a, \forall j \in I_f, \mathbb{A}_{a,f}^{i,j} = \mathbb{A}_{f,a}^{j,i}$. Nous ne détaillerons pas les calculs des sous-matrices de $\mathbb{A}'$, car ils sont similaires aux calculs des éléments de $\mathbb{R}ot$ et $\mathbb{D}iv$, expliqués dans le paragraphe 10.3.1.

Soit $\mathcal{F} \in \mathcal{X}_k$. Afin de définir le produit matrice-vecteur $\mathbb{A}' \underline{E}$, où $\underline{E}$ est le vecteur de $(\mathbb{R}^3)^N$ associé à $\mathcal{F}$, on décompose $\mathcal{F}$ dans B' :

$$\begin{aligned} \mathcal{F} = & \sum_{j \in I_\Omega \cup I_c} \sum_{\beta=1}^3 F_\beta(M_j) \mathbf{v}_{j,\beta} + \sum_{j \in I_f} \left(\sum_{\beta=1}^2 F_{\tau\beta}(M_j) v_j \boldsymbol{\tau}_j^\beta + F_\nu(M_j) v_j \boldsymbol{\nu}_j \right) \\ & + \sum_{j \in I_a} \left(\sum_{\beta=1}^2 F_{\nu\beta}(M_j) v_j \boldsymbol{\nu}_j^\beta + F_\tau(M_j) v_j \boldsymbol{\tau}_j \right), \end{aligned}$$

où $\forall i \in I_a, \forall \alpha \in \{1, 2\}, F_{\nu\alpha}(M_j) = \mathcal{F}(M_i) \cdot \boldsymbol{\nu}_i^\alpha$ et $F_\tau(M_j) = \mathcal{F}(M_i) \cdot \boldsymbol{\tau}_i$.

On considère alors $\underline{E} = (\underline{E}_\Omega, \underline{E}_f, \underline{E}_a, \underline{E}_c)^T$ le vecteur de $\mathbb{R}^{3N}$ associé, dont les sous-composantes $\underline{E}_a$ et $\underline{E}_c$ sont :

$$- \underline{E}_a = (F_{\nu^1}(M_{N-N_{ac}+1}), F_{\nu^2}(M_{N-N_{ac}+1}), F_\tau(M_{N-N_{ac}+1}), \dots, F_{\nu^1}(M_{N-N_c}), F_{\nu^2}(M_{N-N_c}), F_\tau(M_{N-N_c}))^T \in (\mathbb{R}^3)^{N_a},$$

$$- \underline{E}_c = (F_1(M_{N-N_c+1}), F_2(M_{N-N_c+1}), F_3(M_{N-N_c+1}), \dots, F_1(M_N), F_2(M_N), F_3(M_N))^T \in (\mathbb{R}^3)^{N_c}.$$

Le produit matrice-vecteur $\mathbb{A}' \underline{E}$ est ainsi bien défini. Pour construire la matrice de la formulation variationnelle (16.25) discrétisée dans $\mathcal{X}_{\mathcal{H},k}^{0,R}$, on peut alors procéder à l'élimination des conditions aux limites essentielles dans $\mathbb{A}'$. Cela correspond à éliminer d'une part les lignes et les colonnes de $\mathbb{A}'$ dont les éléments dépendent de $\boldsymbol{\nu}$ pour les sommets situés à l'intérieur des faces du bord ($i \in I_f$) ; d'autre part les lignes et les colonnes dont les éléments dépendent de $\boldsymbol{\nu}^{1,2}$ pour les sommets situés

sur les arêtes de $\partial\Omega$ ($i \in I_a$); et enfin toutes les lignes et les colonnes pour les sommets situés sur les coins de $\partial\Omega$ ($i \in I_c$). On appelle $\mathbb{A}'_0$ la matrice ainsi obtenue.

Soit $\mathbb{A}_{\Omega,\tau} \in (\mathbb{R}^{3 \times 3})^{N_\Omega \times N_f}$ la matrice $\mathbb{A}_{\Omega,f}$ dont on a éliminé la colonne dépendant de $\boldsymbol{\nu}$: $\forall i \in I_\Omega, \forall j \in I_f$,

$$\mathbb{A}_{\Omega,\tau}^{i,j} = \begin{pmatrix} (v_j \boldsymbol{\tau}_j^1, v_i \mathbf{e}_1)_{\mathcal{H}} & (v_j \boldsymbol{\tau}_j^2, v_i \mathbf{e}_1)_{\mathcal{H}} & 0 \\ (v_j \boldsymbol{\tau}_j^1, v_i \mathbf{e}_2)_{\mathcal{H}} & (v_j \boldsymbol{\tau}_j^2, v_i \mathbf{e}_2)_{\mathcal{H}} & 0 \\ (v_j \boldsymbol{\tau}_j^1, v_i \mathbf{e}_3)_{\mathcal{H}} & (v_j \boldsymbol{\tau}_j^2, v_i \mathbf{e}_3)_{\mathcal{H}} & 0 \end{pmatrix}.$$

Symétriquement, $\mathbb{A}_{\tau,\Omega} \in (\mathbb{R}^{3 \times 3})^{N_f \times N_\Omega}$ est la matrice $\mathbb{A}_{f,\Omega}$ dont on a éliminé la ligne qui dépend de $\boldsymbol{\nu}$: $\mathbb{A}_{\tau,\Omega}^{i,j} = \mathbb{A}_{\Omega,\tau}^{j,i}$.

On considère $\mathbb{A}_{\tau,\tau}$ la matrice $\mathbb{A}_{f,f}$ dont on a éliminé les termes dépendant de $\boldsymbol{\nu}$, sauf le terme diagonal, pour lequel on a imposé la valeur 1 :

$$\mathbb{A}_{\tau,\tau}^{i,j} = \begin{pmatrix} (v_j \boldsymbol{\tau}_j^1, v_i \boldsymbol{\tau}_i^1)_{\mathcal{H}} & (v_j \boldsymbol{\tau}_j^2, v_i \boldsymbol{\tau}_i^1)_{\mathcal{H}} & 0 \\ (v_j \boldsymbol{\tau}_j^1, v_i \boldsymbol{\tau}_i^2)_{\mathcal{H}} & (v_j \boldsymbol{\tau}_j^2, v_i \boldsymbol{\tau}_i^2)_{\mathcal{H}} & 0 \\ 0 & 0 & \delta_{ij} \end{pmatrix}.$$

En effet, $\mathbb{A}_{f,f}$ correspond à un bloc diagonal de $\mathbb{A}'$. En imposant la valeur 1 sur les termes diagonaux éliminés, on s'assure que la matrice ainsi modifiée soit inversible.

Soit $\mathbb{A}_{\Omega,\tau_a} \in (\mathbb{R}^{3 \times 3})^{N_\Omega \times N_a}$ la matrice $\mathbb{A}_{\Omega,a}$ dont on a éliminé les colonnes dépendant de $\boldsymbol{\nu}^{1,2}$: $\forall i \in I_\Omega, \forall j \in I_a$,

$$\mathbb{A}_{\Omega,\tau}^{i,j} = \begin{pmatrix} (v_j \boldsymbol{\tau}_j, v_i \mathbf{e}_1)_{\mathcal{H}} & 0 & 0 \\ (v_j \boldsymbol{\tau}_j, v_i \mathbf{e}_2)_{\mathcal{H}} & 0 & 0 \\ (v_j \boldsymbol{\tau}_j, v_i \mathbf{e}_3)_{\mathcal{H}} & 0 & 0 \end{pmatrix}.$$

Symétriquement, $\mathbb{A}_{\tau_a,\Omega} \in (\mathbb{R}^{3 \times 3})^{N_a \times N_\Omega}$ est la matrice $\mathbb{A}_{a,\Omega}$ dont on a éliminé les lignes qui dépendent de $\boldsymbol{\nu}^{1,2}$: $\mathbb{A}_{\tau_a,\Omega}^{i,j} = \mathbb{A}_{\Omega,\tau_a}^{j,i}$.

Soit $\mathbb{A}_{\tau,\tau_a} \in (\mathbb{R}^{3 \times 3})^{N_f \times N_a}$ la matrice $\mathbb{A}_{\Omega,a}$ dont on a éliminé la ligne dépendant de $\boldsymbol{\nu}$ et les colonnes dépendant de $\boldsymbol{\nu}^{1,2}$:

$\forall i \in I_\Omega, \forall j \in I_a$,

$$\mathbb{A}_{\Omega,\tau}^{i,j} = \begin{pmatrix} (v_j \boldsymbol{\tau}_j, v_i \boldsymbol{\tau}_i^1)_{\mathcal{H}} & 0 & 0 \\ (v_j \boldsymbol{\tau}_j, v_i \boldsymbol{\tau}_i^2)_{\mathcal{H}} & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}.$$

Symétriquement, $\mathbb{A}_{\tau_a,\tau} \in (\mathbb{R}^{3 \times 3})^{N_a \times N_f}$ est la matrice $\mathbb{A}_{a,f}$ dont on a éliminé les lignes qui dépendent de $\boldsymbol{\nu}^{1,2}$ et la colonne qui dépend de $\boldsymbol{\nu}$: $\mathbb{A}_{\tau_a,\tau}^{i,j} = \mathbb{A}_{\tau,\tau_a}^{j,i}$.

On considère $\mathbb{A}_{\tau_a,\tau_a}$ la matrice $\mathbb{A}_{a,a}$ dont on a éliminé les termes dépendant de $\boldsymbol{\nu}^{1,2}$, sauf les termes diagonaux, pour lesquels on a imposé la valeur 1 :

$$\mathbb{A}_{\tau_a,\tau_a}^{i,j} = \begin{pmatrix} (v_j \boldsymbol{\tau}_j, v_i \boldsymbol{\tau}_i)_{\mathcal{H}} & 0 & 0 \\ 0 & \delta_{ij} & 0 \\ 0 & 0 & \delta_{ij} \end{pmatrix}.$$

En effet, $\mathbb{A}_{a,a}$ correspond à un bloc diagonal de $\mathbb{A}'$. En imposant la valeur 1 sur les termes diagonaux éliminés, on s'assure que la matrice ainsi modifiée soit inversible.

On en déduit que la structure par blocs de la matrice $\mathbb{A}'_0 \in (\mathbb{R}^{3 \times 3})^{N \times N}$ est la suivante :

$$\mathbb{A}'_0 = \begin{pmatrix} \mathbb{A}_{\Omega, \Omega} & \mathbb{A}_{\Omega, \tau} & \mathbb{A}_{\Omega, \tau_a} & 0 \\ \mathbb{A}_{\tau, \Omega} & \mathbb{A}_{\tau, \tau} & \mathbb{A}_{\tau, \tau_a} & 0 \\ \mathbb{A}_{\tau_a, \Omega} & \mathbb{A}_{\tau_a, \tau} & \mathbb{A}_{\tau_a, \tau_a} & 0 \\ 0 & 0 & 0 & \mathbb{I}_{c,c} \end{pmatrix},$$

où $\mathbb{I}_{c,c}$ est la matrice identité de $(\mathbb{R}^{3 \times 3})^{N_c \times N_c}$.

Proposition 16.13 *La matrice $\mathbb{A}'_0$ est inversible.*

La preuve est similaire à celle de la proposition 10.4.

En procédant comme dans le paragraphe 10.4.2, on montre que le problème variationnel (16.25) discrétisé dans $\mathcal{X}_{\mathcal{H},k}^{0,R}$ s'écrit :

Trouver $\mathcal{H}_h \in \mathcal{X}_k$ tel que :

$$\begin{aligned} \forall i \in I_\Omega, \forall \alpha \in \{1, 2, 3\}, \quad & (\mathcal{H}_h, \mathbf{v}_{i,\alpha})_{\mathcal{X}_{\mathcal{H}}^0} = \mathcal{M}_0(\mathbf{v}_{i,\alpha}), \\ \forall i \in I_f, \forall \alpha \in \{1, 2\} \quad & (\mathcal{H}_h, v_i \boldsymbol{\tau}_i^\alpha)_{\mathcal{X}_{\mathcal{H}}^0} = \mathcal{M}_0(v_i \boldsymbol{\tau}_i^\alpha), \\ \forall i \in I_f, \quad & \mathcal{H}_h(M_i) \cdot \boldsymbol{\nu}_i = h(M_i), \\ \forall i \in I_a, \forall \alpha \in \{1, 2\}, \quad & \mathcal{H}_h(M_i) \cdot \boldsymbol{\nu}_i^\alpha = \mathcal{H}^r(M_i) \cdot \boldsymbol{\nu}_i^\alpha, \\ \forall i \in I_c, \quad & \mathcal{H}_h(M_i) = \mathcal{H}^r(M_i). \end{aligned} \quad (16.29)$$

Pour les deux dernières inégalités, on peut procéder comme dans le paragraphe 10.4.3.

Soit $\underline{\mathbf{H}} = (\underline{\mathbf{H}}_\Omega, \underline{\mathbf{H}}_f, \underline{\mathbf{H}}_a, \underline{\mathbf{H}}_c)^T$ le vecteur de $\mathbb{R}^{3N}$ associé dont les composantes sont égales à celles de $\mathcal{H}_h$ dans B' . Considérons le relèvement discret suivant : $\underline{\mathbf{H}}^r = (\underline{\mathbf{Z}}_\Omega, \underline{\mathbf{H}}_\nu, \underline{\mathbf{H}}_{\nu_a}, \underline{\mathbf{H}}_c)^T$, où :

- $\underline{\mathbf{Z}}_\Omega$ est le vecteur nul de $(\mathbb{R}^3)^{N_\Omega}$;

- $\underline{\mathbf{H}}_\nu^r = (0, 0, h(M_{N_\Omega+1}), \dots, 0, 0, h(M_{N_\Omega+N_f}))^T \in (\mathbb{R}^3)^{N_f}$, est le vecteur contenant les valeurs les valeurs de h sur les points des faces de $\partial\Omega$;

- $\underline{\mathbf{H}}_\nu^r = (0, \mathcal{H}^r(M_{N-N_{ac}+1}) \cdot \boldsymbol{\nu}_i^1, \mathcal{H}^r(M_{N-N_{ac}+1}) \cdot \boldsymbol{\nu}_i^2, \dots, 0, \mathcal{H}^r(M_{N-N_c}) \cdot \boldsymbol{\nu}_i^1, \mathcal{H}^r(M_{N-N_c}) \cdot \boldsymbol{\nu}_i^2)^T \in (\mathbb{R}^3)^{N_a}$, est le vecteur contenant les valeurs des composantes de $\mathcal{H}^r$ sur les arêtes du maillage ;

- $\underline{\mathbf{H}}_c^r = (\mathbf{H}_1^r(M_{N-N_c+1}), \mathbf{H}_2^r(M_{N-N_c+1}), \mathbf{H}_3^r(M_{N-N_c+1}), \dots, \mathbf{H}_1^r(M_N), \mathbf{H}_2^r(M_N), \mathbf{H}_3^r(M_N))^T \in (\mathbb{R}^3)^{N_c}$, est le vecteur contenant les valeurs des composantes de $\mathcal{H}^r$ sur les coins du maillage.

Soit $\mathbb{A}_{\tau,f} \in (\mathbb{R}^{3 \times 3})^{N_f \times N_f}$ la matrice $\mathbb{A}_{f,f}$ dont on a éliminé la ligne dépendant de $\boldsymbol{\nu}_i$: $\forall i, j \in I_f$,

$$\mathbb{A}_{\tau,f}^{i,j} = \begin{pmatrix} (v_j \boldsymbol{\tau}_j^1, v_i \boldsymbol{\tau}_i^1)_{\mathcal{X}_{\mathcal{H}}^0} & (v_j \boldsymbol{\tau}_j^2, v_i \boldsymbol{\tau}_i^1)_{\mathcal{X}_{\mathcal{H}}^0} & (v_j \boldsymbol{\nu}_j, v_i \boldsymbol{\tau}_i^1)_{\mathcal{X}_{\mathcal{H}}^0} \\ (v_j \boldsymbol{\tau}_j^1, v_i \boldsymbol{\tau}_i^2)_{\mathcal{X}_{\mathcal{H}}^0} & (v_j \boldsymbol{\tau}_j^2, v_i \boldsymbol{\tau}_i^2)_{\mathcal{X}_{\mathcal{H}}^0} & (v_j \boldsymbol{\nu}_j, v_i \boldsymbol{\tau}_i^2)_{\mathcal{X}_{\mathcal{H}}^0} \\ 0 & 0 & 0 \end{pmatrix}.$$

Soit $\mathbb{A}_{\tau,a} \in (\mathbb{R}^{3 \times 3})^{N_f \times N_a}$ la matrice $\mathbb{A}_{f,a}$ dont on a éliminé la ligne dépendant de ν_i :
 $\forall i \in I_f, \forall j \in I_a$,

$$\mathbb{A}_{\tau,f}^{i,j} = \begin{pmatrix} (v_j \tau_j, v_i \tau_i^1)_{\mathcal{X}_{\mathcal{H}}^0} & (v_j \nu_j^1, v_i \tau_i^1)_{\mathcal{X}_{\mathcal{H}}^0} & (v_j \nu_j^2, v_i \tau_i^1)_{\mathcal{X}_{\mathcal{H}}^0} \\ (v_j \tau_j, v_i \tau_i^2)_{\mathcal{X}_{\mathcal{H}}^0} & (v_j \nu_j^1, v_i \tau_i^2)_{\mathcal{X}_{\mathcal{H}}^0} & (v_j \nu_j^2, v_i \tau_i^2)_{\mathcal{X}_{\mathcal{H}}^0} \\ 0 & 0 & 0 \end{pmatrix}.$$

Soit $\mathbb{A}_{\tau,c} \in (\mathbb{R}^{3 \times 3})^{N_f \times N_c}$, la matrice $\mathbb{A}_{f,c}$ dont on a éliminé la ligne dépendant de ν_i :
 $\forall i \in I_f, \forall j \in I_c$,

$$\mathbb{A}_{\tau,c}^{i,j} = \begin{pmatrix} (v_j \mathbf{e}_1, v_i \tau_i^1)_{\mathcal{X}_{\mathcal{H}}^0} & (v_j \mathbf{e}_2, v_i \tau_i^1)_{\mathcal{X}_{\mathcal{H}}^0} & (v_j \mathbf{e}_3, v_i \tau_i^1)_{\mathcal{X}_{\mathcal{H}}^0} \\ (v_j \mathbf{e}_1, v_i \tau_i^2)_{\mathcal{X}_{\mathcal{H}}^0} & (v_j \mathbf{e}_2, v_i \tau_i^2)_{\mathcal{X}_{\mathcal{H}}^0} & (v_j \mathbf{e}_3, v_i \tau_i^2)_{\mathcal{X}_{\mathcal{H}}^0} \\ 0 & 0 & 0 \end{pmatrix}.$$

Soit $\mathbb{A}_{\tau_a,f} \in (\mathbb{R}^{3 \times 3})^{N_a \times N_f}$ la matrice $\mathbb{A}_{a,f}$ dont on a éliminé les lignes dépendant de $\nu_i^{1,2}$:
 $\forall i \in I_a, \forall j \in I_f$,

$$\mathbb{A}_{\tau_a,f}^{i,j} = \begin{pmatrix} (v_j \tau_j^1, v_i \tau_i)_{\mathcal{X}_{\mathcal{H}}^0} & (v_j \tau_j^2, v_i \tau_i)_{\mathcal{X}_{\mathcal{H}}^0} & (v_j \nu_j, v_i \tau_i)_{\mathcal{X}_{\mathcal{H}}^0} \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}.$$

Soit $\mathbb{A}_{\tau_a,a} \in (\mathbb{R}^{3 \times 3})^{N_a \times N_a}$ la matrice $\mathbb{A}_{a,a}$ dont on a éliminé les lignes dépendant de $\nu_i^{1,2}$:
 $\forall i, j \in I_a$,

$$\mathbb{A}_{\tau_a,a}^{i,j} = \begin{pmatrix} (v_j \tau_j, v_i \tau_i)_{\mathcal{X}_{\mathcal{H}}^0} & (v_j \nu_j^1, v_i \tau_i)_{\mathcal{X}_{\mathcal{H}}^0} & (v_j \nu_j^2, v_i \tau_i)_{\mathcal{X}_{\mathcal{H}}^0} \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}.$$

Soit $\mathbb{A}_{\tau_a,c} \in (\mathbb{R}^{3 \times 3})^{N_a \times N_c}$, la matrice $\mathbb{A}_{a,c}$ dont on a éliminé les lignes dépendant de $\nu_i^{1,2}$:
 $\forall i \in I_a, \forall j \in I_c$,

$$\mathbb{A}_{\tau_a,c}^{i,j} = \begin{pmatrix} (v_j \mathbf{e}_1, v_i \tau_i)_{\mathcal{X}_{\mathcal{H}}^0} & (v_j \mathbf{e}_2, v_i \tau_i)_{\mathcal{X}_{\mathcal{H}}^0} & (v_j \mathbf{e}_3, v_i \tau_i)_{\mathcal{X}_{\mathcal{H}}^0} \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}.$$

Soit $\underline{\mathbf{M}}_{\Omega} \in (\mathbb{R}^3)^{N_{\Omega}}$ la restriction de $\underline{\mathbf{M}}$ aux $i \in I_{\Omega}$, $\underline{\mathbf{M}}_{\tau} \in (\mathbb{R}^3)^{N_f}$ le vecteur tel que :

$$\underline{\mathbf{M}}_{\tau} = (\mathcal{M}_0(v_{N_{\Omega}+1} \tau_{N_{\Omega}+1}^1), \mathcal{M}_0(v_{N_{\Omega}+1} \tau_{N_{\Omega}+1}^2), 0, \dots, \mathcal{M}_0(v_{N_{\Omega}+N_f} \tau_{N_{\Omega}+N_f}^1), \mathcal{M}_0(v_{N_{\Omega}+N_f} \tau_{N_{\Omega}+N_f}^2), 0),$$

et enfin $\underline{\mathbf{M}}_{\tau_a} \in (\mathbb{R}^3)^{N_a}$ le vecteur tel que :

$$\underline{\mathbf{M}}_{\tau_a} = (\mathcal{M}_0(v_{N-N_{ac}+1} \boldsymbol{\tau}_{N-N_{ac}+1}), 0, 0, \dots, \mathcal{M}_0(v_{N-N_c} \boldsymbol{\tau}_{N-N_c}), 0, 0),$$

Soit $\mathbb{A}'_r \in (\mathbb{R}^{3 \times 3})^{N \times N}$ la matrice définie ainsi :

$$\mathbb{A}'_r = \begin{pmatrix} 0 & \mathbb{A}_{\Omega, f} & \mathbb{A}_{\Omega, a} & \mathbb{A}_{\Omega, c} \\ 0 & \mathbb{A}_{\tau, f} & \mathbb{A}_{\tau, a} & \mathbb{A}_{\tau, c} \\ 0 & \mathbb{A}_{\tau_a, f} & \mathbb{A}_{\tau_a, a} & \mathbb{A}_{\tau_a, c} \\ 0 & 0 & 0 & -\mathbb{I}_{c, c} \end{pmatrix}.$$

Posons $\underline{\mathbf{M}}_0 = (\underline{\mathbf{M}}_{\Omega}, \underline{\mathbf{M}}_{\tau}, \underline{\mathbf{M}}_{\tau_a}, \underline{\mathbf{Z}}_c)^T \in (\mathbb{R}^3)^N$, où $\underline{\mathbf{Z}}_c$ est le vecteur nul de $(\mathbb{R}^3)^{N_c}$. On montre de la même façon que dans le paragraphe 4.7.1 que les équations (16.29) se réécrivent sous la forme du système matriciel suivant :

$$\boxed{\mathbb{A}'_0 \underline{\mathbf{H}} = \underline{\mathbf{M}}_0 - \mathbb{A}'_r \underline{\mathbf{H}}^r.}$$

Le calcul de $\underline{\mathbf{M}}_0$ est similaire à celui de $\underline{\mathbf{M}}$. Une fois que $\underline{\mathbf{H}}$ est calculé, on peut écrire l'approximation calculée sur les arêtes du bord, $\underline{\mathbf{H}}_a$ dans la base canonique $(\mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3)$.

Pour cela, à chaque point $M_i, i \in I_a$, on associe $\mathbb{O}_{a,i} \in \mathbb{R}^{3 \times 3}$, la matrice de changement de base telle que :

$$\mathbb{O}_{a,i} = \begin{pmatrix} \tau_1(M_i) & \nu_1^1(M_i) & \nu_1^2(M_i) \\ \tau_2(M_i) & \nu_2^1(M_i) & \nu_2^2(M_i) \\ \tau_3(M_i) & \nu_3^1(M_i) & \nu_3^2(M_i) \end{pmatrix},$$

avec $\boldsymbol{\tau} = \sum_{\alpha=1}^3 \tau_{\alpha} \mathbf{e}_{\alpha}$ et pour $\beta \in \{1, 2\}$, $\boldsymbol{\nu}^{\beta} = \sum_{\alpha=1}^3 \nu_{\alpha}^{\beta} \mathbf{e}_{\alpha}$. Soit $\mathbb{O}_a \in (\mathbb{R}^3)^{N_a \times N_a}$ la matrice diagonale par bloc définie par :

$$\mathbb{O}_a = \begin{pmatrix} \mathbb{O}_{a, N-N_{ac}+1} & 0 & 0 \\ 0 & \ddots & 0 \\ 0 & 0 & \mathbb{O}_{a, N-N_c} \end{pmatrix}.$$

Le vecteur correspondant aux composantes dans la base $(\mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3)^T$ s'écrit alors : $\mathbb{O}_a \underline{\mathbf{H}}_a$.

Le vecteur correspondant aux composantes de $\underline{\mathbf{H}}_f$ dans la base $(\mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3)^T$ s'écrit : $\mathbb{O}_f \underline{\mathbf{H}}_f$.

Chapitre 17

Rappels de notations

17.1 Rappel des espaces fonctionnels et des formes bilinéaires 2D

- **Espace variationnel et forme bilinéaire relatifs à la méthode aux CL naturelles :**

Espace variationnel de $\mathbf{E}$:

$$\mathbf{X}_{\mathbf{E}} = \{ \mathbf{u} \in \mathbf{H}(\text{rot}, \omega) \cap \mathbf{H}(\text{div}, \omega) : \mathbf{u}_\tau \in L^2(\partial\omega) \}.$$

Forme bilinéaire coercitive associée :

$$\begin{aligned} \mathcal{A}_{\mathbf{E}} : \mathbf{X}_{\mathbf{E}} \times \mathbf{X}_{\mathbf{E}} &\rightarrow \mathbb{R} \\ (\mathbf{E}, \mathbf{F}) &\mapsto (\text{div } \mathbf{E}, \text{div } \mathbf{F})_0 + (\text{rot } \mathbf{E}, \text{rot } \mathbf{F})_0 + \int_{\partial\omega} \mathbf{E}_\tau \mathbf{F}_\tau \, d\sigma, \end{aligned}$$

- **Espace variationnel et forme bilinéaire relatifs à la MCS :**

Espace variationnel de $\mathbf{E}$:

$$\mathbf{X}_{\mathbf{E}}^0 = \{ \mathbf{u} \in \mathbf{X}_{\mathbf{E}} : \mathbf{u}_\tau = 0 \}, \quad \|\mathbf{u}\|_{\mathbf{X}_{\mathbf{E}}^0} = (\|\text{rot } \mathbf{u}\|_0^2 + \|\text{div } \mathbf{u}\|_0^2)^{1/2}.$$

Forme bilinéaire coercitive associée :

$$\begin{aligned} \mathcal{A}_0 : \mathbf{X}_{\mathbf{E}} \times \mathbf{X}_{\mathbf{E}} &\rightarrow \mathbb{R} \\ (\mathbf{E}, \mathbf{F}) &\mapsto (\text{div } \mathbf{E}, \text{div } \mathbf{F})_0 + (\text{rot } \mathbf{E}, \text{rot } \mathbf{F})_0. \end{aligned}$$

- **Espace variationnel et forme bilinéaire relatifs à la méthode de régularisation à poids :**

Espace variationnel de $\mathbf{E}$:

$$\mathbf{X}_{\mathbf{E}, \gamma}^0 = \{ \mathbf{u} \in \mathbf{H}_0(\text{rot}, \omega) : \text{div } \mathbf{u} \in L_\gamma^2(\omega) \}, \quad \|\mathbf{u}\|_{\mathbf{X}_{\mathbf{E}, \gamma}^0} = (\|\text{rot } \mathbf{u}\|_0^2 + \|\text{div } \mathbf{u}\|_{0, \gamma}^2)^{1/2}.$$

Forme bilinéaire coercitive associée :

$$\begin{aligned} \mathcal{A}_\gamma : \mathbf{X}_{\mathbf{E}, \gamma}^0 \times \mathbf{X}_{\mathbf{E}, \gamma}^0 &\rightarrow \mathbb{R} \\ (\mathbf{E}, \mathbf{F}) &\mapsto (\text{div } \mathbf{E}, \text{div } \mathbf{F})_{0, \gamma} + (\text{rot } \mathbf{E}, \text{rot } \mathbf{F})_0. \end{aligned}$$

17.2 Rappel des espaces fonctionnels et des formes bilinéaires 3D

- Espace variationnel de $\mathcal{E}$ relatif à la méthode aux CL naturelles :

$$\mathcal{X}_{\mathcal{E}} = \left\{ \mathbf{u} \in \mathcal{H}(\mathbf{rot}, \Omega) \cap \mathcal{H}(\mathbf{div}, \Omega) : \mathbf{u} \times \boldsymbol{\nu}_{|\partial\Omega} \in \mathcal{L}_t^2(\partial\Omega) \right\}.$$

- Espace variationnel de $\mathcal{E}$ relatif à la MCS :

$$\mathcal{X}_{\mathcal{E}}^0 = \left\{ \mathbf{u} \in \mathcal{X}_{\mathcal{E}} : \mathbf{u} \times \boldsymbol{\nu}_{|\partial\Omega} = 0 \right\}, \quad \|\mathbf{u}\|_{\mathcal{X}_{\mathcal{E}}^0} = \left(\|\mathbf{rot} \mathbf{u}\|_0^2 + \|\mathbf{div} \mathbf{u}\|_0^2 \right)^{1/2}.$$

- Espace variationnel de $\mathcal{E}$ relatif à la méthode de régularisation à poids :

$$\mathcal{X}_{\mathcal{E}, \gamma}^0 = \left\{ \mathbf{u} \in \mathcal{H}_0(\mathbf{rot}, \Omega) \cap \mathcal{H}(\mathbf{div}_{\gamma}, \Omega) \right\}, \quad \|\mathbf{u}\|_{\mathcal{X}_{\mathcal{E}, \gamma}^0} = \left(\|\mathbf{rot} \mathbf{u}\|_0^2 + \|\mathbf{div} \mathbf{u}\|_{0, \gamma}^2 \right)^{1/2}.$$

Bibliographie

- [1] A. ALONSO, *A Mathematical justification of the Low-frequency Heterogeneous Time-harmonic Maxwell equations*, Math. Meth. Appl. Sci., 9 (1999), pp. 475–489.
- [2] A. ALONSO AND A. VALLI, *Some remarks on the characterization of the space of tangential traces of $H(\text{rot}; \Omega)$ and the construction of an extension operator*, Manuscripta Math., (1996), pp. 159–178.
- [3] ———, *Unique Solvability for High-frequency Heterogeneous Time-harmonic Maxwell equations via the Fredholm Alternative Theory*, Math. Meth. Appl. Sci., 21 (1998), pp. 463–477.
- [4] M. AMARA AND M. MOUSSAOUI, *Approximation of solutions and singularities coefficients for an elliptic problem in a plane polygonal problem*, tech. rep., École Nationale Supérieure de Lyon, 1989.
- [5] C. AMROUCHE, C. BERNARDI, M. DAUGE, AND V. GIRAULT, *Vector potentials in three-dimensional non-smooth domains*, Math. Meth. Appl. Sci., (1998), pp. 823–864.
- [6] F. ASSOUS AND P. CIARLET, JR., *Modèles et méthodes pour les équations de Maxwell*, Rapport de Recherche 349, ENSTA, Paris.
- [7] F. ASSOUS, P. CIARLET, JR., E. GARCIA, AND J. SEGRÉ, *Time dependent Maxwell's equations with charges in singular geometries*. Soumis à Comp. Meth. Appl. Mech. and Eng., 2005.
- [8] F. ASSOUS, P. CIARLET, JR., AND J. SEGRÉ, *Numerical solution to the time-dependent Maxwell equations in two-dimensional singular domains : the singular complement method*, J. Comput. Phys., (2000), pp. 218–249.
- [9] F. ASSOUS, P. CIARLET, JR., AND E. SONNENDRÜCKER, *Resolution of the Maxwell equations in a domain with reentrant corners*, Modél. Math. Anal. Numér., (1998), pp. 359–389.
- [10] F. ASSOUS, P. DEGOND, E. HEINTZÉ, P.-A. RAVIART, AND J. SEGRÉ, *On finite element method for solving the Three-Dimensional Maxwell Equations*, J. Comput. Phys., (1993), pp. 222–237.
- [11] F. ASSOUS, P. DEGOND, AND J. SEGRÉ, *Numerical approximation of the Maxwell Equations in Inhomogeneous Media by a P^1 Conforming Finite Element Method*, J. Comput. Phys., (1993), pp. 222–237.
- [12] I. BABUSKA, *The finite element method with Lagrangian multipliers*, Numer. Math., (1973), pp. 179–192.
- [13] R. BARTHELMÉ, *Le problème de conservation de la charge dans le couplage des équations de Maxwell et de Vlasov*, PhD thesis, Université Louis Pasteur, Strasbourg, 2005.
- [14] Z. BELHACHMI, C. BERNARDI, AND S. DEPARIS, *Weighted Clément operator and application to finite element discretization of the axisymmetric Stokes problem*, Rapport de Recherche R03029, Laboratoire Jacques-Louis Lions, Paris, 2003.

- [15] F. BEN BELGACEM AND C. BERNARDI, *Spectral element discretization of the Maxwell equations*, Math. Comput., 68 (1999), pp. 1497–1520.
- [16] J.-P. BÉRENGER, *A perfectly matched layer for the absorption of electromagnetic waves*, J. Comput. Phys., (1994), pp. 185–200.
- [17] C. BERNARDI AND E. SÜLI, *Time and space adaptivity for the second-order wave equation*, Math. Meth. Appl. Sci., 15 (2005), pp. 199–225.
- [18] M. S. BIRMAN AND M. Z. SOLOMYAK, *L^2 -theory of the Maxwell operator in arbitrary domains*, Russian Mat. Surveys, (1987), pp. 75–96.
- [19] D. BOFFI, *Three dimensional finite element methods for the Stokes problem*, SIAM J. Numer. Anal., 34 (1997), pp. 664–670.
- [20] D. BOFFI, M. COSTABEL, M. DAUGE, AND L. DEMKOWICZ, *Discrete compactness for the hp version of rectangular edge finite elements*, ICES Report 04-29, The University of Texas at Austin, 2004.
- [21] D. BOFFI, L. DEMKOWICZ, AND M. COSTABEL, *Discrete compactness for p and hp 2D edge finite elements*, Math. Meth. Appl. Sci., 13 (2003), pp. 1673–1687.
- [22] D. BOFFI, M. FARINA, AND L. GASTALDI, *On the approximation of Maxwell's eigenproblem in general 2D domains*, Comput. Struct., 79 (2001), pp. 1089–1096.
- [23] D. BOFFI AND L. GASTALDI, *On the time harmonic Maxwell equations in general domains*, in Numerical Mathematics and Advanced Applications, ENUMATH 2001, Springer-Verlag Italia, 2003, pp. 243–253.
- [24] A.-S. BONNET-BEN DHIA AND E. LUNEVILLE, *Résolution numérique des équations aux dérivées partielles*, Cours MA201, ENSTA, Paris, 2000-2001.
- [25] A. BOSSAVIT, *Électromagnétisme en vu de la modélisation*, vol. 14 of Mathématiques et Applications, Springer-Verlag, 1991.
- [26] W. E. BOYSE, D. R. LYNCH, K. D. PAULSEN, AND G. N. MINERBO, *Node based finite element modelling of Maxwell's equations*, IEEE Trans. Antennas Propagat., (1992), pp. 642–651.
- [27] H. BREZIS, *Analyse fonctionnelle. Théorie et applications*, Masson, Paris, 1983.
- [28] F. BREZZI, *On the existence, uniqueness and approximation of saddle-point problems arising from Lagrange multipliers*, R.A.I.R.O. Anal. Numer., (1974), pp. 129–151.
- [29] A. BUFFA AND P. CIARLET, JR., *On traces for functional spaces related to Maxwell's equations. Part I : an integration by parts formula in Lipschitz polyhedra*, Math. Meth. Appl. Sci., (2001), pp. 9–30.
- [30] —, *On traces for functional spaces related to Maxwell's equations. Part II : Hodge decomposition on the boundary of Lipschitz polyhedra and applications*, Math. Meth. Appl. Sci., (2001), pp. 9–30.
- [31] A. BUFFA, M. COSTABEL, AND M. DAUGE, *Anisotropic regularity results for Laplace and Maxwell operators in a polyhedron*, C. R. Acad. Sci. Paris, Série I, (2003), pp. 565–570.
- [32] —, *Algebraic convergence for anisotropic edge element in polyhedral domains*. A paraître dans Numer. Math., 2004.
- [33] P. G. CIARLET, *The finite element method for elliptic problems*, North-Holland, Amsterdam, 1978.

- [34] P. CIARLET, JR., *Étude de préconditionnements parallèles pour la résolution d'équations aux dérivées partielles elliptiques. Une décomposition de l'espace $L^2(W)^3$* , PhD thesis, Université Paris VI, 1992.
- [35] —, *Calcul des $\beta_{D,N}^{i,j}$* . Communication privée, 2002.
- [36] —, *Augmented formulations for solving Maxwell equations*, Comp. Meth. Appl. Mech. and Eng., 194 (2005), pp. 559–586.
- [37] —, *Formulation variationnelle des équations de Maxwell instationnaire*. Communication privée, 2005.
- [38] P. CIARLET, JR., G. GARCIA, AND J. ZOU, *Résolution des équations de Maxwell dans des domaines prismatiques tridimensionnels*, C. R. Acad. Sci. Paris, Série I, (2004), pp. 721–726.
- [39] P. CIARLET, JR. AND V. GIRAULT, *Condition inf-sup pour l'élément fini de Taylor-Hood P_2 -iso- P_1 , 3-D; applications aux équations de Maxwell*, C. R. Acad. Sci. Paris, Série I, (2002), pp. 827–832.
- [40] P. CIARLET, JR., C. HAZARD, AND S. LOHRENGEL, *Les équations de Maxwell dans un polyèdre : un résultat de densité*, C. R. Acad. Sci. Paris, Série I, (1998), pp. 305–310.
- [41] P. CIARLET, JR. AND J. HE, *The singular complement method for 2D scalar problems*, C. R. Acad. Sci. Paris, Série I, (2003), pp. 353–358.
- [42] P. CIARLET, JR., B. JUNG, S. KADDOURI, S. LABRUNIE, AND J. ZOU, *Fourier singular complement method. Part I : prismatic case*. A paraître dans Numer. Math., 2005.
- [43] P. CIARLET, JR. AND J. ZOU, *Finite element convergence for the Darwin model to Maxwell's equations*, Modél. Math. Anal. Numér., (1997), pp. 213–250.
- [44] —, *Fully discrete finite element approaches for time-dependent Maxwell's equations*, Numer. Math., (1999), pp. 193–219.
- [45] G. COHEN, *Higher-Order Numerical Methods for Transient Wave Equations*, Scientific Computation, Springer-Verlag, Berlin, 2002.
- [46] G. COHEN, P. JOLY, E. ROBERTS, AND N. TORDJMAN, *Higher order triangular finite elements with mass lumping for the wave equation*, SIAM J. Numer. Anal., 38 (2001), pp. 2047–2078.
- [47] G. COHEN AND P. MONK, *Gauss point mass lumping schemes for Maxwell's equations*, Numer. Meth. Partial Diff. Eqns., (1998), pp. 63–88.
- [48] —, *Mur-Nédélec finite element schemes for Maxwell's equations*, Comput. Meth. Appl. Mech. Eng., (1998), pp. 157–171.
- [49] M. COSTABEL, *A remark on the regularity of solutions of Maxwell's equations on Lipschitz domains*, Math. Meth. Appl. Sci., (1990), pp. 365–368.
- [50] —, *A coercive bilinear form for Maxwell's equations*, J. Math. An. Appl., (1991), pp. 527–541.
- [51] M. COSTABEL AND M. DAUGE, *Un résultat de densité pour les équations de Maxwell régularisées dans un domaine lipschitzien*, C. R. Acad. Sci. Paris, Ser. I, (1998), pp. 849–854.
- [52] —, *Singularities of electromagnetic fields in polyhedral domains*, Arch. Ration. Mech. Anal., (2000), pp. 221–276.
- [53] —, *Weighted regularization of Maxwell equations in polyhedral domains, a rehabilitation of Nodal finite elements*, Numer. Math., (2002), pp. 239–277.

- [54] M. COSTABEL, M. DAUGE, AND D. MARTIN, *Numerical investigation of a boundary penalization method for Maxwell equations*, in Proceedings of the 3rd European conference on numerical mathematics and advanced applications, Singapore, 2000, World Scientific, pp. 214–221.
- [55] M. COSTABEL, M. DAUGE, AND S. NICAISE, *Singularities of Maxwell interface problems*, R.A.I.R.O. Modél. Math. Anal. Numér., (1999), pp. 627–649.
- [56] M. COSTABEL, M. DAUGE, AND C. SCHWAB, *Exponential convergence of hp-FEM for Maxwell's equations with weighted regularization in polygonal domains*, Math. Models Meth. Appl. Sci., (2004).
- [57] W. DAHMEN, T. KLINT, AND K. URBAN, *On fictitious domain formulation for Maxwell equations*, Found. Comput. Math., (2003), pp. 135–160.
- [58] R. DAUTRAY AND J.-L. LIONS, *Analyse mathématique et calcul numérique pour les sciences et les techniques*, Masson, Paris, 1988.
- [59] P. DEGOND AND P.-A. RAVIART, *An analysis of the Darwin model of approximation to Maxwell's equations*, Forum Math., 4 (1992), pp. 13–44.
- [60] L. DEMKOWICZ AND L. VARDAPETYAN, *Modeling of electromagnetic absorption / scattering problems using hp-adaptative finite elements*, Comput. Methods Appl. Mech. Engrg, 152 (1998), pp. 103–124.
- [61] P. FERNANDES AND G. GILARDI, *Magnetostatic and electrostatic problems in inhomogeneous anisotropic media with irregular boundary and mixed boundary conditions*, Math. Models Meth. App. Sci., 7 (1997), pp. 957–991.
- [62] R. P. FEYNMAN, R. B. LEIGHTON, AND M. SANDS, *The Feynman Lectures on Physics*, vol. 3, Addison-Wesley, Reading (Massachusetts), 1963.
- [63] E. GARCIA, *Résolution des équations de Maxwell instationnaires avec charges dans des domaines non convexes*, PhD thesis, Université Paris VI, 2002.
- [64] V. GIRAULT, R. H. NOCHETTO, AND R. SCOTT, *Maximum-norm stability of the finite element Stokes projection*. A paraître dans J. Math. Pures Appl., 2005.
- [65] V. GIRAULT AND P.-A. RAVIART, *Finite element methods for Navier-Stokes equations*, Springer Series in Computational Mathematics, Springer Verlag, Berlin, 1986.
- [66] P. GRISVARD, *Elliptic problems in nonsmooth domains*, vol. 24 of Monographs and studies in Mathematics, Pitman, London, 1985.
- [67] ———, *Singularities in boundary value problems*, Research Notes in Applied Mathematics, Springer Verlag, Berlin, 1992.
- [68] C. HAZARD AND S. LOHRENGEL, *A singular field method for Maxwell's equations : numerical aspects for 2D magnetostatics*, SIAM J. Numer. Anal., (2002), pp. 1021–1040.
- [69] E. HEINTZÉ, *Résolution des équations de Maxwell tridimensionnelles instationnaires par une méthode d'éléments finis conformes*, PhD thesis, Université Paris VI, 1992.
- [70] P. HOUSTON, I. PERUGIA, A. SCHNEEBELI, AND D. SCHÖTZAU, *Interior penalty method for the indefinite time-harmonic Maxwell equations*, Numer. Math., (2005), pp. 485–518.
- [71] J. D. JACKSON, *Électrodynamique classique*, Sciences Sup, Dunod, Paris, 2001.
- [72] E. JAMELOT, *Éléments finis nodaux pour les équations de Maxwell*, C. R. Acad. Sci. Paris, Ser. I, 339 (2004), pp. 809–814.
- [73] ———, *Corrigendum à la Note "Éléments finis nodaux pour les équations de Maxwell"*, C. R. Acad. Sci. Paris, Ser. I, 340 (2005), pp. 409–410.

- [74] F. KIKUCHI, *Mixed and penalty formulations for finite element analysis of an eigenvalue problem in electromagnetism*, Comput. Meths. Appl. Mech. Engrg., 64 (1987).
- [75] —, *On a discrete compactness property for the Nedelec finite elements*, J. Fac. Sci. Univ. Tokyo, Sect. IA, Math., 36 (1989), pp. 479–490.
- [76] S. LABRUNIE, *Les équations de Maxwell en géométrie axisymétrique et la méthode du complément singulier*, in Journée Éléments finis vectoriels, Paris, 2004, ENSTA.
- [77] P. LACOSTE, *La condensation de la matrice de masse, ou mass-lumping, pour les éléments finis de Raviart-Thomas-Nédélec d'ordre 1*, C. R. Acad. Sci. Paris, Série I, (2004), pp. 727–732.
- [78] J. LIONS AND E. MAGENES, *Problèmes aux limites non homogènes et applications*, vol. 1, Dunod, 1968.
- [79] S. LOHRENGEL AND S. NICAISE, *Les équations de Maxwell dans des matériaux composites : problèmes de densité*, C. R. Acad. Sci. Paris, Série I, (2000), pp. 991–996.
- [80] P. MONK, *A comparison of three mixed methods for the time-dependent Maxwell's equations*, SIAM J. Sci. Stat. Comput., 13 (1992), pp. 1097–1122.
- [81] —, *Analysis of finite element methods for Maxwell's equations*, SIAM J. Numer. Anal., 29 (1992), pp. 714–729.
- [82] —, *Finite element methods for Maxwell's equations*, Numerical Mathematics and Scientific Publication, Clarendon Press, 2003.
- [83] M. MOUSSAOUI, *Sur l'approximation des solutions du problème de Dirichlet dans un ouvert avec coins*, in Singularities and constructive methods for their treatment, P. Grisvard, W. Wendlan, and J. R. Whiteman, eds., no. 1121 in Lectures Notes in Mathematics, Springer Verlag, Berlin, 1984, pp. 199–206.
- [84] —, *Espaces $H(\text{div}, \text{rot}; \Omega)$ dans un polygone plan*, C. R. Acad. Sci. Paris, Série I, (1996), pp. 225–229.
- [85] C. MÜLLER, *Foundations of the Mathematical Theory of Electromagnetic Wave*, Springer, Berlin, 1969.
- [86] S. A. NAZAROV AND B. A. PLAMENEVSKY, *Elliptic problems in domains with piecewise smooth boundaries*, Walter de Gruyter, Berlin, New York, 1994.
- [87] J.-C. NÉDÉLEC, *Mixed finite elements in $\mathbb{R}^3$* , Numer. Math., (1980), pp. 315–341.
- [88] —, *A new family of mixed finite elements in $\mathbb{R}^3$* , Numer. Math., (1986), pp. 57–81.
- [89] S. NICAISE, *Exact boundary controllability of Maxwell's equations in heterogeneous media and an application to an inverse problem*, SIAM J. Control Optim., (2000), pp. 1145–1170.
- [90] I. PERUGIA AND D. SCHÖTZAU, *The hp-local discontinuous Galerkin methods for low-frequency time-harmonic Maxwell equations*, Math. Comp., (2003), pp. 1179–1214.
- [91] I. PERUGIA, D. SCHÖTZAU, AND P. MONK, *Stabilized interior penalty methods for the time-harmonic Maxwell equations*, Comput. Meth. Appl. Mech. Engrg, (2002), pp. 4675–4697.
- [92] W. RACOWICZ AND L. DEMKOWICZ, *An hp-adaptative finite element method for electromagnetics-part II : Data structure and constraint approximation*, Int. J. Numer. Meth. Engen, (2000), pp. 307–337.
- [93] —, *An hp-adaptative finite element method for electromagnetics-part I : A 3D implementation*, Int. J. Numer. Meth. Engen, (2002), pp. 147–180.
- [94] K. RAMDANI, *Lignes supraconductrices : analyse mathématique et numérique*, PhD thesis, Université Pierre et Marie Curie, 1999.

- [95] G. RAUGEL, *Résolution numérique de problèmes elliptiques dans des domaines avec coins*, PhD thesis, Université de Rennes, 1978.
- [96] P.-A. RAVIART AND E. SONNENDRÜCKER, *Approximate models for the Maxwell equations*, J. Comput. Appl. Math., (1995), pp. 69–81.
- [97] ———, *A hierarchy of approximate models for the Maxwell equations*, Numer. Math., (1996), pp. 329–372.
- [98] P.-A. RAVIART AND J.-M. THOMAS, *A mixed finite element method for the second order elliptic problems*, in Mathematical aspects of the finite element method, A. Dold and B. Eckmann, eds., no. 606, Springer, London.
- [99] ———, *Primal hybrid finite element methods for second order elliptic problems*, Math. Comput., (1977), pp. 391–413.
- [100] J. RENAULT, *Formulaire de Physique*, Dunod Université, Paris, 1980.
- [101] L. SCHWARTZ, *Théorie des distributions*, Hermann, Paris, 1966.
- [102] P. SOLIN, K. SEGETH, AND I. DOLEZEL, *Higher-Order Finite Element Methods*, CRC Press, Boca Raton, Florida, 2003.
- [103] E. SONNENDRÜCKER, J. J. AMBROSIANO, AND S. T. BRANDON, *A Finite Element Formulation of the Darwin PIC Model for Use on Unstructured Grids*, J. Comput. Phys., (1995), pp. 281–297.
- [104] A. TAFLOVE, *Computational electrodynamics : The finite-difference time-domain method*, Artech House, 1995.
- [105] T. VAN AND A. WOOD, *A time marching finite element method for an electromagnetic scattering problem*, Math. Meth. Appl. Sci., (2003), pp. 1025–1045.
- [106] R. VERFÜRTH, *Error estimates for a mixed finite element approximation of the Stokes equations*, R.A.I.R.O., Anal. Numer. R2, (1984), pp. 175–182.
- [107] C. WEBER, *A local compactness theorem for Maxwell's equations*, Math. Models Meth. Appl. Sci., (1980), pp. 12–25.
- [108] K. YEE, *Numerical solutions of initial boundary value problems involving Maxwell's equations in isotropic media*, IEEE Trans. on Antennas and Propagation AP-16, (1966), pp. 302–307.
- [109] J. ZHAO, *Analysis of finite element approximation for time-dependent Maxwell problems*, Math. Comp., 73 (2003), pp. 1089–1105.

RÉSOLUTION DES ÉQUATIONS DE MAXWELL AVEC DES ÉLÉMENTS FINIS DE GALERKIN CONTINUS

Résumé : Les équations de Maxwell se résolvent aisément lorsque le domaine d'étude est régulier, mais lorsqu'il existe des singularités géométriques (coins rentrants en $2D$, coins et arêtes rentrants en $3D$), le champ électromagnétique est localement non borné au voisinage de ces singularités. Nous nous intéressons à la résolution des équations de Maxwell dans des domaines bornés, singuliers, à l'aide de méthodes d'éléments finis continus. En pratique, cela permet de modéliser des instruments de télécommunication comme les guides d'onde, les filtres à stubs.

Nous analysons tout d'abord le problème quasi-électrostatique $2D$, afin de maîtriser la discrétisation en espace. Nous présentons trois méthodes de calcul (formulations augmentées mixtes) qui donnent des résultats numériques très convaincants :

- Une version épurée de la méthode du complément singulier (conditions aux limites essentielles).
- La méthode de régularisation à poids : on introduit un poids qui dépend des distances aux singularités géométriques (conditions aux limites essentielles).
- La méthode avec conditions aux limites naturelles.

Nous étudions ensuite la généralisation de ces méthodes aux domaines $3D$. Nous détaillons la résolution des équations de Maxwell instationnaires en domaines singuliers $3D$ par la méthode de régularisation à poids, et nous donnons des résultats numériques inédits.

Mots-clés : Électromagnétisme, Équations de Maxwell, Singularités, Coins rentrants, Éléments finis, Domaines singuliers, Méthode du complément singulier, Méthode de régularisation à poids.

CONTINUOUS GALERKIN METHODS FOR SOLVING MAXWELL EQUATIONS

Abstract: Maxwell equations are easily solved when the computational domain is regular, but if it presents geometrical singularities (reentrant corners in $2D$, reentrant corners and edges in $3D$), the electromagnetic field is locally unbounded close to these singularities. We are interested in computing solutions to Maxwell equations in singular bounded domains with continuous finite elements. It allows to model communication devices such as wave-guides, stub filters.

We first analyse the $2D$ quasi-electrostatic problem, in order to control space discretization. We present three (mixed) augmented methods, which show very convincing numerical results:

- A refined version of the singular complement method (essential boundary conditions).
- The weighted regularization method: Maxwell-Gauss equation is multiplied by a suitable weight, which depends on the distances to the geometrical singularities (essential boundary conditions).
- The method with natural boundary conditions.

We study then the extension of these methods to $3D$ domains. We give details of the resolution of time-dependent Maxwell equations in singular $3D$ domains with the weight regularization method, and we give original numerical results.

Key words: Electromagnetism, Maxwell equations, Singularities, Reentrant corners, Finite elements, Singular domain, Singular complement method, Weighted regularization method.